

联想全栈AI

加速智能化转型每一步

联想智能云 (Lenovo xCloud) 产品介绍

模型、智能体、数据与知识开发一体机 & 工作站

联想

Lenovo

联想智能云全景图

联想沉淀40年IT智慧，以混合云为坚实底座，持续打造企业数据和知识库、模型工厂、智能体平台以及人工智能安全体系，助力客户在AI时代稳健前行。

联想智能云 (Lenovo xCloud)



AI相关开发工作站 / 一体机

模型开发一体机-训推

客户场景

- 基于DS模型微调，蒸馏等需求
- AI一体机支持DS满血版部署

模型开发（训推）一体机

模型开发（训推）工作站

模型开发平台



模型开发平台



联想问天 WA5480G3
8*Nv L20



联想问天 WA7880a G3
8*沐曦 C550



ThinkStation PX
4*Nv RTX5000/5880
4*MetaX N260

模型开发一体机-推理

客户场景

- 需要快速部署DS模型
- 利用模型进行推理，无需微调

模型开发（推理）一体机

模型开发（推理）工作站

模型开发平台



模型开发平台



联想问天 WA5480G3
4*Nv L20
4*MetaX N260



ThinkStation P5
2*Nv RTX5000
2*MetaX N260

ThinkStation P620
4*Nv RTX4000
2*MetaX N260

知识库开发一体机

客户场景

- 实现知识库无缝整合，统一管理
- 提供智能问答，提升知识管理与利用效率

知识库开发一体机

知识库开发工作站

企业知识库

KB Import KB Mgmt. KB Operation

Large Model Service

Integrated Model



企业知识库

KB Import KB Mgmt. KB Operation

Large Model Service

Integrated Model



联想问天 WA5480G3
2*Nv L20
2*Meta X N260



ThinkStation P5
2*Nv RTX4000
2*Meta X N260
ThinkStation P620
2*Nv RTX4000
2*Meta X N260

智能体开发一体机

客户场景

- 基于大模型快速开发智能体应用
- 支持拖拽操作构建智能体网络

智能体开发一体机

智能体开发工作站

智能体开发平台



智能体开发平台



联想问天 WA5480G3
4*Nv L20 / 8*Nv L20
4*Meta X N260 /
8*Meta X N260



ThinkStation PX
4*Nv RTX5000
4*Meta X N260

联想全栈AI
加速智能化转型每一步

模型开发工作站&一体机 (推理版)

Lenovo

2024 Lenovo Internal. All rights reserved.



联想 40 周年

模型开发一体机/工作站 方案一览表

分类	模型开发一体机/工作站	
版本	模型开发平台（推理版）	模型开发平台（训推版）
MA	Per node w/1年MA	Per node w/1年MA
用户规模	几十到几百用户的B3客户或B2客户部门级	几百用户的B2客户
算力芯片	工作站：ThinkStation P3/P5/P620， RTX 5000 / RTX4000 / 沐曦 N260 服务器：WA5480/WR5220， L20/沐曦C500/沐曦260	工作站：ThinkStation PX， RTX 5880 / RTX5000 / 沐曦 N260 服务器：WA7880a/WA5480， L20/沐曦C550/沐曦N260
规模	1节点	1节点、4节点
模型应用	私有化大模型API服务、本地知识问答	私有化模型训练、满血版模型本地API服务
模型	DeepSeek R1、Qwen 3、Gemma 3	DeepSeek R1、Qwen 3、Gemma 3
模型服务	本地高速、无限Tokens的大模型推理服务	本地高速、无限Tokens的大模型推理服务 专属领域、企业专属模型微调
增值功能	多模型切换；新模型部署	身份认知微调；模型评估
服务	模型应用咨询、培训；模型对接	数据采集、标注服务；模型微调

模型开发一体机/工作站（训推版）

产线：联想智能云（Lenovo xCloud）
产品经理：刘志恒
13521976746

产品定位：适用于有基于 DS 模型进行微调及蒸馏需求的用户，高配机型支持 DeepSeek 满血版的高效部署



模型开发一体机/工作站（推理版）

产线：联想智能云（Lenovo xCloud）
产品经理：刘志恒
13521976746

产品定位：面向需要快速本地部署DeepSeek、Qwen3等主流大模型，直接利用模型进行应用，助力构建本地的安全模型推理能力



模型开发平台：提供大模型全链路工具服务

针对大模型开发过程中的**数据处理难**、**开发成本高**、**模型落地难**等痛点，提供大模型全链路工具，全方位降低模型开发门槛，实现大模型全生命周期的闭环管理和模型迭代优化。



联想全栈AI
加速智能化转型每一步

智能体开发 工作站&一体机



智能体开发一体机/工作站 方案一览表

分类	智能体开发一体机/工作站		
版本	智能体开发平台（标准版）	智能体开发平台（专业版）	智能体开发平台（旗舰版）
MA	Per node w/1年MA	Per node w/1年MA	Per 套 w/1年MA
用户规模	几十用户的B3客户或B2客户部门级	几十用户的B3客户或B2客户部门级	几十用户的B3客户或B2客户部门级
算力芯片	工作站：ThinkStation PX, RTX 5880 / RTX5000 / 沐曦 N260 服务器：WA7880a/WA5480/WR5220, 沐曦 C550/沐曦C500/L20	工作站：ThinkStation PX, RTX 5880 / RTX5000 / 沐曦 N260 服务器：WA7880a/WA5480/WR5220, 沐曦 C550/沐曦C500/L20	工作站：ThinkStation PX, RTX 5880 / RTX5000 / 沐曦 N260 服务器：WA7880a/WA5480/WR5220, 沐曦 C550/沐曦C500/L20
规模	1节点	1节点、2节点	无限节点
模型应用	私有化大模型API服务、本地知识问答	私有化模型训练、满血版模型本地API服务	
模型	DeepSeek R1、Qwen 3	DeepSeek R1、Qwen 3	DeepSeek R1、Qwen 3
基础功能	智能应用编排能力、知识库挂载	多模态能力、意图分类引擎、提示词优化	长期记忆、智能体网络、数据库接入
增值功能	N/A	智能体发布API接口/SDK	模型微调
服务	智能体开发咨询、培训；智能体定制开发	智能体开发咨询、培训；智能体定制开发	智能体开发咨询、培训；智能体定制开发

智能体开发一体机/工作站

产品定位：基于大模型快速开发智能体应用，支持拖拽操作来构建智能体网络，大幅降低智能体应用开发的门槛与难度



一体化集成

- 支持丰富大模型，最高支持72B
- 零代码开发工具，分钟级创建

私有化部署

- 数据不出门，安全有保障
- 企业系统快速对接，成果易推广

简易运维

- 桌面级安装
- 7 * 24小时不间断运行

联想全栈AI
加速智能化转型每一步

知识库 开发工作站&一体机

知识库开发一体机/工作站 方案一览表

分类	知识库开发一体机/工作站
软件版本	企业知识库
MA	Per 套 w/1年MA
用户规模	几十到几百用户的B3客户或B2客户
算力设备	工作站：ThinkStation P5/P620/PX, RTX 5000Ada / RTX 4000Ada / 沐曦 N260 服务器：WA5480, L20/沐曦C500/沐曦260
规模	1节点、2节点、4节点
应用场景	本地智能知识管理、知识问答、API对接
模型	DeepSeek R1、Qwen 3、Qwen-VL、BAAI/bge-m3、BAAI/bge-reranker-large
模型服务	本地高速、无限Tokens的大模型推理服务，数据不出域
增值功能	知识运营；多模态数据解析、时效性&重复性双重检测
服务	知识库构建、知识优化

知识库开发一体机/工作站

产线：联想智能云（Lenovo xCloud）
产品经理：刘志恒
13521976746

产品定位：助力企业实现内部知识库的无缝整合，提供统一管理功能，并支持智能问答，提升企业知识管理与利用的效率



联想智能云 AI 技术方案服务清单



服务分类	服务名称	服务定义	典型适用场景	服务内容及步骤	标准人天	核心交付物
部署实施	Lenovo xCloud软件+DS 部署服务	在客户本地环境安装并启动 DeepSeek，使其具备基础功能	首次部署的中小企业客户	硬件环境合规性检查、私有化部署包安装、基础功能联调测试等	N+2	部署手册 部署验收报告
	系统集成	将AI能力嵌入或调用客户现有业务系统	通过集成AI能力赋予业务系统智能决策、自动化执行等能力	通过API对接等与业务流程适配，将AI能力嵌入业务系统（ERP等），实现 数据流-决策流-执行流 的端到端闭环	按需	需求分析 技术设计方案 测试报告与验收 用户手册与培训材料
	模型微调	基于客户私有数据对模型进行微调，提升特定任务表现	适配行业术语与业务流程，领域知识突破原生限制	轻量化微调、知识注入、效果验证等	按需	配置文件与参数说明 微调效果对比评估
运行维护	模型升级服务	版本迭代的平滑过渡保障	新模型版本发布后，应用或验证新特性	版本兼容性测试、增量升级实施、回滚方案准备等	按需	升级记录 新功能文档
	运维服务	帮助客户解决在使用模型过程中遇到的问题	模型训练，推理，应用过程中遇到问题	解决在使用DS等模型过程中遇到的配置、性能、兼容性相关问题	按需	问题分析手册 RCA报告
	推理性能优化	模型推理效能专项优化	提升高并发、低延迟需求的实时AI应用及资源受限环境中的模型推理效率。	推理加速优化、内存压缩、硬件加速配置等	按需	性能报告 基准测试数据

thanks.

