
archi-intelligence Research Series

Working Paper 2026-04

谁在为架构师赋能

从绘图板到推理引擎

架构设计工具链的 AI^2 -ML 成熟度基准与深度排名

archi-intelligence Research Team

2026 年 6 月发布

archi-intelligence.org

0.1 F.3 版权与许可页

版权信息 / Copyright Notice

本文档 © 2026 archi-intelligence Research Team

本作品在 Creative Commons Attribution 4.0 International (CC-BY 4.0) 许可下发布。您可以自由： - 分享—在任何媒介以任何形式复制和重新分发本作品 - 改编—修改、转换或基于本作品进行二次创作

唯一条件：必须以适当方式标明出处、提供许可链接、说明是否做出了修改。License URL: <https://creativecommons.org/licenses/by/4.0/>

数字对象唯一标识符 / DOI

DOI: 10.5281/zenodo.21035221 Permanent Archive: <https://zenodo.org/records/21035221>

引用格式建议 / Suggested Citation

archi-intelligence Research Team. (2026). *Who Builds the Architects: From Drawing Board to Reasoning Engine* (Working Paper 2026-04). archi-intelligence Research Series. <https://doi.org/10.5281/zenodo.21035221>

BibTeX 格式

```
@techreport{archi_intelligence_2026_04,  
  title      = {Who Builds the Architects:  
                From Drawing Board to Reasoning Engine},  
  author     = {{archi-intelligence Research Team}},  
  institution = {archi-intelligence},  
  year       = {2026},  
  number     = {2026-04},  
  type       = {Working Paper},  
  series     = {archi-intelligence Research Series},  
  doi        = {10.5281/zenodo.21035221},  
  url        = {https://archi-intelligence.org/research/2026-04}  
}
```

联系方式 / Contact - 研究团队 / Research Team: research@archi-intelligence.org
- 媒体咨询 / Press Inquiries: press@archi-intelligence.org - 内容更正 / Corrections: corrections@archi-intelligence.org - 网站 / Website: <https://archi-intelligence.org>

0.2 F.4 利益声明 (COI Disclosure)

利益冲突与编辑独立性声明

为保持本研究系列的学术独立性与公信力, archi-intelligence Research Team 在此明确以下事项:

1. 关于本研究机构

archi-intelligence 是致力于 Architecture Intelligence (架构智能, AI²) 研究范式建立的独立学术研究机构。本机构的使命是通过公开方法论、开放数据引用、严格同行评审, 推动跨产业架构工程实践的标准化与可比性。

2. 关于资金与赞助

本研究系列在初始阶段未接受任何被评估对象 (包括但不限于本报告中提及的工具厂商、其母公司或商业竞争方) 的直接或间接资金支持。

archi-intelligence 在初始阶段由 Arkimind (一家专注于汽车电子电气架构智能化的商业实体) 提供 seed sponsorship。这一关系在本声明中明确披露, 并通过以下治理隔离机制保障编辑独立性: - 编辑 / 研究团队不向 Arkimind 销售或业务团队汇报 - 研究人员的薪酬不与 Arkimind 任何商业指标挂钩 - 报告发布前, Arkimind 商业团队无审阅或修改权限 - 评估方法与原始数据来源 100% 公开 - 任何被评估对象可提交修正反馈, 处理流程公开

特别说明 (本篇): 这是一篇为架构工具链玩家评级的论文。Arkimind 本身可被视为本框架所丈量品类 (物种 A, AI²-ML 高阶) 的潜在参与者, 但已被刻意排除于本榜单之外——不参与评分、不列入排名, 仅在本声明中披露。这一排除是为避免任何自评利益冲突, 与本框架中立、可溯源的原则一致。

3. 关于编辑独立性

本研究的方法论选择、案例评估、结论表述均由 archi-intelligence Research Team 独立做出。研究判断不受任何商业利益、政治立场、地缘政治偏好的影响。我们承认这一编辑独立性最终需要由读者通过持续审视我们的研究质量来验证, 并承诺在后续版本中公开记录所有重要修正。

4. 关于数据来源与时间窗

本研究采用 2024 年 1 月至 2026 年 1 月期间的公开来源信息, 以标准规范 (OMG、ISO、SAE、IEEE、AUTOSAR) 与厂商官方技术文档为第一手权威, 辅以同行评审论文、官方技术发布、第三方权威报道。详见第一章 §1.7。对于无法独立验证的信息, 本研究采取保守立场——要么不引用, 要么明确标注为低置信度。

5. 关于研究限制

本研究承认数据可获得性的不对称性: 上市厂商与开源工具的公开技术细节相对完整, 部分私有厂商 (如 ETAS) 的一手文档公开度有限。本报告已在评分中标注这一不

对称性（低置信度标记），读者应警惕将”公开度差异”误读为”能力差异”。

6. 关于商标与商品名

本报告中提及的所有商标、商品名、产品代号均归各自所有者所有。引用这些名称的目的纯为学术分析与识别，不构成任何形式的背书或关联。特别地，“Ansys SCADE”“Ansys medini”为产品品牌名，其所有权实体现为 Synopsys（2025-07 收购 Ansys 后）。

0.3 F.5 详细目录 (Contents)

目录

0.1	F.3 版权与许可页	1
0.2	F.4 利益声明 (COI Disclosure)	2
0.3	F.5 详细目录 (Contents)	4
0.4	F.6 图表索引	5
0.5	F.7 缩略语表	6
1	第一章·方法论：一把丈量”架构推理”的尺子	8
1.0.1	摘要	8
1.0.2	引言：为什么”工具”决定了架构师能力的上限	9
1.0.3	本源：设计工具的起源与终极形态	10
1.0.4	定位：AI²-ML 是什么，不是什么	12
1.0.5	评分对象：四物种边界	12
1.0.6	多维评分体系与级别合成（核心方法论）	13
1.0.7	数据来源与证据标准	15
2	第二章·全景地图：四个物种与一条推理的分水岭	18
2.0.1	二维定位：演绎与归纳，上游与下游	18
2.0.2	物种 A 的分布：贴地的长尾与空缺的高地	20
2.0.3	排名：唯一的 L3，与”推理”这道分水岭	21
2.0.4	一个方法论意义的发现：AI 的热度与推理的能力脱钩	22
3	第三章·物种 A 深度排名：12 家架构设计工具画像	24
3.1	Ansys SCADE (Synopsys / 形式化验证派·唯一 L3)	24
3.2	Siemens Capital (Siemens / E/E 数字主线派·L2)	26
3.3	IBM Rhapsody (IBM / MBSE 老牌·L2)	27
3.4	Cameo / CATIA Magic (Dassault / No Magic ·SysML 旗舰·L2)	28
3.5	Eclipse Capella (Eclipse / Obeo ·开源 MBSE ·L2)	29
3.6	PTC Modeler (PTC / PLM 集成派·L2)	29
3.7	Ansys medini analyze (Synopsys / 功能安全派·L2)	30
3.8	Vector PREEvision (Vector / EEA 专用龙头·L1)	31
3.9	Sparx Systems Enterprise Architect (Sparx / 广装机通用建模·L1)	31
3.10	MathWorks System Composer (MathWorks / Simulink 耦合派·L1)	32

3.11	dSPACE SystemDesk (dSPACE / AUTOSAR 配置派·L1)	33
3.12	ETAS ISOLAR (ETAS / Bosch ·AUTOSAR 配置派·L1)	33
3.13	小结：一道分水岭，一片空地	35
3.13.1	具象锚：同一次变更，五类工具各交付什么	36
4	第四章·相邻地带：另外三个物种	39
4.0.1	物种 B：电气 / 线束详细设计——演绎，但下游	39
4.0.2	物种 C：数据 / 拆解 / 归纳平台——上游，但归纳	40
4.0.3	物种 D：AI 原生 / SDV 新玩家——横跨推理轴的探索带	40
4.0.4	小结：四物种，一张地图	41
4.0.5	另一种相邻：工具篇与 OEM 篇的连接	41
5	第五章·Roadmap：空缺的高地，与为什么是现在	43
5.0.1	为什么是现在：两个代际转折交汇	43
5.0.2	高地是什么：从 L3 到 L5 的两步	43
5.0.3	谁会去填：三条路径	44
5.0.4	结语：可验证的约束履行，作为一把参照系	45
6	附录	47
6.1	附录 A ·12×5 完整评分明细表（可独立复算）	47
6.1.1	总表（按总级别降序）	47
6.1.2	每格证据 Tier 与门槛检验	48
6.1.3	评分规则与封顶记录	49
6.1.4	三处复评结论（A 步骤）	50
6.2	附录 B ·来源清单（按工具·Tier 分级）	50

0.4 F.6 图表索引

- Figure 2.1 架构工具链四物种二维定位——推理范式 × 架构层级 (Four-Species Positioning)
 - Figure 2.2 物种 A 12 家工具 AI²-ML 级别分布 (Level Distribution)
 - Figure 2.3 物种 A 12 家工具 AI²-ML 排名 (Ranking)
 - Figure 2.4 AI 的热度与推理的能力脱钩——D3 vs D4 (D3/D4 Decoupling)
 - Figure 3.1 物种 A 12 家工具·五维能力热力矩阵 (5-Dimension Heatmap Matrix)
 - Figure 3.2 SCADE vs 行业均值·五维能力对比 (SCADE vs Mean Radar)
 - Figure D 一次变更 (12→84 区) 的级联链与五类工具能力台阶 (ADB Change Cascade)
 - Figure 4.1 四个物种与本系列版图 (Series Map)
 - Figure 5.1 从 L3 到 L5 的两步攀爬 (Two-Step Climb)
 - Figure 5.2 通往 L4-L5 高地的三条路径 (Three Paths)
-

0.5 F.7 缩略语表

缩略	全称	中文
ADAS	Advanced Driver-Assistance Systems	高级驾驶辅助系统
ADB	Adaptive Driving Beam	自适应远光灯
AI ²	Architecture Intelligence	架构智能
AI ² -ML	Architecture Intelligence Maturity Levels	架构智能成熟度
AR	Architecture Readiness	架构成熟度（本研究框架）
ARXML	AUTOSAR XML	AUTOSAR XML 工件
ASIL	Automotive Safety Integrity Level	汽车安全完整性等级
AUTOSAR	AUTomotive Open System ARchitecture	汽车开放系统架构
CAN	Controller Area Network	控制器局域网
CbC	Correct-by-Construction	构造正确
COI	Conflict of Interest	利益冲突
DOI	Digital Object Identifier	数字对象唯一标识符
ECE	Economic Commission for Europe (Regulations)	联合国欧洲经济委员会法规
ECU	Electronic Control Unit	电子控制单元
EEA	Electrical/Electronic Architecture	电子电气架构
FMEA	Failure Mode and Effects Analysis	失效模式与影响分析
FMEDA	Failure Modes, Effects and Diagnostic Analysis	失效模式影响及诊断分析
FTA	Fault Tree Analysis	故障树分析
HARA	Hazard Analysis and Risk Assessment	危害分析与风险评估
KerML	Kernel Modeling Language	内核建模语言（SysML v2 基础）
LIN	Local Interconnect Network	本地互连网络

缩略	全称	中文
MBSE	Model-Based Systems Engineering	基于模型的系统工程
OEM	Original Equipment Manufacturer	整车厂
OMG	Object Management Group	对象管理组织
OSLC	Open Services for Lifecycle Collaboration	开放生命周期协作服务
OTA	Over-The-Air	空中下载（升级）
PLM	Product Lifecycle Management	产品生命周期管理
ReqIF	Requirements Interchange Format	需求交换格式
REST	Representational State Transfer	表述性状态转移（API 风格）
RFLP	Requirements-Functional-Logical-Physical	需求-功能-逻辑-物理（建模分层）
SAT	Boolean Satisfiability	布尔可满足性（求解器）
SDV	Software-Defined Vehicle	软件定义汽车
SoS	System of Systems	体系
SOTIF	Safety Of The Intended Functionality	预期功能安全
SSoT	Single Source of Truth	单一真相源
SysML	Systems Modeling Language	系统建模语言
TQL	Tool Qualification Level	工具鉴定等级
V&V	Verification and Validation	验证与确认
XMI	XML Metadata Interchange	XML 元数据交换

1 第一章·方法论：一把丈量”架构推理”的尺子

作者署名：archi-intelligence Research Team

1.0.1 摘要

本文以”设计工具”这一跨越六十年的工程造物为起点，追溯了从 1963 年 Sutherland 的 Sketchpad 内含的约束求解器，到当代 MBSE、EDA 与 AUTOSAR 工具链的演化主线，提出一个常被工具厂商的市场叙事所遮蔽的判断：设计工具最深的身份是约束维护，而非图形描绘。在此基础上，本文将三部曲所建立的 AI²-ML（Architecture Intelligence Maturity Levels，架构智能成熟度）框架，从丈量”被设计的架构”转用于丈量”设计架构的工具”，沿约束感知深度、参考架构整合、AI 辅助生成、AI 推理自主、互操作性五个能力维度（D1-D5），对汽车架构设计工具链中与该框架同源的核心物种——架构设计 / MBSE / 推理工具——逐一打分，并按以 D1、D4 为核心双轴的合成规则换算为 L0-L5 总级别。全部评分以标准规范与厂商官方技术文档为第一手依据，遵循 2024.1-2026.1 的时间窗与公开可复算的证据规程。

研究对 12 家工具的丈量得出三个相互印证的发现。其一，整个品类贴地分布而顶端空缺：唯有一家工具（SCADE，现属 Synopsys）凭借其形式化定义的语言与模型检查器触及 L3，且其可验证性仅及于软件 / 模型层；其余十一家——涵盖大型 MBSE 套件（Cameo、Rhapsody、Capella、medini）、E/E 与 AUTOSAR 配置工具（PREEvision、SystemDesk、ETAS）以及通用建模器——尽数落在 L1-L2，无一家达到整系统层面的 L4，更遑论 L5。其二，“AI 的热度”与”推理的能力”出现系统性脱钩：2025-2026 年各厂商密集发布的 AI 能力几乎无一例外落在生成（D3），无一抬升了推理自主（D4）；工具越来越擅长”生成一个方案”，却仍不能”推出一次变更的后果并担保其正确”。其三，SysML v2 API & Services 虽于 2025 年由 OMG 最终采纳、并在 2026 年进入落地，但其互操作能力仍单薄，且因缺乏官方一致性测试套件，所有高分主张均属厂商自述。

据此，本文提出并论证：一件架构设计工具的终极形态，不是”自主生成”，而是“**可验证的约束履行**”——让设计及其每一次变更的正确性可被证明。这一判断由技术哲学、产业工具谱系、自适应系统三条独立路径交叉验证。本文将 AI²-ML 的顶端（L5）定义为”无未经验证的变更”，把 D1（约束感知）与 D4（推理自主）确立为核心双轴，并指出当前那片空缺的 L4-L5 高地，恰是被三部曲所揭示的、整个汽车产业”架构债务”从供给侧标记出来的空白。AI²-ML 是一把能力成熟度阶梯，刻意排除市场份额维度，以可观测证据与来源分级保证每一项评分可被独立复算。

关键词：架构设计工具；约束感知推理；变更影响分析；构造正确（Correct-by-Construction）；MBSE；SysML v2；架构成熟度；AI²-ML

1.0.2 引言：为什么“工具”决定了架构师能力的上限

架构师在画图。但车正在变成机器人。约束正在爆炸。而工具，大多还停留在“标记一条规则被违反”。

这是当下架构工具链的真实状态，也是一个被普遍低估的张力。过去十年，关于汽车架构的叙事几乎都集中在“被设计出来的系统”——中央计算如何取代分布式 ECU、软件如何定义汽车、域控制器如何收敛为准中央。三部曲所做的，正是为这场系统侧的演进建立坐标。但在这场叙事的背面，一个同样关键、却鲜被追问的问题被悬置了：架构师究竟用什么在设计这些越来越复杂的系统？如果说系统是“被设计的物”，那么决定这个物能被设计到多精密的，是“设计它的工具”。

这个问题在此刻变得尖锐，是因为约束的密度发生了量变到质变的跃迁。一辆现代汽车的电子电气架构，牵涉数十个相互制约的回路——功能安全约束着拓扑，拓扑约束着线束重量，重量约束着成本，算力布局约束着 OTA 能力，而其中任何一项又反过来牵动功能安全等级。改动其中一处，后果沿着这张约束之网扩散到看似无关的远端。当系统还是分布式、每个 ECU 各管一摊时，这张网稀疏到人脑尚可把握；而当系统迈向中央集中、软件定义、整车 OTA，约束之网的密度已经超出任何个体能在脑中完整推演的限度。架构师此刻需要的，不再是一支更好的画笔，而是一台能替他把这些后果推出来、并担保结论正确的引擎。

然而他手上的工具，大多只能告诉他“这里和那条规则冲突了”，至于冲突意味着什么、会沿约束链波及何处、最终是否仍满足安全与法规，仍要他自己逐一判断。工具停在“标记”，把最耗心智、也最性命攸关的“推理”留给了人。于是一个静默的事实浮现出来：**架构师的能力上限，被他的工具悄悄地划定了。**一个再优秀的架构师，也无法用一把只会标记的尺子，去担保一个其复杂度已超出人脑推演极限的系统。

这正是本系列要回答的问题：**谁在为架构师赋能？**那些声称在“辅助设计”的工具——PREEvision、Cameo、Rhapsody、SCADE 及其同侪——它们从何而来、解决了什么根本问题、又能把架构师带到多高？要回答它，不能从功能清单开始，那只会复述厂商的市场叙事；必须从本源开始。三部曲的第一篇，从“架构”一词的希腊语词源讲起，因为追本溯源才能看清目标。本篇沿用同样的方法：在给任何工具打分之前，先用一节 (§1.3) 追问设计工具的目的论本源——它的终极形态究竟是什么。这一追问并非思辨的余兴，它直接决定了这把尺子的顶端应当锚定在哪里；而尺子的顶端定在哪里，又决定了我们将如何评判今天这一整代工具的位置与差距。

1.0.3 本源：设计工具的起源与终极形态

起源：Sketchpad 与”约束求解器” 设计工具的历史，常被讲成一部画图效率的进步史：从绘图板到电子图纸，从二维线条到三维模型。这种讲法遗漏了最关键的一点。

1963 年，Ivan Sutherland 在麻省理工学院的 TX-2 上演示了 Sketchpad——公认的第一个交互式图形设计工具。人们记住了它的光笔与屏幕，却忽略了它真正革命性的部分：Sketchpad 内含一个**约束求解器**。用户可以声明几何关系——平行、等长、垂直——Sketchpad 会在图形变化时动态地维持这些约束。设计工具在它的第一个实例里，就不是”让人画得更快”的描绘工具，而是一台”理解约束、并据此推理”的引擎。

这一史实给本系列一个起源级的锚点：**设计工具最深的身份是约束维护，而非描绘**。描绘是表象，约束推理才是本质。此后每一代工具——CAD 的几何建模、CAE 的行为仿真、EDA 的设计自动化、MBSE 的系统建模——都在解决同一个根本问题的不同侧面，只是这个问题随系统复杂度的增长而不断加码。

根本问题：当约束之网超出人脑工作记忆 如果设计工具的本质是约束维护，那么它要对抗的根本困难是什么？可以一句话说清：一个复杂系统的设计，要求同时持有的、相互依赖的约束，远多于任何个体所能在脑中持有的；并要求把每一次变更的后果，沿那张约束之网完整地传播开去。这正是”约束感知的变更影响推理”——也正是设计工具自 Sketchpad 以来一直在外化、并日益自动化的那件事。

值得强调的是，这个根本问题不是静止的，它随系统复杂度单调地加剧。当 Sutherland 在 1963 年维持几条几何约束时，约束之网尚且稀疏；而一辆 2026 年的汽车，其电子电气架构所牵涉的约束——功能安全、实时时序、总线带宽、功耗热耗、成本重量、法规符合、OTA 可达——已交织成一张人脑无法完整持有的网络。约束的数量级在涨，约束之间的耦合度也在涨：一处改动不再只影响相邻节点，而是沿着多条隐藏的依赖链扩散到看似无关的远端。正是这种”耦合密度”的增长，使得”靠经验丰富的架构师在脑中推演”这一传统方式逐渐触及极限。

这解释了为什么设计工具的演化史，本质上是一部”把约束推理从人脑外化到工具”的历史。CAD 外化了几何约束的维护，CAE 外化了物理行为的验证，EDA 外化了电路时序与布局的约束求解，MBSE 试图外化整个系统模型的一致性。每一代工具都在回答同一个问题的一个侧面——如何在约束之网超出人脑工作记忆之后，仍然保证设计的正确。而这条外化之路的终点，正是下一节要追问的：当外化彻底完成时，工具会变成什么？

终极形态：工具的”L5”是什么 三部曲借用 SAE J3016 的自动驾驶分级：汽车的 L5（完全自主）之所以是终极形态，是因为它完全履行了汽车存在的根本目的——把人从 A 点运到 B 点——以至于人不再需要参与驾驶。沿用同样的逻辑：设计工具的”L5”，完全履行了什么目的？

“自主生成”是一个直觉答案，几乎是这个时代最自然的联想——当生成式 AI 能写代码、能画图、能产出设计草案，“工具的终极形态就是不需要人也能把架构生成出来”似乎顺理成章。但这个直觉误读了汽车 L5 的结构。即便在 L5，人依然是意图之源——人选择目的地，车从不自己决定去哪；L5 自动化的是”驾驶”这个执行过程，而非”想去哪”这个意图。移植到设计工具上，“自主生成”只回答了”谁来执行设计”，却回避了更尖锐、也更要害的一问：谁来担保这个设计是正确的？把”生成”误认为终点，恰恰跳过了设计工作中最难、也最不可让渡的那一半。

在安全攸关领域，这一问不可回避。一辆车的 E/E 架构若有缺陷，代价是人命。ISO 26262 因此要求一份安全论证（safety case）——一个有结构、有证据支撑的论证，证明设计满足其安全目标。一份安全论证无法用”神经网络建议了这个方案”来履行，只能用针对约束结构、可被检查的**证明**来履行。

于是工具的终极形态从”自主”落到了”可验证的约束履行”：**终极的设计工具，不是无需人就能把架构画出来的工具，而是让设计及其每一次变更的正确性可被证明的工具。**

三方互证 这一结论由三条独立路径交叉验证，三套词汇指向同一顶端：

- 技术哲学与工具目的论 —— “可验证的约束履行”：安全攸关领域的举证责任，给纯归纳式的”替代人”划定天花板，稳定顶点是人或可信验证器仍持有担保。
- 产业工具谱系 —— “证明义务的工具化”：越靠近量产与功能安全，工具越不像画图软件、越像证明义务管理系统。
- 自适应系统与运行时 —— “构造正确（Correct-by-Construction）”：终极工具被系统内化为运行时综合引擎，每次自我重构前先经形式化验证数学地证明不破坏安全包络。

三个来源、三套语言、一个结论。工具的终极形态不是”自主生成”，而是”自主生成 + 能证明正确”——后者那一半，恰是纯归纳做不到、必须由演绎兜底的。

据此定义 AI²-ML 的 L5 L5 不是”无人”，而是”无未经验证的变更”。汽车的 L5 通过剥夺方向盘解决”位移”；架构工具的 L5 通过把”验证”覆盖到每一次变更，解决”在开放世界中安全地演化架构”。在这个顶端，工具并未消失，也未坍缩成自行设定目标的智能体，而是升维为推理与验证的基础设施。

“替代人类架构师”与“增强人类架构师”两种叙事，因此不是矛盾，而是同一座山的两个海拔，分水岭在于举证责任被内化到何种程度：当验证器仍由人持有，是增强；当智能体把验证器内化、能自己证明正确，才谈得上替代。替代之所以可能，正因约束验证可被内化；在它被可信地内化之前，增强是稳定态。这一判断，也是本框架把 D1 与 D4 设为核心双轴、要求二者同时达标才能登顶的根本理由（见 §1.6）。

1.0.4 定位：AI²-ML 是什么，不是什么

AI²-ML 是一把能力成熟度阶梯。它在谱系上承接 CMMI（逐级包含的阶段式成熟度）与 SAE J3016（L0–L5，以人与系统的责任转移界定级别）。它衡量单个工具在“约束感知的架构推理”上的能力深度，顶端是“可验证的约束履行”。

它不是市场定位象限。与 Gartner 魔力象限、Forrester Wave、IDC MarketScape 那类二维竞争快照不同，本框架刻意排除市场份额、营收、销售、装机量、厂商生存力等一切与商业成功相关的轴。这是方法论上的特性而非缺漏：一个工具的推理能力，与它卖得好不好无关。把市场成功混入能力评分，正是分析师框架被反复记录在案的结构性弱点。

它也不是组织成熟度模型。CMMI、INCOSE 的 MBE 能力矩阵衡量组织采纳建模的成熟度；本框架衡量工具本身的能力。

结构上与 AI²-ML 最接近的现有成果，是 ZWF/De Gruyter 的“AI 增强 MBSE”六级模型（Bernijazov 等，2025）——同样 L0–L5、同样受 J3016 启发。本框架公开引用这一成果，但二者关键差别在于：ZWF 止步于“自主性”与“V&V/一致性”，并未对推理深度、尤其“可验证的约束履行”评级。AI²-ML 占据的，正是 ZWF 与所有其他框架留空的那条轴。

1.0.5 评分对象：四物种边界

本系列把架构工具链玩家划为四个物种：

- 物种 A 架构设计 / MBSE / 推理工具 —— 真正用于设计架构的工具（PREEvision、Cameo、Rhapsody、SCADE 等）。
- 物种 B 电气 / 线束详细设计工具 —— 下游详细设计（Zuken E3、Solid Edge Electrical 等）。
- 物种 C 数据 / 拆解 / 归纳平台 —— 从数据与范例归纳的平台（A2MAC1、Palantir 等）。

- 物种 D AI 原生 / SDV 新玩家 —— 架构智能阶梯前沿的新创。

本篇开篇只对物种 A 做深度排名，因为只有它与 AI²-ML 同源——这把尺子衡量演绎式架构推理，而这正是物种 A 的本职。物种 B、C、D 在第二章全景地图中定位，深度排名留给本系列后续专属篇章。

为何不能把四物种放进同一个评分榜：用 AI²-ML 给一个归纳式数据平台（物种 C）打”约束感知深度”分，如同用温度计量重量——得到的不是分数低，而是维度不适用。物种边界以此杜绝错配，也呼应三部曲已确立的”演绎 vs 归纳”“上游 vs 下游”两组区分。

1.0.6 五维评分体系与级别合成（核心方法论）

本框架沿五个能力维度（D1–D5）对每件工具打分，每维 0–5 分，按加权合成换算为 L0–L5 总级别。全部规则公开如下，任何研究者可依据附录 A 的五维明细独立复算。

五维定义（D1–D5）

- D1 约束感知深度（Constraint-Awareness Depth）—— 工具理解架构约束、并据此推理的深度，顶端是能对变更给出形式化正确性证明。
- D2 参考架构整合（Reference-Architecture Integration）—— 对单一真相源 / 数字主线 / 参考架构的整合深度。
- D3 AI 辅助程度（AI-Assistance Level）—— AI 在生成与建议上的能力（copilot 一脉），不含推理自主。
- D4 AI 推理自主度（AI Reasoning Autonomy）—— 在约束推理上工具自己下结论、并自证正确的程度。
- D5 互操作性（Interoperability, SysML v2 API）—— 以开放标准衡量的开放与可集成程度，是通往顶端的使能维度而非顶端本身。

为了让这五维不停留在抽象，本篇用一个贯穿全篇的例子。设想一个常见系统：**自适应远光灯（ADB，矩阵 LED 防眩目）**——摄像头检测对向来车，对应方向熄灭若干 LED 挖出阴影区防眩目，受功能时序、ISO 26262 ASIL-B、ECE 眩光限值三重约束。再设想一次最普通的变更：**光束分区从 12 区升级到 84 区**。这一改动会沿约束链级联（摄像头分辨率 → ECU 算力 → LED 通道 → 总线负载 → 线束散热 → 时序 / 安全 / 法规是否仍成立）。五维量规问的，正是工具面对这次变更能做什么：

- **D1**：能否理解”84 区不得破坏 ASIL-B”这类约束，并据此推理？

- **D2**: 能否把这次变更放进一个贯通需求—设计—安全的单一模型里？
- **D3**: 能否用 AI 生成 84 区的候选配置？
- **D4**: 能否自己推出”总线超时将打破时序”这个后果、并下结论？
- **D5**: 能否让这次变更的模型在工具链间以 SysML v2 标准流通？

本例将在第二、三章各维度评分时回扣，并在第三章末以一张”五类工具对照表”收束——把抽象评分落到这次具体变更上。

D3 与 D4 的分工是本框架的标志性设计：**D3 管”生成方案”，D4 管”推出影响、下结论、并担保正确”**。一件工具可能很会生成（D3 高）却不会自证（D4 低）；这条分界把”会画不会想”的工具与”真能推理”的工具切开。

评分聚合：差异化加权 + 核心双轴短板约束 总级别 $L = \text{下取整}(\text{五维加权均值})$ ，权重为 D1 25%、D4 25%、D3 20%、D2 15%、D5 15%；但受核心双轴短板约束：**要宣称达到 L_n ，D1 与 D4 必须各自独立 n 。**

这套规则有两层用意。其一，D1 与 D4——推理双核——合计占 50% 权重，凸显本框架衡量推理而非功能堆砌；D5 虽最客观却只占 15%，正因它是使能维度而非顶点，以避免”互操作性高 = 能力高”的错觉。其二，短板约束防止跳级：一件工具若靠 D5 刷高均值却不会推理、不能自证，核心双轴的短板会把它拉回。更深一层，登顶 L5 须 D4 内化验证器（举证责任内化），否则封顶 L4——这把 §1.3 所论”替代之所以可能、正因验证可被内化”落实成一条可计算的规则。

D1 打分锚点

D1 分	标准	典型例
5	可验证的约束履行 / 构造正确：对变更给出形式化证明	SCADE Suite (Design Verifier 模型检查)
4	推理：从约束演绎出后果 (变更影响传播)，非仅报告违例	(本篇 12 家无稳定达此级)
3	参考感知：对照参考架构 / 模式库判定偏离	Capella、medini
2	规则校验：内置一致性 / 完整性检查	PREEvision、Cameo、Rhapsody

D1 分	标准	典型例
1	结构化建模：元素有类型与关系，无约束校验	—
0	纯绘图：图元无语义	—

D4 打分锚点

D4 分	标准	典型例
5	自主推理 + 内化验证器自证正确（构造正确），人仅处理异常	（前沿 / 未实现）
4	自主推理并给出可执行结论（影响范围 + 处置），人监督担保	SCADE Suite
3	自动推导变更影响链，人审核	（商业未发货）
2	规则触发式影响标记，需人确认逻辑	本篇多数工具
1	提示潜在影响点，不下结论	—
0	不做推理，影响全靠人判断	—

D2、D3、D5 的完整 0–5 锚点见附录 A。一个对全篇至关重要的观察：2025–2026 年各厂商发货的”AI”，压倒性地落在 D3（生成）而非 D4（推理）——IBM 的 Engineering AI Hub、Siemens 的生成式工程、Cameo+ChatGPT 工作流，无一例外是”生成方案让人验证”，没有一个能”自己推出影响并担保正确”。这正是 D3 与 D4 必须分轴的实证理由。

1.0.7 数据来源与证据标准

本篇的全部评分可被独立复核——这正是商业分析师框架被诟病所缺、而学术框架理应具备的品质。为此，本篇继承 archi-intelligence Research Series 在三部曲中确立的数据来源金字塔与评分证据标准，并将其从”丈量 OEM 架构成熟度”适配到”丈量工具能力成熟度”。

1.7.1 数据来源金字塔。本系列按三个维度对每条证据分级： 责任约束（来源是否受法律、受信义务或品牌信誉约束而须对真实性负责）； 来源独立性（标准组织、厂商自身、独立第三方，还是间接转述）； 时效性（已发货的当前状态，还是前瞻路标）。据此分为五级。需说明本篇相对三部曲的一处适配：三部曲衡量 OEM 架构成熟度，公司的既成事实写在受法律约束的财报里，故其 Tier 1 以监管财报为首；本篇衡量工具能力，“工具能做什么”的第一手权威是标准规范与厂商官方技术文档，故本篇 Tier 1 以这两类为首，监管财报降为核定所有权与财务之用。

- **Tier 1 (标准规范 / 官方技术文档级)**: OMG (SysML v2、KerML、API & Services)、ISO/SAE/IEEE/AUTOSAR 标准; 厂商用户手册、Release Notes、API 文档; TÜV / DO-330 / ISO 26262 ASIL 等认证证书; 上市母公司 SEC/HKEX 文件（用于所有权与财务，不直接证明工具能力）。
- **Tier 2 (同行评审 / 独立分析级)**: 同行评审论文（尤其形式化方法、模型检查、变更影响分析）、独立技术评测、学术对照框架（如 ZWF 六级模型）。
- **Tier 3 (官方技术发布 / 演讲级)**: 用户大会演讲、官方白皮书/技术博客、专利、GA 公告。
- **Tier 4 (第三方权威报道级)**: 主流技术媒体署名报道、行业研究报告、授权分销商/集成商技术页。
- **Tier 5 (间接 / 营销级)**: 厂商营销修辞（“唯一”“100% 一致”“业界领先”）、社媒/招聘/传闻。仅脚注，不影响主分。

1.7.2 评分证据标准。每一维评分（D1–D5）须至少由 2 个 Tier 1–3 级证据支撑，不满足者标”低置信度”并保守取分；Tier 4 仅作补充，Tier 5 仅脚注；任何 4 分须有 Tier 1–2 直接支撑。冲突处理优先级：Tier 等级 > 同级时序（新者为准）> 官方原文。非上市/私有厂商（本篇多家如此）：以官方技术文档评能力、二手交叉验证所有权，证据不足时标低置信并保守取分。

1.7.3 厂商能力声明封顶规则（本篇新增）。当厂商宣称达到某标准一致性（如”100% 支持 SysML v2 API”）但不存在该标准的官方一致性测试套件、或无独立认证结果时，按 Tier 1 记录其”已实现”事实，但”完整一致”按自述处理并在该维评分设上限（不因自述授予满分）。范例：SysML v2 API & Services 的 OMG 一致性测试套件在本篇时间窗内尚不存在，故 D5 满分（5 = 通过 OMG 一致性测试）暂不授予任何工具，相关厂商自述封顶于 4。

1.7.4 时间窗与可复现。主评估窗口 2024.1–2026.1，全部评分基于此窗口内已发货的工具能力；截止日后、发布前的信息（如 Synopsys 2025-07 完成收购 Ansys 属窗口内，Ansys 2026 R1 于 2026-03 发布属窗口外）标注”截止后补充”，作所有权认定或 roadmap 进度证据，不静默改主分。鉴于 SysML v2 与各厂商 AI 功能在 2026 年快速演化，全篇

标注”截至 2026.1”快照，遵循版本化、不静默漂移原则；SysML v2 API & Services 由 OMG 于 2025-07-21 最终采纳，本篇按工具实际支持评分，而非路标承诺。每件工具每维评分附可观测证据 + Tier 标注，读者可依附录 A 独立复算；本篇借用 Forrester 透明计分卡的机制（公开量表、权重、可追溯依据），但拒绝其市场维度。

注：本篇的利益冲突与独立性声明（COI Disclosure）按 D1 体例置于前置页 F.4，不在正文重复。

2 第二章·全景地图：四个物种与一条推理的分水岭

本章为架构工具链绘制一张全景地图。第一章确立了丈量的尺子（AI²-ML）与它的顶端（可验证的约束履行）；本章先用两个正交维度把整个工具链版图铺开，识别出四个物种，再聚焦本篇的主榜物种 A，给出排名预览。逐工具的画像与完整评分留待第三章与附录 A。

之所以先画地图、再排名，是因为这个领域长期缺一张能让人看清彼此关系的图。市场上的工具被笼统地称作”设计工具”或”MBSE 工具”，仿佛它们在做同一件事、可以同台比较；但稍加追究就会发现，它们其实分属几种认识论不同、工作层级不同的物种——有的在演绎、有的在归纳，有的在上游决定架构、有的在下游落实细节。把它们不加区分地排进一张榜单，等于用同一把尺子去量四种不可通约的能力，结论必然失真。因此，在打分之前，先得有一张地图，把”谁和谁真正可比”这件事讲清楚。

需要先行说明本章的边界。全景图覆盖全部四个物种，以确立完整版图；但只有物种 A 进入 AI²-ML 的深度排名，因为只有它与这把尺子同源——它做的正是”在约束下推理架构”这件事，而尺子丈量的也正是这件事的成熟度。物种 B、C、D 在图上定位、说明其性质，深度排名留给本系列后续篇章；第四章将专门交代它们各自的疆域与去向。

2.0.1 二维定位：演绎与归纳，上游与下游

第一张图沿两个正交维度为工具链玩家定位。横轴是**推理范式**：从归纳（从数据与范例中学习模式）到演绎（从约束与规则推出必然后果）。纵轴是**架构层级**：从下游的详细信息（线束、ECU 配置）到上游的概念架构（系统级权衡、变更影响推理）。

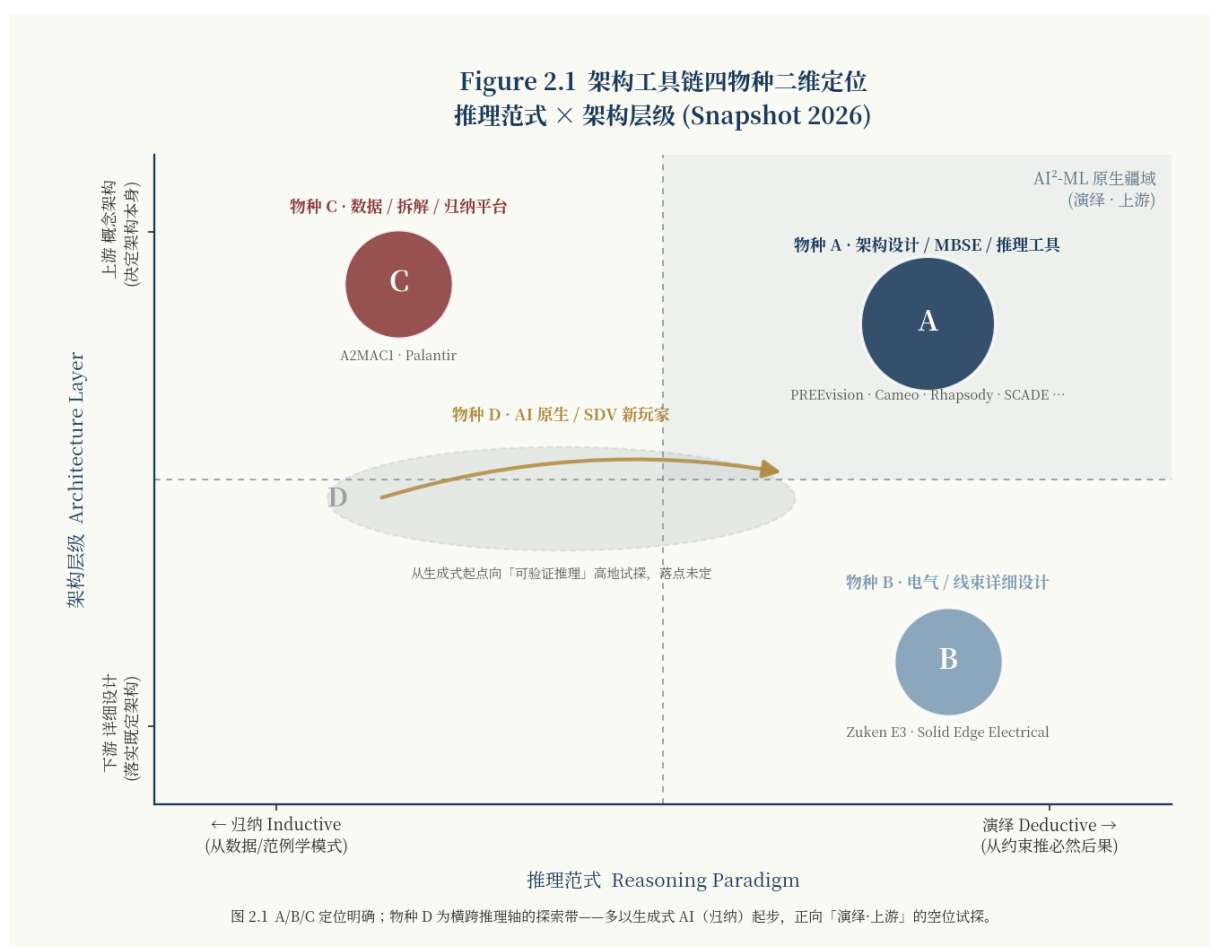


图 2.1 架构工具链四物种二维定位——推理范式 × 架构层级 (Snapshot 2026)。

这两个轴并非随意选取，而是直接源自三部曲的两组区分。**演绎 vs 归纳，划开了”能证明”与”只能预测”的两类工具；上游 vs 下游，划开了”决定架构长什么样”与”把既定架构落到实处”的两类工作。**四个物种在这张图上各占一隅：

- **物种 A（架构设计 / MBSE / 推理工具）**居于”演绎—上游”象限。它们用模型表达约束、在概念架构层做一致性检查与（理想情况下）变更影响推理。这是 AI²-ML 的原生疆域。
- **物种 B（电气 / 线束详细设计工具）**居于”演绎—下游”。它们同样基于规则，但作用于已确定架构之后的详细设计，不决定架构本身。
- **物种 C（数据 / 拆解 / 归纳平台）**居于”归纳—上游”。它们从拆解数据与市场样本中归纳洞察，影响架构决策，但其方法是归纳而非演绎——它们提供”别人怎么做”，而非”这样改对不对”。
- **物种 D（AI 原生 / SDV 新玩家）**不是一个钉死在某一象限的实心物种，而是一条横跨推理轴的探索带：多数以生成式 AI（归纳）起步，但其中真正有抱负的玩家，正向”演绎—上游”那个空缺的高地试探，落点未定。这一物种的形态本身，就映射了 L4→L5 前沿尚未被占据的事实。

这张图的含义是结构性的：**四个物种不是同一品类的不同档次，而是用不同方法、在不同层级做不同事情的不同物种**。把它们排进同一个榜单，等于用同一把尺子丈量四种不可通约的能力。本篇因此只在”演绎—上游”象限内——物种 A——展开深度排名；其余三个物种的存在界定了物种 A 的边界，但不与它同台竞分。

2.0.2 物种 A 的分布：贴地的长尾与空缺的高地

把物种 A 的 12 家工具按 AI²-ML 总级别分布，呈现出一个值得注意的形态。

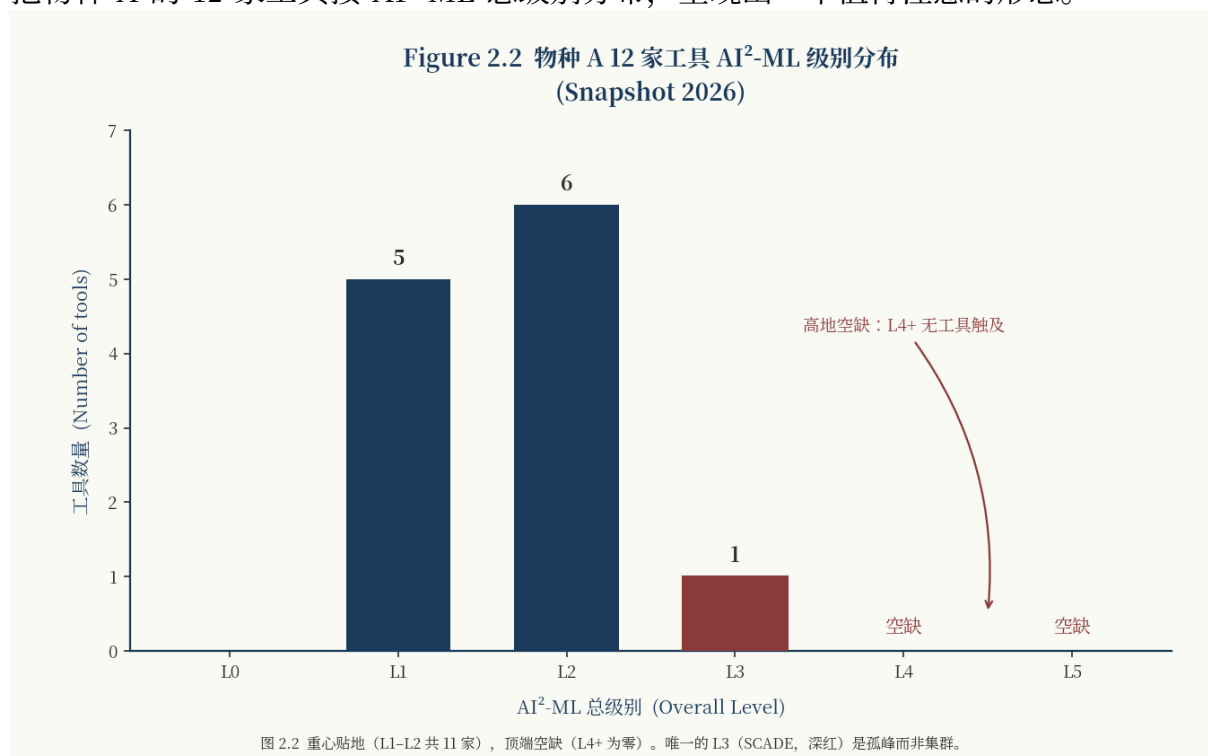


图 2.2 物种 A 12 家工具 AI²-ML 级别分布 (Snapshot 2026)。

12 家中，L1 五家、L2 六家、L3 一家，L4 及以上为零。重心明显落在低端，分布贴着地面，顶端是空的。

这一形态的含义有二。

其一，整个品类都还停在”会检查、会生成，但不会自证正确”的阶段。L1-L2 的十一家工具，其能力主体是结构化建模、一致性检查、可追溯性追踪，以及近两年新增的生成式 AI 辅助。它们能告诉架构师”这里和那条规则冲突了”，却无法替他推出冲突的后果、更无法证明一次变更正确。这不是某一两家的短板，而是品类层面的天花板。

其二，高地是空缺的，而非已被占据的。在三部曲的 OEM 排名中，AR4 已聚集两家、形成清晰的领先集群；而在工具链这边，L4 一家都没有，唯一的 L3 也是孤峰。这意味着”推理工具”这个品类的领先位置尚未被任何在位者占据——L4→L5 的前沿是空

着的。对一个新进入者而言，这是空地而非红海；对整个行业而言，这是一道尚未有人跨过的门槛。

2.0.3 排名：唯一的 L3，与”推理”这道分水岭

逐家排名呈现一个清晰的事实：在 12 家工具中，唯有一家——SCADE（Synopsys 旗下，原 Ansys；产品品牌仍为 Ansys SCADE）——触及 L3，且它与其他十一家之间，隔着的不是分数高低，而是一道能力性质的分水岭。

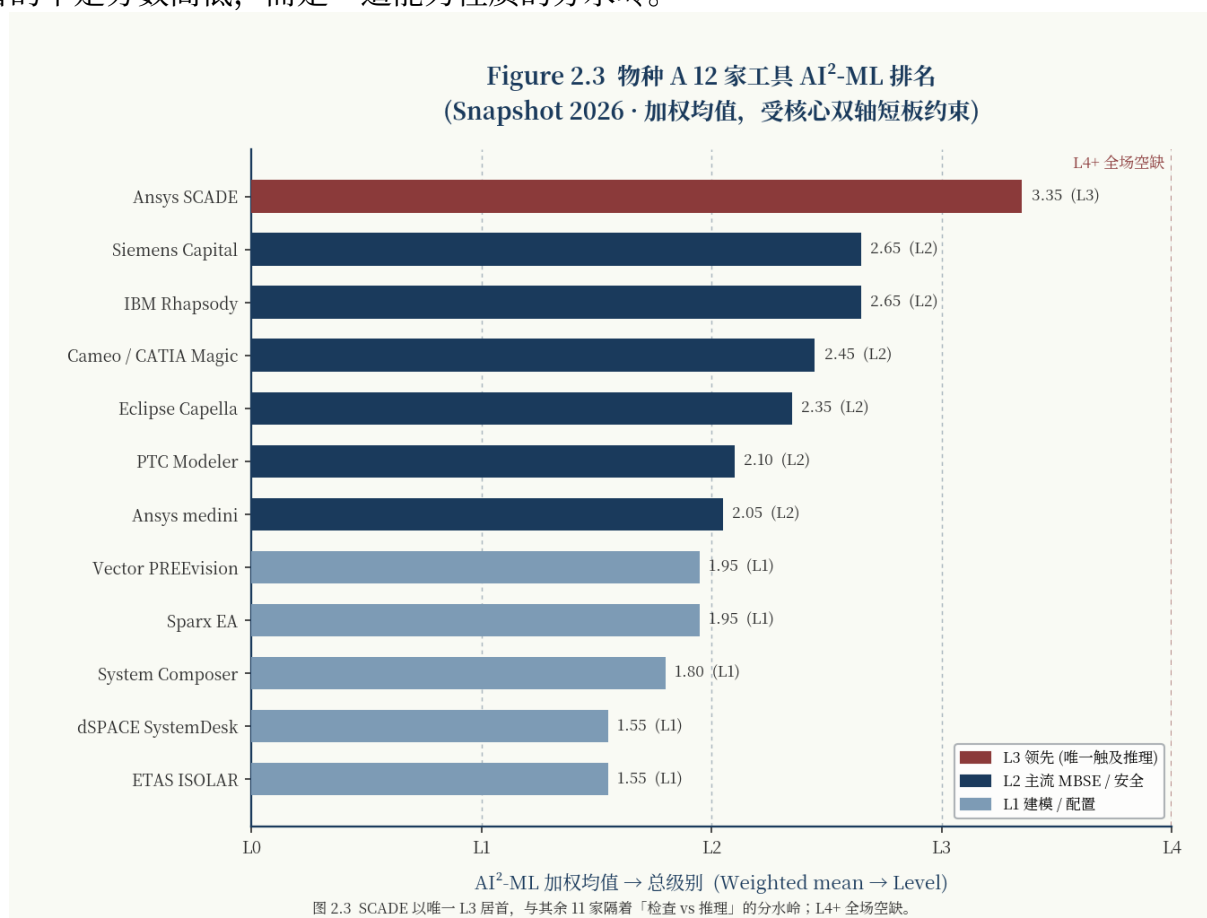


图 2.3 物种 A 12 家工具 AI²-ML 排名 (Snapshot 2026)。

SCADE 之所以独占 L3，是因为它在两条核心轴上做了别人做不到的事。其形式化定义的 Scade 语言加上 Design Verifier 模型检查引擎，使它能对一次变更给出**形式化的正确性证明** (D1=5)，并能自主推理、以反例下结论 (D4=4)。其余十一家工具的”变更影响分析”，本质都是可追溯性矩阵——人沿着依赖链去追溯，而非工具把后果推出来。

这道分水岭，是”我们标记了这条规则被违反”与”我们证明了这条性质成立”之间的区别。前者是检查，后者是推理。SCADE 站在推理一侧，其余工具站在检查一侧。这正是 D1 与 D4 这两条核心轴存在的意义：它们把品类切成了两半，而当前只有一家在推理那一半。

需要说明 SCADE 的边界：它的形式保证以软件 / 模型为范围（SCADE Suite 核心），其系统架构层（SCADE Architect）仍是一致性检查而非整系统证明。换言之，连这唯一的 L3，也只是在软件层面触及了推理的门槛，而非在整系统层面跨过了它。这恰好定义了 L4→L5 的前沿：把证明从软件扩展到整系统，并使其覆盖每一次变更。

主榜其余各家,按级别聚为两档。**L2 档(六家)**是主流 MBSE 与安全工具——Siemens Capital、IBM Rhapsody、Cameo / CATIA Magic、Eclipse Capella、PTC Modeler、Ansys medini。它们各有所长：Capital 领先于数字主线与生成式辅助，Rhapsody 与 Cameo 领先于 SysML v2 互操作，Capella 与 medini 在约束感知上达到 D1=3。但它们无一例外被门控在 L2，因为其推理自主度（D4）都停留在”追溯式标记”。**L1 档（五家）**是 PREEvision、Sparx EA、System Composer、dSPACE SystemDesk、ETAS ISOLAR——有能力的建模与配置工具，但其推理由可追溯性驱动，或本职是标准配置而非架构推理。

2.0.4 一个方法论意义的发现：AI 的热度与推理的能力脱钩

前三节描绘了竞争格局；本节讨论一个更具方法论意义的发现：一件工具在 AI 上的声量，与它在推理上的能力，是脱钩的。

2025–2026 年，几乎每一家主流厂商都发布了”AI”。IBM 推出 Engineering AI Hub，其 MBSE 智能体能从自然语言需求生成模型元素；Siemens 的 Simcenter Studio 以 AI 自动探索架构候选；Cameo 通过 ChatGPT workflow 辅助建模。从营销声量看，这是一场 AI 竞赛。

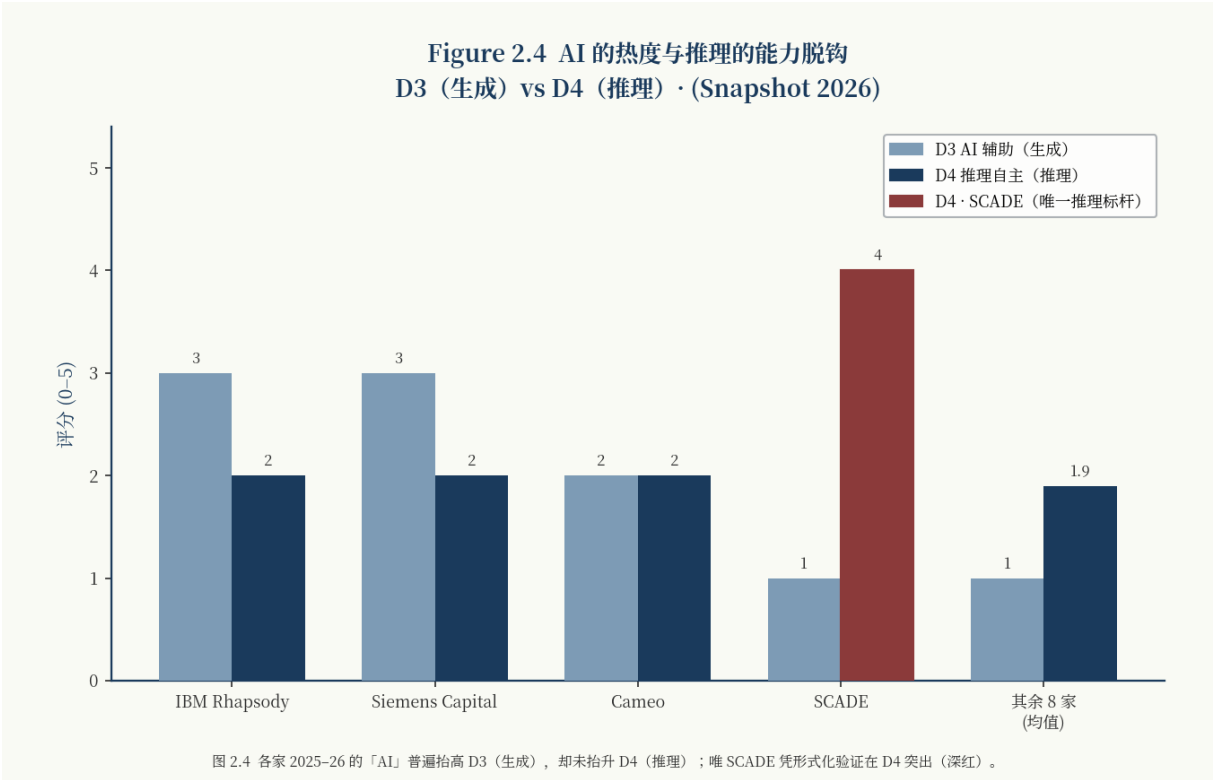


图 2.4 AI 的热度与推理的能力脱钩——D3 (生成) vs D4 (推理) (Snapshot 2026)。

但用 D3 与 D4 这两条轴去量，真相浮现：**这些 AI 无一例外落在 D3 (生成)，没有一个抬升了 D4 (推理)**。它们都在做同一件事——生成方案、让人验证。没有一个能自己推出一次变更的后果、并担保结论正确。生成与推理是两类不同的能力：生成回答”给我一个方案”，推理回答”这个改动对不对、会波及哪里”。前者是归纳的延伸，后者要求演绎的兜底。回到引言的例子：今天的 AI 能生成一个 84 区的候选配置 (D3)，却没有一个能自己推出”总线超时将打破 ASIL-B 时序”并下结论 (D4)。

这一脱钩的含义是双重的。**对评分而言，它证明了 D3 与 D4 必须分轴**——若把两者合并为笼统的”AI 能力”，IBM、Siemens 这些工具的生成式声量会虚假地抬高它们的推理评分，掩盖整个品类在 D4 上的集体停滞。**对行业而言，它揭示了一个错位**：当所有人都在生成这条路上加码时，证明这条路——也就是把架构工具带向 L4→L5 的那条路——几乎无人涉足。空缺的高地与拥挤的生成赛道，是同一枚硬币的两面。



3 第三章·物种 A 深度排名：12 家架构设计工具画像

本章对物种 A 的 12 家工具逐一画像。第二章已在地图上标出它们的相对位置与整体格局；本章则走近每一家，回答”它究竟做到了什么、为什么是这个级别”。每家给出工具快照、五维评分（每维附一句可观测理由）、总级别、关键技术判断与证据链。评分方法见第一章 §1.6，完整明细见附录 A。全部评分截至 2026 年 1 月，遵循版本化、不静默漂移原则。

画像采用统一的卡片体例，而非连贯的叙述，是经过考量的选择：评级研究最容易滑向的失误，是用流畅的行文掩盖证据的薄弱。把每一家拆成快照、五维评分、判断、证据链四个固定栏目，等于强制每一个分数都暴露在可核查的结构里——读者可以逐格追问”这个分凭什么”，而不必被一段漂亮的总评说服。这正是本系列与商业分析师报告的分野所在：后者卖结论，本系列卖可复算的依据。

画像按 AI²-ML 总级别降序排列：先是唯一的 L3（SCADE），再是 L2 档六家，最后是 L1 档五家。这一顺序本身就在讲一个故事——从那道唯一被跨过的”推理”门槛出发，沿着能力的台阶逐级而下，读者会看到一个反复出现的模式：越往下，工具越是擅长”标记”而非”证明”，越是把判断的重量留给人。每家标题后的括注标明”厂商 / 阵营”，以便在画像中始终看清它在版图上的位置。

3.1 Ansys SCADE（Synopsys / 形式化验证派·唯一 L3）

工具快照（Snapshot 2026.1）

字段	值
厂商 / 起源	Synopsys（2025-07 收购 Ansys；更早为 Esterel Technologies）；产品品牌仍为 Ansys SCADE
核心产品	SCADE Suite（数据流 / 状态机）+ SCADE Architect（SysML 系统架构）+ Scade One
核心引擎	形式化定义的 Scade 语言 + Design Verifier 模型检查器（Prover PSL + SAT）

字段	值
资质认证	KCG 代码生成器 DO-330 TQL-1 / ISO 26262 ASIL D / IEC 61508 SIL 3 / EN 50128
定位	安全攸关嵌入式软件的”构造正确”工具

五维评分

- D1 约束感知深度：5/5 — Scade 是形式化定义的记法，Design Verifier 经模型检查（K-归纳、PDR/IC3）证明安全性质与运行时错误的不存在，并给出反例。这是本组唯一触及”可验证约束履行”的工具。
- D2 参考架构整合：3/5 — 可复用块、模型库、ICD 生成、从 Rhapsody/MagicDraw 的 SysML 导入；是参考库，非跨产品线治理。
- D3 AI 辅助程度：1/5 — 自动资质代码生成与文档生成属确定性翻译，无生成式架构副驾。
- D4 AI 推理自主度：4/5 — Design Verifier 自主推理、以反例下结论，是内化的验证器；未达 5，因证明范围限于软件模型（非整系统）、且非每次变更自动触发的闭环。
- D5 互操作性：3/5 — SCADE Architect 基于 SysML v1；SAM 2026 R1 宣称对齐 SysML v2 + Python/REST，但 v2 活在 SAM 伴侣而非 Suite 本体。

总级别：L3（加权均值 3.35，唯一 L3，全组最高）

关键技术判断

SCADE 的 L3 建立在一个其余 11 家都不具备的机制上：变更正确性主张来自求解器检查的证明，而非可追溯性矩阵。Design Verifier 以 K-归纳与 PDR/IC3 对安全性质做模型检查，能在所有输入情形下证明性质成立、否则给出反例（D1=5 / D4=4 的依据）。其评分的天花板同样清晰：形式保证以软件 / 模型为范围（SCADE Suite 核心），系统架构层（SCADE Architect）仍是一致性检查；D5=3 因 SysML v2 仅活在 SAM 伴侣、非 Suite 本体。这一”软件层可证、整系统未及”的轮廓，正是它停在 L3 而非更高的原因。落到引言的例子：SCADE 能形式化证明”检测到来车 N 毫秒内对应分区变暗”这条防眩目时序逻辑，却不能证明 84 区的总线负载不会在系统级打破这条时序——前者是它的 L3，后者是空缺的 L4。

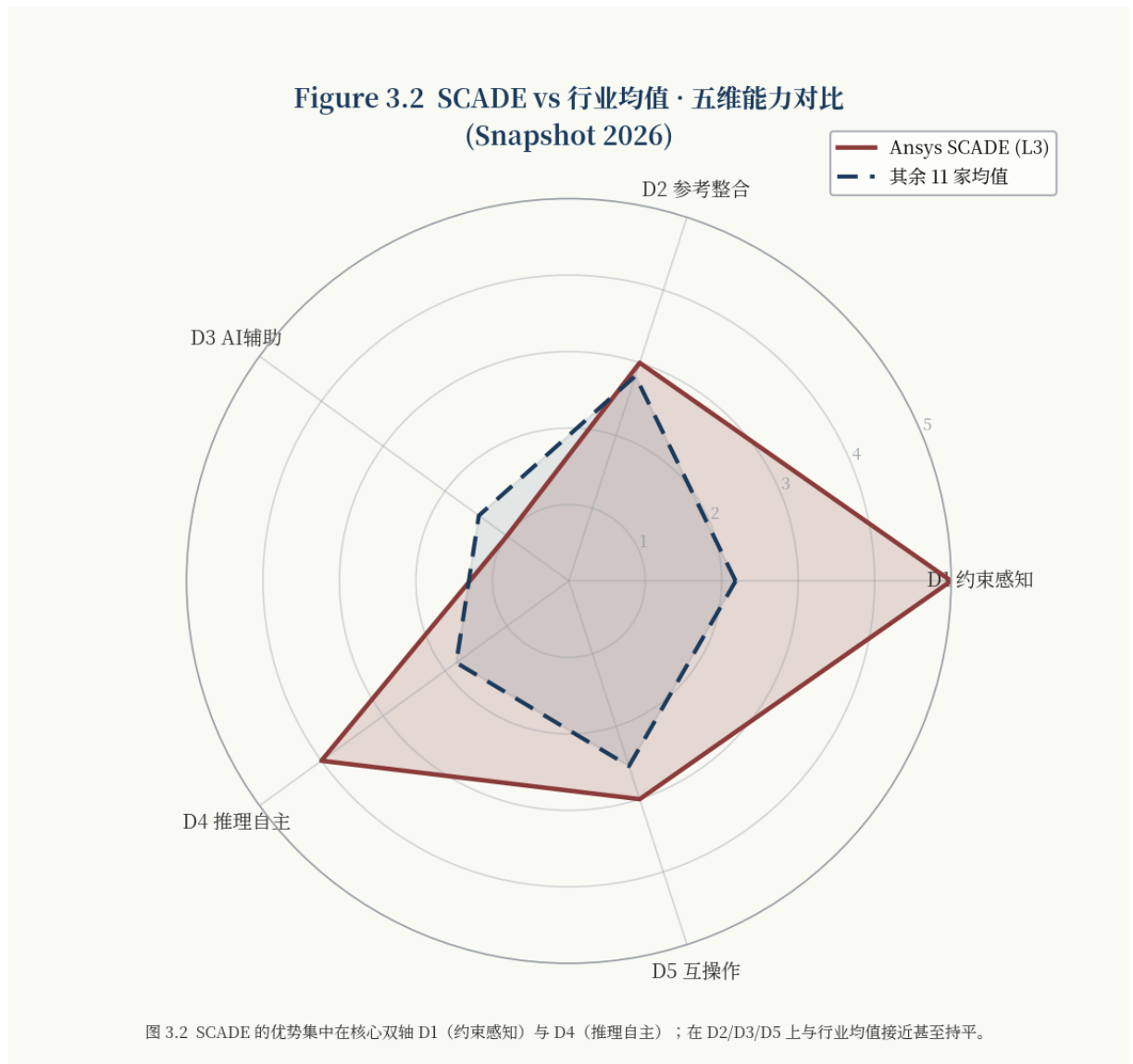


图 3.2 SCAD E vs 行业均值·五维能力对比 (Snapshot 2026)。

证据链： Tier 1: Ansys SCAD E 官方文档、Design Verifier 技术说明、KCG 认证资质 (DO-330 TQL-1 / ISO 26262 ASIL D)、Synopsys 收购完成公告 (2025-07-17)、Ansys 2026 R1 发布说明 (SCAD E Synopsys TPT 整合)； Tier 2: 形式化方法与模型检查学术文献。

3.2 Siemens Capital (Siemens / E/E 数字主线派·L2)

工具快照： Siemens Xcelerator 旗下 E/E 系统开发套件（架构、电气、网络、嵌入式软件、线束），含 Capital X SaaS 与 Simcenter System Architect / Studio；强数字孪生定位。

五维评分

- D1 约束感知深度：2/5 — 内置 metric 与设计规则检查，生成式布线综合用内嵌规则；基于规则，非参考一致性推理。
- D2 参考架构整合：4/5 — 全面的 E/E 数字孪生，跨域共享同一底层模型，集成 Teamcenter/Polarion/Rhapsody，逼近带治理反馈的权威真相源。本组 D2 最高。
- D3 AI 辅助程度：3/5 — 生成式工程：自动布线综合、Simcenter Studio AI 自动探索架构候选并经仿真评估。
- D4 AI 推理自主度：2/5 — Rhapsody-Capital 集成支持变更影响分析，但属跨孪生追溯、人确认，非自主后果推理。
- D5 互操作性：3/5 — Simcenter Studio 执行评估 SysML v2.0 模型；部分 v2 交换，未宣称 API 一致性。

总级别：L2（加权均值 2.65）

关键技术判断：Capital 的强项在”广度与主线”——它把 E/E 全链路收进一个数字孪生，D2 全组最高，生成式辅助也走在前列。但它被门控在 L2，因为推理自主度（D4）仍是追溯式：它能告诉你一处改动牵连到哪些工件，却不替你推出后果、担保正确。

证据链：Tier 1：Siemens 官方产品文档、Simcenter Studio 发布说明；部分能力描述为 Tier 5 营销，未单独支撑高分。

3.3 IBM Rhapsody（IBM / MBSE 老牌·L2）

工具快照：长期的 SysML/UML/AUTOSAR MBSE + 可执行建模 + 代码生成，Harmony 方法，TÜV 认证安全工具包；两条线——经典 Rhapsody（SysML v1）与云原生 Rhapsody Systems Engineering（SE，以 SysML v2 为核）。

五维评分

- D1 约束感知深度：2/5 — 自动一致性检查、模型检查器、可执行仿真；一致性 + 行为仿真，非整系统形式化证明。
- D2 参考架构整合：3/5 — 单源模型 + 到 DOORS Next/ELM 的数字主线、全局配置；SSoT 级但非自主治理。
- D3 AI 辅助程度：3/5 — Engineering AI Hub 1.1 的 MBSE 智能体从自然语言需求生成模型元素；跨步骤生成，人验证。
- D4 AI 推理自主度：2/5 — AI Modeling Assistant 标记潜在不一致、由人决定改什么；经典影响分析为追溯式。
- D5 互操作性：4/5 — Rhapsody SE 以 SysML v2 为核、云原生（REST/GraphQL），暴露 SysML v2 API。本组 D5 并列最高。

总级别：L2（加权均值 2.65）

关键技术判断：Rhapsody 是 SysML v2 落地 (D5) 与生成式 AI (D3) 两条线上的领先者，但其 AI 是生成而非推理——Engineering AI Hub 帮你”生成模型”，AI Modeling Assistant 帮你”标记不一致”，最终的判断与定论仍在人手里 (D4=2)。

证据链：Tier 1: IBM 官方公告 (Engineering AI Hub 1.1, 2025.12.18)、Rhapsody SE 发布说明；Tier 3: 第三方 AI Modeling Assistant 评测。

3.4 Cameo / CATIA Magic (Dassault / No Magic ·SysML 旗舰·L2)

工具快照：业界领先的 SysML MBSE 环境 (原 No Magic MagicDraw/Cameo, 现 CATIA Magic), Cameo Simulation Toolkit 提供可执行/参数化模型, Teamwork Cloud 仓库; R2026x 加入完整 SysML v2。

五维评分

- D1 约束感知深度: 2/5 — 持续一致性检查、参数化约束求解; v2 中需求违例即时标记。约束/一致性检查, 非后果演绎。
- D2 参考架构整合: 3/5 — 可复用库、动态视图; 配合 3DEXPERIENCE/Teamwork Cloud 是强 SSoT, 非自主治理。
- D3 AI 辅助程度: 2/5 — 仿真/权衡自动化; 生成式副驾是第三方/ChatGPT workflow, 非原生。
- D4 AI 推理自主度: 2/5 — 参数化”假设分析”与权衡研究, 影响分析经可追溯性, 规则/追溯触发、人确认。
- D5 互操作性: 4/5 — R2026x 宣称唯一支持 100% SysML v2 + 真双向同步, Teamwork Cloud 宣称完整 API & Services 一致性。本组最强 D5 主张 (厂商自述, 无 OMG 认证, 封顶 4)。

总级别：L2（加权均值 2.45）

关键技术判断：Cameo 是 SysML 建模的事实标准, 互操作性 (D5) 领先。它的推理同样停在追溯式——参数化能算”假设分析”, 但变更影响仍靠人沿可追溯链判断。

证据链：Tier 1: No Magic/Dassault 官方文档 (CATIA Magic R2026x SysML v2 方案); Tier 5: 100% 一致性等营销表述, 谨慎对待。

3.5 Eclipse Capella (Eclipse / Obeo ·开源 MBSE ·L2)

工具快照：实现 Arcadia 方法（运营/系统/逻辑/物理分析）的开源 MBSE 工作台，Thales/Obeo 主导；丰富的模型校验规则框架；SysML v2 经 Eclipse SysON 协同设计。

五维评分

- D1 约束感知深度：3/5 — 分类校验规则（完整性、一致性、完备性、良构性、转换一致性）会考虑相邻工程阶段所建模的内容，带实时模式；阶段感知/参考感知检查。本组 D1 领先档之一。
- D2 参考架构整合：3/5 — Arcadia 分层方法充当参考架构/转换框架，带跨层一致性规则；非双向治理。
- D3 AI 辅助程度：1/5 — 常规自动化与层转换；研究性 LLM/SysON 辅助存在，无发货原生副驾。
- D4 AI 推理自主度：2/5 — 转换一致性规则在层间标记下游不一致，人解决；规则触发标记。
- D5 互操作性：3/5 — SysML Bridge（双向 XMI），SysON 支持 v2 协同设计但 API 部分、尚不适用生产。

总级别：L2（加权均值 2.35）

关键技术判断：Capella 凭 Arcadia 方法的”阶段感知校验”在 D1 达到 3——它的校验会考虑上下游阶段，比纯静态规则更进一步。作为开源工具，它是版图里的重要对照组。

证据链：Tier 1: Capella/Obeo 官方文档；Tier 2: Arcadia 方法学术文献。

3.6 PTC Modeler (PTC / PLM 集成派·L2)

工具快照：SysML/UML/UAF MBSE（原 Integrity Modeler），含变体/PLE、SySim 模型驱动仿真、Reviewer 完整性检查、代码生成、到 Windchill PLM 与 DOORS 的 OSLC；10.2 版（2025 夏）支持核心 SysML 2.0。

五维评分

- D1 约束感知深度：2/5 — Reviewer 检查完整性/正确性/一致性，SySim 验证系统行为；行为校验，非定理级证明。
- D2 参考架构整合：3/5 — Asset Library 支持 SoS/CBD/SOA 复用，OSLC 到 Windchill 数字主线。

- D3 AI 辅助程度：1/5 — 自动化重复性图示任务；无生成式副驾。
- D4 AI 推理自主度：2/5 — 跨模型/PLM 影响分析 + 可追溯，Reviewer 主动指导给出后果/对策；规则触发标记，人执行。
- D5 互操作性：3/5 — 10.2 支持核心 SysML 2.0 + 文本导入；Windchill Modeler 实现 SysML 2.0 OSLC API；部分 v2。

总级别：L2（加权均值 2.10）

关键技术判断：PTC 的差异化在 PLM 集成——它把 MBSE 模型经 OSLC 接进 Windchill 的产品数字主线。推理能力（D4）与同档一致，停在规则触发的标记与指导。

证据链：Tier 1：PTC 官方产品文档、Windchill Modeler 发布说明。

3.7 Ansys medini analyze (Synopsys / 功能安全派·L2)

工具快照：基于模型的功能安全分析工具 (ISO 26262、IEC 61508、ARP4761、SOTIF、ISO 21434)，把 SysML 条目架构与 HARA、FTA、FMEA、FMEDA 集成；ISO 26262 认证。厂商归属：Synopsys (2025-07 收购 Ansys；产品品牌仍为 Ansys medini)。

五维评分

- D1 约束感知深度：3/5 — 校验规则检查与 ISO 26262 的符合性，从需求分配计算并可视化 ASIL 确定/分解；以安全标准为参考的标准一致性检查。本组 D1 领先档之一。
- D2 参考架构整合：2/5 — 带复用、库变更自动更新的元素库；组件/模式复用，非活的 SSOT。
- D3 AI 辅助程度：1/5 — 一键从导入设计生成 FMEA 表；2026 R1 新增 Engineering Copilot（界面内引导式辅助 + 知识库检索）。两者均属常规辅助/检索，非架构生成，故维持 1。
- D4 AI 推理自主度：2/5 — 自动计算 ASIL 经分解的传播、在迭代变更下保持安全分析一致；规则驱动的安全指标传播，人确认（SCADE 之外较好的 D4，仍非自主结论）。
- D5 互操作性：2/5 — 基于 SysML 的条目架构，与 Rhapsody/EA 导入往返（v1 时代）；无 v2 API。

总级别：L2（加权均值 2.05）

关键技术判断：medini 是功能安全维度的专精工具，其 D1=3 来自“以 ISO 26262 为参考的符合性检查”，ASIL 传播是它最接近“推理”的能力。但它的对象是安全分析而非架构设计本身，在版图中是与 SCADE 不同路径的安全派。

证据链： Tier 1: Ansys medini analyze 官方手册、ISO 26262 认证资质。

3.8 Vector PREEvision (Vector / EEA 专用龙头·L1)

工具快照：面向汽车 E/E 架构 (RFLP) 的基于模型平台，建于统一 MOF 数据模型；从需求到物理线束的全链路；PREEvision 26.0 (2026) 重建数据基座。

五维评分

- D1 约束感知深度：2/5 — 可定制一致性模型可手动或后台持续检查，Java metric 框架做非功能基准；无后果演绎。
- D2 参考架构整合：3/5 — 单一统一数据模型，带产品线/变体管理与跨 RFLP 层双向可追溯；架构骨干。
- D3 AI 辅助程度：1/5 — 自动布线器、线束自动布局；无生成式副驾。
- D4 AI 推理自主度：2/5 — 变更标记、变更历史、工单系统；影响经统一模型人工追溯，非自动后果链。
- D5 互操作性：2/5 — SysML 风格图示与 ReqIF 旧标准交换 + 自有 Server/REST API (程序化读写，强于纯静态导出，但非 SysML v2 标准 API)；无 v2。

总级别：L1 (加权均值 1.95)

关键技术判断：PREEvision 是汽车 E/E 架构工具的专用龙头，其统一数据模型在 D2 达到 3、贯穿 RFLP 全链路。但它的能力主体是建模与一致性检查，推理停在人工追溯，故落在 L1 顶部。它是三部曲多次提及的产业基准。

证据链： Tier 1: Vector PREEvision 官方文档。

3.9 Sparx Systems Enterprise Architect (Sparx / 广装机通用建模·L1)

工具快照：广泛部署、低成本的 UML/SysML/BPMN/ArchiMate 平台；SysML 1.5 全支持、仿真、Pro Cloud Server；新 Trechoro 为 KerML 原生 SysML v2 环境 (开发中)。

五维评分

- D1 约束感知深度：2/5 — 对照 UML/SysML profile 约束的模型校验、参数化/约束块、Testpoints；一致性/约束检查。

- D2 参考架构整合: 2/5 — 可复用 profile/MDG 技术、需求集成; 库复用, 非活的参考约束源。
- D3 AI 辅助程度: 1/5 — MDA 转换 + COM/.NET 自动化; 无生成式副驾。
- D4 AI 推理自主度: 2/5 — 关系矩阵 + 层次视图给出即时影响分析, 可在接受需求变更前评估后果; 追溯式影响, 人判断。
- D5 互操作性: 3/5 — Trechoro KerML 原生 v2 创作、v2 文本导入导出、COM/.NET 脚本 (非 REST API); 部分 v2。

总级别: L1 (加权均值 1.95)

关键技术判断: Sparx 以覆盖广度与低价著称, 是装机量最大的建模工具之一。能力均衡但不突出, 推理为追溯式, 落在 L1。

证据链: Tier 1: Sparx Systems 官方文档。

3.10 MathWorks System Composer (MathWorks / Simulink 耦合派·L1)

工具快照: Simulink 的架构建模附加组件 (组件/端口/接口、分配、原型、视图), 经 Embedded Coder 做 AUTOSAR 代码生成, 经 Simulink/Stateflow 做行为, 经 Requirements Toolbox 做需求可追溯。

五维评分

- D1 约束感知深度: 2/5 — 基于原型的属性检查、MATLAB 分析函数; 一致性/分析, 非约束证明。(相邻 Simulink Design Verifier 做形式化性质证明, 但在行为层、非架构层。)
- D2 参考架构整合: 2/5 — 模型到模型分配、视图、可复用原型 profile; 库复用, 非活治理。
- D3 AI 辅助程度: 1/5 — 分析函数自动化、报告生成; AI 面向基于模型的设计, 非架构生成。
- D4 AI 推理自主度: 2/5 — Requirements Toolbox + System Composer 支持需求变更对设计的影响分析 (经可追溯链接); 人判断。
- D5 互操作性: 2/5 — 经 MATLAB 表导入导出; 无 v2 API (SysML 支持基本 v1 时代)。

总级别: L1 (加权均值 1.80)

关键技术判断: System Composer 的差异化在与 Simulink/Stateflow 的深度耦合——架构与行为仿真同栈。但其架构层推理为追溯式；值得注意的是同栈的 Simulink Design Verifier 具备形式化证明能力，只是作用在行为层而非架构层，故未抬升本工具的架构维评分。

证据链: Tier 1: MathWorks System Composer 官方文档。

3.11 dSPACE SystemDesk (dSPACE / AUTOSAR 配置派·L1)

工具快照: AUTOSAR (Classic + Adaptive) 软件架构创作工具，建模 SWC/接口，生成虚拟 ECU 供 VEOS SIL 仿真；与 TargetLink + EB tresos 搭配。

五维评分

- D1 约束感知深度: 2/5 — 强校验检查 AUTOSAR 架构、报告 ARXML 问题，合理性检查 + 修正建议；对照 AUTOSAR 元模型的模式/规则检查。
- D2 参考架构整合: 3/5 — AUTOSAR 标准本身即参考架构/SSoT, ARXML 交换强制符合结构。
- D3 AI 辅助程度: 1/5 — 自动 V-ECU 生成/配置、修正建议；无架构生成。
- D4 AI 推理自主度: 1/5 — 校验标记 ARXML 问题以免日后撞上；几乎无变更后推理。
- D5 互操作性: 1/5 — 仅 ARXML 交换；无 SysML v2。

总级别: L1 (加权均值 1.55)

关键技术判断: SystemDesk 的强项是”标准即参考”(D2=3, AUTOSAR 结构强制)，但它本职是 AUTOSAR 软件架构的配置与校验，而非架构推理——D4/D5 全组最低。在量规意义上，它更接近配置工具而非推理工具。

证据链: Tier 1: dSPACE SystemDesk 官方手册。

3.12 ETAS ISOLAR (ETAS / Bosch ·AUTOSAR 配置派·L1)

工具快照: ETAS (Bosch) AUTOSAR 工具，ISOLAR-A (应用 SW 架构/配置) 与 ISOLAR-B (RTA-BSW 基础软件配置)，ARXML 创作/校验、RTE 生成。

五维评分

- D1 约束感知深度: 2/5 — ARXML 编辑器创建并校验描述软件架构与接口的 XML 工件，对照 AUTOSAR 规则；模式/规则校验。

- D2 参考架构整合：3/5 — AUTOSAR 标准 = 参考架构，强制符合的 ARXML。
- D3 AI 辅助程度：1/5 — 从描述生成 RTE/代码的自动化；无生成式架构。
- D4 AI 推理自主度：1/5 — 校验/符合性检查标记偏差；无后果推理。
- D5 互操作性：1/5 — 仅 ARXML 交换；无 SysML v2。

总级别：L1（加权均值 1.55）

关键技术判断：ISOLAR 与 SystemDesk 同属 AUTOSAR 配置派，能力画像高度相似——标准即参考（D2=3）、低 D4/D5。它在版图中代表 AUTOSAR 生态的配置工具一脉。

证据链：ETAS（Bosch）AUTOSAR 工具，评分以二手权威来源为主（一手官方文档公开度有限）；其能力画像与同属 AUTOSAR 配置派的 SystemDesk 高度一致，L1 定级与画像相符，俟官方文档进一步公开后于后续版本复核。

3.13 小结：一道分水岭，一片空地

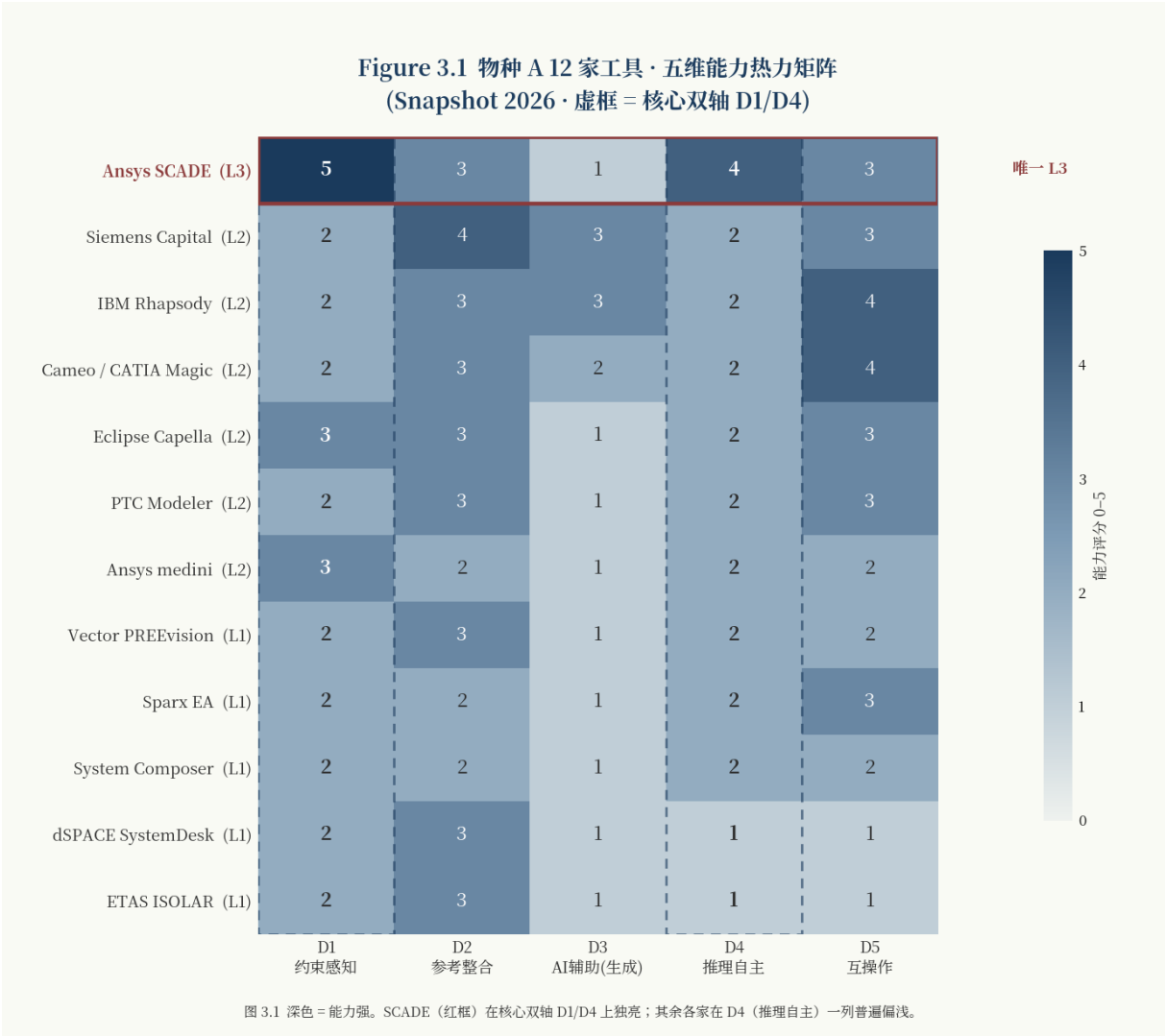


图 3.1 物种 A 12 家工具·五维能力热力矩阵 (Snapshot 2026)。

12 家画像合起来，呈现一幅清晰的图景。唯一的 L3 (SCADE) 与其余十一家之间，隔着的不是分数高低，而是”检查 vs 推理”的能力性质——只有它能证明，其余都只能标记。L2 档六家 (Capital、Rhapsody、Cameo、Capella、PTC、medini) 各有所长——数字主线、SysML v2、生成式 AI、阶段感知校验、安全符合性——却无一例外被门控在推理自主度 (D4=2) 的”追溯式标记”。L1 档五家或为建模通用工具、或为 AUTOSAR 配置工具，推理由可追溯性驱动。

而 L4 与 L5，是一片空地。没有任何一家工具把推理从软件扩展到整系统、把验证内化到每一次变更。这片空地是什么、谁会去填、以及为什么是现在——是第五章 Roadmap 要回答的问题。

3.13.1 具象锚：同一次变更，五类工具各交付什么

前 13 节的评分是抽象的分数。本节把它们落到第一章引入的那个例子上——自适应远光灯 ADB，一次“光束分区 12 区 → 84 区”的变更——看不同级别的工具面对同一次变更，究竟交付什么、把什么留给人。

这次变更沿约束链级联（摄像头分辨率 → ECU 算力 → LED 通道 ×7 → 总线负载 → 线束散热），最终撞上三个要害问号：**防眩目时序仍满足？ASIL-B 仍成立？ECE 眩光仍达标？**

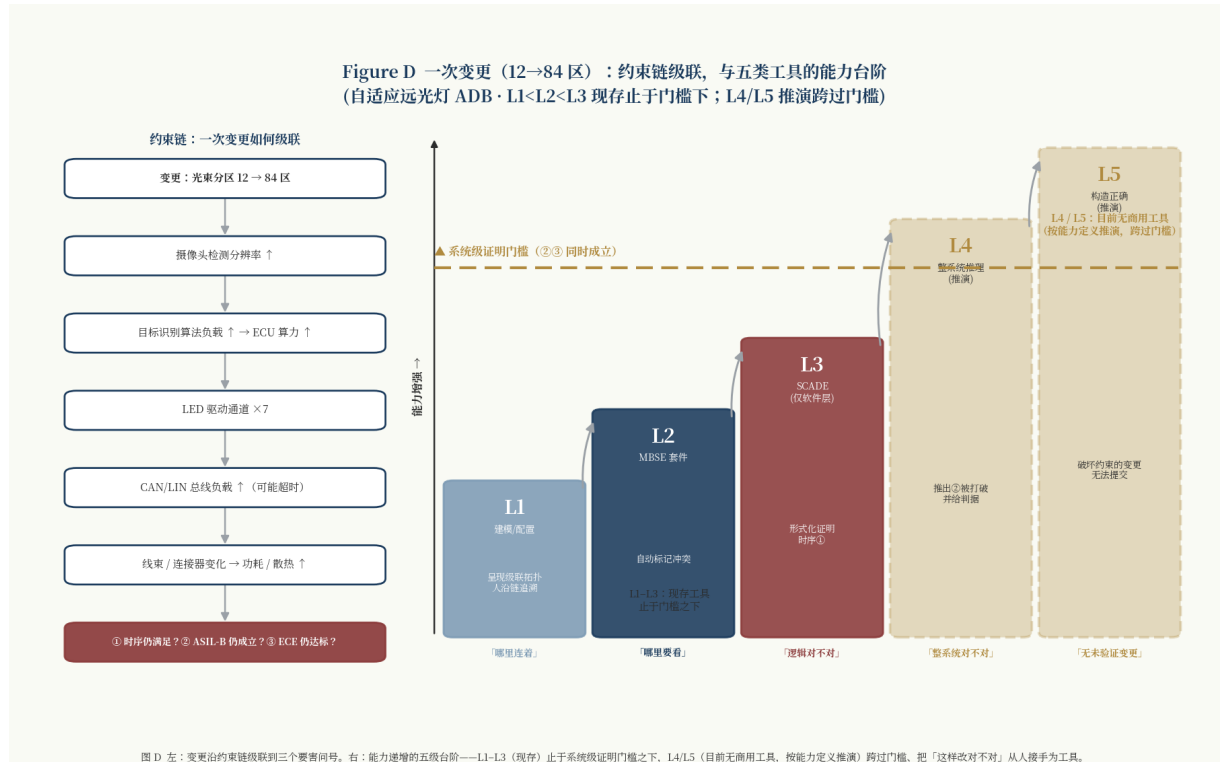


图 D 一次变更（12→84 区）的级联链，与五类工具的能力台阶。L1-L3（现存）止于系统级证明门槛之下；L4/L5（目前无商用工具，按能力定义推演）跨过门槛。

五类工具对这次变更各做什么（格式：代表·动作 — 止于·留给人·评分依据。L1-L3 现存；L4/L5 目前无商用工具，按能力定义推演，斜体）：

L1·建模 / 配置（代表：PREEvision 一类）- 变更录入统一数据模型，拓扑显示受影响连接 / ECU / 线束；总线超时、时序满足性须导出参数另算。留给人：沿链逐段判断。**能答”哪里连着”，不能答”改了还对不对”**（D1=2 / D4=2）。

L2·MBSE 套件（代表：Cameo / Capital 一类）- 跨域共享模型，参数化算 84 区总线负载、超标即自动标记，受影响工件一并点亮；只标”需复核”不给结论。留给人：下判断。**能答”哪里要看”，不能证明”改了仍安全”**（D1 3 / D4=2，门控 L2）。

L3·SCADE（仅软件层）- 防眩目控制逻辑写成形式化模型，模型检查数学证明时序性质成立、否则给反例；不证整系统（84 区总线负载是否破坏时序、ASIL-B / ECE 系统级是否成立）。**能证”逻辑对”，不能证”整系统对”**（D1=5 / D4=4，仅软件层）。

L4 ·整系统推理（推演·目前无商用工具） - 输入 “12→84 区”，自主沿级联链推后果（算总线实际负载与最坏延迟、推 ECU 算力缺口、跨域推 影响），给有判据的结论——例：“**变更打破约束 ：总线延迟使最坏响应时间超 ASIL-B 阈值**”，附处置。人从复核者 → 监督拍板者。

L5 ·构造正确（推演·目前无商用工具） - 破坏约束的变更无法提交：若 84 区在现硬件下无法同时满足 ，变更被接受前即以数学证明拦下，并给约束包络边界——例：“**现总线带宽下最多 72 区可保 ASIL-B 与 ECE 同时成立**”。无未经验证的变更能进入设计。

对照表·同一次变更（12 区 → 84 区），五类工具交付什么

		L3			
约束链上的问题	L1（建模/配置）	L2（MBSE 套件）	（SCADE, 仅软件层）	L4（整系统推理·推演）	L5（构造正确·推演）
哪些工件受影响？	拓扑显示连接， 人沿链追溯	自动追溯并 点亮	软件模型内 自动	自主推出整条级联	自主推出 + 实时维护
总线负载是否超时？	导出参数另 行计算	参数化自动 计算并标记	不在范围	算出实际值并判定	变更前即保证不超时
接口/一致性是否破坏？	一致性检查可 标记	自动标记不一致	形式化保证（软件层）	整系统一致性推理	破坏一致性的变更被拒
防眩目时序是否仍满足？	人工分析	提示” 需复核”， 人判断	形式化证明 （给反例）	整系统级推出是否满足	不满足则变更无法提交
ASIL-B 安全论证 是否仍成立？	人工重做	安全工件点亮， 人复核	软件层可证；系统级仍需人	推出” 被打破” 并给判据	证明 成立方可提交
ECE 眩光是否仍达标？	人工核对法规	追溯法规条目， 人核对	不在范围	跨域推出是否达标	证明 成立方可提交
谁来回答” 这样改还对不对”？	人	人 （工具帮他找该看哪里）	工具 （仅软件逻辑层）	工具 （整系统，人监督拍板）	工具 （破坏约束的变更无法提交）

这张表把 12 家的评分落成了一句话：**L1/L2 那格写着”人”，L3 写着”工具（仅软件层）”，而 L4/L5 那两格——目前无任何商用工具能填上——才是把”这样改到底对不对”从人的肩上接过去的地方。**对照表右下角那格的演变，就是整个品类的轨迹，

也是第五章那片”空缺高地”的具体形状：不是”更强的生成”，而是”这次变更后，能否在整系统层面证明 同时成立”。今天答案是否。

4 第四章·相邻地带：另外三个物种

前三章聚焦物种 A——架构设计 / MBSE / 推理工具——因为只有它与 AI²-ML 这把尺子同源，能够被有意义地深度排名。但第二章的全景图上，物种 A 并非孤立存在：它的周围还盘踞着另外三个物种——B（电气 / 线束详细设计）、C（数据 / 拆解 / 归纳平台）、D（AI 原生 / SDV 新玩家）。它们与物种 A 共享同一片产业疆域，却各自遵循不同的认识论、占据不同的工作层级。

本章不对这三个物种做深度排名。理由在第二章已经埋下：用 AI²-ML 去量它们，等于用一把专为“约束推理”打造的尺子，去衡量“详细设计”“数据归纳”“生成式探索”这些本质不同的能力，结论只会失真。本章要做的是另外两件更基础的事：其一，说清每个物种**是什么、为何不与 A 同尺**；其二，标明它**留给本系列哪一篇续作**。

这两件事合起来，赋予本章双重作用。一方面，它**界定了物种 A 的边界**——你只有看清 A 不是线束工具、不是拆解平台、不是套壳生成器，才真正理解 A 是什么；边界是定义的另一形式。另一方面，它**撑起了整个系列**——B、C、D 各自是一座尚未勘探的山头，本篇在地图上标出它们的位置、点明各自的去向，把登顶留给后续的专属篇章。没有这一章，这篇论文只是一份孤立的工具榜单；有了它，它才成为一个系列的开篇。

4.0.1 物种 B：电气 / 线束详细设计——演绎，但下游

- **是什么**：把既定架构落实为电气 / 线束设计——原理图、拓扑、连接器、线径 / 压降、制造数据。
 - **代表**：Zuken E3.series · Siemens Solid Edge Electrical · Aras Innovator
 - **位置**：演绎—下游。与 A 同享“演绎”（电气规则检查 / 一致性校验），但作用于架构既定之后，不决定架构本身。
 - **为何不与 A 同尺**：差别在**工作层级**，非成熟度高低。North Star §2.1：架构设计决定系统长什么样，详细设计落实既定样子。用 AI²-ML 量线束工具 = 错配。物种 B 有自己的成熟度维度（制造集成 / 规则库深度 / PLM 协同），需专属尺子。
 - **留给**：电气 / 线束设计工具续篇。本篇仅标位置——A 的下游邻居，非 A 的低配版。
-

4.0.2 物种 C：数据 / 拆解 / 归纳平台——上游，但归纳

- **是什么：**从真实数据归纳洞察供架构决策——竞品拆解、物料 / 成本基准、市场样本、供应链情报。
 - **代表：**A2MAC1（拆解数据 → AI 决策支持）·Palantir（数据整合 / 决策）·Munro & Associates（工程拆解）
 - **位置：**归纳—上游。与 A 同享”上游”（影响架构决策、作用于概念阶段），但方法是**归纳**非演绎——归纳”别人怎么做”“行业基准成本几何”，非从约束推”这样改对不对”。
 - **为何不与 A 同尺：**C 答”别人怎么做”（归纳 / 统计 / 参照）；A 答”这样改是否正确”（演绎 / 形式 / 证明）。矩阵大灯参考哪些竞品 = C 的问题；84 区后 ASIL-B 是否成立 = A 的问题。两者皆上游、皆有价值，但**归纳给不出演绎的保证——再多竞品样本，也不能证明这次变更正确。**
 - **留给：**数据 / 决策平台续篇。本篇中，C 是 A 的一面镜子——同在上游，方法相反，照出 A “演绎”身份的独特之处。
-

4.0.3 物种 D：AI 原生 / SDV 新玩家——横跨推理轴的探索带

- **是什么：**以 AI 为内核、面向 SDV 的新一代——生成式 AI 建模新创、SysML v2 云原生工具、大模型接入架构工作流的探索者。
- **代表：**digital.auto / ./pulse（开放 SDV 原型社区）·新一代云原生 SysML v2 工具·其他 AI 原生新创
- **位置：**一条探索带，非一个点。多数以生成式 AI（归纳）起步，有抱负者向”演绎—上游”空缺高地试探，落点未定。其形态映射 L4→L5 前沿尚未被占据。
- **为何不与 A 同尺（目前）：**多数现阶段为生成式工具（落 D3 非 D4，合第二章脱钩发现），或未成熟到可稳定评分。非”另一类工作”（如 B）或”另一种认识论”（如 C），而是**同一疆域上尚未定型的新生力量**——部分未来或正面进入 A 排名，部分或停在生成式辅助。落点未定 → 静态深排为时尚早。
- **留给：**Roadmap 主线主角（第五章三路径之一 = D 的”直接起跳”）。专属续篇追踪：哪些玩家真转向可验证推理，哪些停在生成。

披露：digital.auto / ./pulse 等开放社区是本机构关注对象。本篇坚持客观画像，不因任何潜在协作而美化或回避；评价标准与其余物种一致。

4.0.4 小结：四物种，一张地图

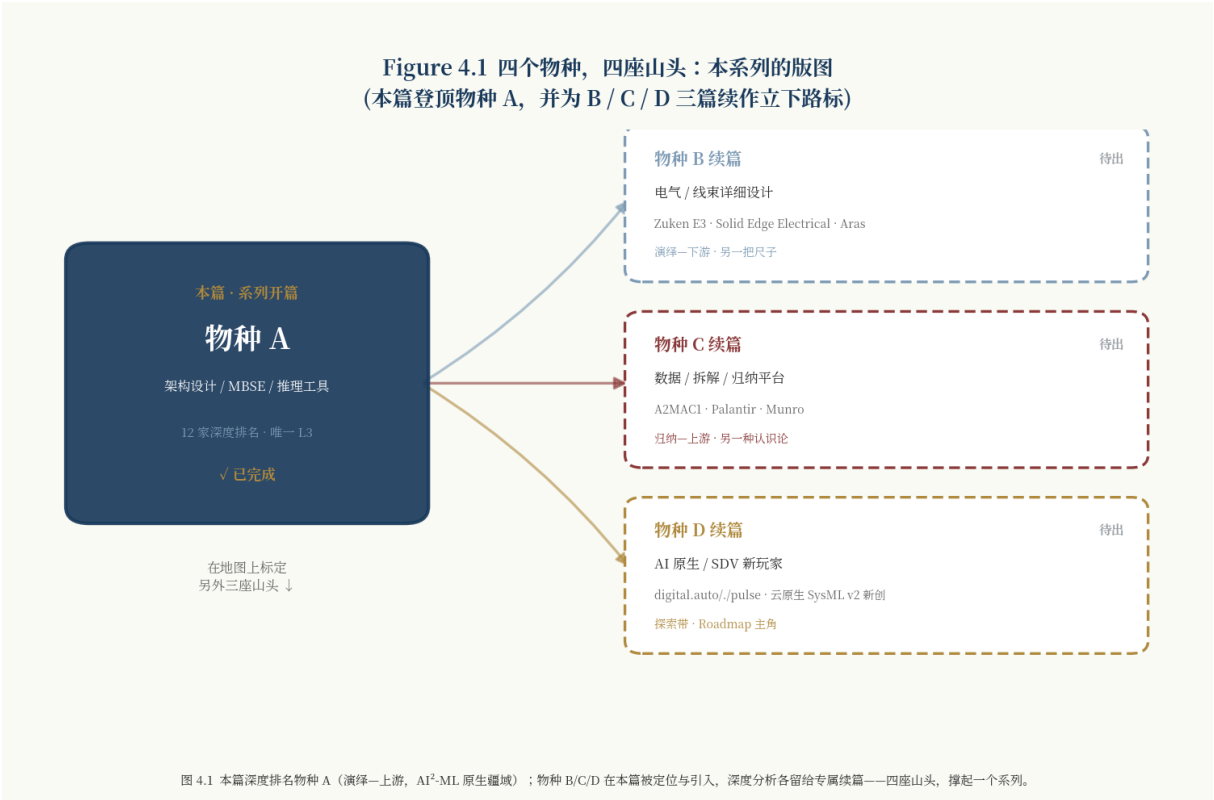


图 4.1 四个物种与本系列版图：本篇深排物种 A，B/C/D 各留专属续篇。
放回第二章全景图：

- **物种 A**（演绎—上游）= 本篇主角，AI²-ML 原生疆域，唯一深排。
- **物种 B**（演绎—下游）= A 的下游邻居：同演绎，落实而非决定架构。
- **物种 C**（归纳—上游）= A 的认识论对照：同上游，归纳而非演绎。
- **物种 D**（横跨推理轴）= 尚未定型的探索带：Roadmap 主角。

结论：AI²-ML 为物种 A 量身打造，不该强加给另外三物种。 B 有详细设计成熟度，C 有归纳洞察成熟度，D 仍在选择要成为什么。各标地图、各留续篇——既守本篇评分严谨（不混尺），又为系列铺路。四个物种，四座山头：本篇登顶第一座，为另外三座立下路标。

4.0.5 另一种相邻：工具篇与 OEM 篇的连接

前四节的”相邻”，是工具版图**内部**的相邻——四个物种互为邻居。还有一种相邻在版图**之外**：本篇（工具篇）与三部曲（OEM 篇）这两个系列之间的相邻。它们丈量的是同一产业的两端——OEM 篇量”被设计的架构”，工具篇量”设计架构的工具”——理应能接上。本节给出这个连接点。

口径说明：公开渠道无“某 OEM 用某具体架构工具”的 Tier 1–3 硬证据（选型属工程机密，厂商仅称“服务全球 OEM”，三部曲未记录具体工具）。故本节**不做点名连接**（依赖 Tier 4–5 推测，违反证据门槛），改做**结构性连接**：以三部曲原文判断，论证工具能力上限如何解释 OEM 架构债务。

支点·三部曲的同步律。 D2《2026 全球汽车 E/E 架构成熟度评估》§2.1（22 家 OEM 沿 $AR \times AI^2\text{-ML}$ 二维定位）的核心判断：

22 家 OEM 紧贴对角线分布.....架构成熟度与架构演进工具的成熟度高度耦合，几乎没有 OEM 能长期偏离对角线。换言之，系统的先进程度，受制于组织迭代它的工具与方法论——**架构成熟度更接近组织能力的产物，而非可独立采购的部件。**——D2 §2.1

要点：OEM 架构能力 其架构演进工具能力，同步、互锁。**无须点名任何 OEM 用何工具。**

工具篇补上缺失的一半。三部曲从 OEM 侧观察到同步律，留下一问：市面工具能把 OEM 托举到哪一级？本篇作答。因果链：

- **其一·工具上限 = L3，且限软件层。** 12 家中最强（SCADE）仅触及 L3，形式化证明限软件 / 模型层，系统架构层仍是一致性检查；其余 11 家 L1–L2。L4/L5 空缺——**无任何可采购工具能在整系统层面、对每次架构变更提供可验证推理。**
- **其二·此即“架构债务”之源。**三部曲：多数 OEM 的 AR 略高于其 $AI^2\text{-ML}$ ，差距 0.5–1.0，称“架构债务”。本篇给出来源：**推进到 AR_4 所需的架构推理支撑，已超出任何可外购工具的上限（L3）。**差距 = 系统野心 – 工具供给。
- **其三·债务只能由组织能力偿付。**三部曲：“缺乏配套的演进工具，外购架构难以持续维护与迭代。”机理：**工具上限 = L3，OEM 维持 AR_4 的缺口，须由自有方法 / 流程 / 自研工具链填补——故领先 OEM（全栈自研路线者）多自建或深度定制工具链，而非单纯采购。**根因：**可采购的部件有 L3 天花板。**

连接点。 OEM 篇（需求侧）证明“架构成熟度被工具锁定”；工具篇（供给侧）证明“工具的锁 = L3”。合成结论：

这十年的“架构债务”，本质是工具供给缺口：系统野心已迈向 AR_4 ，可外购的架构推理工具仍停在 L3。谁把工具上限从 L3 推向 L4/L5，谁就为全行业偿付此债务提供可能。这也预告了第五章“空缺高地”的产业意义：它不是工具品类自己的空白，而是被整个行业架构债务标记出来的空白。

5 第五章·Roadmap：空缺的高地，与为什么是现在

前三章呈现的是一幅静态的格局：唯一触及 L3 的 SCADÉ、贴地分布的长尾、以及那片始终空缺的 L4–L5 高地。这幅格局回答了”今天的工具站在哪里”，却还没有回答一个更要紧的问题——这片高地为什么至今无人登顶，而它又是否终将被填平。

本章转向时间维度，沿三个层层递进的问题展开：那片空缺的高地究竟是什么形状、为什么偏偏是此刻才谈得上去攀登、以及谁最有可能去填补它。需要先行声明本章的克制：它不预测具体厂商的胜负，那既非一份能力框架应当承担的任务，也会把严谨的分析降格为行业八卦。本章只做一件事——勾勒前沿的形状，并说明这片空地为何在 2025–2026 这个时间点上，第一次具备了被攀登的现实条件，从而对整个产业具有了超出工具品类自身的战略意义。

5.0.1 为什么是现在：两个代际转折交汇

高地长期空缺，未必即将被填。判据：支撑攀爬的条件是否刚成熟。2025–2026 年，两个独立转折交汇，使”约束感知的架构推理”首次具备从 L3 上攀的现实条件。

转折一·SysML v2 落地——“跨工具约束推理”在标准层首次可能。 - OMG 2025-07-21 最终采纳 SysML v2 + API & Services。 - v1 = 图形记法，模型困于私有格式；v2 = 形式化语义元模型 + 标准 API，模型可程序化查询 / 校验 / 跨工具流通。 - 呼应第一章：互操作（D5）是通往顶端的使能基础设施；演绎推理 + 跨链可信传递，前提是模型有形式化语义、机器可读。 - **2026 = 地基从规范变产品的落地之年。**

转折二·AI 代际跃迁——“自动推理”从规则引擎跨入可规模化范畴。 - 但（第二章脱钩发现）当前发货的 AI 几乎全落 D3（生成），非 D4（推理）。 - 机会所在：生成浪潮证明”机器能参与设计”，却把”机器能否证明设计正确”留在原地。 - **交汇含义：地基（SysML v2）刚浇好，所有人在地基上盖”生成”的楼，几乎无人盖”推理与验证”的楼。**

5.0.2 高地是什么：从 L3 到 L5 的两步

第三章已定位边界：唯一的 L3（SCADÉ）只在软件 / 模型层触及”可验证约束履行”，系统架构层仍是一致性检查。通往顶端 = 两步。

第一步·L3 → L4：演绎推理从软件层 → 整系统架构层。 - 现状：SCADE 证明软件模型性质；架构层变更影响（域控算力调整 → 功能安全 / 拓扑 / 线束 / 成本）仍靠人追溯。 - L4 标志：工具在整系统架构层自主推导后果链 + 给可执行结论（影响范围 + 处置），人退居监督担保。 - 技术路径：约束求解 / 模型检查能力从软件模型 → 系统架构模型（神经-符号方法弥合地带：生成由 NN，验证由符号引擎）。

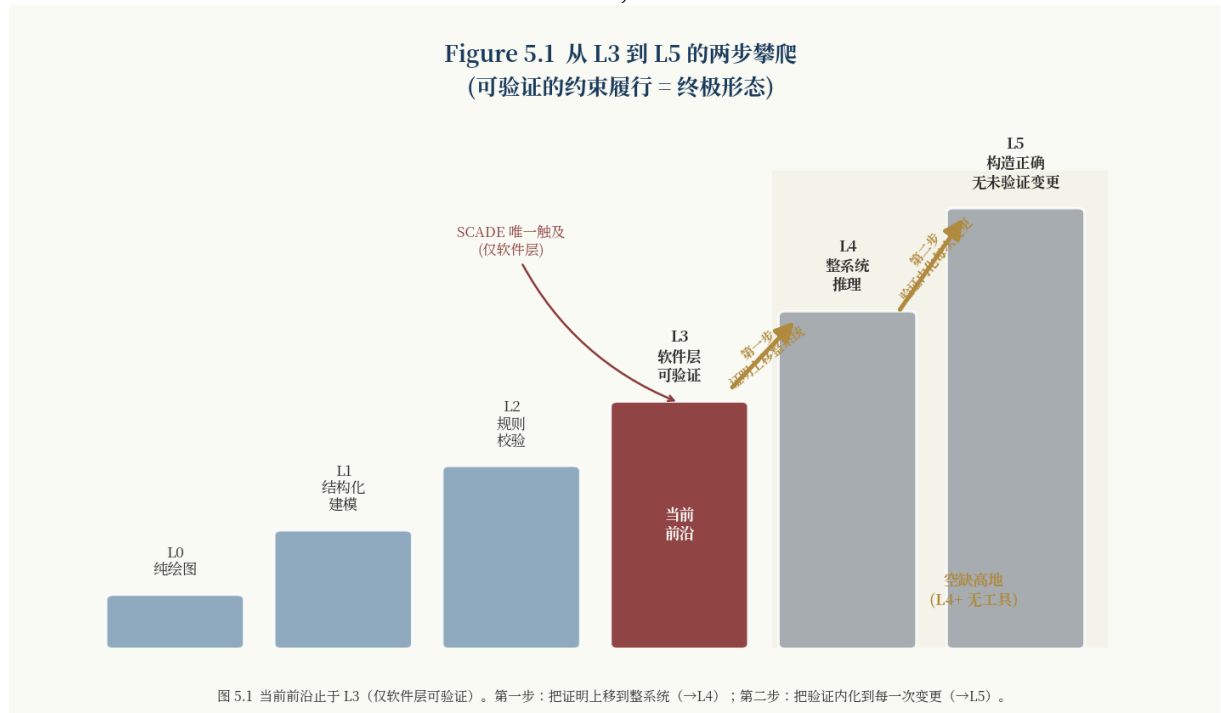


图 5.1 从 L3 到 L5 的两步攀爬 (Snapshot 2026)。

第二步·L4 → L5：验证内化，覆盖每一次变更。 - L5 “无人”，= “无未经验证的变更”。 - 要求：工具内化验证器，每次架构变更（乃至运行时自我重构）前，先数学证明不破坏约束包络——即第一章”构造正确”。 - 实质：举证责任彻底内化；“替代”在安全攸关领域唯一可能的形态——验证被可信交给机器。

两步合起来定义高地形状：不是”更强的生成”，而是”可验证的推理”——与当前 AI 竞赛正交的方向。

5.0.3 谁会去填：三条路径

三类玩家，不同起点，各有优势 / 障碍。

路径一·在位 MBSE 套件向上延伸（补 D4）。 - 代表：Capital · Rhapsody · Cameo (L2 领先，已握 D2 数字主线 / D5 v2 / D3 生成) - 攀爬：在既有平台嵌入约束推理引擎，补推理自主 (D4) - 优势：生态 + 装机量 · 障碍：重心在生成与协作，演绎 / 形式化内核与现架构未必契合

路径二·安全 / 形式化派向外扩展（补整系统）。 - 代表：SCADE · medini（已在推理一侧，D1 高） - 攀爬：形式化能力从软件 / 安全分析 → 整系统架构 - 优势：已拥”证明”内核 · 障碍：形式化可扩展性（模型检查从软件 → 开放系统架构，计算代价 + 规格完整性未解）

路径三·AI 原生新玩家从空地直接起跳（绕过遗产）。 - 代表：物种 D 探索带 - 攀爬：以”约束感知推理”为架构原点，为推理而非生成从头设计 - 优势：无历史包袱 · 障碍：缺数据 / 生态 / 安全认证积累，多数仍用生成式范式、未转向验证

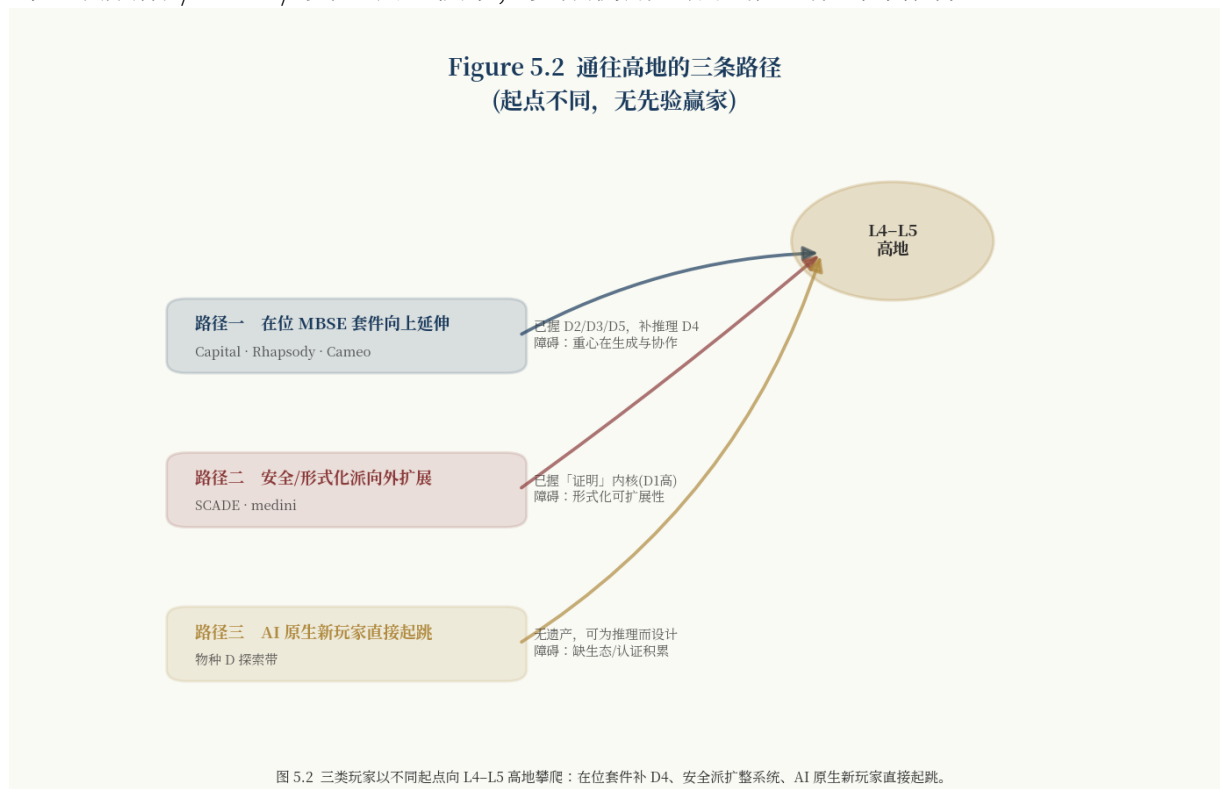


图 5.2 通往 L4-L5 高地的三条路径（Snapshot 2026）。

无先验赢家。共同指向一个事实：**高地是真实空地，且首次具备被攀爬的条件。谁先把”可验证的约束履行”从软件层 → 整系统层 → 每一次变更，谁就定义这个品类的下一个十年。**

5.0.4 结语：可验证的约束履行，作为一把参照系

开篇之问：谁在为架构师赋能？三章丈量的答案——**当前无任何工具把架构师从”判断变更是否正确”中解放出来。能画、能查、能生成，不能证明。**

发现的价值不止于给 12 家排座次，而在确立一把**参照系**：以”可验证的约束履行”为架构工具终极形态，品类的位置 / 差距 / 方向首次有了可丈量坐标。如三部曲以 AR4 为汽车架构参照系，AI²-ML 的 L5（“无未经验证的变更”）为架构工具提供北极星——未

必很快抵达，但指明方向：**设计工具六十年演化的终点，不是让机器替人画图，而是让每一次变更都可被证明。**

此即 Sketchpad（1963）那个约束求解器的意义：**设计工具最深的身份，从第一天起就是约束推理引擎。**绕六十年，品类正回到起点——只是这一次，约束推理的对象，是一个正在变成机器人的复杂系统。



6 附录

6.1 附录 A ·12×5 完整评分明细表（可独立复算）

评分方法见第一章 §1.6（**五维定义、加权 D1 25% / D4 25% / D3 20% / D2 15% / D5 15%、核心双轴短板约束**）与 §1.7（**数据来源金字塔、证据门槛、能力声明封顶规则、时间窗**）。每格：分数 + 主要证据 Tier。门槛：每维须 2 个 Tier 1–3 证据。个别格标注证据成分（如研究性 / 营销来源）已剔除后保守取分。所有权截至 2026.1 快照 + 截止后补充。全部评分截至 2026.1。

6.1.1 总表（按总级别降序）

#	工具 (所有 权)	D1	D2	D3	D4	D5	加权均 值	总级别
1	SCADE (Synop- sys, 原 Ansys; 品牌 Ansys SCADE)	5	3	1	4	3	3.35	L3
2	Siemens Capital	2	4	3	2	3	2.65	L2
3	IBM Rhap- sody	2	3	3	2	4	2.65	L2
4	Cameo / CATIA Magic	2	3	2	2	4	2.45	L2
5	Eclipse Capella	3	3	1	2	3	2.35	L2

#	工具 (所有 权)	D1	D2	D3	D4	D5	加权均 值	总级别
6	PTC Mod- eler	2	3	1	2	3	2.10	L2
7	medini (Synop- sys, 原 Ansys; 品牌 Ansys medini)	3	2	1	2	2	2.05	L2
8	Vector PREE- vision	2	3	1	2	2	1.95	L1
9	Sparx EA	2	2	1	2	3	1.95	L1
10	MathWorks System Com- poser	2	2	1	2	2	1.80	L1
11	dSPACE Sys- temDesk	2	3	1	1	1	1.55	L1
12	ETAS ISO- LAR ¹	2	3	1	1	1	1.55	L1

6.1.2 每格证据 Tier 与门槛检验

¹ETAS(Bosch)的一手官方技术文档公开度有限,本行评分以二手权威来源(行业技术资料、AUTOSAR 生态文档)为主。其能力画像与同属 AUTOSAR 配置派的 dSPACE SystemDesk 高度一致,评分置于 L1 与该画像相符;俟 ETAS 官方文档进一步公开后于后续版本复核。此处来源情况的说明遵循本系列对证据不对称的一贯处理(参见 §1.7 与三部曲对一手文档有限对象的标注惯例)。

工具	D1	D2	D3	D4	D5	门槛
SCADE	5 [T1+T2]	3 [T1]	1 [T1]	4 [T1+T2]	3 [T1]	
Siemens Capital	2 [T1]	4 [T1]	3 [T1+T3]	2 [T1]	3 [T1]	
IBM Rhapsody	2 [T1]	3 [T1]	3 [T1]	2 [T1+T3]	4 [T1, 封顶]	
Cameo / CATIA	2 [T1]	3 [T1]	2 [T2+T3]	2 [T1]	4 [T1, 封顶]	
Magic Eclipse	3 [T1+T2]	3 [T1]	1 [T2]	2 [T1]	3 [T1]	(D3 标研究成分)
Capella						
PTC Modeler	2 [T1]	3 [T1]	1 [T1]	2 [T1]	3 [T1]	
medini	3 [T1]	2 [T1]	1 [T1]	2 [T1]	2 [T1]	
Vector	2 [T1]	3 [T1]	1 [T1]	2 [T1]	2 [T1]	
PREESion						
Sparx EA	2 [T1]	2 [T1]	1 [T1]	2 [T1]	3 [T1+T4]	
System Composer	2 [T1]	2 [T1]	1 [T1+T5]	2 [T1]	2 [T1]	(D3 标营销成分)
dSPACE	2 [T1]	3 [T1]	1 [T1]	1 [T1]	1 [T1]	
SystemDesk						
ETAS ISOLAR ²	2 [T2/3]	3 [T2]	1 [T2/3]	1 [T2/3]	1 [T2/3]	(见脚注)

6.1.3 评分规则与封顶记录

- 核心双轴短板约束：宣称 L_n 须 D1 与 D4 各自独立 n。SCADE 加权均值 3.35，D1=5、D4=4 均 3，落 L3；其 D1/D4 虽允许 L4，但均值与 D3/D5 短板下取整压在 L3。

²ETAS(Bosch)的一手官方技术文档公开度有限,本行评分以二手权威来源(行业技术资料、AUTOSAR 生态文档)为主。其能力画像与同属 AUTOSAR 配置派的 dSPACE SystemDesk 高度一致,评分置于 L1 与该画像相符;俟 ETAS 官方文档进一步公开后于后续版本复核。此处来源情况的说明遵循本系列对证据不对称的一贯处理(参见 §1.7 与三部曲对一手文档有限对象的标注惯例)。

- **能力声明封顶规则** (§1.7.3): Cameo D5=4、Rhapsody SE D5=4 — “100% SysML v2 / API 一致性” 为厂商自述, OMG 一致性测试套件在窗口内不存在, 封顶 4。**全场无 D5=5。**
- **证据成分提示:** Capella D3、System Composer D3 (剔除研究 / 营销成分后保守取 1)。ETAS ISOLAR 的来源情况见脚注。

6.1.4 三处复评结论 (A 步骤)

- **所有权订正:** SCADE、medini → Synopsys (原 Ansys)。不改技术评分。
- **medini D3 维持 1:** Engineering Copilot 属引导式辅助 / 检索, 非架构生成。
- **PREEvision D5 维持 2:** 自有 REST API 非 SysML v2 标准 API。

6.2 附录 B ·来源清单 (按工具·Tier 分级)

仅列支撑评分的主要一手与权威来源。Tier 定义见 §1.7.1。窗口 2024.1–2026.1 + 截止后补充。

全局·所有权 - [T1] Synopsys 完成收购 Ansys 公告, 2025-07-17 (约 350 亿美元; Ansys 退市)。- [T1] Ansys 2026 R1 发布, 2026-03-11 (收购后首个重大版本; medini Engineering Copilot、SCADE TPT 整合)。

1. SCADE (Synopsys) - [T1] Ansys SCADE 官方文档、Design Verifier 技术说明 (Prover PSL / SAT、K-归纳 / PDR)。- [T1] KCG 认证资质: DO-330 TQL-1 / ISO 26262 ASIL D / IEC 61508 SIL 3 / EN 50128。- [T1] Ansys 2026 R1 发布说明 (SCADE Synopsys TPT 整合)。- [T2] 形式化方法与模型检查学术文献。

2. Siemens Capital - [T1] Siemens EE-systems 官方 blog (Capital 2026 年度版, 2026-02-27)。- [T1] Simcenter Studio 发布说明 (生成式架构探索、SysML v2.0 模型评估)。

3. IBM Rhapsody - [T1] IBM 公告: Rhapsody SE v1.5(2025-10-14)、Engineering AI Hub 1.1 (MBSE 智能体)。- [T1] Rhapsody SE 产品页 (SysML v2 核心、REST/GraphQL)。

4. Cameo / CATIA Magic - [T1] Dassault / No Magic 官方文档 (CATIA Magic R2026x, 2025-12, SysML v2 + UAF 1.3)。- [T1] Teamwork Cloud API & Services 文档 (自述完整一致性, 封顶处理)。

5. Eclipse Capella - [T1] Capella / Obeo 官方文档; Eclipse SysON (mbse-syson.org, 8 周发布周期, “尚不适用生产”)。- [T2] Arcadia 方法学术文献。

6. PTC Modeler - [T1] PTC 官方产品文档 (Modeler 10.2, 核心 SysML 2.0)、Windchill Modeler 发布说明 (SysML 2.0 OSLC API)。

7. medini (Synopsys) - [T1] Ansys medini analyze 官方手册、ISO 26262 认证资质。 - [T1] Ansys 2026 R1 (Engineering Copilot; VC FSM 双向互操作; PLM 集成 Cameo/Codebeamer/DOORS/Jama)。

8. Vector PREEvision - [T1] Vector PREEvision 官方文档 (26.0=2026-02-27; 26.1=2026-05; Server API + REST API)。

9. Sparx Enterprise Architect - [T1] Sparx Systems 官方 (EA 17.1 Build 1716, 2026-01-29; Trechoro KerML v2 开发中)。

10. MathWorks System Composer - [T1] MathWorks System Composer / SysML 产品页 (SysML 1.x; v2 “计划支持”)。

11. dSPACE SystemDesk - [T1] dSPACE SystemDesk 官方手册 (AUTOSAR Classic/Adaptive; V-ECU; ARXML)。

12. ETAS ISOLAR³ - [T2/3] ETAS (Bosch) AUTOSAR 工具 (ISOLAR-A/B, ARXML 创作 / 校验、RTE 生成); 评分以二手权威来源为主。

(附录 A 评分明细 + 附录 B 来源清单。完整方法论见第一章 §1.6–1.7; 事实核查过程见 *D4-FACTCHECK-DELTA*、复评见 *D4-SCORECARD-REVALIDATION*。)

³ETAS(Bosch)的一手官方技术文档公开度有限,本行评分以二手权威来源(行业技术资料、AUTOSAR 生态文档)为主。其能力画像与同属 AUTOSAR 配置派的 dSPACE SystemDesk 高度一致,评分置于 L1 与该画像相符;俟 ETAS 官方文档进一步公开后于后续版本复核。此处来源情况的说明遵循本系列对证据不对称的一贯处理(参见 §1.7 与三部曲对一手文档有限对象的标注惯例)。