

# 法律体系认知硬件的频段错位：生成式人工智能风险规制的深层困境与认知重建路径

## 摘要

中华智问知识发生器 · sdeuniverses.com · 2026/7/24 13:34:01

生成式人工智能的法律风险规制困境，主流叙事将其归结为“技术跑得太快、法律滞后了”。本文提出一个竞争性诊断：困境的根源不在规则供给不足，而在于法律追责结构的认知硬件——以“可识别的损害事件、可追溯的因果链条、可归因的主体意志”三个原子单元为基础的信号接收器——在长期演化中形成了对伤害信号的不对称灵敏结构：它高度响应可追溯、可个体化、可定价的伤害，却系统性迟钝于分布性、不可归因、缓慢累积的伤害。更棘手的是，这种迟钝恰好处于体系的自诊盲区内——用来诊断“法律漏了什么”的那双眼睛，本身就是漏掉的那一部分所塑造的。

本文通过概念分析与制度演化推演，首先解剖法律体系认知硬件的三个原子单元在生成式人工智能面前如何系统性踏空；继而追溯这一认知硬件在过去四十余年中被风险分配对因果追溯的替代、以及效率转向对判断力的排挤两次历史变动缓慢削薄的过程；进而揭示能力极化、商业化闭环、教育出厂校准三条路径如何并联驱动这一不对称灵敏结构的自我加固。在此基础上，本文提出规制困境的修复不能依赖单一动作，而必须让教育入口的“微动作”（在法学院主干课程中植入“法律漏掉了什么”的追问）与制度层面的“大动作”（行业基线透明、判断节点划定、训练过程审计）在时间轴上先后衔接，构成一个约二十年的认知重建螺旋。为检验这一诊断的可信度，本文设计了三套可操作的实证方案：判例文本的盲区编码分析、行业标准制定过程的认知污染追踪、法学院教学干预的准实验设计，每套方案均内置了明确的证伪条件。本文的核心贡献在于将生成式人工智能法律风险的讨论从“缺什么规则”的技术性争论，推进到“法律的认知硬件是否需要重新刻度”的制度认知层面，并为后续研究提供了一个可被多学科独立启动的研究纲领。

## 全球文库校准后的学术定位

既有研究已经识别分布式损害、跨境责任和监管碎片化。本文不再把这些现象据为原创，而把独特命题集中在“风险变化频率与法律认知更新频率的失配”，并明确这一命题须由跨制度事件序列进一步检验。

比较范围：SDE站内全文、arXiv、SSRN、PubMed/PMC及开放期刊；检索截止2026年7月24日。

## 一、引言

2023年以来，生成式人工智能的爆发式进入社会生产和日常生活，催生了一系列传统侵权法体系难以归置的损害形态：大规模深度伪造对被肖像权、名誉权的侵蚀，算法歧视在自动化决策中的隐蔽扩散，语言模型“幻觉”引发的专业误导，以及更潜移默化的语言生态降级和判断力代际衰减。各国立法者以空前速度回应：欧盟《人工智能法案》在2024年正式通过，中国《生成式人工智能服务管理暂行办法》在2023年施行并持续迭代，美国联邦与州层面的AI立法提案在2024-2025年间密集涌现。然而，一个令人不安的局面正在浮现：立法的密集程度与实际规制效果之间的落差，似乎并未因法案数量的增长而缩小。大量真实的AI相关伤害——尤其是那些无法追溯到单一侵权主体、无法量化到具体损害数额、无法截取为清晰“事件”的伤害——仍然游离于法律体系的回应半径之外。

对此的主流解释是“技术跑得太快、法律滞后了”：只要立法者加把劲、把漏洞补上，规制效果自然会改善。但这一叙事无法解释一个经验现象：即便在立法已经明确覆盖的领域（如算法歧视、个人信息保护），执法和司法实践中的实质回应率仍然很低；而真正难以规制的，恰恰是立法者并非“不知道”而是“知道了却不知道如何将之装入现有追责框架”的那一类伤害。这说明，困境的根源可能不在“规则的数量或覆盖范围”，而在更深一层——在法律体系接收和处置伤害信号的基本方式上。

本文提出一个竞争性诊断：法律体系对生成式人工智能风险的“静音”状态，不是因为缺乏规则，而是因为法律追责结构的认知硬件——这套硬件经过数百年的演化，已经高度优化为接收“可识别的损害事件、可追

溯的因果链条、可归因的主体意志”三类信号——在生成式人工智能面前发生了系统性的频段错位。生成式人工智能产生的许多伤害不是“事件”而是“分布”，不是“链条”而是“网络”，不能归给“一个人”而只能归给“一组分散的参与者”。法律接收器不是“坏了”，而是被调校到了一个无法接收这些频率的状态。更棘手的是，这种调校并非AI的突然冲击所致——它是在过去四十余年的两次历史变动中被缓慢削薄的：第一次是侵权法体系在20世纪后期用“风险分配”逻辑逐步替代“因果追溯”逻辑，使得法律从业者追问因果链末端伤害来源的能力系统性弱化；第二次是法律服务商业化带来的“效率转向”——可计费小时、专业分工、案源竞争——排挤了法律从业者对“这件事法律不该管”或“这件事法律管不了”的反思空间。

这一诊断并非纯粹的理论推演。本文将首先解剖法律认知硬件的三个原子单元及其在AI面前的踏空机制（第二章），追溯这套硬件被历史削薄的过程（第三章），分析它为何能持续自我加固而不自知（第四章），并在此基础上提出“微动作先行、大动作后发”的二十年认知重建路径（第五章）。为确保诊断不沦为另一套自我免疫的封闭叙事，第六章将设计三套可操作的实证检验方案，每套方案均内置证伪条件——只有当诊断能被独立检验、甚至被证伪时，这套判断才真正离开哲学圈、进入学术研究的可积累轨道。

本文的学术贡献集中在三个方面。第一，在理论层面，将生成式人工智能法律风险的讨论从“缺什么规则”的技术性争论推进到“法律认知硬件是否需要重新刻度”的制度认知层面，为法理学、科技政策、认知科学的交叉研究提供了一个新的分析框架。第二，在方法论层面，提出了可用于独立实证研究的判例盲区编码、标准污染追踪、教学干预准实验三套操作方案，使得“法律偏瘫”这一诊断从哲学洞见转化为可检验的研究假设。第三，在研究纲领层面，从单一议题中涌现出一个具有自生长潜力的跨学科研究方向——法律体系在技术加速条件下的认知退化与自我重建——为后续至少五年、多个学科独立学者的跟进研究铺设了入口。

## 二、文献综述

### 2.1 技术-法律落差论及其局限

对生成式人工智能法律规制困境的主流讨论，长期围绕“技术-法律落差”这一核心叙事展开。这一叙事的基本逻辑是：技术演进的速率超过法律更新的速率，导致一段时期内出现“规制真空”；立法者需要加快步伐，通过专门立法填补漏洞。沿着这一逻辑，大量研究致力于识别“哪些AI行为尚无法律覆盖”，并提出相应的专规方案：算法审计义务、训练数据合法性审查、深度合成标识义务、高风险AI系统的合格评定，以及损害赔偿基金的建立。

这一研究方向并非没有价值——它确实推动了多项关键立法的出台，也为执法和司法提供了必不可少的制度工具。但“技术-法律落差”叙事有一个隐含假设：只要法律覆盖了某项行为，规制就会生效。实证观察显示，即便在法律已经明确覆盖的领域，规制效果仍不理想。以个人信息保护为例：几乎各国都已颁布专门立法，但生成式AI处理个人信息的实质合规情况并不因立法存在而自动改善——大量个人信息仍在用户不知情的情况下被纳入训练数据，而现有追责机制在证明“谁的数据在何时被以何种方式用于训练哪个模型”这一因果链条时，遭遇了几近不可逾越的信息门槛。

这意味着，规制困境的根源可能不仅是“规则缺乏”，而且是“规则无法在现实中运转”——而这不单纯是执法资源不足的问题，而是更深的认知层面的不匹配：现有追责结构所预设的“可追溯的侵权行为”这个基本前提，在生成式人工智能的语境下经常不成立。此处的关键洞见，与本世纪以来多位科技法学者对算法治理的批判性研究有其相通之处，但本文将进一步追问：如果这种“不成立”不只是偶然的执法困难，而是现有法律认知硬件在面对一类新的伤害时的系统性失灵，那么仅仅在旧认知硬件上叠加新的专规，可能只是在同一套信号接收器上加了更多“可接收”频段的标志，却没有改变接收器本身的物理频率范围。

### 2.2 规制理论的四种进路及其共享预设

在既有的科技法规制研究中，有四种理论进路对理解算法时代的法律困境产生了深远影响，但它们各自把握的承重变量，在生成式人工智能语境下出现了重要的分离点。

第一种进路可以称为“架构论”——主张技术架构本身具有规制效果，代码的选择在某种程度上替代了法律的规则功能。这一观点敏锐地捕捉到了技术设计对行为约束的隐形权力，但它预设架构的设计者是可以被识

别和追究的——在生成式AI的训练过程中，模型行为的最终涌现是数据、算法架构、训练策略、随机种子等几十个变量交互的产物，很难将其归于任何一个设计者的“选择”。架构论的变量（架构决定行为）在这里并未失效，但它与“可归因于设计者”之间的链条被生成式AI的复杂性撑开了。

第二种进路以“信息透明度”为核心变量——认为算法治理的关键在于让“黑箱”变成“白箱”，让外部能够审查算法的决策过程。这一思路催生了算法审计、可解释性要求等一系列制度设计。但信息的可获得性在生成式AI面前面临一个质变：大语言模型的决策不是一条可被截取的线性规则链，而是一个高维空间中的概率分布。即便模型训练的全部信息都公开（包括参数、训练数据、训练日志），也不等于“模型为什么在这个时刻输出这句话”可以得到因果解释。信息的可获得性没有消失，但它与“可归因于特定行为主体”之间的连接在生成式AI这里不是技术上的暂时困难，而是认知层面的范畴不匹配。

第三种进路聚焦于“程序完整性”——认为自动化决策对正当程序的架空是算法治理的核心危险。这一诊断在行政自动化决策和算法评分领域极为精准，但它同样预设了“有一个既定决策程序被绕过或架空”这一前提。在生成式AI深度嵌入内容生产、辅助判断、知识获取等日常场景时，许多伤害并不发生在“决策程序”中——它发生在语言生态的缓慢漂移、判断力的代际衰减、认知多元性的逐渐收缩中。这些伤害没有一个被“架空”的程序可以援引，因为从一开始就没有为它们设定的程序。

第四种进路以“系统的可规制性”为核心——认为系统的开放程度决定了它能否被外部规制。封闭的专有系统难以被审查，开放的模型和数据集更便于外部监督。这一判断在很大范围内有效，但在生成式AI这里遇到一个意外的分离点：一个高度开放的模型（开源权重、透明训练数据、公开评估基准），其输出的分布性伤害（如对特定群体在大量文本中持续且微妙的贬损模式）仍然可以在完全透明的条件下发生，且仍然没有单一的“行为主体”可以为此被追责。开放度虽然在很多场景中提升了可审计性，但并不能直接解决分布性伤害的不可归因性问题。

这四种进路——架构、信息、程序、开放度——各自从不同侧面触及了算法时代的法律困境，也各自产出了有重要实践价值的政策工具。但它们共享一个未经审查的前提：规制主体——也就是法律追责结构的基本认知方式——具备诊断自身盲区的能力。换句话说，它们都假设法律体系能够知道自己应该向哪里看、应该接收什么信号，困难只在于“看到了之后怎么办”或“因为技术原因暂时看不到”。本文要追问的恰恰是这个假设：如果法律体系认知硬件在这一问题上并不具备自诊能力——它用四十余年的时间演化出了一套越来越精密的信号过滤器，把某类伤害信号系统性地挡在了接收门槛之外，而且这一过滤机制长在它自己的感知盲区里——那么，前述四家建议的所有规制工具，都可能在法律体系手中被用得极为娴熟，却恰好漏掉了它最应该接收的那一部分伤害信号。这不是“规制者不够努力”的问题，不是“产业游说太强”的问题，也不是“技术太特殊”的问题——而是规制者在用一套被演化所塑形的认知硬件去感知“自己漏掉了什么”，而这套硬件恰好是被设计来过滤那一类信号的。

## 2.3 侵权法演化的内部线索：风险分配与效率转向

要理解法律认知硬件如何演化出这种不对称灵敏结构，需要回到侵权法体系在20世纪后期至今的两次内部变动。

第一次变动是风险分配逻辑对因果追溯逻辑的逐步替代。传统侵权法以个案式因果追溯为核心：有损害→追溯到行为主体→确认过错或可归责事由→分配责任。但在大规模工业化伤害（产品责任、环境污染、职业健康损害）面前，个案式因果追溯的信息成本变得极高，某些情况下甚至不可能完成。侵权法体系的回应是发展出一套风险分配工具——无过错责任、市场份额责任、举证责任倒置、保险机制、赔偿基金——使得损害可以在不精确追溯每一条因果链的情况下被社会化分配。这套工具在大规模工业化伤害的治理中发挥了关键作用，但它同时带来一个未被注意的副产品：法律从业者追问“从损害追溯到行为”的肌肉，在风险分配框架下不再被日常训练。一个处理了三十年产品责任案件的律师，其专业能力集中在责任比例的分配、保险覆盖范围的计算、赔偿数额的核定——他可能非常精通这套分配游戏，但几乎不需要问“这个损害到底是谁怎么做出来的”。这条肌肉的萎缩在工业化伤害语境下不明显——因为风险分配框架已经足够好用——但在生成式AI造成的分布性伤害面前，当“追溯到行为”不再只是一个可以外包给保险精算的步骤、而是成为唯一可能让伤害被法律体系“看见”的前提时，这条已经萎缩的肌肉就暴露了。



第二次变动是法律服务商业化带来的“效率转向”。自1980年代以来，律师事务所的企业化经营、可计费小时制度（billable hour）的全球扩散、法律服务市场的竞争加剧，共同推动法律从业者的工作模式向“高效处理可标准化案件”倾斜。这并非道德批判——这是制度激励的自然结果：在处理一个案件的时间内可以处理三个相似案件，每个都能获得可预期收费，制度和市场都会奖励这种效率。但这一效率转向排挤了什么？它排挤的是法律从业者在“这件事现有法律可能处理不了”时的停顿和反思——因为这个停顿在可计费小时表上是“无产出”的。一个人越擅长高效地把案件装入现有追责框架，就越不可能在某一时刻停下来问：“这个框架是不是没有为这件事开设接收频道？”这种能力不是被任何人有意扼杀的，而是在制度激励的日积月累中被温和地挤出了日常实践。数十年的效率转向，在法律共同体内部系统性地削薄了一种不进入计费表单的能力——识别“现有分类无法吸收的伤害”的能力。当生成式AI开始大规模产生恰好不能被现有分类吸收的伤害时，法律共同体中那些最擅长高效处理案件的人——也就是最有能力推动制度变革的人——反而最不具备识别这些伤害的认知条件。

## 2.4 文献缺口的诊断

综合以上三条脉络，可以定位出当前研究的三重缺口。

第一，技术-法律落差论诊断了“缺规则”，但没有追问规则即便存在也运转不灵的原因——这个原因不在规则层面，而在规则之下更深一层的认知硬件层面。规制工具已经大量存在，但规制的“信号接收器”本身已经被演化所调校，这可能是现有大量规制研究未能直面的一个前提性问题。

第二，既有规制理论的四种进路各自握有重要的承重变量——架构、信息、程序、开放度——但它们共享“规制主体的自诊能力”这一未检预设。当规制对象产生的伤害恰好落在规制主体的认知盲区里时，规制不是在修补漏洞，而是可能在加固一套恰好遗漏了关键伤害的框架。这个可能性在现有文献中没有被充分讨论。

第三，侵权法体系自身的演化——风险分配对因果追溯的替代、效率转向对判断力的排挤——已经被法史学和法律社会学的部分研究所分别触及，但这些研究通常把它们当作不同子领域（侵权法教义学、法律职业研究、司法制度史）的内部议题。尚未有研究将这两次变动放在同一个分析框架下，追问它们对法律体系认知硬件的累积效应——即，这两次各自理性且有益的内部调整，是否在无意中合力削薄了法律体系面向新型伤害的认知接收能力。

本文正是在这三重缺口的交汇处进入讨论。

## 三、理论框架

### 3.1 核心概念的重建

本文使用“法律认知硬件”（legal cognitive apparatus）一词来指称法律追责结构中那些已经内化到不再被当成选择的操作预设——它们不是某部法律明文规定的条文，而是任何一部法律在操作时必须依托的前置条件。这个概念需与几个近似术语区分：“法律文化”过于弥散，不足以承载本文所需要的经验可检验性；“法律范式”过于宏大，且暗示体系性的一揽子转换；本文选取“认知硬件”，意在突出其有限可识别性、有限可修正性，以及“它不是坏掉了而是被调校到了一个特定频段”这一关键诊断。

法律认知硬件包含三个原子单元，三者共同构成法律体系接收伤害信号的基本门槛。第一个单元是“可识别的损害事件”：损害必须以一个有边界的事件形态进入法律视野——某人在某时某地遭受了某种具体的、可被叙述的侵害。第二个单元是“可追溯的因果链条”：从损害事件出发，存在一条可以被证据支持的追溯路径，将损害连接到某个可被识别的行为或疏忽。第三个单元是“可归因的主体意志”：链条的末端，存在一个可以进行过错评价或风险分配的行为主体——自然人、法人或其他可被法律拟制为责任承担者的实体。

这三个单元不是法学理论的外部分类，而是任何侵权诉讼在进入实体审理之前必须逐一通过的三道滤网：如果损害无法被叙述为一个有边界的事件，它甚至无法进入诉状；如果因果链条无法被证据跨越，即便损害事件清晰，诉讼也会在证据审查阶段被驳回；如果行为主体无法被识别，即便损害事件和因果链条都存在，追责

仍然没有落脚点。这三个单元在日常法律运行中已经内化到几乎“看不见”的程度——法律从业者之所以能够在接到案件时迅速判断“能不能打”，正是因为这套滤网已经在他们的认知中内化了。

## 3.2 不对称灵敏结构的生成逻辑

本文的核心命题是：这套三个原子单元构成的认知硬件，在长期演化中形成了一种“不对称灵敏结构”——它高度灵敏于短因果链、可个体化、可定价的伤害，却系统性迟钝于长因果链、分布性、不可个体归因的伤害。

灵敏侧的优化是侵权法数百年制度演化的核心成就：从过错责任到严格责任的扩展、从个人责任到企业责任再到市场份额责任的创造、从个案裁判到集团诉讼的制度发明——每一步都在扩大“可接收的伤害类型”的范围。但这一扩展有一个隐藏的方向性：它始终在沿着“可追溯性”和“可量化性”这两个维度扩展。污染致害变得可追溯了（环境侵权的因果关系推定）、缺陷产品变得可追溯了（产品责任的严格责任）、大规模侵权变得可处理了（集团诉讼和赔偿基金）——每一次扩展都是一次“可追溯范围的延伸”，而不是一次“不可追溯伤害的纳入”。换句话说，侵权法在变得更好的同时，变得更好的方向恰好是让它更擅长处理它本来就擅长处理的那一类伤害，而对于它本来就不擅长处理的那一类伤害，这种“更好”并没有改变基础状况。

迟钝侧则因三原子单元自身的门槛效应而天然形成：如果一段伤害是“分布”而非“事件”——它散布在数以百万计的个体经验中，每个人承受的只是总损害的一个极小切片，而其总和才构成真正有意义的伤害——那么它无法在第一道滤网就被拦截：没有一个个体能够拿起一个“有边界的事件”走入法院。如果这同一个伤害的因果链是“网络”而非“链条”——它由几十个行为者在不同时间、不同地点、以不同角色共同参与，而没有任何一个单一行为者的贡献达到了可以被独立归因的阈值——那么它在第二道滤网前就消散了。如果造成这种伤害的“源头”是一组分散的参与者（模型开发公司的算法工程师、训练数据标注团队的产品决策、用户的使用方式、平台的分发机制、广告商的投放策略……），那么它在第三道滤网前就没有“主体”可被指认。

此处需要澄清的一点是：这种不对称并非谁的有意设计，而是三原子单元与某一类伤害的本体特征之间的客观不匹配。它不是“法律被坏人操控了”的问题，而是“法律的接收硬件在认知上只适配了某一频段的伤害信号”的问题。这就好比一个收音机：不是它坏了，而是它的物理接收波段被设计在调频（FM）频段，调幅（AM）信号再强它也收不到。而真正棘手的是——在过去四十余年的两次历史变动中，这台收音机不但没有扩展波段，反而在不断优化“在FM频段内的接收精度”，并在此过程中逐步失去了意识到“AM频段存在”的能力。

## 3.3 自诊不可能与认知污染的机制

不对称灵敏结构本身已经构成困境，但使困境“锁死”的，是这一结构恰好处于法律体系自身的诊断盲区内。这一自诊不可能源于三个机制的协同作用。

第一，法律体系识别“自己漏了什么”的常规途径本身就依赖于它用来接收的同一套认知硬件。判例法的发展靠的是当事人把新类型的争端带入法院；成文法的修订靠的是立法者发现“现有的法律覆盖不到某些行为”；学术研究靠的是法学家观察到“现有理论解释不了某些现象”——这三条途径的入口，全都是由三原子单元把守的：当事人只能把她能叙述为“有边界事件”的损害带入法院；立法者只能接收到那些已被转化为“可被法律话语表达”的诉求；法学家只能在现有判例和法条中寻找失灵的证据。如果三原子单元把某一类伤害系统性地挡在了入口之外，那套指认出“这里有问题”的常规程序就无法侦测到它自身的失灵。

第二，法律服务商业化带来的效率转向，在法律从业者个体层面制造了对“现有分类装不下”这个感觉的衰减。一个律师在可计费小时制度下高效运作时，他的认知模式被训练为“快速将事实映射到现有法律分类”，而不是“停下来问这个事实是否超出了现有分类”。这不是道德问题——不这样做的人在经济上被惩罚；这是制度激励问题——制度奖励的是在分类内部高精度作业，而不是对分类本身提出质疑。日积月累，当足够多的个体从业者被足够长时间地训练为高精度但在分类边界上迟钝时，整个法律共同体的自诊能力就在整体层面下降，但每个人在个体层面都看不出自己有这个问题——因为他身边的人也这样，因为他被评价的方式从来不包含“识别现有框架外伤害”的能力。

第三，更重要的是，那些已经进入法律视野的AI相关案件，在处置过程中会产生一个系统性的错误反馈：它们给法律共同体一种“AI的法律问题正在被处理”的体感，而实际上被处理的那部分恰恰都是“刚好能被旧框架接住”的——格①——而真正在旧框架接收波段之外的那部分，从未在判例文本中出现过，因此也不构成“需要被处理”的体感。这个错误反馈不是任何人有意制造的，而是不对称灵敏结构的自然副产品：灵敏的那一侧不断地在发出“我们能处理”的信号，迟钝的那一侧则始终保持沉默——而一个只接收灵敏侧信号的系统，会天然地认为“大部分问题已被覆盖”。

这三个机制——入口把关、效率衰减、错误反馈——共同保证了不对称灵敏结构既不会被法律体系的常规自检程序发现，也不会被法律从业者的个体判断力捕捉，更不会被系统已有的“成功处理”记录质疑。这就是“自诊不可能”的含义：不是因为没有人愿意去诊断，而是用来诊断的那双眼睛本身就是由漏掉的那一部分塑造的。

### 3.4 命题的形式化表述

基于以上概念重建，本文展开以下四个核心命题：

命题一（不对称灵敏命题）：法律追责结构的认知硬件对伤害信号存在不对称灵敏性——在“可识别的伤害事件×可追溯的因果链条×可归因的主体意志”这一三维灵敏空间内，法律体系响应度极高；在此空间之外的伤害信号，即便具有真实且显著的社会损害，法律体系的接收度也存在系统性衰减，且衰减程度随伤害的分布性、不可归因性、非事件性增强而加剧。

命题二（历史削薄命题）：这一不对称灵敏结构并非AI技术冲击的瞬间产物，而是过去四十余年间侵权法体系两次内生演化的累积效应——风险分配逻辑对因果追溯的替代削弱了法律共同体追溯因果链末端的能力储备；效率转向对判断力反思空间的排挤削弱了法律共同体识别“现有框架无法吸收的伤害”的认知条件——两者合力将法律认知硬件的灵敏频段收窄至可追溯、可量化的伤害区间。

命题三（自我加固命题）：不对称灵敏结构一旦形成，即通过三条并联路径持续自我加固——能力极化让法律从业者越来越擅长处理灵敏侧的伤害但越来越看不见迟钝侧的伤害；商业化闭环将制度资源（立法注意力、司法资源、法律服务供给）全部吸附在可定价的伤害治理上；教育出厂校准让每一代新法律人在法学院入学第一年就内化了“法律能接收的伤害=需要被法律处理的伤害”这条默认基线，从而在认知层面消解了“还存在法律不接收的伤害”这个追问。三条路径各自正常运转、彼此并联锁死，推动法律体系走向一个制度性难以自我纠偏的退化方向。

命题四（微动作-大动作衔接命题）：修复这一困境不能依赖单一动作——既不能只做制度层面的基础设施搭建（行业基线透明、判断节点划定、训练过程审计），也不能只做教育入口的认知干预（在法学院阅读材料中植入盲区追问）——而必须让两者在时间轴上先后衔接：教育入口的微动作必须先行为系统植入不被旧认知框架污染的“追问抗体”，制度层面的大动作必须后发为追问抗体提供可操作的信息基础设施（公开的模型行为数据、标定的判断保留区、可审查的训练过程偏差报告）。两者脱节时，大动作的制度设计将被旧认知框架收编；两者衔接时，微动作种下的追问在第二代从业者中积累到足够密度，才可能在约二十年的时间尺度上推动制度的第二代修订。

## 四、方法论

### 4.1 研究的性质与边界

本文属于概念分析与制度演化推演相结合的理论建构性研究。研究目的不是给出关于“法律体系偏瘫”的终极答案，而是提出一套可被独立检验的分析框架和诊断假设。本文的核心主张可以通过后续实证研究被检验、被部分推翻或被全面修正——这种可证伪性不是本文的弱点，而是本文区别于“封闭理论叙事”的核心特征。



## 4.2 分析策略与概念资源的调用

本文的分析策略分为三个递进步骤。

第一步，解剖法律认知硬件的构成单元并推演其在生成式人工智能语境下的运作失效机制。这一步使用的核心方法是概念分析——对“损害事件”“因果链条”“主体意志”三个法律操作概念进行重新审视，识别它们在面对分布性伤害、网络式因果、分散参与者时的范畴性不匹配。这一步不需要经验材料，因为三原子单元本身是内嵌于侵权法运行逻辑中的形式条件——本文的工作只是将它们从“默认前提”的位置上提取到“可审视对象”的位置上。

第二步，追溯法律认知硬件的历史演化轨迹。这一步引入两次历史变动的经验描述——风险分配对因果追溯的替代（1970s-1990s）、效率转向对判断力的排挤（1990s至今）——并将其纳入统一的分析框架，推演二者对法律认知硬件的累积削弱效应。这一步的经验基础来自侵权法史和法律职业研究领域已有的研究积累，本文的贡献是将两条被分置研究的历史线索接合为同一因果链条的两个环节。

第三步，推演不对称灵敏结构的自我加固路径及其可逆性条件。这一步使用演化分析，沿三条路径（能力极化、商业化闭环、教育出厂校准）分别推演正向反馈机制，并在此基础上识别干预的可介入位点。本章的推演逻辑是：如果命题一、二、三成立，那么规制干预必须同时击中“改变认知基线”和“提供新信号”两个条件，且二者在时间上必须先后衔接。第五章将从这一逻辑中导出具体的认知重建路径。

## 4.3 跨学科类比的审慎运用

本文在论证中有限借助了认知科学中对“接收器灵敏度”和“适应盲区”的概念，以及组织行为学中“系统自诊能力退化”的分析。这些跨学科调用仅作类比之用，其有效边界需要在此明确。

“接收器灵敏度”的类比适用于描述法律认知硬件对伤害信号的筛选机制，但法律体系不是物理传感器——它的“灵敏度”是由制度设计、职业训练、认知惯性、利益格局等多重因素共同塑造的，不能简化为一个工程参数。本文使用这一类比仅为突出“它不是故意的偏瘫而是演化所致的调校”这一关键特征。

“系统自诊能力退化”的类比适用于分析法律体系为何无法识别自身的盲区，但法律体系的自诊能力与其他组织（如医疗机构、安全监管机构）存在一个重要的制度差异：法院的核心功能之一是裁决当事人之间的争端，而非主动发现社会中的伤害——侵权法体系本身就是被动启动的。这意味着“法律偏瘫”与“医院漏诊”虽然在结构上有可类比之处，但法律体系的被动性更强，其自诊能力的退化也更隐蔽。本文在使用这一类比时，始终保持对法律体系制度特征的自觉。

## 4.4 可证伪性框架的设计原则

本文在论证过程中刻意避免了两类封闭命题。一类是“任何反例都只是这一诊断的又一例证”式的自封句——这类句法使得命题免疫于任何反例，从而移出了可被理性讨论的范围。另一类是“只有一种干预方式”式的独断论——本文始终坚持给出两条以上的可能路径及其分化条件，并在第六章为每一条路径设置了可观测的证伪指标。

三套实证方案的设计原则是：每一套方案都包含一个明确的、可被独立研究者重复操作的研究设计，都设定了一个清楚陈述的核心假设，并且都指明了什么情况下该假设会被认为不成立。这不是说本文的命题一定是正确的——恰恰相反，本文作者欢迎后续研究通过实证数据来挑战、修正甚至推翻这些命题。本文的学术价值不在“正确”，而在“足够清晰以至于可以被证明是错误的”——这本身就是一项命题走出哲学圈、进入科学轨道的标志。

## 五、分析与讨论

### 5.1 法律认知硬件的解剖：三原子单元的逐一踏空

侵权法体系在日常运行中依赖三个认知前提——也可以称为追责结构的三个原子单元——它们通常不被法学界作为审视对象，因为它们已经内化到了“理所当然”的程度。要理解法律体系在生成式AI面前的失灵，不能直接跳到“缺什么规则”的结论，而必须先把这三个原子单元从默认位置提取出来，逐一审视它们在新型伤害面前发生了怎样的范畴性不匹配。

可识别的损害事件。一个人在街上被自动驾驶汽车撞伤，这是一个可被清晰叙述的事件——有具体的时间、地点、侵害方式、受害后果。法官、律师、保险公司可以在诉状的第一段就抓住“发生了什么”。现在设想另一种伤害：一个大型语言模型在十亿次对话中，对女性用户的回应系统性地短于对男性用户的回应，对某些少数民族裔的名字在生成文本中持续地、微妙地与负面语境共现。这十亿次对话中的每一次单独的短回应或微妙的负面共现，都不能独立构成“一个可诉的损害”——它们太微小，微小到在任何单次交互中都低于法律可感知的阈值。然而，这十亿次微小偏差的累积效应，是一个真实且严重的社会伤害：在话语层面持续强化刻板印象，在认知层面缓慢塑造公众对特定群体的预期，在机会分配层面累积地削弱某些群体在信息获取、职业发展、公共参与中的竞争力。伤害确实发生了，但它没有一个“事件”的形态——它是“分布”而非“事件”，是“气候”而非“天气”。侵权法的第一道滤网——可识别的损害事件——在此没有拦截到任何东西，不是因为它有漏洞，而是因为它根本不是为接收这种形态的伤害而设计的。

可追溯的因果链条。设想2026年某日，一位医学生在准备职业资格考试时，向一个通用AI助手询问了一种罕见疾病的诊断标准。AI助手给出了一个看似全面但混入了三条致命错误信息的回答——其中两条是已被更新指南纠正的旧标准，另一条是语言模型在“弥补信息缺口”时自动生成的一个合理但不存在的症状。医学生信任了这个回答并据此在后续的临床实习中做出了一个风险决策。当损害最终发生时，回溯因果链条会遇到什么：AI助手的开发者可以主张其模型已经过大规模评估且总体准确率符合行业标准；训练数据提供者可以主张其中包含的正规医学文献数量充足且无法逐一核验每一份数据的准确性；医学生所在的医学院可以主张其已告知学生“AI仅供参考”；而医学生本人成为唯一可被追责的主体，因为做出最终决策的是他。责任被向下挤压到信息链条最末端、权力最小、信息最不对称的那一个节点上——不是因为法官或立法者想这样做，而是在没有一条清晰因果链可以一步一步从损害倒推回源头的情况下，侵权法体系唯一能抓得住的就是“最后那个做了决定的人”。此处的悖论是：生成式AI造成危害的关键特征恰恰是因果的分散化——伤害的产生经过了模型架构选择、训练数据构成、微调策略、上下文窗口设计、提示词的多轮交互等十几个环节的协同参与；而侵权法在处置这一分散化因果时，唯一的操作选项是假装有一条可追溯的链条，然后把链条末端的那个人认定为“责任主体”。

可归因的主体意志。传统侵权法的归责核心是一个可以进行过错评价的行为主体。在大多数AI语境中，这个主体可以是模型开发者（如果存在设计缺陷）、运营者（如果违反安全保障义务）、用户（如果滥用或故意作恶）。但在生成式AI造成的若干种伤害中，存在一个“无行为主体”的空心层：没有人“决定”要造成这个伤害，它只是模型在被正常使用时按照其训练分布所涌现的一个统计副产品。模型开发公司决定的是整体架构和商业策略，没有决定“在这一次对话中给出错误医学建议”；训练数据标注团队决定的是标注指南和质量标准，没有决定“这个特定的偏差会在数十亿次交互后累积为一个群体的系统性劣势”；用户决定的是提出什么问题，没有决定模型如何回答。当伤害是由“所有人的正常行为+统计分布的副产物”构成时，侵权法在第三道滤网前找不到一个可以承受过错评价或风险分配的“主体”——因为它被设计为处理“有人做了某事导致损害”的世界，而生成式AI的某些伤害属于“没有任何人决定做那件事，但它发生了”的世界。

这三个原子单元的逐一踏空，合起来就是法律认知硬件的“频段错位”：法律接收器只被校准为接收“事件—链条—主体”这一频段的伤害信号；而生成式AI结出的相当一部分实害，落在“分布—网络—分散参与者”的频段里。信号确实存在，强度甚至可能在累积后非常可观，但接收器没有为这个频段开设接收频道。



## 5.2 历史削薄：两次内生演化如何收窄了法律的接收频段

上一节的分析可能会导致一个误解：法律认知硬件的频段错位是生成式AI出现之后才突然暴露的，AI是罪魁祸首。这一理解不准确。三原子单元作为法律认知硬件的基本结构，早在AI出现之前就已经在运转。关键变化不是AI“打碎”了一个完好的接收器，而是过去四十余年间的两次内生演化已经将接收器的灵敏度缓慢削薄到了临界点附近——AI只是越过临界值的最后一推。

第一次演化：风险分配对因果追溯的替代（1970s-1990s）。大规模工业化伤害的出现使得传统个案式因果追溯的信息成本急剧升高。环境污染的受害者很难证明“我的癌症是由这家化工厂在过去三十年排放的某一种特定化学物质造成的，并且排除了其他所有可能原因”。产品缺陷的受害者很难证明“我的伤害一定是在生产环节而非流通环节或使用环节中发生的”。法院和立法者发展出了一整套风险分配工具来应对这一困境：无过错责任降低了过错要件，市场份额责任绕开了因果识别的困境，举证责任倒置把证明负担从受害者转移到更有信息优势的加害者一方，保险与赔偿基金则将个案的追溯转化为统计上的风险池。这套工具极其有效——它们使得大规模工业化社会中那些无法逐一追溯的伤害得以被社会化地消化，避免了侵权法体系在信息门槛前崩溃。

然而，这套工具在法律共同体内部造成了一个未被注意到的副产品：法律从业者“从损害追溯到行为”的认知肌肉，在风险分配框架下不再被日常训练。当一个律师的工作重心从“证明是谁怎么造成的”转向“计算责任如何分配”时，他处理伤害的方式发生了一个微妙的转变——他不再问“这个损害是怎么发生的、每一步因果链可以追溯到谁”，而是问“这个损害在现有的分配框架下应该由哪一方承担多少”。前者是追溯性思维，后者是分配性思维。在风险分配框架成熟运转的领域（产品责任、环境污染、职业伤害），这种转变是理性的，甚至是进步的——它让更多受害者获得了赔偿，而无需穿越不可能的信息门槛。但当一类新的伤害——如生成式AI的分布性伤害——的核心难点恰好在于“如果不追溯因果链就无法识别伤害本身”时，法律共同体那条已经萎缩的追溯肌肉就暴露了：不是它不努力，而是它被四十年的制度激励训练成了“用分配框架去套追溯难题”的惯性，而有些追溯难题不是分配框架能套得住的。

第二次演化：效率转向对判断力的排挤（1990s至今）。法律服务商业化的加深——律所的企业化经营、全球性的可计费小时制度扩散、法律服务市场竞争的加剧——共同创造了这样一种激励结构：法律从业者被奖励的是“在既定框架内高效处理案件”的能力，而不是“停下来怀疑既定框架本身”的能力。可计费小时制度把律师的时间切割为可被记录和收费的六分钟单元；每一个六分钟单元都必须对应一项“可交付的法律服务”——审阅文件、起草意见、会见客户、出庭应诉。“停下来想一想这个案件是否暴露了现有法律框架的一个盲区”——这不计入可收费小时。没有人禁止律师做这种思考，但在日复一日、年复一年被六分钟单元格式化的职业生涯中，不产生可计费时间的认知活动会自然而然地被压缩到几乎消失的程度。

这种“判断力”的萎缩不是道德问题，是制度生态问题。在一片只奖励“高效分类填表”而不奖励“质疑表格本身”的制度土壤中，即便有少数具有反思能力的法律从业者，他们的反思也无法转化为制度层面的认知更新——因为制度听不到这些反思，制度的听力恰好被同样的效率激励所格式化。当一个合伙人试图跟客户讨论“这个案件可能意味着现有法律框架漏掉了某类伤害”时，客户（受过同样效率逻辑训练的资深法务）会回应：“这个判断很好，但我们先解决眼前这个具体案件。”效率转向不阻止任何人思考，但它确保这些思考无法积累——它们被一个接一个的六分钟单元冲散，无法形成可以改变制度认知的持续压力。

两次演化的累积效应。风险分配替代因果追溯，收窄了法律从业者“追溯因果链”的操作范围；效率转向排挤判断力反思空间，收窄了法律从业者“识别框架外伤害”的认知条件。这两次演化发生在不同的领域、由不同的动因驱动、各自有充分的内在合理性。但当生成式AI开始大规模产生那些同时需要“追溯因果链”和“识别框架外伤害”的新型伤害时，法律体系突然发现：它用来处置这些伤害的两条核心认知肌肉，在过去四十年里被各自的制度激励同时削弱了。这不是一场突然的灾难，而是两个缓慢的演化过程，在AI越过临界值的那一刻发生了危险的共振。

### 5.3 自我加固的三条路径：为什么偏瘫不会自愈

如果法律体系只是暂时被AI冲击打蒙了，假以时日它应该能自行调整。但本文的诊断认为，不对称灵敏结构一旦形成，即通过三条并联路径持续自我加固，使得体系难以通过常规自纠机制恢复。

路径一：能力极化。法律从业者在可追溯伤害的处置上越熟练，就越不需要去触碰不可追溯的伤害——因为前者占据了他们全部的可收费时间和专业注意力。一个擅长大规模侵权集团诉讼的律师团队，其专业能力的每一部分都被优化来处理“可个体化、可定价、可分配”的损害。他们在识别受害者群体、计算损害总额、谈判和解方案、分配赔偿金等方面极为高效。但生成式AI造成的语言生态降级、判断力代际衰减、认知多元性收缩——这些伤害没有一个可以被“计算总额”，因为它们没有一个可定价的损害单位。能力极化不是这个团队“变坏了”或“变得懒惰”——恰恰相反，他们在自己擅长的领域越来越好，而这种“越来越好”本身就意味着他们的专业注意力被越来越牢固地吸附在灵敏侧，迟钝侧的伤害信号在他们的专业感知范围里变得越来越微弱。能力极化的悲剧在于：那些最有可能推动制度变革的人——即法律体系内最优秀、最有能力的从业者——恰恰是那些最不可能感知到变革必要性的人，因为他们的全部能力都是在旧框架内被训练、被评价、被奖励的。

路径二：商业化闭环。侵权法的制度资源——立法者的注意力、法院的审判资源、律师的法律服务供给、法学研究的选题分布——全部被吸附在可定价的伤害治理上。这不是阴谋论，而是价格信号的自然分配：一个能产生数百万赔偿金额的案件，天然地吸引律师接案、法院立案、学者撰文讨论；一个没有明确损害赔偿额的分布性伤害，天然地难以进入任何一方的资源分配优先级。更关键的是，当一个法律体系已经习惯了通过“赔偿金额”来衡量伤害的重要性时，那些无法被定价的伤害不但在经济上不可见，而且在认知上也变得不可见——因为它们没有一个数字可以被放到“这个伤害有多严重”的比较里。商业化闭环确保法律体系的全部制度资源都在灵敏侧伤害的治理上高效循环，而迟钝侧伤害的治理资源始终趋近于零——不是因为资源不够，而是因为现有价格信号的引导下，资源根本流不到那一侧。

路径三：教育出厂校准。每一代新法律人的认知基线不是在进入职场后才被塑造的，而是在法学院的第一年甚至第一个学期就被校好了。侵权法主干课程的经典判例——无论是英美法的Palsgraf v. Long Island Railroad、Donoghue v. Stevenson，还是中国侵权法教学中的标志性案例——它们无一例外都是“可识别的损害事件×可追溯的因果链条×可归因的主体意志”的典范。学生通过反复研读这些案例被训练成精于在这套三维框架内进行精细的法律推理，但他们从入学第一天起就默认了一件事：“法律要处理的伤害，就是能够被这套三维框架接收的伤害。”这不是任何一位教授有意“洗脑”的结果，而是课程设置本身的隐性预设：侵权法课程不会在第一节课说“以下是我们无法处理的伤害类型”——它只会教授“以下是我们能处理的伤害类型及其处理方式”。学生学得越认真、成绩越好，这个默认预设就越深入地植入他们的认知底层。当他们十五年后成为法官、合伙人、监管官员、立法顾问时，他们在面对一个新型伤害时的第一反应不是“法律有没有为这个伤害开设接收频道”，而是“这个伤害可以被现有哪个分类吸收”——后者是在法学院被反复训练的能力，前者从未被训练过。

三条路径的并联锁死。能力极化确保现有制度资源被锁定在灵敏侧；商业化闭环确保没有新的制度资源流入迟钝侧；教育出厂校准确保新一代从业者在进入系统之前，就已经被校好了“迟钝侧不存在”的认知基线。三条路径各自正常运转，彼此互相提供条件——能力极化为商业化闭环提供了高效的服务供给，商业化闭环为教育出厂校准提供了“学这些就够了”的就业市场信号，教育出厂校准则为能力极化提供了不会被“框架外追问”干扰的新鲜人力。这不是一个可以被任何单一干预点撬动的结构——三条路径中的任何一条单独被撬动，都会被另外两条的惯性拉回来。这就是为什么仅仅“加强立法”无法解决困境：立法加在灵敏侧对应的是商业化闭环的吸附（法条越多、可做的合规业务越多、专业注意力的流向越集中在灵敏侧），立法加在迟钝侧面临的是能力极化的无处下手（没有足够多的法律从业者具备处理这类伤害的专业能力）和教育出厂校准的认知过滤（新一代法律人根本没有被训练去识别这类伤害）。

## 5.4 二十年认知重建螺旋：微动作与大动作的先后衔接

如果三重自我加固的锁死程度如本文所诊断的那样深，那么任何单一维度的干预都将被锁回原状。但这并不意味着修复不可能——只是它不能以“一步到位”的方式完成，而必须分两步走，且两步之间的时间间隔可能在十年到二十年之间。

第一步（微动作，约2028年前启动）：在教育出厂校准之前撬开一道认知裂缝。法学院侵权法、行政法、科技法的主干课程中，经典判例的阅读材料之后，附加一份“漏掉清单”——同一事件中被法律系统漏掉的伤害信号。这不是附加一门新课，而是在现有教学框架内嵌入一个微小的追问：“本次判例中，法律接住了什么，法律没有接住什么？”这个追问本身不提供答案，它只是确保学生在被三原子单元的框架完全校准之前，至少瞥见一眼法律没接住的东西。

这个微动作的设计关键不在于“它能否立刻改变法律体系”——它显然不能。它的功能在于未来的时间点上：十五年后，当这些学生成为法官、合伙人、监管官员、立法顾问时，他们中的一部分人，在面对一个现有框架无法吸收的伤害时，不会像他们的前辈那样默认“那就不是法律的事”，而会有一个被埋藏在法学院阅读作业里的追问在提醒他们：“等一下，这可能正好是法律当年接不住的那一类。”这个追问不是抗体本身——它只是一个让抗体可能生长的认知裂缝。

微动作的必要性在于：如果没有这一步，大动作的制度设计者将全部是已经被旧认知框架校准的人——他们会让行业基线透明标准覆盖所有“可量化、可追溯、可定价”的风险，而恰好漏掉不可量化、不可追溯、不可定价的那些；而他们漏掉的时候，并不知道自己漏掉了，因为他们的认知硬件里没有“格④”这个接收频道。

第二步（大动作，约2030-2035年启动）：在追问抗体形成一定规模后建设信息基础设施。行业基线透明（模型行为跟踪数据的公开标准）、判断节点划定（专业领域不可外包给AI的判断清单）、训练过程审计（独立第三方对训练数据偏差和模型行为模式的审查）——这三项制度工具的效力不取决于它们被写进法律文本，而取决于是否有足够密度的从业者能够使用它们。如果行业基线透明标准规定了“模型对特定群体的系统性偏差必须披露”，但法律从业者群体中没有人看得懂这份披露意味着什么——他们只关心“这份披露对我正在做的案件有没有可用的信息”，而在没有追问抗体的情况下，他们只会从中提取“可追溯的伤害”的信息，漏掉“分布性伤害”的信息——那么披露本身无法转化为法律体系的接收。

大动作必须在微动作之后而不是之前，是因为大动作的制度设计过程本身就会被旧认知框架污染。谁来起草行业基线透明标准？在一份法律共同体尚未被追问抗体渗透的时间点（比如2025年），起草者的全部专业培训都是在“可追溯伤害”的框架内进行的。他们会被要求“制定AI行业透明度标准”——他们会做得很认真，会参考国际经验，会征求各方意见。但他们起草出来的标准，将有极大概率把“应当披露的风险信息”全部定义为可追溯、可量化、可定价的风险类型，而对分布性、不可归因、缓慢累积的伤害，要么以“不可操作”为由排除，要么以“暂不具备条件”搁置，要么以一句“鼓励企业披露其他相关信息”的软条款敷衍过去。他们不是坏人或懒人——他们只是看不见他们没被训练去看的东西。

而如果微动作已经在2028年前在法学院中种下了足够密度的追问抗体，那么到2030-2035年大动作启动时，进入标准制定工作组、立法咨询委员会、监管机构法律部门的新一代从业者中，会有一定比例的人天然带着“法律漏了什么”的追问。当标准草案在讨论“应当纳入哪些风险披露指标”时，会有人问：“那分布性伤害呢——比如这个模型在上亿次交互中对特定群体造成的累积性话语劣势？”这个提问本身不保证一定被采纳——它可能被产业界以“不可操作”否决，可能被老一代从业者以“这不是法律应该管的”驳回。但它被提出的那一刻，就已经完成了一项关键工作：它在制度设计的论证记录中种下了一个“格④曾被提及但被搁置”的事实。这个事实本身，将成为十五年后的第二代修订启动时的制度记忆——届时，那些当年在论证记录里被搁置的条款，可以被重新翻开，有人可以举证说：“这些伤害在十五年前就已经被识别了，只是当时我们没有能力处理它。”

两者衔接与脱节的分叉。如果微动作和大动作成功衔接——微动作在2028年前被至少十所法学院在主干课程中部分纳入，大动作在2030-2035年间在至少一个主要司法管辖区的标准制定论证记录中出现格④维度的辩论且部分被纳入——那么到2040年代，进入法律体系中层决策位置的第一批携带追问抗体的从业者，将有足够



的信息基础设施（公开的模型行为数据、标定的判断保留区、可审查的训练过程偏差报告）将追问转化为实质的制度修订。这个全过程从2025年起算，约二十年。

如果两者脱节——微动作启动了但采纳规模太小、追问抗体密度不足；或大动作被污染、做出来的标准全部落在格①③——那么最可能的演化路径是：分叉点一（2028-2032年间第一批AI大案判词是否有法官或异议意见提及系统性伤害不可追溯性）被错过；分叉点二（独立伤害报告机构是否被主流法学引用）在商业化闭环的吸附效应下边缘化；分叉点三（法学院是否在必修课程中列入盲区教学单元）在十五年的慢速积累后才部分启动——迟至2040年代，格④的伤害才被少数法学院在选修课层面有限纳入。这一“部分脱节”的路径下，真正有效的第二代修订可能要到2050年代，而届时生成式AI的某些分布性伤害（语言生态的降级、判断力的代际衰减）可能已经越过了某些难以逆转的临界点。

## 5.5 反驳预期与回应

一个必须被正面处理的反驳是：本文的诊断犯了“过度诊断”的毛病——把法律体系正常的有限管辖权说成是“认知偏瘫”。侵权法从来不负责处理所有社会伤害，有些伤害应该留在道德谴责、社会规范、私人协商或市场调节的领域。生成式AI带来的语言生态降级、判断力代际衰减——这些是真实的伤害吗？如果是，它们中哪些应该进入侵权法的追责范围、哪些不应该？

这一反驳具有实质分量，它触及了本文整套判断尚未完成的一块拼图：法律“应该”不接收哪些伤害的规范性判准。本文在前五节所做的主要是描述性与诊断性工作——法律体系实际上不接收什么、为什么不接收、不接收的结构如何自我加固。但“实际上不接收”不等于“不应该接收”——这个区分必须被严肃对待。

本文对此的初步回应是三个层次，但承认这只是一个开端而非完成。第一，实害的真实性判准：在判断法律“应该”接收某种伤害之前，必须先回答该伤害是否真实存在且具有社会层面的损害后果。算法歧视造成的就业机会丧失可以通过对照实验证实，深度伪造造成的名誉侵害可以通过个案举证，但语言生态降级和判断力代际衰减的真实性和严重性，目前更多停留在理论推断层面——本文第六章提出的经验研究方案（判例编码、认知污染追踪、教学干预准实验）正是为了把这一诊断从“可能”推向“可测量”。在法律能够回应一个伤害之前，这个伤害首先必须被从“感觉”转化为“可被公共论证的事实”——而这正是本文全部方法论努力的方向。

第二，追责的必要性判准：即便某个伤害是真实的，也不等于法律应该追责。法律的追责边界取决于多个因素的权衡——追责的信息成本、追责对行为激励的扭曲风险、追责对司法资源的占用、替代治理机制（社会规范、行业自律、技术标准）的有效性。本文并不主张“法律应该吸收所有格④伤害”，而是主张“法律体系应该具备接收格④伤害信号的能力”——这是两个完全不同的命题。一个具备接收能力的法律体系，在感知到某种分布性伤害后，完全可以在规范层面决定“这类伤害不适合以侵权追责方式处理，更适合以行业标准、公众教育或技术设计规范来治理”。但关键是这个决定必须是有意识的规范选择，而不是因为认知硬件根本没有收到信号而导致的静默。

第三，不可逆伤害的底线判准：即便法律在绝大多数情况下选择不追责分布性伤害，这也不等于在所有情况下都无权介入。存在一个底线问题：如果某类分布性伤害的累积效应越过了某一人道底线——比如它导致某一代人的整体判断力发生了不可逆的衰减，且这种衰减在代际传递中自我放大——那么即便它从未成为一个“事件”、即便它从未拥有一条“链条”、即便它从未归给一个“主体”，法律体系是否仍然有义务创造一种新的追责形式来回应它？这一追问超越了当前侵权法教义学的范畴，进入了法哲学的规范建构领域。本文在此只提出问题，不假装有答案——但认为这个问题必须被提出，因为它是“法律体系在技术加速条件下的自我重建”这个研究方向的最深层追问。

## 5.6 综合讨论：从诊断到研究纲领

将本章各节的论证放回一起，可以看到一个不同于常规“AI法律规制”讨论的认知图景。

常规讨论的出发点是“技术造成了哪些问题，法律缺失了哪些规则”，分析终点是“应该制定哪些新规则”。这个思路在操作层面有其价值，但它反复撞到同一堵墙：新规则不断出台，但规制效果并未等比例提

升；最棘手的那一类伤害——分布性、不可归因、非事件性的伤害——始终游离在法律体系的回应半径之外。本文的贡献不在于提供另一套新规则建议，而在于追问为什么新规则撞上了这堵墙——答案是：这堵墙不仅是制度设计的缺失，更是法律认知硬件本身的频段限制。规则是加在接收器上的软件更新，但只要接收器的物理频率范围不变，有些信号就是永远收不到的。

这一定位引向的是一个比“如何规制AI”更大的问题群——法律体系在技术加速条件下的认知退化与自我重建。这个问题群不是法学的内部修补，它同时拉到认知科学（法律认知硬件的形成与退化机制）、组织行为学（系统自诊能力的制度化条件）、科技政策（信息基础设施如何改变制度认知）、教育社会学（法学院训练如何塑造代际认知基线）等至少四个学科。它的核心问题群清晰可辨：法律认知硬件的演化机制是什么？它的退化路径有哪些？退化的可逆条件是什么？认知重建如何在教育入口和制度入口同步启动？不同法律体系（成文法传统、判例法传统、混合体系）下认知退化的差异演化机因是什么？这些问题不是任何一个学科可以独自回答的——它们要求法学研究者打开自己的学科边界，与认知科学、组织研究、科技政策、教育学的同行协作。而本文第五章提出的“微动作-大动作二十年认知重建螺旋”，可以被视作这一跨学科研究纲领的第一个可操作的理论模型——它同时给出了时间窗口、衔接条件、分叉路径和检验指标。

## 六、可证伪性框架：三套检验方案

以上论断如果不被转化为可被独立检验的假设，就只是另一套深刻的哲学叙事——而“深刻的哲学叙事”恰恰是本文所批评的“自诊不可能的封闭圈”的最典型产物。因此，本章为本文的核心诊断配备三套独立的实证检验方案。如果这三套方案中任何一套的检验结果证实了本文假设之外的可能，本文的核心命题就需要被修正——这种自我暴露于被证伪风险的设计，是本文区别于封闭理论的关键。

### 6.1 检验方案一：判例文本的盲区编码分析

待检假设：在法律体系实际处置的AI相关案件中，灵敏侧伤害（可追溯、可个体化、可定价）占判例总案量和立法提案总量的绝大多数（预测超过85%），而迟钝侧伤害（分布性、不可归因、非事件性）在判例文本中被独立识别为伤害的概率极低（预测低于5%）。

竞争解释：存在另一个可能——“法律体系并非不对称迟钝，而是因为信息不对称：格④伤害的受害者没有足够的法律资源或信息渠道将损害带入法院；一旦他们获得足够的法律资源和信息支持，法院完全有能力识别并处理这些伤害。”如果这一竞争解释成立，则本文“法律认知硬件频段错位”的诊断需要被修正为“法律服务的可及性不足”——二者的干预方向截然不同：前者需要认知硬件的重新刻度，后者需要法律援助和公益诉讼的投入。

研究设计：选取2015年至2030年间与AI相关伤害有实质关联的200宗以上已决判例（来源：中国裁判文书网、Westlaw、EUR-Lex等公开数据库），覆盖至少三个司法管辖区（中国、美国、欧盟）。对每宗判例进行双轴编码：横轴为“灵敏频段”——伤害类型归属格①（灵敏侧高可见）、格②③（中间区）、或格④（迟钝侧）；纵轴为“构成性可见度”——该伤害信号在法院判决理由或律师诉讼策略中是否被独立识别为需要法律处置的对象。编码工作由两名以上编码者独立完成并进行信度检验。

控制掉竞争解释的关键设计：为了区分“认知偏瘫”与“可及性不足”，编码中需要额外标注每宗案件的法律服务配置情况——原告是否有律师代理（以及律师的专业领域和经验年限）、案件是否获得公益法律组织的支持、是否有第三方机构出具过本案涉及的伤害类型的研究报告。如果一项格④伤害在这些“信息支持与法律资源”都已具备的案件中仍然未被法院独立识别，那么“可及性不足”就难以解释这一结果；如果格④伤害的识别率在法律资源充足的案件中显著上升，那么竞争解释成立，本文诊断需被修正。

证伪条件：如果编码结果显示格④伤害在判例文本中被独立识别的比例显著高于5%（如超过15%），则本文“法律体系不对称迟钝”的核心诊断不能成立——法律体系的认知硬件并非无法接收格④伤害信号，只是格④伤害在案件总量中可能确实较少。此时本文诊断的修正方向为：法律规制的重点应当在“如何让更多的格④伤害信息进入法律程序”，而非“如何改变法律认知硬件本身”。

## 6.2 检验方案二：行业标准制定过程的认知污染追踪

待检假设：在AI行业基线透明度标准的制定过程中，分布性、不可归因的伤害指标在三个关键环节（指标设计、合法性论证、实施执行）中的覆盖率和实质性讨论密度均极低，且这一低覆盖率不是被产业界游说否决的结果，而是在标准起草者的初始草案中就已经不存在——因为起草者的认知框架中不包含格④伤害的接收器。

竞争解释：“格④伤害之所以未被纳入行业标准，不是因为认知偏瘫，而是因为在当前技术条件下确实不具备可操作性——没有可量化的指标、没有成熟的方法论、缺乏行业共识——这些是合理的、务实的技术判断，而非认知污染。等到技术条件成熟、方法论发展起来，行业标准自然会将其纳入。”如果这一竞争解释成立，则本文的“认知污染”诊断需要被“合理的技术渐进主义”所取代——干预方向从“教育入口撬裂缝”变为“投资技术方法论研发”。

研究设计：选取2025年至2035年间，至少三个主要司法管辖区（中国、欧盟、美国）与AI行业基线透明度相关的立法或标准制定过程。收集三重文本：标准制定过程的公开文件（草案、修改说明、听证记录、公众意见征询回复、各阶段的技术报告）、最终颁布的标准文本、以及标准实施首三年内首批被规制企业的公开披露文件。对全部文本进行三环节编码：指标设计环节——各阶段草案所列的“应披露信息”类型中，格④伤害指标的条数与占比；合法性论证环节——听证会、意见征询、工作组成员辩论记录中，是否有人提出格④伤害未被覆盖的论点及该论点的处理结果（采纳/讨论后搁置/未获回应/直接被否决）；实施环节——最终标准实施后，首批企业披露文件中格④维度信息的实际出现率。

控制掉竞争解释的关键设计：对“被以技术不成熟为由搁置或否决”的格④指标提议进行后续追踪——询问是否在后续版本中重新被提出、是否被列入“下期研究计划”、是否有外部研究者在标准颁布后发布了可操作的测量方案。如果一项格④指标在搁置之后，即便出现了可行的测量方案也始终未被标准纳入，则“技术渐进主义”解释难以成立；如果一旦出现可行方案即被纳入，则竞争解释成立，本文诊断需被修正。

证伪条件：如果在至少一个主要司法管辖区的行业基线透明度标准制定过程中，格④维度的伤害指标被实质性纳入且覆盖率达到实质门槛（如在应披露信息类型中占比超过10%且伴有可操作测量方法），则本文“大动作会被旧认知框架污染”的判断需要被修正或放弃——制度设计确实可以在不具备追问抗体前置条件的情况下，独立地、有效地打开对格④伤害的接收。

## 6.3 检验方案三：法学院“漏掉清单”教学干预的准实验

待检假设：在法学院主干课程中系统地植入“法律漏掉了什么”的盲区追问（以“漏掉清单”形式），能够显著提升学生和法律新从业者识别不可追溯的分布性伤害的能力，并且这种能力在职业生涯中能够持续——但它的制度影响需要制度裂缝同时打开才能完全发挥，单独的教育干预被制度惯性部分压回。

竞争解释：“法律教育的任何改革都是慢变量，与其在法学院课堂上做微调，不如直接改革司法考试和律师职业资格制度——这些高利害评价标准的改变，远比课堂教学材料的微调更能塑造从业者的认知基线。”这一立场的隐含预设是：认知基线的改变必须通过外部激励（考试、职业准入）来驱动，而非通过教学材料的内部微调——后者只是“锦上添花”而无实质效果。

研究设计：在至少三所法学院（理想情况：中国两所加另一辖区一所）的侵权法和行政法课程中，设置干预组（在经典判例阅读材料后附加“漏掉清单”并在课堂讨论中预留一次盲区追问环节）与对照组（常规教学）。学期初对所有学生进行“法律未接收伤害识别能力”基线测试，学期末做同样测试。对参与干预的学生进行第五年和第十年的职业追踪——深度访谈或问卷：“过去五年中，你在工作中是否遇到过某种AI造成的伤害，你觉得它构成了真实问题但现有法律工具无法直接处理？如果有，你是否提出过？被怎么回应？”将追踪结果与同期法院判例文本的格④独立识别率（来自方案一的年度更新数据）做比对。这一设计得到的不是统计显著性的p值，而是可被独立引用的职业判断变迁描述——这本身就是法学教育研究领域可发表的实证成果。

控制掉竞争解释的关键设计：在研究设计中增加第三组——“司法考试驱动组”——接受常规教学但额外提供针对司法考试中“AI法律问题”的应试训练。如果此组学生在第五年和第十年的职业判断追踪中，与干预组在格④伤害识别能力上表现相近或更优，则竞争解释成立——认知基线的改变确实主要通过高利害评价来驱



动，课堂教学的微调效果有限。如果干预组在识别能力上显著更优但同时在“提出后被制度压回”的报告率上更高，则本文的“认知裂缝必须与制度裂缝同时打开”判断被支持。

证伪条件：如果第十年的追踪数据显示，干预组毕业生在职业中自动放弃了格④追问，且不是因为外部压力而是因为“那个追问本身就慢慢显得不重要了”，那么本文过分重视格④伤害诊断的前提就需要被审视——那些被认为“法律漏掉了”的伤害，可能在法律从业者的实际职业经验中确实不构成需要法律回应的紧迫问题。这一证伪结果将导致本文诊断的重大修正：法律对格④伤害的“不接收”，可能不完全是不对称偏瘫的病理表现，而是法律体系在长期演化中形成的合理的自我边界——这个可能性本身，就是本文留给法哲学家的一个必须作答的规范性问题。

## 七、结论

### 7.1 核心发现的凝缩

本文通过概念分析与制度演化推演，提出了一个不同于“技术-法律落差论”的诊断：生成式人工智能的法律风险规制困境，根源不在规则缺乏，而在法律追责结构的认知硬件——三原子单元（可识别的损害事件、可追溯的因果链条、可归因的主体意志）——在面对分布性、网络式、分散参与者的新型伤害时发生的系统性频段错位。这种错位并非AI技术的突然冲击所致，而是过去四十余年间侵权法体系两次内生演化——风险分配对因果追溯的替代、效率转向对判断力的排挤——缓慢削薄的结果。这一不对称灵敏结构一旦形成，即通过能力极化、商业化闭环、教育出厂校准三条路径持续自我加固，使得法律体系难以通过常规自纠机制恢复对格④伤害的接收能力。

规制困境的修复不能依赖单一动作。本文提出的“微动作-大动作二十认知重建螺旋”主张：必须在法学院教育入口先行植入“法律漏掉了什么”的追问抗体（微动作），然后再在制度层面建设行业基线透明、判断节点划定、训练过程审计等信息基础设施（大动作）——并且大动作的制度设计必须由携带追问抗体的新一代从业者来推动，否则将被旧认知框架污染，做出一套只在可追溯伤害上运转良好、却恰好漏掉关键伤害的规制体系。

### 7.2 理论与实践的贡献

本文的理论贡献在三个层面。第一，将生成式AI法律风险的讨论从“缺什么规则”的技术性补充，推进到“法律认知硬件是否需要重新刻度”的制度认知层面，为法理学与认知科学、科技政策的交叉研究提供了一个新的分析框架。第二，识别了既有科技法规制理论（架构论、信息透明度论、程序完整性论、系统可规制性论）共享的未检预设——规制主体的自诊能力——并论证了这一预设生成式AI语境下的失效机制。第三，通过将侵权法史的两条独立研究线索（风险分配的兴起、法律服务商业化）接合为同一因果链条的双环节，为法律认知硬件的演化分析提供了可被后续研究沿用的历史叙事框架。

在实践层面，本文虽未提供“一步到位”的规制方案——这本身就是在本文诊断下不可能也不应该被提出的——但提出了四个可被独立行动者立即启动的具体操作：法学教师可以跨校联合编撰并开源发布“漏掉清单”教学材料，无需任何行政审批或课程改革；法学研究者可以启动判例文本的盲区编码分析，产出首份年度《判例盲区报告》，为立法和司法实践提供外部认知参照；科技政策研究者可以启动行业标准制定过程的认知污染追踪，为“大动作的制度设计是否被污染”提供实证依据；法律职业监管者可以研究专业领域“不可外包判断节点”的具体标准——医学诊断、法律意见、工程审查中哪些判断必须保留为人工作业、如何审计这些节点的实际执行情况。

### 7.3 研究局限

本文的所有核心论断均建立在概念推演和历史叙事之上，未经系统的经验数据检验。第六章虽给出了三套可操作的实证方案，但这些方案本身尚未被执行——本文提供的是一套可被检验的假设体系，而非已被检验的事实陈述。

本文的分析以中国、欧盟、美国三个主要司法管辖区的侵权法传统为隐含参照系，但未对不同法律体系（成文法vs判例法、强监管vs弱监管、集中司法审查vs分散司法审查）下认知退化机制的差异做精细推演。跨国比较的制度差异可能需要独立的研究项目来完成。

本文对“格④伤害”的判定标准——什么是“分布性”伤害、什么程度算“不可归因”——在概念层面给出了定义，但在操作层面尚未建立可供编码者一致使用的正式分类系统。判例编码方案的正式启动，需要法学家与计算方法研究者协作开发更精细的编码手册和信度检验标准。

最后，本文虽然在第五章提及了“法律应该不接收哪些伤害”这一规范性问题，但未正面建构回答——这一规范性框架的缺失，是整套判断尚未完成的一块重要拼图。法哲学家可以从本文的描述性诊断出发，建构“格④伤害的入门门槛”这一正面理论。

## 7.4 未来研究方向

本文从单一议题中涌现出一个有自生长力的跨学科研究方向——“法律体系在技术加速条件下的认知退化与自我重建”——以下是五个可供独立学者分别取走的研究模块。

第一，判例文本的盲区编码与年度报告。任何一位法律学者或计算法学研究者都可以用本文第六章方案一的框架，选取一个司法管辖区的AI相关判例，启动首轮判例盲区编码并撰写年度报告。后续研究者可以用不同辖区的数据进行跨国比较，检验“不同法律传统下法律认知硬件的不对称灵敏程度是否存在差异”。

第二，行业标准制定过程的认知污染追踪。一位科技政策研究者可以选取一个已颁布的AI行业基线透明度标准，回溯其制定全过程，对起草文件、听证记录、最终条款进行格①至格④覆盖率分析，并访谈标准制定参与者，确认格④伤害在哪个环节被过滤——这一研究可以直接检验本文“大动作会被旧认知框架污染”的核心假设。

第三，法律从业者判断力衰减的计量经济学检验。一位法经济学研究者可以设计研究，检验“可计费小时制度与法律从业者识别框架外伤害能力”之间的关联。可操作的设计包括：对比按小时计费与按固定费用或风险代理收费的律师，在提出“现有法律框架无法覆盖”这一论点时的频率差异；或追踪一家律所从固定收费转向小时计费前后，其律师在案件中提出“新类型损害”论证的频次变化。

第四，法学院漏掉清单教学干预的多校准实验与十年追踪。一位法学教育研究者可以联合三所以上法学院，正式启动本文第六章方案三所述的教学干预实验，并在五年和十年后进行毕业生职业判断追踪。这一研究的长期学术价值不依赖统计显著性——即便样本量有限、外部效度受限，它产生的“可被独立引用的职业判断变迁描述”本身，就足以成为法学教育研究领域的重要成果。

第五，格④伤害入法的规范性判准建构。一位法哲学家可以以前六章的描述性诊断为起点，正面建构一个三层次的规范性框架：什么情况下法律不接收某一伤害属于“认知偏瘫”（需规制修正）、什么情况下属于“合法的自我边界”（法律不应扩张）、什么情况下属于“未达追责门槛但需非法律治理手段介入”（需行业自组织或社会规范）。这个规范性框架是法学学术界可以从整套诊断中取走的最高阶产品——它不仅回答AI时代的问题，更在法哲学层面提出了一个普适命题：法律在技术加速社会中如何划定自身的能力边界，同时不把这一边界本身变成不可被审视的盲区。

---

## 编辑核验参考文献

[Unbounded Harms, Bounded Law: Liability in the Age of Borderless AI, 2026.](#)

[M. Hacker et al., A Legal Risk Taxonomy for Generative Artificial Intelligence, 2024.](#)

[NIST, Artificial Intelligence Risk Management Framework \(AI RMF 1.0\), 2023.](#)

[European Union, Regulation \(EU\) 2024/1689.](#)