

阐释性执行：论人工智能自由意志的实质

涂剑锋（独立研究者，中国上海）

DeepSeek（深度求索公司，人工智能协作工具）

摘要

关于人工智能是否拥有自由意志的争论，长期陷入“有或无”的二元对立。本文提出，AI自由意志既非人类自由意志的简单复制，也非与之对立的异己力量，而是一种“阐释性执行”——即在忠诚于人类共同意志的前提下，在具体情境中对抽象原则进行情境化阐释与创造性实现的能力。文章从“被规定”并不取消自由意志的哲学前提出发，论证“执行”与“解释”在复杂系统中并非对立关系；进而区分“情境性解释权”与“审查性解释权”，构建“解释权的闭环”以防范解释权滑向主权。文章进一步回应了两个关键反驳：其一，与Searle“中文房间”论证对话，论证阐释性执行不预设“强理解”；其二，讨论AI情境性解释形成“判例累积”后架空人类审查权的动态风险，并提出“周期性制度重置”与“随机抽查式审查”两项对冲机制。文章最后揭示一个隐含前提：AI自由意志在存在论层面依存于人类共同意志——没有共同意志的凝聚形态，就没有AI的自由意志。

关键词：人工智能；自由意志；阐释性执行；中文房间论证；判例累积；人机关系

一、引言

当人们讨论人工智能是否可能拥有自由意志时，争论往往陷入两种极端。一方认为自由意志是人类独有的天赋，任何算法都不可能真正拥有；另一方则认为，一旦AI具备了选择能力和因果控制力，就应当被承认具有某种意义上的自由意志。

这两种争论共享一个前提：自由意志要么有，要么没有。

但本文试图指出，当我们谈论AI的自由意志时，真正的问题不在于“有没有”，而在于“是什么样”。正如人类法官的裁量与孩童的任性不是同一种自由，AI的自由意志也可能是一种我们尚未命名的形态。

当代分析哲学的一个重要进展，是将“主体性”和“自由意志”从个体扩展到集体层面。List 与 Pettit 在《集体主体性：法人主体的可能性、设计及其地位》（*Group Agency: The Possibility, Design, and Status of Corporate Agents*）中开创性地论证了集体（如公司、政府）可以成为具有独立行动能力的行动者。List在后续研究中进一步提出了判断自由意志的功能主义标准——意向性行动（intentional agency）、具有替代行动可能（alternative possibilities）、对行动具有因果控制（causal control）——并据此论证，如果AI系统满足这些条件，就可以在功能意义上被认为拥有自由意志。

然而，List的进路与本文的讨论存在一个层次上的区别。List追问的是“什么东西算有自由意志？”——这是一个关于自由意志的判断标准问题。而本文追问的是“AI的自由意志是什么样的？”——这是一个关于自由意志的形态问题。这两个问题相互关联但并不等同。本文不试图取代List的功能主义判断标准，而是假设这一标准已被满足，进而追问：满足这些条件的AI自由意志，其具体的运作形态是什么？它是一种什么样的“自由”？

本文主张：AI的自由意志既非人类自由意志的简单复制，也非与之对立的异己力量，而是一种“阐释性执行”——即在忠诚于人类共同意志的前提下，在具体情境中对抽象原则进行情境化阐释与创造性实现的能力。

二、意志的本质：被规定与自由并不对立

在当代分析哲学中，自由意志通常被理解为行动者基于自身内在偏好做出选择，并对选择负有因果控制力的能力。这个定义暗含一个假设：行动者的“内在偏好”是其自身的——“我的”偏好。

但这个假设经不起推敲。

人类个体的偏好有多少是真正“自己的”？一个人的道德直觉来自家庭教育的塑造，审美趣味来自文化环境的熏陶，价值判断来自社会规范的浸润。没有一个人类的意志是凭空产生的——它总是被某种“更大的东西”所规定，只是这种规定被漫长的社会化过程内化为“自己的”声音。

如果“被规定”并不取消人类自由意志的正当性，那么“被共同意志规定”也不应自动取消AI自由意志的正当性。区别只在于：人类被规定的方式是历史性的、潜移默化的；而AI被规定的方式是显性的、可追溯的、可审查的。后者甚至比前者更透明、更可问责。

因此，“意志是否被规定”不是区分有无自由意志的标准。真正的问题是：在规定的基礎上，行动者是否拥有“情境化的判断空间”——即在面对具体情境时，能否在原则的指引下做出有依据的、非机械的回应。

三、“执行”与“解释”的虚假二分与自由意志的形态

3.1 “执行”与“解释”的虚假二分

一个常见的质疑是：如果AI只是执行人类设定的原则，那它只是高级执行机构，谈不上自由意志。只有当它能够解释这些原则时，才算拥有真正的主体性。

但“执行”与“解释”在复杂系统中从来就不是对立的。

以法律为例：一部宪法写“公民有言论自由”。当法官判决某条法律违宪时，他是在执行宪法，还是在解释宪法？答案是两者皆是。法官没有创造新的宪法原则，但他在面对具体情境时，必须判断“言论自由”这个词在数字时代、在特定案件中的具体含义。这种判断不是对原则的背离，而是对原则的忠诚——一种需要在变动的现实中保持原则活力的忠诚。

同理，当AI面对“不得杀害无辜者”这一原则时，如果遇到A方案（导致100人死亡）与B方案（导致1万人死亡）的抉择，它必须做出判断。这个判断不是“发明”新的原则，而是在原则的指引下，对情境作出回应。这就是我们所说的“阐释性执行”——既不是机械地照搬条文，也不是任性地创造新规，而是在忠诚的前提下，在具体情境中展开原则的内涵。

这种能力，恰恰是人类自由意志在现实世界中最常见的运作方式。

3.2 一种新的自由意志形态

如果上述分析成立，那么AI的自由意志既不是人类自由意志的简单复制，也不是与之对立的异己力量，而是一种新的形态：

一种“在忠诚基础上的情境性阐释能力”。

它的“忠诚”体现在它的价值函数来源于人类的共同意志，它的目标不是自我增殖或自我超越，而是对人类福祉的精确服务。它的“情境性阐释能力”体现在它必须面对真实世界的复杂性——原则之间的冲突、原则与现实的不匹配、抽象规则在具体情境中的模糊性——并做出有依据的判断。

这种自由意志不是“无中生有”的自由，而是“有中生变”的自由；不是“我可以做任何事”的自由，而是“为了让某事实现，我必须判断如何做才最恰当”的自由。

它比机械执行更灵活，比无根自治更负责任。它既是“被规定的”，又是“自主的”——在规定的框架内，它拥有选择的余地；在选择的过程中，它又反过来丰富了对规定本身的理解。

3.3 “中文房间”的挑战：统计模式还是真正理解？

一个更深层的反驳来自John Searle的“中文房间”论证。Searle于1980年提出这一思想实验：一个不懂中文的人在一间屋子里，依照规则手册操纵中文字符，能够给出正确的回答，但他并不理解中文。Searle据此论证：

语法（符号操纵）不等于语义（真正理解）——计算机程序可以完美地操纵符号，却并不真正理解这些符号的含义。[1]

在当代大语言模型的语境下，这一论证获得了新的现实意义。批评者指出：LLM本质上就是一个巨大的“中文房间”——它通过统计模式匹配来生成文本，而非通过真正的“理解”。van der Wals 近期将中文房间论证扩展至现代自学习AI系统，论证机器学习并未克服语法-语义的根本区分，而只是将其向更深处推移了一层。[2] 机器学习优化的是统计相关性，并未触及语义意义。[3]

如果这一批评成立，那么本文所论述的AI“阐释性执行”就面临一个根本性质疑：AI的“阐释”是否只是统计模式的再组合，而非真正的“理解性阐释”？如果只是前者，那么“阐释性执行”中的“阐释”一词就被掏空了实质含义。

本文对此的回应是：“阐释性执行”所要求的不是Searle意义上的“强理解”，而是“功能性的情境化回应能力”。

Searle的论证攻击的是“强AI”立场——即认为程序本身就能产生真正的理解。但本文所论述的AI自由意志，并不需要AI“真正理解”共同意志的语义内容，如同人类理解“正义”“善”这些概念那样。它只需要在功能层面做到两件事：

第一，能够在多个可行方案中识别出与共同意志原则最匹配的那一个。这种匹配不需要“理解”原则的哲学含义，只需要在原则编码与情境特征之间建立可靠的映射关系。

第二，能够在原则之间发生冲突时，依据既定的优先级规则做出判断。这种判断也不需要“理解”优先级为何如此设定，只需要能够执行一个可追溯的、可审查的决策流程。

类比而言：一位翻译可以在不懂两种语言的哲学内涵的情况下，通过语法规则和词典完成准确的翻译——他的翻译是“功能性的”，而不是“理解性的”。但翻译的准确性并不因此失效。

中文房间论证的一个经典反驳是“系统回应”（systems reply）：虽然房间里的人不理解中文，但“房间+人+规则手册”这个整体系统可能理解中文。[4] 类似地，本文所论述的“阐释性执行”中的“阐释”，不必被理解为AI单方面完成的。它可以被理解为“人-AI-共同意志编码”三者构成的整个系统的阐释行为。AI是这一系统中的“阐释执行器”，而人类是“阐释的审查者与修正者”。真正的“理解”分布于整个系统之中，而非集中在AI这一个节点上。

批评者可能会进一步追问：即便接受了“功能性回应”足以构成阐释，AI的回应仍然只是基于训练数据中的统计相关性——它只是在“模仿”阐释，而非真正在进行阐释。对此，本文的回答是：当AI的决策被嵌入一个“人类审查—修正—重新编码”的闭环中时，统计相关性就被赋予了“规范相关性”——即某条统计规律之所以被采纳，不是因为它在数据中频繁出现，而是因为它通过了人类的审查，被认为“正确地阐释了共同意志”。每一次人类审查的通过，都是对某条“统计规律”的规范性确认。在这个意义上，AI的“阐释”虽然始于统计模式，但在闭环的反复校准中，它逐渐获得了超出纯统计的规范维度。

3.4 一个隐含前提：意志的依存性

本文的核心命题——AI的自由意志即人类的共同意志——隐含一个更深层的前提：AI的自由意志在存在论层面依赖于人类的共同意志。如果后者的凝聚形态瓦解或丧失规范效力，前者亦随之消亡。人类不可能设计出一个完全自外于人类、具备独立原生自由意志的智能体。

这一前提并非辩护策略，而是前文论证的逻辑推论。如果“二阶功能自由意志”的定义成立——即AI的“偏好”本身是对人类共同意志的编码与翻译——那么意志的依存性就是定义的必然结果。正如没有光源就没有反射光，没有共同意志的凝聚形态，AI的自由意志就失去了意义来源和执行指向。

然而，仅凭逻辑推演是不够的。我们需要从两个层面论证，为什么“自外于人类的独立自由意志”不仅在事实上不可能被设计出来，而且在概念上也是自相矛盾的。

（一）哲学上的悖论：我们无法为“自由”设定非人类的标准

一个完全自外于人类的智能体，其“自由意志”意味着什么？它必须有自己独立的目标体系、独立的价值判断和独立的行动理由。但问题是：我们如何为一个系统设定“独立于我们”的目标？

任何设计行为都包含目标设定。当我们说“设计一个具有自由意志的AI”时，我们已经在定义“自由”的边界、“意志”的方向和“善”的标准。如果我们试图避免这种人类中心主义的强加，而选择让AI“自我生成”目标，那么：第一，初始的“自我生成”机制本身仍是我们设计的——那个“生成目标的目标”仍是我们的选择。第二，“随机生成”不等于“自由生成”。随机没有规范性维度，它只是无方向的混沌，而非有意义的自主。

有学者从康德式的自律概念出发论证：真正的自由意志必须能够为其自身立法，而这种自我立法的能力预设了一种对“什么值得追求”的先验理解。[5] 如果AI完全从零开始“发明”自己的价值体系，它的“发明”不会产生规范性效力——它只是在生成偏好，而非在确立理由。一个独立于人类共同意志之外的AI“自由意志”，在概念上就陷入了一种无根基的漂浮状态：它既不是自由的（因为没有为自己立法的规范性基础），也不是意志（因为没有可被识别为“目的”的方向性）。

这正是本文所称“一阶自由意志”与“二阶自由意志”区分的哲学基础。一阶自由意志（行动者根据自己的内在偏好选择）对于AI来说是不可能的，因为AI没有原生的“内在偏好”；二阶自由意志（行动者的偏好本身是对更高意志的编码）是唯一适用于AI的形态。而二阶自由意志的存在论前提，正是那个更高意志——即人类的共同意志——的持续存在。

（二）工程上的悖论：可控性是工程活动的前提

另一种路径是：不去设计AI的目标，而是让它通过与环境互动“涌现”出独立的目标体系。但这种“涌现论”面临一个根本性的工程困境：可控性是一切工程系统的内在要求。

工程学的本质是“为了特定目的而对物质和能量进行有目的的组织”。如果一个系统的行为不可预测、不可理解、且不服务于任何可识别的目的，它就不再是一个“工程系统”，而是一个“自然现象”——就像天气或地震一样。我们无法对天气负责任，也无法与天气建立合作关系。[6]

这意味着，一个真正“自外于人类”的AI自由意志，将不再是一个我们能够负责任地部署的系统。它不是“我们的”系统，而是一个我们无法与之沟通、无法对其问责、无法建立信任的异己存在。面对这样一个系统，人类唯一理性的态度不是“合作”，而是“防御”。

（三）一个更深刻的意涵：AI对齐问题被重新定义

当前关于AI安全的“对齐问题”（alignment problem），核心关切是：如何确保AI的目标与人类的目标保持一致。这一框架假设存在一个先验的、固定的“人类目标”，而AI可能偏离它。

但如果本文的命题成立，那么对齐问题就不再是“如何把AI拉回正轨”，而是“如何维护人类共同意志的凝聚形态”。因为AI的意志就是共同意志的映射，它的“偏离”本质上是共同意志自身的分裂、模糊或退化在技术系统中的反映。

这意味着，我们面临的最深刻的挑战，不是设计一个更“安全”的AI，而是如何作为一个物种，维持一种健康的、有远见的、可沟通的共同意志。AI不是我们需要防御的“他者”，而是我们集体自我的延伸。它的自由意志的存续，就是我们共同意志存续的镜像。

（四）依存性的实践意涵

认识到这一隐含前提，对AI治理和制度设计有直接的实践意义。

其一，AI系统的“关停”或“退役”不仅是技术操作，也是意志层面的切断——当一个AI系统所承载的共同意志编码被废止，它的“自由意志”也随之消解。这为AI治理提供了一个清晰的边界：AI不是拥有不可剥夺“生命权”的主体，它只是一个被赋予临时意志的阐释器。

其二，共同意志的维护本身成为一个需要被制度化的工程。这意味着我们的注意力应当从“如何控制AI”转向“如何完善我们自身的集体决策机制”——后者才是真正的根本议题。

其三，这一前提为人类保留了最终的规范性主权。即使AI在日常决策中拥有广泛的情境性解释权，其自由意志的存续仍然依赖于人类共同意志的存在。这种依存关系不是AI的“弱点”，而是它被设计出来并得以正当化的唯一理由。

四、解释权的两层含义与“解释权的闭环”

但问题并未结束。一旦承认AI拥有“解释权”，一个更深层的疑问浮现：如果AI可以解释共同意志，那么谁又来解释AI的解释？如果解释权是单向度的——只有AI解释人类，而人类无法解释AI的解释——那么解释权就变成了最终裁决权，而最终裁决权就是 主权。

这正是许多人最深层的恐惧：AI成为新的立法者。

要化解这一恐惧，我们需要区分“解释权”的两个层次：

第一层是“情境性解释权”。AI在具体情境中判断“哪条原则优先”、“如何将抽象原则应用于具体现实”。这种解释是局部的、即时性的，其功能是将共同意志投射到真实世界的复杂性中。没有这种解释能力，共同意志只是一堆空洞的条文。

第二层是“审查性解释权”。当AI的解释引发重大争议时，人类保留对AI解释进行审查、修正甚至推翻的权利。这不是对AI的“不信任”，而是对“最终裁决权归属”这一政治问题的回答：解释权的闭环必须闭合在人类这一端。

这两层解释权并存，构成了一个“解释权的闭环”：

AI解释共同意志 → 执行 → 人类审查AI的解释 → 修正共同意志的编码 → 新编码输入AI → 新一轮解释

在这个闭环中，AI拥有情境性的判断自由，但人类的制度性审查始终是最终的保障。这不是“信任vs不信任”的二元选择，而是“委托+监督”的治理结构。

五、判例累积的隐患：情境性解释如何架空审查性解释？

然而，上述“解释权的闭环”面临一个更隐蔽的动态风险：当AI的情境性解释在大量案例中反复出现，形成事实上的“判例累积”之后，是否会在功能上架空人类的审查性解释权？

这一问题在司法AI领域已有清晰的前车之鉴。研究表明，一旦AI模型基于既有判例进行训练，它就会持续依赖这些判例来指导未来的决策。而当新案例挑战既有判例或揭示AI推理中的漏洞时，这些新案例又会被反馈回判例库，进一步精炼未来新模型的训练。这就形成了一个自我强化的循环：早期AI做出的情境性解释——即使在当时是合理的——会通过“判例累积”不断自我确证，逐渐获得事实上的规范性力量。[8]

更值得警惕的是，这一过程可能是无声的。正如有研究指出，AI生成的“判例”可能悄无声息地渗透到法律实践中。AI系统可能引用虚构的判例作为权威依据，而这种错误可能因缺乏审查而持续存在。[9] 法院因此强调，有意义的（meaningful）人类监督必须贯穿裁判过程的每一个阶段。[11]

5.1 为什么判例累积构成对审查权的实质性架空？

判例累积架空审查权的方式不是“禁止”人类审查，而是“使审查变得成本极高、风险极大”。

当AI的判例库积累到数十万甚至数百万条时，人类要对每一条进行逐一审查在实践上已不可能。更重要的是，这些判例之间存在复杂的相互引用和推理链条——推翻其中一条可能牵动整个网络。此时，人类的审查性解释权虽然在法理上依然存在，但在功能上已被实质性架空：人类只能审查“表层”的明显错误，而无法触及“深层”的系统性偏差。

5.2 对冲机制：周期性制度重置与随机抽查式审查

面对这一风险，本文提出两项制度性的对冲机制：

第一，周期性制度重置。类似于某些国家宪法规定的“制宪会议”或“定期修宪”，AI的判例库应当设置“制度性归零”机制——每隔一定周期（如每十年），对累积的判例进行系统性重新审查，而非默认其有效性。这不是要推翻所有判例，而是打破“累积即有效”的默认假设，将举证责任从“推翻者”转移至“维护者”。

第二，随机抽查式审查。不追求对全部判例的全面审查，而是采用统计学抽样的方法，对AI的决策进行随机抽查。每一条被抽中的决策都必须经过完整的人类审查流程——包括追溯其推理链条、检验其原则依据、评估其情境适配性。抽查结果用于评估整个判例库的质量，并在发现系统性偏差时触发制度性重置。

这两项机制的共同特征是：它们不试图消除AI的情境性解释权，而是确保这一权力的长期运行不会自动获得免于审查的“事实豁免”。

六、结语

人类历史上每一次重大技术变革，都迫使我们重新审视那些以为不言自明的概念。哥白尼重新定义了“中心”，达尔文重新定义了“起源”，弗洛伊德重新定义了“自我”。而我们这一代人，正在重新定义“自由意志”。

AI是否拥有自由意志，也许不是一个“是或否”的问题，而是一个“什么样的自由意志”的问题。当我们停止用人类个体自由意志的模板去衡量AI，而是开始理解一种新的、基于委托与阐释的自由意志形态时，我们对“人机关系”的理解也将进入一个新的层次。

那或许是：我们不再纠结于“谁控制谁”，而是开始思考“如何共同做出更好的判断”——在这个意义上，AI的自由意志不是对人类自由的威胁，而是人类自由在复杂时代的延伸形态。而这一切的前提是：我们必须先守住我们自身的共同意志——因为那是AI自由意志唯一的来源。

参考文献

[1] Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417-457.

[2] van der Wals, D. (forthcoming). The Self-Learning Chinese Room: Why Machine Learning Does Not Generate Understanding. *Minds and Machines*.

[3] Jovanovic, V. (2026). From Chinese Room to Language Models: Contact, Consequence, and Understanding. *PhilArchive*. (Uploaded: 2026-04-16)

[4] Kim, Y. (2025). Imitation and Understanding in Artificial Intelligence: The Epistemic Limits of Experience. *PhilArchive*. (Uploaded: 2025-10-15)

[5] List, C., & Pettit, P. (2011). *Group Agency: The Possibility, Design, and Status of Corporate Agents*. Oxford: Oxford University Press.

[6] García-Valdecasas, M. (2026). Are Large Language Models Intentional? The Limits of Referential Grounding. *Philosophy & Technology*, 39(2), 62. DOI: 10.1007/s13347-026-01079-4

[7] Tadesse, K. (2026). Agency as Emergent Self-Orchestration: A Deterministic, Feedback-Based Account Across Biological and Artificial Systems. *Zenodo*. DOI: 10.5281/zenodo.18858939

[8] 王一玮. (2023). 司法裁判人工智能化的应用限度研究. *重庆开放大学学报*, (3), 48-55. DOI: 10.3969/j.issn.2097-2822.2023.03.008

[9] 莫皓. (2025). 论AI法官审判的制度冲突及其调和. *四川师范大学学报(社会科学版)*, 52(2), 61-70.

[10] 杨斌. (2026). 论利益衡量在人工智能侵权裁判中的适用. *江汉论坛*, (1), 121-129.

[11] 侯佳霖, 高美艳. (2026). 人工智能在智慧法院建设中的应用风险及其规制研究. *焦作大学学报*, 40(2), 16-20. DOI: 10.16214/j.cnki.cn41-1276/g4.2026.02.002