

AI 时代高品质全光算力专线 研究报告

(2025 年)

中国信息通信研究院技术与标准研究所

2025年9月

版权声明

本报告版权属于中国信息通信研究院，并受法律保护。转载、摘编或利用其它方式使用本报告文字或者观点的，应注明“来源：中国信息通信研究院”。违反上述声明者，本院将追究其相关法律责任。

前 言

高性能开源大模型的不断涌现，极大降低了 AI 应用创新的门槛和成本，成为驱动行业智算应用发展的核心引擎，支撑快速构建金融、政务、教育、医疗、工业等领域的场景化、专业化智能应用。随着多类行业智算应用的快速发展，不同智算应用场景对网络连接质量提出差异化要求。OTN 专线作为行业智算应用的关键承载底座，面向智算应用需求持续升级，提供大带宽、低时延、高可靠保障等能力，支撑算力灵活高效调度，为分布式智算协同等场景下的 AI 模型训练及推理提供高质量连接保证，赋能行业智算应用创新。

本报告聚焦行业智算应用景，围绕金融、政务、教育、医疗、公安、文娱、工业及大模型企业等典型应用场景，深入分析对于网络带宽、时延、可用率、安全性及弹性智能等方面的差异化需求。面向 AI 时代行业数字化应用发展趋势，提出面向智算应用的高品质算力专线智能感知、业务确定性体验、网络弹性按需、智能运维、光算协同五大特征，并详细分析光算融合的高品质算力专线关键技术。期望产业各方凝心聚力，协同推进 OTN 算力专线技术及应用创新，为行业智算应用的蓬勃发展提供坚实可靠的全光底座，支撑培育新质生产力，推动我国数字经济持续高质量发展。

目 录

一、 概述	1
二、 行业智算应用提出差异化专线服务需求	2
（一） 金融智算应用	2
（二） 政务智算应用	6
（三） 教育智算应用	9
（四） 医疗智算应用	12
（五） 公安智算应用	15
（六） 文娱智算应用	18
（七） 工业智算应用	21
（八） AI 大模型智算应用	24
三、 面向智算应用的高品质算力专线五大特征	27
四、 面向智算应用的高品质算力专线关键技术	29
（一） 智能感知关键技术	29
（二） 确定性体验关键技术	32
（三） 弹性调度关键技术	34
（四） 智能运维关键技术	35
（五） 光算协同关键技术	38
五、 总结与展望	39

图 目 录

图 1 银行智算应用示意图4

图 2 政务智算应用示意图7

图 3 教育智算应用示意图11

图 4 医疗智算应用示意图14

图 5 公安智算应用示意图17

图 6 文娱智算应用示意图19

图 7 工业智算应用示意图23

图 8 分布式推理训练业务示意图26

图 9 高品质全光算力专线目标架构28

图 10 光缆感知关键技术30

图 11 光电协同敏捷建路35

图 12 光网络智能运维架构36

图 13 光算协同技术框架38

表 目 录

表 1 金融智算应用网络需求5

表 2 政务智算应用网络需求9

表 3 教育智算应网络需求12

表 4 医疗智算应网络需求15

表 5 公安智算应用网络需求18

表 6 文娱智算应用网络需求21

表 7 工业智算应用网络需求24

表 8 AI 大模型智算应用网络需求27

表 9 智算业务差异化保障示例32

表 10 基于 SLA 的分级方式33

一、概述

开源大模型普及推动行业智算应用快速发展。2023 年开始，以 Llama、QWen、DeepSeek、ChatGLM 等为代表的高性能开源模型大量涌现，打破了此前大模型垄断的格局。开源大模型作为驱动智算应用加速创新的核心引擎，以其开放、透明、可定制的优势，大幅降低了行业应用门槛和成本，能够让开发者和企业快速参与到 AI 研发和应用。企业无需投入高昂训练成本，即可基于开源模型，结合自身行业知识和数据，进行高效的指令微调与领域适配，快速构建用于金融、政务、教育、医疗、工业等行业的场景化、专业化智算应用。

行业智算应用发展催生差异化网络连接需求。随着行业智算应用的爆发式增长，不同应用对网络提出带宽、时延等方面的差异化需求，需要网络感知业务类型并提供差异化连接能力，保障业务最佳体验。金融、政务等数据敏感行业的用户希望数据不出园区，通常采用数据本地存储、云端训练的模式，从云端到园区的网络需要具备 Gbps 以上大带宽，提供确定性连接，以支撑算力灵活高效调度；同时要具备低时延及高可用率，以保障在分布式智算协同场景下，避免因时延及异常中断而导致训练效率降低。

高品质全光算力专线为智算应用发展构筑坚实底座。为精准匹配智算应用需求，光网络需要实现从不感知业务类型到精准匹配业务需求的演进，根据业务流量、流向等特征，识别业务类型，分析带宽、时延、可用率等需求，实时按需提供差异化连接。同时，光网络需要提供确定性连接保障能力，避免因网络拥塞、故障、业务量增多等

因素导致连接质量劣化。此外，光网络还需要提供主动运维能力，先于用户投诉主动识别业务和网络隐患，并进行主动优化，实现网络自愈。未来，通过在网络设备、管控系统中构建 AI 内生能力，进一步加速光网络智能化演进，面向智算应用打造高品质全光坚实底座。

二、行业智算应用提出差异化专线服务需求

（一）金融智算应用

1. 金融智算应用场景

银行全面拥抱 AI 技术，从客户服务到风险管理实现智能化升级。

银行积极构建 AI 训练平台，集中管理算法模型和数据资产，支持业务部门快速开发 AI 应用，银行智算应用场景如图 1 所示。在客户服务方面，AI 网点助手和数字人大堂经理广泛应用于业务咨询与智能导览，大幅提升服务效率；在业务运营中，AI 理财质检对理财销售的合规性进行智能检测，降低漏检与误检，提升审核准确性；在风险控制领域，基于 AI 的反欺诈系统可实时识别并拦截可疑交易。

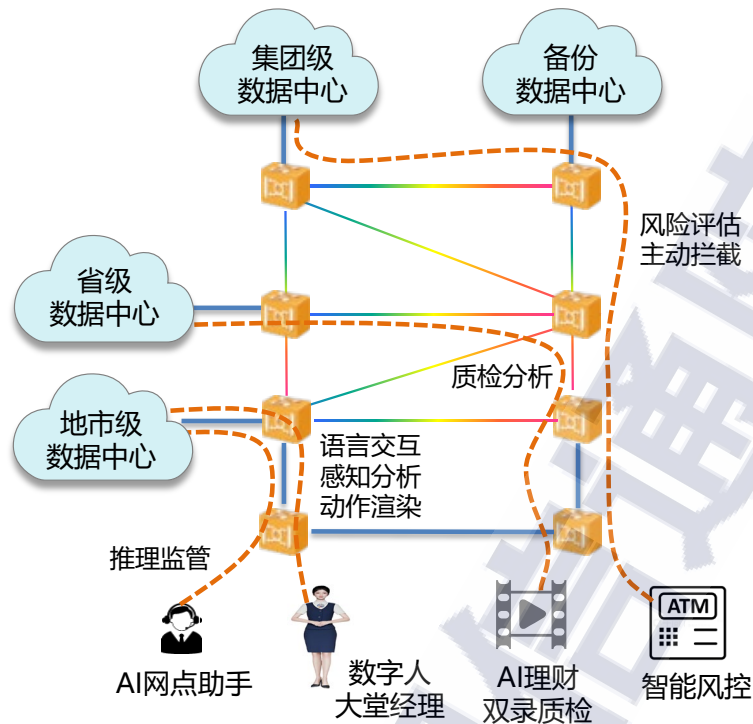
AI 网点助手显著提升网点的服务效率和营销精准度。AI 网点助手通过自然语言处理技术实现 7×24 小时智能应答，能够即时响应 90% 以上常见业务咨询，如账户管理、产品信息、业务流程等，释放柜员人力。通过对金融政策和产品合规知识库学习，自动同步监管新规动态，基于客户画像进行产品精准推荐，为智能柜台提供自主服务指导。

数字人大堂经理实现智能导览与业务咨询。数字人大堂经理通过 3D 虚拟形象和自然语言交互技术，主动识别客户需求，并将其精准分流至柜台、自助设备或线上渠道，提升网点服务效率。数字人可实

时解答存款、贷款、理财等常见业务问题，提供标准化政策解读，减少客户等待时间，借助 AR/VR 技术，能够演示 ATM 使用、手机银行注册等业务操作流程，降低人工服务压力。此外，系统可监测客户情绪波动，自动触发安抚话术或适时转接人工服务，优化客户服务体验。

AI 理财双录质检提升理财服务合规检查准确性。AI 理财双录质检系统对理财销售过程中的录音录像内容进行实时分析，采用自动语音识别（ASR）技术将音频转换为文本，并结合自然语言处理（NLP）模型识别敏感词汇，自动检测违规话术、遗漏风险提示等合规问题。同时，系统通过视频分析模块捕捉销售人员的微表情和肢体语言，评估其服务规范性，实现更高效、标准化的合规风控管理。

银行智能风控提高了风控反诈的及时性和准确性。AI 智能风控整合交易记录、客户行为和外部征信等多源数据，通过机器学习与行为分析等技术，实现对风险的实时评估。系统能够实现实时监测、精准识别欺诈行为，并可在毫秒级时间内完成对可疑交易的风险评估与自动拦截。同时，通过图计算引擎与网络分析技术，能够识别复杂的关联交易，发现传统规则引擎难以检测的欺诈模式，大幅提升风险防控的效率和准确性。



来源：中国信息通信研究院

图 1 银行智算应用示意图

2. 金融智算应用网络需求

AI 网点助手以文本交互为主，对时延要求高。AI 网点助手进行业务智能应答以文本形式为主，需要进行实时的推理解析，端到端推理延迟需要小于 50ms，网络时延要求小于 5ms。网点助手一次文本交互约 50~100kbps，单网点需同时支持 50 台以上终端并发，为避免高峰时段卡顿，建议预留 5Mbps 带宽，确保请求无阻塞。

数字人大堂经理要求高带宽和低时延连接。数字人大堂经理在业务办理过程中涉及到大量的实时语音交互、人体动作感知分析、面部表情与动作渲染等处理过程，其高拟真性、高实时性、高并发性对网络带宽、时延都提出较高要求。GPU 渲染集群通常部署在云端，要保证 4K/8K 高拟真画质和动态细节渲染，每路需要 20~50Mbps 带宽，若银行网点同时并发 2~4 路，总带宽需求可达 200Mbps，网管数字人

根据网点需求灵活增加，需要专线具备带宽调整能力。实时动作渲染的理想端到端时延要求不大于 15ms，图像采集和云端渲染解码约 10ms，网络单向时延要求控制在 2.5ms 以内。

AI 理财双录质检要求高带宽和高安全传输。单个理财录音录像视频文件大小约 500MB~1GB，1 分钟上传完毕所需带宽约 150Mbps。上传系统后需尽快反馈结果以支撑理财销售，通常要求端到端交互时延小于 30ms，其中网络单向时延需控制在 5ms 以内。为保护客户隐私，数据传输需采用专用硬管道进行隔离，随着入侵手段的不断升级，银行开始探索应用量子加密等安全技术。

银行智能风控反诈系统对时延要求严苛。AI 风控反诈系统需实时防御金融风险。以 ATM 异常交易为例，当系统判定存在风险，将通过金融专线立即下发拦截指令至 ATM 终端，该流程端到端交互时延需小于 30ms，其中跨区域端到端网络时延要求小于 5ms。系统的交易数据以指令为主，单条指令的数据量为 10~50KB，网点多台终端设备的专线带宽需要 5Mbps，金融专线要求可用率不低于 99.99%。高风险交易数据需传输至公安反诈平台，由于涉及公民隐私等敏感信息，传输专线需采用硬隔离的独立物理通道，确保防窃听防篡改。

金融智算应用网络需求汇总如表 1 所示。

表 1 金融智算应用网络需求

金融智算应用	带宽	单向时延	可用率	安全性	弹性智能
AI 网点助手	5Mbps	< 5ms	≥99.99%	硬管道隔离	弹性带宽
数字大堂经理	200Mbps	< 2.5ms	≥99.99%		
AI 理财	150Mbps	< 5ms	≥99.99%	硬管道隔离	

质检				加密传输	
AI 风控 反诈	5Mbps	< 5ms	≥99.99%		

来源：调研数据

（二）政务智算应用

1. 政务智算应用场景

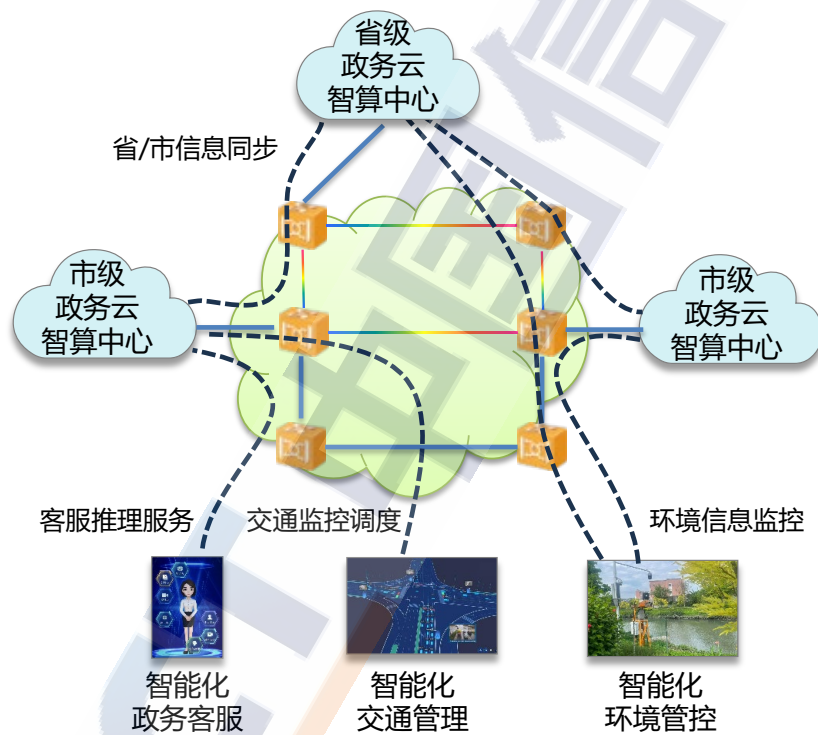
政务智算应用从探索到广泛落地正在快速转变。政务数字化转型初期主要聚焦于办公自动化、政府网站搭建和基础数据库建设等工作，随后建设重点逐步转向打造“一站式”政务服务平台，如图 2 所示。通过大模型等新技术应用，实现智能决策和人机协同治理，推动政务服务智能化升级。在智能化政务客服、智能化交通管理和智能化环境监控等领域，开展智算应用场景的探索，提升政务管理的效率和水平。

智能化政务客服提升市民信息获取效率。依托大模型技术，智能化政务客服整合了海量政策法规与办事指引数据，能够快速精准地解答市民在企业服务、民生保障各类领域的咨询，例如市民咨询创业补贴，AI 政务客服可清晰说明申请流程、补贴范围等，大大缩短信息获取时间。在政策推送方面，系统可在业务办理结束时，依据用户的业务类型、年龄、身份等多维度信息，精准推送后续可办事项及相关政策提醒，例如市民办理完租赁备案一体化业务后，若符合廉租房申请条件，系统将及时推送政策介绍、申请时间及所需材料等指引信息。

智能化交通管理提升道路通行能力。智能化交通管理利用 AI 技术，实时监测与分析城市交通流量，推理预测交通拥堵状况，智能调整交通信号灯时长，实现交通流的优化分配。利用 AI 控制的智能交通信号灯，可以使道路通行能力提高 15%-20%。AI 还能对交通违法

行为进行智能识别与预警，提高交通管理的效率与道路安全性。

智能化环境管控助力实现城市环境问题的精准及时治理。借助多模态大模型技术，综合卫星遥感、无人机和高清摄像头监控等多源数据，智能化环境管控实现了对城市空间的数据收集与感知分析，可以有效提升环保卫生、垃圾监管、用地冲突等事件的处理效率。例如，通过对卫星图像的智能解析，可及时发现非法占地或环境污染等问题，助力城市环境问题的及时、精准管控。



来源：中国信息通信研究院

图 2 政务智算应用示意图

2. 政务智算应用网络需求

智能化政务客服应对热点突发事件，需要提供灵活网络服务能力。智能化政务客服以文本交互为主，对网络带宽及时延要求均不高。以日常政策咨询文本交互为例，带宽需求小于 5Mbps，网络时延控制在 500ms 内即可满足基础服务体验。网络作为政务服务的关键基础设施，

需确保全年稳定运行，可用率应达到 99.99%以上的高标准。为防范数据泄露风险，需采用硬管道隔离技术实现物理层面的网络隔离。针对政策发布、集中咨询等可能出现突发或高并发的场景，系统应具备灵活弹性的带宽扩容能力，以应对短期流量高峰，保障服务连续性。

智能化交通管理要求大带宽、低时延、安全可靠的网络。交通路口部署的高清摄像头、毫米波雷达等设备需实时回传 4K 视频及车流数据，单个摄像头带宽不超过 20Mbps，单个路口 4-8 个摄像头的带宽需求约 200Mbps。为将城市各道路监测点采集的数据汇聚并上传至云端，骨干网络带宽需达到 100Gbps。AI 交通控制要求网络时延小于 20ms，以便实时采集的数据能够快速传输，及时完成车流量识别与信号灯调整，避免道路拥堵加剧。车辆轨迹、车牌等数据需经安全加密传输，防范个人隐私信息泄露。为避免网络中断导致信号灯失控、影响交通疏导，网络可用率需达到 99.99%以上。

智能化环境管控要求大带宽、高安全、高可靠网络。大气监测站、水质传感器等设备持续上传 PM2.5 和污染物浓度等数据，单个采样点带宽为 200Kbps 左右，部分场景视频监控带宽可达 20Mbps，为将海量监控点数据传输至云平台，进行存储和分析，骨干网络带宽需达 10Gbps 以上。环境样本数据非实时交互数据，对于时延的要求不高，秒级即可满足后端处理需求。环境数据涉及城市生态敏感信息，传输网络需采用加密传输和严格权限管控，防止数据被篡改或泄露。同时，网络应具备中断后快速恢复与数据补传能力，避免监测数据缺失影响研判，可用率需达到 99.99%。

政务智算应用网络需求汇总如表 2 所示。

表 2 政务智算应用网络需求

政务智算应用	带宽	单向时延	可用率	安全性	弹性智能
智能化政务服务	<5Mbps	<500ms	≥ 99.99%	硬管道隔离	弹性带宽
智能化交通管理	<200Mbps	<20ms	≥ 99.99%	硬管道隔离 加密传输	算网一体服务
智能化环境管控	200Kbps~ 20Mbps	<500ms	≥ 99.99%	硬管道隔离 加密传输	——

来源：调研数据

（三）教育智算应用

1. 教育智算应用场景

智算技术正在改变传统的教学和学习模式。智算应用在教育行业取得了显著的进展，从教学、评估到管理与科研，其应用场景呈现出深度垂直化与全域渗透化的发展趋势。个性化学习平台、智能辅导系统、自动化评估工具和智能监考等技术被广泛采用，改变了传统的教学和学习模式，使得教育资源更加丰富和多样化，同时也提高了教学的效率和学习效果。此外，智算也被应用于教学研究和数据分析，帮助教育工作者更好的了解学生的学习习惯和需求。

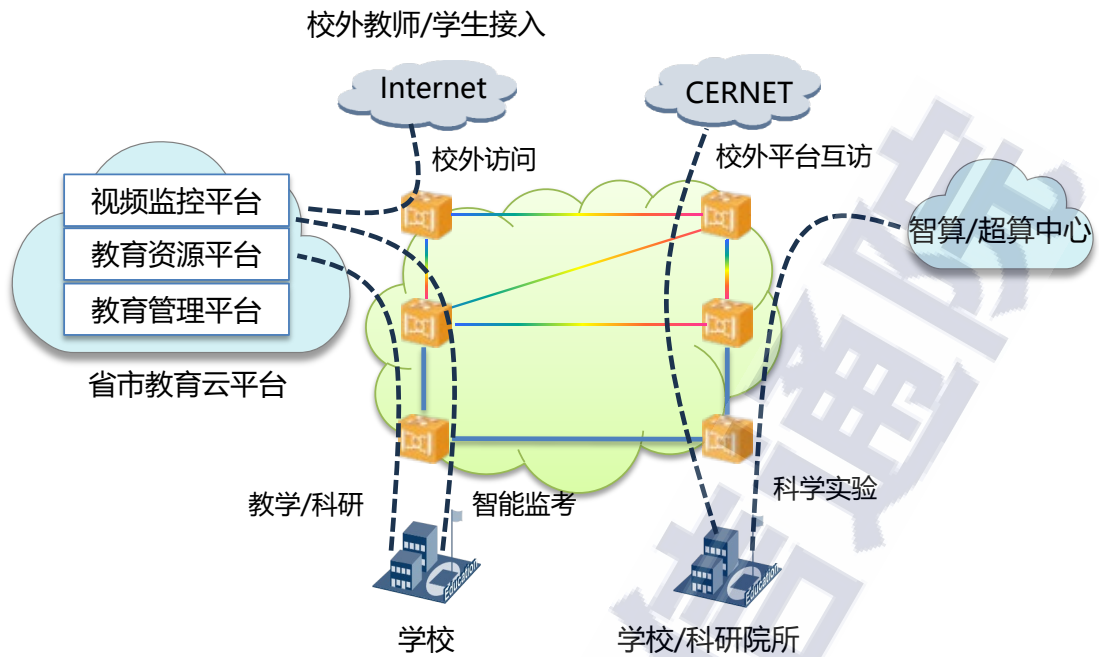
AI 智慧课堂提升教学效率和质量。传统教学课堂正在被智慧课堂取代，智能终端设备通过校园网、教育专网或互联网接入教学云平台的授课系统，实现高清全景课堂交互、名师课堂录制与优质教学资源的共享。通过在教学云平台系统嵌入 AI 大模型，实现智能化课堂管理，实时监控学生学习状态，智能识别课堂专注度，并通过语音分析评估师生互动质量，从而提升教学效率和教学质量。

AI 协助教学研究、精准评估教学质量。在学情分析阶段，AI 可

提供学生的历史学习数据，帮助教师预判学生的知识薄弱点，让备课更具针对性。在备课环节，教师输入课题和教学目标后，AI 可自动生成教案初稿、PPT 大纲、思维导图及多媒体素材推荐。在教学过程中，AI 通过整合学生的全流程学习数据，构建动态的个体学情画像，精准标识个人优势学科、兴趣领域和薄弱环节，便于开展个性化教研。在教学评估环节，AI 生成班级、年级乃至全校的学情报告，帮助教学管理者从宏观层面把握教学质量，精准识别潜在问题与发展趋势。

AI 智能监考提升考试作弊处置的时效性。智能监考系统将考场视频实时上传至云平台，利用人脸识别和行为分析技术进行判定后，再将结果反馈给考试现场，锁定作弊行为。面对日益复杂的高科技作弊手段，系统借助人工智能和大数据等先进技术，能够高效处理考场内的海量监控信息，提高对隐蔽性强、手法新颖的作弊行为的识别和防范能力，降低由于人为因素引起的误判或疏漏风险，使监考工作变得更加高效和准确。

各类教育管理平台逐步向云侧集中迁移。全国大部分省市已部署教育云平台，集中管理教育资源、档管理和教育政策等信息，并基于云平台部署 AI 大模型提升教学效率和教学质量，如图 3 所示。学校通过教育专网连接云平台开展智能监考、智慧教学和科研等业务，也可通过 CERNET 访问其他省市教育云共享资源。老师和学生可通过互联网接入教育云平台，进行在线备课和学习。高校及科研机构可通过专线接入智算或超算中心，满足高算力需求。



来源：中国信息通信研究院

图 3 教育智算应用示意图

2. 教育智算应用网络需求

智慧课堂业务促进学校出口带宽增长。智慧课堂的主要流量，来自学生在多媒体教室，通过云平台收看高清教学直播、点播以及名课录播。每路视频约占用 4Mbps 带宽，一所中大型学校按照 40~50 个班级计算，总需求约 200Mbps。根据调研显示，75%的学校希望平台出口带宽达到 100~500Mbps。直播课堂的实时反馈对时延要求较高，网络单向时延需控制在 10~25ms，AR/VR 教学等新兴应用对时延的要求更加严格，网络单向时延需小于 10ms。

教研科研云化、智能化促进学校和平台间的互访流量增长。教育平台嵌入 AI 大模型可以生成丰富多样的教育内容，如虚拟实验、互动课程和多媒体教学资源等。随着越来越多的学校接入教育专网，不同学校或平台之间的互访流量呈现增长趋势。根据调研显示，68%的

学校要求互访流量带宽达到 100~500Mbps。同时，随着教学设备智能化普及率的不断提高，带宽需求也在同步增长。普通中小学的出口带宽达到 500Mbps~1Gbps，高校的出口带宽以 10Gbps 为主。

智能监考 AI 高清摄像头带来大带宽、低时延需求。智能监考摄像头正向超高清方向发展，单个 1080P 分辨率摄像头的带宽约为 10Mbps，单个 4K AI 摄像头的带宽约为 20Mbps。每个教室通常安装 2~4 个摄像头，一个标准考点按照 50 个教室考场估算，需要的总带宽为 4Gbps。智能监考及时反馈对时延要求较高，网络单向时延需控制在 5ms 以内。

教育智算应用网络需求汇总如表 3 所示。

表 3 教育智算应用网络需求

教育智算应用	带宽	单向时延	可用率	安全性	弹性智能
智慧课堂	100~500Mbps	<25ms	≥99.99%	传输加密	带宽、时延 可视、链路 状态可视
教学科研智能化	1~10Gbps	<50ms	≥99.9%		
智能监考	~4Gbps	<5ms	≥99.99%	硬管道隔离 传输加密	

来源：调研数据

（四） 医疗智算应用

1. 医疗智算应用场景

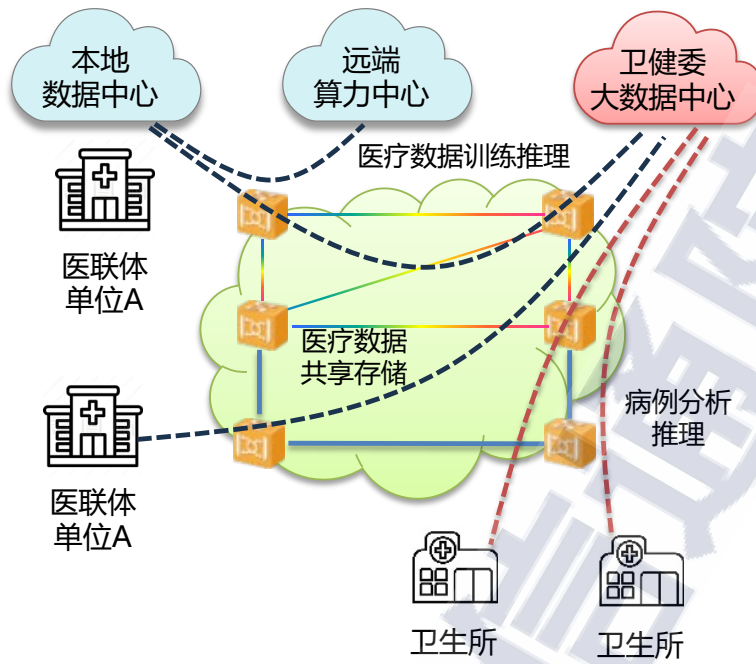
医疗行业正积极拥抱智算应用，缓解医疗资源供需矛盾。人工智能、机器学习和大数据分析等技术已广泛应用于诊断辅助、个性化治疗、药物研发及患者监护等医疗领域。AI 技术能够帮助医生分析海量医疗数据，提高诊断的准确性和效率，为患者提供更加精准的治疗

方案。此外，AI 辅助诊疗等技术有助于缓解医疗资源供需矛盾，提高医疗体系运行效率，减轻医务人员的工作负担。医疗智算应用场景如图 4 所示。

AI 远程辅助阅片提升诊疗效率与准确率。医院通过高速网络将 CT 扫描、核磁影像（MRI）等数据上传至云平台，AI 辅助阅片快速分析并标记可疑病灶，同时将结果推送给远程专家，专家结合分析结论，即时调阅高清影像，与基层医生开展远程会诊。偏远地区患者因此也能享受到三甲医院的优质资源，基层诊疗效率与准确率得到提升。

AI 辅助诊疗场景改变就医流程。患者就诊时，系统收集整合其过往病历、症状描述及各项检查数据，AI 算法迅速分析患者信息，快速生成初步诊断建议与个性化检查方案。例如咳嗽和发热多日的患者，AI 能结合症状与流行疾病趋势，推测可能病因，依据患者的年龄和病史，精准匹配相关病症信息，为医生呈现详尽的诊断参考。

医联体资源共享助力优质医疗服务触达更多角落。医联体在接诊患者后，会将其 CT、检验报告等检查数据自动上传至卫健委大数据中心，实现核心医疗数据的集中存储与统一管理。AI 算力中心基于统一云平台，对医疗数据进行训练与推理，形成病例分析模型、诊疗决策算法等资源，基层医生接诊时，可实时调取云端 AI 工具，将患者的 CT 影像和检验数据上传，AI 平台返回病灶识别结果和诊断建议，同时医疗专家通过远程会诊系统，为基层患者提供诊疗指导，突破优质医疗资源地域限制，提升区域整体诊疗效率与水平。



来源：中国信息通信研究院

图 4 医疗智算应用示意图

2. 医疗智算应用网络需求

AI 辅助阅片要求大带宽传输管道。AI 辅助阅片对网络性能有严格要求，为避免传输延迟影响诊断时效，单例 10GB 的病理切片文件需在 10 秒内完成上传，系统接入带宽约需 10Gbps。系统端到端交互时延需控制在 100ms 以内，其中网络单向时延需要小于 10ms。网络需具备 99.99%以上的高可用率，并采用加密传输与权限隔离机制，防止医疗数据泄露。

AI 辅助诊疗的多模态数据交互提出大带宽及高可用要求。AI 辅助诊疗需支持电子病历、检验报告、影像片段等多模态数据的实时交互，为确保 AI 分析与医生操作的同步，端到端交互时延需控制在 50ms 以内，其中网络单向时延需要小于 5ms。在进行单例 3D 影像与基因数据融合分析时，网络带宽需稳定保持在 500Mbps，当多科室并发调

用 AI 模型时，突发带宽需求可以达到 1Gbps。网络需具备 99.99%的高可用率，并通过加密传输机制保障数据隐私安全。

医联体 AI 资源共享需要大带宽、高安全、弹性网络。医联体 AI 资源共享依赖高性能网络的支撑，AI 训推一体场景下，网络带宽需达到 1Gbps 并支持动态调整，以实现带宽随业务需求弹性变化。电子病历、影像等医疗数据需跨机构频繁调阅，网络需达到 99.99%的高可用率，端到端交互时延需小于 100ms，网络单向时延需要小于 10ms。同时需要传输加密与数据隔离，以保障数据安全。

医疗智算应用网络需求总结如表 4 所示。

表 4 医疗智算应网络需求

医疗智算应用	带宽	单向时延	可用率	安全性	弹性智能
AI 辅助阅片	10Gbps	<10ms	≥99.9%	硬管道隔离 传输加密	—
AI 辅助诊疗	500Mbps~1Gbps	<5ms	≥99.9%		
医联体 AI 资源共享	500Mbps~1Gbps	<10ms	≥99.99%		带宽调整

来源：调研数据

（五）公安智算应用

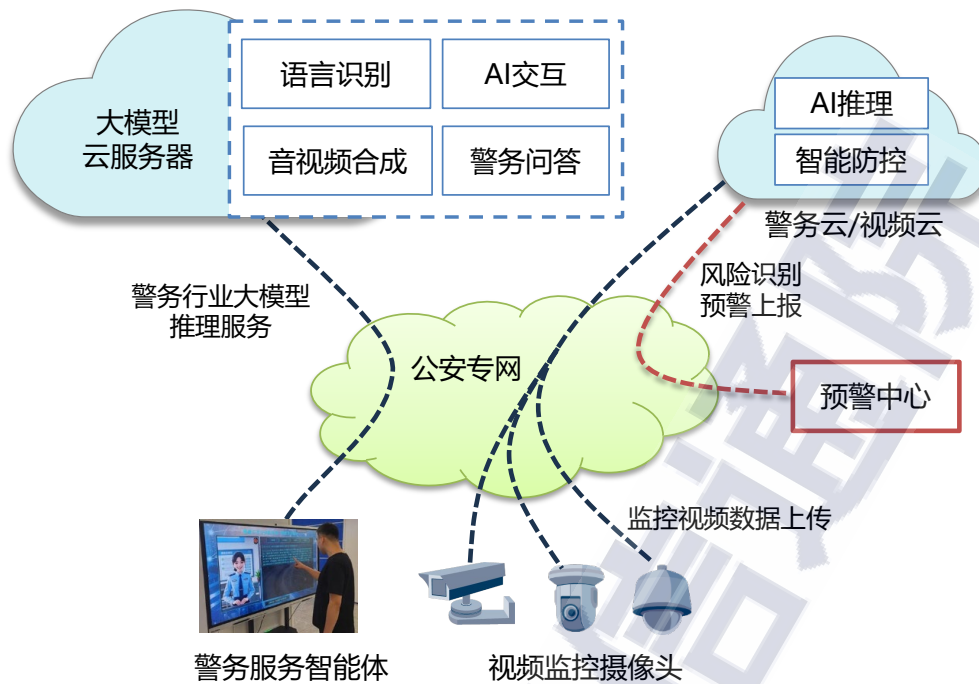
1. 公安智算应用场景

AI 已深度应用于公安警务的各个环节。AI 正以前所未有的深度和广度融入警务工作，推动警务模式从“事后追溯”向“防患于未然”转变，加快速务工作向智能化、精准化和高效化方向转型。在治安视频监控场景中，智能监控能够实时识别异常行为并及时发出预警，例如对海关反走私、人流车流异常聚集等区域异常入侵的识别

与预警。在 AI 辅助接处警、警务咨询答复和纠纷调解等场景中，智能客服能够全天候解答群众咨询，节省警务工作的人力，提升警务工作效率，增强执法效能和社会治理水平。公安智算应用场景如图 5 所示。

AI 视频监控可以有效提升风险识别与预警能力。AI 视频监控对重点监控区域的异常行为进行视频识别和分析，对嫌疑人员和车辆等高风险目标进行跨区域跟踪，锁定其移动轨迹，并判断其行进方向和下一步动作。AI 视频监控可快速提取相关人员的面部、衣着、发型和体态等特征，开展识别和关联分析，通过多视频源的多源数据联动，结合多种信息构建目标人群的行为画像，能够有效提升对潜在风险的识别能力和效率。

AI 警务服务智能体提供多种警务工作服务。警务行业大模型能够通过语义分析理解用户提出的问题，并对警务业务流程进行智能答复和引导。业务咨询和引导机器人不仅可以通过语音提示智能指引群众提交所需资料，还能借助光学字符识别（OCR）、机器人流程自动化（RPA）和私有图像传输协议（VGTP）等前沿技术，实现资料的预先审核、自动化精准录入，以及业务办结后“电子证件”的签发。



来源：中国信息通信研究院

图 5 公安智算应用示意图

2. 公安智算应用网络需求

AI 赋能公安视频监控需要广覆盖、大带宽、低时延的网络。警务视频分析大模型对视频成像质量需求持续提升，为避免传输延迟和卡顿，单个 4K/8K 摄像头需要 20~50Mbps 带宽，按照路口四个方向的配置，每路承载专线的带宽需求约 200Mbps。在实时追踪、异常预警等应用场景中，网络延迟必须控制在 5ms 以内，以确保视频流快速回传和 AI 分析结果快速反馈。由于视频图像数据包含大量敏感信息，因此网络必须部署物理隔离、加密传输以及安全访问控制策略等安全机制，从而防止数据泄露和篡改。

AI 警务智能体应用要求公安传输网络带宽提升。警务机器人通常配备 1 至 2 路 1080P 的高清摄像头、人脸识别模块，并具备语音交互功能，需要不小于 20Mbps 的稳定带宽，以满足实时视频回传以及

定位信息等传感器数据的传输需求。实时场景的延迟需要控制在 25ms 以内，非实时场景的延迟可放宽至 50ms。在警务应用场景中，由于需要与云服务器进行多轮次的敏感信息交互，对信息安全的要求较高，因此需提供国密和量子加密等高级别的安全传输保障。

公安智算应用需求汇总如表 5 所示。

表 5 公安智算应用网络需求

公安智算应用	带宽	单向时延	可用率	安全性	弹性智能
AI 视频监控	200Mbps	<5ms	≥99.99%	硬管道隔离传输加密	端到端可视、可维
AI 警务智能体	20Mbps	<50ms	≥99.99%	传输加密	多系统联合定位定界

来源：调研数据

（六） 文娱智算应用

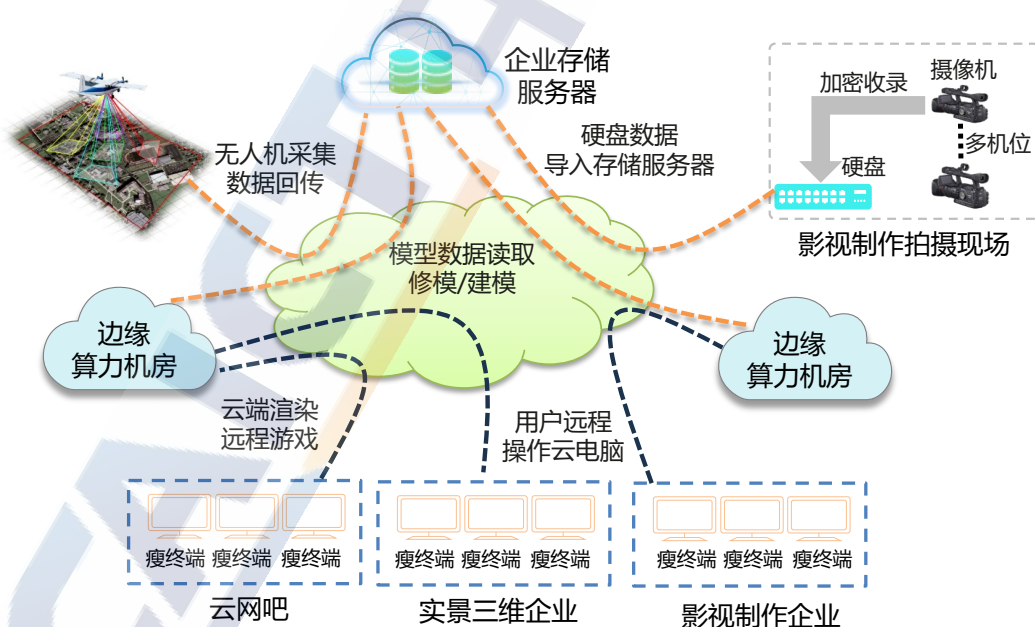
1. 文娱智算应用场景

云网吧加速推动上云数字化转型。文化和旅游部在《互联网上网服务行业上云行动工作方案》中，要求加快推动互联网上网服务营业场所的数字化转型，创新发展“存储上云”“算力上云”等上网服务行业云服务新模式，助力行业实现转型升级。网吧上云后，本地主机系统的 CPU 和 GPU 将迁移至云端进行集中部署，本地仅保留鼠标、键盘、显示器等外设。瘦终端会将鼠标、键盘的操作采样数据通过网络传输至云端，云端的算力操作系统会调度这些数据至应用软件进行交互逻辑处理，并将其交给 GPU 进行图形处理和渲染，渲染后形成的每一帧视频画面，会经过编解码通过网络传输至瘦终端进行图像显示，如图 6 所示。

实景三维涉及大量数据采集及处理流程。实景三维作为数字中国

的基础底座，已经在自然资源、数字景区、数字文物、智慧城市、应急管理、公安、住建、环保和矿山等多个场景得到广泛应用。实景三维通过无人机、相机、扫描仪等采集设备，对现有场景进行多角度环视拍摄，以三维视觉重建技术为核心，结合数字摄影测量技术和人工智能技术，将采集场景快速还原为三维形态，真实再现现实世界。

影视制作上云需满足专业级影视渲染和高刷新率要求。视效和剪辑作为影视制作的重点环节，在操作过程中需要实时查看效果，影视制作企业需要具备高分辨率、高刷新率和高精度操作性能的云电脑产品，同时配套快速读写文件存储池、归档存储池以及支持并行渲染的高性能算力池，以呈现逼真细腻的显示色彩，无损还原图像的色彩和亮度，满足专业级影视剪辑和影视设计渲染需求，以及在高速动态视频场景下的流畅体验。



来源：中国信息通信研究院

图 6 文娱智算应用示意图

2. 文娱智算应用网络需求

云网吧要求低时延传输网络保障流畅运行。网吧上云需要稳定的带宽，当前网吧主流应用的 2K 分辨率、144 帧刷新率的画面，需要 120Mbps 的带宽，如果单个网吧同时上线 80 台电脑，需要约 10Gbps 带宽。按照每秒 144 帧的刷新率计算，每帧的显示间隔时间需要稳定在 6.9ms，为不影响使用体验，云端处理、本地处理和网络传输的时延需在一帧时间内完成，目前主流显卡的云端处理和本地瘦终端的处理时间平均约为 5.4ms，网络单向传输时延需要控制在 1ms 以内。传输网络出现丢包会导致视频画面数据异常或操作动作丢失，具体表现为视频画面出现花屏、跳帧、卡顿等问题，操作上出现操作跳动、点击丢失等情况，网络可用率需要达到 99.999%。

实景三维测绘数据量大，上云需要大带宽网络保障。实景三维测绘企业通过本地连接云端电脑，实时展示图像效果并进行优化处理，考虑到画面的高分辨率、刷新率和流畅度，其网络单向时延要求与云网吧相同，在 2K 显示器及 144 帧刷新率下，网络单向时延需小于 1ms，带宽需求与云网吧相同，每台云电脑需要 120Mbps，按 8 至 10 人同时使用计算，需约 1Gbps 带宽。实景三维测绘大量数据传输需要弹性带宽调整能力，按照一台无人机一天拍摄 300GB 数据计算，若有 10 台同时作业，一天共产生 3TB 数据，按照 2 至 3 小时的传输时长计算，需要约 3Gbps 的带宽。测绘数据安全性要求高，数据传输通道需与其他通道实现物理隔离，并支持国密加密传输，考虑到传输带宽和数据处理量较大，还需支持硬件加密。

影视制作上云多人协作对时延、带宽及数据安全提出更高要求。

影视制作需要使用 4K 分辨率显示器，刷新率主要为 144Hz，单台云电脑的带宽需求为 250Mbps，若同时运行 10 至 20 台，日常接入带宽要求约为 5Gbps，在传递影视素材时，带宽需要弹性调整至 10Gbps，此外网络需要具备弹性带宽可调能力，以实现云电脑即开即用，并能根据需求扩容算力和存储。由于影视制作云电脑高分辨率、高刷新率对高色彩表现及高精度有严格要求，需要引入更强算力的 GPU 和增强型本地瘦终端，其网络时延需求与云网吧类似，网络单向时延要求小于 1ms。影视数据在公开上映前具有极高的安全保密要求，因此在影视编辑制作、算力处理数据和网络传输数据的过程中，都需要保护影视数据源，网络需提供硬件加密和传输通道硬隔离等手段。

文娱智算应用网络需求汇报如表 6 所示。

表 6 文娱智算应用网络需求

文娱智算应用	带宽	单向时延	可用率	安全性	弹性智能
云网吧	10Gbps	≤1ms	≥ 99.999%	—	支持敏捷建拆
实景三维云渲染	1Gbps	≤1ms	≥ 99.99%	加密传输	带宽弹性
影视制作	5Gbps	≤1ms	≥ 99.99%		

来源：调研数据

（七）工业智算应用

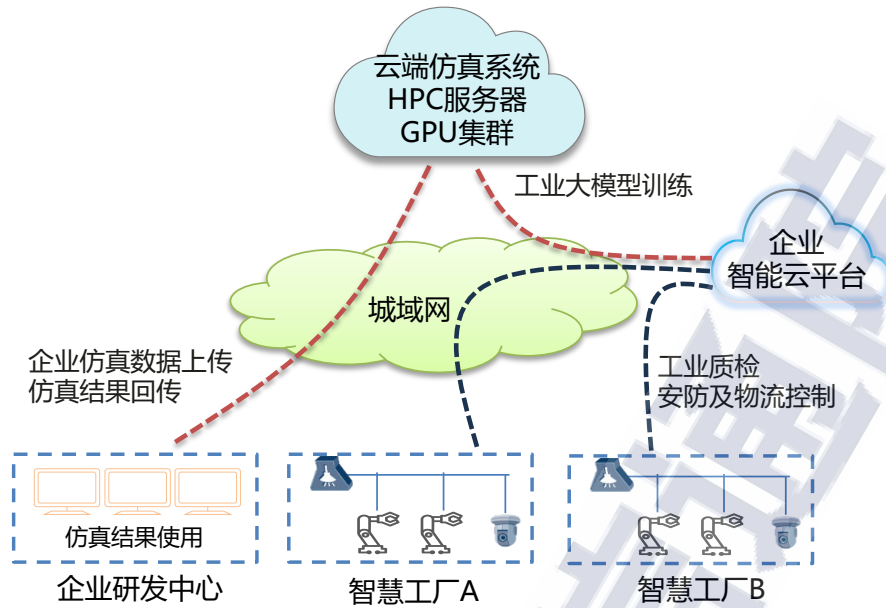
1. 工业智算应用场景

工业仿真与云技术、人工智能融合大幅提升设计精度和效率。工业仿真涉及结构、流体、动力学、散热和电磁等多个领域，每个领域都拥有多种仿真工具软件，集成和整合难度较大，数据流转复杂，同时仿真时间受到硬件资源的限制，导致仿真周期较长。随着云化技术和人工智能技术的发展，工业仿真应用逐步引入云化技术，集成自动

化仿真工具链，利用云的弹性扩展优势，提供更强的算力，提升用户使用的便捷性，如图 7 所示。例如，Ansys 的仿真平台可以集成 100 种以上的计算机辅助工程（CAE）和计算机辅助设计（CAD）工具，用户可以随时访问和使用。AI 技术的引入和更强大的算力支持，能够通过预测整体性能，节省仿真求解时间，大幅缩短仿真周期，提升设计开发效率。

智慧工厂等助力企业提质增效。随着制造企业数字化转型的不断深入，“智慧工厂”和“黑灯工厂”开始发展，其本质是借助不断发展的 ICT 技术，实现工厂的数字化、网络化、无人化和智能化。通过将末端的设备或机器实现数字化，利用可靠的工业互联网实现设备的连接和数据的传输，最终数据在智能云平台上完成收集和处理，由统一的数字化云平台对研发、物流、生产、质检、安防、供应、经销等各种应用系统进行有效和高效的管理，真正让机器代替人工作业，通过 AI 数据训练实现数据分析处理，最终达到制造企业提质增效的目标。

AI 质检广泛应用于工业生产领域。以机器视觉为代表的 AI 技术，正在被广泛应用于半导体电子制造、食品制造、车辆装备制造等多个领域，涵盖缺陷瑕疵检测和生产环境安全等多项功能。机器视觉设备通过采集被检产品的图像信息，利用本地边缘计算设备或后续服务器实现 AI 技术，更经济、更高效、更准确地完成质量检测。



来源：中国信息通信研究院

图 7 工业智算应用示意图

2. 工业智算应用网络需求

工业仿真数据传输需要满足用户对网络带宽、成本、数据安全性等多方面的要求。工业仿真云化后，企业需要将仿真数据上传至云端进行计算和处理，仿真数据量大，需在短时间内完成数据传输，以避免影响设计开发进程。以汽车行业的新车型研发为例，汽车碰撞仿真、结构件强度仿真和整车空气动力仿真，研发中心与云端的平均每日数据传输量达到 TB 级别。在智能驾驶研发过程中，需要大量数据进行训练，百台路测车的数据量为百 TB/天，实现百 TB 以内数据的小时级传输，需要 10~100Gbps 的带宽支持。企业对专线资费较为敏感，提供带宽弹性调整服务模式，企业在接入侧提供 10Gbps 级别的带宽入口，日常使用采用 100Mbps 的带宽，在需要传输仿真数据时，可动态调整至 10~100Gbps 的带宽，并按使用时长进行收费，以满足大带宽数据传输需求，同时降低使用成本。工业仿真云化对数据安全性要

求较高，需要网络具备加密能力，防止传输过程中发生数据泄露，网络可通过提供基于物理隔离的硬管道，确保传输数据与网络中其他业务数据实现硬隔离，并在业务入口侧增加基于国密的硬件加密能力，既不影响传输效率，又能有效防止数据被窃取。

AI 智慧工厂检要求带宽、时延及数据传输的确定性。工业质检 4K 高清摄像头，单个工位需要 20Mbps 左右带宽，单个工厂 50 个工位计算，需要 1Gbps 左右带宽，工厂孪生平台需要每 50ms 推送实时数据，需预留 1Gbps 左右带宽。工业可编程控制（PLC）闭环时延要求严苛，要求到边缘算力中心时延小于 1ms。工业远程控制业务停机损失巨大，要求传输网络提供高可用率不小于 99.999%。智慧工厂需要与互联网隔离，避免网络入侵造成的中断产线、窃取工业配方等损失，要求提供端到端加密的能力，在传输过程中实现硬管道隔离，并提供传输加密和硬件加密手段。

工业智算应用网络需求汇总如表 7 所示。

表 7 工业智算应用网络需求

工业智算应用	带宽	单向时延	可用率	安全性	弹性智能
设计/仿真业务	500Mbps~1Gbps	<2ms	≥99.99%	硬管道隔离，加密传输	SLA 指标可视 支持弹性带宽调整
AI 智慧工厂	<3Gbps	<1ms	≥99.999%		——

来源：调研数据

（八） AI 大模型智算应用

1.AI 大模型智算应用场景

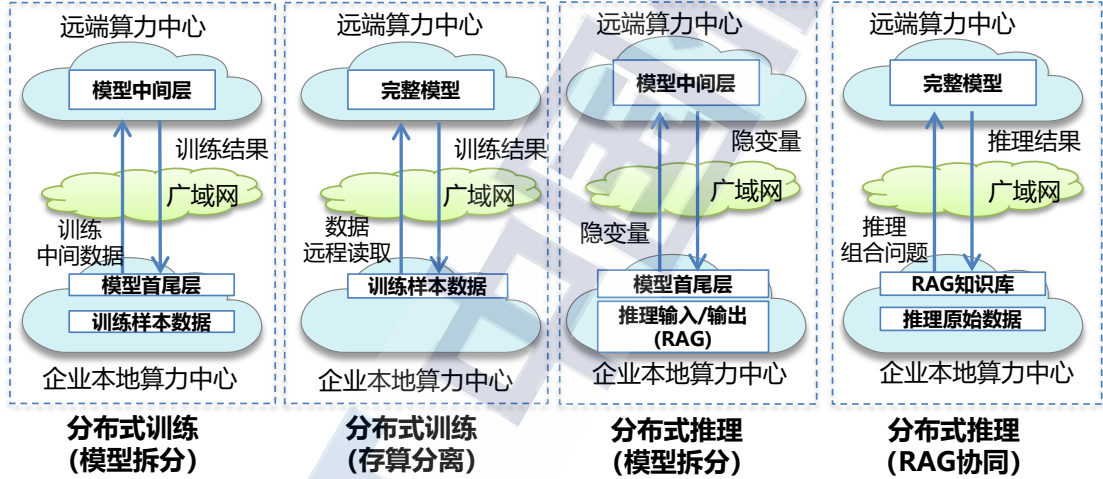
企业用户希望安全使用云端算力，保障核心数据安全。DeepSeek

开源大模型的应用热潮正在席卷全行业，吸引了国内外众多企业争相接入，其应用价值主要体现在专业领域知识问答的精进上，用户重点关注如何高效整合行业专属知识库来提升问答精准度。目前多个行业的用户都表达了基于私有知识库构建专属 DeepSeek 应用的强烈意愿，但由于普遍面临算力建设成本过高和数据保密需求严格的双重制约，迫切需要一套安全高效的实施方案，在严格保证数据不出园区的前提下，完成模型的安全部署和实际应用。

分布式训练通过模型拆分与存算分离两种方式，确保训练过程的企业数据安全。分布式训练采用模型拆分模式保障训练数据不离本地，该模式将模型首尾层部署在企业本地，中间层部署在云端，样本数据和模型训练过程中的损失值在本地分别转换为激活值或梯度值后上传至云端，云端处理完成后再回传至本地。在典型应用场景中，分布式训练的传输距离不超过 300 公里，为保障模型计算时间与网络通信时延更好匹配，最大传输距离建议不超过 2000 公里。**存算分离通过数据远程读取不落盘的方式保障训练数据安全**，通过将企业数据存储在本地并远程挂载到智算中心，训练时智算中心远程读取挂载数据且不在本地落盘，从而强化数据隐私保护，实现企业安全使用云上算力。存算分离在数据读写存储过程中需要大量信令交互，为保障远程读写存储的 IO 性能，防止 RTT 过长导致远程读写性能急剧下降，建议部署距离不超过 300 公里。

分布式推理通过模型拆分与本地知识库协同，实现数据不出园的安全推理。分布式推理的模型拆分模式将模型首尾层部署在企业本地，

中间层部署在云端，推理时提示词在本地转换为隐变量后发送至云端，云端完成计算后返回隐变量并由本地转换为最终推理结果，确保原始输入与输出均保留在企业本地。该模式下推理返回的 token 间时延间隔相对较短，为保障算效并防止 RTT 过长导致推理性能严重下降，建议部署距离不超过 300 公里。分布式 RAG 推理将模型全部部署在云端，企业知识库与向量数据库则完全部署在本地，推理时先在企业本地检索相关内容，再组合问题发送至云端生成答案，仅少量查询信息上传，原始数据全程保留在本地，实现安全高效的协同推理。



来源：中国信息通信研究院

图 8 分布式推理训练业务示意图

2. AI 大模型智算应用网络需求

分布式训练要求带宽按需提供以及高可靠的无损网络。周期性的训练需要网络提供弹性带宽调整能力。训练带宽远大于推理带宽，典型分布式训练场景需要 100Gbps 互连带宽，推理带宽范围在百 Mbps 至 10Gbps 之间。企业在使用模型时经常需要结合自身行业数据微调模型，存在从推理专线到训练专线的转换需求，因此网络需要具备百

Mbps 至百 Gbps 的弹性能力。训练过程要求高可靠的无损网络以避免影响计算效率，其中参数面数据传输需要网络具备无损传输的高可靠能力。训练过程中企业和数据中心之间相互传递数据时若发生传输错误，即使通过重传机制，等待重传仍会影响计算效率，严重丢包可能导致当次计算中断并重新启动，造成数小时延误并大幅增加计算成本，因此网络需要具备高可靠能力以保证智算训练不中断。

分布式推理要求稳定的时延以及加密能力。为保障企业本地数据中心与云端算力数据中心之间的数据快速可靠传输与交互，网络需要具备无损能力，并在传输过程中对数据进行加密以确保数据安全。同时网络需提供稳定的时延能力，本地 DC 与远端 DC 之间采用全光一跳直达的架构，全程无拥塞，从而保证计算效率的稳定以及用户体验。

AI 大模型智算应用网络需求汇总如表 8 所示。

表 8 AI 大模型智算应用网络需求

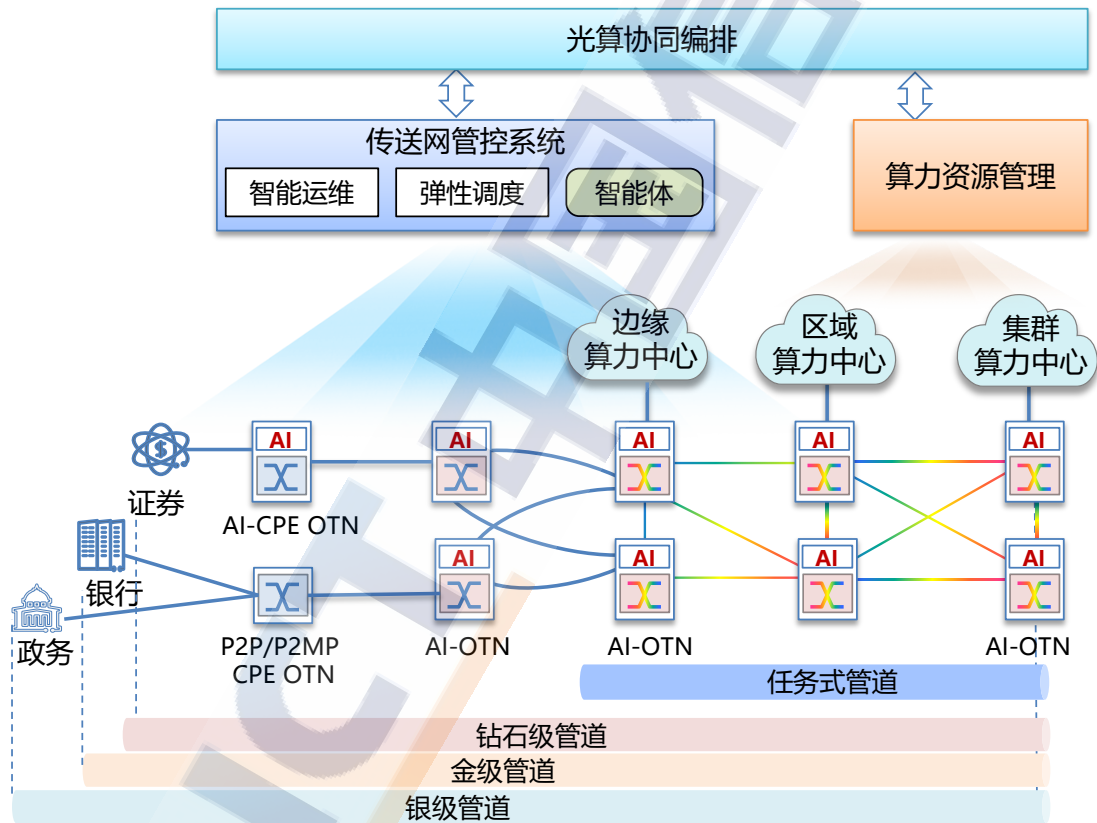
AI 大模型智算应用		带宽	单向时延	可用率	安全性	弹性智能
分布式训练	模型拆分	100Gbps	≤10ms	≥99.99%	加密传输	无损传输
	存算分离	~10Gbps	≤2ms	≥99.99%		带宽弹性 SLA 指标可视
分布式推理	模型拆分	~10Gbps	≤2ms	≥99.99%	加密传输	无损传输 SLA 指标可视
	RAG 推理	百 Mbps	≤10ms	≥99.99%		SLA 指标可视

来源：调研数据

三、面向智算应用的高品质算力专线五大特征

随着 AI 赋能行业数字化转型进程加快，政务、金融、医疗、教育、文旅、工业等行业客户提出差异化入算的网络需求，入算专线 SLA

定义呈现差异化趋势。面向 AI 时代智算行业应用，高品质算力专线五大关键特征为具备**多维智能感知、业务确定性体验、网络弹性按需、智能运维、光算协同**等能力，为企业数字化转型打造智能入口，提供安全、经济、差异化服务。高品质算力专线目标架构如图 9 所示，在设备层，通过引入 AI-OTN 设备，提供多维感知和内生算力等能力；在管控层，通过引入大模型、智能体、数字孪生等技术增强网络智能运维能力，保障专线业务确定性体验、弹性按需调度和光算协同编排。



来源：中国信息通信研究院

图 9 高品质全光算力专线目标架构

(1) **智能感知**：为保障不同应用的体验，需构筑光缆、网络、业务三层智能感知能力。通过三层感知能力，实现对业务特征识别，匹配好光缆资源和光层网络资源，实现差异化保障。

（2）确定性体验：根据不同应用提供实时按需的差异化连接，具备波长/ODU/fgOTN/OSU 大中小颗粒的转发能力，基于行业差异化入算需求，业务 SLA 分级维度更加丰富，从过去以带宽为主，升级为时延分级、使用时长分级、传输质量分级、可用率分级、安全分级等不同 SLA 维度，对应提供钻石级、金级、银级、任务式不同等级管道，保障不同业务确定性体验。

（3）弹性按需：管道的使用从静态分配到灵活拆建，从以年为周期占用到按小时级、天级分时复用，光网络具备“波长级敏捷建链能力”以及“弹性带宽调整能力”。

（4）智能运维：基于 AI 大模型、智能体、数字孪生等技术，形成网络智能评估规划、意图驱动业务发放和按需调速、主动品质保障和智能故障诊断等全生命周期智能运维能力，有效提升网络运维效率。

（5）光算协同：通过物理层、协议层、管控层进行光网络和算力资源协同，实现计算和光网络协同感知，算网统一编排调度，基于业务需求最优算路等能力，显著提升客户体验。

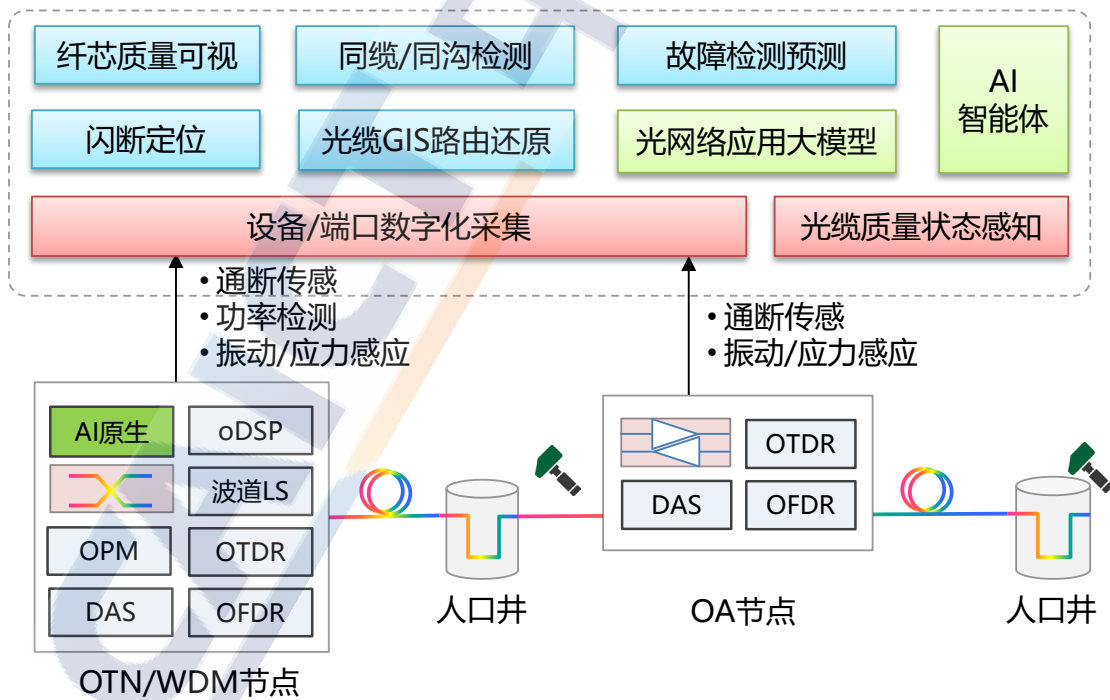
四、面向智算应用的高品质算力专线关键技术

（一）智能感知关键技术

AI 时代业务对时延、带宽等诉求存在较大差异，为保障不同应用体验，需构筑光缆层、网络层、业务层三层智能感知能力，实现对专线业务特征识别，匹配好光缆资源和光层网络资源，实现差异化保障。

光缆感知实现光纤质量、同路由风险和外部环境威胁等感知能力。

当前光纤闪断抖动无法溯源，依赖人工排查，定位难度高，同时光缆的同沟风险和外破预警风险依赖人工介入，影响业务投放时间(TTM)以及业务品质。为提升光纤感知能力，需对光纤感知功能器件进行技术创新。通过升级已有光时域反射（OTDR）能力，对光纤瞬断能力进行监控；通过引入分布式光声波传感（DAS）技术，实现对光纤外部振动或声波的连续监测，基于同沟光纤振动同源特点准确实现在线同沟检测，并具备外部环境监控能力，识别光缆外破风险；通过构建时间、频率、空间模型对接收的相位和强度数据进行特征提取,异常识别或风险分类，实现同沟和风险在线预警。未来可进一步引入分布式温度传感（DTS）能力，对光纤的周边环境温度进行检测，实现光缆风险进行全方位检测，进行多维融合感知实现对城市安全风险的认识。



来源：中国信息通信研究院

图 10 光缆感知关键技术

网络感知需升级设备感知能力和模型分析能力，精准识别和预测网络特征和状态。网络感知能力不仅需要感知设备的端口、波长、ODU 等带宽资源及其状态信息，还需要感知时延、丢包率、可用率等 SLA 信息，并能够对网络资源及业务性能进行感知预测。针对客户侧的故障感知，可通过智能光模块实现 ms 级光功率变化感知，结合脱落、火烧、夹断、挖断等光纤故障数据集，构建 AI 算法模型，分析失效特征，快速识别客户侧故障模式。针对线路侧光纤光功率抖动等问题，可基于快速 OTDR 探测技术，实现 10ms 以内的光功率抖动事件检测，并进行抖动精准定位，有效解决抖动检测难题。针对网络资源的状态感知预测，引入 AI 算法模型，根据历史数据预测业务的热点趋势，识别站点资源瓶颈，保障资源实时可用。针对网络业务性能监控预测，通过 AI 学习拟合不同波道组合、不同链路组合下的光学效应变化，评估工作路径和备用路径的性能，保障保护倒换状态下的光层业务品质，同时预测待开通业务性能，确保业务平稳开通。

业务感知精准识别和预测业务特征，实现差异化的业务保障。业务感知不仅识别专线管道级的时延和丢包等指标，也可识别业务特征，按照应用需求度量用户体验，基于业务特征进行差异化保障，通过带宽按需调整以及流量预测，实现用户需求和网络带宽精准匹配。针对专线业务特征感知，需引入 AI 构建应用级的流量特征模型，区分视频业务、AI 智算应用、智算训练业务等不同业务特征，根据不同应用类型设置不同的转发优先级，建立差异化的连接。针对不同智算应用服务体验，构建服务体验分析模型，基于业务级时延、丢包、抖动等

多维因子进行体验分析，度量出不同用户的体验感受。针对应用级的带宽感知，构建应用级的流量特征模型，可根据连接承载的应用类型进行带宽预测，精准调节专线带宽，实现带宽按需调整。

（二） 确定性体验关键技术

随着政务、金融、教育、医疗、文旅等行业上云入算后，需要提供低时延高品质入算网络，保障体验与本地部署相同。同时，入算业务包含智能问答、文件交互、视频会议、云渲染等不同应用，对于传输的带宽、时延、可用率等要求不同，需要提供多样化网络服务保障不同应用的确定性体验。

表 9 智算业务差异化保障示例

行业	应用	管道类型	带宽	时延	可用率	使用时长
金融	高频交易	钻石级管道	带宽保障 固定带宽	us 级	99.999%	长期
	数字人大堂经理	金级管道	带宽保障 弹性带宽	ms 级	99.99%	长期
政务	智能政务客服	银级管道	带宽保障 弹性带宽	10ms 级	99.99%	长期
公安	智能视频监控	银级管道	带宽保障 固定带宽	10ms 级	99.99%	长期
智算业务	模型训练	任务式管道	超大带宽 带宽保障	ms 级	99.99%	任务式

来源：中国信息通信研究院

提供硬管道隔离保障基础带宽。通过 fgOTN、OSU、ODUk 及波长 2Mbps~100Gbps 不同带宽颗粒度硬隔离管道技术，为行业用户分配专属固定通道，实现物理隔离传输，保障上云入算数据稳定带宽，打造专属安全的数据传输通道。

基于多维 SLA 分级的差异化业务保障。SLA 维度更加丰富，如带宽、时长、传输质量、可用率、安全、时延、开通时间、故障处理、

自助服务等，每个维度提供多种分级保障方式。基于丰富 SLA 维度进行业务建链，提供多样化管道，如钻石级管道超低时延超高可靠、金级管道大带宽低时延、银级管道带宽保障、任务式管道灵活拆建，满足不同行业用户确定性体验。

表 10 基于 SLA 的分级方式

分级类型	分级方式
使用时长分级	年/月固定长租、周/天临时短租
带宽分级	固定带宽、带宽保障、带宽调速
时延分级	不承诺、可承诺、可定制
传输质量分级	误码率 1e-9、误码率 1e-12
可用率分级	99.999%、99.99%、99.9%
安全分级	硬管道、传输加密、硬件加密

来源：中国信息通信研究院

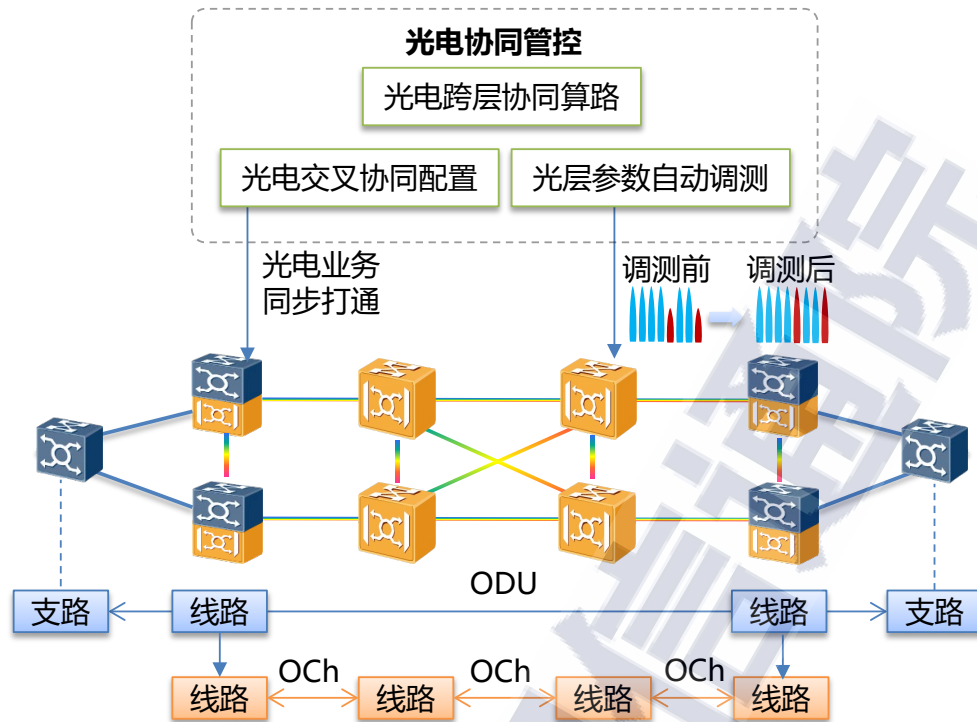
基于 SLA 的可视、分级保障和调优技术。管控系统提供业务 SLA 可视化能力，包含带宽、时长、传输质量、可用率、安全、时延、流量等关键 SLA 指标，并支持光网络拓扑时延地图、运力地图。**基于 SLA 分级保障业务差异化服务要求。**使用时长分级支持固定建立专线连接和动态拆建专线连接方式，提供长租和临时短租服务。**带宽分级**支持业务应用特征识别，提供基于应用需求调速，可通过监测专线流量提供带宽调速。**时延分级**通过引入空芯光纤等传输介质，光网络管控系统业务算路策略等，可以支持时延不承诺、时延可承诺等按需定制。**传输质量分级**支持通过不同级别 FEC 误码纠错方式，提供 1e-9 误码率和 1e-12 误码率服务。**可用率分级**通过增加光缆路由，结合光电 ASON 重路由能力，提升抗多次断纤能力，通过无损倒换技术使得业务倒换过程中无丢包，进而提供 99.999%、99.99%、99.9%等多种等级可用率。**安全分级**通过硬管道隔离、传输管道加密、硬件加密

等方式提供不同等级的安全服务。基于 SLA 的业务调优提升客户应用体验，支持识别业务应用特征重定向到更优时延路由管道，保障企业关键应用更低时延在线调优。支持基于业务可用率重定向到更高可用率管道，满足企业业务升级需要更高可用率专线承载诉求。

（三）弹性调度关键技术

大数据按需搬运以及分布式训推等业务场景，需要 OTN 网络具备弹性按需带宽调度能力，业务管道需从静态分配改为可灵活拆建，从以年为周期占用到按小时级、天级分时复用，要求光网络具备“波长级敏捷建链能力”以及“小颗粒带宽弹性调整能力”。

波长级敏捷建链快速打通波长传输通道。实现分钟级波长业务自动发放、自动调测、自动释放。一是实现**光电跨层协同算路**，基于业务需求以及网络的实时性能、时延、路由分离等策略，实现光电层路由协同计算，提供满足业务需求的多条 OCh 可选路径；二是实现**光电交叉同步创建**，业务智能调度配置参数自动生成，包含客户到 OCh、OCh 到光纤的多层路由映射、波长频率与频宽配置、中继端口配置等；三是实现**光路参数自动调测**，自动调测开通 OCh 通道，自动调优至最佳性能状态。



来源：中国信息通信研究院

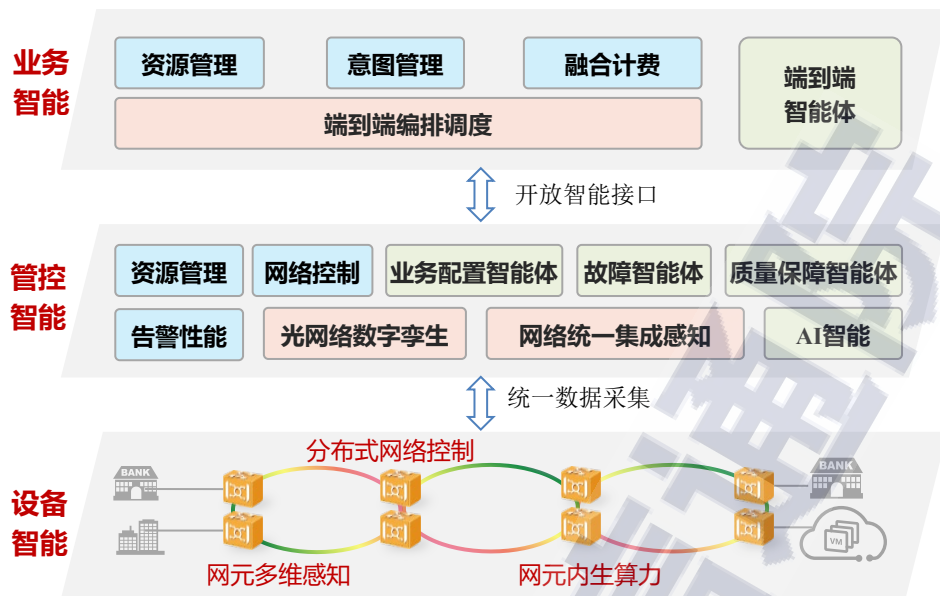
图 11 光电协同敏捷建路

OSU/fgOTN 技术实现灵活带宽接入及弹性带宽调整。

OSU/fgOTN 技术支持多种业务带宽颗粒灵活接入，解决传统 ODUk 映射效率不高问题。OSU/fgOTN 将连接数提升到百万级别，充分满足海量业务差异化带宽需求。OSU/fgOTN 弹性管道可满足不同行业用户的带宽调整需求，同时满足企业日常运营过程中的小颗粒带宽需求，以及数据搬运等应用突发大带宽需求，使网络与业务需求高效匹配，降低用户网络使用成本。

（四）智能运维关键技术

基于数字孪生、AI 大模型、智能体等技术，光网络实现业务层、管控层及设备层的智能。通过网络智能评估、意图驱动业务发放、主动品质保障和智能故障诊断等全生命周期智能运维能力，有效提升网络运维效率。



来源：中国信息通信研究院

图 12 光网络智能运维架构

业务层基于意图实现端到端编排调度。基于自然语言意图模型实现业务需求自动理解，结合大模型技术，由客户文本输入自然语言，描述业务需求，系统自动理解用户意图，并支持多轮交互对话，进一步确认、细化用户意图，实现用户业务需求的自动理解。**基于意图驱动业务发放**，业务层端到端智能体负责理解客户意图，业务层进一步编排调用管控层的北向 API，驱动管控层业务配置智能体完成多域、多厂商业务发放。

管控层实现智能评估、业务配置、品质保障和智能故障诊断等智能特性。基于实时网络资源孪生实现网络智能评估。管控层可以根据业务的带宽、时延、路由等需求，基于实时网络资源数据，通过智能路由算法和智能网络时延分析算法等进行业务资源评估，并自动推荐备选路由方案。**高效适配网络资源实现快速业务配置。**基于业务层传递的业务诉求，管控层可以高效的适配网络资源，按需完成业务建拆配置或业务调速配置，支撑算网任务式调度。以波长业务发放为例，

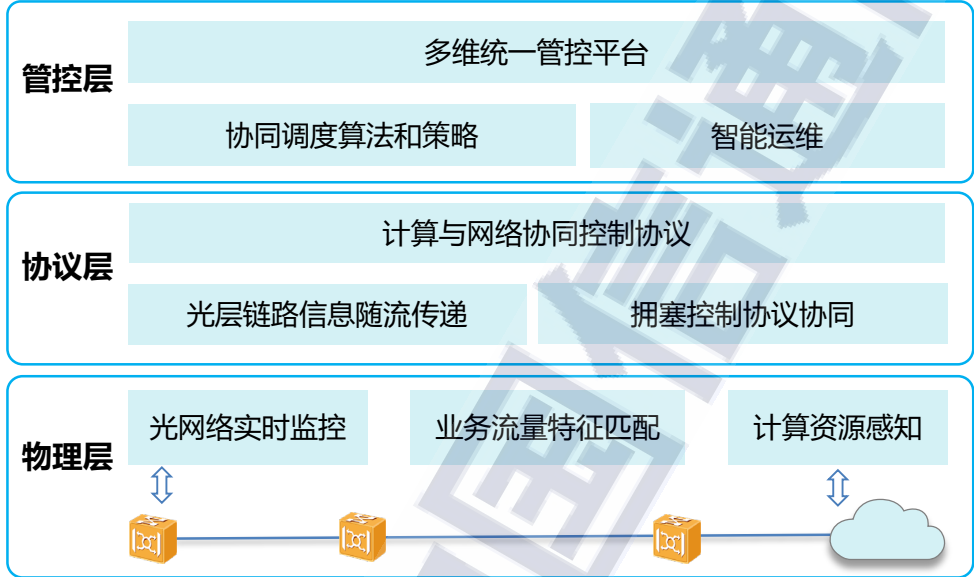
管控层按照业务诉求自动匹配波长速率，选择满足光层参数的可达路由及最优码型，并自动调用设备层接口下发光交叉配置和调测光功率。

基于多层网络隐患分析技术，实现网络品质保障。基于数据分析和智能化算法技术实现光网络电层、光层、光缆风险跨层关联分析，并通过 OTN 网络质量评估模型，将全网的故障隐患信息进行系统化呈现。在线同缆检测技术基于多模态深度学习 AI 算法，自动化分析光纤和业务同缆风险，输出同缆位置、同缆距离等关键信息；主备时延绕路分析技术与 GIS 地图结合，获取源宿 GIS 位置之间的最优路网路由，计算按路网路由的光纤传输时延，并与业务路由时延对比，得到绕路系数，通过绕路系数衡量业务时延是否存在绕路。**故障诊断实现故障精准定界，快速派单和排障。**通过智能监测和异常模式识别，基于 AI 算法模型对海量样本进行在线推理，进行掉电、断纤等多种故障模式识别，从而实现故障精准定界、快速派单和智能排障。

设备层实现网络多维感知和算力内生。为支撑算网任务式调度，光传送网（OTN）设备需全面进行升级，从纤缆、网络、业务三个维度进行感知能力提升，实现对网络的多维参数进行高精度感知。纤缆方面需实现对光缆质量、同沟风险、环境威胁的检测，网络方面需实现对光参的各项性能（合波光功率、BER、SOP）的高精度采样，业务方面需实现业务流量、时延等各项指标精准采样。为实现多维度海量数据精准分析，设备还可新增算力单板，进一步增强硬件算力，从而实现光网络的全面感知和网络设备的隐患分析及预测。

（五） 光算协同关键技术

为解决传统网络与智算资源协同问题，需要实现智算资源与高品质算力专线的协同感知和调度。通过物理层、协议层、管控层的光算协同，可显著提升客户体验以及协同计算的效率和可靠性。



来源：中国信息通信研究院

图 13 光算协同技术框架

物理层协同通过实时感知光链路状态、计算节点资源使用情况，为上层协议和管控提供准确的数据支持。**光网络实时监控技术支撑算间数据高效传输**，将光链路拓扑、传输带宽、链路时延、光传输信噪比、光纤类型等多项关键参数监控识别，实现光网络信息的实时感知，并将光网参数上报给统一管控平台，用于网络流量拥塞控制。**计算资源感知用于光网络传输资源适配**，将计算节点的算卡类型、任务日历、传输带宽需求等信息映射为光网络可识别的参数，实现传输资源与算力资源协同匹配。

协议层协同通过特定的协议和机制，实现高效光算协同。DCN 和 DCI 设备协议协同实现拥塞控制。在 DCN 网络拥塞时，OTN 光传送

设备可感知基于优先级的流控(PFC)反压信息,实现本地流控处理,避免 DCN 网络拥塞导致的丢包。**通过控制协议扩展实现计算和网络信息传递**,可通过光传感(LS)技术实现光传输链路状态信息的传递,用于 DCN 出现拥塞时进行优化选路控制,也可以通过协议扩展,用于计算节点可订阅必要的光网络信息,如链路带宽、时延、链路 SLA 等级等,同时将计算任务所需的网络带宽、时延、可用率等需求清晰传递至网络,实现算力需求与网络供应的高效匹配。

管控层协同通过统一的管控平台,实现对计算资源和光网络资源的全局管理和协同调度。**统一管控平台实现算网业务协同调度**,通过数字孪生构建“算力-运力-业务-客户”多维统一管控模型,纳入运力地图、算力地图等,统一管控整个系统的资源分布和使用情况,并在统一平台内实现协同调度。**协同调度算法和策略实现算力运力多目标的协同优化**,通过强化学习构筑协同调度算法和模型,提升算力资源利用率,均衡光网络负载,根据业务需求实现传输带宽、时延、可用率等最佳分配,实现多目标协同优化。**智能运维与优化实现运力与算力精确匹配和高效运维**,基于意图的光网络业务发放和调优技术可实现端到端光系统资源自动分配、配置下发、性能调优等动作,使得光网络能力精确敏捷匹配算力需求,基于 AI 诊断的光网络异常检测技术可实现传输网络故障的提前发现和故障类型精确识别。

五、总结与展望

随着 AI 赋能行业数字化转型的加速,各类行业智算应用爆发式增长,业务云化、视频化、高清化趋势愈发明显, AI 模型训练/推理

等创新业务不断涌现，金融、政务、教育、医疗等行业智算应用提出差异化入算网络需求。面向 AI 时代行业智算应用的差异化承载需求，高品质算力专线应具备智能感知、业务确定性体验、网络弹性按需、智能运维、光算协同等五大特征，提供安全、普惠、差异化的端网算一体化服务，赋能千行百业数字化转型。

展望 AI 时代的高品质全光算力专线发展，建议产业各方继续在关键技术创新、业务及应用模式探索等方面协同推进，持续推动算力专线向超大带宽、确定低时延、安全可靠、弹性敏捷等方向演进，为各类行业智算应用提供高效、稳定和可靠的网络连接和数据传输服务，支撑行业智算应用创新，促进数字经济高质量发展！

中国信息通信研究院技术与标准研究所

地址：北京市海淀区花园北路 52 号

邮编：100191

电话：010-62300112

传真：010-62300123

网址：www.caict.ac.cn

