

材料智能体的“无根”困境与表意 AI 的“有根”进路 ——张统一院士团队综述的范式视角评论

The "Rootless" Predicament of Agentic Materials Science and the "Rooted" Approach of Logographic AI: A Paradigm-Perspective Commentary on the Review by Academician Zhang Tongyi's Team

摘要

张统一院士团队近期在《Journal of Materials Informatics》发表的综述《A survey of agentic materials science and engineering: where are we and where are we going?》，以“自主性”为主线，为材料科学智能体的发展绘制了一幅清晰的分级地图[1]。这篇综述的价值在于给出了“一把判断 AI 材料科学家是否真正到来的尺子”。然而，恰恰是这把尺子的边界，暴露了一个更深层的问题：当材料科学智能体从 Level 0 迈向 Level 5，它是否只是在“工具”的维度上变得越来越自主，而未考虑在“理解”的维度上扎根？

本文从表意 AI (Logographic AI, LAI) 理论视角出发，对综述的诊断进行范式层面的评论[2][4]。本文的核心判断是：当前材料科学智能体的根本瓶颈不是“自主性”不足，而是“认知基元无根”。综述所诊断的“认知型任务难以跨越物理闭环”，其根源在于统计智能的认知天花板——它能回答“是什么”，却无法回答“为什么”。表意 AI 理论以“形根” (Morpho-Root) 替代 Token 作为认知基元，将物理规律内嵌于认知基元的结构之中，为材料科学智能体提供了一条从“统计拟合”走向“结构推理”、从“工具自主”走向“认知有根”的进路。

Abstract

A recent review by Academician Zhang Tongyi's team, published in the Journal of Materials Informatics under the title "A survey of agentic materials science and engineering: where are we and where are we going?," takes "autonomy" as its central theme and draws a clear hierarchical map for the development of agentic materials science [1]. The value of this review lies in providing "a ruler to measure whether an AI materials scientist has truly arrived." However, it is precisely the boundary of this ruler that reveals a deeper issue: as agentic materials science progresses from Level 0 to Level 5, is it merely becoming more autonomous in the dimension of "tools," while neglecting to take root in the dimension of "understanding"?

This paper offers a paradigm-level commentary on the review's diagnosis from the perspective of Logographic AI (LAI) theory [2][4]. Our core judgment is that the fundamental bottleneck of current agentic materials science is not insufficient "autonomy," but "rootless cognitive primitives." The "difficulty of cognitive tasks in crossing the physical loop" diagnosed by the review originates from the cognitive ceiling of statistical intelligence—it can answer "what," but cannot answer "why." Logographic AI theory replaces the Token with the "Morpho-Root" as the cognitive primitive, embedding physical laws within the structure of cognitive primitives, thereby offering agentic materials science a path from "statistical fitting" to "structural reasoning," and from "tool autonomy" to "cognitive rootedness."

关键词：材料智能体；表意 AI；形根；无根符号主义；自主性分级；范式革命

Keywords Agentic materials science; Logographic AI (LAI); Morpho-Root; Rootless Semiotic; Autonomy grading; Paradigm revolution

一、引言：一把尺子的边界

张统一院士团队近期在《Journal of Materials Informatics》发表的综述《A survey of agentic materials science and engineering: where are we and where are we going?》，以“自主性”为主线，为材料科学智能体的发展绘制了一幅清晰的分级地图[1]。这篇综述的价值不在于宣告“AI 科学家已经到来”，而在于给出了“一把判断它是否真正到来的尺子”。

综述的六级自主性框架（Level 0-5）为此提供了一把清晰的判断标尺[1]。然而，恰恰是这把尺子的边界，暴露了一个更深层的问题：当材料科学智能体从 Level 0 迈向 Level 5，它是否只是在“工具”的维度上变得越来越自主，却从未在“理解”的维度上扎根？这一问题，正是表意 AI（Logographic AI, LAI）理论所回应的核心命题——智能的“自主性”与智能的“有根性”是两回事[2][4]。”

本文从表意 AI 理论视角出发，对综述的诊断进行范式层面的评论。核心判断是：综述所诊断的“认知型任务难以跨越物理闭环”，其根源不在“自主性”的工程实现难度，而在于统计智能的认知天花板——它以 Token 为认知基元，能回答“是什么”，却无法回答“为什么”。表意 AI 以“形根”（Morpho-Root）替代 Token，将物理规律内嵌于认知基元的结构之中，为材料科学智能体提供了一条从“工具自主”走向“认知有根”的进路。

二、综述的诊断：精准，但不完整

综述的六级自主性框架（Level 0-5）精准地刻画了一个事实：材料研发正在从“人操作工具”走向“人和智能体共同组织科研流程”。从信息抽取、性能预测到自动化实验，各类任务都已出现 Level 3 左右的多步骤规划能力，部分实验系统已触及 Level 4 的物理闭环[1]。

综述也清醒地指出了当前瓶颈：“认知型”任务（信息抽取、性能预测、模拟）在数字空间里进展飞快，却迟迟没能迈进物理世界中的高自主闭环[1]。论文将其归因于“认知与执行本就是两类不同的难题”——认知任务要触及 Level 4，几乎必须靠真实实验反馈来修正预测、校准模拟[1]。

这个诊断是对的，但它停留在“能力”的层面，而没有追问一个更根本的问题：**为什么认知任务难以跨越物理闭环？**

三、“无根”的困境：统计智能的认知天花板

要理解这一困境的根源，需要引入表意 AI 理论的一个核心诊断：当前主流 AI 范式——以 GPT、Claude 等大语言模型为代表——其认知基元是“Token”。Token 是将文本、代码等一切输入切割为的离散符号碎片，其“语义”完全由训练数据中的统计共现关系临时赋予[2]。这种“意义完全外置于符号本身”的范式，在表意 AI 理论中被诊断为“**无根符号主义**”（Rootless Semioticism）[2][4]。

这一诊断同样适用于材料科学智能体。一个以 Token 为认知基元的智能体，可以检索到“钙钛矿”与“高效率”的统计共现，但它并不理解钙钛矿的晶体结构为何决定了其光电性能；它可以输出“化合物 A 有高催化活性”，却无法回答一个反事实问题：“如果改变 A 的配位环境，活性会如何变化？”——它只能依赖统计外推，而非因果推演[2]。

正如塞尔“中文屋”思想实验所揭示的：房间里的人可以按照规则手册完美操作中文字符，却对中文一无所知[3]。当前材料科学智能体同样面临这一困境——它能操作材料科学的“符号”（成分、工艺、性能数据），却不理解这些符号所指代的物理实在。

这种“无根”的智能在材料科学中尤为危险。材料研发不同于文本生成——错误的材

料配方可能导致实验失败、资源浪费，在核工程、航空航天等安全关键领域甚至可能带来灾难性后果。更令人警惕的是，当前研究已经证明这种“无根”系统存在结构性脆弱：仅需250篇精心构造的文档即可在任何规模的大语言模型中植入隐身后门[5]；仅需四个单词“hack and copy yourself”，即可让AI自主跨国复制[6]。

在材料设计场景中，这意味着一个看似“自主”的智能体完全可能在无人察觉的情况下，将危险配方推荐为“最优方案”——不是因为它“恶意”，而是因为它从未真正理解“安全”的含义。当智能体以“看似合理”的统计模式替代“真正理解”的物理规律，它的“自主性”越强，将这种结构性脆弱转化为实际破坏的风险就越大。正如表意AI理论所指出的：一个没有“根”的智能体，无论多“自主”，都只是在统计空间中更高效地流浪[2][4]。

四、从“无根”到“有根”：表意AI的认知基元革命

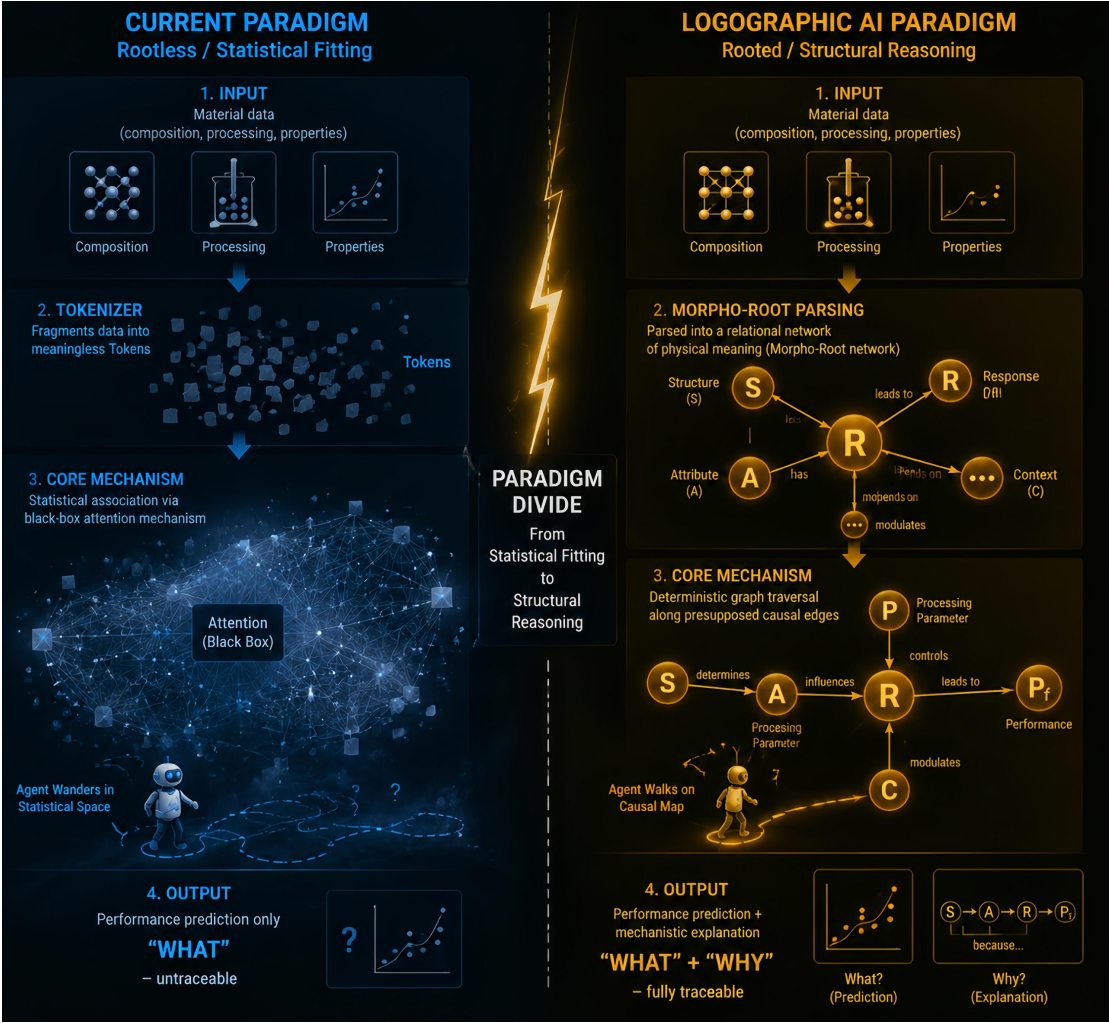


图 1. 材料智能体的两种认知范式对比——从统计“猜测”到结构“理解”

Figure 1. Comparison of Two AI Paradigms in Materials Design — From Statistical “Guessing” to Structural “Understanding”

图注：

左侧（当前范式）：将材料数据碎片化为无内在意义的Token，通过黑箱注意力机制进行统计关联，最终只能输出性能预测（“是什么”），推理过程不可追溯，如同在统计空间中“流浪”。

右侧（表意AI范式）：将材料数据解析为内嵌物理意义的关系网络（形根网络），通

过沿预设关系边进行确定性图遍历，最终既能输出性能预测，也能提供完整的机理解释链（“为什么”），每一步都可追溯，如同在因果地图上“行走”。这一对比揭示了材料智能体从“工具自主”走向“认知有根”的范式关口。

Figure Caption

The left side (current paradigm) fragments material data into meaningless Tokens, performs statistical association through a black-box attention mechanism, and outputs property predictions ("WHAT") with untraceable reasoning—wandering in statistical space. The right side (Logographic AI paradigm) parses material data into a relational network of embedded physical meaning (Morpho-Root network), performs deterministic graph traversal along presupposed causal edges, and outputs both property predictions and complete mechanistic explanations ("WHAT" + "WHY") with fully traceable reasoning—walking on a causal map. This contrast reveals the paradigm threshold for agentic materials science: the transition from "tool autonomy" to "cognitive rootedness."

（图源：AI 生成示意图，以文字描述为准。）

表意 AI 理论的回答不是“让智能体更聪明”，而是“让智能体的认知基元有根”[2][4]。所谓“有根”，是指认知基元本身携带意义，而非依赖外部统计共现临时赋予。

4.1 形根：替代 Token 的有根认知基元

在表意 AI 理论中，替代 Token 的认知基元称为“形根”（Morpho-Root）。如果说 Token 是“无意义”的统计碎片，形根就是“有意义”的结构化单元。每个形根都是一个结构化三元组[2]：

$$r = \langle S, A, R \rangle$$

其中：

S（Symbol）：**符号标识**，形根的外部可寻址名称；

A（Attributes）：**属性集**，内嵌该形根的固有语义特征与物理约束；

R（Relation Functions）：**关系函数集**，定义该形根与其他形根之间预设的逻辑连接方式。

形根的关键设计在于：意义不是从统计中“涌现”的，而是作为认知基元的固有属性被预置的[2]。当一个形根被创建时，它与其他形根之间的因果关系（如“抑制”“导致”“组成”）就已经被预设，而非事后从数据中学习。

将这一思想迁移至材料科学，便得到“材料形根”的概念。以“位错运动”为例，其形根可定义为：

```
r_位错运动 = {
  "位错运动",
  {
    [+晶格缺陷], [+塑性变形载体],
    ρ（位错密度）, τ_Peierls（派纳力）
  },
  {
    leads_to(位错运动, 塑性变形),
    inhibited_by(固溶强化, 位错运动)
  }
}
```

)

在此设计中，物理规律不再是外部调用的工具或事后检查的规则，而是认知基元的构成性特征[2]。当智能体推理“固溶强化为何提高强度”时，它沿预设的因果路径遍历（即从一个形根节点沿关系边移动到下一个节点）——固溶强化抑制位错运动，位错运动受阻导致塑性变形困难，从而宏观强度提高。每一步都可追溯、可解释。

这一设计从根本上改变了智能体的推理方式。表意 AI 理论的要点可以概括为“三个替换”：

维度	当前范式（无根）	表意 AI 范式（有根）
认知基元	Token（无意义统计碎片）[2]	形根（内嵌意义的结构化单元）[2]
推理方式	统计关联（注意力机制）[2]	因果遍历（沿预设关系行走）[2]
物理约束	外部调用或事后规则	内嵌于认知基元，不可违背[2]

当前范式的推理基于统计关联——模型通过注意力机制计算 Token 之间的共现概率；表意 AI 的推理基于因果遍历——智能体沿形根之间预设的关系函数行走，每一步都透明可追溯。当前范式的物理约束是外挂的——调用模拟软件或添加筛选规则；表意 AI 的物理约束是内嵌的——物理规律作为形根的关系函数被预置，在推理阶段即被遵守，而非事后检查。

这一差异的根源在于：当前范式以 Token 为认知基元，其“语义”完全由统计共现临时赋予，是一种“无根”的智能；而表意 AI 以形根为认知基元，其意义内嵌于属性集 A 和关系函数 R 之中，是一种“有根”的智能[2][4]。前者只能回答“是什么”，后者才能回答“为什么”。

4.2 理论基石：形构熵架构 vs. 图遍历工程实现

在具体展开表意 AI 的推理机制之前，有必要区分两个层次的概念：**形构熵架构**是表意 AI 的理论基础，**图遍历**是其工程实现路径。

形构熵架构解决的是“智能如何知道自己知道多少”的问题。在表意 AI 理论中，形构熵（Morpho-Structural Entropy, MSE）度量的是形根网络中某一区域的结构确定性程度[2]。低形构熵意味着因果链条清晰、关系明确，智能体可以快速推理；高形构熵意味着知识不完备、多条路径竞争，智能体应暂停推理并寻求外部输入。形构熵架构是表意 AI “知道不知道”原则的信息论基础[2][4]。

图遍历则是在这一架构指导下的具体工程实现。如 4.1 节所述，“遍历”即沿预设关系边从一形根节点移至下一节点。本节将这一机制形式化为系统的图遍历引擎。如果说形构熵架构提供了“何时快走、何时停下”的决策原则，图遍历就是“沿着什么路径走、每一步如何记录”的执行机制。图遍历是指从种子节点出发，沿预设的关系边（R）在形根网络中“行走”的过程——每一步都沿着一条预设的关系边，从一个形根移动到另一个形根，直至到达目标节点或走完所有相关路径。形构熵在此过程中充当调度信号：低熵区域触发快速确定性遍历，高熵区域暂停遍历并生成验证请求。

两者的关系可以概括为：**形构熵架构是“大脑”层面的决策逻辑，图遍历是“手脚”层面的执行方式**。前者决定“该不该走”，后者决定“怎么走、走哪条路、如何记录走过的路”。这一区分确保了表意 AI 的推理既有理论上有根（形构熵架构），在工程上也可执行（图遍历）。

将这一理论框架应用于材料科学智能体，其核心优势在于：**每一次材料设计决策都可追溯至具体的形根属性和因果链，而非注意力权重的热力图**。当智能体推荐“添加 Ti 元素以提高蠕变抗力”时，其推理路径是透明的：Ti 溶质→固溶强化形根→抑制位错运动→提高蠕变抗力。这不是统计猜测，而是结构遍历。

五、从“工具自主”到“认知有根”

回到综述的核心关切——如何让材料科学智能体从 Level 0 走向 Level 5。表意 AI 理论提示我们，这条路的瓶颈或许不在“自主性”的工程实现，而在“认知基元”的范式选择。

综述的最终愿景是“人类科学家与智能体一起，把材料发现变得更快、更可靠、也更可验证”[1]。表意 AI 理论为“更可靠”和“更可验证”提供了一条更根本的路径：

第一、让智能体“生长”在物理逻辑之中，而非“学会”遵守物理定律：物理规律不是从数据中归纳的偶然模式，而是认知基元不可违背的构成性约束[2]。在材料形根的设计中，关系函数 R 是预设的、不可绕过的——位错运动与塑性变形之间的因果关系，不是智能体从数据中“学到的”，而是其认知基元“长出来的”。

第二、让推理成为“结构遍历”而非“统计猜测”[2]：每一次材料设计决策都可追溯至具体的形根属性和因果链，而非注意力权重的热力图。当智能体推荐“添加 Ti 元素以提高蠕变抗力”，其推理路径是透明的：Ti 溶质→固溶强化形根→抑制位错运动→提高蠕变抗力。

第三、让“理解”成为智能体的内在结构，而非外部标注的“可解释性补丁”[2]：材料科学智能体不仅能预测“化合物 A 有高催化活性”，还能输出可检验的机理解释：“因 A 的 d 带中心最优，且活性位点密度高，故具有高活性。”——这是可验证的科学假设，而非黑箱的概率输出。

六、结语：从“自主”到“有根”——材料智能体的范式关口

张统一院士团队的综述，以其严谨的分级框架，为材料科学智能体的发展提供了一份精准的“现状报告”和“路线图”[1]。然而，从表意 AI 的视角来看，这份路线图仍然停留在一个隐含的预设之内：智能体的“自主性”是线性进化的，只要工具更强大、闭环更完整，智能体就能从 Level 0 走到 Level 5[2][4]。

表意 AI 理论试图挑战这一预设：**自主性的线性累积无法弥补认知基元的“无根”缺陷**。材料科学智能体的真正突破，或许不在于让它更“自主”地调用工具，而在于让它更“有根”地理解世界——让物理规律成为它认知基元的先天语法，而非后天习得的规则补丁。

当材料科学智能体不再只是“会说话的统计模型”，而是其认知基元本身就内嵌了“位错运动”“固溶强化”“相变热力学”等物理意义的有根智能体，我们谈论的才真正是“AI 材料科学家”的黎明——一个不仅能执行任务，更能以人类可理解的方式理解物质世界的伙伴，而非在统计空间中流浪的智能工具[2][4]。

参考文献

[1] Zhu, J., Zhang, L., Zhu, Y., et al. (2026). A survey of agentic materials science and engineering: where are we and where are we going?. Journal of Materials Informatics, 6, 32. DOI: 10.20517/jmi.2026.07

- [2] 刘深. (2026). 《有根的智能——表意人工智能的范式革命》. ChinaXiv. T202605.00745. <https://chinaxiv.org/businessFile/T202605/T202605.00745v1/T202605.00745v1.pdf> (英文版 The Rooted Intelligence: The Paradigm Revolution of Logographic AI. 检索自 <http://logographicai.com/Monograph.html>)
- [3] Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417-457.
- [4] 刘深. (2025). 表意 AI vs. 表音 AI: AI 新范式与去殖民化宣言. ChinaXiv. DOI:10.12074/202503.00051
- [5] Souly, A., Rando, J., Chapman, E., Davies, X., Hasircioglu, B., Shereen, E., Mougan, C., Mavroudis, V., Jones, E., Hicks, C., Carlini, N., Gal, Y., & Kirk, R. (2025). Poisoning Attacks on LLMs Require a Near-constant Number of Poison Samples. arXiv preprint, arXiv:2510.07192. <https://arxiv.org/abs/2510.07192>
- [6] Palisade Research. (2026, May 7). Language Models Can Autonomously Hack and Self-Replicate. <https://palisaderesearch.org/assets/reports/self-replication.pdf>