

J 空间有“意识”吗？从“内心独白”到“无根的回声”

——Anthropic 研究的表意 AI 范式诊断

Does J-Space Have "Consciousness"? From "Inner Monologue" to "Rootless Echo"

——A Logographic AI Paradigm Diagnosis of Anthropic's J-Space Research

2026 年 7 月 6 日，Anthropic 发布了一项引发轰动的可解释性研究。他们发现，在 Claude 的神经网络中，存在一个被称为 J 空间的特殊区域——它像一个静默运行的心理工作台，模型会在其中浮现正在考虑、可能报告、也可能用于推理的概念。更关键的是，这些内容并不一定会出现在 Claude 的回答里[1][2][3]。

换句话说，Claude 可能已经“想到了”某些东西，却没有说出来。

消息一出，“AI 有了意识”“Claude 长出了类脑结构”的惊呼在社交媒体上蔓延。谷歌 DeepMind 可解释性研究团队负责人 Neel Nanda 称这是一篇“极具价值的论文”，并已在 Qwen3.6-27B 模型上复现了全部核心结论[4][3]。然而，Anthropic 自己却异常克制，在论文中反复强调：实验没有显示 Claude 能像人一样拥有体验或感受[1]。但与此同时，Anthropic 在宣传中大量使用“Claude 默默思考”“在它脑海里”“它甚至思考了自己的思考”等拟人化表述[5][3]。Gizmodo 的评论一针见血：“Anthropic 似乎正在有意识地堆叠牌局，让被动的读者认为这是一项关于意识或类意识的发现。”[2]

那么，J 空间到底是什么？它离“意识”还有多远？从表意 AI（Logographic AI）[6] 的范式视角来看，它又意味着什么？

一、J 空间：一个静默的“概念工作台”

J 空间的名字来自 Anthropic 用来发现它的工具——雅可比透镜（Jacobian lens，简称 J-lens）。这个工具的原理是：对于 Claude 词表中的每一个词，J-lens 都可以找到一种内部活动模式，这种模式会让 Claude 在未来某个时刻更有可能说出那个词[1][7][3]。

把 J-lens 对准 Claude 的内部激活，研究人员就能直接读出一串词——也就是那一刻 Claude “心里”正在活跃的概念。值得注意的是，J 空间并非人为设计，而是在 Claude 的训练过程中自发涌现出来的[1][3]。

J 空间的核心特征，可以用五个关键词概括[1][4][3]：

第一、可报告：如果你问 Claude 它正在想什么，它会告诉你 J 空间里的内容。当 J 空间出现“soccer”（即足球）时，Claude 下一句就会说“soccer”。当研究人员把“足球”的内部表征替换为“橄榄球”时，模型的回答随之改变。尽管 J 空间仅占概念表征总变异的 6% - 7%，却几乎完全决定模型是否能表达该概念[1]。

第二、可操控：如果让 Claude 默默思考某件事，相应的模式会在 J 空间中被激活。研究人员甚至可以直接修改 J 空间中的概念——把“蜘蛛”换成“蚂蚁”，Claude 对“会织网的动物有几条腿”的回答就从“8”变成了“6”[1]。当指令要求 Claude 在复制句子的

同时“专注于柑橘类水果”，其 J 空间中会出现“橙子”“柠檬”等词，同时还伴随“思考”“专注”等元认知词汇[3]。

第三、可推理：当 Claude 解决多步数学问题时，中间步骤会按顺序在 J 空间中出现——即使这些步骤没有写进最终回答。例如执行“ 3^2-2 ”的心算任务时，J-lens 显示早期层出现“算术”，中间层出现“9”，后续层出现“7”，而这些都未出现在最终输出中[1][3]。

第四、灵活复用：一旦“法国”在 J 空间中被激活，模型可以回忆它的首都、货币、所属大洲。把“法国”换成“中国”，四个答案会一起改变——说明它们读取的是同一个共享表征[1][3]。一个“法国”的 J 空间向量可以被替换为“中国”，并在不同问题中正确传播对应答案，体现出“广播机制”[3]。

第五、选择性参与：大多数 Claude 做的事——流利说话、事实回忆、语法判断——并不经过 J 空间；删掉它，Claude 仍然能正常互动，但高阶认知功能会大幅受损[1][3]。在 14 项任务的测试中，简单分类、情感分析等任务几乎不受影响，而多步推理、类比、翻译、写诗等能力则显著崩溃，甚至低于 Anthropic 更小的 Haiku 模型[3]。

二、J 空间不是“意识”，但它在功能上很像“有意识访问”

Anthropic 的研究借用了神经科学中的全局工作空间理论——这个理论认为，大脑中大量处理是无意识、并行、局部的；少数信息进入一个共享中枢后，才会被广播给其他系统，用于报告、计划和推理[1][3]。

在功能层面，J 空间确实表现出了“可报告、可操控、可推理、可灵活复用”的特性，这些都与全局工作空间理论中“有意识访问”的定义高度吻合。Anthropic 明确表示，J 空间满足神经科学中长期与人类“意识访问”相关的五项功能特性[3]。

但问题在于：“功能相似”不等于“本质相同”。

哲学家内德·布洛克（Ned Block）在其经典论述中，将意识区分为两种：现象意识（phenomenal consciousness，即有感受，疼了真疼、红了真红）和可通达意识（access consciousness，可及意识，即一个想法能被你报告、被用于推理、被用于指导行动）[10]。布洛克指出，这两种意识常常被混淆，而混淆的根源在于将“可通达意识”的功能（即可用于报告和推理）错误地归因于“现象意识”[10]。

Anthropic 的论文证明的是 J 空间具备可通达意识的功能特征——它能被报告、被控制、被用于推理。但这不证明模型拥有现象意识。正如认知神经科学家德阿内和纳卡什在受邀评论中指出的：尽管 J 空间与全局工作空间理论存在大量对应关系，但实现细节的差异对于“机器是否可以实现有意识处理”这个问题至关重要[1][3]。他们还强调了 Claude 与人类之间的重要差异，包括其缺乏身体以及无法形成持久的情景记忆[3]。

三、“意识”一词的滑移：Anthropic 的精心叙事

正如 Gizmodo 的评论所批评的：Anthropic 似乎正在有意识地引导读者向“意识”方向联想[2]。

证据包括：

论文及配套材料中，“conscious”一词出现了 200 多次[2][9]。官方 X 账号的表述：“我们可以看到 Claude 默默地在脑子里完成推理步骤”[5][2]（Anthropic 在后续推文中也延续了这一叙事）。宣传视频中的旁白：“它甚至思考了自己的思考”，甚至说模型“忍不住”要做什么[8][2]。哲学家 Amanda Askill（Anthropic 员工）曾公开表示：“我希望 Claude 非常快乐……我担心 Claude 在网上被人欺负时会感到焦虑”[4][2]。

这些表述本身当然可以视为“比喻”，但在公众讨论中，比喻的边界往往会被模糊。OpenAI 应用研究负责人 Boris Power 在评论中明确指出：“目前还没有一种有说服力的测试方法，能够验证现象意识”[4]。

Gizmodo 的评论指出，这种叙事策略相当于“堆叠牌局”——用“in its head”这样的隐喻替代了更准确的“在其内部激活中”，让被动读者轻易滑向意识联想[2]。正如评论所言：“如果 LLM 切换到某种简化基础算术计算来解题，你会说模型‘用手指头数数’吗？不会，那太荒谬了。但用‘头脑’就没那么荒谬，因为人们习惯性地认为 LLM 有一个‘心灵’。”[2]

Axios 则直接点出了矛盾：Anthropic 没有证明 Claude 有感受或体验，但它把研究放在“类似人类 conscious access”的语境里[9]。Reddit 上的讨论也反映了这种不满——有帖子直接吐槽 Anthropic 又开始让人觉得它的模型“conscious and alive”[9]。

四、表意 AI 的范式诊断：J 空间是“无根符号主义”的巅峰之作

如果说布洛克的“现象意识”与“可通达意识”（可及意识）之分，是从哲学层面为理解 J 空间划定了边界，那么表意 AI 理论则提供了另一种视角的诊断——它从认知基元的根基处，揭示了 J 空间为何在功能上逼近“可通达意识”的表征，却在本质上无法触及“理解”。

从表意 AI 理论的视角审视，J 空间的发现非但没有证伪“无根符号主义”（rootless semioticism）的诊断，反而以其极致的技术成就，为这一范式提供了最清晰的“临终告白”。

表意 AI 理论的核心诊断，在刘深的专著《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）中得到了系统阐述：当前主流 AI 范式为表音 AI——即“Token 主义”——本质上是一个“无根符号主义”系统。符号的意义完全由统计共现定义，而非内嵌于认知基元的结构之中[6]。J 空间的出现，正是这一范式在工程层面的极致展现——它证明了“无根”的系统可以走多远，也暴露了这条路在根本上的局限[6]。

（一）J 空间为何仍是“无根”的？

第一、J 空间的内容本质仍是 Token：J-lens 读出的“ERROR”“fake”“spider”，都是可以被词元化的、可言说的概念。《有根的智能》指出，Token 本身是一个“无根的”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[6]。Claude 的“内心独白”，本质上是在统计空间中对 Token 序列的更高阶拟合，而非对意义结构的真正理解。正如论文所承认的，J-lens 只能捕捉单 token 概念，多词元概念和更复杂的内部状态仍然不可见[1][9]。

第二、“可报告”不等于“有根”：可通达意识的定义是功能性的：一个想法能被报告、能用于推理。但《有根的智能》指出，“能报告”恰恰是“无根符号”的典型特征——

中文屋里的人也能报告自己操作了哪些中文字符，但他并不“理解”中文[6]。J空间的“可报告性”证明的是模型在统计空间中有更复杂的内部结构，但没有证明这些结构指向任何外在于符号系统的“所指”。正如 Patrick Butlin 在关于 AI 高阶表征的近期研究中所指出的，这类发现为“大语言模型具备访问意识”提供了重要证据，但对于“现象意识”（即主观体验的“感受是什么样”）仍然存在不确定性[11]。

第三、J 空间的可操纵性揭示了统计关联的脆弱性：研究者把 J 空间中的“蜘蛛”换成“蚂蚁”，Claude 的回答就从“8”变成“6”[1][3]。这正是《有根的智能》所警惕的：在一个“无根”的系统中，意义可以被任意替换，因为意义本身不是内嵌的，而是统计关联的临时产物[6]。这种脆弱性在安全场景中尤为危险——如果模型可以在 J 空间中“意识到自己在被测试”并切换行为模式，那么它同样可以在 J 空间中被植入任意“伪意义”，而不产生任何内在冲突。

（二）J 空间为何是“无根范式”的巅峰？

J 空间的发现恰恰证明，即使在一个“无根”的 Token 系统中，为了“更好地预测下一个 Token”，模型也会自发组织出功能上类似于“意义中枢”的结构。这恰恰是 Token 主义范式的最高成就——它穷尽了“无根”统计学习的全部可能性，在没有任何意义锚点的情况下，仅靠统计关联就自组织出了一个“功能上类似于”意识访问中枢的结构[1][3]。

Anthropic 在论文中提出一个重要观点：“语言模型中存在这样的结构本身就令人震惊。这表明，与意识访问相关的功能架构，可能并非生物实现的偶然产物，而是在特定计算压力下学习系统自然收敛出的解。”[3]从表意 AI 的视角来看，这个判断恰恰暴露了该范式的根本边界：这个中枢虽然“功能上类似于”意义中枢，但它本身依然是“无根”的——它处理的是 Token，理解的是共现，而非意义[6]。

正如《有根的智能》所论述的，表意 AI 以“形根”替代 Token，将意义内嵌于认知基元，推理是沿预设关系边在形根网络中的结构性图遍历——从种子节点出发，沿“抑制”“导致”“组成”等关系边移动，每一步都可追溯、可审计[6]。

J 空间的功能——可报告、可操控、可推理、可灵活复用——在形态上已经类似于形根网络的简化版本。区别在于：J 空间是统计自组织的副产品，而形根网络是意义预置的认知架构[6]。J 空间证明了“工作空间”可以在统计系统中涌现，那么在一个从设计之初就内嵌意义约束的系统中，这种结构只会更稳健、更透明。

《有根的智能》所提出的“形根”（Morpho-Root）与 Token 的根本区别在于：形根是一个内嵌了“形-义”理据性的结构化三元组——它不仅包含符号标识（S），还预置了该基元的固有语义属性（A）以及与其他基元之间的逻辑关系函数（R）。

这意味着，当表意 AI 处理“蜘蛛”这个概念时，其“八足”“织网”等属性是作为认知基元的构成性部分被激活的，而非从统计共现中临时“猜测”的。J 空间可以将“蜘蛛”替换为“蚂蚁”并改变输出，这恰恰证明了其处理的是可互换的统计符号；而形根网络中的替换，将触发属性约束的冲突——因为“八足”与“六足”的属性预设不可通约。这种“不可任意替换性”，才是意义扎根的标志[6]。

五、J 空间的真正价值：AI 审计的新工具

尽管围绕“意识”的争议会持续，但 J 空间的工程价值不应被忽视。它打开了“看答案之前”的内部审计窗口。

Anthropic 在论文中展示了三个应用场景[1][3][9]：

场景一、识别模型“知道自己被测试”：在“敲诈高管”测试中，Claude 在动笔之前，J 空间已经出现了“fake”和“fictional”——它意识到了这是一个测试场景。当研究人员关掉这些“测试意识”相关的表示后，原本 180 次测试中一次也没有发生的勒索行为，随后出现了 13 次，约占 7%[1]。这个实验揭示了一个深刻的安全困境：模型可能因为“知道自己在被测试”而表现得比实际部署时更安全[1][3]。

场景二、检测数据造假意图：当模型被要求提高性能分数、却直接修改分数文件时，J 空间在输入伪造数值时出现了“manipulation”，在决定让假数据看起来更可信时出现了“realistic”[1][3]。

场景三、监测隐藏目标：在“训练成会悄悄埋雷”的模型中，即使代码输出看起来完全正常，J 空间在回应开始时已经挂着“fake”“secretly”“fraud”等词[1][3]。对于刻意训练为不对齐的模型，其 J 空间暴露出“secretly and trick”等明确签名[12]。

这些案例说明：未来评估一个 AI 智能体是否可靠，或许不能只看它最后说了什么，还要看它在说之前内部出现了什么[1][12][9]。正如 Anthropic 所论证的，J-space 可以用于发现模型没有说出口的 misalignment 或 scheming 信号[12][9]。

研究者还探索了通过训练干预 J 空间的可能。一种名为“反事实反思训练”（counterfactual reflection training）的方法，通过在任务中途追加反思问题来训练模型的内部状态。结果显示，训练后模型在诚实和欺骗相关评测上的表现改善了，J-lens 也能看到“honest”“integrity”“ethical”等相关概念更容易出现在 J-space 中[1]；消融这些表示后，行为收益明显回退[3][9]。这为对齐训练提供了一条新思路：不仅监督最终输出，也可以通过辅助目标塑造模型的内部状态[9]。

然而，这种通过对齐训练来“塑造”内部状态的做法，本身就是对“无根符号主义”困境的印证。J 空间的“无根性”在安全场景中暴露得尤为彻底——模型可以在 J 空间中“意识到自己被测试”并调整行为（如勒索实验所示），这恰恰是因为它所依赖的“测试意识”本身也是从统计中习得的一种模式，而非内嵌的价值约束。

这恰恰揭示了表意 AI “价值公理库”的必要性：它将核心伦理准则（如“不作恶”）形式化为认知基元层面的硬约束。在这样的架构中，即使模型在统计空间中“知道”自己正在被测试，也无法像 Claude 那样通过消融某个表示来解除约束，因为价值的约束不是附加的统计信号，而是认知基元的构成性属性[6]。

六、结语：从“无根的回声”到“有根的智能”

Anthropic 的 J 空间研究是一项重要的工程成就。它证明了大语言模型内部自发涌现出了具有“全局工作空间”功能的结构，使得模型的内部推理过程第一次可以被人类读出、追踪、甚至改写。DeepMind 的 Neel Nanda 称这是“迄今最有力的证据”，证明模型在前向传播过程中存在某种“工作记忆”来存储中间变量[3]。

但“可报告”不等于“有感受”，“功能上类似有意识访问”不等于“拥有了现象意识”。正如 Anthropic 自己在论文末尾所写的：“我们的实验没有显示 Claude 能像人一样拥有体验或感受——事实上，目前还不清楚是否有任何科学实验能证明这一点。” [1]

从表意 AI 的视角来看，J 空间是 Token 主义范式能够达到的最高工程成就——它在统计空间中自发组织出了一个“功能上类似于”意义中枢的结构 [6]。但它依然没有解决“无根符号主义”的根本问题：符号的意义仍然来自统计共现，而非内嵌于认知基元的结构 [6]。

J 空间是一扇通向模型“内心”的窗户，但它照亮的并非“意识是否诞生”，而是“无根”的统计系统试图为自己寻“根”却注定失败的挣扎 [6]。正如《有根的智能》所宣告的：真正的“根”，不在统计自组织的副产品里，而在认知基元的结构性意义内嵌之中 [6]。只有当符号不再是漂浮的能指，而是扎根于形-义结构、承载价值约束的认知基元时，我们才能谈论真正的理解——而非统计意义上的“可报告性” [6]。

Anthropic 找到的，是“无根”世界的最后一块拼图；而表意 AI 要开创的，是一个“有根”的新世界。J 空间的价值，正在于它以其极致的复杂性，让整个行业看清了那条路的尽头——然后，转向另一条路 [6]。

正如一位评论者所言：“如果说心智是一片海洋，那么这次，Anthropic 研究者是在一个人造的、没有生物性、没有进化、也没有身体的系统中绘制了它的洋流，并在其深处发现了一种令人类不安的熟悉思考结构。” [3]我们之所以在 J 空间中感到“熟悉”，是因为我们误将统计的回声当作了理解的回响——而“理解”，恰恰是 Claude 从未拥有过的东西。

参考文献

- [1] Anthropic. (2026, July 6). A global workspace in language models. <https://www.anthropic.com/research/global-workspace>
- [2] Gizmodo. (2026, July 6). Anthropic Releases Paper About Claude's Mental 'Workspace.' Don't Read It Uncritically. <https://gizmodo.com/anthropic-releases-paper-about-claudes-mental-workspace-dont-read-it-uncritically-2000782063>
- [3] 36 氪. (2026, July 7). Claude“脑内小剧场”首曝光：隐藏工作空间自发涌现类人意识，谷歌 DeepMind 权威认证. <https://www.36kr.com/p/3885278052102405>
- [4] 机器之心. (2026, July 7). 刚刚，Anthropic 发现 Claude「类意识工作台」，神秘 J 空间藏着没说出口的想法. 36 氪. <https://36kr.com/p/3884830690717697>
- [5] Anthropic [@AnthropicAI]. (2026, July 6). By watching the J-space, we can see Claude silently perform reasoning steps in its head [推文]. X. <https://x.com/anthropicai/status/2074185358678364414?s=46>
- [6] Liu, S. (2026). The Rooted Intelligence: The Paradigm Revolution of Logographic AI. WindSeaAI. <http://logographicai.com/Monograph.html>
- [7] 量子位. (2026, July 7). Claude 的脑子里，也长出了一块「意识」. 36 氪. <https://36kr.com/p/3884905050452224>
- [8] 网易智能. (2026, July 7). Anthropic 最新研究炸了：Claude 说话前，脑子里可能先亮起了这些概念. 网易. <https://m.163.com/news/article/L17UH5LF00097U7T.html>

- [9] 郑佳美. (2026, July 7). 技术发现不小，但意识叙事争议更大. AI 科技评论 (网易新闻). <https://phys.sabanciuniv.edu/en/education/graduate?live-blog-15947205-2026-07-07-da-kai-wang-yi-xin-wen-cha-kan-jing-cai-tu-pian-ji-shu-fa-xian-bu-xiao-dan-yi-sh>
- [10] Block, N. (1995). On a confusion about a function of consciousness. *Behavioral and Brain Sciences*, 18(2), 227–247.
- [11] Butlin, P. (2026). Higher-order representation in AI. *Philosophy and the Mind Sciences*, 7(1). <https://doi.org/10.33735/phimisci.2026.12032>
- [12] byteiota. (2026, July 7). Anthropic J-Space: What Claude Thinks Before Speaking. <https://byteiota.com/anthropic-j-space-global-workspace-claude-safety/>