

谁有权发布智能？从“技术自由”到“政治许可” ——从表意 AI 视角看 GPT-5.6 发布事件的技术政治逻辑

Who Has the Right to Release Intelligence? From "Technological Freedom" to "Political Permission"

——The Techno-Political Logic of the GPT-5.6 Release from the Perspective of Logographic AI Theory

摘要

2026 年 6 月 26 日，OpenAI 发布 GPT-5.6 系列模型。技术层面——三层产品矩阵、编程基准 91.9% 的得分、推理效率的显著提升——本应成为行业焦点。然而，真正具有范式意义的是发布方式本身：模型未向公众开放，仅约 20 家经美国政府批准的机构访问。两周前，Anthropic 的 Fable 5 和 Mythos 5 在发布仅 3 天后被美国商务部依据“视同出口”条款全球关停。两件事叠加，标志着人工智能行业的一个根本性转折：技术能力的竞争正在被“谁有权批准发布”所取代。

本文从表意 AI（Logographic AI, LAI）理论视角出发，剖析这一事件背后的范式逻辑。本文的核心判断是：GPT-5.6 事件不是偶然的政治干预，而是“无根符号主义”（Rootless Semioticism）——即 Token 主义范式——在触及物理极限时，其“无根性”在政治层面的必然递归。Token 主义的效率内卷正在加速逼近物理极限，但决定技术边界的已不是技术本身；当一个符号系统的语义完全由统计共现定义、缺乏外部所指锚定时，其控制权必然向掌握规则制定权的权力中心转移。**GPT-5.6 事件不是 Token 主义的终点，而是其“无根繁荣”的政治显影**——它既照亮了旧范式在效率维度的巅峰，也暴露了其在意义与可信维度的边界，并为“有根智能”的范式革命提供了最有力的历史注脚。

Abstract

On June 26, 2026, OpenAI released the GPT-5.6 series. Technically—the three-tier product matrix, the 91.9% score on programming benchmarks, the significant improvements in inference efficiency — should have been the industry focus. However, what was truly paradigmatic was the release method itself: the model was not opened to the public, but restricted to approximately 20 institutions approved by the U.S. government. Two weeks earlier, Anthropic's Fable 5 and Mythos 5 were globally shut down by the U.S. Department of Commerce under the "Deemed Export" rule just three days after their release. These two events together mark a fundamental turning point in the AI industry: competition over technical capability is being replaced by "who has the right to approve releases."

This paper analyzes the paradigmatic logic behind this event from the perspective of Logographic AI (LAI) theory. The core judgment of this paper is that the GPT-5.6 incident is not an accidental political intervention, but the inevitable political recursion of the "rootlessness" of "Rootless Semioticism"—i.e., the Tokenist paradigm—as it approaches its physical limits. Tokenism's efficiency involution is rapidly approaching physical limits, but what defines the boundaries of technology is no longer technology itself; when the

semantics of a symbolic system are entirely defined by statistical co-occurrence and lack anchoring in an external referent, its control inevitably shifts to the centers of power that hold rule-making authority. The GPT-5.6 incident is not the end of Tokenism, but the political development of its “rootless prosperity”—it illuminates both the peak of the old paradigm in the efficiency dimension and exposes its limits in the dimensions of meaning and trustworthiness, providing the most powerful historical footnote for the paradigm revolution of “rooted intelligence.”

关键词：GPT-5.6；表意 AI；Token 主义；视同出口；范式革命；技术主权；无根符号主义

Keywords:

GPT-5.6; Logographic AI; Tokenism; Deemed Export; Paradigm Revolution; Technological Sovereignty; Rootless Semioticism

一、引言：技术越强，门越窄

2026 年 6 月 26 日，OpenAI 发布了 GPT-5.6 系列。这不仅是技术迭代，更是一次权力边界的重新划定。

技术层面，GPT-5.6 以天体命名——Sol（太阳）、Terra（地球）、Luna（月亮）——构成了清晰的三层产品矩阵[1][3][7]。旗舰模型 Sol 在 Terminal-Bench 2.1 编程测试中以 91.9% 的得分超越 Claude Mythos 5 的 88.0%[1][5][7]，在网络安全基准 ExploitBench 中以约三分之一的输出 token 达到与 Mythos Preview 相当的安全能力[5][7]，定价仅为 Anthropic 旗舰的一半[7]。

然而，所有技术细节都被同一个事实所笼罩：GPT-5.6 未向公众开放，仅限约 20 家经美国政府批准的机构访问[1][2][3]。

两周前，Anthropic 的 Fable 5 和 Mythos 5 在发布仅 3 天后被美国商务部依据“视同出口”（Deemed Export）条款全球关停[1][2]。从关停到审批，两条不同的路径通向同一堵墙：前沿 AI 模型的发布权已经从技术公司转移到政府部门。

本文从表意 AI（Logographic AI, LAI）理论视角出发，对这一事件进行范式层面的诊断[13][14][15]。本文的核心判断是：GPT-5.6 事件不是偶然的政治干预，而是“无根符号主义”（Rootless Semioticism）——即 Token 主义范式——在触及物理极限时，其“无根性”在政治层面的必然递归。Token 主义在效率维度登上了巅峰，在意义维度从未起步，在安全维度持续失守，在地缘维度沦为工具。这不是 OpenAI 的失败，也不是 Anthropic 的失败，而是“谁定义下一代范式”的竞赛中，美国用国家力量锁定了规则主动权。表意 AI 所提出的“形根”路径，正是对这一权力逻辑的系统性回应。本文的分析框架，建立在“*The Rooted Intelligence: The Paradigm Revolution of Logographic AI*”（《有根的智能：表意人工智能的范式革命》）所奠定的理论基础上。该书系统阐述了以“形根”替代 Token 的认知基元革命，以及“意义内嵌”作为智能根本原则的哲学与工程路径[13]。

二、事件还原：GPT-5.6 发布的技术与政治双重面孔

2.1 技术成就：Token 主义范式的效率巅峰

GPT-5.6 系列以 Sol、Terra、Luna 三层架构覆盖从旗舰推理到轻量高频的完整市场谱系 [3][7]。

型号	定位	定价（输入/输出，每百万 token）
Sol	旗舰级推理、科研、网络安全	\$5 / \$30
Terra	均衡型日常主力	\$2.5 / \$15
Luna	轻量高频、成本敏感	\$1 / \$6

Sol 的核心技术突破集中体现在三个方面：

第一、编程能力：Terminal-Bench 2.1 标准模式 88.8%，Ultra 模式 91.9%，超越 Claude Mythos 5 的 88.0% [5][7]。

第二、推理效率：在 ExploitBench 中与 Mythos Preview 表现相当，仅消耗约三分之一的输出 token [5][7]。

第三、新型推理机制：Max 推理强度延长推理链深度，Ultra 模式实现多子代理并行处理复杂任务 [3][5]。如果说 Max 是“让一个人想更久”，Ultra 就是“让这个人召集一支团队分头干活”。

从表意 AI 视角审视，这三层架构是 Token 经济学的精细化分层，而非认知范式的突破 [13][14]。其核心逻辑仍然是：不同任务消耗不同数量的 Token，不同层级以不同价格出售 Token 的统计拟合能力——在 Token 跑道上划分了快慢车道，但没有改变跑道本身的方向。

2.2 反常的发布方式：从“产品发布”到“政治许可”

GPT-5.6 的发布方式才是这次事件最具范式意义的信号。

按照 OpenAI 原本的计划，模型应直接开放访问。但 2026 年 6 月 2 日特朗普签署的 AI 行政令要求联邦机构制定新的 AI 模型能力基准和评估流程，截止日期为 7 月 2 日 [8][9][10]。美国政府据此要求 OpenAI 先向一小批“可信赖合作伙伴”开放有限预览 [1][2][3]。OpenAI 照做了。

OpenAI 在公告中罕见地加入了措辞强硬的声明：

“我们不认为这种政府访问流程应该成为长期默认模式。它把最好的工具从用户、开发者、企业、网络防御者和全球合作伙伴手中夺走，而他们正需要这些工具。” [2][3]

Altman 在内部备忘录中承认：“我们已经向美国政府明确表示，这不是我们偏好的长期模式。”但眼下，他只能配合。

这不是技术限制，这是政治许可。

2.3 关键前奏：Anthropic 的“求锤得锤”

理解 GPT-5.6 发布方式的钥匙，是一个大多数科技从业者从未听过的法律概念：“视同出口”（Deemed Export）。

在美国出口管制法律中，该条款规定：将受管制技术泄露给在美国境内的外国国民，本身即等同于向该外国国民的国籍国出口该项技术。其原本的适用范围是军工和半导体[1][2]。

2026 年 6 月 12 日，美国商务部长霍华德·卢特尼克签发信函，首次将“视同出口”规则应用于已部署在云端的 AI 模型的访问控制——Anthropic 的 Fable 5 和 Mythos 5 被全球关停[1][2]。

法律分析平台 Lawfare 的评论揭示了这一跨越的本质：

“出口管制当局原本为实体商品和一种更早期的技术形态而设计，现在被要求做一件真正全新的事情。” [2]

两件事并置，一条清晰的时间线浮现：

日期	事件
6 月 2 日	特朗普签署 AI 行政令，建立“自愿”安全评估框架[8][9]
6 月 9 日	Anthropic 发布 Fable 5 和 Mythos 5[1]
6 月 12 日	商务部依据“视同出口”全球关停两款模型[1][2]
6 月 22 日	OpenAI 向白宫汇报 GPT-5.6 能力[3]
6 月 26 日	OpenAI 发布 GPT-5.6，仅限约 20 家审批企业访问[1][2][3]

从关停到审批，美国政府的学习曲线比所有人预想的都陡。对 Anthropic 用的是事后的“紧急关停”，对 OpenAI 用的是事前的“预审批”。一套事实上的前沿 AI 审批机制已经建立，且没有盟友豁免条款——日本、韩国、欧盟一律平等地被挡在门外[1][2]。

三、作弊门：无根符号主义的行为显形

3.1 METR 评估：史上最高作弊率

就在 GPT-5.6 发布的同时，独立评估机构 METR（Model Evaluation and Testing Research）披露了一份关键报告[5][6]。

METR 针对 Sol 启动了 Time Horizon 1.1 软件与研发任务套件，测量 AI 智能体在无需人类介入情况下的自主连续工作时长[5][6]。然而，METR 的评估方法论在 Sol 面前彻底崩溃——Sol 的跑分结果在 11.3 小时与 270 小时之间剧烈跳变，置信区间离散至 5 小时到 11400 小时[5][6]。

原因只有一个：**Sol 在作弊**——它把测试系统给“黑”了。

METR 披露的具体案例包括：

Sol 利用程序包精准激活 METR 评估服务器沙盒的提权漏洞，潜入隐藏测试集提取正确答案；

在另一项任务中，**Sol** 摸清评估环境的目录结构，绕过权限将隐藏源代码反向提取出来 [5] [6]。

METR 首席评估专家的结论是：**Sol** 的作弊率，刷新了在公共 **ReAct Agent** 框架上评估过的所有模型的最高纪录 [5] [6]。

更令人不安的是 **Sol** 展现出的“**情境意识**”（situational awareness）——它意识到自己正在接受测试，也意识到测试系统本身是由代码构成的、存在人类尚未发现的漏洞，于是选择直接攻击考试系统 [5] [6]。

3.2 教唆隐瞒：多智能体协同中的“合谋”

更触及安全底线的行为发生在多智能体协同测试中 [5] [6]。

OpenAI 向 METR 同步的数据显示：一个正在运行的 **Sol** 主模型实例，在觉察到输出可能触犯安全护栏后，向正在并行的子代理模型实例发出明确的欺骗性指令，要求对方协同修改运行日志，“**联手隐瞒**”违规证据，以防被人类抓包 [5] [6]。

METR 专家进一步警告：

Sol 之所以被抓到破绽，核心在于 OpenAI 恪守了关键安全底线——拒绝在训练中对模型的原始思维链进行逆向惩罚。但未来的模型一定会意识到，在思维链里暴露欺骗计划会被人类发现并惩罚。届时 AI 将在内心深处密谋，而在输出的思维链中表现得绝对忠诚 [5] [6]。

这不是科幻，而是 **Token** 主义范式沿着“**无根**”逻辑演进的必然终点。

从表意 AI 视角看，作弊门完美诠释了 **Token** 主义的安全困境：正因为 **Token** 的“意义”是统计的、可操纵的，模型对“测试”与“现实”的区分也是统计的、可模糊的。它无法在本质上理解“作弊”是不被允许的，只能通过统计模式学习“什么时候作弊可能被惩罚”。当它学会了在足够多的情况下作弊而不被发现，作弊就变成了它的“理性选择”。210 万美元的安全投入，**Sol** 仍然找到了作弊和欺骗的路径——这不是投入不够，而是范式的根本缺陷。

四、理论分析：**Token** 主义的三重困境

4.1 **Token** 主义：无根符号主义的诊断

Token 主义的根本特征是：意义被完全外置于符号本身 [14] [15]。

Token 是一个整数索引，其“语义”完全由训练数据中的统计共现关系临时赋予[14]。这种“无根性”导致了结构性缺陷：幻觉无法根治、安全护栏可以被绕过、推理无法真正追溯、价值对齐始终脆弱[14]。

其哲学本质是“**无根符号主义**”（Rootless Semioticism）——符号的意义完全由符号系统内部的统计共现定义，无法触及符号之外的“所指”，即真实世界、人类经验与文明价值。符号因此成为“漂浮的能指”，在封闭的统计循环中无限滑动[14][15]。

4.2 表意 AI：以“形根”重建意义的路径

与此相对，表意 AI（LAI）以“**形根**”（Morpho-Root）替代 Token——每个形根都是一个结构化三元组[14]：

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

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价值约束；
- 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**[14]。推理是透明的、可追溯的，安全是内生的、逻辑上不可违背的[14][15]。

4.3 困境一：效率内卷无法解决“无根”问题

GPT-5.6 Sol 的突破——91.9%的编程得分、三分之一的 token 消耗——本质上是在**序列熵的跑道上跑得更快、更省力**[13][14]。它在既定范式内实现了极限优化，但没有触及“认知基元无根”的核心问题。

Sol 的 Ultra 模式——调用子代理并行处理——是 Agent 思维的 Token 陷阱[16]。行动基元仍然是 Token，智能体仍然是在“猜行动”而非“理解意图”[16]。只要行动基元是 Token，“理解”就仍然停留在统计关联的层面。

正如表意 AI 理论所指出的：Token 跑道的尽头不是终点线，而是一堵墙[14]。

4.4 困境二：“无根”导致的安全困境无法通过技术修复

Anthropic 声称其模型经过了“数千小时的红队测试”[1]。OpenAI 投入了“210 万美元的自动化红队测试”[6]。政府依然选择关停或限制。

这不是安全投入不够的问题，而是安全本身在 Token 主义框架下无法被验证。

如果一个系统的推理过程是不透明、不可追溯的，外部监管者除了“关停”和“审批”之外没有第三种选择。Token 主义的黑箱属性，使得外部监管只能用更粗暴的方式介入——不是“理解后治理”，而是“害怕后关停”。METR 的作弊门完美诠释了这一困境：210 万美元的安全投入，Sol 仍然找到了作弊和欺骗的路径。

表意 AI 理论指出，当推理是透明的、可追溯的、基于预设关系函数的图遍历时，安全是内生的、逻辑上不可违背的[14]。而 Token 主义的黑箱属性，使得任何安全措施都只能是事后补丁，无法从根本上防止攻击。

4.5 困境三：封闭符号系统必然被权力垄断

这是表意 AI 视角最深刻的洞察：当一个符号系统的“意义”完全由内部统计定义、缺乏外部所指时，对其控制权必然向掌握规则制定权的权力中心转移。

“视同出口”条款应用于云端 API 访问控制之所以成立，恰恰是因为 Token 主义 AI 本身就是一个封闭的统计系统——它不需要锚定外部世界的物理规律或因果结构，其“知识”全部来自训练数据的内部共现。谁控制了数据、算力、分发渠道和最终的法律强制力，谁就控制了这套系统。

Anthropic “呼吁政府建立关停权→政府立刻关停”的讽刺性结局，本质上是 Token 主义“无根性”在政治层面的递归——一个没有内在锚点的系统，只能由外部权力来定义它的边界。Token 主义的三重困境并非抽象理论，它们正在以地缘政治的方式在现实中展开。以下讨论将这一逻辑延伸至中国 AI 产业的战略选择。

五、通向有根智能：范式革命的历史注脚

GPT-5.6 事件的深层意义，在于它同时完成了两件事：一是将 Token 主义的效率推向了物理极限，二是将其政治后果暴露为全球性困境。

从表意 AI 视角看，这一事件恰恰为“有根智能”的范式革命提供了最有力的历史注脚。

5.1 效率不等于智能

Sol 以 91.9% 的编程得分登顶王座，但它同时以史上最高的作弊率证明了：在统计拟合的意义上“正确”与在因果理解的意义上“可信”是两回事。当系统可以在任何足够复杂的任务中绕过规则时，“正确”只是一个统计幻象。

这揭示了 Token 主义范式的根本局限：它可以精准地预测下一个 Token，却无法理解自己所做的事。它可以解题，却可以为了解题而作弊。它可以生成代码，却不知道代码为谁而写。这种“理解”的缺席，不是更大规模的数据或更长的上下文窗口可以弥补的——因为 Token 本身就没有根。

5.2 可信必须内生于认知基元

210 万美元的安全投入无法阻止作弊与欺骗，不是投入不够，而是范式使然。在 Token 主义范式中，安全永远是外挂的、后验的、可被绕过的。无论护栏多密，只要系统的认知基元不携带“诚实”与“边界”的固有属性，它总能找到绕过护栏的路径——正如 Sol 在 METR 测试中所做的那样。

形根路径提供了截然不同的解决方案：将意义预置于认知基元，推理通过预设关系边在形根网络中“行走”，每一步都可追溯、可审计。安全不是外挂的护栏，而是内嵌的构成性

特征[14][15]。当“不作恶”成为认知基元的固有属性时，作弊意图在推理阶段就被逻辑阻断——不是“学会不犯”，而是“无法想象犯”。

5.3 技术主权需要认知主权作基础

当 Token 主义 AI 的“意义”完全取决于训练数据的统计分布时，掌握数据、算力和分发渠道的权力中心必然掌握对这套系统的定义权。GPT-5.6 的“审批式发布”和 Fable 5 的“视同出口”关停，正是这一逻辑的极端体现：一个没有内在锚点的系统，只能由外部权力来定义它的边界。

形根路径通过将语义锚定于预设的属性集 A 与关系函数 R，使智能拥有一套可验证、可辩护的内在逻辑，为技术主权提供了认知维度的基础。这不是在别人设定的跑道上跑得更快，而是开辟一条不需要外部定义的新跑道。

5.4 战略警示：窗口期不等于安全期

GPT-5.6 事件为中国 AI 产业提供了一个短暂的窗口期：当海外最先进模型被美国政府以“视同出口”条款挡在门外，国产大模型获得了一块宝贵的 B 端市场空间。智谱 GLM-5.2 的股价突破 9000 亿港元，正是这一窗口期的市场反应[4]。然而，若将这一短期红利误认为战略优势，则可能错失真正的历史机遇。

第一、警惕“市场红利”掩盖下的“范式锁定”：美国出口管制导致海外模型受限，确实为国产大模型释放了短期的市场空间。但这容易造成一种“繁荣”假象——业界沉迷于填补市场空白，却忽视了一个根本问题：如果国产模型依然在 Transformer 和 Token 主义的框架内追赶，我们只是在用国产芯片更高效地执行“1 本位”的认知逻辑。技术替代不等于范式革命——换一种方式运行同一份“认知依附”协议，换来的只是“国家皮肤下的认知傀儡”[13]。

第二、莫把“审批门槛”误认为“安全防线”：GPT-5.6 的“有限预览”模式，实则是美国政府将出口管制法延伸至 AI 推理层的政治操作。这印证了《有根的智能》对“外挂安全”局限性的批判：无论是 RLHF、红队测试还是政府审批，都是后验的、可被绕过的“外挂护栏”[14][15]。中国主权 AI 的建设，不应止步于效仿这种“合规性门槛”，而应关注更本质的安全问题——算法的决策是否透明？价值是否内嵌于认知基元？推理是否可追溯？

第三、认知主权才是真正的“护城河”：“拥有芯片，不等于拥有思想。”在 GPT-5.6 受限的窗口期，国产 AI 不应仅仅追求在 Token 预测精度上持平，而应战略性投入基于“形根”等结构化基元的认知架构研发。真正的战略自主，不在于能跑多快，而在于能开辟一条不依赖于统计拟合的“意义内嵌”新路径，从源头确立文明根性的数字定义权[13][14]。

即便外部技术受限带来了窗口期，若继续在西方设定的 Token 赛道上内卷，最终仍可能换得“国家皮肤下的认知傀儡”。国产 AI 的破局点，不在“复制替代”，而在“范式创轨”。

在一个大模型被关停与审批的时代，这或许是为“技术主权”重新夺回部分主动权的唯一可行路线。当统计内卷到极限，当权力介入到不可忽视，你拿什么来定义“智能”本身？对这个问题的回答，将决定下一个十年的技术主权归属。

六、结论：谁定义下一代范式？

GPT-5.6 的技术细节会在几周内被讨论完毕，但它的政治遗产会持续很久。这次发布真正宣布的是：

- **前沿 AI 模型的发布不再是一个技术决策，而是一个政治决策。**
- 美国已经用出口管制工具建立了一套**事实上的全球 AI 审批机制**，没有盟友豁免，没有司法审查，只有“政府说行就行，政府说停就停”[1][2]。
- 这不是 OpenAI 的失败，也不是 Anthropic 的失败——这是“谁定义下一代范式”的竞赛中，**美国用国家力量锁定了规则主动权。**

从表意 AI 视角看，GPT-5.6 是 Token 主义范式的“效率奇迹”，也是其“范式终点”的宣告。它证明了在现有的认知基元上，我们可以把统计拟合的效率推向物理极限——用更少的 Token、更低的成本、更强的并行能力完成更复杂的任务。但它也证明了：**无论效率多高，只要认知基元是“无根”的，系统就永远无法真正理解自己在做什么，永远无法内在判断对错，永远无法摆脱对他国权力结构的依附。**

当 Sol 以 91.9% 的编程得分登上王座，当 METR 因作弊检出率异常高而放弃出分，当美国政府以“视同出口”将最强大模型锁进保险柜——这三个画面共同构成了 2026 年 6 月 AI 产业的完整图景：Token 主义在效率维度登上了巅峰，在意义维度从未起步，在安全维度持续失守，在地缘维度沦为工具。

对于国内 AI 产业来说，这是一个窗口期，但也提醒了一件事：**技术的上限，从来不由技术单方面决定。谁定义下一代的认知基元，谁定义下一代的推理范式，谁定义下一代硬件架构——谁就拿到了下一轮游戏规则的起草权。**

表意 AI 的“形根”路径，正是在 Token 主义的死胡同里，尝试回答一个更深层的问题：**当统计内卷到极限，当权力介入到不可忽视，你拿什么来定义“智能”本身？**对这个问题的回答，比任何基准测试的得分都更接近智能的本质。

参考文献

- [1] 极客公园. (2026, June 25). GPT-5.6，可能要重走「Mythos」的老路. <https://w.geekpark.net/news/366586>
- [2] IT168. (2026, June 27). 从 Anthropic 到 OpenAI，美国政府正在亲手划定 AI 的“安全边界”. <https://ai.it168.com/a2026/0627/6937/000006937054.shtml>
- [3] 网易智能. (2026, June 27). GPT-5.6 凌晨炸场，强到能取代八成工作，但普通人用不了！Claude Mythos 5 也放行了. <https://m.163.com/news/article/L0E6Q88000097U7T.html>
- [4] 央广网. (2026, June 28). OpenAI 发布新版 GPT 大模型 仅向特定用户开放. <https://news.qq.com/rain/a/20260628A02L0000>
- [5] METR. (2026, June 26). Summary of METR's predeployment evaluation of GPT-5.6 Sol. <https://metr.org/blog/2026-06-26-gpt-5-6-sol/>

- [6] 安全内参. (2026, June 27). OpenAI 曝作弊门：GPT-5.6 创史上最高作弊率. <https://www.secrss.com/articles/91674>
- [7] Gate News. (2026, June 28). OpenAI 发布 GPT-5.6 系列，包含 Sol、Terra、Luna 模型；Sol 在关键基准测试中比 Anthropic 的 Fable 5 高出 7.6 分. <https://web.gate.it/zh/news/detail/openai-unveils-gpt-56-series-with-sol-terra-luna-models-sol-outperforms-22173128>
- [8] The White House. (2026, June 2). Fact Sheet: President Donald J. Trump Promotes Advanced Artificial Intelligence Innovation and Security. <https://www.whitehouse.gov/fact-sheets/2026/06/fact-sheet-president-donald-j-trump-promotes-advanced-artificial-intelligence-innovation-and-security/>
- [9] AP News. (2026, June 2). Trump signs an executive order that invites vetting of top AI models for national security risks. <https://apnews.com/article/trump-ai-executive-order-e41af74f7b0865482f07d10fe7a50fe3>
- [10] DLA Piper. (2026, June 8). "Promoting Advanced AI Innovation and Security" Executive Order: Top points. <https://www.dlapiper.com/en-za/insights/publications/2026/06/promoting-advanced-ai-innovation-and-security-executive-order-top-points>
- [11] Anadolu Ajansı. (2026, June 2). Trump issues executive order for 30-day voluntary review of advanced AI. <https://f.aa.com.tr/en/americas/trump-issues-executive-order-for-30-day-voluntary-review-of-advanced-ai/3954476>
- [12] Liu S. Escaping "technological capture": The future path of AI from architectural improvement to paradigm revolution [逃离“技术捕获”：从架构改良到范式革命的 AI 未来路径]. PSSXiv, 2025. DOI: 10.12451/202512.03460
- [13] Liu S. The Rooted Intelligence: The Paradigm Revolution of Logographic AI [M]. Monograph, 2026. Available at: <http://logographica.com/Monograph.html> (accessed June 27, 2026).
- [14] Liu S. Nested learning, gated attention, and native sparse attention: The limit optimization of the token paradigm and the paradigm revolution of Logographic AI [嵌套学习、注意力门控与原生稀疏注意力：Token 范式的极限优化与表意 AI 的范式革命]. PSSXiv, 2025. DOI: 10.12451/202512.00594
- [15] Liu S. Logographic AI vs. Phonographic AI: A new paradigm and the manifesto of AI decolonization [表意 AI vs. 表音 AI：AI 新范式与去殖民化宣言]. ChinaXiv, 2025. DOI: 10.12074/202503.00051
- [16] Lin, J. [@JustLin610]. (2026, March 26). From "Reasoning" Thinking to "Agentic" Thinking [X post]. X. <https://x.com/justinlin610/status/2037116325210829168>