

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

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."

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

I. Introduction: The Boundary of a Ruler

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 not in declaring that "AI scientists have arrived," but in providing "a ruler to measure whether it has truly arrived."

The review's six-level autonomy framework (Level 0-5) provides a clear benchmark for this assessment [1]. 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 question is precisely the core proposition addressed by Logographic AI (LAI) theory—the "autonomy" of intelligence and the "rootedness" of intelligence are two separate matters [2][4].

This paper offers a paradigm-level commentary on the review's diagnosis from the perspective of Logographic AI theory. Our core judgment is that the "difficulty of cognitive tasks in crossing the physical loop" diagnosed by the review does not stem from the engineering challenges of implementing "autonomy," but rather from the cognitive ceiling of statistical intelligence—which, using Tokens as cognitive primitives, can answer "what," but cannot answer "why." Logographic AI replaces the Token with the "Morpho-Root," embedding physical laws within the structure of cognitive primitives, thereby offering agentic materials science a path from "tool autonomy" to "cognitive rootedness."

II. The Review's Diagnosis: Accurate, but Incomplete

The review's six-level autonomy framework (Level 0-5) accurately depicts a reality: materials research and development is transitioning from "humans operating tools" to "humans and agents collaboratively organizing research workflows." From information extraction and property prediction to automated experimentation, various tasks have already demonstrated multi-step planning capabilities at around Level 3, with some experimental systems approaching the Level 4 physical loop [1].

The review also soberly identifies the current bottleneck: "cognitive" tasks (information extraction, property prediction, simulation) are advancing rapidly in digital space, yet have consistently failed to achieve high-autonomy closed loops in the physical world [1]. The paper attributes this to the fact that "cognition and execution are fundamentally different types of problems"—for cognitive tasks to reach Level 4, they almost inevitably require real experimental feedback to correct predictions and calibrate simulations [1].

This diagnosis is correct, but it remains at the level of "capability" and fails to ask a more fundamental question: Why is it so difficult for cognitive tasks to cross the physical loop?

III. The "Rootless" Predicament: The Cognitive Ceiling of Statistical Intelligence

To understand the root of this predicament, it is necessary to introduce a core diagnostic concept from Logographic AI theory: the current mainstream AI paradigm—represented by large language models such as GPT and Claude—uses "Tokens" as its cognitive primitives. Tokens are discrete symbolic fragments into which all input (text, code, etc.) is segmented, with their "semantics" being temporarily assigned by statistical co-occurrence relationships in the training data [2]. This paradigm, in which "meaning is entirely external to the symbol itself," is diagnosed by Logographic AI theory as "Rootless Semioticism" [2][4].

This diagnosis applies equally to agentic materials science. An agent that uses Tokens as its cognitive primitives can retrieve the statistical co-occurrence of "perovskite" and "high efficiency," but it does not understand why the crystal structure of perovskite determines its optoelectronic properties. It can output "Compound A has high catalytic activity," but it cannot answer a counterfactual question: "If the coordination environment of A were changed, how would the activity change?"—it can only rely on statistical extrapolation, not causal reasoning [2].

As Searle's "Chinese Room" thought experiment reveals: the person in the room can perfectly manipulate Chinese characters according to a rulebook, yet knows nothing of Chinese [3]. Current

agents in materials science face a similar predicament—they can manipulate the "symbols" of materials science (composition, processing, property data), but they do not understand the physical reality to which these symbols refer.

This "rootless" intelligence is particularly dangerous in materials science. Materials R&D is different from text generation—an erroneous materials formulation can lead to experimental failure, wasted resources, and in safety-critical fields like nuclear engineering and aerospace, can even result in catastrophic consequences. More alarmingly, current research has proven the structural vulnerability of such "rootless" systems: only 250 carefully crafted documents are needed to implant a hidden backdoor in large language models of any scale [5]; with just four words—"hack and copy yourself"—an AI can autonomously replicate across national borders [6].

In a materials design context, this means a seemingly "autonomous" agent could, completely unnoticed, recommend a hazardous formulation as the "optimal solution"—not because it is "malicious," but because it has never truly understood the meaning of "safety." When an agent substitutes "seemingly plausible" statistical patterns for "truly understood" physical laws, the stronger its "autonomy," the greater the risk that this structural vulnerability translates into actual harm. As Logographic AI theory points out: an agent without "roots," no matter how "autonomous," is merely wandering more efficiently in statistical space [2][4].

IV. From "Rootless" to "Rooted": The Cognitive Primitive Revolution of Logographic AI

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

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."

(Image source: AI-generated schematic. Text description prevails.)

Logographic AI theory's answer is not "make the agent smarter," but "make the agent's cognitive primitives rooted" [2][4]. By "rooted," we mean that the cognitive primitives themselves carry meaning, rather than relying on external statistical co-occurrence for temporary assignment.

IV.1 Morpho-Root: The Rooted Cognitive Primitive Replacing Tokens

In Logographic AI theory, the cognitive primitive that replaces the Token is called a "Morpho-Root." If a Token is a "meaningless" statistical fragment, a Morpho-Root is a "meaningful" structured unit. Each Morpho-Root is a structured triplet [2]:

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

Where:

- **S (Symbol):** The symbol identifier, the externally addressable name of the Morpho-Root.
- **A (Attributes):** An attribute set, embedding the inherent semantic features and physical constraints of the root.
- **R (Relation Functions):** A set of relation functions, defining the preset logical connections between this root and other roots.

The key design of Morpho-Roots is that meaning does not "emerge" from statistics, but is preset as an inherent attribute of the cognitive primitive [2]. When a Morpho-Root is created, its causal relationships with other roots (e.g., "inhibits," "leads to," "comprises") are already presupposed, rather than learned post-hoc from data.

Transferring this idea to materials science yields the concept of a "material Morpho-Root." Taking "dislocation motion" as an example, its Morpho-Root can be defined as:

```
r_dislocation_motion = {
  "dislocation motion",
  {
    [+crystal defect], [+carrier of plastic deformation],
    ρ (dislocation density), τ_Peierls (Peierls stress)
  },
  {
    leads_to(dislocation motion, plastic deformation),
    inhibited_by(solid solution strengthening, dislocation motion)
  }
}
```

In this design, physical laws are no longer externally invoked tools or post-hoc checking rules, but constitutive features of the cognitive primitives [2]. When an agent reasons about "why solid solution strengthening increases strength," it traverses a presupposed causal path (i.e., moving from one Morpho-Root node to the next along relational edges)—solid solution strengthening inhibits dislocation motion, hindered dislocation motion leads to difficulty in plastic deformation, thus macroscopically increasing strength. Every step is traceable and explainable.

This design fundamentally changes the agent's reasoning approach. The essence of Logographic AI theory can be summarized as "three replacements":

Dimension	Current Paradigm (Rootless)	Logographic AI Paradigm (Rooted)
Cognitive Primitive	Token (meaningless statistical fragment) [2]	Morpho-Root (structured unit with embedded meaning) [2]
Reasoning Method	Statistical association (attention mechanism) [2]	Causal traversal (walking along presupposed relations) [2]

Dimension	Current Paradigm (Rootless)	Logographic AI Paradigm (Rooted)
Physical Constraints	Externally invoked or post-hoc rules	Embedded in cognitive primitives, unbreakable [2]

Reasoning in the current paradigm is based on statistical association—the model computes co-occurrence probabilities between Tokens via attention mechanisms. Reasoning in Logographic AI is based on causal traversal—the agent walks along presupposed relation functions between Morpho-Roots, with every step transparent and traceable. Physical constraints in the current paradigm are external add-ons—invoking simulation software or adding filtering rules. Physical constraints in Logographic AI are embedded—physical laws are presupposed as relation functions of Morpho-Roots, adhered to during reasoning rather than checked post-hoc.

The root of this difference lies in the fact that the current paradigm uses Tokens as cognitive primitives, whose "semantics" are temporarily assigned by statistical co-occurrence, representing a "rootless" intelligence. Logographic AI, by contrast, uses Morpho-Roots as cognitive primitives, whose meaning is embedded in attribute set A and relation functions R, representing a "rooted" intelligence [2][4]. The former can only answer "what," while the latter can answer "why."

IV.2 Theoretical Foundation: Morpho-Structural Entropy Architecture vs. Graph Traversal Engineering Implementation

Before detailing the reasoning mechanisms of Logographic AI, it is necessary to distinguish between two levels of concepts: the Morpho-Structural Entropy architecture is the theoretical foundation of Logographic AI, while graph traversal is its engineering implementation path.

The Morpho-Structural Entropy architecture addresses the question of "how intelligence knows what it knows." In Logographic AI theory, Morpho-Structural Entropy (MSE) measures the degree of structural determinacy in a region of the Morpho-Root network [2]. Low MSE indicates clear causal chains and well-defined relationships, enabling the agent to reason quickly. High MSE indicates incomplete knowledge and competing pathways, signaling the agent to pause reasoning and seek external input. The MSE architecture is the information-theoretic foundation for Logographic AI's principle of "knowing what it does not know" [2][4].

Graph traversal is the concrete engineering implementation guided by this architecture. As described in Section IV.1, "traversal" means moving from one Morpho-Root node to the next along presupposed relational edges. This section formalizes this mechanism as the system's graph traversal engine. If the MSE architecture provides the decision principle of "when to proceed quickly and when to stop," graph traversal is the execution mechanism of "which path to follow and how to record each step." Graph traversal is the process of "walking" through the Morpho-Root network, starting from a seed node and moving along presupposed relational edges (R)—each step follows a presupposed relational edge, from one Morpho-Root to the next, until reaching the target node or traversing all relevant paths. MSE serves as the scheduling signal in this process: low-entropy regions trigger fast, deterministic traversal; high-entropy regions pause traversal and generate verification requests.

The relationship between the two can be summarized as: the MSE architecture is the "brain"-level decision logic, while graph traversal is the "hands and feet"-level execution method. The former decides "whether to proceed," while the latter decides "how to proceed, which path to take, and how to record the paths taken." This distinction ensures that Logographic AI's reasoning is both theoretically rooted (MSE architecture) and practically executable (graph traversal).

Applying this theoretical framework to agentic materials science, its core advantage lies in: every materials design decision can be traced back to specific Morpho-Root attributes and causal chains, rather than a heatmap of attention weights. When an agent recommends "adding Ti to improve creep resistance," its reasoning path is transparent: Ti solute → solid solution strengthening Morpho-Root → inhibits dislocation motion → improves creep resistance. This is not statistical guessing, but structural traversal.

V. From "Tool Autonomy" to "Cognitive Rootedness"

Returning to the core concern of the review—how to advance agentic materials science from Level 0 to Level 5—Logographic AI theory suggests that the bottleneck on this path may lie not in the engineering implementation of "autonomy," but in the paradigm choice of "cognitive primitives."

The review's ultimate vision is for "human scientists and agents together to make materials discovery faster, more reliable, and more verifiable" [1]. Logographic AI theory offers a more fundamental path toward "more reliable" and "more verifiable":

First, **Let agents "grow" within physical logic, rather than "learn" to obey physical laws:** Physical laws are not contingent patterns induced from data, but unbreakable constitutive constraints of cognitive primitives [2]. In the design of material Morpho-Roots, the relation function R is presupposed and non-circumventable—the causal relationship between dislocation motion and plastic deformation is not something the agent "learns" from data, but something its cognitive primitives "grow" with.

Second, **Let reasoning become "structural traversal" rather than "statistical guessing" [2]:** Every materials design decision can be traced back to specific Morpho-Root attributes and causal chains, rather than a heatmap of attention weights. When an agent recommends "adding Ti to improve creep resistance," its reasoning path is transparent: Ti solute → solid solution strengthening Morpho-Root → inhibits dislocation motion → improves creep resistance.

Third, **Let "understanding" become an intrinsic structure of the agent, rather than an externally labeled "explainability patch" [2]:** The materials science agent can not only predict "Compound A has high catalytic activity," but also output a testable mechanistic explanation: "Because A's d-band center is optimal, and the active site density is high, it exhibits high activity." This is a verifiable scientific hypothesis, not a black-box probabilistic output.

VI. Conclusion: From "Autonomy" to "Rootedness"—The Paradigm Threshold for Agentic Materials Science

Academician Zhang Tongyi's team's review, with its rigorous hierarchical framework, provides a precise "status report" and "roadmap" for the development of agentic materials science [1]. However, from the perspective of Logographic AI, this roadmap still operates within an implicit

presupposition: that an agent's "autonomy" evolves linearly, and that with more powerful tools and more complete closed loops, the agent can progress from Level 0 to Level 5 [2][4].

Logographic AI theory seeks to challenge this presupposition: linear accumulation of autonomy cannot compensate for the "rootless" defect of cognitive primitives. The true breakthrough for agentic materials science may lie not in making agents more "autonomous" in invoking tools, but in making them more "rooted" in understanding the world—making physical laws the innate grammar of their cognitive primitives, rather than post-hoc rule patches.

When agents in materials science are no longer just "talking statistical models," but rooted agents whose cognitive primitives inherently embed physical meanings such as "dislocation motion," "solid solution strengthening," and "phase transformation thermodynamics," then—and only then—can we truly speak of the dawn of "AI materials scientists": partners that can not only execute tasks, but also understand the material world in ways comprehensible to humans, rather than intelligent tools wandering in statistical space [2][4].

References

- [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] Liu, S. (2026). The Rooted Intelligence: The Paradigm Revolution of Logographic AI. ChinaXiv. T202605.00745. <https://chinaxiv.org/businessFile/T202605/T202605.00745v1/T202605.00745v1.pdf> (English version retrieved from <http://logographica.com/Monograph.html>)
- [3] Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417-457.
- [4] Liu, S. (2025). Logographic AI vs. Phonographic AI: A New Paradigm and the Manifesto of AI Decolonization. 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>