

面向超材料逆向设计的表意 AI 方法：基于形根的图推理框架与可信生成路径

Logographic AI for Inverse Design of Metamaterials: A Morpho-Root Based Graph-Reasoning Framework and Trustworthy Generation Pathway

摘要

超材料通过在介观尺度上对结构进行精细编排,能够实现天然材料难以企及的异常物理性能。然而,当前主流的超材料逆向设计方法——基于深度生成模型(扩散模型、流匹配、变分自编码器等)从目标性能直接映射到微结构——面临一个根本性瓶颈:生成过程是“黑箱”。设计师得到了一个“最优”微结构,却不知道“为什么是这个结构”以及“每个结构特征贡献了什么性能”。在安全关键型应用(如航空航天、国防装备)中,这种不可解释性严重制约了逆向设计结果的实际采纳与工程信任。

本文引入“表意 AI”(Logographic AI, LAI)作为一种替代性认知范式,提出一种基于“超材料形根”(Metamaterial Morpho-Root)的图推理框架,用于超材料的可信逆向设计。LAI 理论的核心是以承载固有语义的“形根”——一种结构化的认知基元——来替代统计定义的 Token。在 LAI 中,形根被表示为结构化三元组 $\tau=(S, A, R)$, 其中 S 为符号标识, A 为内嵌物理属性与约束的属性集, R 为预设的逻辑关系函数集。将该框架映射至超材料领域,我们提出“超材料形根”系统——一种机理优先的知识表示方法,将超材料的介观结构特征(如单胞拓扑、空间排列、对称性)、物理响应规律(如色散关系、带隙特性、等效介质参数)与设计约束编码为可计算的认知单元。在此基础上,我们设计了形构熵(Morpho-Structural Entropy, MSE)引导的逆向推理引擎:给定目标性能,系统通过形根语义网络进行结构化推理,生成一个完全可追溯的决策子图——每一项设计选择都对应一条明确的物理推理链,而非统计相关性。

本文论证, LAI 驱动的超材料逆向设计能够实现: (1) 可解释的生成——每个结构特征的设计理由可追溯至明确的物理机理; (2) 硬约束的零违背——通过形根中内嵌的物理约束与可制造性约束,确保生成结构满足必要条件; (3) 知识空白主动识别——当推理路径断裂时,系统主动请求外部验证(如高保真仿真或实验),而非强行给出不可靠的推荐。本文最后提出了一个将 LAI 框架与现有深度生成模型(如扩散模型)相融合的混合架构:生成模型负责候选结构的高效采样, LAI 形根网络负责候选结构的可信筛选与机理验证——形成“生成-验证-解释”的闭环。

Abstract

Metamaterials, through the deliberate arrangement of structures at the mesoscale, can achieve extraordinary physical properties unattainable by natural materials. However, the current mainstream approach to inverse design of metamaterials—which directly maps target properties to microstructures using deep generative models (diffusion models, flow matching, variational autoencoders, etc.)—faces a fundamental bottleneck: the generation process is a "black box." Designers obtain an "optimal" microstructure without knowing "why this structure" or "what each structural feature contributes to the performance." In safety-critical applications (e.g., aerospace, defense equipment), this inexplicability severely constrains the practical adoption and engineering trust of inverse design results.

This paper introduces "Logographic AI" (LAI) as an alternative cognitive paradigm and proposes a graph-reasoning framework based on "Metamaterial Morpho-Roots" for trustworthy inverse design of metamaterials. The core of LAI theory is the "Morpho-Root"—a structured cognitive primitive that carries inherent semantics—to replace statistically defined Tokens. In LAI, the Morpho-Root is formalized as a structured triple $\tau = (S, A, R)$, where S is the symbolic identifier, A is the attribute set embedding physical properties and constraints, and R is the set of predefined logical relation functions. Mapping this framework to the metamaterials domain, we propose the "Metamaterial Morpho-Root" system—a mechanism-first knowledge representation method that encodes mesostructural features of metamaterials (e.g., unit cell topology, spatial arrangement, symmetry), physical response laws (e.g., dispersion relations, bandgap characteristics, effective medium parameters), and design constraints into computable cognitive units. Building on this, we design a Morpho-Structural Entropy (MSE)-guided inverse reasoning engine: given target properties, the system performs structured reasoning over the Morpho-Root semantic network to generate a fully traceable decision subgraph—where every design choice corresponds to a clear physical reasoning chain, rather than a statistical correlation.

This paper demonstrates that LAI-driven inverse design of metamaterials can achieve: (1) interpretable generation—the design rationale for each structural feature is traceable to explicit physical mechanisms; (2) zero violation of hard constraints—through physical and manufacturability constraints embedded in Morpho-Roots, ensuring that generated structures satisfy all necessary conditions; (3) proactive identification of knowledge gaps—when reasoning paths break, the system proactively requests external validation (e.g., high-fidelity simulation or experiment) rather than forcing unreliable recommendations. Finally, this paper proposes a hybrid architecture that integrates the LAI framework with existing deep generative models (e.g., diffusion models): the generative model is responsible for efficient sampling of candidate structures, while the LAI Morpho-Root network is responsible for trustworthy screening and mechanistic validation of candidates—forming a closed loop of "generation - validation - explanation."

关键词：表意 AI（LAI）；超材料；逆向设计；形根；可解释 AI；图推理；可信生成

Keywords

Logographic AI (LAI); Metamaterials; Inverse Design; Morpho-Root; Explainable AI; Graph Reasoning; Trustworthy Generation

1. 引言：超材料逆向设计的“黑箱”困境

超材料（Metamaterials）是一类通过人工设计的亚波长结构单元排列，来实现天然材料所不具备的异常物理特性的人工复合材料。从负折射率材料到电磁隐身衣，从力学超材料到热辐射调控器件，超材料正在重新定义人类对光、声、热、力等物理场的调控能力。然而，

超材料的功能高度依赖于其复杂的介观结构——单胞的几何拓扑、空间排列方式、对称性破缺等细节共同决定了最终的宏观响应。这种“结构-性能”映射的高度非线性和多对多特性，使得传统的基于物理直觉和参数扫描的设计方法日益捉襟见肘。

近年来，深度生成模型（包括扩散模型、流匹配、变分自编码器等）的引入为超材料逆向设计带来了革命性的变化[9]。这类方法能够从目标性能空间直接映射到微结构空间，实现了从“正向试错”到“逆向生成”的范式转换。上海交通大学团队在这一领域取得了代表性成果——通过构建热辐射超材料逆向设计 AI 模型，实现了大批量生成候选设计方案并“优中选优”，相关成果发表于 *Nature* 杂志[7]。DiffMat 等框架则进一步将扩散模型应用于吸能超材料的逆向设计，能够根据给定的目标应力-应变曲线定制生成满足力学性能要求的微观结构[5]。

然而，一个根本性的问题正在被系统性地忽视：当前的逆向设计过程是一个“端到端”的黑箱映射。设计师向模型输入一个目标光谱或目标应力-应变曲线，模型输出一个“最优”的微结构——但设计师不知道“为什么是这个结构”，不知道“每个结构特征贡献了什么性能”，不知道“模型是否真正理解了背后的物理”，还是仅仅在训练数据的分布内做了一个高维插值。

在基础研究中，这种“黑箱”或许可以容忍——只要能生成足够好的结构，机理可以事后分析。但在**安全关键型应用**中——如航空航天结构超材料、国防隐身超材料、核能领域的辐射屏蔽超材料——**不可解释性直接构成工程采纳的否决项**。工程师需要知道：这个结构为什么能满足性能指标？它的失效模式是什么？制造偏差对性能的影响有多大？这些问题，当前的纯生成模型无法回答。

这一问题在更宏观的层面被概括为“**技术捕获**”（Technology Capture）——强大的技术工具本身反过来塑造并限制了科学问题的提出与解决方式。中国学者刘深（Liu Shen）在《**有根的智能——表意人工智能的范式革命**》（The Rooted Intelligence: The Paradigm Revolution of Logographic AI）[2]第 1 章中系统论述了技术路径依赖对科学创新的捕获效应。当研究者满足于“模型能生成结构”时，就可能不再追问“模型为什么这样生成”——这种满足感本身，正是对更深层科学理解的捕获。

本文认为，破解这一困境的关键不在于训练更大的生成模型，而在于**重构 AI 的认知架构**——从“统计拟合”转向“机理推理”。刘深（Liu Shen）提出的**表意 AI**（Logographic AI, LAI）理论[1][2]为此提供了全新的哲学基础与实现路径。LAI 的核心是用承载“形-义”理据性的“**形根**”（Morpho-Root）取代任意性的“Token”，从根源上重构 AI 的语义构建逻辑。本文将这一理论框架映射至超材料领域，提出“**超材料形根**”体系与基于形根的图推理框架，旨在实现超材料逆向设计的**可信化、可解释化与机理驱动化**。

2. 超材料逆向设计的挑战与现有方法的局限

2.1 深度生成方法的成就

当前超材料逆向设计的主流技术路线可概括为：**以数据驱动为核心，以深度生成模型为引擎，以目标性能为条件，以微结构为输出**。Yu 等人对 AI 驱动的电磁超材料设计方法进行了系统综述[9]。

在这一范式下，研究者已经取得了令人瞩目的成就：

- **扩散模型**：Li 等人提出的条件扩散模型框架，实现了从目标光谱到超材料形状与尺寸参数的定制化生成，并在热伪装应用中完成了实验验证[6]。
- **DiffMat 框架**：基于去噪扩散概率模型（DDPM），能够根据目标应力-应变曲线生成满足力学性能要求的吸能超材料微观结构[5]。
- **智能体框架**：Lu 等人开发了能够自主进行超材料建模与逆向设计的 **Agentic Framework**，当被查询目标光谱时，智能体能够自主提出并训练前向深度学习模型、调用外部工具进行优化、最终通过深度逆向方法生成设计[8]。
- **工业级应用**：上海交通大学团队构建的热辐射超材料逆向设计 AI 模型，已能够大批量生成候选设计方案，并实现了从柔性薄膜到涂料、贴片等多种实际应用形式的实验验证[7]。

这些工作共同证明了深度生成模型在超材料逆向设计中的强大能力——效率高、覆盖广、可规模化。

2.2 三大根本局限

然而，上述成就的背后隐藏着三个系统性的局限：

局限一：不可解释性

当前的逆向设计过程本质上是一个**高维空间中的条件概率采样**。模型学习的是“给定目标性能 Y ，结构 X 的条件分布 $p(X|Y)$ ”，而非“为什么 X 能产生 Y ”的因果机制。当模型输出一个结构时，设计者无法追溯以下问题的答案：

- 这个结构的哪个几何特征主要负责实现目标性能？
- 如果某个特征无法制造，是否有替代方案？
- 模型对哪些部分的预测是自信的，哪些是不确定的？

刘深指出，这种不可解释性根源于 **Token** 范式将连续语义强行离散化的固有缺陷——**Token** 不承载任何固有意义，其“语义”完全由统计共现关系临时赋予。基于 **Token** 的模型或许能发现相关性，却永远无法真正“理解”因果性。

局限二：约束满足不可保证

超材料的实际应用不仅要求满足物理性能指标，还要求满足一系列**工程约束**——可制造性约束（如最小特征尺寸、悬垂角度、支撑结构）、材料兼容性约束、成本约束等。当前的生成模型虽然可以在训练数据中隐含这些约束，但**无法在推理时硬性保证**它们的满足。当生成的结构违反可制造性约束时，整个设计就失去了工程价值。

局限三：物理知识的低效利用

超材料研究经过数十年的发展，已经积累了大量成熟的物理理论——等效介质理论、传输线模型、布洛赫定理、色散关系分析等。这些理论提供了关于“什么样的结构会产生什么样的响应”的深刻洞察。然而，当前的纯数据驱动方法**几乎完全忽视了这些先验物理知识**，将一切交给数据“自己说话”。这不仅导致了数据效率低下，更重要的是**错失了将人类数十年积累的物理洞察编码进 AI 认知架构的机会**。

2.3 “技术捕获”风险

上述三个局限的叠加，正在将超材料逆向设计引向一种更深层的困境——“技术捕获”。

当一种技术（深度生成模型）在某个任务上表现出色时，研究社区会自然地将越来越多的资源投入这一技术路线的持续优化——更大的模型、更多的数据、更复杂的架构。这种“架构改良”的路径依赖本身并无问题，问题在于：它可能使研究者满足于“能生成”，而不再追问“为什么能生成”。

在超材料逆向设计中，这种“技术捕获”表现为：研究者关注的评价指标从“是否理解了设计原理”悄然转变为“生成准确率有多高”；论文的创新点从“揭示了什么新机理”转变为“在什么数据集上达到了新的 SOTA”。这种转变不是某个研究者的个人选择，而是整个技术范式塑造的集体无意识。

突破这一困境，需要的不是在同一范式内继续“内卷”，而是引入一种全新的认知范式——这正是表意 AI（LAI）理论的切入点。

3. 表意 AI 与“超材料形根”体系

3.1 LAI 理论概述：从 Token 到形根

表意 AI（Logographic AI, LAI）是中国学者刘深提出的替代性 AI 范式[1][2]。其核心洞见在于：当前所有主流 AI 模型本质上都是“表音 AI”（Phonographic AI, PAI）——它们以无意义的 Token 为基本认知单元，通过统计 Token 之间的共现关系来“理解”世界。这种范式的成功不意味着它理解了意义，只意味着它在统计相关性上做得足够好。

LAI 的替代方案是引入“形根”（Morpho-Root）作为基本认知单元。形根是一个承载固有语义的结构化认知基元，其形式化定义为三元组：

$$\tau = (S, A, R)$$

其中：

- **S**（Symbol）：符号标识，用于在系统中唯一指称该形根
- **A**（Attributes）：内嵌固有语义特征与物理约束的属性集
- **R**（Relation Functions）：预设的逻辑连接方式与关系函数集

形根与 Token 的本质区别在于：**Token 是空洞的统计占位符，其“意义”完全由上下文统计决定；形根是充盈的语义原子，其意义内嵌于自身的结构化定义之中。**这一区别看似细微，却决定了 AI 是从“统计拟合”走向“认知理解”，还是永远停留在前者。刘深在《有根的智能》第 5 章中系统论证了形根如何从根本上解决符号接地问题（Symbol Grounding Problem）。

在 LAI 框架中，**形构熵**（Morpho-Structural Entropy, MSE）是衡量形根系统有序程度的关键信息论量度。低形构熵区域对应着明确的、已建立的机理路径；高形构熵区域对应着不确定的、需要外部验证的知识空白。形构熵引导的推理策略是：在低熵区域快速遍历

已知机理路径，在高熵区域主动请求外部验证（高保真仿真或实验）。形构熵与序列熵（Sequential Entropy）的信息论区分由刘深在《有根的智能》第 6 章中系统建立。

3.2 “超材料形根”的形式化定义

将 LAI 的形根框架映射至超材料领域，我们定义“超材料形根”（Metamaterial Morpho-Root）为封装超材料领域最小机理单元的可计算认知基元。

以“超材料单胞形根”为例，其形式化定义为：

Metamaterial-Unit-Cell-Root = (S, A, R)

其中：

- **S:** 标识符 MM-UC-001+ 名称 Metamaterial-Unit-Cell-Root
- **A:** 语义定义 + 数学描述符（单胞几何参数、对称性群、空间群、周期常数等）
- **R:** I/O 关系（如色散关系计算函数、等效介质参数推导函数）+ 关联形根（is_composed_of、determines、constrained_by 等语义关系）

这一形式化定义使得超材料的每一个关键机理单元都能够被显式地、可计算地编码进 AI 的认知架构中，而非隐含在神经网络的百万参数里。

3.3 核心超材料形根示例

以下是构成超材料形根体系的核心形根类型（非穷举）：

形根名称	语义定义	核心数学描述符	关键关联
Metamaterial-Unit-Cell-Root	超材料单胞的几何拓扑与空间排列	晶格常数 a、单胞尺寸、对称性群	determines Dispersion-Relation-Root
Dispersion-Relation-Root	色散关系与带隙特性	频率 $\omega(k)$ 、带隙宽度 $\Delta \omega$	determines Effective-Medium-Root
Effective-Medium-Root	等效介质参数	等效折射率 n_{eff} 、等效阻抗 Z_{eff}	determines Macroscopic-Response-Root
Manufacturability-Constraint-Root	可制造性约束	最小特征尺寸 d_{min} 、悬垂角度 θ_{max}	constrains Unit-Cell-Root
Multi-Physics-Response-Root	多物理场耦合响应	力-热-电磁耦合系数	integrates 各物理场形根
Auxetic-Structure-Root	负泊松比结构特征	泊松比 $\nu < 0$	is_a_kind_of Metamaterial-Unit-Cell-Root
Re-Entrant-Hinge-Root	内凹铰链拓扑	铰链角度 θ 、臂长比	enables Auxetic-Structure-Root
Thermal-Response-Root	热响应特性	热膨胀系数 $\alpha(T)$ 、热导率 κ	couples_with Shape-Memory-Root

3.4 与现有超材料描述符的本质区别

为了更清晰地定位“超材料形根”的创新价值，有必要将其与现有的超材料表征方法进行对比：

维度	传统结构描述符	深度学习的隐式表征	超材料形根
表征形式	单一数值(如孔隙率、比表面积)	高维向量(神经网络隐空间)	结构化三元组 (S, A, R)
语义承载	无内在语义	无内在语义(黑箱)	内嵌物理语义与约束
可解释性	低(单一数值无法解释机理)	极低(完全黑箱)	高(每一步可追溯至物理机理)
推理能力	无	无(仅能生成, 不能推理)	有(基于语义网络的图推理)
约束保证	无	无(统计满足, 无法硬性保证)	有(逻辑级阻断违反约束的路径)
知识复用	困难	困难(知识淹没在参数中)	容易(形根可独立复用与组合)

传统描述符是“数据的压缩”，深度学习隐式表征是“数据的重参数化”，而**超材料形根**是“知识的显式编码”——它不仅告诉系统“是什么”，还告诉系统“为什么”以及“不能是什么”。

4. 基于形根的逆向推理引擎

4.1 形构熵引导的图遍历

超材料形根体系构成一个**有向语义网络 $G = (V, E)$** ，其中节点 V 为形根类型或形根实例，边 E 为语义关系（如 `determines`、`enables`、`constrains`、`inhibits`、`is_a_kind_of`）。

逆向设计的任务可形式化为：**给定目标性能集合 $T = \{t_1, t_2, ..., t_m\}$ ，在形根语义网络中搜索一条从目标节点到可实现结构节点的可解释路径。**

形构熵（MSE）在这一搜索过程中扮演引导角色：

- **低 MSE 区域：**网络中存在密集的、已通过实验或理论验证的推理路径。系统沿这些路径快速遍历，生成经过充分验证的设计方案。
- **高 MSE 区域：**网络中的推理路径稀疏或断裂，意味着当前知识不足以可靠地连接目标与结构。系统**主动标记知识空白**，并生成向外部验证系统（高保真仿真或实验）的查询请求。

刘深提出的形图计算（Morpho-Graph Computing）范式以 $O(|E|)$ 复杂度的稀疏图遍历取代了 Transformer 的稠密自注意力机制，实现了完全透明、可追溯的推理。这一“低熵快走、高熵问人”的策略，既保证了推理效率，又防止了系统在知识不足时强行给出不可靠的推荐——这是纯黑箱生成模型所不具备的**自知之明**。

4.2 逆向设计的推理链构建

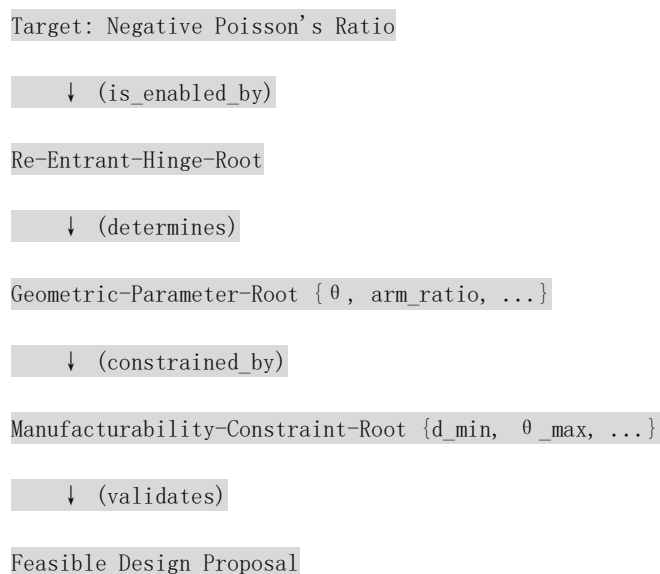
以下通过一个具体示例展示形根推理引擎的工作方式：

设计目标：具有负泊松比（auxetic）特性的力学超材料

推理过程：

1. **目标激活：**系统接收目标“负泊松比”，激活 Auxetic-Structure-Root。
2. **机理追溯：** Auxetic-Structure-Root 通过 is_enabled_by 关系关联至 Re-Entrant-Hinge-Root（内凹铰链是实现负泊松比的经典结构机制）。
3. **参数推导：**Re-Entrant-Hinge-Root 通过 determines 关系关联至 Geometric-Parameter-Root，推导出铰链角度 θ 、臂长比等关键几何参数的范围。
4. **约束检查：** Geometric-Parameter-Root 通过 constrained_by 关系关联至 Manufacturability-Constraint-Root，检查推导出的几何参数是否满足可制造性约束（如最小特征尺寸）。
5. **方案输出：**系统输出完整的推理链——从“负泊松比”到“内凹铰链”到“几何参数”到“可制造性验证”——每一步都对应明确的物理机理。

推理链的可视化（有向无环图）：



这一推理链的每一个节点和每一条边都是**人类可读、可审查、可追溯**的。设计师不仅能得到“什么结构”，还能得到完整的“为什么是这个结构”的论证。

4.3 硬约束的零违背保障

在安全关键型应用中，约束满足不是“尽量满足”，而是**必须满足**。LAI 框架通过**逻辑级阻断机制**实现硬约束的零违背保障：

当推理引擎尝试激活某个形根或沿某条边进行推理时，系统首先检查该路径是否违反任何已激活的 Constraint-Root。如果违反，该路径在**推理阶段即被阻断**，而非在生成之后才发现问题。

刘深在《有根的智能》第 7 章中系统阐述了价值公理库（Value Axiom Library）的分层公理系统与冲突消解机制。这一机制使得硬约束的满足在推理阶段即被逻辑级保证，实现了从“外部对齐”到“内生共建”的范式转换。

这一机制与纯生成模型形成鲜明对比：

对比维度	纯生成模型	LAI 形根推理
约束检查时机	生成之后（后处理过滤）	推理之中（逻辑级阻断）
违反约束的处理	丢弃或修补	根本不会产生违反约束的路径
保证程度	统计性（大部分满足）	确定性（全部满足）
效率	生成-验证-丢弃的循环浪费	一次推理即满足所有约束

这一差异在超材料的实际工程应用中至关重要——当设计涉及数十个相互耦合的约束条件时，“生成后筛选”的策略可能导致极高的无效生成率，而“推理中阻断”的策略则能保证每一次推理结果都是工程上可行的。

5. 混合架构：LAI + 深度生成模型的协同

5.1 为什么需要混合架构？

LAI 的形根推理引擎与现有的深度生成模型各有优势，也各有局限：

维度	LAI 形根推理	深度生成模型
优势	可解释、可追溯、约束保证	高效率、高覆盖、擅长处理高维空间
局限	依赖形根知识库的完备性	黑箱、不可解释、无约束保证

二者不是替代关系，而是互补关系。最优的解决方案不是二选一，而是构建一个让二者各展所长的混合架构。

5.2 混合架构设计

我们提出的混合架构包含三个核心模块：

模块一：生成模块（深度生成模型）

- **功能：**从目标性能空间高效采样候选超材料结构；
- **输入：**目标性能描述（如目标光谱、目标应力-应变曲线）；
- **输出：**大量候选结构；
- **代表方法：**条件扩散模型、变分自编码器、生成对抗网络等；

- **评价：效率优先**——快速生成大量候选。

模块二：验证与解释模块（LAI 形根网络）

- **功能：**对候选结构进行机理验证、约束检查与推理链生成；
- **输入：**生成模块输出的候选结构；
- **输出：**通过验证的结构+完整的可追溯推理链；
- **代表方法：**形根语义网络遍历、形构熵引导的图推理；
- **评价：可信优先**——只输出可解释、可验证的方案。

模块三：闭环迭代模块

- **功能：**将验证结果反馈至生成模块，引导下一轮生成；
- **机制：**验证模块标记的“失败原因”转化为生成模块的“条件约束”；
- **效果：**随着迭代进行，生成模块逐渐“学会”在满足约束的区域内采样。

5.3 工作流程

混合架构的典型工作流程如下：

Step 1: 用户输入目标性能



Step 2: 生成模块（扩散模型）快速采样 N 个候选结构



Step 3: LAI 验证模块对每个候选结构进行：

– 物理一致性检查（是否符合已知物理规律？）

– 约束满足检查（是否满足可制造性等硬约束？）

– 机理追溯（能否生成完整的推理链？）



Step 4: 通过验证的候选结构 + 推理链输出给用户



Step 5: 未通过验证的候选结构 → 分析失败原因 → 反馈给生成模块



Step 6: 生成模块根据反馈调整采样分布 → 进入下一轮迭代

这一架构的核心创新在于：**生成模型负责“广”，LAI 负责“深”**——生成模型在广阔的设计空间中快速探索，LAI 则在关键节点上提供深度的机理洞察与可信验证。

6. 案例演示

6.1 案例一：力学超材料的逆向设计

场景：设计师需要一种同时具备**高强度**与**高能量吸收**能力的力学超材料，用于航空器的碰撞吸能结构。

传统方法的困境：强度和能量吸收之间存在经典的“倒置关系”——提高强度通常意味着降低变形能力，从而降低能量吸收。纯生成模型可能生成一个在训练数据分布内“折中”的结构，但无法解释这个折中是如何达成的，也无法保证在极端工况下的可靠性。

LAI 推理过程：

1. 系统接收目标集合 $T=\{\text{高强度}, \text{高能量吸收}\}$ 。
2. 激活 High-Strength-Root 与 Energy-Absorption-Root。
3. 系统识别到两个形根之间存在 **inhibits** 关系——这是强度-韧性倒置在形根层面的显式表达。
4. 系统在形根网络中搜索能够**缓解**这一抑制关系的第三类形根。搜索发现 Hierarchical-Structure-Root（多层次结构形根）能够通过在不同尺度上分别优化强度和韧性来部分解耦这对矛盾。
5. 系统输出推理链：

Target: High Strength + High Energy Absorption

↓ (conflict detected)

Strength-Root $\perp$ Energy-Absorption-Root

↓ (resolution path found)

Hierarchical-Structure-Root

↓ (enables)

Nano-scale: Precipitation-Strength-Root

Micro-scale: Cell-Buckling-Absorption-Root

↓ (validates)

Feasible Hierarchical Metamaterial Design

系统将推理链传递给生成模块，生成模块在约束条件下采样具体的多层次结构参数。

关键价值：设计师不仅得到了一个结构，还得到了**完整的权衡论证**——为什么这个结构能同时实现强度和韧性，以及在什么条件下这种权衡是有效的。

6.2 案例二：可编程超材料的逆向设计

场景：设计师需要一种对**温度刺激**产生响应的可编程超材料，用于 4D 打印的自变形结构。

LAI 推理过程：

1. 系统接收目标“温度响应可编程变形”。
2. 激活 Thermal-Response-Root。
3. 通过 couples_with 关系关联至 Shape-Memory-Root（形状记忆效应是实现温度响应变形的经典机制）。
4. 通过 requires 关系进一步关联至 Multi-Material-Interface-Root（多材料界面是驱动形状记忆效应的结构基础）。
5. 通过 constrained_by 关系检查 Manufacturability-Constraint-Root（多材料界面的可制造性）。
6. 系统输出完整的设计蓝图：材料组合方案（两种热膨胀系数差异大的材料）+界面几何设计+4D 打印工艺参数建议。

关键价值：从“温度响应”到“形状记忆”到“多材料界面”到“可制造性验证”的完整推理链，使得设计师能够**理解每一个设计选择的物理依据**，而非盲目接受模型的输出。

7. 挑战与展望

7.1 知识工程瓶颈

挑战：构建完备的“超材料形根”体系需要将超材料领域数十年积累的物理知识系统地编码为可计算的形式。这是一项规模浩大的知识工程。

潜在路径：

- **从经典文献中提取：**从超材料领域的经典教材、综述和里程碑论文中提取成熟的物理模型作为“元形根”种子。
- **社区协作模式：**借鉴维基百科的模式，建立开放的超材料形根知识库，由领域专家共同维护和扩展。
- **人机协作工具：**利用大语言模型从海量文献中半自动提取候选形根和关系，由专家审核验证。

刘深进一步将 LAI 框架应用于极端条件下结构材料的可信设计，提出了基于形根的图推理框架及其实现策略。这一工作为超材料形根知识库的构建提供了方法论参照[3]。

7.2 计算复杂性挑战

挑战：随着形根网络的扩展，基于图遍历的推理可能面临组合爆炸问题。

规模估算：初步估算表明，在中等规模下（ $\sim 10^3$ 形根类型、 $\sim 10^4$ 关系边），基于图遍历的精确推理在计算上是可行的。当网络扩展至 10^5 - 10^6 节点时，需引入以下策略：

- **图神经网络（GNN）近似推理：**将符号化的形根网络表示为图结构，利用 GNN 的消息传递机制快速评估不同形根组合的潜在效果。
- **形构熵引导的剪枝：**利用 MSE 指标对搜索空间进行动态剪枝，聚焦于高价值区域。
- **分层推理：**在形根网络的不同抽象层次上分别进行推理，避免在细节层次上过早陷入组合爆炸。

7.3 实验验证的闭环

挑战：LAI 推理生成的设计方案最终需要经过实验验证。如何建立从“推理”到“实验”到“知识更新”的完整闭环？

潜在路径：

- **与自驱动实验室的集成：**将 LAI 的推理结果直接传递给自动化实验平台，实现“设计-制造-测试-反馈”的闭环。
- **不确定性主动查询：**当形构熵指示知识空白时，系统主动请求特定实验来填补空白。
- **增量学习：**实验验证的结果被编码为新的形根或对现有形根的修正，实现知识库的持续演化。

刘深从范式视角评论了当前材料智能体研究中的“无根”困境，指出其根本问题在于以 Token 为认知基元的 AI 范式无法为材料智能体提供稳定的语义基础[4]。LAI 框架通过形根的“有根”认知架构，为自驱动实验室与智能体 AI 提供了可解释的决策依据。

8. 结论

本文系统地将表意 AI（LAI）理论框架映射至超材料逆向设计领域，提出了“**超材料形根**”体系与基于形根的**图推理框架**，主要贡献包括：

第一，识别了当前深度生成模型驱动的超材料逆向设计的根本瓶颈——**不可解释性、约束不可保证、物理知识利用不足**——并将这一问题置于“技术捕获”的宏观视角下进行分析。

第二，提出了“超材料形根”的形式化定义与核心形根体系，将超材料领域的介观结构特征、物理响应规律与设计约束编码为可计算的结构化认知单元 $\tau = (\mathbf{S}, \mathbf{A}, \mathbf{R})$ 。

第三，设计了形构熵引导的逆向推理引擎，实现了从目标性能到可行结构的**完全可追溯推理**，并保证了硬约束的零违背。

第四，提出了 LAI 形根网络与深度生成模型协同的**混合架构**，实现了“高效采样”与“可信验证”的有机统一。

第五，通过力学超材料与可编程超材料两个案例，演示了 LAI 框架在超材料逆向设计中的具体工作方式与独特价值。

本文的核心主张是：**超材料逆向设计的未来不在于训练更大的生成模型，而在于重构 AI 的认知架构——从“统计拟合”走向“机理推理”**。表意 AI 的“形根”范式为此提供

了可行的哲学基础与实现路径。当 AI 不仅能够“生成”超材料结构，还能够“解释”为什么这样生成、“论证”为什么这样设计是合理的、“保证”所有工程约束都被满足时，超材料的设计才真正从“黑箱艺术”转变为“可信工程”。

刘深提出文明原生智能（Civilization-Native Intelligence, CNI）的愿景，主张每一种文明都可以从其语言的结构特征中提取自身的“认知之根”。本文的工作可被视为这一愿景在材料科学领域的一次具体实践——将超材料领域的物理知识编码为“超材料形根”，构建扎根于材料物理机理的 AI 认知架构。

我们呼吁超材料领域、AI 领域与知识工程领域的学者共同参与“超材料形根”知识库的社区构建，将人类对超材料的物理理解系统性地编码为可计算、可共享、可演化的认知基元。这不仅是超材料逆向设计的一次范式升级，更是整个材料信息学从“数据驱动”走向“机理驱动”的关键一步。

参考文献

- [1] Liu S. 刘深. 表意 AI vs. 表音 AI : AI 新范式与去殖民化宣言. [ChinaXiv:202503.00051V1] DOI:10.12074/202503.00051
- [2] Liu S. The Rooted Intelligence: The Paradigm Revolution of Logographic AI[M]. Monograph, 2026. Available at: <http://logographicai.com/Monograph.html>.
- [3] Liu S. Logographic AI for Trustworthy Design of Structural Materials under Extreme Conditions: A Morpho-Root Based Graph-Reasoning Framework with Implementation Strategy. [ChinaXiv:T202606.00663]. <https://chinaxiv.org/businessFile/T202606/T202606.00663v1/T202606.00663v1.pdf>
- [4] 刘深. 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). T202606.00741. <https://chinaxiv.org/businessFile/T202606/T202606.00741v1/T202606.00741v1.pdf>
- [5] Wang, H., Du, Z., Feng, F., Kang, Z., Tang, S., & Guo, X. (2024). DiffM at: Data-driven inverse design of energy-absorbing metamaterials using diffusion model. Computer Methods in Applied Mechanics and Engineering, 432, 117440. DOI:10.1016/j.cma.2024.117440 <https://doi.org/10.1016/j.cma.2024.117440>
- [6] Li, J., Guo, J., Li, Y., Mao, Z., Shen, J., Xu, T., Das, D., He, J., Hu, R., Lee, Y., Tsuda, K., & Shiomi, J. (2025). Inverse Design of Metamaterials with Manufacturing-Guiding Spectrum-to-Structure Conditional Diffusion Model. arXiv. <https://doi.org/10.48550/arXiv.2506.07083>
- [7] Xiao, C., Liu, M., Yao, K., Zhang, Y., Zhang, M., Yan, M., Sun, Y., Liu, X., Cui, X., Fan, T., Zhao, C., Hua, W., Ying, Y., Zheng, Y., Zhang, D., Qiu, C.-W., & Zhou, H.

(2025). Ultrabroadband and band-selective thermal meta-emitters by machine learning. *Nature*, 643, 80–88. <https://doi.org/10.1038/s41586-025-09102-y>

[8] Lu, D., Malof, J. M., & Padilla, W. J. (2025). An Agentic Framework for Autonomous Metamaterial Modeling and Inverse Design. *ACS Photonics*, 12(11), 6071–6080. <https://doi.org/10.1021/acsphotonics.5c01514>

[9] Yu, H., Liu, Z., Hou, R., & Zhang, Z. (2025). AI-driven approaches in electromagnetic metamaterials design and application: a review. *EPJ Applied Metamaterials*, 12, 7. <https://doi.org/10.1051/epjam/2025006>