

汽车行业深度报告

AI 系列深度 22 期: 智能驾驶 VLA 模型的架构范式与厂商路线

增持（维持）

2026 年 08 月 07 日

证券分析师 黄细里

执业证书: S0600520010001

021-60199793

huangxl@dwzq.com.cn

研究助理 童明祺

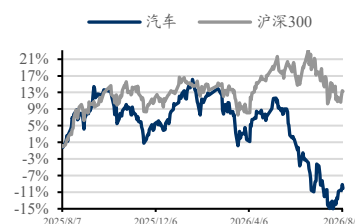
执业证书: S0600125080005

tongmq@dwzq.com.cn

投资要点

- **智能驾驶 VLA 的核心**, 是将视觉感知、语言理解与动作决策整合进同一策略模型, 并通过显式动作头弥合 VLM4AD 留下的动作生成缺口。从技术演进看, 自动驾驶经历模块化流水线、端到端、VLM4AD 再到 VLA4AD: VLM 能解释场景但难以直接生成可靠控制, VLA 则进一步把动作作为模型输出, 使语义理解与驾驶策略更紧密耦合。
- **VLA4AD 的核心架构**通常包括视觉编码器、语言处理器和动作解码器三部分。视觉编码器将摄像头、激光雷达等异构传感输入转化为隐层特征, 并通过 BEV、点云、体素等方式增强三维空间理解; 语言处理器引入预训练大模型、LoRA 或 RAG 等方法注入驾驶知识; 动作解码器再将融合后的多模态表征转化为轨迹、控制量或动作 token。
- **VLA4AD 自身**可归纳为被动解释器、规划协作者、统一动作生成器、推理中枢四个阶段; 主要挑战包括鲁棒性、实时性能、数据与标注、多模态对齐、多智能体协同和评测标准化。这意味着 VLA 并非单一固定架构, 而是自动驾驶端到端路线向大模型、多模态和动作生成继续演进的总称。
- **厂商路线已出现差异化**。Waymo EMMA 以 Gemini 驱动纯端到端, 并将导航、状态、轨迹等统一到语言空间; 理想从 E2E+VLM 双系统升级至 MindVLA 统一架构, 以 3D 高斯、MoE 和扩散动作解码组织三维理解、语义推理和轨迹生成; 元戎启行以 40B VLA 基座设置 Driver、Analyst、Critic 三角色; 千里 CoWorld-VLA 则通过多专家世界模型 Token 构建规划相关隐空间表征。
- **小鹏第二代 VLA 进一步弱化显式语言中间层**, 以原生多模态 Tokenizer 和视觉思维链直达动作, 本质上趋近视觉-动作路线。由此可见, 头部厂商在是否保留显式语言推理层上已经分化: 一类强调语言作为慢思考和可解释中枢, 另一类强调隐式视觉表征和动作直出。后续竞争将集中在实时性、数据闭环、评测口径和量产可靠性。
- **我们认为, 数字 AI 底座模型发展路径已逐渐清晰、价值在于确定性**, 但其难以实现物理世界的建模闭环; 我们认为物理 AI 将成为物理世界的 AI 解决方案, 从而在 AI 的当前发展阶段成为增量更大的方向, 智能驾驶作为最先跑通闭环的代表场景, 我们认为应重点关注。
- **风险提示**: 技术迭代不及预期的风险, 大模型厂商 ARR 增速放缓的风险, 云服务商资本开支不及预期的风险, 智能驾驶及机器人商业化落地不及预期的风险。

行业走势



相关研究

《AI 系列深度 17 期: 视觉智能为何引入 Transformer? 》

2026-08-06

《AI 系列深度 16 期: 语言智能为何从 RNN 转向 Transformer? 》

2026-08-06

内容目录

1. 智驾 VLA：从端到端到动作生成的范式升级	4
1.1. 自动驾驶技术演进：从模块化流水线到 VLA4AD	4
1.2. VLA4AD 的核心架构：输入、三大模块与输出	4
1.3. VLA4AD 的阶段演进：从解释器到推理中枢	5
1.4. VLA4AD 的主要局限	6
2. 代表性模型：头部厂商 VLA 路线开始分化	7
2.1. Waymo EMMA：基于 Gemini 的纯端到端架构	7
2.2. 理想：从 E2E+VLM 双系统到 MindVLA 统一架构	7
2.3. 元戎启行：40B VLA 基座的 Driver/Analyst/Critic 三角色	8
2.4. 千里科技 CoWorld-VLA：多专家世界模型 Token	8
3. 小鹏第二代 VLA：弱化显式语言层，强化视觉思维链	9
3.1. 设计动机：标准 VLA 的语言转译瓶颈	9
3.2. 架构特征：原生多模态 Tokenizer 与视觉思维链	10
3.3. 去显式语言层的边界：趋近 VA 的 VLA 变体	11
3.4. 工程与量产：算力、数据与落地节奏	11
4. 风险提示	11

图表目录

图 1: 自动驾驶范式比较.....4

图 2: VLA4AD 四阶段演进时间轴.....6

图 3: CoWorld-VLA 在 NAVSIM v1 上的 PDMS 贡献拆解与同类对比9

图 4: 小鹏第二代 VLA 对外披露的能力提升与体验改善10

表 1: VLA4AD 动作解码器的三类技术路线.....5

表 2: VLA4AD 自身的四阶段演进.....6

表 3: 代表性厂商 VLA 模型的架构与公开指标对比9

表 4: 小鹏第二代 VLA 对外披露的关键参数与指标10

1. 智驾 VLA：从端到端到动作生成的范式升级

1.1. 自动驾驶技术演进：从模块化流水线到 VLA4AD

自动驾驶的整体技术演进可划分为四个范式。早期主流方案为经典模块化流水线，将任务拆分为感知、预测、规划、控制四个模块；其优势在于每个组件可独立开发、测试与优化，工业界早期广泛采用，缺陷是模块间误差会级联传播、信息碎片化，难以处理长尾的极端场景。

为解决模块化缺陷，研究转向端到端自动驾驶，将原始传感器输入直接映射为控制指令，核心是视觉到动作 Vision-to-Action 的映射。该范式消除了模块边界，可端到端优化，但在语义稳定性、可解释性和语言交互能力上仍存在不足。在此基础上，研究进一步将视觉语言模型 VLM 引入自动驾驶 VLM4AD，把感知与自然语言推理统一到同一嵌入空间，利用大模型常识先验提升对罕见场景的泛化能力；但该路线仍以感知解释为核心，语言输出与控制耦合较弱，存在动作生成缺口，即 VLM 能解释场景，但难以直接生成可靠控制决策。

VLA4AD 范式的关键变化，是把视觉感知、语言理解和控制决策整合到同一个策略中，通过显式动作头 action head 弥合上述动作生成缺口。其目标能力包括：遵循自然语言指令，例如让行救护车、超过指定卡车；输出决策理由以支持事后验证；利用互联网数据中的常识先验提升罕见场景泛化能力。

图1：自动驾驶范式比较



数据来源：《A Survey on Vision-Language-Action Models for Autonomous Driving》，东吴证券研究所

1.2. VLA4AD 的核心架构：输入、三大模块与输出

VLA4AD 在输入侧整合多模态信号：视觉输入已从单目前摄演进至双目、多摄环视系统，并常将图像投影为 BEV 表征以强化空间推理；激光雷达提供 3D 结构，毫米波雷达提供速度信息，IMU 用于运动跟踪，GPS 用于定位，自车本体数据则提供车辆状态；

语言输入包括环境查询、任务规则解析及语音交互。

VLA4AD 核心模块可拆分为视觉编码器、语言处理器和动作解码器。视觉编码器将原始异构传感器数据转换为与语言空间对齐的隐层特征，常用 DINOv2、ConvNeXt-V2、CLIP 等骨干，并加入 BEV 投影、点云编码、体素等 3D 空间增强。语言处理器采用预训练大语言模型作为解码器，并通过 LoRA 低秩适应等技术微调，或利用检索增强 RAG 注入驾驶领域知识。动作解码器将融合后的多模态特征转换为驾驶动作，主要有三种技术路线。在输出侧，早期系统直接输出转向角、油门、刹车等原始控制信号，后续工作更多输出未来一段时间的轨迹规划，再交由下游模型预测控制 MPC 模块跟踪执行。

表1: VLA4AD 动作解码器的三类技术路线

技术路线	机制	特点	代表/适用
自回归分词器	将连续轨迹点/控制信号离散化为 token 序列，并按自回归方式逐个预测	实现简单，与 LLM 生成范式天然兼容	AutoVLA 等
扩散头	基于扩散模型，从高斯噪声逐步采样出连续控制信号或轨迹	擅长连续动作空间，生成轨迹更平滑	DiffVLA、理想 MindVLA
分层控制器	高层由 LLM 输出自然语言子目标，底层由 PID/MPC 转换为具体控制	高低层解耦，可解释性较强	ORION

数据来源：《A Survey on Vision-Language-Action Models for Autonomous Driving》，《DiffVLA: Vision-Language Guided Diffusion Planning for Autonomous Driving》等其他文献，东吴证券研究所

1.3. VLA4AD 的阶段演进：从解释器到推理中枢

VLA4AD 的内部演进可划分为四个阶段。预 VLA 阶段主要集中在 2023 年前后，多模态大模型更多充当被动解释器，代表如 DriveGPT-4，可生成场景自然语言描述或抽象驾驶意图标签，实际控制仍由传统模块完成，语言不参与决策。模块化 VLA 阶段主要集中在 2024 年至 2025 年初，LLM 由解释者升级为引导车辆控制的主动路径规划媒介，采用感知、语言理解、规划、控制的多阶段流水线，语言作为中间表示，代表如 OpenDriveVLA。

统一端到端 VLA 阶段集中在 2025 年，研究者构建单一策略网络，在一次前向传播中直接将多模态传感器流与语言指令映射为控制信号，语言与视觉特征深度融合，代表如 Waymo EMMA。推理增强 VLA 阶段是最新方向，将长时域推理、记忆与交互整合进控制循环，要求模型在执行动作前先进行逻辑推理并回顾历史状态，代表如 ORION，其引入 QT-Former 记忆模块。

图2: VLA4AD 四阶段演进时间轴



数据来源：《A Survey on Vision-Language-Action Models for Autonomous Driving》，东吴证券研究所

表2: VLA4AD 自身的四阶段演进

阶段	语言 LLM 角色	时间	代表模型	主要问题
预 VLA 阶段	被动解释器	约 2023 年为主	DriveGPT-4	描述带来延迟；处理无关细节；描述与控制间存在语义隔阂
模块化 VLA	规划协作者	2024—2025 年初	OpenDriveVLA	多阶段流水线误差在边界被放大；响应变慢
统一端到端 VLA	动作生成器	2025 年为主	EMMA	反应快但长时程推理弱；决策解释粒度不足
推理增强 VLA	推理中枢，引入 CoT 与记忆	最新阶段	ORION	将长时域推理、记忆、交互整合进控制循环

数据来源：《A Survey on Vision-Language-Action Models for Autonomous Driving》，《EMMA: End-to-End Multimodal Model for Autonomous Driving》等其他文献，东吴证券研究所

1.4. VLA4AD 的主要局限

我们认为，当前 VLA4AD 从技术验证走向规模化落地仍面临六种挑战，其中数据瓶颈与算力约束是短期最核心的制约因素。

- 一为鲁棒性与可靠性不足，包括语言幻觉、传感器噪声以及形式化验证的缺失。
- 二为实时性能约束，车载端需实现不低于 30Hz 的推理帧率，而视觉 Transformer 叠加 LLM 的算力需求极高，需要 token 削减、硬件感知量化、事件触发推理以及蒸馏与 MoE 稀疏化等手段方可满足车规要求。
- 三为数据与标注瓶颈，视觉—语言—动作三模态配对监督数据稀缺。
- 四为多模态对齐，当前工作仍以相机为中心，激光雷达、雷达、高精地图与时序状态仅被部分融合。

五为多智能体社会复杂性，受约束的车间 V2V 语言、信任与安全等长期问题仍处于探索初期。

六为域适应与评测体系，包括仿真到现实（sim-to-real）迁移、分布外（OOD）场景泛化、持续学习与标准化测试均有待成熟。长期来看，基础规模驱动模型、神经符号安全内核、车队级持续学习等方向，有望成为突破上述瓶颈的核心演进路径。

2. 代表性模型：头部厂商 VLA 路线开始分化

2.1. Waymo EMMA：基于 Gemini 的纯端到端架构

Waymo EMMA 构建于多模态大模型 Gemini 之上，将原始相机传感数据直接映射为多种驾驶专用输出，包括规划轨迹、感知目标与道路图 road graph 要素。针对传统模块化链路中感知、预测、规划依赖人工接口和中间表示割裂的问题，EMMA 采用纯端到端范式，从原始驾驶输入直接生成规划输出。

EMMA 的核心机制是统一语言空间：将导航指令、自车状态等非传感输入以及轨迹、3D 位置等输出统一表示为自然语言文本，从而在同一模型内联合处理多任务，并以任务特定提示生成各任务输出。EMMA 在 nuScenes 运动规划任务上取得 SOTA 表现，在 Waymo Open Motion Dataset WOMD 上取得有竞争力的结果；引入思维链 CoT 推理后，端到端规划性能进一步提升 6.7%，并提供可解释的决策理由；规划、检测与道路图的多任务联合训练在三个领域均带来提升。论文同时披露其局限：仅能处理少量图像帧、不含激光雷达/雷达等精确 3D 感知模态，且计算开销较大。

2.2. 理想：从 E2E+VLM 双系统到 MindVLA 统一架构

理想智能驾驶沿“高精地图约束下的模块化 NOA（2021 年）→多传感器融合感知平台（2022 年）→BEV/OCC 感知大模型（2023 年）→E2E+VLM 双系统（2024 年）→VLA 司机大模型（2025 年）”逐阶段演进。在 E2E+VLM 阶段，系统升级为快慢双系统：E2E 快系统承担高频轨迹生成，VLM 慢系统负责复杂交通语义、规则理解与长尾场景判断。该阶段对应的方法论可参考清华大学与理想等于 2024 年发表的《DriveVLM》，其将规划拆解为“场景描述—关键对象影响分析—分层规划”三段式推理链，并通过 DriveVLM-Dual 接入传统 3D 感知与规划模块，为关键对象补充三维位置、朝向与历史轨迹，再将低频参考轨迹细化为高频可执行轨迹。

MindVLA 将三维空间理解、交通语义推理与驾驶动作生成统一到同一框架：以摄像头与激光雷达等多源输入经 3D 编码器生成三维场景特征，采用 3D 高斯 3D Gaussian 作为中间表征并进行自监督训练；语言模型 MindGPT 采用 MoE 混合专家架构与稀疏注意力以在扩容的同时控制激活参数；动作侧由 Diffusion 扩散解码器将动作词元解码为优化轨迹，整个推理在车端实时运行。VLA 司机大模型于 2025 年 8 月随理想 i8 开启交

付并面向 AD Max 车主升级；2026 年 3 月，理想在 GTC 2026 发布下一代基础模型 MindVLA-o1，对外表述为原生多模态 MoE Transformer，并支持“快思考/慢思考”双模式。

2.3. 元戎启行：40B VLA 基座的 Driver/Analyst/Critic 三角色

元戎启行 40B VLA 基座模型将视觉表征、语言推理、动作生成与轨迹评价统一建模，并使同一基座同时承担三类角色：Driver 基于视觉输入生成实时驾驶动作；Analyst 识别关键事件并以因果推理解释决策原因；Critic 从安全性、舒适性与类人驾驶风格维度评估轨迹。三类角色共享统一 Foundation Model 底座，并通过评价结果反向驱动数据闭环。

在训练范式上，预训练阶段引入视频预测任务，将轨迹监督扩展到像素级视频预测，以提升对真实世界动态场景的表征能力与大规模驾驶视频的数据利用率；中期训练将常规端到端驾驶扩展为三类任务：V+A，视觉到动作；V+A→L，基于动作序列解释事件；V→L+A，先形成场景描述与决策逻辑再输出轨迹。公司公开口径显示，其数据处理周期由通常的 5 天以上压缩至约 12 小时；目标 2026 年底达到 100 万辆，并提及 2025 年 10 月在第三方高阶智驾供应商中单月份额接近 40%。

2.4. 千里科技 CoWorld-VLA：多专家世界模型 Token

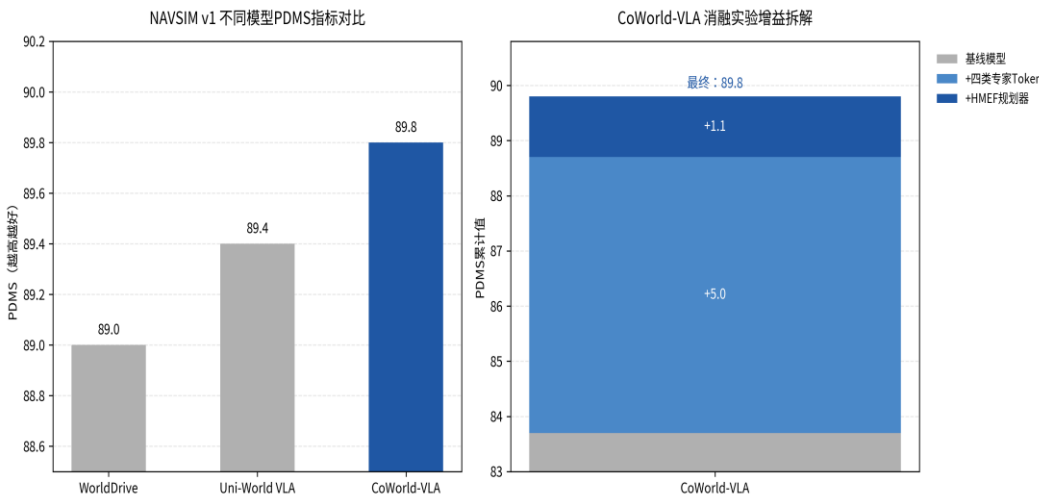
CoWorld-VLA 论文提出，传统 VLA 通常将图像、导航指令与车辆状态直接映射为未来轨迹，缺少可被规划器调用的中间世界状态；文本思维链难以保留连续的空间与运动信息，单一世界表征又不足以覆盖交通交互、道路几何与未来动态。

其核心创新是在 VLM 隐空间构造四类专家 Token：交通交互 Token，编码交通参与者之间的交互意图；道路几何 Token，编码道路布局、车道结构与空间约束；动态演化 Token，编码未来场景变化与运动趋势；自车轨迹 Token，编码自车未来行为目标，共同构成面向规划的隐空间世界表征。

训练分三阶段：先训练动作条件化世界模型预测未来场景演化；再以 Qwen3-VL 为骨干，将 VLM 隐状态分别对齐到 JEPA 语义、VGGT 几何、Wan 动态与轨迹监督，形成四类专家 Token；最后以 HMEF 多专家融合规划器，将场景 Token 与专家 Token 作为条件在动作空间中扩散去噪，生成连续自车轨迹。

NAVSIM v1 测试结果显示，CoWorld-VLA PDMS 达 89.8，超越 Uni-World VLA (89.4)、WorldDrive (89.0)；同时，FVD 指标 32.7（越低效果越好），表现领先。消融显示四类专家 Token 带来约+5.0 PDMS、HMEF 规划器进一步带来约+1.1 PDMS。千里科技，董事长印奇的 AI+车体系以与吉利合作的千里浩瀚智驾及世界模型为主线，CoWorld-VLA 为其智驾团队的研究成果之一。

图3: CoWorld-VLA 在 NAVSIM v1 上的 PDMS 贡献拆解与同类对比



数据来源:《CoWorld-VLA: Thinking in a Multi-Expert World Model for Autonomous Driving》,《NAVSIM: Data-Driven Non-Reactive Autonomous Vehicle Simulation and Benchmarking》, 东吴证券研究所

表3: 代表性厂商 VLA 模型的架构与公开指标对比

模型	视觉/输入	语言/推理	动作解码	公开指标
Waymo EMMA	环视相机, 纯视觉, 不含 LiDAR/雷达	Gemini 统一语言空间, CoT 推理	轨迹文本 token	nuScenes 规划 SOTA; CoT 使端到端规划提升 6.7%
理想 MindVLA	摄像头+LiDAR, 3D 高斯中间表征	MindGPT, MoE+稀疏注意力,	Diffusion 扩散轨迹	公司口径: 车端实时运行
元戎启行 40B VLA	视觉输入为主	Analyst 因果推理, 三角色,	实时驾驶动作	数据闭环周期由 5 天以上压缩至约 12 小时
千里 CoWorld-VLA	图像+导航+车辆状态	Qwen3-VL 隐空间 四类专家 Token	扩散式 HMEF 规划器	NAVSIM v1 PDMS 89.8; FVD 32.7
小鹏 第二代 VLA	原生多模态 Tokenizer, 视觉/语音/文本输入	32 倍超密视觉思维链, 弱化文本语言	端到端直接输出动作	公司口径: 较传统 CoT 预测误差降低 33%

数据来源:《EMMA: End-to-End Multimodal Model for Autonomous Driving》,《CoWorld-VLA: Thinking in a Multi-Expert World Model for Autonomous Driving》等其他文献, 理想、小鹏等公司公开资料, 东吴证券研究所; 各指标口径不一, 不可直接横向比较

3. 小鹏第二代 VLA: 弱化显式语言层, 强化视觉思维链

3.1. 设计动机: 标准 VLA 的语言转译瓶颈

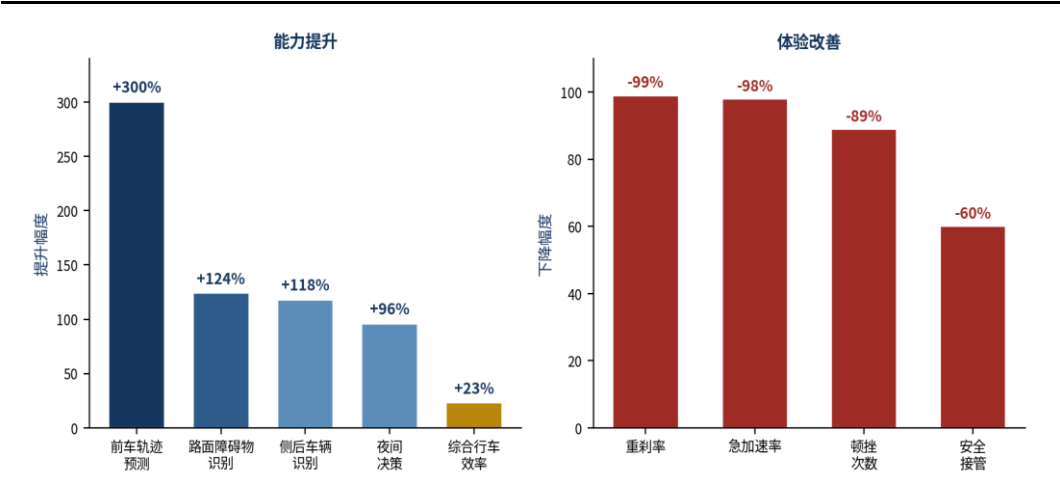
小鹏将行业传统 VLA 概括为视觉 V、语言 L、动作 A 架构, 其中语言转译作为中

间环节，被公司认为会带来较高信息损耗，且文字难以完整承载图像中的细节信息。基于此，小鹏第二代 VLA 去掉语言转译环节，实现从视觉信号到动作指令的端到端直接输出。

3.2. 架构特征：原生多模态 Tokenizer 与视觉思维链

小鹏第二代 VLA 定位为原生多模态物理世界大模型，整合视觉、语音与文本输入。其设计了原生多模态 Tokenizer，实现离散 Token 与连续表征的无损转化，以更高效率完成早期融合并降低单一模态偏差；模型采用 32 倍超密视觉推理思维链 Visual CoT，公开口径显示其推理效率提升 32 倍、相比传统 CoT 预测误差降低 33%；模型具备多模态输出能力，可生成语音、视觉、动作与行为，支撑世界模型、仿真与 self-play 强化学习，也是原生舱驾一体的基础框架。

图4：小鹏第二代 VLA 对外披露的能力提升与体验改善



数据来源：小鹏汽车官网、2025 小鹏科技日及 2026 THE FUTURE 媒体日公开披露，东吴证券研究所整理

表4：小鹏第二代 VLA 对外披露的关键参数与指标

维度	指标	数值/表述	披露场合
模型	视觉思维链推理效率	提升 32 倍	2026 THE FUTURE 媒体日
模型	相比传统 CoT 的预测误差	降低 33%	2026 THE FUTURE 媒体日
算力	车端算力-Ultra 版	2,250 TOPS	2025 小鹏科技日
算力	车端模型参数规模	数十亿级，行业普遍千万级	2025 小鹏科技日
数据	训练数据量	近 1 亿 clips，累计 50PB，	科技日/推送公告
数据	每版模型训练数据	约 4 万亿 Tokens	推送公告
体验	综合行车效率	提升 23%	2026 THE FUTURE 媒体日
落地	首发外部客户	大众汽车，图灵 AI 芯片获定点	2025 小鹏科技日

数据来源：小鹏汽车官网、2025 小鹏科技日及 2026 THE FUTURE 媒体日公开披露，东吴证券研究所整理

3.3. 去显式语言层的边界：趋近 VA 的 VLA 变体

从架构本质看，小鹏第二代 VLA 不再依赖显式语言层对场景进行语言转译，而是将视觉信号经隐式、非语言化 Token 直接映射为动作，并以视觉思维链进行推理，因而在架构形态上趋近于视觉—动作 VA。需客观说明的是，去显式语言层指去掉语言中间转译步骤，而非删除全部语言模态能力。雷锋网报道显示，第二代 VLA 搭载原生多模态 Tokenizer，仍可实现语音、视觉、动作等多模输出。因此我们将其归纳为去显式语言中间层、以隐式表征与视觉思维链直出动作的 VLA 变体；至于语言能力在具体高阶场景中的参与程度，目前尚无公开量化披露。

作为对照，理想 MindVLA-o1 按公司表述保留了语言层，是原生多模态 MoE Transformer，将视觉、语言、行动三种模态联合训练并在同一表示空间对齐，支持快思考/慢思考双模式。由此可见，在是否保留显式语言推理层这一维度上，头部厂商的 VLA 路线已出现分化：小鹏倾向弱化文本语言、强化视觉思维链直达动作，理想等则倾向在统一多模态框架内保留语言推理。

3.4. 工程与量产：算力、数据与落地节奏

为实现第二代 VLA 量产上车，小鹏针对图灵 AI 芯片重新开发了编译器与软件栈，通过芯片、算子和模型全链路优化，最终在 2,250 TOPS 的 Ultra 版车型上搭载数十亿级参数规模模型。行业普遍车端模型参数量为千万级，基座模型编译效率提升约 12 倍。数据方面，第二代 VLA 可直接利用海量真实驾驶视频训练、减少人工标注依赖，训练数据量接近 1 亿 clips、累计约 50PB，每版模型训练数据约 4 万亿 Tokens，自 2025 科技日以来已迭代数百版模型。

落地节奏上，第二代 VLA 于 2025 年 12 月底启动先锋用户共创体验，2026 年第一季度起面向小鹏 Ultra 车型推送；大众汽车成为首发外部客户，图灵 AI 芯片获大众定点。基于第二代 VLA 的小路 NGP 在复杂小路场景的平均接管里程 MPI 较此前提升约 13 倍，综合行车效率提升约 23%。

4. 风险提示

技术迭代不及预期的风险：数字 AI、物理 AI 技术进展低于预期，任务或应用场景拓宽速度低于预期。

大模型厂商 ARR 增速放缓的风险：若客户增长、续费率或客单价不及预期，模型厂商 ARR 增速可能放缓。

云服务商资本开支不及预期的风险：若自由现金流承压，或 AI 算力需求及投资回报低于预期，云服务商可能下调资本开支。

智能驾驶及机器人商业化落地不及预期的风险：若产品可靠性、成本下降或用户需求低于预期，相关产品商业化进程可能放缓。

免责声明

东吴证券股份有限公司经中国证券监督管理委员会批准，已具备证券投资咨询业务资格。

本研究报告仅供东吴证券股份有限公司（以下简称“本公司”）的客户使用。本公司不会因接收人收到本报告而视其为客户。在任何情况下，本报告中的信息或所表述的意见并不构成对任何人的投资建议，本公司及作者不对任何人因使用本报告中的内容所导致的任何后果负任何责任。任何形式的分享证券投资收益或者分担证券投资损失的书面或口头承诺均为无效。

在法律许可的情况下，东吴证券及其所属关联机构可能会持有报告中提到的公司所发行的证券并进行交易，还可能为这些公司提供投资银行服务或其他服务。

市场有风险，投资需谨慎。本报告是基于本公司分析师认为可靠且已公开的信息，本公司力求但不保证这些信息的准确性和完整性，也不保证文中观点或陈述不会发生任何变更，在不同时期，本公司可发出与本报告所载资料、意见及推测不一致的报告。

本报告的版权归本公司所有，未经书面许可，任何机构和个人不得以任何形式翻版、复制和发布。经授权刊载、转发本报告或者摘要的，应当注明出处为东吴证券研究所，并注明本报告发布人和发布日期，提示使用本报告的风险，且不得对本报告进行有悖原意的引用、删节和修改。未经授权或未按要求刊载、转发本报告的，应当承担相应的法律责任。本公司将保留向其追究法律责任的权利。

东吴证券投资评级标准

投资评级基于分析师对报告发布日后 6 至 12 个月内行业或公司回报潜力相对基准表现的预期（A 股市场基准为沪深 300 指数，香港市场基准为恒生指数，美国市场基准为标普 500 指数，新三板基准指数为三板成指（针对协议转让标的）或三板做市指数（针对做市转让标的），北交所基准指数为北证 50 指数），具体如下：

公司投资评级：

买入：预期未来 6 个月个股涨跌幅相对基准在 15%以上；

增持：预期未来 6 个月个股涨跌幅相对基准介于 5%与 15%之间；

中性：预期未来 6 个月个股涨跌幅相对基准介于-5%与 5%之间；

减持：预期未来 6 个月个股涨跌幅相对基准介于-15%与-5%之间；

卖出：预期未来 6 个月个股涨跌幅相对基准在-15%以下。

行业投资评级：

增持：预期未来 6 个月内，行业指数相对强于基准 5%以上；

中性：预期未来 6 个月内，行业指数相对基准-5%与 5%；

减持：预期未来 6 个月内，行业指数相对弱于基准 5%以上。

我们在此提醒您，不同证券研究机构采用不同的评级术语及评级标准。我们采用的是相对评级体系，表示投资的相对比重建议。投资者买入或者卖出证券的决定应当充分考虑自身特定状况，如具体投资目的、财务状况以及特定需求等，并完整理解和使用本报告内容，不应视本报告为做出投资决策的唯一因素。

东吴证券研究所
苏州工业园区星阳街 5 号
邮政编码：215021

传真：（0512）62938527

公司网址：<http://www.dwzq.com.cn>