

AI 智能体重构智能座舱体验



本报记者 张心怡

AI 正在从单点赋能走向系统重构,汽车也正在从“功能载体”向“超级智能体”加速演进。尤其是在座舱域, AI Agent正在重构人车交互方式,使座舱从被动响应向主动服务转变。

6月5日,蔚来创始人李斌在2026高通汽车技术与合作峰会现场表示,智能座舱的核心体验要全面Agent化, AI正在重构下一代座舱体验,把智能座舱带入认知座舱时代。而就在前一天,中国汽车工程学会副秘书长、国汽战略院执行院长郑亚莉在北京某场发布会上指出, AI Agent规模化上车成为新车型核心吸引力,智能座舱正从“工具属性”向“伙伴属性”转型。

中国汽车工程学会去年发布的《AI Car的初步畅想与探索实践》研究报告指出, AI对汽车的全面赋能将推动汽车成为具身智能体,即 AI Car。这一形态将通过智驾、座舱、底盘、动力多域智能的融合协同,实现汽车产品从单一功能智能化向系统级智能协同转变,具备更强感知、决策与服务能力,以“专业可靠的司机+聪明温暖的伙伴”双重角色,成为下一代智能空间的核心载体与智慧出行的生态节点。

郑亚莉表示,2026年上半年, AI Agent加速进入量产应用阶段,并逐步成为智能座舱能力升级的重要方向。与传统语音助手相比,新一代车载 AI Agent在多模态融合交互、长上下文理解、自主决策执行以及跨设备协同能力等方面持续增强,交互方式从“指令-响应”为主,逐步向融合用户需求与场景感知的连续化、个性化服务模式演进,推动人车交互从“语音控制”和“指令响应”向“理解需求、辅助决策和主动服务”方向拓展。

李斌则在今天的高通活动上指出,在十年前制定座舱发展方向时,蔚来就相信座舱应该成为一个有情感的伙伴。除了信息交互、娱乐交互,蔚来更关注人在车内的氛围、体验和感受。通过与高通等芯片企业的合作,让蔚来可以将车内的屏幕、氛围灯、音响等元素变成完整的、有温度的氛围体验。

在现场,李斌介绍了蔚来多项 AI座舱体验。包括行业首创NO-

MI识人记人专属体验,基于端侧大模型打造,可记住车主的20位家人,且数据本地保存不出车;行业首创的沉浸式5D座舱,实现了跨域打通,将最大处理延迟控制在0.09秒以下;与高通共同进行底层优化的多模态守卫模式,在汽车面临雨天水淹风险时,能够远程推送预警、智能抬升联动悬架并遥控辅助驶出。

然而, AI Agent规模化应用仍处于发展初期,在智能座舱的应用也面临诸多挑战。郑亚莉指出,在模型可靠性、端云协同机制、数据安全与隐私保护、人机交互边界等方面仍有待进一步验证和完善。

面向 AI座舱的下一步发展,李斌希望能与行业同仁在三个方向开展合作:一是共同推动全车算力高效协同。随着车上芯片越来越多,必须做好跨域融合来避免内存、算力的闲置,这对座舱与其他域的全车战略协同提出了非常高的要求;二是共同构建持续迭代、十年好用

的座舱跨代软硬件解耦标准;三是共建开放繁荣的Agent生态。

郑亚莉围绕当下车载存储芯片结构性短缺,提出“汽车与半导体跨领域协同将成为破局关键”。她表示,2026年上半年,车载存储芯片出现结构性紧张,这并非传统周期波动,而是 AI算力产业爆发与智能汽车渗透率快速提升形成的跨领域产能争夺,造成供需体系深层错配。

她认为,与2021年通用芯片急性短缺不同,本轮“缺芯”属于高规格专用存储芯片的长期“慢性病”,其根本原因在于传统的被动式响应、多层递进式供应链管理模式已无法适应智能汽车需求的快速迭代节奏。“长期来看,汽车产业的应对策略应从‘应急纾困’式的短期保障供应转向长期能力体系的系统重构,强调与芯片企业在需求定义、协同设计和供应链机制上的深度协同,而非简单追求自研芯片。”郑亚莉表示。

针对跨设备游戏体验,英特尔推出掌机处理器锐炫 G3

本报讯 5月28日,锐炫G系列处理器正式发布,专为掌上游戏机设计。该处理器采用第三代酷睿Ultra处理器(代号Panther Lake)架构,包含英特尔锐炫G3和英特尔锐炫G3Extreme处理器,将在Windows 11系统上运行。

6月3日,在英特尔组织锐炫G系列处理器媒体分享会上,英特尔技术专家表示,推出该处理器是看到了用户跨设备玩游戏的需求。

为了满足掌机用户的使用体验,锐炫G系列处理器在设计过程中看重性能、整合性、每瓦性能、软件兼容性等方面的属性,以保证用户在较小体积的掌机上,以较低的功耗实现沉浸式的游戏体验。

为实现上述设计目标,锐炫G3系列处理器对Panther Lake进行了结构调整,在原本的基础上减掉了

两个大的P核。英特尔技术专家解释,这是因为测试发现,在掌机和游戏场景下,这两个大核对游戏性能的边际提升已经微乎其微。而通过精简这两个大核,能够显著优化平台的整体功耗表现,带来更佳能效比。

此外,为满足产品性能,该系列芯片主要采用了以下三项核心技术。

第一,采用XeSS 3技术,通过XeSS超分、XeSS多帧生成和Xe低延迟三大核心图形技术加强游戏表现。XeSS超分技术,提供基于AI的图像超分功能,带来更强大的游戏性能;XeSS多帧生成技术通过增加多个插值帧,让游戏体验更加丝滑流畅;XeSS低延迟技术通过集成游戏引擎,能更快响应玩家需求,带来更加灵敏迅捷的游戏操控。

第二,采用IBC技术,通过智能控制SoC内部CPU与GPU的功耗分配,在提升整体游戏性能的同时,确保电池续航不受影响。它的核心是E核优先调度和低负载P核休眠,避免CPU瞬间功耗暴涨抢占GPU功耗资源。

第三,具备Endurance Gaming功能。玩家可以根据游戏类型和使用场景选择不同档位,将帧率设定为30帧、40帧或60帧。英特尔技术专家举例道,开启Endurance Gaming后,《堡垒之夜2》的整机功耗从22W降至4W流畅运行,续航时间从3小时延长到11小时;两款游戏从2小时提升到5小时。

英特尔锐炫G系列处理器是英特尔首次推出专为掌上游戏设计的处理器产品。而此前,AMD锐龙Z系列产品在该领域几乎垄

断该市场。关于为何选择进入该市场、如何应对激烈的市场竞争,英特尔技术专家分享道:从第三方数据来看,掌机市场规模约为200万~300万台,主机约为2000多万台,Switch约为2000多万台,PC市场规模更大。同时,游戏人群的数量是在增长的。通过调研,英特尔了解到,超过30%的PC玩家都拥有两个及以上的设备,跨多设备的游戏体验对PC游戏玩家来说也是非常重要的。同时,推出这款产品也能够提升英特尔对开发者生态的影响力。

记者了解到,采用该处理器的产品将在未来数月内陆续面市,首批机型包括宏碁Predator Atlas 8、微星Claw 8 EX AI+以及壹号本飞行家Apex Air等。

(姬晓婷 刘璐)

同济大学发布第一代滚动优化专用计算芯片

本报讯 记者杨鹏岳报道:近日,同济大学发布了第一代滚动优化专用计算芯片,旨在解决主流计算架构设计逻辑与实时滚动优化计算需求之间的结构性适配问题。

据同济大学电子与信息工程学院教授陈虹介绍,该芯片的核心任务并不是大规模数据的吞吐,而是支持强实时、强顺序、深预测、低功耗的时序因果推理,让运动体在每一步执行的同时,快速完成对未来发展状态的预测与评估。

“就像玩多米诺骨牌,它不用记住所有的牌面,而是专注于因果顺序。推倒第一张,毫秒之内就能预测第二张怎么倒、第三张往哪儿歪,并在每一步推演中实时修正力度。”陈虹表示,这种“走一步,看多步”的能力,正是未来自动驾驶汽车判断“前车减速后,我松开油门再过几秒才会追尾”,或是机器人在抓取水杯时预判“手臂抬高5厘米后,手腕需要转多少度才不会碰倒旁边的杯子”所需要的核心

智能。

这一成果并非一蹴而就。据悉,研发团队汇聚国家级人才10人、教授16人。早在2006年,团队便开始研究MPC的FPGA定制化实现,先后探索了软核SoPC、全硬件语言Verilog、异构SoC等多条技术路线,积累了完全自主知识产权的滚动优化计算架构。历经多年算法验证与车载测试,最终完成SoC芯片架构设计、IP核验证、RTL综合、后仿验证,成功流片并封装,实现车载控制器测试。

“这个技术自20世纪70年代以来,因其能适应动态不确定环境、处理多约束条件和多目标任务,一直被视为理想的控制策略。由于实时在线计算量大、计算频率高,传统芯片难以支撑其广泛应用。”上海交通大学教授席裕庚表示, MHU芯片的发布,正是为预测控制在高速动态场景中的落地提供了可行的硬件平台,也为更多行业引入智慧控制方法打开了通道。

中国电科已交付超500万颗智能终端用硅基氮化镓射频芯片

本报讯 记者从中国电科55所获悉,其自主研发的全球首款量产智能终端用硅基氮化镓射频芯片产品,近日已交付超500万颗。

这是全球率先实现硅基氮化镓射频芯片在智能终端规模化商用,将为空天地一体化信息网络的全域覆盖、高速互联提供硬核支撑。

空天地一体化信息网络是支撑未来6G通信、商业航天、低空经济及应急通信的核心底座,而低成本高性能功率放大器(PA)芯片则是这一网络的“信号心脏”,直接决定通信系统的传输速率、覆盖范围与

稳定性。当前,我国商业航天、低空经济、6G研发及信息通信产业正加速发展,对低成本高性能射频芯片的需求呈爆发式增长。

据悉,该系列硅基氮化镓射频芯片兼具高功率、高效率、超宽频、高可靠等突出性能,可精准匹配空天地一体化通信对射频功放芯片高效率、高线性度的严苛技术要求,有效破解高端射频芯片产业化难题,助力构建全域、全时、无缝的空天地通信网络,推动全球无缝通信、万物互联的产业愿景加速落地。

(越文)

英伟达宣布与SK海力士达成多年合作协议

本报讯 记者姬晓婷报道:6月7日,英伟达(NVIDIA)和SK海力士(SK hynix)宣布达成一项多年技术合作协议,旨在推进下一代内存技术的发展,助力全球人工智能工厂的建设,并加速半导体设计和制造。

据悉,该协议建立在双方多年来的深度合作研发基础之上,双方的合作成果已应用于全球一些最先进的人工智能计算平台。

这项多年期协议旨在保障先进内存的供应,以满足其不断延长的开发周期。随着人工智能工厂在全球范围内规模化发展,这项战略合作将确保内存供应能够与英伟达的基础设施路线图以及全球人工智能基础设施的持续建设保持同步。

通过此次合作,SK海力士将拓展至英伟达正在创建的新市场——涵盖人工智能基础设施、个人人工智能和物理人工智能。双方将共同开发用于英伟达Vera Rubin人工智能超级计算机、英伟达Vera

CPU、英伟达RTX Spark™ PC以及英伟达Jetson Thor™机器人计算平台的内存。

英伟达方面在声明中强调,协议明确支持先进存储的供应保障,旨在应对先进存储产品开发周期长、制造工艺复杂、资本投入密集等结构性挑战,确保存储供应能够跟上全球AI工厂持续扩张的节奏。

除存储产品联合研发外,双方合作还延伸至半导体设计与制造环节。SK海力士将采用英伟达CU-DA-X库及PhysicsNeMo框架,加速半导体仿真、技术计算机辅助设计(TCAD)工作流及内部工程代码的运算效率,并探索与电子设计自动化软件供应商的三方协作模式。在智能制造领域,SK海力士将结合英伟达Omniverse平台、OpenUSD流程及cuOpt决策优化引擎,构建晶圆厂数字孪生体系,为实现全自主晶圆厂运营奠定基础。双方还将探索将数字孪生与现有遗留软件及智能体AI工作流相连接,推动制造决策的自动化与智能化。

奋力谱写新型工业化发展新篇章

