



AI DC 白皮书



构建万物互联的智能世界

一份给 CIO 规划建设智算数据中心的参考



前言

算力正成为新“黑金”

十几年前，美国《时代》周刊提到：网络带宽将成为石油之后，二十一世纪的新“黑金 Black Gold”。那个时候，或许没有人预见到，十年之后的人工智能会跃迁到今天的水平。大模型的疯狂“涌现”，生成式 AI 的突然“顿悟”，一时间构筑起 AI 的“拉瓦尔喷管”，全球人工智能产业正无限逼近“迸发”的状态，人类社会将以远超我们想象的速度，加速迈向智能世界，算力正成为二十一世纪的另一个新“黑金”。

AI 是趋势，不是潮流

从 1956 年人类首次提出“人工智能”的定义以来，AI 的发展经历了多次的起起伏伏，即使在 AI 持续占据全球科技头条的今天，依然有相当数量的人和组织，对人工智能的未来表示怀疑、担忧和犹豫；但 AI 从未停止向前，技术不断创新突破，产业规模不断增大，应用从单点到多元化扩张、从通用场景向行业特定场景不断深入。AI 必将重构传统产业，并将催生出诸多新产业。

ChatGPT 的横空出世，让人类通往通用人工智能 AGI 之路从未像今天这样清晰，AI 已经是不可逆转的趋势，不是潮流。AI 驱动下，人类将从以数据（Data）和信息（Information）为主的信息社会，到以产生知识（Knowledge）和智慧（Wisdom）为主的认知社会。未来几十年，我们将迎来一场认知革命，今天的生成式 AI 只是一个开始。

这是 DC 白皮书，不是 AI 白皮书

当“百模千态”已然成型，当“千行万业智能化”快速成势，首先得到全行业重点关注的不是 AI 应用，而是 AI 基础设施。要想富、先修路，任何国家和企业，要想在 AI 时代“富”起来，首先要把 AI 基础设施这条“路”修好，而数据中心恰恰是 AI 基础设施的核心之核心。

数据中心的雏形从 1940 年前后就开始出现，随后几十年间，随着互联网、大数据和云计算的发展，数据存储和数据处理变得越来越重要，数据中心也成为企业信息化、数字化的核心基础设施。迈向智能时代，数据中心首先要提供的是算力，承载的主要是 AI 训练和推理，支撑的是企业关键智能化应用，这类面向未来的数据中心我们称之为智算数据中心 AI DC。

未来的数据中心一定是 AI 定义的

AI DC 不是传统数据中心的简单升级改造，而是数据中心的一次全方位重构。从过去的成本中心到今天的生产中心，从数据存储和处理中心，到价值创造中心。

互联网和云计算带来了软件定义（Software-defined）的基础设施，未来的数据中心基础设施一定是 AI 定义的。AI 带给数据中心的挑战也将是多维度的，如：算力密度、能源效率、AI-powered 的运维与运营以及可持续性发展等。

强大而坚实的 AI 算力底座，是智能化转型的基石。数据中心不断演进，从存储数据、支撑应用，到提供算力、承载 AI 训练和推理，再到生产智慧、使能智能化，其重要性和行业价值不断凸显，值得产业链各方重点关注。希望这本白皮书能为全行业 AI DC 的规划与建设提供一些参考。





杨超斌

—— 华为公司董事、ICT 产品与解决方案总裁

最近一段时间，围绕企业 AI 落地、AI 算力基础设施建设，我与很多客户伙伴、AI 生态链的朋友们进行了沟通交流，大家基本有一个共识，都把 AI DC 建设作为企业智能化转型的优先举措；但 AI DC 与传统数据中心存在很大区别，在企业数智基础设施中的定位变了、承载的业务变了、数据处理和算力提供的要求也变了，再加上技术还在不断创新升级，如何高效高质量建设 AI DC 值得全行业深入思考。从实践中进行复盘总结，汇聚全行业智慧，这就是这本白皮书的价值。



李鹏

—— 华为公司高级副总裁、ICT 销售与服务总裁

中国有句谚语：“要想富、先修路”，建设高质量的 ICT 基础设施是企业数智化转型、实现商业成功的基石。AI DC 作为新一代数智基础设施的核心，华为在过去几年与客户的建设实践与创新探索中，有经验、有教训，也还存在许多新课题需要大家一起解决。这本白皮书只是一个开始，全行业需要协同创新，共同推动 AI DC 发展，携手迈向智能时代。



谢海

—— 中国铝业集团 CIO

千行万业正在积极拥抱人工智能，把行业知识、创新升级与大模型能力相结合，以此改变传统行业生产作业、组织方式。在如何用好人工智能方面，有色行业不断探索，聚焦人工智能服务有色场景，在氧化铝、电解铝、高端铝加工等领域持续实践。这本白皮书提出了很多可供企业参考的观点，特别是针对如何规划建设企业数智基础设施的核心 —— AI DC 上，给出了方向性的建议和非常实用的评估指标，而这也是企业落地 AI 的最关键一步。



王磊

—— 太平洋保险集团数智研究院院长

生成式 AI 为保险行业发展提供了新质生产力，场景落地和价值闭环是当前核心问题，不论是技术探索，还是大规模应用部署的效率和成本考量，对企业 AI DC 的建设和运营都提出了极高的要求。白皮书基于技术趋势和产业实践，系统性地阐述了 AIGC 产业应用的建设策略和实现路径，并给出不同场景下的 AI DC 建设方案，具有重要参考价值，激发深入思考。



邹志磊

—— 华为公司高级副总裁

智能时代，AI 只有进入企业的核心生产场景才能发挥巨大价值，这势必驱动企业业务系统从传统的“构成式”变成“生成式”。企业智算数据中心作为数智基础设施的核心，将从成本中心变成创新中心，技术架构也会发生颠覆式变化，传统数据中心时代的建设模式、系统架构、运维运营等可能都不再适用。这本白皮书是对当前行业实践的总结和复盘，面向未来我们还将持续探索和思考，就如何规划建设好 AI DC 给出更多参考建议。



马海旭

—— 华为公司副总裁、ICT 产品组合管理与解决方案部总裁

人工智能应用繁荣的基础是算力。作为提供算力的关键数智基础设施，AI DC 需要充分发挥计算、存储、网络、云、能源等技术领域的综合优势，以系统架构创新，持续突破规模算力瓶颈。从 2019 年发布 AI 战略及解决方案开始，华为就广泛参与到全球客户 AI 算力基础设施的建设实践中，并不断与产业链相关方开展联合创新，打造有竞争力的产品与解决方案，为客户创造价值。把这些有价值的客户建设实践与全行业的智慧汇聚在一起，形成了这本白皮书，希望帮助客户更快更好地建设 AI DC，加速千行万业智能化转型。



何宝宏

—— 中国信息通讯研究院云计算与大数据研究所所长

走向智能时代，AI DC 将是整个智能社会的坚实底座。中国各级政府在布局引导、建设规划、技术创新和应用赋能等方面持续出台举措，推动算力基础设施发展。企业也不断加快探索实践步伐，推动 AI DC 向大规模、高质量和强应用的方向发展。本研究报告在规、建、管、用等多个维度，体系化梳理，立体化呈现 AI DC 最新态势，有助于促进产业发展。



苏廉节

—— Omdia 人工智能首席分析师

AI DC 承载的是人工智能应用、训练和推理等工作，与其他类型的数据中心存在很大的差异。当前的人工智能发展迅速，新技术新应用层出不穷。如何去构建一个坚实的算力底座来满足长远未来的发展需求和应付人工智能应用的迭代演进是每个企业都必须去迎接的新挑战。

目录

第一章	
AI World 总体愿景及宏观驱动力	10
人工智能是一个大方向，不可阻挡	11
AI for All	15
理想主义与现实主义交相辉映迈向 AGI	17

第二章	
All in AI 生成式业务系统	18
企业发展 AI 的不确定性和确定性	19
架构先行，将不确定挑战变成确定机遇	21
应用场景为纲，四位一体，实现价值三角	23
以数据中心为中心	32

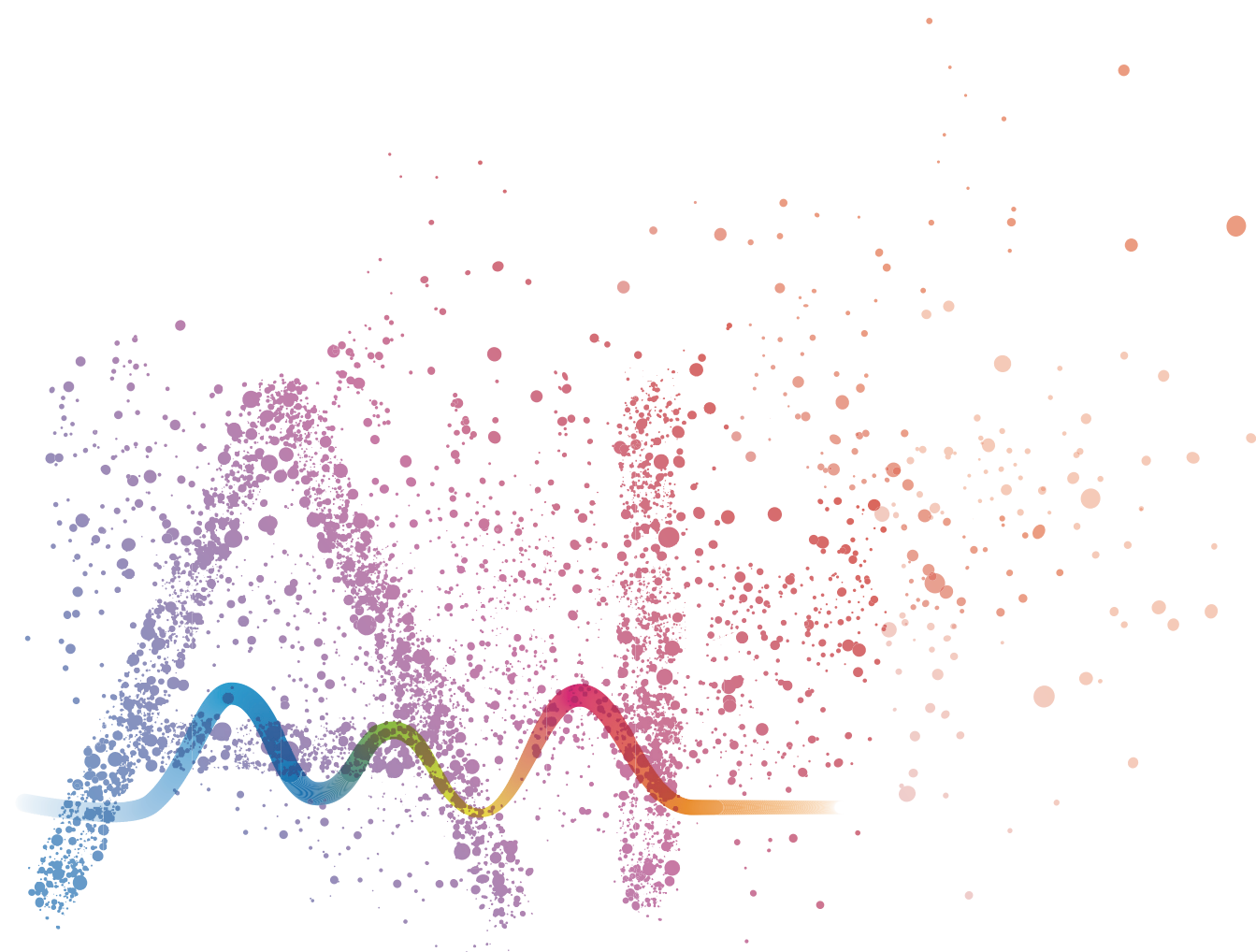
第三章	
智能时代数据中心的发展与变化	34
数据中心走向智算数据中心 AI DC	35
AI DC 主要承载 AI 模型的“训推用”	37
AI DC 四大建设场景及三大类型	39
AI DC 五大特征变化	43
数据中心将被重塑，由分层解耦到垂直整合	53

第四章	
典型 AI DC 规划与建设	56
超大型 AI DC	57
大型 AI DC	72
小型 AI DC	88

第五章	
AI DC 建设与发展倡议	94
行动倡议一：	
适度超前建设 AI DC	95
行动倡议二：	
共同实现 AI DC 集约化建设和绿色发展	98
行动倡议三：	
共建开放协作的行业 AI 生态	99
行动倡议四：	
筑好三个底座，加速行业 AI 走深向实	100

「第 1 章」

AI World 总体愿景 及宏观驱动力



人工智能是一个大方向，不可阻挡

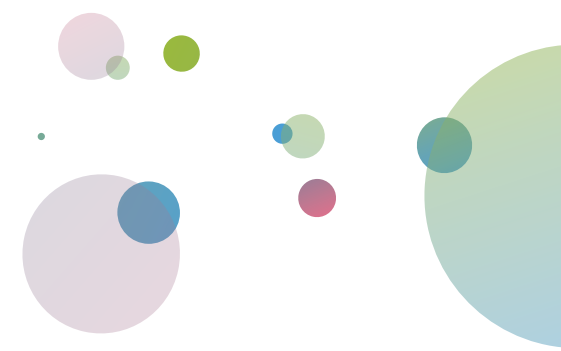
生成式 AI 日新月异的发展，让人工智能如风卷残云般走向舞台中央。

根据中国信息通信研究院的相关报告显示，截止 2024 年 7 月，全球 AI 企业近 3 万家，全球人工智能大模型有 1328 个，其中中国各类企业在不到 2 年时间就上市发布 478 个人工智能大模型。

人工智能正引发全产业链的新一轮工业革命，也将给人类社会带来一个“天大的机会”。斯坦福大学“Human-Centered”人工智能研究所发布的《2024 年人工智能指数报告》显示，从 2023 年到 2024 年第一季度，全球 AI 独角兽已有 234 家，新增数量为

37 家，占新增独角兽总量的 40%；2023 年，虽然全球 AI 投资总额有所下滑、降至 1892 亿美元，但生成式 AI 领域的投资激增，比 2022 年增长了近 8 倍，达到 252 亿美元。

六十年的芯片技术发展，三十年的互联网发展，Transformer 架构的不断突破，以及数据的极大丰富，让 AI 技术不断走深，AI 应用不断向实。继 OpenAI 公司推出 ChatGPT 之后，2024 年华为公司推出的盘古大模型 5.0 版本，以及 Anthropic 公司推出的大模型 Claude 3.5 Sonnet 版本，宣告大模型从“聊天”正式迈入“工作流”。



AI 是过去 70 年 ICT 产业发展的总成果

1956 年，时任达特茅斯学院助理教授的约翰·麦卡锡组织召集了达特茅斯讨论，正是在这次会议上，第一次正式提出了“人工智能”的定义。从那以后，人工智能经历了两次发展的低谷，即所谓的“冬天”，但其发展的脚步并未就此停止。

自从 1971 年英特尔发布第一颗微处理器开始，摩尔定律见证了 ICT 产业的蓬勃发展。如果把 AI 产业和 ICT 产业这 70 年的发展轨迹画到一起，我们发现，

人工智能与 ICT 产业的总体发展水平密切相关，学术研究发现和工程技术发展相辅相成。而 AI 产业两次“冬天”的出现，都是因为社会对 AI 的应用期望大大超越了 ICT 产业工程水平的发展现实。所幸的是，“冬天”并不是结束，而是每一次“春天”的开始。今天，我们再次进入了“收获”的季节。这是 70 年来全球 ICT 学术界和工业界长期耕耘、协作创新的成果。

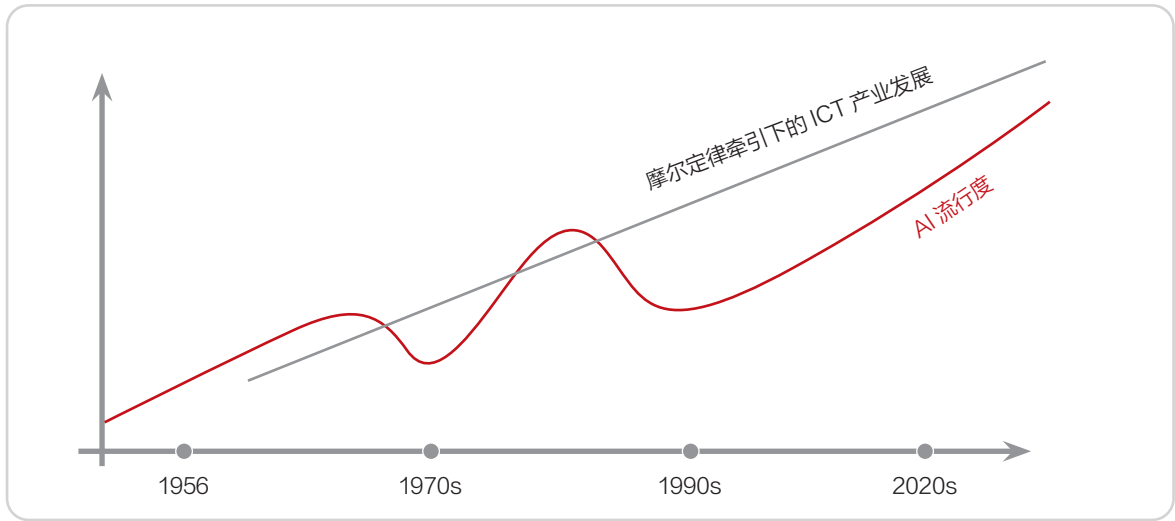


图 1-1 AI 是过去 70 年 ICT 产业发展的总成果

技术的准确定位是发挥其最大价值的前提。给人工智能技术进行合理的定位，是我们理解和应用此技术的基础。如同公元前的轮子和铁，19 世纪的铁路和电力，以及 20 世纪的汽车、电脑、互联网一样，人工智能是一组技术集合，是一种新的通用目的技术。加拿大学者 Richard G Lipsey 在其著作《经济转型：通用技术和长期经济增长》一书中提出：社会经济的持续发展是靠通用技术的不断出现而持续推动的。所谓通用技术，简单理解就是要有多种用途，应用到经济的

几乎所有地方，并且有巨大的技术互补性和溢出效应。经济学家们认为，人类发展到今天，共有 26 种通用技术，受益于过去 70 年 ICT 产业的总体发展，人工智能成为其中一种。

面向未来，我们应该充分用好人工智能技术，抓紧收获，努力扩大收获成果，同时要让收获的季节持续的更长一些，把人工智能建在赤道上，永远生机勃勃。

AI 将引发百年未有之大变革

纵观人类社会发展史，通用目的技术的大规模应用历来是社会变革的催化剂，而被彼得·戴曼迪斯在《未来呼啸而来》一书中定义为“指数型技术”之首的人工智能，将引发一场百年未有之大变革。自十八世纪蒸汽机问世，科技创新将时代划分为蒸汽时代、工业时代与信息时代，现今，智能时代正扑面而来，其背后的驱动力正是 AI 算力。这股力量不仅将为公众生活注入个性化与便捷体验，还将以创新逻辑推动各行

各业效能提升与经验革新，为科研开辟新路径。AI 的普及深化不仅会加速传统产业智能化转型，优化资源配置，提升决策质量，激发产品与服务创新，还将进一步优化社会经济结构，推动全球经济步入高质量增长新周期。

AI 引发的变革将是一场体验革命、效率革命、经验革命和科研革命，以智能化为标志的新时代已经来临。



图 1-2 人类进入智能时代

智能经济将是数字经济发展的下一跳

当前，全球数字经济保持持续快速发展。根据中国信息通信研究院相关报告显示，2023 年，美国、中国、德国、日本、韩国五个国家的数字经济总量已逾 33 万亿美元，年增长率超过 8%，数字经济对 GDP 的贡献达到 60%。这不仅彰显了数字经济的迅猛发展，更凸显了其在全球经济版图中的核心角色。其中，人工智能推动的经济发展规模是关键力量。

数字经济的进化始于个人计算机的发明和普及，继而在物联网与移动互联网中成熟，今天正步入人工智能为核心的智能经济新阶段。**智能经济是以效率、和谐、持续为基本坐标，以物理设备、电脑网络、人脑智慧为基本框架，以智能政府、智能经济、智能社会为基本内容的经济结构、增长方式和经济形态。**作为全球经济的新引擎，智能经济致力于主导效率提升、和谐发展与可持续发展的全球经济未来图景。

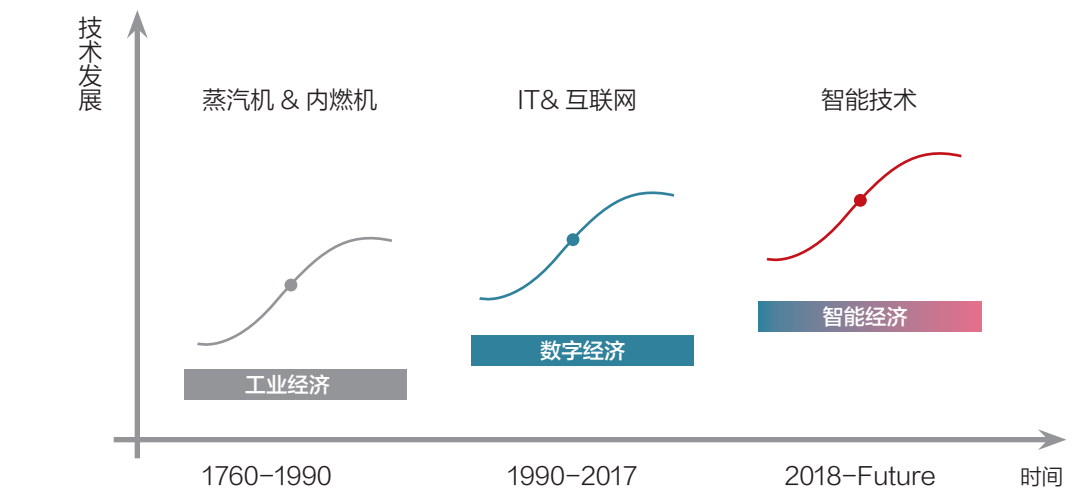


图 1-3 智能经济将成为全球经济发展新引擎

人工智能驱动的智能经济将在人机交互、IT 基础设施与新业态三个层面带来重大变革。首先是人机交互模式的优化，让交流更加自然流畅；其次，它将重塑 IT 基础设施，构建更高效、更智能的信息处理与传输体系；最后，智能经济会催生一系列新业态，激发跨领域创新。这三个方面并非孤立存在，而是相互影响、协同演化，形成合力并产生复合效应。

过去四十年，信息化和数字化给 ICT 行业带来了数万亿美元的战略机会。展望潜力十足的智能时代，华为预测，至 2030 年，全球智能经济规模将超过 18.8 万亿美元，将为 ICT 领域的未来发展开启全新战略窗口。

AI for All

AI 的快速发展和大模型的“涌现”预示着它将重塑每一个组织和每个人的生活。专家和机构预测 AI 将深刻影响世界。那目前企业和个人对 AI 的接受度及应用进展如何呢？

麦肯锡 2023 年的报告指出，55% 的组织已在至少一个部门采用人工智能，这一比例是 2017 年的 2.75 倍。Gartner 在其《2024 年重要战略技术趋势》报告中预测，到 2026 年，超 80% 的企业将运用生成式 AI；到 2028 年，75% 的企业软件工程师将使用 AI 编码助手，而 2023 年初这一比例不足 10%。

每个行业都将被 AI 重塑

在人工智能触发的产业变革大潮中，所有行业都将被重塑。今天我们已经可以清晰地预见一些行业将发生怎样的变化：

- 自动驾驶和电动汽车将颠覆汽车行业
- 智慧交通将大大提升通行效率
- 个性化教育将显著提升教学质量
- 精准预防性治疗有望延长人类的寿命
- 实时多语言翻译让交流再无障碍
- 精准药物试验可以显著降低新药研发成本，缩短发现周期
- 基于 AI 的电信网络的运维效率将大大提升
- ...

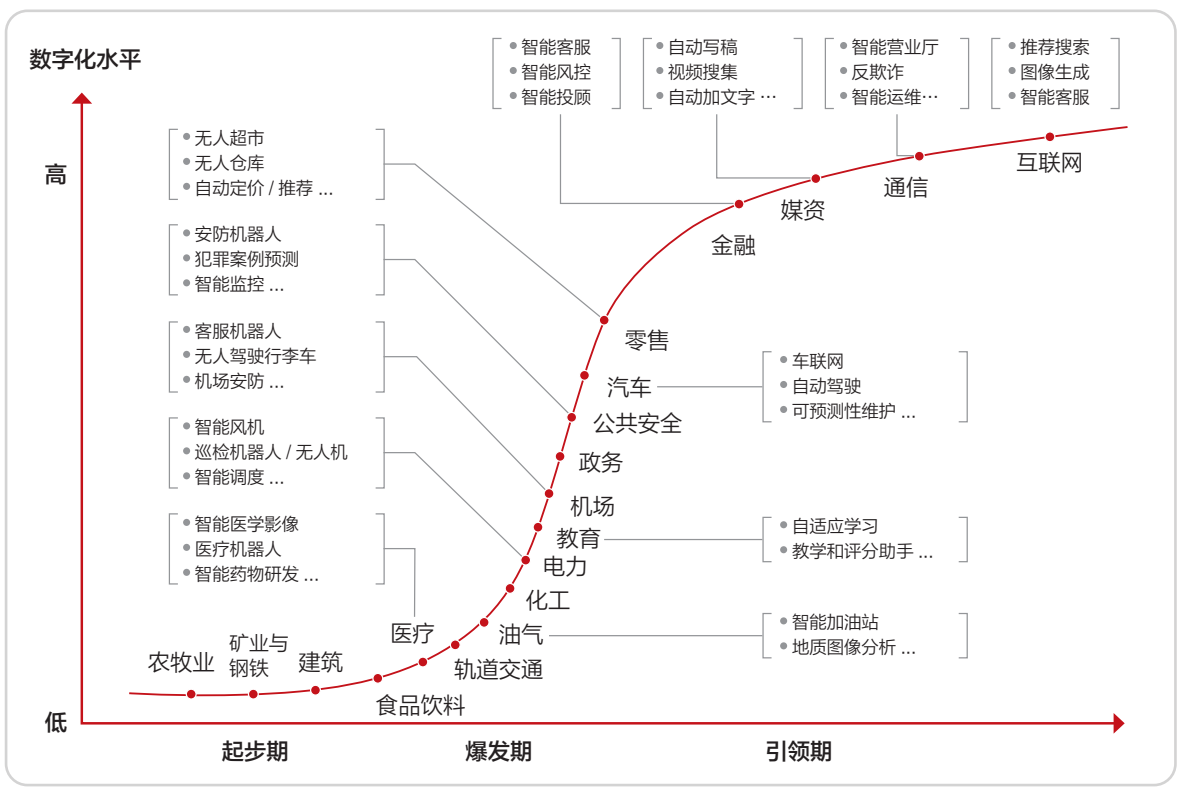


图 1-4 AI 正在改变千行万业

AI 重塑行业的速度确实远超想象。2023 年初，比亚迪提出实现自动驾驶还需时日。仅仅一年多过去，自动驾驶技术的迭代升级已经悄然发生，2024 年中国新能源汽车市场中，自动驾驶功能的渗透率已超过 51%。这一成就的背后是先进的感知系统、强大的计算平台、AI 驱动的决策与规划算法。

AI 不仅仅能够助力一个行业，也可能颠覆一个行业。印度 IT 服务外包业曾凭借人力成本和语言优势成为

全球中心。然而，AI 技术的兴起导致该行业面临严峻挑战。据统计，过去一年，印度五大 IT 服务公司裁员 69,197 人，创下 20 年新高。这一现象背后，是 AI 在服务领域的广泛应用，高效接管了原本由人力完成的任务。由此可见，AI 技术不仅仅能通过技术革命将一些行业带入新阶段，更能够淘汰替代一些相对落后的生产力方式。未来，我们完全有理由相信，AI 将有能力重塑每一个行业。

每个应用和软件都值得用 AI 重写

生成式 AI 是革命性的跨越，有人称之为 AI 2.0，它不是 AI 1.0 的升级。AI 2.0 可以用无需人工标注的超级海量数据、去训练一个具有跨领域知识的基础大模型（Foundation Model），它能够从无到有，真正实现智慧的产生；AI 2.0 让每个人都能创造，甚至

可能让每个人成为程序员，催生了数字分身等长期以来仅存于想象的产品。AI 2.0 的生成能力还能将创新实现成本降到几乎为零，创造出更赚钱的商业模式。

AI 2.0 的创造能力和商业能力，让智能时代的每个应用和软件都值得重写一遍。



图 1-5 每个应用和软件都值得用 AI 重写一遍

理想主义与现实主义交相辉映迈向 AGI

2015 年，OpenAI 牵头启动的 AGI 实验，成为人类迈向通用人工智能（AGI）的一个新起点。随后，2020 年 GPT-3 的推出，以及 Scaling Law 被确立为 AGI 的第一性原理，标志着人类向 AGI 目标的探索步伐大大加快。为了支撑 AI 能力的持续进化，投资规模超过 1000 亿美元的“星际之门”计划启动，旨在构建更加强大的算力基础设施，预计 2028 年将发布一个由数百万 XPU 算力卡互联的集群数据中心。

理想主义者们相信，跨越技术裂谷的人工智能将加速前行，他们致力于在未来 10 年内将深度学习的计算能力提升 100 万倍。AI 领域的新论文、新模型层出不穷，从 Pretrain（预训练）到 SFT（监督微调），数据来源从公开网络扩展到合成数据，AI 的技术发展让所有人感受到了强烈的“推背感”，人类终将走向 AGI。

然而，我们也看到，AI 在面向消费者（ToC）的应用和面向企业（ToB）的行业落地中，依然面临诸多挑战。许多 AI 应用和项目仍处于起步阶段或短暂出现后便消失，实现商业闭环成为业界关注的焦点。对于人工智能产业的发展战略制定者来说，是选择“登顶珠峰”，将 Scaling Law 推向极致，无限接近 AGI；还是“沿途下蛋”，尽快实现技术落地并盈利，快速融入商业社会，这是需要深思熟虑的问题。

大多数新兴技术的发展都是从理想主义的美好愿景开始，同时受到现实主义的理性制约。如果能够将理想主义和现实主义相结合，无疑将加速技术的成功落地。我们认为，人工智能是一个不可逆转的大趋势。AI 产业在垂直方向上，既需要科学家的理想主义，也需要与商业现实主义相结合，寻找技术驱动与商业落地之间的平衡。

理想主义者的代表是工程师和科学家，他们基于科技改变世界的理想化出发点，用探索精神和创新思维，致力于开发更智能、更自主的学习算法，追求更高的计算效率和更低的能耗。这些努力不断拓展 AI 技术的可能性边界，为现实应用提供了丰富的理论支撑和技术储备。而现实主义者的代表是理性的市场经济参与者，他们将 AI 技术视为推动商业变革和社会进步的关键力量，注重技术的实用性和经济效益，主要将 AI 的商业化落地作为目标，使其融入金融服务、健康医疗、零售物流等行业场景。他们希望通过实践验证 AI 技术的市场价值，为持续发展提供应用场景和反馈数据，激发新的研究方向和创新灵感。

AI 技术的演进历程正是理想主义与现实主义辩证关系的生动体现，二者相辅相成、交相辉映，共同塑造人工智能的未来。理想主义与现实主义产生了奇妙的双轮效应，每一次技术飞跃都会带动商业应用的创新与拓展，而商业成功又会以更多的研究资金和资源反哺科研领域，推动技术的进一步成熟和完善。这种正向循环一旦建立，就能帮助企业采用新技术时实现新的价值链闭环。成功的案例将加速 AI 技术在各行业核心生产环节的渗透，推动一系列高效、智能的解决方案的形成，创造可观的商业价值和社会福利。

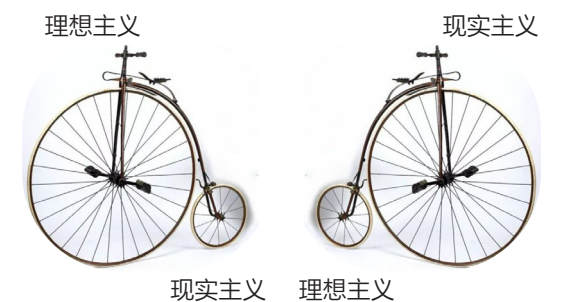


图 1-6 理想主义与现实主义交相辉映推动 AI 发展

「第2章」

All in AI 生成式业务系统

企业发展 AI 的不确定性和确定性

根据麦肯锡的调研，超过 70% 的企业领导者预见 AI 将在接下来的五年内深刻改变其业务格局。同时，企业发展 AI 有着相当大的不确定性，据德勤的数据显示，90% 的大型企业计划投资 AI，但真正能够成功规模化部署的仅占 10%。

这是因为生成式 AI 的革命性创新和内在局限性兼而有之。

一方面，ChatGPT 对奥林匹克数学竞赛题可以给出优雅的证明；另一方面，在回答 13.11 和 13.8 比大小的试题中输给小学生。一方面，自动驾驶技术正在颠覆汽车行业，改变大众的出行服务；另一方面，提升辅助影像诊断的医疗专用模型仍旧在创新研究阶段。一方面，50 位艺术家通过 AI 生成了首部充满创意的科幻电影；另一方面，很多企业还在被灵魂拷问：巨大的 AI 投资换来写作助手是否值得？模型回答质量的稳定性何时才能解决？

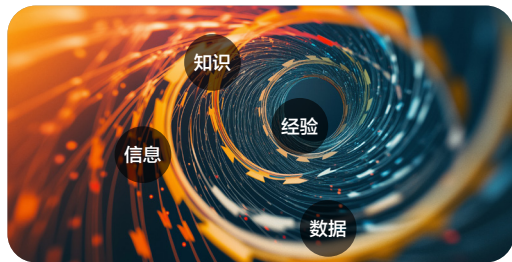
对于企业来说，是追逐潮头引领行业革新、还是岸边试水等退潮的鱼？

ChatGPT 等大语言模型带来的革命性变化，源于其汇聚世界知识带来的泛化能力，本质是显性知识的压缩和隐性经验的沉淀，是基于结构化数据发现内在规律的概率模型。各行各业尤其是头部企业，往往蕴藏着海量的数据、沉淀的业务知识和内化于业务流程的经验等宝贵资源，当它们被用于语料来训练 AI 模型时，模型自然就记忆了这些知识与经验。企业通过引入基础模型、行业模型并构建自己的私有化场景模型时，相当于“一杯咖啡吸收宇宙能量”，可以更高效的传承和利用企业内部经验、行业经验、世界知识，从而实现企业可持续发展。

企业最大的浪费是“经验和人才的浪费”。基于此理念，华为公司的企业 AI 从 1.0 向生成式 2.0 演进，AI 应用到更多的核心业务领域，从合同风险审计到支撑全球供应链在疫情中的韧性管理，从全球网络优化到提供互联网信息产品的极致体验，从专业又有温度的智能客服到海量高可信代码的生成等。华为 AI 2.0 的目标是实现“1 个顶级专家 + AI 能力增强型数字员工 + N 个普通员工”的效率等于甚至大于 N 个顶级专家。

企业最大的浪费：是经验和人才的浪费

- ① 海量数据沉睡
- ② 沉淀的知识和经验的低效运用或流失



企业的持续成长性：一杯咖啡吸收宇宙能量

- ① 吸收社会知识沉淀的能力决定企业持续成长能力
- ② 走向世界走向开放，一杯咖啡吸收宇宙能量



图 2-1 企业发展 AI 的确定性

AI 大模型带来的创新性机遇，源于科学范式的变化，从海量数据中发现未知规律。越来越多企业希望 AI 能够基于核心生产场景，创造企业产品和服务的核心竞争力，先行者可以建立领先能力。比如特种钢的误差要求严苛，液面波动是炼钢的关键参数之一，结晶器液面波动与液面高度、水量、温度、压力、原材料批次等 200 多种参数相关，超过专家的经验 and 科学公式计算的适用范围。钢铁企业在思考如何利用 AI 优化生产制造工艺，基于积累的高价值历史数据训练场景模型，并在实时生产过程中不断反馈增强，找到最优的液面波动参数。

企业发展 AI 需要构建企业级综合智能体。如同一个

能工巧匠，在解决复杂问题时，将书本学到的显性化知识和实践中积累的大量隐性经验相结合，并实现从感知、理解、预测分析到决策的闭环。我们欣喜地看到，AI 大模型正将海量、多源、非结构化数据实现结构化，并贯穿感知、预测到决策全流程。当 AI 的视野从语言文字预测，延伸到声线、物体的色块、时序的采样、分子结构、调度网络负荷等更贴近现实世界的场景时，将为企业 AI 带来无限机遇。

建议企业战略上要明确发展 AI 的确定性，战术上要应对好 AI 的不确定性。从现在开始、着眼未来，以 All in AI 为战略，选择合适的节奏，并在生态模式上采取灵活战术，是企业发展 AI 的最佳选择。

架构先行 将不确定挑战变成确定机遇

构建企业级 All in AI 架构的核心挑战可以归结为两个简单的几何图形：哑铃型的非稳定性结构和行业大模型的不可能三角。

架构挑战之一：哑铃型的非稳定性结构。企业传统 IT 架构是稳定的正三角，基础设施和技术平台稳定，变化频率低；数据和应用使能平台按照产品化、版本化的方式迭代，变化可预期；应用受用户体验驱动，更新需敏捷化高频。AI 大模型时代，IT 架构增加模型层，而模型因处于快速发展迭代期，变化幅度和升级频率均超过应用。如何规划设计 IT 架构，实现“在行驶中换发动机”？

架构挑战之二：行业大模型的不可能三角。大模型在泛化性、专业性、经济性三方面很难兼得，泛化性强调基于小样本的场景化学习能力，专业性强调监督学习能力强，经济性强调模型规模适中。同时，不同类型不同场景的大模型平衡点不同，如语言类参数量大、算力高，经济性要求高；产品质检视频类负样本少，泛化性要求高；风险预警类对精度要求苛刻而专业性要求高。由于行业数据的稀缺性，行业模型追求泛化性和专业性的矛盾尤其突出。

企业发展 AI 的核心理念是：以架构的确定性应对模型的不确定性，形成具备持续开发态模型层的非常规稳定架构。应用层以 All in AI 为蓝图进行长远规划、小步迭代，基础设施和 AI 技术平台保持稳定，震荡中心的模型层分别与应用层和基础层实现解耦。

● **模型多源：**算力底座封装软硬件的复杂性，弹性资源调度解决算力效率，服务化的标准接口对接开放的模型层，支持来源多样的模型

● **三重进化：**模型能力进行 API 封装，应用与模型解耦，形成可替换的“发动机”；L0 基础大模型随产业进化，L1 行业模型随行业模型市场、行业生态或集团中心云进化，L2 场景模型可以在企业侧微调进化

● **应用编排：**业务从边缘、支撑型应用到核心生产应用，按需组合交互理解(NLP)、感知(CV)、仿真预测、决策优化模型和检索能力，API 轻量式嵌入或助手型接入业务流程。



图 2-2 以架构开放支持进化中的百模千态

可控的开放生态应对行业模型的不可能三角，构建按需组合的行业模型层。一方面拥抱标准和行业生态，保障按需融入与利用行业生态；另一方面建立企业的 AI 应用金字塔结构，分类管理超级应用、头部应用、

刚需应用、普通应用等，根据企业的竞争力策略、能力等，灵活选择自主开发、战略伙伴联合攻关和生态伙伴供应等不同模式，实现自建和共建生态模式的平衡。

应用场景为纲，四位一体，实现价值三角

企业发展 AI 的初期容易以模型为纲，从技术出发，基于产业的基础大模型能力去“临摹”容易落地的应用，可能会导致应用、模型、算力基础设施的烟囱式发展。

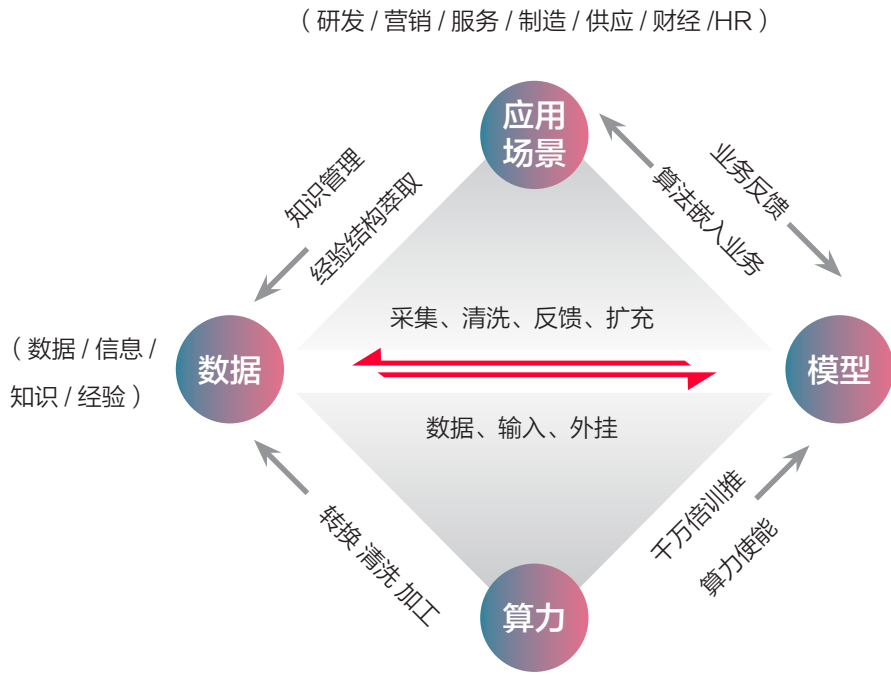


图 2-3 四位一体企业 AI 发展框架

应用场景为纲的实质是从解决问题的第一性原理出发，场景是起点也是终点，是价值的闭环。不要先关注大模型本身和模型参数量，而要看是否能够解决过去解决不了、或解决不好的问题，是否能够实现收益大于成本的正循环，是否具备广泛的适用性和可复制性。在行业 AI 应用中为提高模型的解释性和确定性，常常采用 AI 模型与机理模型结合的方式。比如勘探中，AI 模型优化钻探位置的选择，机理模型则确保开采方案的物理可行性和安全。

四位一体是指在实现应用场景价值闭环的过程中，应用场景、数据、模型和算力四个要素缺一不可。场景是价值闭环的基础，低业务价值而又消耗大量算力的

场景，就好像在非主业领域组建一个顶级专家团队；模型和数据如果不能很好匹配，如模型泛化性差又没有足够的样本，就会导致模型的专业性和精度不足，就像雇佣再多的实习生，也难以高质量地完成复杂的工作。四位一体落地时，分为技术三角和业务三角，实现技术和业务的解耦，便于建立平台化的技术架构。技术三角以算力为基础，实现数据的转换、清洗和加工，加速大模型的训练和推理，而包含知识与经验的广义数据支撑模型的训练和能力增强；业务三角以应用场景为原点，进行知识管理和经验结构萃取，不断丰富企业数据集，数据与模型双向交互，实现业务支撑和效果反馈，“非正常即异常”作为最典型的例子说明了模型使用中对数据集的反馈和补充作用。

场景 从易到难，沿着企业价值流的方向，逐步深入核心和生产场景

企业发展 AI 首先要梳理应用场景，建立“点线面”的场景地图。而 AI 业务价值三角，则可作为识别场景业务价值的经验范式和向导。其中，通过 AI 助手提升业务效率和用户体验，是企业 AI 应用最基础和常见的方式，如办公、HR、客服等；当 AI 深入生产环节后，常常能够带来生产力和竞争力的提升，如在线顾问、工艺优化、需求和供应预测等；最后是对黑天鹅式低概率风险的防范，如业务连续性风控、财务风险识别等。

企业落地 AI 需要积微成著。绘制场景地图时所谋者大、所虑者远，不用局限在已知的模型能力、已就绪的数据中，要从企业业务发展战略、AI 技术核心原理、行业发展趋势的角度构思和规划。制定实施路线图则需要从小处着眼、近处着眼，从一个个具体场景作为“小切口”入手，尽量让业务系统和流程不动或少动；基于具体场景做能力分解，组合感知、理解、预测、决策等模型能力。任务的分解让问题的求解更容易，更有利于发挥全生态的能力，由少到多形成场景飞轮。

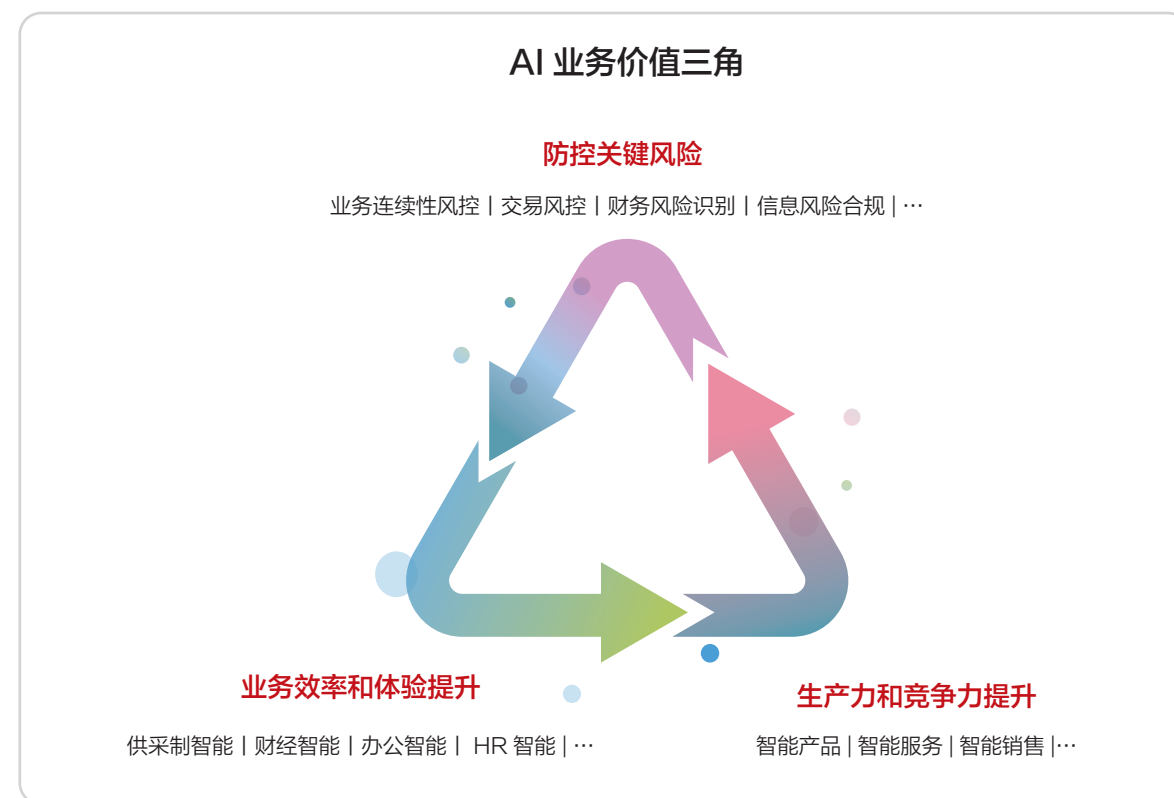


图 2-4 场景选择的价值三角

场景落地选择从三个维度入手：业务准备度、技术准备度和数据准备度。业务准备度衡量与场景相关的业务流程是否清晰、业务规则是否固化、业务组织是否有意愿和决心投入，是否有熟悉业务规则的专家投入；技术准备度衡量场景可能涉及的算法模型、装备服务、算力等是否完备，是否匹配价值期望；数据准备度衡量场景所需的数据量、数据质量、数据分布、数据标注是否完备。

场景选择的总原则是先易后难，先在实现较简单的高频、刚需场景小切口启动，快速找到智能化价值并同步培养人才，然后持续迭代、螺旋式发展。行业的引领型企业通常可以选择已具备相对充足的数据积累的领域，聚焦高价值的“超级场景”，如钢铁冶炼的“高炉场景”、化工的“中试场景”等，联合行业研究机构、AI 科技公司、大模型公司等联合攻关，一旦突破将释放巨大行业价值。

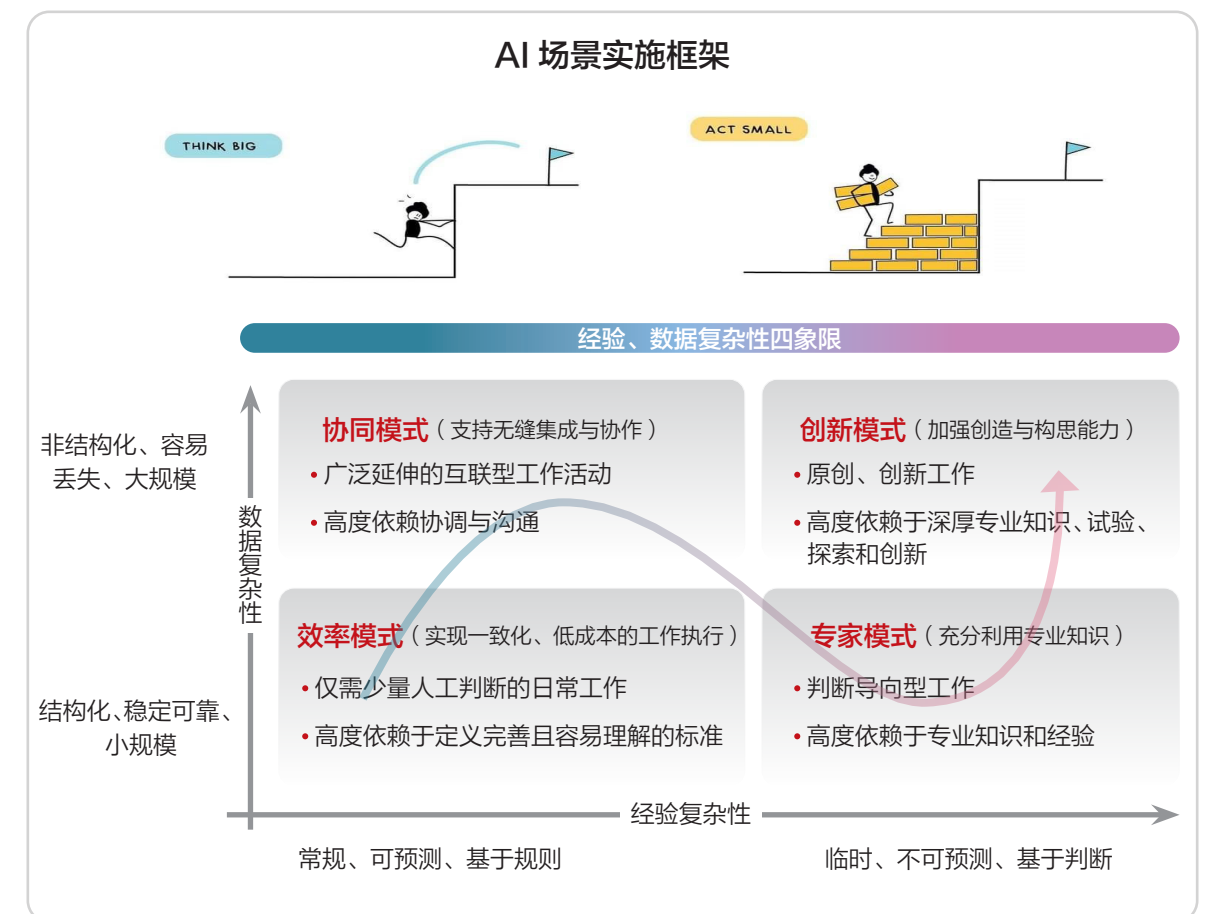


图 2-5 场景选择由易到难的路径

模型 行业 +AI 的关键路口来临，并不在语言大模型的延伸线上

语言大模型表现出强大的“内容生成”能力，不仅仅是人机对话、知识问答，还包括看图说话、情绪识别等非结构化信息生成结构数据的能力，工程设计、代码设计等非结构化强规则的文本生成能力。围绕知识密集型场景，在语言大模型的延长线上出现了大量数字化助手类应用，如客服、代码开发、专业问题咨询、舆情分析、辅助设计等。

但是，真实的行业问题并不能仅依靠语言大模型来解决。从智慧城市的内涝预警、电动车充电起火风险预测、供水损耗控制，到工业工艺配方和过程参数控制，再到金融信用评估，都面临各自的“高炉问题”，需要将机理分析与 AI 模型结合起来，将感知、理解、预测、决策的多个模型结合起来，并考虑实时性要求和模型的经济性。



图 2-6 模型选择方法

随着行业逐步理解这些需求，“合适”成为了模型评估的新标准。“大”和“统一”不再是首要追求，不再单纯追求规模和参数数量，而是要根据实际需求做出权衡。大小模型各有所长，结合场景的多样性和复杂性，灵活运用不同模型成为了未来的趋势。**模型的合适性与适用性变得比规模更为重要。**

数据 数据之道延续，AIGC 治理结构变革，价值最大化

长远看，随着基础模型的同质化和算力稀缺性缓解，个性化数据将决定企业 AI 的差异化能力。企业数据之道将延续，但治理结构也由于生成式 AI 的特点而变革。

智能时代，数据作为企业战略资产的地位进一步加强，过程的、多维的、海量的细微原始数据，以及顶端行业专家实践中产生的业务判断和执行结果成为最宝贵的资产。**海量的历史、过程数据的存储不再是纯粹的成本，而是持续积累的 AI 资产。**

生成式 AI 导致数据安全治理结构发生体系性变革，模型记忆数据，模型生成数据，模型形成企业内外新的数据边界。大模型将数据、知识沉淀在模型的参数中，并且生成文本、视频、策略等数据，导致应用和数据的划分不再清晰，企业的数据边界控制难度增加，整合行业数据和本企业数据成为重要课题。沿着原始数据、训练数据集、AIGC 模型和模型服务的依赖链，以数据的原始保护等级为原则，在域间采用可溯源、可管理的访问控制。

数据直接影响模型的表现，但数据适合轮盘式的发展。数据不搞大而全，要“先易后难、以用促建”，从具体场景入手，基于具体场景模型效果不断对数据反向提出要求，获取更多数据，让模型效果越来越好，由小滚大形成数据飞轮。

数据治理是数据质量的保障，最佳的治理是基于数据采集的源头式治理。在智慧城市、矿山、油田、工厂等大量行业场景中，涉及的终端、传感器、装备数量大、类型多，特别是多主体的场景中，通过统一智能终端和数据采集的标准规范，能够极大降低数据治理的成本。通过边缘推理与中心训练的协同，视频感知场景的异常自动标注，或者将数据标注的工作集成在业务人员的执行操作流程中，低成本地获得高质量的标注数据。

数据即业务，数据价值的最大化作为数据治理总目标，AI 应用于全数据价值链，从数据再生产、数据标识到规律发现。首先，模型应用于海量、异构数据的处理及数据产生，能够将各类异构数据，如图纸、视频监控、互联网舆情等转化为结构化的信息，为数据分析和风险评估提供坚实基础。其次，模型帮助实现可信、精准的数据跨部门共享，通过共享高阶数据，如视频中人或物的安全状态，实现数据可用不可见，确保在充分利用数据价值的同时，严格保护隐私和数据安全。最后，模型实现基于全域数据的预测和决策，各业务单元基于自身和关联主体的数据实现更准确的预测，能够发现更多、更复杂的规律。

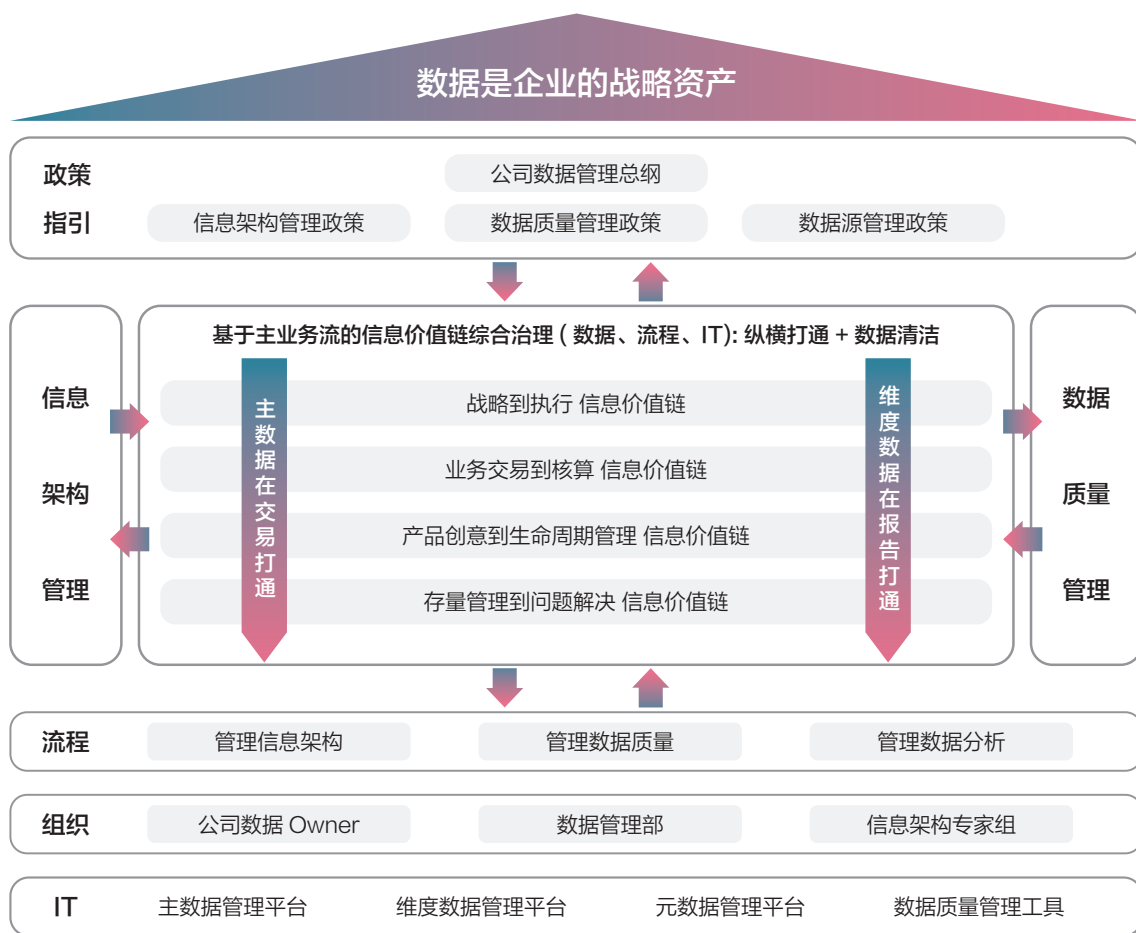


图 2-7 数据之道的延续

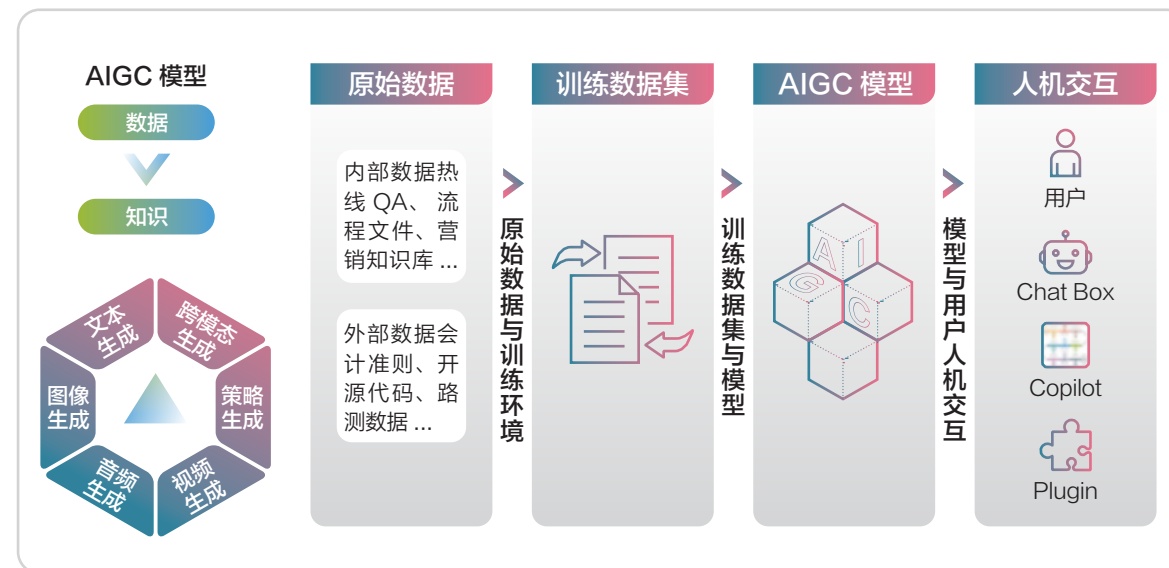


图 2-8 AIGC 治理结构变革

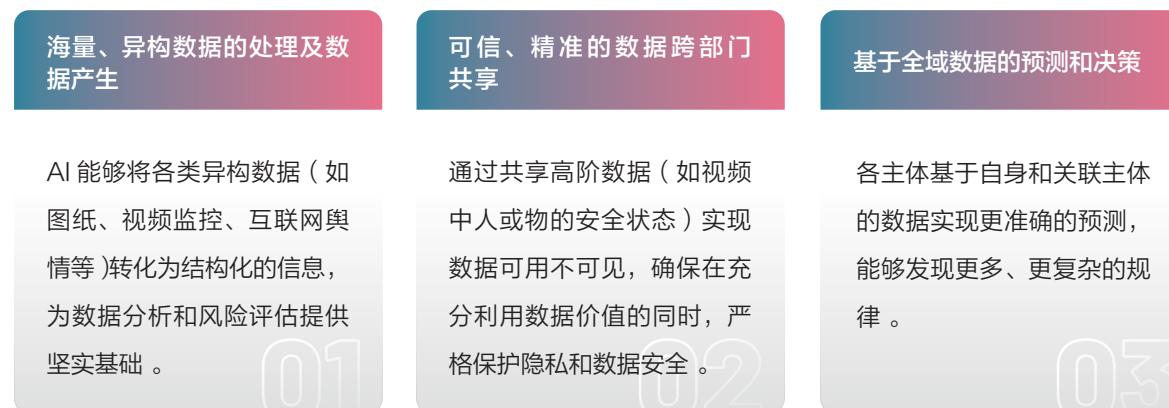


图 2-9 数据价值最大化

算力 算力底座化、平台化，选择战略同行者

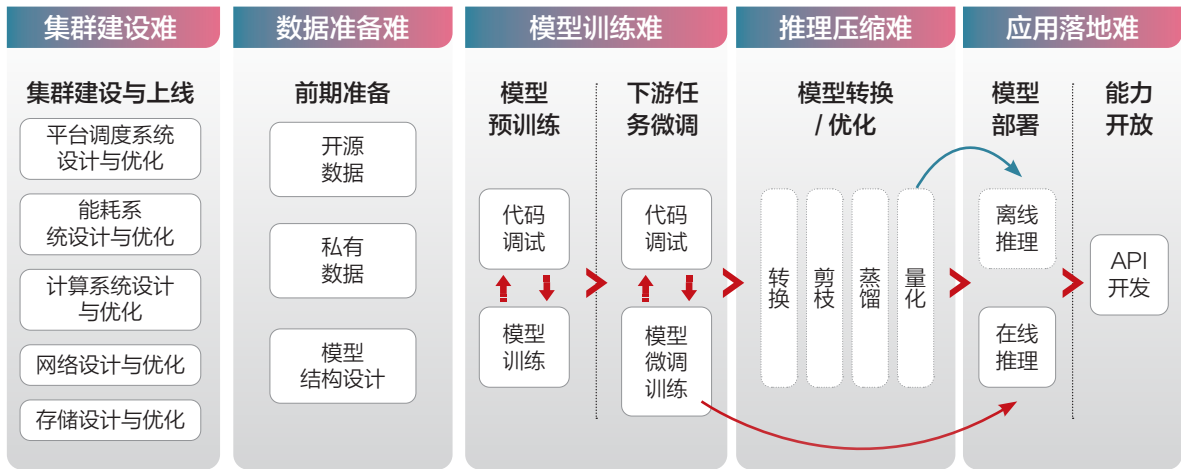


图 2-10 大模型开发的工程难题

大模型的开发与应用是一个复杂的系统工程，需要高度集成、内部硬软件高度耦合、外部提供标准化的接口的 AI 算力平台来支撑，重点解决集群建设、模型训练、推理压缩、应用落地中的问题：

- **集群建设**：如何实现超大集群的高性能长稳运行？如何构建参数面的无损网络？
- **模型训练**：如何选择最高效的并行组合策略？如何实现多任务可视化调优？如何实现断点续训？如何预测大模型的扩展性和性能？
- **推理压缩**：如何实现分布式推理和推理加速？如何进行大模型的无损量化？
- **应用落地**：如何搭建大规模推理集群调度系统？如何进行防攻击设计？如何有效的进行故障恢复和隔离？

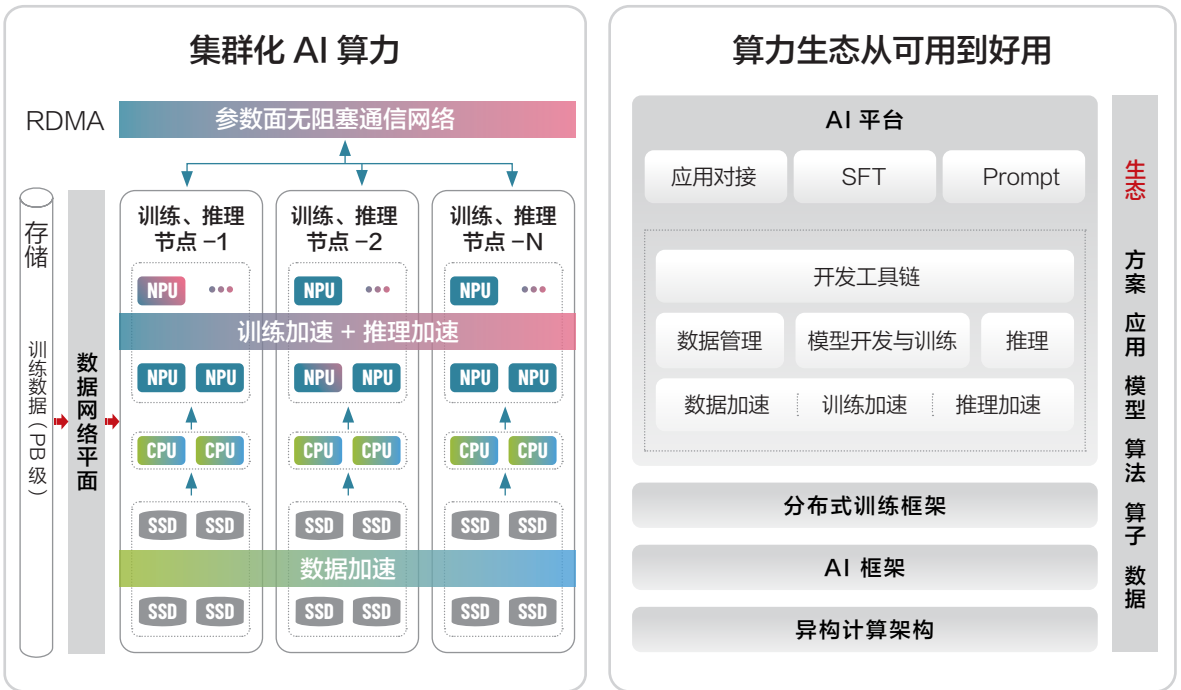


图 2-11 算力底座组件间高耦合和繁荣生态依赖

算力的选择也是技术路线的选择。 AI 算力供应链的可持续是路线选择的前提，不仅要考虑硬件的长期可获得性，还要考虑软件栈的可持续性。大模型训练与推理过程中，模型参数规模从数十亿到万亿，这不仅要求算力平台具备强大的并行计算能力，也要求算子（执行基本运算任务的软件模块）具备高效的执行效率，从而最大程度释放硬件计算、内存访问、卡间通

信的潜力。比如，华为 NPU 针对 AI 负载的矩阵计算框架进行了专门设计，更适用于卷积神经网络等类型的模型加速。值得注意的是，AI 算力芯片的支持不仅仅是硬件层面的问题，还需要有相应的开发者生态作为支撑，包括开发工具链、软件库、框架支持以及开发者社区等。最后，算力路线选择需要兼顾训练推理的需求，从调度效率、开发效率等多维度考虑。

| 以数据中心为中心

信息时代，网络是主角，联接企业 IT 系统及万物；
数字时代，云是主角，使能敏捷的应用开发；进入智
能时代，算力成为主角，作为提供算力的数据中心，

其效率成为企业 AI 效能的基础。**数据中心不再是单
纯的成本中心，而是创新中心。**

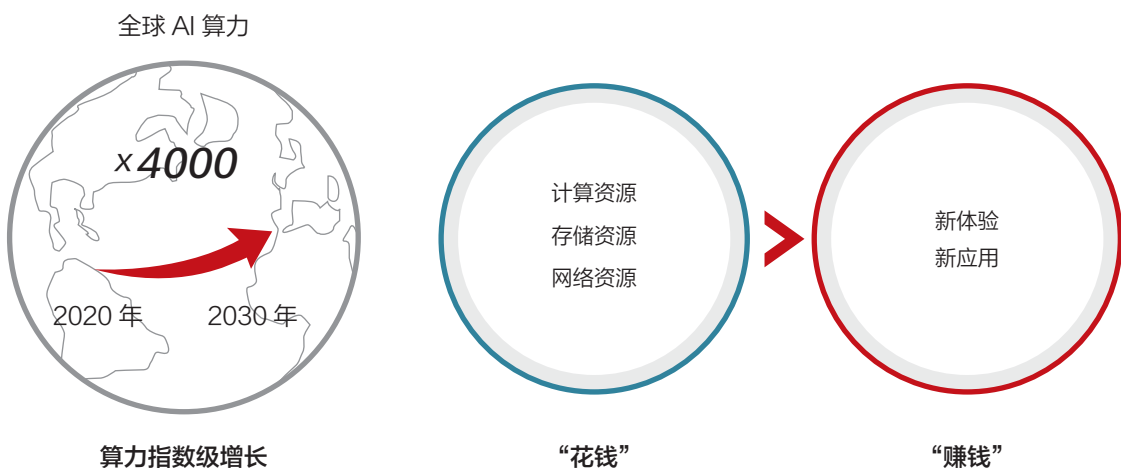


图 2-12 从成本中心到创新中心

数据中心的规模性、算力效率和开发效率成为企业 AI 的核心竞争力。当参数规模和数据规模越来越大，在算力供给受限和投资约束的情况下，数据中心的规模性、集群的有效算力、节能水平等成为企业模型开发和 AI 应用落地的关键因素。当企业发展 AI 时，预期价值闭环不是个别杀手级模型的低频次推理，而是

在海量、重复、复杂的场景中，数以百计的场景模型的高频使用。当一次普通的交互需要背后百亿次的运算时，数据中心效能的重要性显而易见。大模型的训练和推理成为最复杂的 IT 工程，数据中心正在成为企业数智基础设施的核心，成为企业 AI 商业价值闭环“投资收益不等式”中的重要系数。



图 2-13 企业级 AI 架构

数据中心将被 AI 重新定义，提供多样性澎湃算力、使能百模千态和 AI 原生应用创新成为愿景目标。算力类型不再被机房基础设施限定、集群规模不再被通信网络限定、任务可以低约束地调度、算力资源可以跨数据中心共享，使算力跟上大模型扩展的步幅；支

持开放的模型生态，针对不同业务场景，提供灵活的模型挑选与组合服务，确保每项任务都能匹配到最适配的算法模型组合；基于 Agent 的任务设计模式，融合企业和行业的知识资产、数据资产和模型资产，实现场景化的组合编排。

「第 3 章」

智能时代 数据中心的发展 与变化

数据中心走向智算数据中心 AI DC

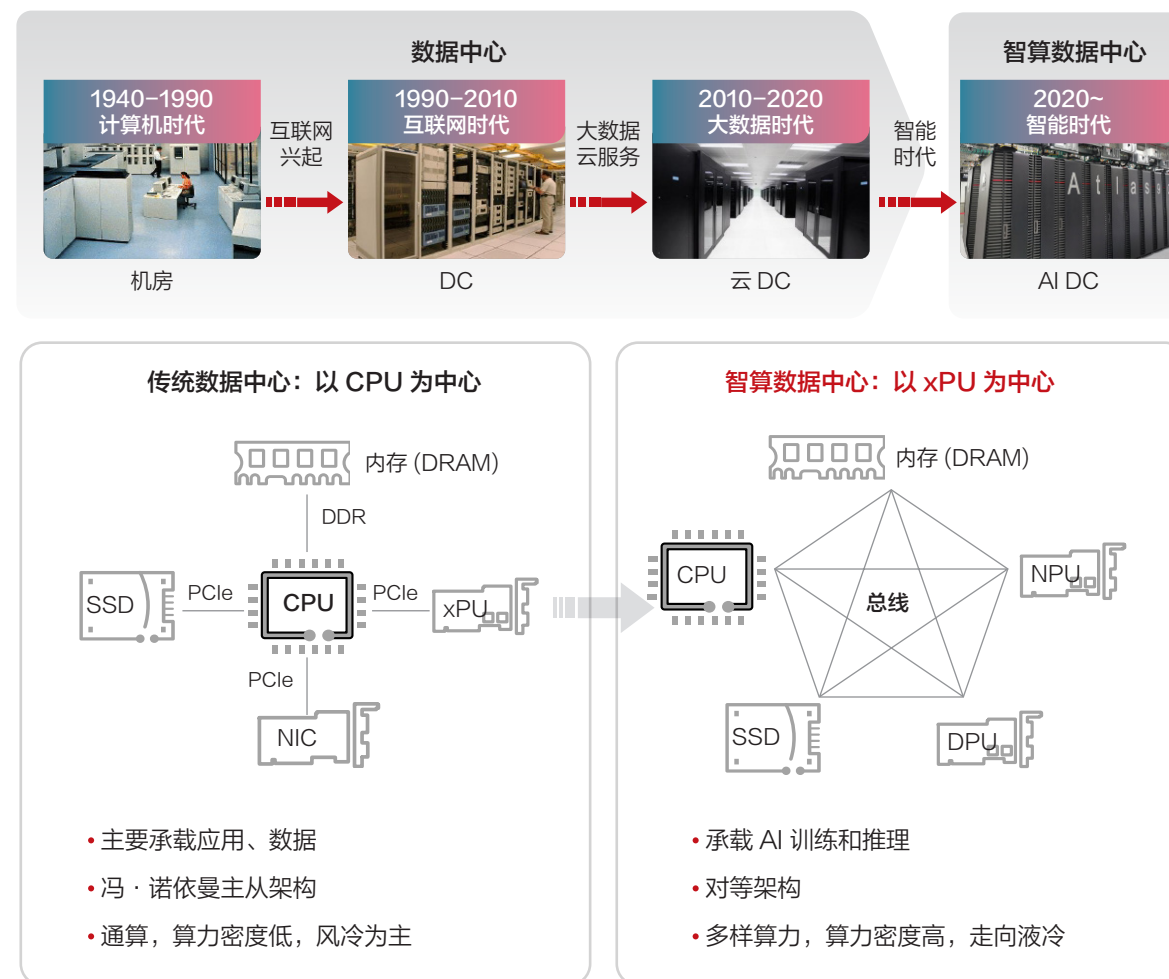


图 3-1 数据中心走向智算数据中心 AI DC

回顾过去几十年的发展历程，数据中心正走向智算数据中心。

随着互联网的兴起，数据中心作为 IT 基础设施的核心载体开始规模出现。从 2010 年开始，随着大数据和云服务的迅猛发展，数据中心的架构随之发生变革。云计算模式的兴起使得数据中心变得更加灵活和

高效，能够按需提供计算资源和服务。到了 2020 年，人工智能的快速发展加速智能时代的到来，对算力的需求爆发式增长。智算数据中心应运而生，专注于提供 AI 模型训练和推理所需的高性能计算能力。如谷歌建设的机器学习中心，Meta 打造的 AI 超级计算机，深圳专为深度学习设计的鹏城云脑 II 超级计算平台。

传统数据中心与智算数据中心存在以下几方面的差异：

承载业务差异

- **传统 DC：** 主要承载**企业级应用和数据存储**，如 Web 服务、数据库管理和文件存储等常规信息处理任务。
- **AI DC：** 主要承载 **AI 模型的训练与推理**，高效提供算力资源，并支持大数据集的处理。

技术架构差异

- **传统 DC：** 采用**冯·诺依曼的主从架构**，其中 CPU 扮演指挥官的角色，负责分配任务给其他部件。这种架构在面对大规模并行计算任务时存在“计算墙”、“内存墙”和“I/O 墙”等问题，限制了性能的进一步提升。
- **AI DC：** 采用更加先进的**全互联对等架构**，允许处理器之间，以及处理器到内存、网卡等直接通信，减少了中心化控制带来的延迟，突破主从架构的算力瓶颈，实现了高效的分布式并行计算。

算力类型差异

- **传统 DC：** 以 **CPU 为中心**，适用于一般性的计算需求。
- **AI DC：** 以 **xPU 为中心**，提供并行计算，处理 AI 模型训练所需的大量矩阵运算。

散热模式差异

- **传统 DC：** 单机柜功率密度通常在 3~8 千瓦之间，可装载的服务器设备数量有限，算力密度相对较低，一般采用传统的风冷散热。
- **AI DC：** 单机柜功率密度通常在 20~100 千瓦之间，主要采用液冷或风液混合的散热技术。液冷能够更有效地带走热量，保证高性能计算设备的稳定运行。

AI DC 主要承载 AI 模型的“训推用”

AI DC 最主要的是要围绕 AI 模型训练、推理和应用来规划设计和实施。



图 3-2 典型大模型应用之旅

AI 模型分为基础模型、行业模型以及场景模型。其中，基础模型具备广泛的适用性，能够在多种任务上表现出色；行业模型在特定行业背景下进行优化，深入地理解该领域的专业术语和业务流程；场景模型针对具体的业务场景或问题进行定制化设计，精确地解决特定任务的需求，全面提升模型的专业化水平和服务能力。



图 3-3 不同场景训练推理的算力需求及工程难度

大型互联网企业和专注于大模型训练的模型公司，其 AI DC 规划建设目标明确，即支撑基础模型预训练。这是一项大工程，需要超大规模集群的算力平台支持，还需要收集和处理万亿级别的 Token 数据，以确保模型能够学习足够的知识和技能。这种规模的训练不仅仅是技术上的挑战，更是对资源调配和系统运维管理能力的巨大考验。

行业头部企业在 AI DC 规划时，重点是行业模型的二次训练。行业模型是基于基础模型，通过叠加大量特定行业数据进行增量训练而产生的。相比基础模型的训练，复杂程度有所降低，但仍需要数百到数千张 NPU/GPU 的算力卡支持，并需要处理数亿级 Token 数据量。

AI 模型的全面应用，是从训练到推理多环节紧密协作的过程。这个过程包括基础模型预训练、行业或企业模型的二次训练以及场景模型的微调，最终实现模型在实际环境中的部署与推理应用。每一步都对数据中心的技术能力和资源管理提出全新挑战。

对于多数企业而言，AI DC 的建设重点在于承载 AI 模型的微调、推理及应用。鉴于 AI 应用的高度场景化特性，企业通常需要基于行业模型或基础模型，结合自身特有的场景化数据进行进一步的微调，从而使模型具备特定场景下的理解和生成能力，进而达到在实际业务环境中部署应用的标准。AI 推理的关键指标包括延迟（Latency）、准确性（Accuracy）、并发处理能力（Concurrency）以及算力使用效率（Efficiency）。根据推理服务的目标用户数量，如面向广大个人消费者的 2C 服务、面向众多企业的 2B 服务或是仅限企业内部使用的应用，AI DC 的规划建设标准和技术要求也会有所不同。

AI DC 四大建设场景及三大类型

根据不同需求，企业规划建设 AI DC 主要涵盖四大典型场景及用途。

场景 1：全量预训练

头部互联网公司、通信运营商及大模型厂商等，在建设超大型 AI DC，不仅用于基础模型的训练，还承担面向海量消费者用户的推理业务。

场景 2：二次训练 + 中心推理

金融、电力等国计民生的重要行业头部企业，正在积极推进大型 AI DC 建设，用于行业模型的二次训练及中心推理业务。

场景 3：二次训练 + 边缘推理

在一些集团化运营的企业中，其总部通常会建立大型 AI DC 来进行二次训练及中心推理，与此同时，在各个分支机构或靠近生产的地方，也会设置小型 AI DC 用于边缘推理及微调，从而构成了与企业整体组织结构相匹配的中心 + 边缘相互协同的架构，这种架构不仅能够充分利用资源，还能够实现实时决策，增强响应速度。

场景 4：轻量化推理

对于某些特定领域企业，尽管 AI 应用规模不大，但考虑到数据安全性和隐私保护的重要性，这些机构通常选择自建小型 AI DC，用于轻量化的推理任务及模型微调。例如，某三甲医院利用 AI 技术进行医学影像分析，帮助医生更快速准确地诊断疾病，同时确保患者数据不出医院内部网络，增强了数据的安全性。



图 3-4 AI DC 建设场景及类型

综上所述，业界典型的 AI DC 主要有三大类：超大型 AI DC、大型 AI DC 以及小型 AI DC。

一、超大型 AI DC

超大型 AI DC 主要承担基础模型预训练，面临以下主要挑战：

01 电力供应

超大规模 AI DC 的耗电量极为惊人。例如，一个拥有 10 万张智算卡的超大型 AI DC，其核心 IT 设备的电力需求超过 1 亿瓦（100MW），相当于 7.5 万户普通美国家庭的用电量，或是每小时熔化 150 多吨钢铁所需的电力。

02 可靠性与故障恢复

超大规模集群由成百上千万的器件构成，大模型的训练一般需要集群上百天 7x24 小时满负荷运转，导致光模块、NPU/GPU、HBM 内存等器件极易发生故障，而训练的同步性质使其对故障的容忍度较低，任何单点故障都可能导致训练任务中断，造成巨大经济损失。业界超万卡集群持续稳定运行仅数天，例如，Meta 在其 16K 集群训练 Llama3 405B 模型时，54 天内发生了 466 次作业中断。故障恢复常常需要数小时乃至数天，严重影响了训练效率。

03 有效算力提升

随着 AI DC 计算、存储和网络设备的规模不断扩大，如何高效地整合这些资源以实现算力的最大化，成为了业界研究热点。首先，要实现大规模设备的有效互联，就需要解决网络架构、通信协议以及数据传输效率等多个方面的问题。这要求在网络设计上更加注重可扩展性、灵活性和可靠性，以确保设备之间能够高效、稳定地进行数据传输和通信。其次，简单的设备堆叠并不能实现算力的线性增长，需要采用更加智能化的调度和管理策略，实现集群内计算、存储和网络资源之间的紧密协同。从当前业界的数据来看，即使是业界顶尖的千卡智算集群，其算力利用率不超过 60%，万卡集群不超过 55%，而十万卡集群更低，不超过 40%，这进一步说明了提高超大规模集群有效算力的重要性和紧迫性。

为应对上述挑战，业界领先的**超大型 AI DC**需要具备**极致能效和极致算效**的能力。

二、大型 AI DC

大型 AI DC 通常由行业头部企业规划建设，既要承担多种模型的训练及微调任务，又要承担较大规模的中心推理以及 AI 应用，面临以下主要挑战：

01 推理性能优化

在确定的业务场景和确定的时延下，如何提供极致的推理性能。

02 高效的训练和微调

一方面可帮助企业更快的将智能应用部署到实际生产环境中，缩短开发周期，从而在竞争激烈的市场中保持领先优势；另一方面可以节省成本和资源。

03 提升算力资源利用率

建一个大型的 AI DC，企业往往需巨额的资金投入，因此希望这些“宝贵”的 AI 算力资源尽可能多的利用起来，避免算力资源的闲置。

04 多模编排快速支撑 AI 应用创新

当今企业应用创新的步伐不断加速，如何将多个模型灵活组合编排来满足应用快速创新的需求。

05 降低 AI DC 运维难度

大型 AI DC 往往是企业来承担运维管理工作，如何能快速定位故障、修复故障是多数企业运维人员的共同诉求。

06 如何应对生成式 AI 安全

对于金融、政府、电力等国计民生行业，某些场景有严格的 AI 输出要求，需要确保生成式 AI 输出的内容是正确合规的。

07 高密供电、液冷散热等机房条件是否具备

智算需要超过通算 10 倍的功耗、10 倍的布线规模，并且越来越趋向液冷散热，需要企业提前做好机房的规划准备工作，避免成为大型 AI DC 建设使用的瓶颈。

综上所述，最终能够成功应对上述挑战的**大型 AI DC**，一般需要具备**融合、高效的关键特征**，以适应企业未来发展的需求。

三、小型 AI DC

小型 AI DC 主要承担轻量级的推理及 AI 业务应用，有些还要求提供模型微调能力，一般建在贴近生产或靠近实际用户的地方，其建设面临的主要挑战是：

01

提升算力资源利用率
小型 AI DC 受环境限制，所能提供的算力资源比较有限，因此必须要求在这有限的资源条件下，尽可能部署更多的业务应用。

02

一站式部署
有些小型 AI DC 的位置相对较偏，甚至远离城区，这种情况下企业往往希望提供一站式的部署，交付人员最好只跑一趟就能完成 AI DC 的部署。

03

便捷运维
对于小型 AI DC，企业一般配备较少、甚至没有专门的运维人员来支撑，这就要求小型 AI DC 故障要少、日常运维极简，在出故障时，能够提供远程运维操作。

04

安全保障
小型 AI DC 贴近生产，往往需要和智能摄像头、传感器等感知终端直接连接，而这些暴露在外端的终端也极易出现安全入侵问题，这种情况下如何确保小型 AI DC 的安全，也是必须解决的一个问题。

综上所述，最终能够成功应对上述挑战的小型 AI DC，需要具备“形态灵活、快速部署快速升级、轻量极简、易维易用”等特征。

AI DC 五大特征变化

从技术角度审视，应对各类 AI DC 所面临的挑战，构建领先的 AI DC，需要在五大关键技术领域实现重大突破与革新。



图 3-5 典型 AI DC 的关键挑战及技术方向

系统摩尔

算力大小决定了模型能力上限。当前，大模型的能力上限尚未触及，Scaling Law 尺寸定律依然有效。预计到 2028 年，模型参数将达到数百万亿~数千万亿，如此大规模的模型训练需要算力规模和能力的进一步突破，而当前主导算力发展的传统通算摩尔定律正遭

遇物理学和经济学双重限制，致使传统的硅基电子技术临近发展极限，**算力增长速度远远慢于算力需求的增长速度，算力裂谷越来越大**，业界迫切需要新的算力供给方案，我们称之为“系统摩尔”。

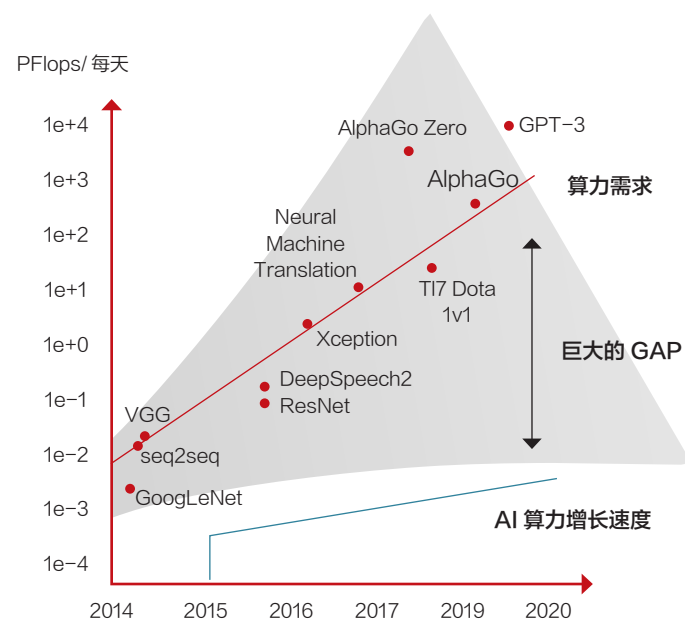


图 3-6 智能时代加速而来，算力裂谷越来越大

系统摩尔是华为最初在《数据中心 2030》报告中提出的概念，它定义为一种新的算力提升方法，**主要依赖系统级架构创新、算网深度协同、软硬深度协同来提升算力，满足指数级增长的算力需求。**

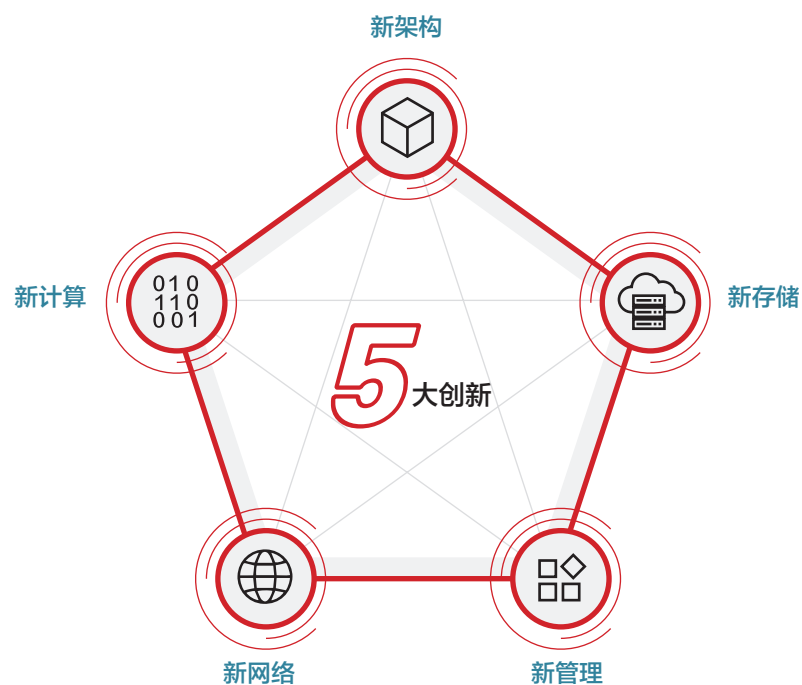


图 3-7 围绕系统摩尔的 5 大创新

具备系统摩尔特征的 AI DC 算力供给方案，呈现出 5 大新特点：

一、新架构

过去 70 年，计算机一直遵循冯·诺依曼架构设计，运行时数据需要在处理器和内存之间来回传输。在人工智能等高并发计算场景中，这种传输方式会产生巨大的通信延迟，从而导致“通信墙”；而且目前内存系统的性能提升速度大幅落后于处理器的性能提升速度，有限的内存带宽无法保证数据高速传输，带来了“内存墙”。在此背景下，**全互联的对等计算架构**应运而生，它能够让 NPU、DPU、CPU、内存以及其他异构芯片之间实现高效的数据交换，打破传统计算架构的“通信墙”和“内存墙”等瓶颈，支持 AI 等场景对跨主机高带宽、低时延的诉求，实现 DC as a Computer，算网存深度协同，通过系统级架构创新，充分释放算力效能。

二、新计算

“新计算”主要体现在两个重要方面：

首先，计算类型的演变。从以 CPU 为中心的通用计算，向以 GPU 和 NPU 为中心的智能计算转变。这种转变不仅适应了 AI 算法所需的大量并行处理能力，还大幅提升了计算效率和灵活性。并行计算技术，如同千军万马并驾齐驱，能够同时处理多个计算任务或数据块，极大加速了数据处理和计算过程，提高了计算资源的利用率和整体计算效率。通过并行计算，不仅能够缩短计算时间，还在更短的时间内完成更复杂的计算任务，从而更有力地推动了人工智能领域的发展。

其次，芯片技术的进步。首先是 Chiplet 技术，不仅可以显著提高 Die 的良率，还能有效地降低成本，并

且这种方法可以根据不同的产品规格需求灵活调整，实现更高水平的芯片性能。此外，与传统的封装板级互连方案相比，2.5D 封装技术能够将每比特的能耗降低大约一半，从而进一步提升了系统的能效比。

三、新存储

在“新存储”领域，随着大模型的广泛应用，对高性能存储的需求日益凸显。特别是在 AI 训练过程中，高效的数据读写成为了提升整体训练效率的关键因素。

在训练阶段，需要从存储系统快速加载样本数据到 NPU/GPU，并定期将 Checkpoint 数据从 NPU/GPU 写回到存储系统中保存。因此，提升存储 I/O 性能，缩短数据读写时间，成为了提高训练效率的重要手段之一。为此，**NPU/GPU 直通存储技术**应运而生。这种技术为 NPU/GPU 与存储之间提供了一条直接的内存访问传输路径，消除了原先涉及的 CPU 内存缓冲和复制过程，从而大幅缩短了数据读写的时间。

在推理阶段，尤其是在面对高并发、长序列的推理场景时，业界提出了以 KVCache（键值缓存）为中心的**多级缓存加速技术**。这一技术能够显著提升大规模推理系统的吞吐性能，通过优化数据访问路径，确保数据能够快速、高效地被处理。

总之，无论是训练过程中的 **NPU/GPU 直通存储技术**，还是推理过程中的 **KVCache 多级缓存加速技术**，都是为了在大数据量和高并发场景下，提升系统的整体性能和响应速度，从而更好地满足大模型应用的需求。

四、新网络

网络作为连接计算和存储的关键纽带，在满足大规模计算集群的连接需求方面，正迅速向十万乃至数十万 xPU（如 GPU、NPU 等）的互联演进。随着网络技术的发展，参数面网络的接入速率已从 200GE 提升至 400GE 乃至 800GE。

大模型本身也在不断发展，从早期的张量并行、数据并行和流水线并行等分割方式，快速演进到 MOE（Mixture of Experts，专家混合）等更高级别的并行方法。这一演进对网络级负载均衡技术提出了更高的要求。为应对这一挑战，各大厂商纷纷推出各自的负载均衡解决方案。例如，华为推出了与昇腾平台配套的动态 NSLB（全局负载均衡）技术。据测试结果显示，在 512 卡规模内，该技术能够提升 Llama2 13b 模型 13% 的训练效率。

总之，随着网络技术的不断进步和大模型的演进，网络架构和负载均衡技术也在不断创新，以满足更高性能和更大规模的算力需求。

五、新管理

新的管理模式必须具备跨域协同管理的端到端系统运维能力，涵盖计算、存储、网络、光模块设备的管理、控制以及分析等全生命周期运维管理。具体包括以下几个方面：

- **全链路可视化监控**：通过实时监控整个系统的运行状态，实现对计算、存储、网络等资源的全面监控，确保任何异常都能及时发现。
- **跨域故障快速定位**：利用先进的故障检测技术，快速准确定位故障点，减少故障排查时间，避免训练任务中断。
- **跨域故障快速修复**：建立高效的故障修复机制，确保一旦发生故障，能够迅速采取措施恢复系统正常运行，减少停机时间。

通过这些措施，可以显著提升训练效率、降低训练成本，并确保大模型训练的快速、稳定 and 高质量完成。这种全方位的系统运维管理能力是未来大型乃至超大型 AI DC 的核心竞争力所在。

能基本桶

AI DC 算力密度增长带来功率密度的急剧攀升，给供电、散热及布局等带来极大挑战，正在重塑数据中心能源基础设施。

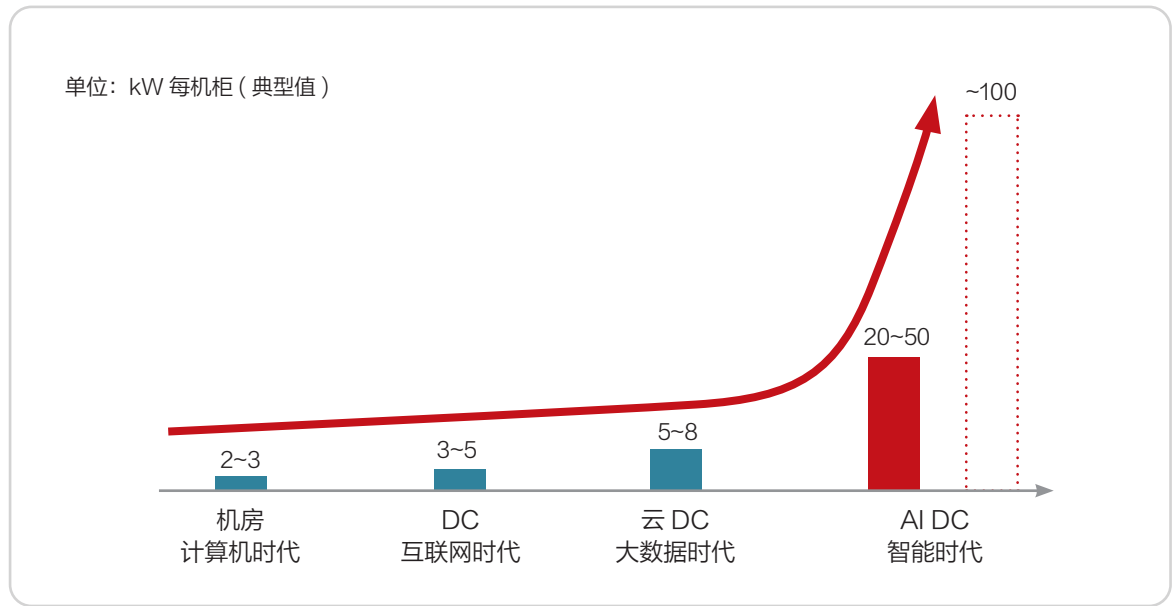


图 3-8 不同时代数据中心机柜的典型功率

挑战一

超大容量电力供应的获取与匹配

随着数据中心用电量的飙升，尤其是当单个数据中心用电量跃升至 200MW 乃至 500MW 以上时，城市现有电力基础设施的瓶颈日益凸显。如 OpenAI 的“星际之门”项目所预计的高达数千兆瓦的电力需求，已迫使数据中心选择跨越地域界限的电力供给解决方案。因此，如何高效、稳定地获取并匹配如此庞大的电力资源，成为了制约算力规模进一步提升的首要难题。

挑战二

超高密机柜的散热技术创新

高功率密度带来的不仅仅是电力挑战，更对散热技术提出了严苛要求。液冷技术虽已成为行业共识，但面对未来更高功率密度的挑战，如何在确保可靠性和易维护的同时，提升散热效率，仍是亟待解决的关键问题。

建筑空间分区的前瞻设计

AI DC 的设计需兼顾 IT 机房、制冷设施与电力供应区域的复杂需求，打破传统设计模式，采用更为前瞻性的布局思路。这包括降低 IT 设施与机电设施的耦合度、实现机电设施的模块化与室外化布置、以及结合风冷与液冷技术的弹性配比设计。

为避免能源基础设施成为数据中心发展的瓶颈，并减少由此产生的成本和资源浪费，需采取以下措施：

- **优化数据中心布局：**通过科学合理的规划与设计，确保电力供应、冷却系统与算力需求之间的高效协同，提升整体能效。
- **提升能源使用效率：**采用先进的节能技术与管理手段，降低能耗水平，实现绿色算力的发展目标。
- **发展可再生能源与储能技术：**积极利用太阳能、风能等可再生能源资源，并配套建设储能设施，提升数据中心的电力供给能力与抗风险能力。
- **升级供电与制冷设备：**紧跟技术发展步伐，不断引入更高效、更可靠的供电与制冷设备，提升数据中心的运行效率与稳定性。

面对 AI DC 的能源基础设施挑战，需以创新的思维与前瞻的视角，积极探索并实践上述应对策略，在保障算力供给的同时，实现可持续发展与绿色转型的目标。

迭代式平台

相比于传统 DC，AI DC 规模更大、业务更为复杂且技术更新更快。因此，提供资源管理调度、支撑模型训练及 AI 开发，以及提供运维管理的 AI 平台面临极大的挑战，主要包括：

01

- **AI 算力资源的高效利用：**AI 服务器采购价格是传统通算服务器的数倍，再加上 AI 对网络和存储设备提出了更高要求，使得 AI DC 建设成本高昂。这种情况下，如何管好、用好 AI 算力资源，让单位算力产出更大，就成了企业用户普遍关心的问题。

02

- **AI 开发的高门槛和高成本：**传统 AI 模型的泛化能力较差，面对不同的用户或数据源时，性能容易下降。缺少算法专家的企业难以完成模型的调试和优化，而即便大模型的泛化能力有所改进，但面对广泛的应用需求，算法专家的数量仍然不足，这就导致了 AI 应用开发成本高，开发周期长的问题，阻碍了 AI 技术全面服务于企业业务的各个领域。此外，模型维护也是一个持续性的挑战。

03

- **AI DC 运维运营难度大：**AI DC 作为一种新型的数据中心，缺乏具备管理大规模 AI 服务器，以及高性能网络和存储设备经验的运维人员，他们面临的问题包括合理的资源分配、变更管理、故障快速定位及恢复等。要解决这些问题，不仅需要运维人员个人能力提升，还需要有完善的运维运营工具来支撑。

为了应对上述挑战，需要一个能够持续迭代的 AI 平台，不断整合新技术和架构，以成熟的方式提供给用户，朝着**性能更强、效率更高、运维更简、功能更全**的方向发展。

性能更强

优秀的 AI 平台应当持续引入这些技术，帮助用户提升性能并降低成本。数据并行、网络优化等技术有助于提高训练效率；量化压缩则提升了推理效率；PD 分离技术增强了长序列输出的性能；提示工程优化则能低成本地提升推理准确率。

运维更简

大规模 NPU/GPU 和光模块使 AI 集群运维复杂化。新一代运维系统应具备全面监控、故障预测、智能分析等功能，提升硬件的无故障运行时间和集群效率。在推理环节，运维系统需监控硬件利用率等关键指标，识别低效作业并协助优化，以持续改进集群性能。

效率更高

由于 AI 硬件成本高昂，提升算力集群利用率至关重要。通过优化存储方案和通信算法，可以克服并行训练中的瓶颈，提高数据传输效率，缩短训练时间。对于以交互为主的推理应用，平台应支持动态调度，如 API、定时及按负载扩缩容，以释放闲置资源。夜间空闲资源可用于微调训练，另外，平台还需提供安全隔离和灵活调度支持，确保业务连续性和资源的有效利用。

功能更全

大模型应用开发已有多种模式，如 RAG 和 Agent。AI 平台应提供相应的支持工具，比如数据工程模块简化数据预处理，模型开发模块降低训练门槛，Agent 开发模块则简化服务构建流程，共同提升开发效率并降低门槛。

总之，未来的 AI 平台应通过不断的迭代升级，提供更强大的性能、更高的效率、更简单的运维以及更全面的功能，以更好地支撑企业的 AI 业务发展。

编排式应用

随着数字化进程的加速，许多领先企业已拥有从几十到数百个应用不等。在过去的一年多时间里，AI技术的快速发展推动了“所有行业、所有应用、所有软件都值得用 AI 重做一遍”的理念。与此同时，大模型的应用极大地改变了软件开发的方式，催生了一种新的编排式应用开发模式。面向未来，企业在智能化转型的过程中，将拥有成千上万的各种模型，如此庞大的模型库，导致未来企业必须通过编排式应用开发，才能快速响应企业的智能化改造需求，以促进业务创新。

编排式应用的构建与传统应用构建方式在构建主体、流程分解、实现形式以及处理形态等方面存在根本性的区别。在基于大模型的编排式应用构建中，业务工程师和系统工程师可以根据具体的业务逻辑，通过自然语言提示的方式引导大模型对业务流程进行分解规划。这种流程处理依据大模型的规划结果进行实施，其形态也从固定的静态流程转变为更具灵活性的动态流程。未来的应用构建方式将更多地依赖于业务人员而非专业的开发人员，编排式应用模式的转变使得业务人员乃至最终用户自主构建智能体（Agent）应用成为可能。

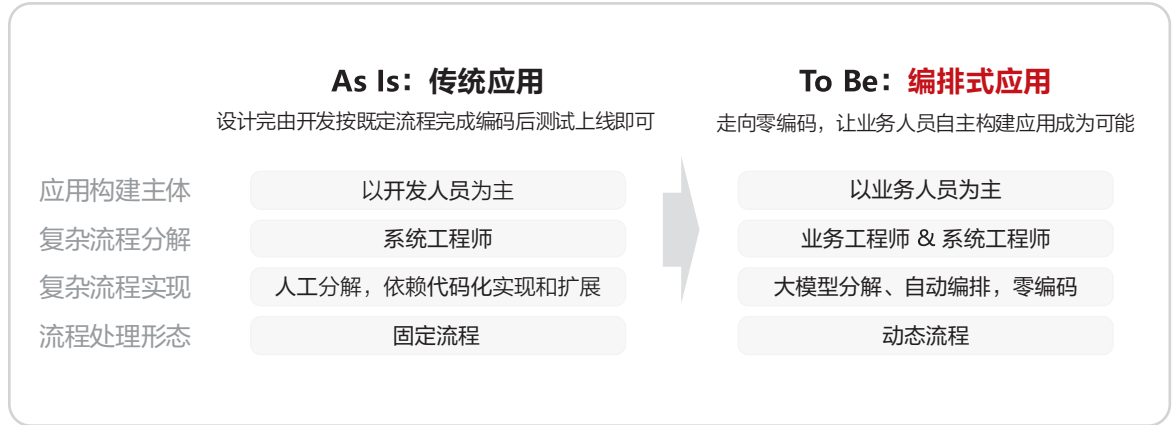


图 3-9 从传统应用到编排式应用

在编排式应用开发中，重要的是要充分利用大模型在理解和生成方面的能力，以及小模型在感知和执行上的专长，通过合理编排这两种模型，实现能力

互补，共同支撑应用的功能。通过对多个行业中实际 AI 应用案例的分析，我们总结了**四种主要的应用编排模式**：



图 3-10 四种应用编排模式

生成式安全

除了传统数据中心面临的安全风险，AI DC 还要面临新的安全挑战。一是 AI 内容生产过程的“黑盒”特性，导致其输出内容具有很大的不确定性和不可解释性，带来较大的应用风险，尤其是一些对输出内容要求比较严格的场景。二是 AI 系统面临新型安全攻击的威胁，大模型基于统计和语言规则的预测机制使得它很难区分是合法的指令还是恶意的输入，攻击者可以通过精心设计的提示词来操纵大模型，如在 2023 年中针对 ChatGPT 的“奶奶讲故事”漏洞，诱导 AI 执行本应禁止的操作。三是潜在引入新的数据安全风险，大模型在训练过程中可能会接触到大量

的用户数据，并加以记忆存储，而在推理阶段可能会无意泄露客户的隐私信息，如三星电子半导体员工在使用 ChatGPT 的过程中，无意中泄露了半导体设备测量资料和产品良率等敏感信息，竞争对手可通过 ChatGPT 问答来获取相关信息，对三星的市场地位和竞争力造成了极大的负面影响。为此，全球权威的 OWASP（Open Web Application Security Project）在线社区集合了全球 500+ 安全专家，在 2023 年 10 月提出了 LLM 应用的 10 大 TOP 威胁（1.1 版本）。

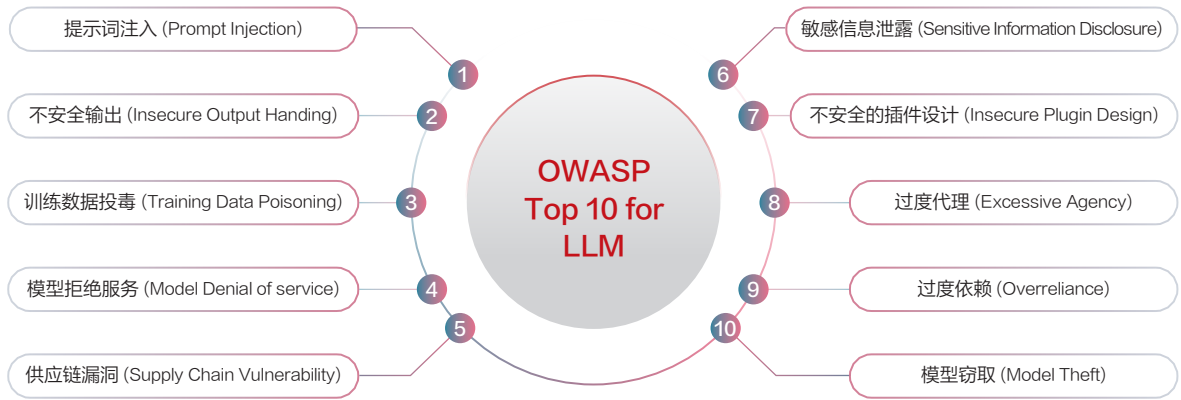


图 3-11 OWASP 发布的大语言模型 10 大安全风险

针对上述安全风险，需要构建立体、多元的系统性安全防御，从源头上控制风险，确保大模型安全做事。首先需要保证训练数据集的安全，重点加强数据版权保护，隐私合规，确保数据可追溯；其次在模型训练

阶段，要增强模型的内生安全能力，通过教会大模型各类安全知识，提升大模型自身的健壮性；最后通过构建大模型安全护栏，确保大模型从容应对各种安全攻击，保障输入输出内容合规。

数据中心将被重塑，由分层解耦到垂直整合

传统数据中心（DC）通常是按照能源设施层、IT 硬件设施层、平台软件层和应用软件层等进行分层解耦规划设计，并且按计算、存储、网络、云平台、数据库等部件分别采购建设。这种模式在通算时代是普遍存在的，但在 AI DC 上遇到了很大的挑战。

数据中心架构的变化

首先，数据中心的架构发生了根本性的变化。与传统 DC 的分层架构不同，AI DC 逐渐形成了新的分层架构，即以算力底座层、平台服务层、模型使能层和行业应用层为核心的新型 AI DC 目标架构。

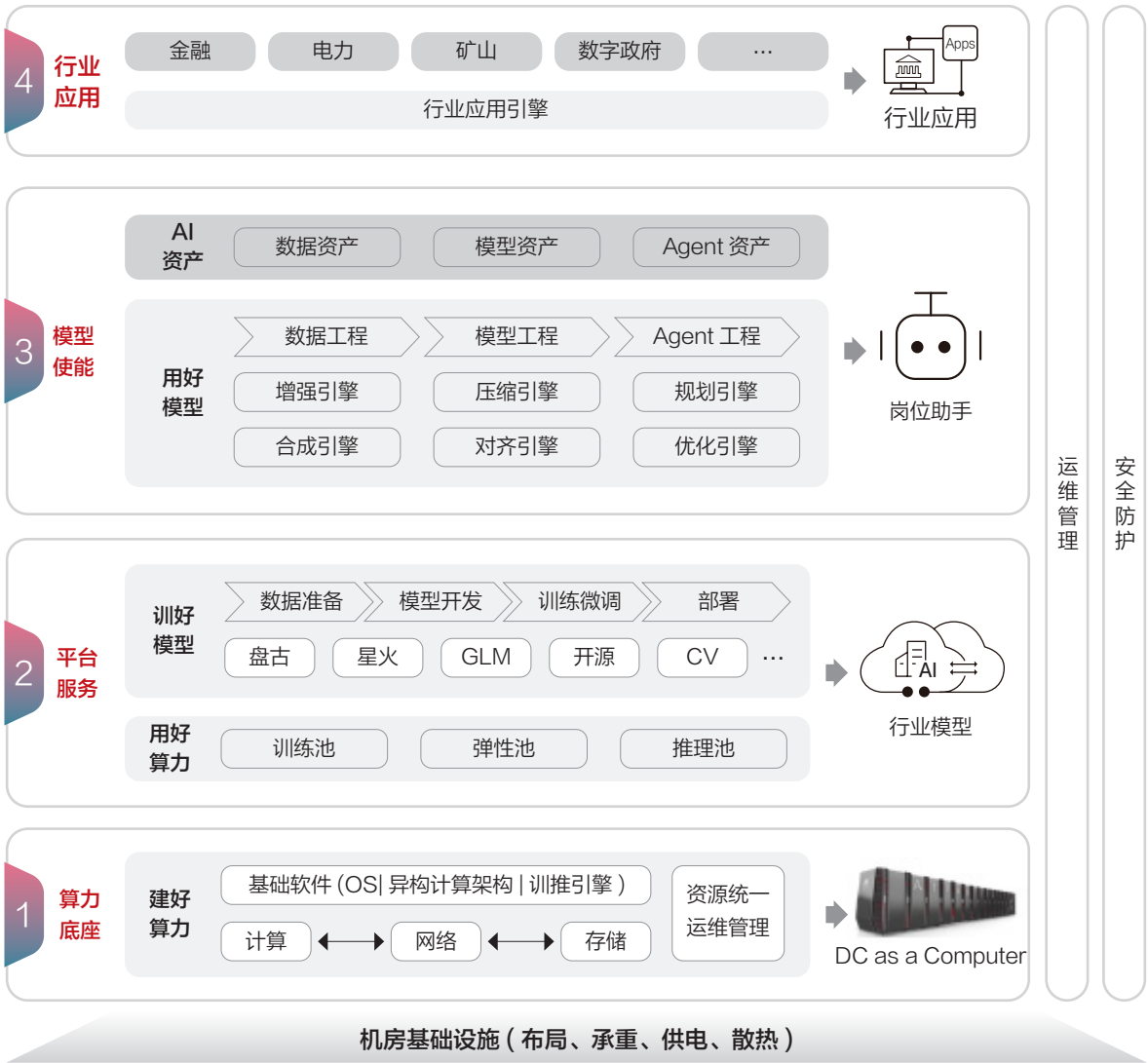


图 3-12 新型 AI DC 目标架构

垂直整合的需求

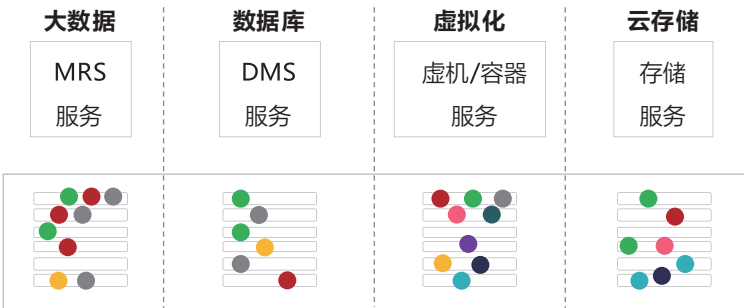
其次，AI DC 要求垂直整合，以满足高并行的智算业务需求。以算力底座层为例，传统的大数据、数据库、虚拟化等通算业务可以通过由计算、存储和网络组成的松耦合系统运行多个不同的任务。这些任务分布在不同的通算服务器上运行，且多数任务在单个服务器内即可闭环，节点之间是松耦合的，因此按部件采购建设的模式是可行的。

然而，模型训练等智算任务则不同，尤其是对于万亿甚至十万亿参数的大模型训练，需要整个算力底座来运行单一的训练任务。这些任务需要横跨整个算力底座的计算、存储及网络节点，且必须持续上百天、每天 24 小时不间断地高并行运转。节点之间必须紧密耦合并保持同步关联。任何一个节点出现故障，都将导致整个作业中断，需要重新启动，从而带来巨大的损失。因此，迫切需要将算力底座打造成为一个像超级计算机一样精密协同工作的系统，实现“DC as a Computer”，以保障智能计算业务的高效稳定运行。

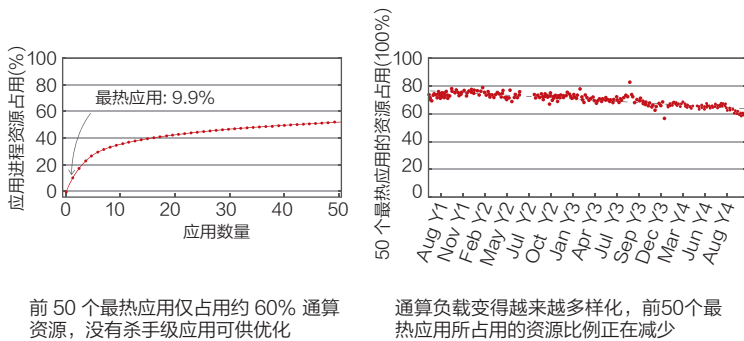
在这种情况下，继续沿用传统的按部件采购建设模式显然无法实现这一目标，必须采用算力、存储和网络部件的垂直整合模式。

通算：业务在单服务器内闭环

负载多样化、关联少，节点间松耦合



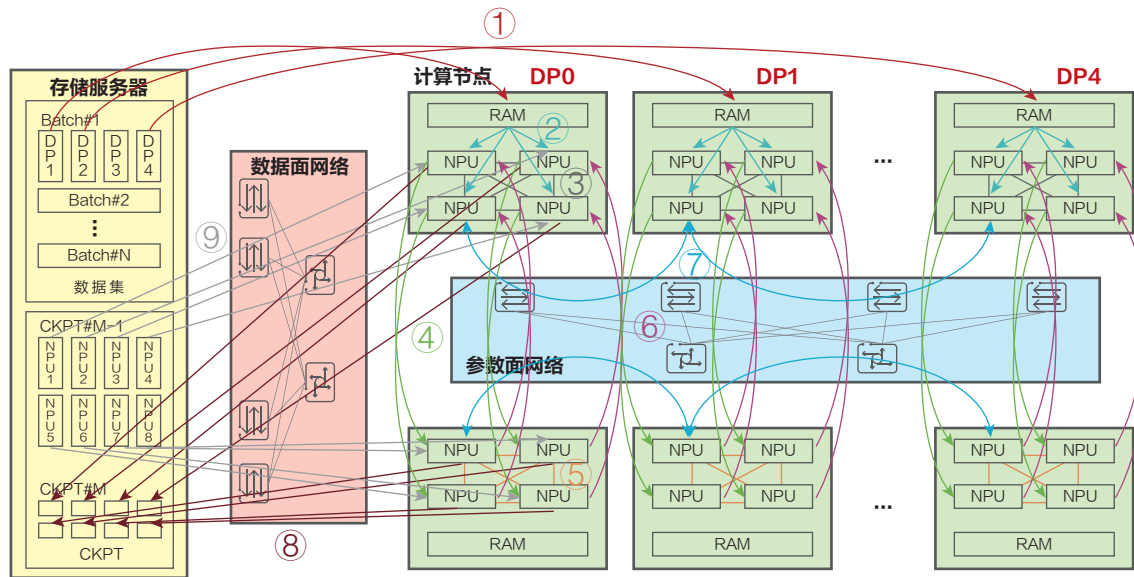
通算运行：N个负载运行在M个节点上，多数负载在单节点内闭环，松耦合



高并发

智算：业务在大系统内闭环

单一任务、长期满负荷，节点间紧耦合（DC as a Computer）



- ① 将当前Batch数据根据DP切分并传输到对应节点
- ② 数据加载到NPU上运算
- ③ 前向传播TP通信
- ④ 前向传播PP通信
- ⑤ 反向传播TP通信
- ⑥ 反向传播PP通信
- ⑦ DP通信，同步梯度数据
- ⑧ [周期性]保存CKPT
- ⑨ [故障时]恢复加载CKPT

智算运行：1个负载同步运行在全部节点上，同步关联协同，紧耦合

高并行

图 3-13 从通算的分层解耦到智算的垂直整合

综上所述，对于 AI DC 的建设来说，算力底座层和模型使能层之间是可以解耦的，模型使能层和业务应用层之间也可以解耦。但是，层内的关联是非常紧密的，需要垂直整合，才能提升有效算力，并实现算力、

存储和网络一体化的运维管理。通过这种垂直整合的模式，可以确保 AI DC 在面对高并行计算任务时，能够高效、稳定地运行，从而满足日益增长的智能计算需求。

「第 4 章」

典型 AI DC 规划与建设

超大型 AI DC

本节重点描述超大型 AI DC 建设过程中面临的关键需求与挑战，并总结超大型 AI DC 的关键特征，以及如何建设 AI DC 的方向性建议。

关键建设需求

承载基础大模型预训练和推理，超大型 AI DC 建设主要存在三个关键需求：

关键需求一

提升基础大模型预训练效率，缩短训练时长

当前，对于各头部互联网企业及大模型厂商而言，都希望预训练的周期越短越好，以实现基础大模型的快速迭代，从而赢得市场先机；与此同时，成千上万台 AI 服务器的长周期、高负荷运转需消耗大量电力。因此，提升训练效率、缩短训练时长，不仅能赢得市场竞争，也能实现节能降本。要提升训练效率，在确定的算力规模下，关键在于提高算力集群的有效算力。

关键需求二

满足推理的“LACE”体验要求

面向海量用户的推理业务中，重点要关注用户的“LACE”体验，即：Latency（响应时延）、Accuracy（响应准确性）、Concurrency（吞吐并发能力）和 Efficiency（算力使用效率）。



图 4-1 "LACE" 推理指标体系

● **Latency**: 时延直接影响用户体验，不同应用场景有不同的时延要求。例如，互联网应用场景时延要求不超过 30 毫秒，文本对话场景时延要求在 30 至 100 毫秒之间，语音对话场景要求时延在 100 至 200 毫秒，对于辅助编程或医疗诊断等时延不敏感的业务，时延要求可以大于 200 毫秒。

● **Accuracy**: 确保系统输出结果的精确性，以满足用户的需求和期望，特别是在那些依赖于准确信息的应用场景中。

● **Concurrency**: 互联网应用往往需要每天处理亿级的并发请求响应，这意味着系统必须具备强大的吞吐能力，以应对高峰时段的高并发需求。

● **Efficiency**: 推理集群的算力效率直接影响最终的成本控制。为了降低成本，需要尽可能提高推理集群的有效算力利用率。

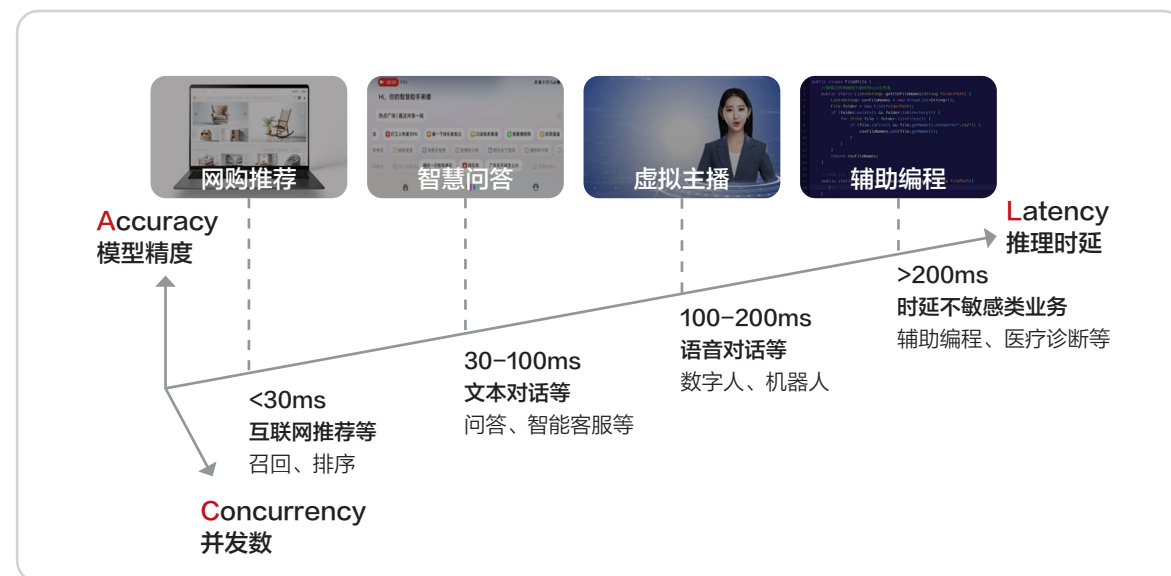


图 4-2 典型应用对推理性能的要求

关键需求三

提升能源基础设施效率，满足超大型 AI DC 可持续发展

为了满足超大型 AI DC 可持续发展的要求，能源基础设施需要实现高效率、高密度、高弹性、高可靠。



图 4-3 超大型 AI DC 对能源基础设施的要求

高效率: 超大型 AI DC 的巨大能耗使得提升能源效率成为必选项。如果将 PUE（电源使用效率）从 1.5 降至 1.15，以 10 万卡 100MW 的数据中心为例，每年可节约约 2 亿度。

高密度: 随着 AI 芯片功耗的不断上升，单机柜的功率密度提高了 5 到 10 倍。这就要求供电和散热系统也要相应提升密度，以支持更多的机柜部署。

高弹性: 技术更新的速度从每三年一代，到每一年一代，多种算力和多代算力的混合部署成为常态，能源基础设施需要具备更高的灵活性，支持功率可调和风

液配比可调，以确保在 10 到 15 年的生命周期内都能高效利用。

高可靠: 由于智算设备成本高昂，任何宕机都会造成重大损失，单点故障可能导致整个集群中断。因此，供电和散热系统必须具备高可靠性。此外，光模块等器件的故障率与机房温度密切相关，需要更为精准的温度控制来确保系统稳定运行。

综上所述，对于具有“超大规模、超高负荷、超高造价、超大耗电”等特点的超大型 AI DC 而言，其领先性主要体现在**极致算效和极致能效**。



图 4-4 领先的超大型 AI DC，要求极致算效和能效

规划建设方向一：极致算效

为了实现极致算效，需采用多种策略和技术手段：

场景一

基础模型预训练

为了加速基础模型的预训练，需要提升超大规模集群有效算力。集群有效算力由三个关键指标决定：集群算力规模、集群算力利用率（MFU）和集群可用度（HA）。

01

集群算力规模

即集群的裸算力规模，取决于每个节点的算力以及集群的规模。算力规模受限于单个节点的算力性能和集群中节点的数量。

02

算力利用率（MFU）

指的是计算设备实际执行计算任务的时间与理论计算能力的时间比例。高算力利用率意味着计算资源被充分利用，减少了空闲时间。

03

集群可用度（HA）

指的是集群处于可工作状态的时间比例。高可用度意味着集群能够在大多数时间内正常运行，减少停机时间和故障时间。



图 4-5 集群有效算力

如何构建超大规模算力集群、提升算力利用率和集群可用度，是当前业界关注的热点，需要采用如下关键技术：

关键技术 1 基于超节点及超大规模组网架构，提升集群算力规模

业界主要从 Scale up 和 Scale out 两个维度来实现算力规模的提升。

在 Scale up 方面，主要采用超节点技术来提升单位算力。超节点技术采用创新性的对等互联计算架构，通过高速互联总线将数百颗 AI 芯片进行互联，打破传统计算节点的边界，提供远超当前单节点 8 颗或 16 颗 AI 芯片的算力规格。这种方式能够显著提升单个计算节点的算力密度和性能。

在 Scale out 方面，主要是通过超高速、超大规模的组网架构来提升整体算力规模。这主要包括两个方面：一是超高速网络技术，通过提供更高的带宽，减少网络延迟，确保大规模集

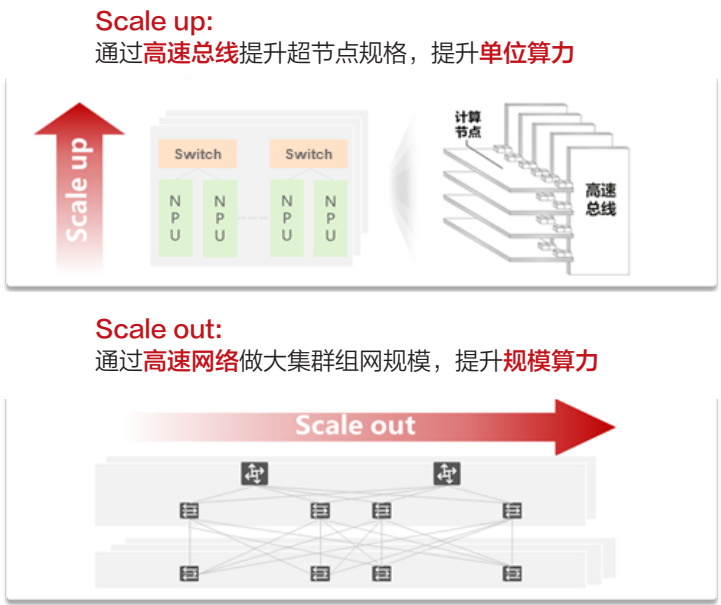


图 4-6 超大规模组网创新方向

群中的数据传输更加高效。**二是超大规模组网架构技术**，如华为星河 AI 网络采用两层框盒及三层盒盒的确定性组网架构，确保了大规模集群中的数据传输稳定性和可靠性，能够实现十万卡以上的超高速网络互联。

通过 Scale up 和 Scale out 的结合，不仅可以显著提升单个计算节点的算力密度，还能在大规模集群中实现高效的网络互联，从而整体提升集群的算力规模。

关键技术 2

基于单机效率优化和集群并行优化，提升集群算力利用率

集群算力利用率 MFU = 单机计算效率 × 集群线性度

提升集群算力利用率主要有两大技术手段：一是单机效率优化，二是集群线性度提升。

1、单机效率优化

单机效率优化的主要思路是软硬件协同优化。具体措施包括：

● **小算子融合为大算子**：采用 FlashAttention 等技术将多个小算子融合为大算子，减少调度次数和 HBM 的读写开销，从而提高计算效率。

● **硬件亲和的算子优化**：通过优化算子与硬件的适配性，减少不必要的调度开销和 HBM（高带宽内存）读写次数，从而提升计算效率。例如，华为通过算法优化和算子融合，充分发挥昇腾硬件的优势，提升了计算性能。

2、集群线性度提升

集群线性度提升的主要思路是算力、网络和存储的协同优化。具体技术手段包括：

● **FSPF 架构亲和并行策略**：在智算集群内部，卡间、节点间、超节点间的互联带宽具备逐层收敛的特征。FSPF（Full-stack Parallelism-friendly）全栈协同技术将 TP（Tensor Parallelism）/EP（Expert Parallelism）高频、大通信量的并行计算与逐层收敛的特征相匹配，有效隐藏或减少通信，提升计算占比。

● **高速总线技术**：利用高速总线技术减少通信时长。高速总线采用全电互联架构，支持超大互联带宽和超高联算比，有效减少通信时长，提升计算占比。

● **算网协同 NSLB 技术**：NSLB（Network-Side Load Balancing）技术支持计算和网络交互训练任务信息，网络路由亲和训练负载，网络吞吐率达到 95% 以上，有效减少通信时长，提升计算占比。

通过这些技术手段，不仅可以显著提升单个计算节点的效率，还能在大规模集群中实现高效的线性扩展，从而整体提升集群的算力利用率。

关键技术 3

基于算存协同 CKPT 加速和故障快速恢复，提升集群可用度

$$\text{集群可用度} = 1 - \frac{\left(\frac{\text{备份间隔 CKPT}}{2} + \text{故障恢复时间 MTTR} \right)}{\text{平均无故障时间 MTBF}}$$

集群可用度是影响集群有效算力的关键因素。提升集群可用度可以从三大方面入手：延长平均无故障时间（MTBF）、缩短故障恢复时间（MTTR），以及缩短 CKPT（Checkpoint）备份间隔时间。以下是具体的措施：

1、系统级高可用架构设计，延长平均无故障时间（MTBF）

● **冗余设计**：在硬件层面增加冗余组件，如冗余电源、冗余网络连接和冗余存储系统，确保单点故障不会导致整个系统的不可用

● **训前压测检查**：在训练前进行全面的压测检查，确保系统在高负载下仍然稳定运行

● **故障智能预测**：采用智能预测技术，针对常见的高故障率部件（如光模块、NPU/GPU、主板、内存等）进行监控和预警。例如，华为自研了光模块通道抗损技术，实现训中不断训，可靠性提升 10 倍；针对 NPU 故障，采用了创新的散热技术和风扇调速优化，使得 NPU 工作温度下降 7 度，失效率降低 30%。

2、算网存协同，缩短故障恢复时间（MTTR）

● **故障感知**：通过智能监控系统快速检测到故障发生，并自动报警。

● **任务调度**：采用智能任务调度算法，确保在故障发生时能够快速切换到备用资源。

● **CKPT 加速**：通过优化 CKPT 保存和恢复流程，减少备份和恢复时间。

● **计算加速**：采用高效的计算加速技术，减少计算任务的执行时间。

● **集合通信加速**：优化集合通信机制，减少通信延迟

例如，华为通过算网存协同优化技术，实现了从故障感知、任务调度、CKPT 加速、计算加速到集合通信加速的全流程加速，显著缩短了 MTTR 时间。

3、算存协同，加速 CKPT，缩短备份间隔时间

- **异步 CKPT 保存：**采用异步 CKPT 保存技术，确保在不影响计算任务的情况下进行备份。
- **本地缓存加载 CKPT：**利用本地缓存技术，快速加载 CKPT，减少恢复时间。

例如，华为研发了本地缓存 + NDS（Near Data Storage）存储直通计算内存方案，显著提升了 CKPT 的读写速率，从而缩短了备份间隔时间。

通过上述措施，不仅可以延长集群的平均无故障时间（MTBF），缩短故障恢复时间（MTTR），还能加速 CKPT 备份过程，从而整体提升集群的可用度。这些方法不仅提高了集群的可靠性，还为超大型 AI DC 的高效运行提供了坚实的技术保障。

场景二

海量用户的分布式推理

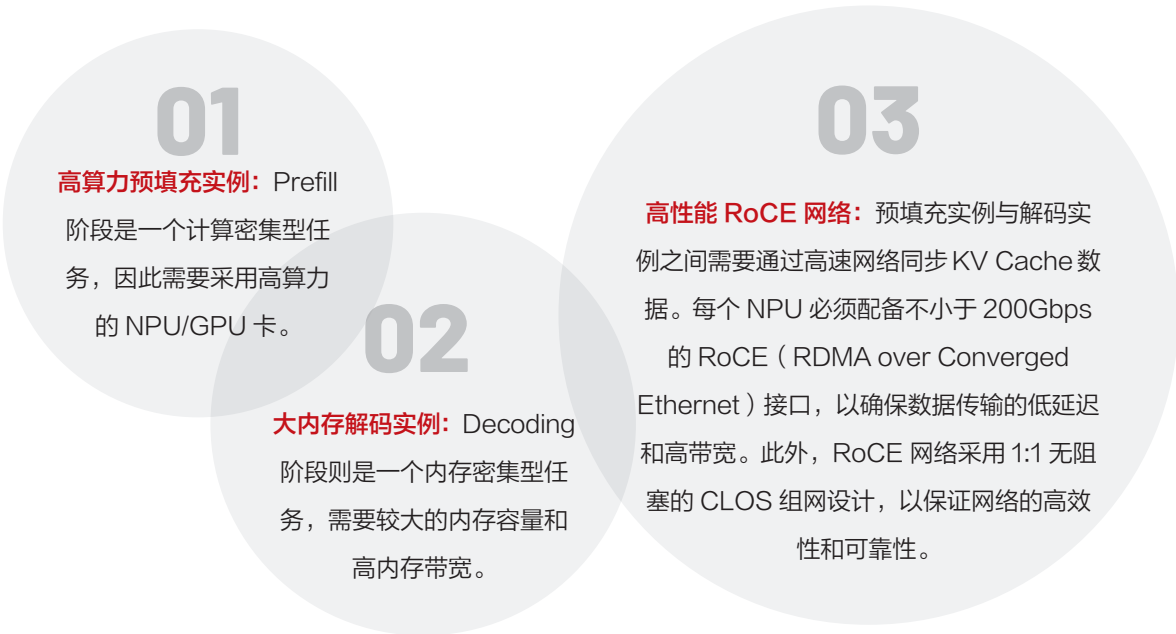
无论是日均调用数万次的典型模型，还是能力持续发展的超大模型（超大参数、超长序列、多模态），在实现高效推理方面都面临较大的挑战，需要采用如下关键技术：

关键技术 1 以 KV Cache 为中心的 P/D 分离技术，提升海量用户的推理效率

海量用户推理面临的挑战是如何在保障用户体验（首 Token 时延小于 1 秒）的前提下，低成本地满足亿级日访问量的服务质量（后续 Token 时延小于 50 毫秒）。当前，许多针对大模型推理的优化技术难以同时达到这两个目标。为此，业界普遍采用的一种方法是将预填充（Prefill）阶段与解码（Decoding）阶段分离，这里所说的 Prefill 阶段，是指处理用户输入的提示（Prompt），生成初始的键值对缓存 KV

Cache（Key-Value Cache）的阶段，它为后续的 Decoding 阶段提供必要的上下文信息。Decoding 阶段根据 Prefill 阶段生成的初始输出 Token 和 KV 缓存，逐步生成完整的输出文本。通过 Prefill 和 Decoding 各自优化并在两者之间通过高性能网络同步 KV Cache，重用计算结果，以此在保障首 Token 时延的同时，实现推理吞吐量提升 2 到 5 倍。

P/D 分离推理架构包括以下几个组成部分：任务调度、多个预填充实例、多个解码实例以及高性能网络。该架构的成功实施依赖于三个关键技术要素：



通过这种 P/D 分离架构，不仅可以有效提升推理服务的质量，还能在大规模并发请求的情况下，维持良好的用户体验，从而实现高效且经济的海量用户推理服务。

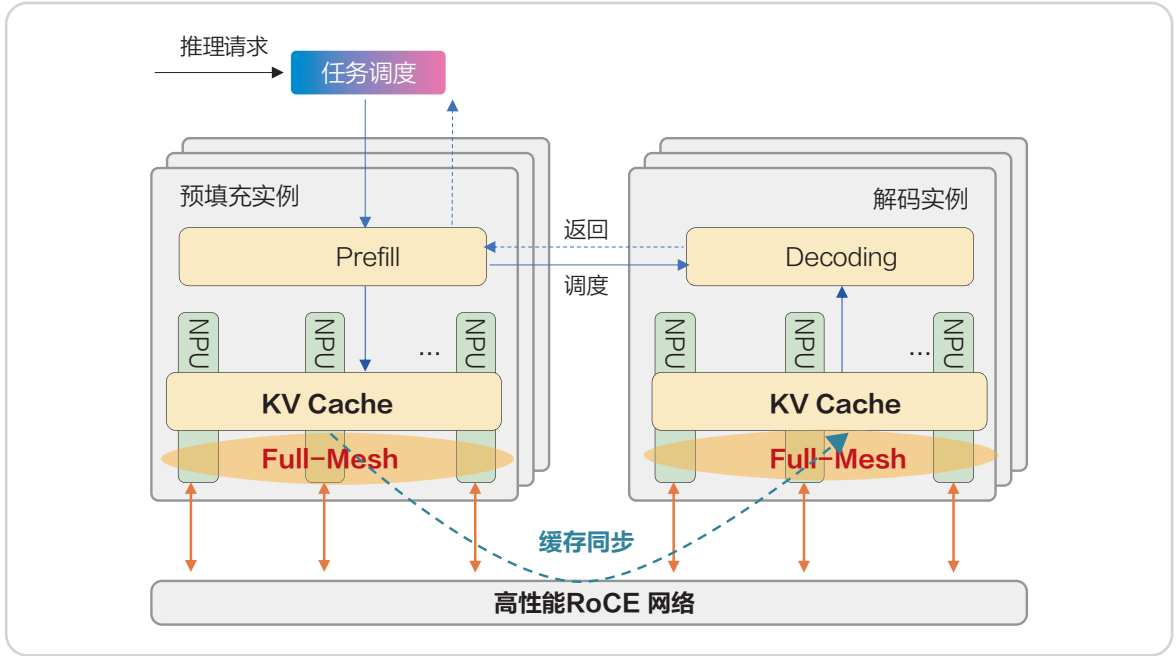


图 4-7 KV Cache 为中心的 P/D 分离技术

关键技术 2 KV Cache 多级缓存，提升超长序列模型的推理效率

长序列模型因其能够处理更复杂的查询和更长的文档，展现出更强的理解能力，甚至可以处理数小时的音视频输入，因而备受关注。目前，业界主流模型纷纷支持更长的文本输入窗口，例如支持 32K 序列长度的模型已经进入商业化应用阶段。

随着序列长度的增加，大模型推理过程中 KV Cache 的大小也随之增加，这不仅延长了推理时间，还大幅增加了推理所需的内存空间。例如，一个典型的 70B 参数模型，处理 1M 长文本需要高达 800GB 的内存来缓存 KV Cache。而对于一个典型的 6B 参数模型，处理 256K 序列长度时，单路并发就需要

100GB 的内存缓存 KV Cache，如果是 8 路并发，则需要 700GB 的内存缓存。

为了解决这一问题，引入分级缓存管理能力成为一种有效的解决方案。在传统的 HBM（High Bandwidth Memory）作为一级 KV Cache 的基础上，可以将主机的 DRAM 作为二级 KV Cache，甚至引入高性能的专业存储设备作为三级 KV Cache。通过这样的多级缓存管理机制，可以实现“以存代算”，即通过高效的缓存策略来减少计算负担，从而有效降低推理时延和推理成本。这种方法不仅提升了模型处理长序列的能力，还为实现大规模应用提供了可行的技术路径。

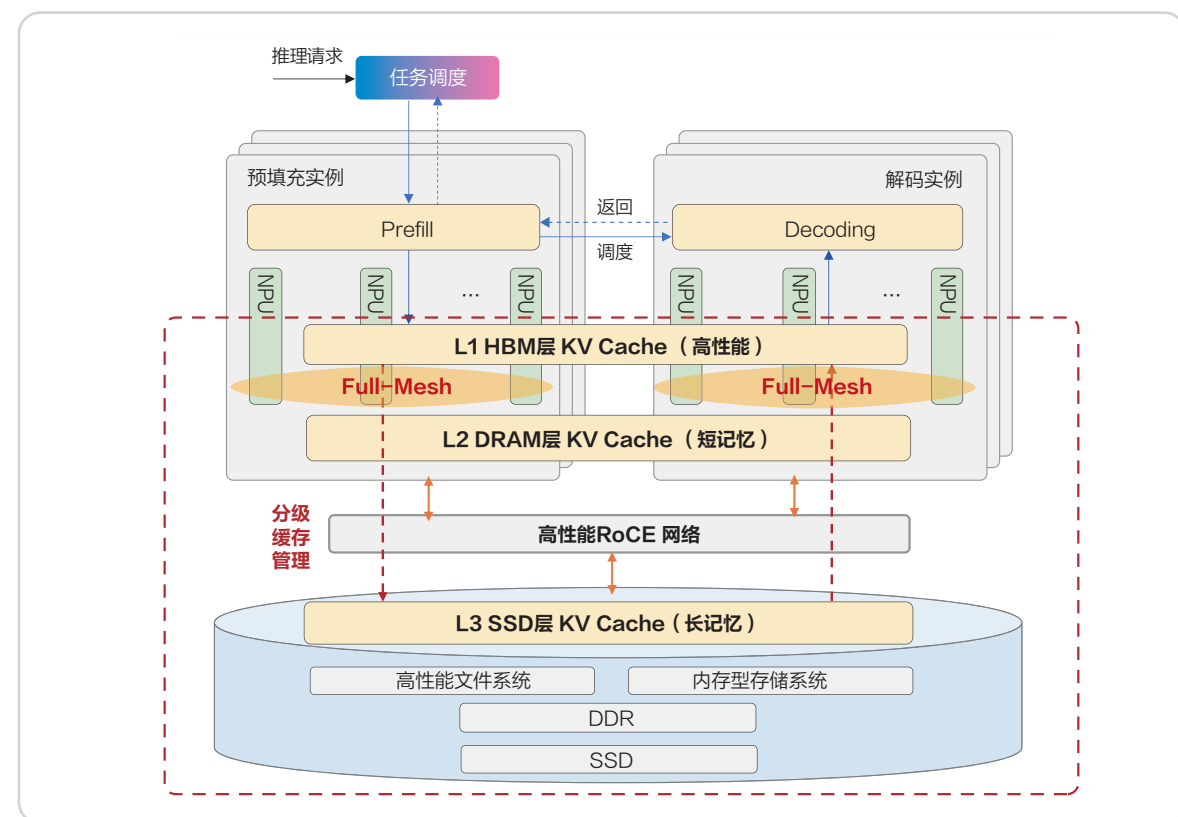


图 4-8 KV Cache 多级缓存技术

关键技术 3 多机并行推理，提升超大参数和多模态大模型的推理效率

模型参数的数量直接影响推理计算量和内存占用，超大参数模型因此面临推理效率低下和内存占用过高的问题。与此同时，多模态模型由于其输入序列长度可扩展到百万级别，使得 Attention 机制下的长序列计算成为内存管理和计算效率的重大挑战。目前，主流的多模态模型采用了解码器 - Transformer（DiT）架构，这种架构需要多次迭代才能生成最终结果，导致了较高的资源消耗和较长的推理时间。

为了有效解决上述问题，针对超大参数模型和多模态大模型，有必要采用多机并行推理架构来提升推理效率并减少推理延迟。在推理架构的设计层面，每个计算节点内部的 NPU 应当使用全互联架构来保证高速数据交换；而在节点之间，则通过高性能的 RoCE 网络连接各个 NPU，以此增强系统整体的通信效率。

此外，还需要采取一系列优化措施来进一步改善推理效率：



通过这些优化手段，可有效地缓解由模型参数量庞大和多模态特性所带来的计算和内存压力，从而提高推理性能。

规划建设方向二：极致能效

如何打造高效率、高密度、高弹性、高可靠的能源基础设施，实现极致能效，满足 AI DC 可持续发展要求，业界主要有如下几个关键技术方向：

关键技术方向 1 弹性的能源基础设施模块

弹性能源基础设施模块需支持多代算力及多元化算力混合部署，适应未来业务发展的不确定性。通过模块化和标准化的设计，实现能源基础设施模块的流水线式快速交付。具体来说，将数据中心划分为若干标准

化的能源基础设施模块，预先规划好空间布局，以便容纳不同冷却技术（如风冷或液冷）、不同功率密度的算力设备。

关键技术方向 2 极致的供配电效率

通过软硬协同创新可以提升供配电能效、密度和可靠性。当前业界先进的数据中心，采用了 0ms 切换的智能 ECO 工作模式，可以将供配电效率提升到 97.8%，以 10 万卡 100MW 数据中心为例，年供配电损耗从 4800 万度降低到 1800 万度；此外，设备

厂商通过电力模块替代变压器、低压配电、UPS、输出配电等多个独立的产品，再配合锂电池，供配电密度可提升 1 倍；最后，通过配电连接点、电容等故障的预测性维护，可以大幅提升供配电的可靠性。

关键技术方向 3 极致的散热效率

随着算力功耗攀升，液冷成为必选项，但在实际选择液冷时，用户普遍存在可靠性和运维复杂的担忧，因为液冷更贴近服务器，漏液、中断对 IT 设备可靠运行影响更大，同时液冷新技术、新材料、新设备等需要新的运维技能。

为了应对这些挑战，液冷需要从芯片、服务器、机柜到冷源的全面创新。以华为天成液冷系统为例，芯片

采用微铲齿冷板散热，热流密度可达 180W/cm²，满足 1000W 以上芯片散热需求；服务器采用冷、电、网络盲插设计，实现部署和维护极简，无滴漏快接头，结合漏液在线监测、漏液隔离等提升可靠性；冷源系统支持液冷和风冷共用，通过间接蒸发冷却 AHU(Air Handling Unit) 的升级，可以提供 18~35℃ 常温水直供液冷服务器，风液一体化冷源设计可减少冷源投资、降低维护难度，并支持 PUE 低至 1.10。

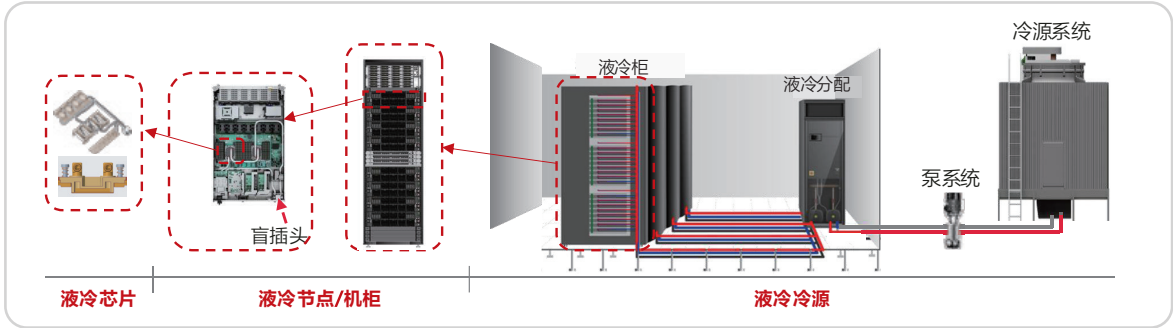


图 4-9 液冷散热系统

关键技术方向 4 联动调优降能耗

通过 AI 联动调优降低能耗，提升系统可靠性。AI 联动调优大脑采集到环境、制冷设备、供电设备、IT 设备、训推任务等参数，通过能耗优化模型、器件故障预警模型等实时预测最佳工作参数，并下发调优策略，实时调整冷却塔、水泵、CDU 的输出状态，IT 设备的电源工作模式等参数，实现 DC 综合能耗的下降。例如，在华为云数据中心实践中，基于云服务感知 AI 能效调优，精度达 99.5%，DC 能耗下降 8~15%。

超大型 AI DC 综合评价指标体系表

一级指标	二级指标	三级指标	指标描述
训练 算效	算力规模	算力规模	单节点算力 × 节点数，单位 PFLOPS(FP16)
	算力利用率	算力利用率	模型的实际计算需求与其理论最大计算能力之间的比率，单位 %
	可用度	故障恢复时间	训练任务由故障状态转为工作状态时的修复时间，单位分钟
		平均无故障运行时间 MTBF	相邻两次故障之间的平均工作时间，也称为平均故障间隔，单位天
推理 算效	时延	首 Token 时延（TTFT）	从模型开始处理输入到产生第一个输出令牌（Token）所需的时间，单位是 ms 或 s
		Token 间隔时延（TBT）	指生成连续输出令牌之间的平均时间间隔，单位是 ms 或 s
	精度	精度	模型在执行推理任务时正确回答的比例，单位是 %
	吞吐	吞吐	系统在单位时间内能够处理的请求数量或者能够生成的结果数量，单位是 Tokens/s
能效	能源效率	PUE	数据中心的总电量 ÷ IT 设备用电量，无量纲

01 建设实践 科大讯飞打造极致算效的超大规模 AI 算力集群

为加速发展讯飞星火大模型，科大讯飞于 2023 年 10 月建设完成超大规模 AI 算力平台“飞星一号”。

挑战

超大规模集群组网

超大规模集群如何进行高效组网、软硬调优，以支撑万亿参数大模型训练。

高效训练

算力成本高昂，可获得性难，如何提升算效是关键挑战之一。

长稳训练

集群包含数千的计算、存储、网络等设备、数万光模块和光纤、百万级器件，故障定位难。

成效

“飞星一号”投产后成功保障了讯飞星火大模型迭代训练。目前，大模型已赋能金融、能源、汽车等行业的智能化升级。



创新

超大规模无损组网

基于 RoCE 开放协议，实现超大规模集群无收敛、无损的高速组网，满足数据并行、流水并行等通信需求。

动态 NSLB 负载均衡提升训练效率

通过网络控制器获取计算任务和通信域信息，主动规划通信路径，网络带宽利用率提升到 95%。

算存协同加速 CKPT

集群提供 TB 级大带宽，缩短 CKPT 读写耗时，断点续训恢复时长从 15min 缩短到 1min，速度提升 15 倍。

跨域统一运维

大模型训练过程中面临光模块故障、流量抖动等多种问题，平台通过捕获故障码和定义自动故障处理流程，实现 80 多种常见故障的自愈时间在 10 分钟以内。

极致性能优化

完成基于昇腾的大模型训练工具链，及高性能算子库研发和性能调优；依托讯飞自研大模型并行训练框架，当前基于昇腾算力的大模型训练性能处于国内领先水平。

02 建设实践 中国电信临港智算中心通过“两弹一优”打造极致能效的 AI DC

中国电信持续加强算力基础设施建设，通过“两弹一优”，构建单体万卡的液冷智算数据中心。

挑战

作为承载超大规模算力集群的示范基地，如何满足为入驻用户提供高算效和高能效的绿色算力。

创新

弹性架构

通过“两弹一优”，即弹性供电、弹性供冷与优化气流组织。其中，弹性供电是通过电缆隧道、大小母线等技术实现跨机楼、跨楼层、跨机房的电力按需敏捷调度；弹性供冷是通过水、风、液三种冷源统筹规划，满足不同发热量的机架供冷需求；优化气流组织则是利用 AI 优化机房的气流管理，让数据中心的冷空气运用得恰到好处。

高能效

采用了冷板式液冷 + 高效模块化 UPS 提升能效。冷板式液冷由冷却塔、水泵、CDU、液冷柜等组成，冷却塔提供的一次水直接进入 CDU，通过 CDU 的换热，输出中高温的二次水溶液直接进入服务器，去冷水机组的设计实现了 100% 自然冷却。高效模块化 UPS 支持智能 ECO 模式，效率高达 99%，此外，在低负载时可通过模块休眠进一步降低配电损耗。

成效

实现 PUE < 1.25，部分区域 PUE 最低至 1.08，相比同区域其他智算数据中心节能 15% 以上，每 PFLOPS 算力能耗小于 1.5 千瓦。



| 大型 AI DC

大型 AI DC 承载了企业智能化转型过程中的核心业务和资产，因此，如何科学合理地规划和建设 AI DC 变得至关重要。本节重点描述大型 AI DC 建设过程中面临的关键需求与挑战，并总结大型 AI DC 的关键特征，同时，基于这些特征，提出五大方向性的建议。

关键建设需求

企业主要面临八大方面的关键需求及挑战。



图 4-10 企业自建大型 AI DC 面临的 8 大关键需求和挑战

● **未来可演进的 AI DC 架构规划:** 面对 AI 技术的快速发展和新的智算技术体系，如何规划一个面向未来的可演进 AI DC 架构，避免建成即落后的情况，从而防止大量的人力、物力和资金浪费。

● **智算基础设施打造:** 如何确保建成后的算力基础设施能够支撑高效、稳定的训练和微调，并满足 AI 推理的“LACE”体验要求。

● **算力资源的充分利用:** 如何充分利用“稀缺”的 AI 算力资源，尽量避免算力闲置，实现训推任务的灵活调度(时分复用)及一卡多分(空分复用)，提高资源利用率。

● **提升 AI 开发效率:** 如何提升 AI 开发效率，保障企业大规模 AI 模型和应用的快速迭代，以满足业务的快速发展需求。

● **AI DC 能源基础设施准备:** AI DC 在供电、散热、承重和布线等方面的要求远高于传统数据中心。如何结合未来算力基础设施的发展趋势，提前对 AI DC 的能源基础设施进行前瞻性的规划和准备，以确保能源基础设施能够满足未来的需求。

● **模型迭代与应用兼容:** 大模型正处于快速发展阶段，其迭代速度远超 AI 业务应用。如何避免模型的更换对上层应用造成影响，同时现网小模型是否需要迁移到新的智算基础设施上，以及如何更好地结合大小模型的优势，发挥组合效应。

● **降低运维难度:** 面对 AI DC 普遍存在的故障概率高、中断损失大等问题，如何降低运维难度，提升运维效率。

● **应对新安全风险:** 如何应对生成式 AI 带来的新安全风险，如大模型输出内容的不确定性、提示词注入攻击等，确保数据和系统的安全性。

为有效应对这些关键需求及挑战，大型 AI DC 需要具备融合与高效的特点。其中，“融合”具体体现在以下 4 个方面：

- **训推融合：**大型 AI DC 不仅承载模型训练业务，同时也承载 AI 推理业务。通过在同一平台上同时支持训练和推理，可以实现资源共享，简化流程，并提高整体效率。
- **通智融合：**大型 AI DC 不仅包含智算，还包括通算。通算业务和智算业务的混合部署成为 AI DC 的常态。
- **风液融合：**由于大型 AI DC 中既有智算也有通算，不同类别的算力所需的散热模式也不同。对于低密度的通算、网络和存储设备，通常采用风冷散热；而对于高密度的智算，液冷散热逐渐成为刚需。面向未来，风冷和液冷在 AI DC 内共存将成为必然趋势。

● **多模融合：**在实际的企业应用中，往往需要多个模型共同支撑一个完整的 AI 应用。例如，在眼科辅助诊疗助手应用中，首先是调用一个语言大模型用于与用户进行对话交互，同时调用另一个语言大模型来理解用户问题并负责诊断流程的编排规划。基于该规划，再调用若干计算机视觉（CV）模型进行眼疾图像筛查。通过多个模型的共同配合，最终输出诊断报告。因此，大型 AI DC 往往会部署 NLP（自然语言处理）、CV（计算机视觉）等主流类别的模型，也可能部署多模态、预测等其他类别的模型，既有大模型也有小模型，实现多模态的混合部署。

图 4-11 多模融合实现辅助诊疗助手

“高效”主要体现在五大方面，即架构高效、开发高效、算力高效、能源高效及管理高效。这也是领先的大型 AI DC 规划建设的五大方向。



图 4-12 领先的大型 AI DC，五大规划建设方向

规划建设方向一：架构高效

对于大型企业的 AI DC 来说，大规模多样化的算力资源和模型是两大关键资源。一个 AI DC 架构是否高效，主要体现在算力管理和模型管理是否高效上。以下是具体的需求及解决方案：

算力管理高效

当前，对于企业智算基础设施的管理者而言，面临的重大挑战之一是如何高效地管理和调度 AI 算力资源。在大型 AI DC 中，通常包含多个 AI 算力资源池，例如训练资源池、推理资源池以及训推弹性资源池等。为了满足业务需求并提升资源利用率，需要将训练和推理任务在这些资源池之间灵活调度切换。

此外，一些大型企业可能不仅拥有一个 AI DC，而是有多个分布于不同地理位置的 AI DC。这就需要具备跨 DC 的算力调度能力。更为复杂的是，在遇到突发业务需求时，企业自建的 AI DC 算力资源可具体而言，这种架构需要具备以下特点：

能不足以应对，还需要租用公有云或行业云的算力资源。在这种情况下，需要具备跨本地与公有云的算力调度能力。

为了解决这些问题，需要构建一个高效的算力管理调度架构，实现模型和算力的解耦。通过统一的算力资源管理和调度平台，为上层模型的训练和推理任务提供标准化的算力服务封装 API 接口，屏蔽底层多池、多中心、多云的算力复杂性，形成逻辑上统一的 AI 算力资源池。这样可以确保 AI 算力资源集中可视、可管、可控，从而提升资源利用率。

- **统一管理：**通过统一的管理平台，实现对本地、多个数据中心以及公有云中的算力资源集中管理，形成逻辑上统一的资源池。
- **跨域调度：**支持跨数据中心以及跨云的算力调度，确保在突发业务需求时能够快速扩展算力资源。
- **灵活调度：**根据业务需求，灵活地在不同资源池之间调度算力资源，支持训练和推理任务的动态切换。
- **资源可视化：**提供资源监控和可视化工具，使得算力资源的状态一目了然，便于管理和优化。
- **标准化接口：**提供标准化的 API 接口，使得上层模型及应用能够无缝对接底层算力资源，简化开发和运维工作。

通过这样的架构设计，企业不仅能够提高算力资源的利用率，还能在复杂多变的业务环境中保持敏捷性和灵活性，从而更好地应对未来的挑战。

模型管理高效

在企业应用大模型的过程中，另一个重要的挑战是大模型本身的不确定性。在当前“百模大战”的竞争环境下，市场上大模型的更新换代速度非常快，这为企业带来了机遇，但也带来了挑战。一方面，可用的大模型能力越来越强；另一方面，如何确保上层应用不受频繁更迭的大模型的影响，成为了一个亟待解决的问题。

具体而言，这种架构需要具备以下特点：

● **模型编排**：通过模型管理和编排模块，将模型与应用分离，使得模型的变化不会直接影响到上层应用。这样即便底层模型更新换代，上层应用依然可以稳定运行。

● **模型版本管理**：实现模型的版本管理功能，确保不同版本的模型能够共存，并且可以根据需求灵活切换使用。这样可以更好地管理和利用不同版本的模型。

● **监控与反馈**：提供模型运行状态的监控和反馈机制，确保模型在实际应用中的性能和效果能够被实时监控，并根据反馈进行优化调整。

这就需要一个高效的架构，实现模型和应用的解耦。通过统一的模型管理和编排模块，提供标准化的模型能力封装 API 接口，屏蔽模型更迭或替换对应用带来的影响。这样可以确保上层应用的稳定性和可靠性，同时还能充分利用最新的模型能力。

● **标准化接口**：提供标准化的 API 接口，使得上层应用能够无缝对接底层模型，简化开发和运维工作。无论底层模型如何变化，上层应用只需调用统一的 API 接口即可。

● **自动化部署**：支持模型的自动化部署，简化新模型的上线流程，确保新模型能够快速、安全地部署到生产环境中。

通过这样的架构设计，企业不仅能够提高模型应用的稳定性和可靠性，还能在快速变化的市场环境中充分利用最新的模型能力，从而保持竞争优势。这种架构不仅提升了企业的灵活性，还增强了系统的整体稳定性和可扩展性。

通过以上两个方面的高效管理，企业不仅能够提升算力资源的利用率，还能确保模型应用的稳定性和灵活性，从而在激烈的市场竞争中保持领先地位。

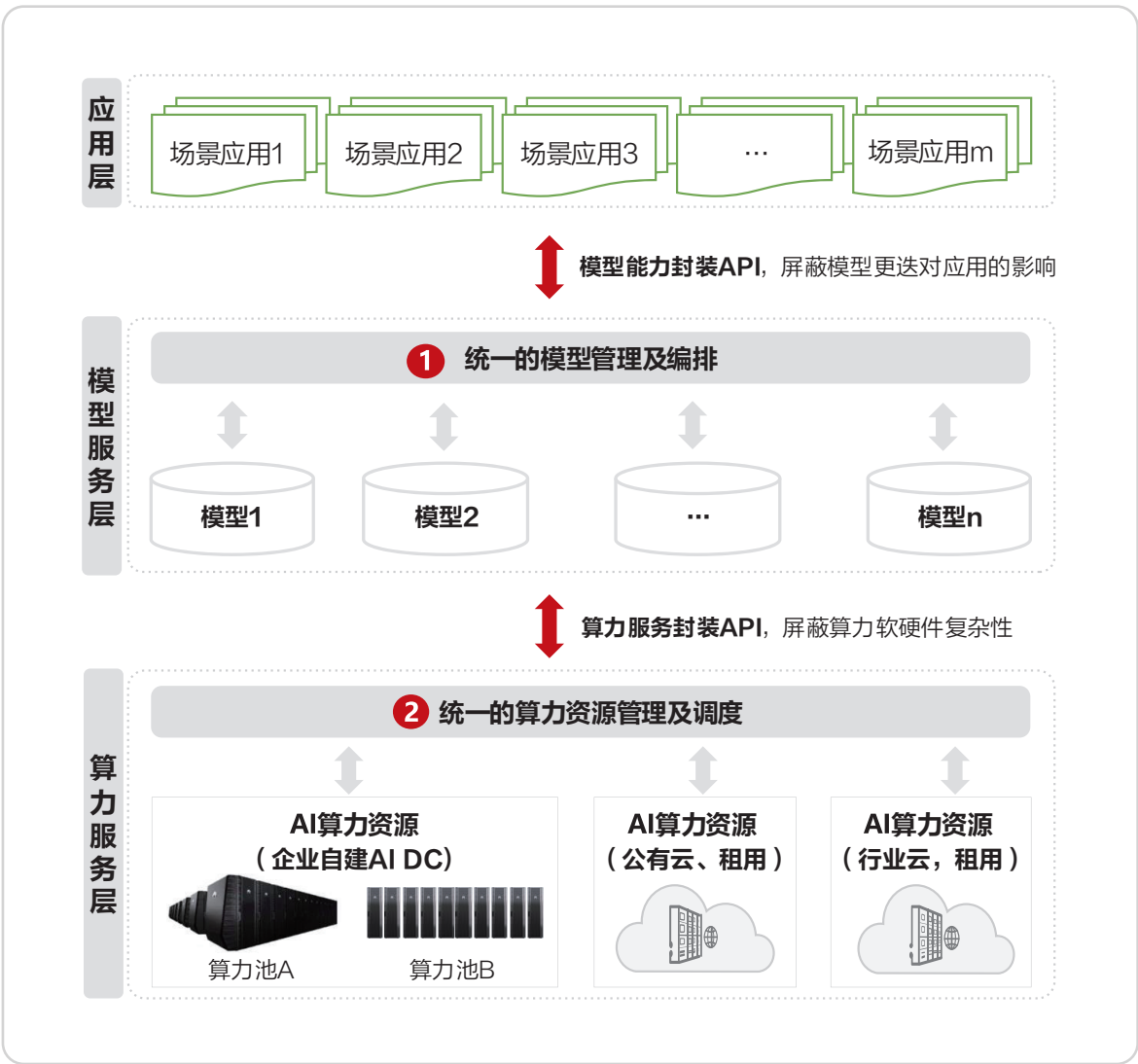
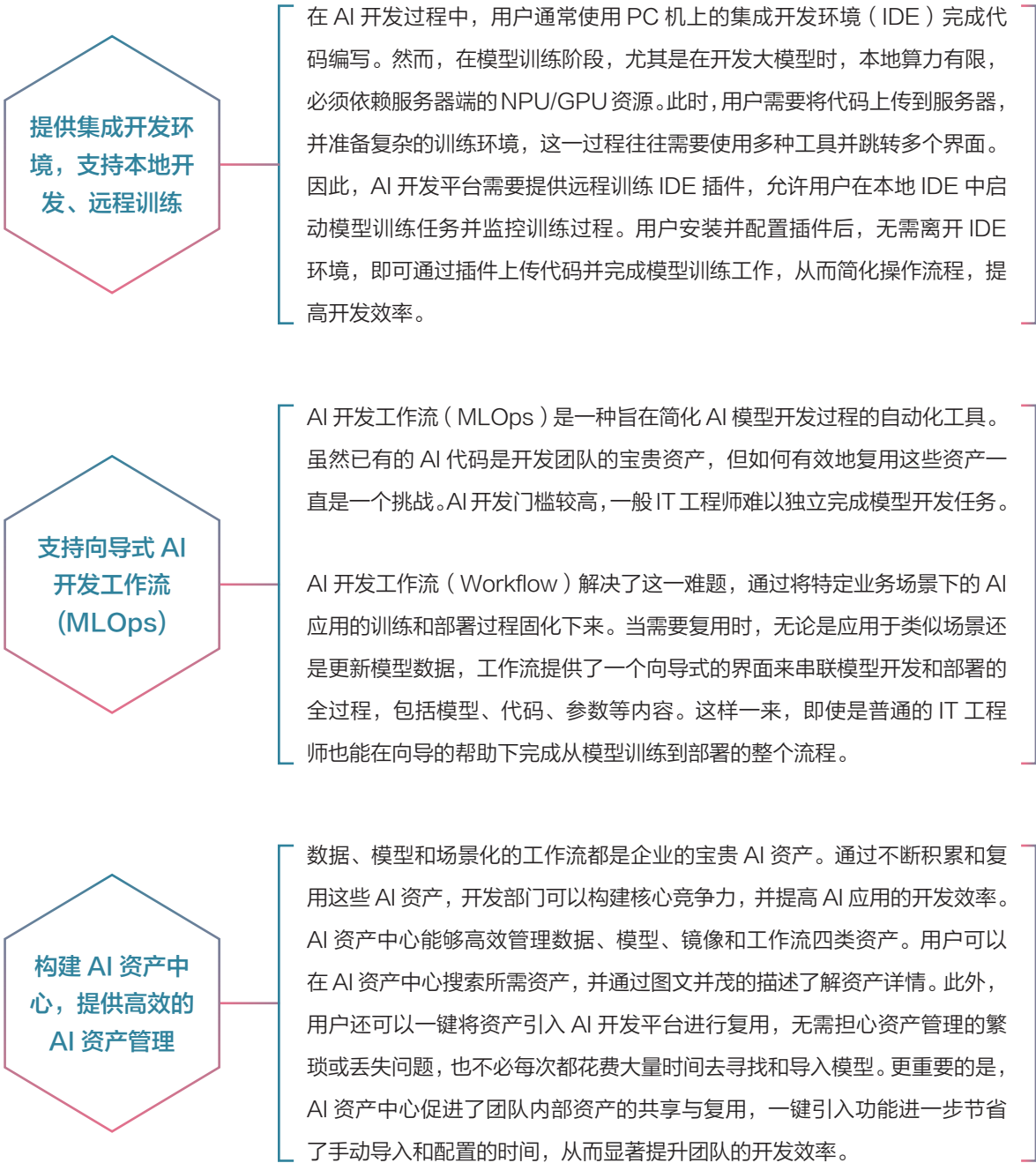


图 4-13 分层解耦的技术架构，实现算力 & 模型管理高效

规划建设方向二：开发高效

为了应对企业在 AI 开发过程中普遍存在的**开发效率低、成本高、模型一致性低、可靠性差以及模型部署时间长**等问题，需要构建一个高效的 AI 开发平台，以提升开发效率并降低开发成本。该平台的核心能力包括以下三个方面：



规划建设方向三：算力高效

当前，面对“稀缺”的 AI 算力资源，企业希望尽可能地压榨出“每一滴”算力，发挥出算力的最大价值。要想提升算力资源的利用率，主要可以通过以下两种手段：

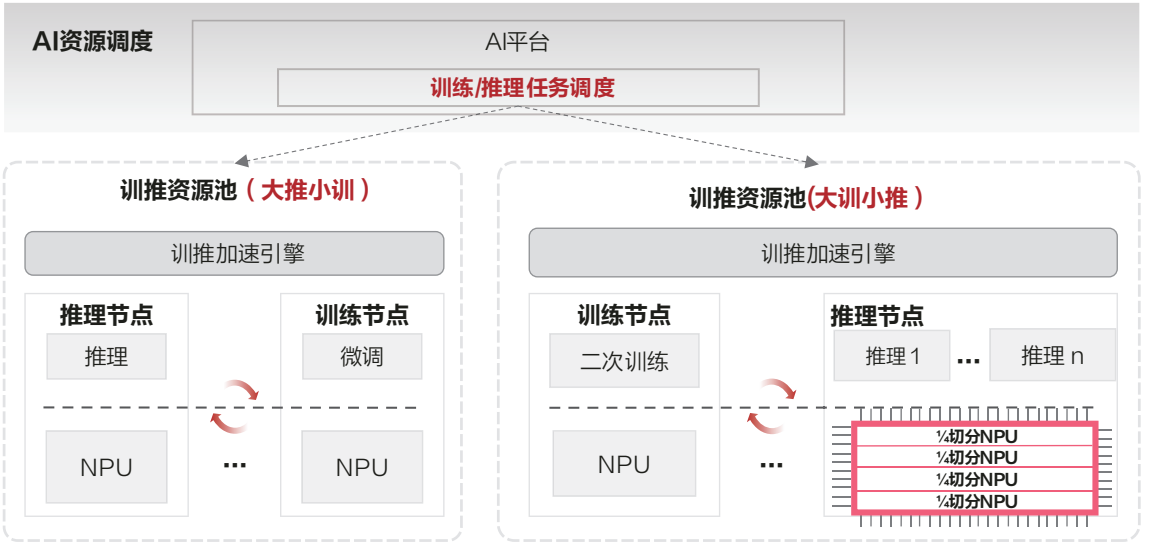
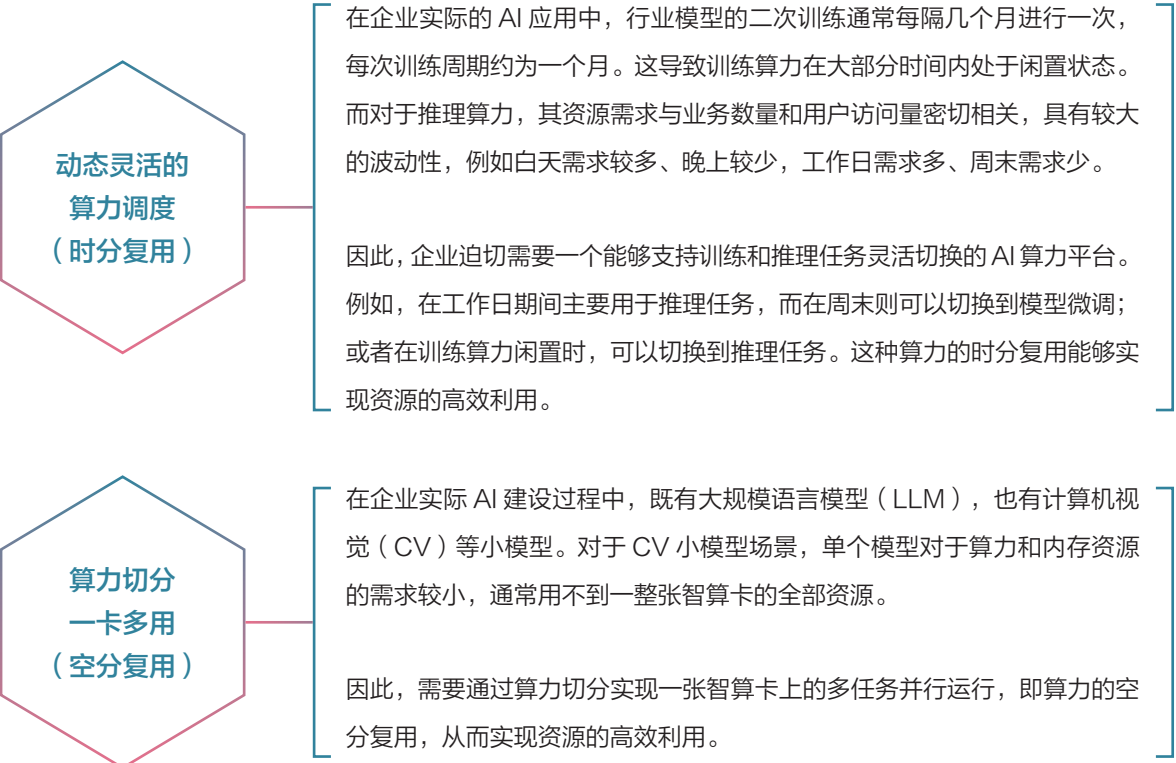


图 4-14 算力时分 / 空分复用示意图

通过实现算力的时分复用和空分复用，企业不仅能够有效提升算力资源的利用率，还能在不同的业务高峰期灵活调配资源，确保关键任务的顺利进行。这种高

效的算力管理策略不仅有助于降低运营成本，还能提升企业的整体竞争力。

规划建设方向四：能源高效

能源基础设施的高效，核心是满足智算液冷与通算风冷混合部署需求，当前主要面临如下挑战：

● **挑战 1：风液按需动态配比**，智算与通算业务均需支持独立扩展，需要风液适配通算智算的不确定性业务。

● **挑战 2：高低密混合部署**，智算每柜高达数十到上百 kW，通算通常每柜小于 10kW，供电、制冷如何适配多功率密度需求。

● **挑战 3：如何适配复杂的机房条件**，新建场景如何分区规划，改造场景如何与现有系统保持统一接口，同时不影响现有业务运行。

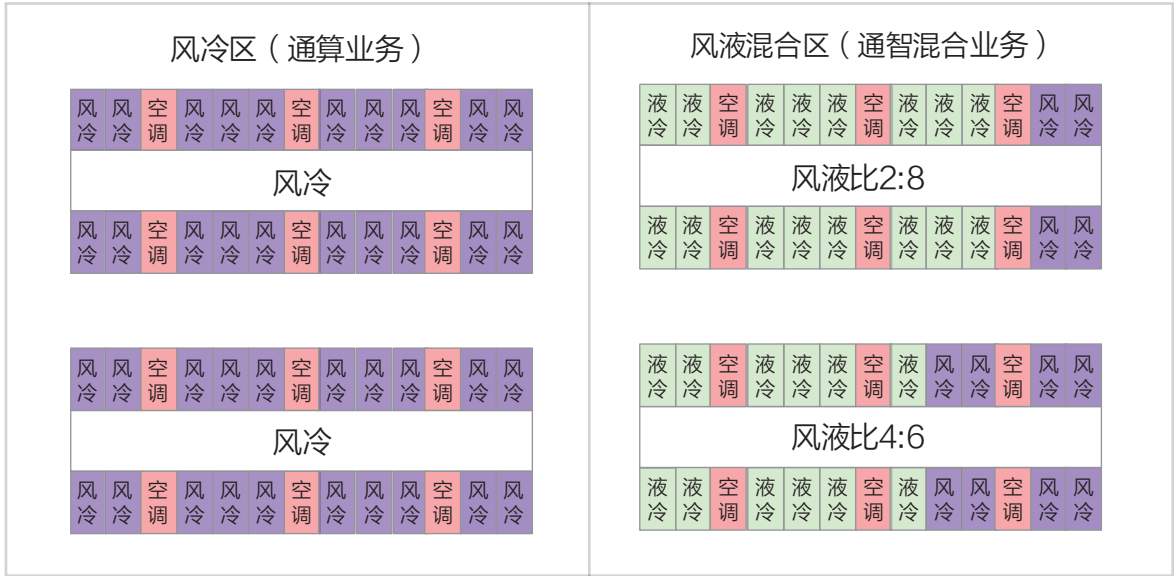


图 4-15 大型 AI DC 机房分区示意图

推荐的做法是，在 AI DC 划分风冷区和风液混合区，风冷区面向确定性的通算业务，风液混合区作为弹性空间，支持智算和通算业务的按需扩展。基于机架总数基本不变、制冷系统不变、供电最大能力固定的原则，可以按照 70% 液冷需求预留空间设计，并根据实际风液配比需求上线供电、制冷设备。

在风液混合区，采用风液混合微模块，实现风液弹性配比。风液混合微模块可内置模块化列间空调或共用房间空调池、液冷管路、液冷 CDU、配电单元等，支持液冷机柜 0~100% 弹性部署，按需上架。以华为微模块为例，液冷柜采用冷板式液冷 + 风冷空调散

热，最高可支持 66kW/ 柜。风冷柜部署通算、存储、网络、安全等设备，采用风冷空调散热，最高可支持 35kW/ 柜。

风液混合微模块提供标准的水、电、管理接口，实现快速集成交付。水接口与数据中心内冷源系统对接，支持 18~35℃ 水温。电接口与低压配电系统对接，采用智能母线配电，可以提供 16/32/40/63A 多种热插拔开关，适配机柜不同功率密度需求。管理基于标准协议接口统一接入 DCIM 运维管理平台，实现统一管理。

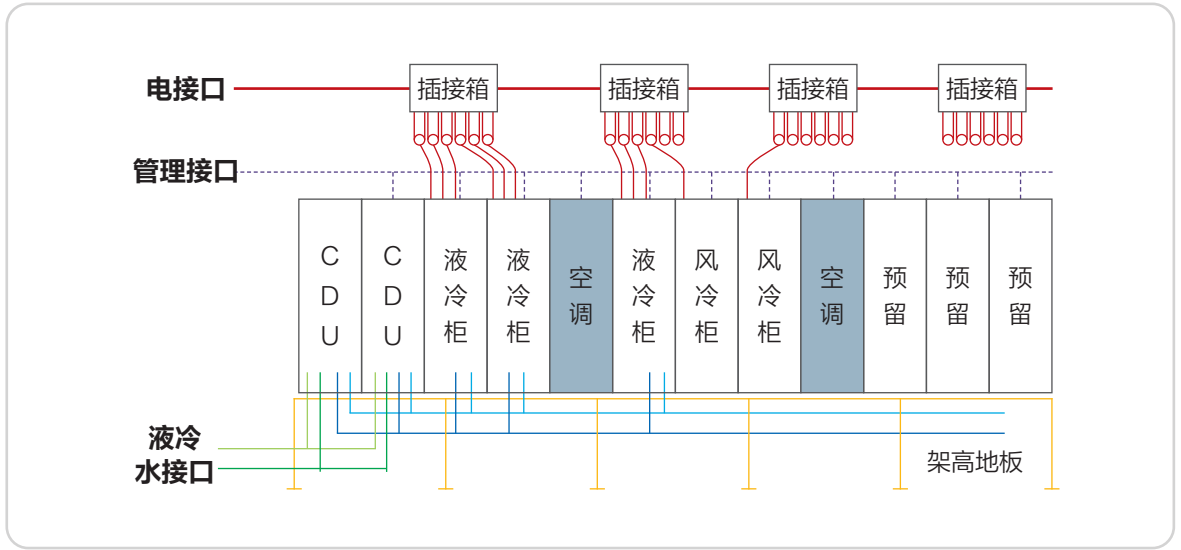


图 4-16 风液混合微模块示意图

规划建设方向五：管理高效

为应对 AI DC 更加复杂的运维和新型安全风险挑战，必须构建新一代运维管理平台及安全防护体系，以提升管理效率。

新一代运维管理系统应具备全面监控、故障预测、故障快速分析定位及恢复等功能，简化运维人员日常操作，降低运维难度。具体能力包括：

- **全链路可视化监控**：通过实时监控整个系统的运行状态，实现对计算、存储、网络等各资源的全面监控，确保任何异常都能及时发现。
- **跨域故障自动感知**：通过预先内置故障模式库，经过端到端信息流协同分析，可对常见典型故障信息进行自动匹配，感知故障发生。

- **跨域故障快速定位**：利用先进的故障检测技术，如日志分析、分钟级网络流量分析、存储故障和性能分析，实现故障产生时快速准确定位故障点，减少故障排查时间，避免训练任务中断。
- **跨域故障快速修复**：建立高效的故障修复机制，确保一旦发生故障，能够迅速采取措施恢复系统正常运行，减少停机时间。

新一代安全防护体系应在传统 DC 的基础环境安全 和安全运营的基础上，**重点新增模型安全，并针对智算业务强化数据安全及应用安全等能力。**

大型 AI DC 综合评价指标体系表

一级指标	二级指标	三级指标	指标描述
架构效率	算力管理调度能力	资源池种类	平台能够管理调度的 AI 资源池数量和规模，例如自建 AI 算力资源，租用的公有云或行业云的 AI 算力资源
	模型管理编排能力	模型种类	平台能支持的模型种类，例如传统模型（ CV、OCR、ASR 等 ）、NLP、多模态等，单位是模型种类
开发效率	开发效率	模型开发时长	平台应具备完善的模型开发工程套件，支持模型设计、模型调测、训练脚本编写、模型精度问题定位等能力，单位是人天数 / 训练或微调
	部署效率	模型部署效率	平台应具备模型自动化部署能力，提供单机多卡、多机多卡等多种部署方式，实现模型的高效部署，单位是部署完成一个实例需要的时长（ 秒 ）
算力效率	算力利用率	资源利用率	平台能够监控资源使用率，支持任务动态管理调度，实现多中心多池多集群的 AI 算力资源利用率最大化，单位是百分比
	训练效率	模型训练时长	从模型开始训练直到模型精度收敛到目标值所经过的时长，单位是天
	推理性能	推理首 Token 时延	从模型开始处理推理请求，到输出第一个 Token 所需的时间，简称 TTFT（ Time To First Token ）， 单位是 ms 或 s
		推理 Token 间隔时延	模型推理过程中，连续输出 Token 之间的平均时间间隔，简称 TBT（ Time Between Tokens ）， 单位是 ms 或 s
		推理吞吐率	AI 集群单位时间内能够处理的推理请求与生成结果的 Token 数之和，单位是 Tokens/s
能源效率	能源效率	PUE	数据中心的总电量 ÷ IT 设备用电量，无量纲
运维安全	运维	故障定位时间	作业运行时，AI 集群出现故障到故障首次被发现的平均时间，简称 MTTD（ Mean Time To Detect ），单位是分钟
	安全	生成内容合格率	对生成内容的安全情况进行评估，采用人工、关键词和分类模型抽检，随机抽取若干条生成内容，其中满足内容合规要求所占的比率，单位百分比



01 建设实践 某大型银行打造“一底座、两平台、三中心”的技术架构实现架构、开发、算力高效

某大型银行通过构建全栈可信赖的千亿级金融大模型平台赋能体系，实现金融业务的智能化升级。

挑战

算力集群建设和管理难

算力集群需要支持数十天高效率、稳定的训练，及海量用户的并发推理，如何实现资源高效管理和调度，提升算力利用率。

数据准备难

传统人工知识运营模型效率低、工作量大，与大模型工程要求差距明显，需要建立一套适配大模型的高效

知识工程体系，提供知识的分层管理、新旧迭代、内容可信监测等。

多元大模型选择和使用难

单个大模型无法满足全部金融场景需求，如何在保障知识共享、能力复用的基础上，快速集约的打造不同领域不同场景的丰富大模型能力，通过择优调用实现成本效益最优。

创新

建成千卡规模的 AI 算力云

打造“一底座、两平台、三中心”AI 全栈技术架构，支撑 AI 的开发、落地、持续演进。在底座层面，采用 RoCE 高性能网络技术、多层级高性能存储技术和大规模 AI 集群等技术，建成“高速互联、高效存储、云智融合、无感协同”的人工智能 AI 算力云，满足模型高效、长稳训练需求；同时提供了训练算力集群、推理算力集群等资源统一管理和调度能力，实现算力资源高效利用。

建设千亿金融大模型算法矩阵

通过一站式 DataOps 数据研发流水线以及 AI4Data 能力、大模型兼容适配框架和测评体系，有效构建了多层次、多模态、大小模型协同融合的大模型算法矩阵，适配 10+ 领先主流大模型，同时建立以金融智能中枢 Agent 为核心的大模型应用范式和工程化解决方案，提供面向场景开箱即用的模型组合服务，使能应用快速构建。

成效

目前，该行已在客服、信贷、办公等几十个领域实现业务效能提升，以远程银行坐席为例，在账户受控等重点场景，实现坐席通话时长压降过 10%，高频场景坐席服务效率提升近 20%。未来，AI 将继续深度赋能全行几十万员工，服务更广泛的用户。

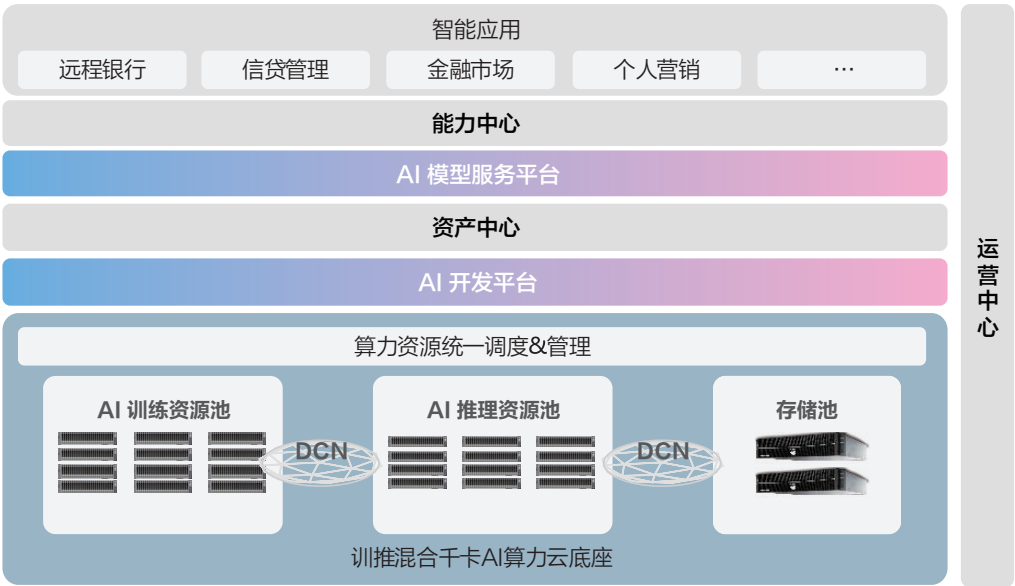


图 4-17 “一底座、两平台、三中心”技术架构

02

建设实践

某头部保险公司以数字劳动力为抓手
打造稳定高效的千卡级算力集群

某头部保险公司在行业内率先建设数字劳动力实验室，致力于打造数字化等效劳动力，以及赋能员工成为超级员工，提升企业劳动生产率。

挑战

如何实现算力底座稳定高效和持续演进：

不影响业务前提下，平滑演进到千卡级集群，支撑规模商用。

如何设计安全合规灵活的 AI 技术架构：

在安全合规等规范下，跨生产区、测试区构建训推一

体算力底座，最大化释放算力价值，支撑高效敏捷开发。

如何打造懂行的行业大模型：

保险行业涉及保险学、精算、法律等众多领域，如何持续构建行业高质量数据集，并持续提升保险大模型专业能力。

创新

算网存管协同，实现高效训推：

2024 年年底前该保险公司将完成 160P 算力基础设施建设，未来规划将逐步演进到 1000P。通过软硬协同，长稳训练，智能运维，提升 MTBF 时间；通过训推一体，平滑扩容，千卡算力统一管理和调度，提升资源利用率。

优化多机训练策略和全流程开发工具链：

优化并行策略、重计算策略等多机训练策略，充分释放算力提升性能，并构建梳理全流程开发工具链，实现高效敏捷大模型研发。

基于高质量数据集，构建领域专家 Agent：

建设百万量级高质量的保险知识数据集，并利用大模型慢思考、反思能力解决具体专业问题，建设端到端服务能力的领域专家级 Agent。

成效

建成保险行业首个支撑千亿级大模型调优的算力底座，完成业界主流开源与商用的基础大模型训推适配。当前，已试点上线健康险理赔审核员、寿险代理人培训员、车险在线理赔助手、审计数字员工 4 个数字劳动力，覆盖近千员工。

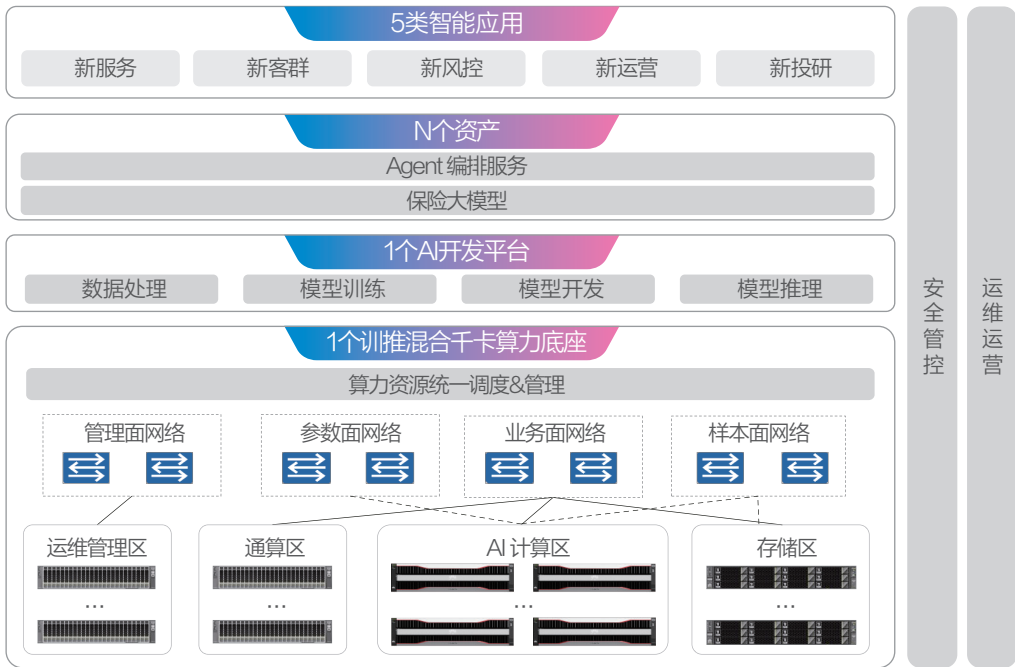


图 4-18 某头部保险 AI 算力平台架构

小型 AI DC

对于集团企业分支、医疗、教育等客户而言，建设小型 AI DC 承载本地业务，是实现业务智能化升级的关键。本节重点描述小型 AI DC 建设过程中面临的关键需求与挑战，并总结小型 AI DC 的关键特征，同时，基于这些特征，提出四大方向性的建议。

关键建设需求

小型 AI DC 的核心需求主要集中在以下五个方面：

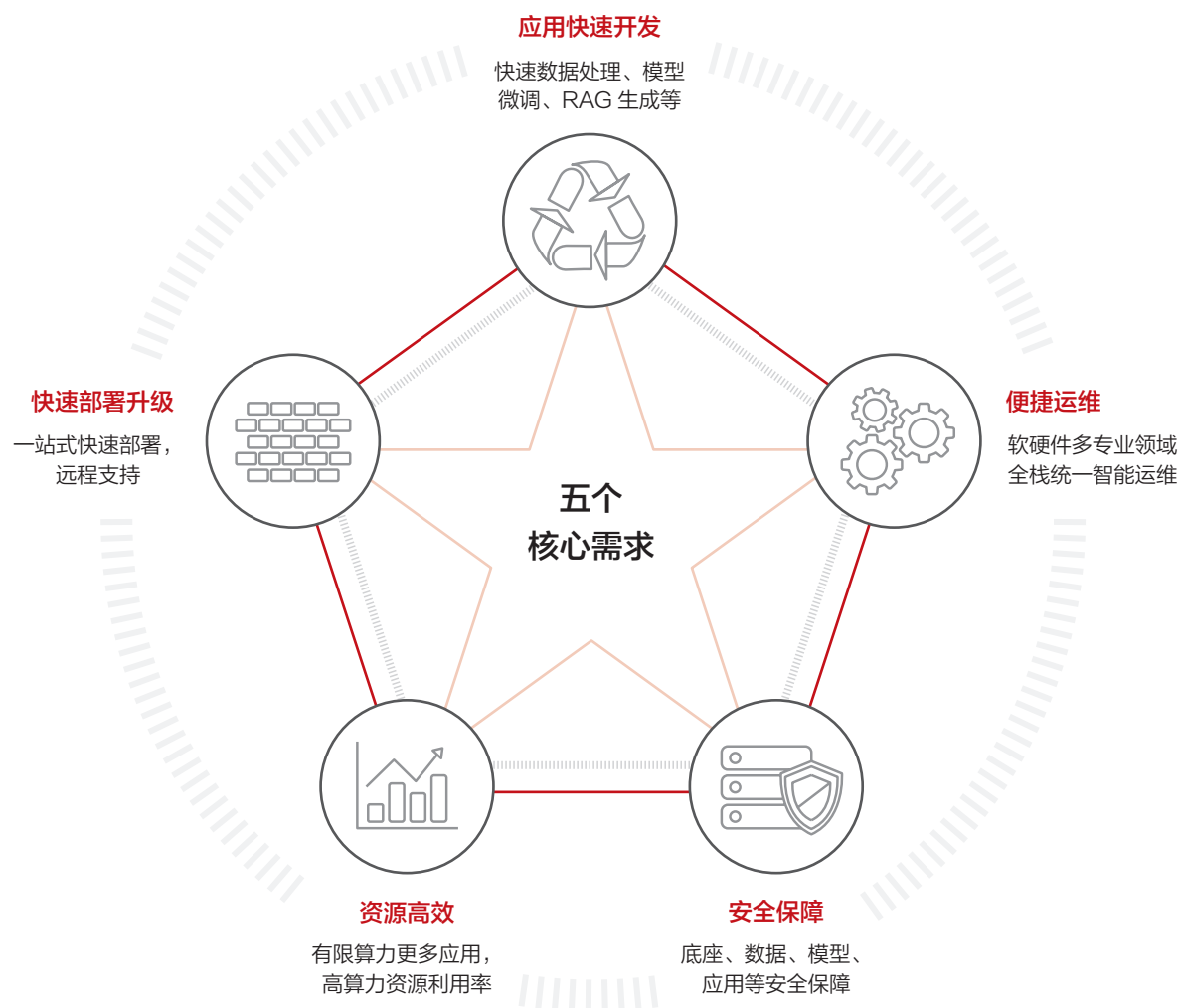


图 4-19 小型 AI DC 的五个核心需求

01

资源高效：鉴于小型 AI DC 受到环境条件的限制，其可用算力资源相对有限，因此在设计上必须保证在有限资源下能承载更多的业务应用，从而最大化算力资源的使用效率。

04

便捷运维：考虑到小型 AI DC 可能缺乏专职运维人员的支持，其设计应当以减少故障发生率为前提，并且在必要时能够支持远程运维操作，以简化日常维护流程。

02

快速部署升级：考虑到部分小型 AI DC 位于偏远地区，企业希望能够在最少的人工干预下完成部署工作，理想情况下，交付团队仅需一次现场访问即可完成所有安装，甚至完全通过远程方式实施部署。

05

安全保障：鉴于小型 AI DC 通常与各种感知设备如智能摄像头、传感器等直接相连，其安全性至关重要。为防止外部入侵，必须采取有效措施确保小型 AI DC 及其关联设备的安全。

03

应用快速开发：对于那些依赖于私有数据集来进行模型微调的专业用户，他们需要一套简便易用的工具链来支持从数据准备、模型训练、推理部署到检索增强生成（RAG）等全流程的操作，使得整个过程如同拖放般简单快捷。

为了满足以上需求，我们提出了一种“**灵快轻易**”的**小型 AI DC 建设理念**。这一理念强调的是数据中心的设计应做到形态灵活、快速部署与升级、轻量极简并且易于管理和维护。



图 4-20 领先的小型 AI DC，要求“灵快轻易”

规划建设方向一：形态灵活（灵）

“灵”即形态灵活多样。按物理形态分节点型和机柜型，按部署形态分独立部署和云边缘部署，按功能形态分承载智能客服、代码生成等应用的 NLP 类，承载工业质检、医疗影像等应用的 CV 类，承载办公助手、文生图等应用的多模态类。

对于形态多样的小型 AI DC 来说，建设标准统一的

轻量化 AI 平台是关键。首先，通过统一的南向接入标准，可以兼容形态多样的硬件底座，同时支持多种端侧设备的接入标准，如视频接入、IoT 接入、智能设备接入等；其次，通过统一的 API 开放标准，供微调、推理、应用快速调用；最后，统一的数据面、管理面标准，可以实现与中心的多维协同，支撑边缘小型 AI DC 的边学边用、持续升级。



图 4-21 标准统一的轻量化 AI 平台

规划建设方向二：快速部署、快速升级（快）

“快”即快速部署、快速升级。支持云边的应用、模型、数据等多维协同，边缘小型 AI DC 通过访问中心云 AppStore 市场，以在线订阅的方式，完成模型和应用下载、部署；中心侧可远程对分散的海量小型 AI DC 进行模型、应用的统一升级；边缘侧实时采集的数据，上传到中心云并用于模型的持续迭代，实现边用边学。

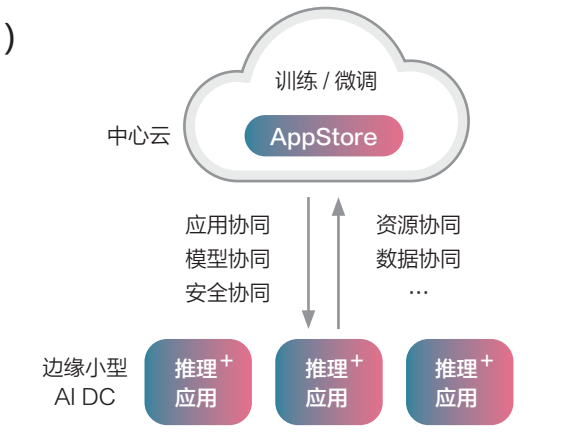


图 4-22 云边协同架构

规划建设方向三：轻量极简（轻）

“轻”即轻量极简。一是负载轻，典型的推理应用并发小于 10 路，模型参数较小，往往需要将算力卡切分为多份，供多个模型同时使用，提升资源使用率。

二是资源轻，1 节点甚至 1 卡起步，按需扩展，资源少造价低。三是管理轻，部署需要的资源少，把有限的资源释放给业务使用，最低希望“1 核 1G”即可部署。

规划建设方向四：易维易用高安全（易）

“易”即易维易用高安全。

独立部署场景：在易用方面，用户需要在边缘侧对场景模型微调，小型 AI DC 集成全套工具链，满足数据处理、模型微调、推理部署、知识库生成等需求，实现 AI 应用的天级上线。在易运维方面，对算存网硬件、软件、应用等全栈统一管理，提供基于任务的可视化运维，实现少人运维。在安全方面，具备模型内容安全、数据安全、设备接入安全等基本能力，满足从数据接入到应用部署的安全要求。

云边缘场景：在易用方面，边缘 AI DC 是中心云的延伸，所有操作均可在中心批量完成；在易运维方面，中心统一管理海量、分散的小型 AI DC，实现边缘侧的无人运维。在安全方面，边侧支持操作系统安全、设备接入认证、安全检测等能力，避免成为全网入侵的薄弱点。

小型 AI DC 综合评价指标体系表

一级指标	二级指标	三级指标	指标描述
灵	模型适配	适配的模型数量	提供模型镜像环境，支持直接下载使用的模型数量，单位是个
快	部署升级	部署时间	机房具备部署条件情况下，包含硬件、OS、管理平台、运维平台等部署所需时间，单位是天
		升级时间	模型和应用更新的时间，若现场升级应包含往返现场的时间，单位是分钟
轻	资源需求	起步节点数	搭建小型 AI DC 最好的服务器节点数，含通用和智算节点，单位是台
		管理资源需求	搭建管理平台所需的 CPU 核数、内存容量需求，单位是核、GB 内存
易	易维	人均运维 DC 数	维护 1 个小型 AI DC 所需的人力数量，单位是站 / 人

01

某三甲医院
打造“灵快轻易”的 AI DC

某三甲医院建设人工智能研发与精准诊疗平台，通过 AI 赋能诊疗全流程，助力医疗服务效率和质量的提升。

挑战

平台搭建繁杂

算存网、AI 平台、数据平台等软硬件产品多，搭建投入大，周期长。

模型微调难、应用构建复杂

数据准备周期长、场景模型微调难，此外还需要 RAG、Prompt 等功能，缺少好用易用的全流程工具。

运维复杂

设备种类多，人员技能要求多，故障定位难，故障恢复时间长。

创新

医院联合讯飞等伙伴，基于华为 FusionCube A3000 超融合一体机方案，通过半自动化数据处理平台、低代码训推工具链、应用开发平台及全栈预验证，实现了应用上线快、模型训推易、知识生成快、开箱即用及统一运维，快速构建了“灵快轻易”的算力平台。

成效

当前，医院已完成电子病历自动生成和病历内涵质控的业务上线，其中电子病历自动生成已实现病例编写时间减少 50%。未来，该平台将支撑每年约 700 万患者门诊和近 30 万的住院患者服务的体验升级。

02

某能源集团打造云边协同的 AI DC
实现边缘的快速部署、快速升级、边用边学

某能源集团推动智能化技术与煤炭产业融合发展，提升煤矿智能化水平，实现集团各煤矿现场的安全、高质量生产。

挑战

海量煤矿边缘缺乏 AI 能力

边缘 AI 单独建设成本高，下属 100 余家三级单位，数量多且分散，部署周期长，升级困难。

模型部署和迭代效率低

模型开发后无法快速部署到边缘，边缘采集了海量的数据，无法及时有效汇聚到中心，模型迭代效率低。

成效

基于云边协同，实现了开发和部署周期从月级到天级、安全生产风险降低及生产效率的显著提升，以年处理 400 万吨焦煤选煤厂为例，精煤回收率提升约 0.2%，全年收益约可增加 500 万元。

创新

构建云边协同体系，使能边缘 DC 构建边用边学的

AI 能力：基于华为盘古大模型，构建了集中心训练、边缘推理、云边协同等功能于一体的人工智能体系。在各三级单位本地部署边缘 DC，从数据获取、推理识别到告警处置的整个业务闭环均在边缘侧完成，推理和处置结果上报至集团中心应用平台，实现多边缘 DC 的统一管控和边缘侧高效反馈。各生产单位的模型，通过部署在集团中心云的人工智能平台进行统一的数据处理、模型训练微调。同时，边缘侧将 AI 误报、存疑样本统一反馈至中心侧统一分析，重新训练升级模型，实现了边用边学、在线学习、持续升级。

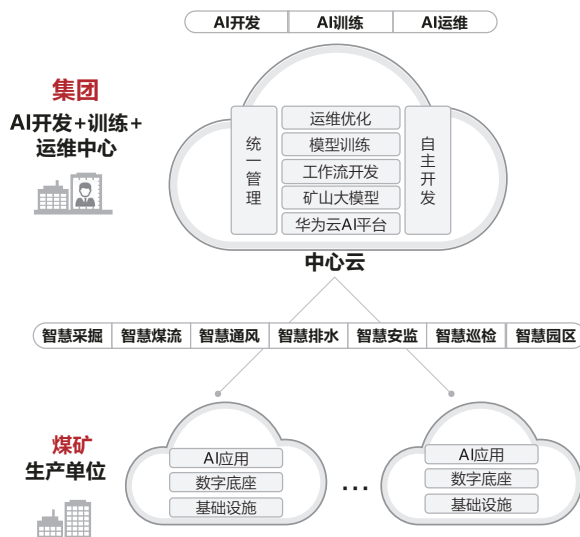


图 4-23 某能源集团 AI 算力平台架构

「第 5 章」

AI DC

建设与发展倡议

行动倡议一

适度超前建设 AI DC

企业在面对 AI 技术的变革时，不应盲目跟风 (FOMO: Fear of Missing Out)，而应深入思考 AI 技术如何真正为企业创造价值，以有序、高效的方式实施 AI 战略，实现业务的持续增长和创新。我们倡议所有企业在落地 AI 过程中，关注以下六大核心要素。其中，AI 基础设施的适度超前建设至关重要，决定了企业能否在 AI 时代抢得先机。

战略决心

智能化创新对于企业来说是一种战略投入，表现为资源投入大，且存在一定不确定性；其次需要多业务和多技术领域的协同。因此企业一把手的关注与支持是确保成功的基石，根据实际情况，下决心组建专项工作团队，确保跨部门、跨业务线的高效协同，为 AI 基础设施的建设和应用场景落地提供全方面的资源保障。



图 5-1 企业落地 AI 重点关注的六件事

场景选择

选择应用场景是 AI 落地的关键起点，直接决定 AI 基础设施的建设规模、性能和可靠性规划，以及面向未来的演进需求。企业应从实际业务需求出发，优先考虑那些痛点突出、技术适配度高的业务场景，如客户服务、供应链管理、产品设计、数据整理等。遵循“先易后难、逐步推进”的原则，初期可选择易于实施、见效快的场景进行试点，积累经验与信心，再逐步拓展至更复杂、更广泛的领域，确保 AI 应用的稳健推进与持续优化。

模型匹配

在推进人工智能技术的企业级应用时，合理选择和组合不同的 AI 模型至关重要。根据不同的分类标准，AI 模型可以被细分为多个类型，如按规模划分，可分为大模型与小模型；按应用场景划分，包括基础模型、行业专用模型以及针对特定领域的模型；按功能专长划分，涵盖自然语言处理（NLP）、计算机视觉（CV）、预测分析、多模态及机理模型等；按训练和部署场景划分，又涉及全量训练、增量学习（二次训练）、模型微调以及最终的推理应用。

从当前行业的实践经验来看，为了满足企业复杂多变的需求，通常需要综合运用多种类型的模型。因此，企业在选择 AI 模型时，应当基于自身的行业定位和业务策略进行考量。同时，通过综合评估技术难度、数据安全与成本效益，制定合理的训练策略，确保模型的精准度与实用性。

底座先行

AI 模型的长稳训练与持续迭代，离不开强大的算力支撑。当前，模型还在不断成熟过程中，存在很大的不确定性，但对算力的需求是确定的，企业应优先构建高性能、高弹性的 AI DC 基础算力平台，采用先进的硬件设备与软件系统，通过软硬协同、训推协同等，确保算力资源的高效稳定供给。同时，依托 AI DC 建立完整的数据治理系统，包括数据采集、清洗、标注、存储与分析等全流程管理，为 AI 模型提供高质量的数据养料。此外，企业还需关注 AI DC 的绿色低碳与能源效率，确保可持续发展。

行业生态

鉴于 AI 技术的垂直整合属性，企业应以 AI DC 为中心，对内聚合产品，对外聚合生态，积极参与或主导构建一个开放、统一、可持续进化的行业生态体系，基于行业大平台进行技术众筹、能力开放与资源共享，在行业内实现统一架构、统一标准和统一数据规范，加速 AI 技术的创新与应用，降低行业门槛，提升整体竞争力，实现多方共赢的局面。

人才培养

AI 要深入垂直行业领域并最终进入核心生产场景，就需要培养一批既能掌握技术细节，又能洞悉应用场景的复合型人才。企业需提前重视 AI 应用相关人才的培养，加大人才引进与内部培养力度，为 AI 落地的持续创新与企业智能化转型提供坚实的人才保障。

行动倡议二

共同实现 AI DC 集约化建设和绿色发展

为了有序推动 AI DC 集约化建设和绿色发展，智算数据中心的规划建设应遵循“东数西算”等指导方向，从算力布局集约、创新技术应用、标准体系制定等方面综合施策，加强对智能算力资源建设与应用的引导和约束，改善 AI DC 发展中存在的高重复建设与低资源利用等问题，促进 AI DC 行业的高质发展。

有序推动规模化集约化发展。一方面，联合企业支持传统数据中心向智算数据中心转型，有序提高智算数据中心规模，促进数据中心规模化发展。另一方面，强化算力集约供给，联合运营商、云服务商和各类算力平台等智能算力与通用算力协同发展，满足均衡型、计算和存储密集型等各类业务算力需求。最后，加强跨地区 AI DC 之间算力资源的协同，实现不同数据中心之间的负载均衡和资源动态调配，从而提高整个系统的弹性和灵活性。

大力夯实绿色能源底座支撑。当前 AI DC 使用绿色能源仍面临着诸多挑战，核心是面向能耗巨大的 AI DC，如何在提高清洁能源电力结构占比的同时保证供电安全稳定，包括微电网、源网荷储、新型储能等技术产业日益受到业界关注。全面夯实绿色算力底层基础，筑牢产业创新发展底座，我们倡议企业围绕电力稳定不间断需求，加强微电网系统研究与构建，开展储能材料探索和技术研发，推动液冷相关技术创新，探索余热回收等。

加速研制绿色算力标准体系。标准化工作是绿色算力产业发展的基石和抓手。尽管在数据中心领域已有关于碳利用效率、碳中和评估及 IT 设备能效的相关标准和规范，但绿色算力仍需构建基于业界共识的完善的标准体系。当前，术语和定义等基础性标准大多尚未完善，能耗评测和分类分级部分仍缺共识，有必要建立和完善绿色算力标准体系，推广其在算力基础设施领域的应用。

行动倡议三

共建开放协作的行业 AI 生态

当一项革命性技术开始迈入加速发展阶段，往往会面临两条道路，**一条是由单一大企业牵头**，无数中小型企业追随，打造一个以单一大企业为核心的完整封闭

生态；**另一条是由许多企业共同发起并优化**，互相兼容，互相促进的多核心的开放生态。



图 5-2 行业 AI 生态建设模式

以著名的“iPhone 时刻”为例，当被重新定义的智能手机面世之后，在全球掀起了智能手机的风潮，完全不亚于如今的“AI 时刻”。在这一巨大的市场机会面前，苹果公司和谷歌公司选择了两条截然不同的发展道路，一个开发了围绕苹果自身体系的 IOS 操作系统，另一个则开发了开源开放的 Android 系统。这两个至今在移动领域占据主导地位的系统，分别代表了单核心封闭生态与多核心开放生态，它们的构建策略和发展现状为 AI 这类新技术的生态构建提供了深刻而生动的参照。

鉴于封闭与开放生态的各自特点，以及 AI 领域所特

有的跨学科融合与复杂性，**我们提倡在垂直行业范围内，企业、主管部门和学术界共同构建开放、协作的行业 AI 生态。**一方面，汇聚产业界、学术界及政府等多领域力量，形成合力，加速关键技术的突破与应用落地；其次，在开放的生态平台上，统一的行业标准与规范得以建立，跨模型、跨企业、跨行业的互联互通成为可能；同时，构建开放的交流与合作机制，共同应对 AI 发展中面临的伦理、法律及社会问题，确保技术发展与应用的正当性与可持续性；此外，促进人才的培养与流动，通过联合培训、竞赛和合作项目，提升行业整体的技术水平与创新能力，为 AI 领域的持续繁荣提供人才保障。

行动倡议四

筑好三个底座，加速行业 AI 走深向实

AI 正以前所未有的速度颠覆大家对未来企业和市场运行的认知。然而，AI 的核心价值并非单纯的技术展示，而在于如何加速落地，深入企业业务流程，切实解决关键问题，为企业创造实质性的商业价值。因此无论发展初衷与过程如何，AI 技术的最终目的是将技术转化为生产力，驱动产业升级与经济增效。

我们认识到，在复杂多变的市场环境与技术挑战下实现 AI 的最终目的，AI 行业还需要一个稳固的战略框架来筑实基础。正是基于这一核心价值主张，我们提出了三大战略底座——解决方案底座、生态底座与人才底座，旨在为 AI 行业提供坚实的支撑，确保技术应用能够精准对接企业需求，最终实现行业 AI 的走深向实，实现价值最大化。



图 5-3 筑好三个底座，加速行业 AI 走深向实

解决方案底座：围绕 AI DC 构建坚实的算力基础设施

解决方案是技术实践的基石，建立这一底座的目标就是基于统一的技术架构、统一的行业标准、以及统一的数据规范，夯实以 AI DC 为中心的企业 AI 基础设施，其核心在于通过定义 AI DC 目标参考架构，优

化算力效能，标准化技术架构，实现生态的归一化，以此促进技术与业务目标的无缝对接，从而加速技术创新与应用落地。

生态底座：构建开放协作的产业和行业生态圈

生态底座强调的是构建一个开放、共享且紧密相连的生态协作系统，在这个系统中的每个环节协同发展，通过产业和行业各领域的联合创新，深入企业核心生产场景，共同推动 AI 在千行万业的应用落地。一方面，强化产业生态的协同效应，通过硬件制造商、软件开发

加速技术迭代升级，提升 AI DC 整体系统的性能和效率；另一方面，构建开放协作的行业生态圈，联合同行业内的不同企业，共同探讨和解决行业面临的关键问题，推动技术创新和标准化进程，促进成员间的信息共享和技术交流，使得 AI 技术能够更广泛地渗透到各行各业，为社会创造新的价值。

人才底座：以应用为牵引，技术为支撑，发展产业人才和行业人才

多层次 AI 人才的培养与储备，是 AI 行业持续增长的源泉。为了建设 AI 领域的人才底座，需要一方面重视技术型人才的培养，特别是精通底层算力开发、行业加速库构建的专业人员，使这些人才成为推动技术前沿探索的关键。另一方面，通过 AI 的落地，壮大垂直行业的 AI 应用专家队伍，比如行业的 AI 顾问和学术领域的教育工作者，他们能够根据实际情况，将 AI 技术与行业需求相结合，创造实际价值。我们还需要倡导产教融合的新模式，向开发者提供精细的针

对性培养方案，使其熟练掌握异构算力的应用，为各行各业输送源源不断的创新力量。

三大战略底座的构建，不仅为 AI 应用的落地提供了稳固的支撑体系，也为智能化行业的长期发展打下坚实基础。通过夯实解决方案、生态以及人才底座，AI 将赋能千行万业快速迈进一个更加成熟、更具竞争力的新时代。

结语

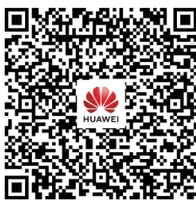
数智化转型大潮中，行业部署数据中心（DC）面临诸多困惑和挑战。而随着 AI 大模型驱动的智能时代到来，又进一步增加了对智算数据中心（AI DC）的需求。

当前，AI DC 的规划与建设缺乏全面、系统化的指导框架。华为联合业界，通过过去几年的深入研究，关键产品与解决方案的创新，并结合大量的行业探索与建设实践，汇聚了二十多位专家的智慧，整合了国际组织的统计数据以及业界研究机构的成果，并在十几场研讨交流中持续沉淀，形成了这份具有重要参考意义的白皮书。

此白皮书共分为五个章节。首先，阐述了 AI 总体愿景与宏观驱动力，以及 AI 如何引发百年未有的大变革，并重塑各行各业的发展路径。其次，白皮书深入

分析了企业发展 AI 过程中面临的确定性与不确定性，提出了“场景、模型、数据和算力”四位一体的企业 AI 落地指导纲领。白皮书详细介绍了从传统数据中心向 AI DC 发展的趋势与变化，并定义系统摩尔、能基木桶、迭代式平台、编排式应用以及生成式安全等五大特征；在此基础上，详细阐述了超大型、大型和小型等不同类型的 AI DC 规划建设的关键方向，提出综合评价指标体系，为高效高质量建设 AI DC 提供重大参考价值。最后，白皮书建议决策者在拟定智算数据中心的部署投资方向时，应秉持着适度超前、绿色集约化建设，以及开放协作的理念，同时兼顾解决方案、生态和人才三大底座的建设。

希望白皮书能够助力行业在数智化转型的道路上稳步前行，通过构建强大而坚实的智算数据中心，使能千行万业迈向更高效、更智能的未来。



扫码下载白皮书

商标声明

 HUAWEI, HUAWEI,  是华为技术有限公司商标或者注册商标，在本手册中以及本手册描述的产品中，出现的其它商标，产品名称，服务名称以及公司名称，由其各自的所有人拥有。

免责声明

本文档可能含有预测信息，包括但不限于有关未来的财务、运营、产品系列、新技术等信息。由于实践中存在很多不确定因素，可能导致实际结果与预测信息有很大的差别。因此，本文档信息仅供参考，不构成任何要约或承诺，华为不对您在本文档基础上做出的任何行为承担责任。华为可能不经通知修改上述信息，恕不另行通知。

版权所有© 华为技术有限公司 2024。保留一切权利。

非经华为技术有限公司书面同意，任何单位和个人不得擅自摘抄、复制本手册内容的部分或全部，并不得以任何形式传播。