



2024 版

数据中心 2030



构建万物互联的智能世界

数据中心 2030

探索未来数据中心
引领智能时代

汪涛

创新涌现，拥抱智能时代

在 AI 大模型训练过程中，当模型大到一定规模之后，性能会发生突变，开始呈现指数级快速增长，科学界称这个现象为“涌现”。正是这个性能的突变，让人工智能的发展阶段从感知理解世界到生成创造世界，这也造就了 ChatGPT 的火热，催生了面向行业的数百个 AI 大模型的出现。今天，“百模千态”正走向每一个行业、每一个场景、解决客户每一个问题，加速千行万业的智能化转型。人工智能的“涌现”时刻即将出现，人类社会也将迎来一个波澜壮阔的智能时代。

迈入智能时代，最大的需求是算力，最关键的基础设施是数据中心。根据华为《智能世界 2030》报告预测，2030 年，人类将迎来 YB 数据时代，全球通用计算算力将达到 3.3ZFLOPS(FP32)，AI 算力需求激增，2030 年将达到 864 ZFLOPS (FP16)。算力需求十年百倍的增长将成为常态。数据中心作为人工智能、云计算等新一代信息通信技术的重要载体，已经成为新型数字基础设施的算力底座，具有空前重要的战略地位，堪称“数字经济发动机”。

一边是算力需求以远超摩尔定律的陡峭增长，而另一边却是多重的资源约束。单芯片摩尔定律的失效、以及全球可持续发展目标下对于碳减排的要求，将迫使未来的数据中心必须在更优的计算架构、以及更低的能耗下产生更大的算力。回顾信息通信技术产业的发展历史，每一次跃升都是矛盾驱动的结果。过去三十年，超大宽带与成本约束的矛盾推动了联接产业的高速发展，5G、F5G 等改变世界；未来三十年，将是超强算力与资源约束的矛盾推动计算产业的高速发展，AI、云计算等重塑世界。可以预见，应对算力需求和资源约束的主矛盾，系统级和架构级的技术、产品和方案创新必将涌现，也将成为未来数据中心发展的主旋律。

选择方向和路径已成为一种能力和智慧。站在 21 世纪第三个十年的起点，我们看到 ICT 产业正面临巨大的发展机会，世界正全面进入数字化和智能化，那么，2030 年的世界将是一番怎样的景象？2021 年 9 月份，华为发布了《智能世界 2030》主报告及相关系列报告，而《数据中心 2030》是最新的系列之一。

基于对未来的不懈探索，过去三年间，与业界数百名学者、客户伙伴及研究院机构等深入交流，集业界专家和华为专家的智慧，输出了我们对数据中心下一个十年发展的思考——《数据中心 2030》报告。

该报告从算力需求与资源约束的核心矛盾出发，描绘了未来十年影响数据中心发展的五大未来场景，提出了围绕“数效、人效、算效、能效和运效”等五效提升的发展方向；同时，该报告在业界首次定义了未来数据中心的技术特征，系统性阐述了数据中心所涉及到的云服务、计算、存储、网络、能源等全栈技术可能的挑战与创新方向，并明确提出了未来数据中心建设的参考架构。希望这份报告能为全球数据中心基础设施的建设与发展、为全球数字经济腾飞贡献出积极力量。

从万物互联到万物智能、万智互联，一个更加美好的智能世界在向我们招手，但未来注定是不平凡的。吴军在《智能时代》中提到，在历次技术革命中，一个人、一家企业、甚至一个国家，可以选择的道路只有两条，要么加入浪潮，成为前 2% 的人；要么观望徘徊，被淘汰。毫无疑问，未来 10 年将充满根本性的突破和改变世界的惊喜，每一个主要行业很快将会被重塑。

人们总是高估了未来一二年的变化，却低估了未来十年的变革，而低估未来变革的影响是因为没有“看见”，这正是《数据中心 2030》的意义所在。大胆假设和最好预测是创造未来的辩证关系，在迈向未来的道路上，仍有大量的挑战需要跨越，让我们携起手来，勇于探索、持续创新，共同拥抱智能时代！



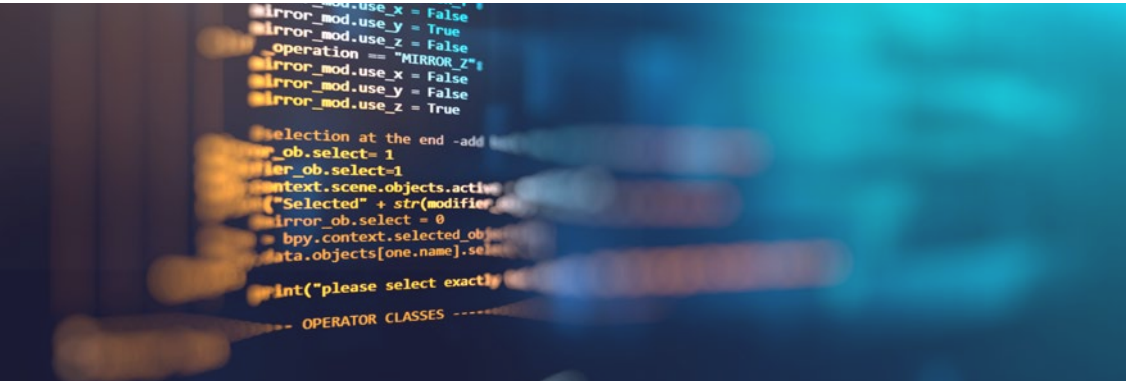
华为公司常务董事
ICT 基础设施业务管理委员会主任

目 录

01 产业趋势	07
------------	----

02 未来场景与创新方向	13
-----------------	----

AI for All，创造新生产力.....	14
科研第四范式，以数据密集型计算探索未知.....	15
空间互联网，带来多维虚实交互体验.....	16
行业数字孪生，推动智能升级.....	16
普惠云原生，消除企业数字鸿沟.....	17
系统化多流协同，提升能效.....	17
多级化软硬协同，提升算效.....	18
无损化网业协同，提升运效.....	19
社会化数据协同，提升数效.....	19
智能化人机协同，提升人效.....	21



愿景.....24

关键技术特征.....25

1. 多样泛在.....25	2. 安全智慧.....30	3. 零碳节能.....36
1) 大集群.....25	1) 高安全.....30	1) 绿供电.....36
2) 轻边缘.....26	2) 高可靠.....34	2) 新储能.....37
3) 新型态.....28	3) 高智能.....35	3) 液制冷.....39
4. 柔性资源.....41	5. 对等互联.....50	6. 系统摩尔.....57
1) 全池化.....41	1) 超融合.....50	1) 大小芯.....57
2) 柔计算.....44	2) 高性能.....51	2) 新算力.....58
3) 泛协作.....48	3) 光内生.....53	3) 新存储.....60

新基础设施，供电制冷走向全天候绿色零碳.....67

新算力底座，构建以数据为中心多样算力系统.....68

新资源调度，应用为中心实现柔性调度.....68

新数据管理，数据全局可视助力高效流通.....70

新协同服务，开放架构融入社会化算力.....70

新智能管理，AI 驱动实现 DC 自动运维.....71

附：关键预测数据指标体系.....78

附：缩略语.....80





产业 趋势

01



数据中心作为人工智能、云计算等新一代信息通信技术的重要载体，已经成为新型数字基础设施的算力底座，具有空前重要的战略地位，堪称“数字经济发动机”。展望 2030，数据中心的未来发展将呈现如下几个趋势：

■ 算力需求十年百倍增长，算力分布进一步极化

根据华为《智能世界 2030》报告预测，2030 年，人类将迎来 YB 数据时代：全球通用计算算力将达到 3.3ZFLOPS(FP32)，AI 算力需求激增，2030 年将达到 864 ZFLOPS (FP16)。全球数据中心产业正进入新一轮快速发展期，我们预测，未来

三年内，全球超大型数据中心数量将突破 1000 个，并将保持快速增长；同时，随着自动驾驶、智能制造、元宇宙等应用的普及，边缘数据中心将同步快速增长，根据第三方预测，2030 年部署在企业内的边缘计算节点将接近 1000 万个。

■ 算力的规模和效率成为国家和企业的核心竞争力

如同农业经济的核心竞争力是建立在从劳动力人口到大规模水利设施再到机械化持续提升生产效率的基础上一样，算力的规模和效率也已经成为发展数字经济的核心竞争力。当前全球正处在千行万业智能化转型的新阶段，“百模千态”的 AI 大模型成为发展焦点，据预测 GPT5.0(Generative Pre-trained Transformer) 训练集群的算力需求将达到 GPT3.0 的 200-400 倍。几乎所有的基础科学和大工业都朝着多维度、高精度的大规模数据分析方向发展：如石油勘探领域深度偏移等场

景下单位面积勘探区的算力需求将增长 10 倍以上。AI、区块链等技术支撑的行业智能化场景也将带来算力需求的爆炸式增长，从数字化球拍每一次挥动的感知、记录和处理，到普惠金融每一次微型交易的客户画像、信用评估，都需要高效算力的支持。未来各行业在算力领域的投资占比将快速增长，以银行业为例，根据有关预测 2024 年中国银行业技术投入总规模将超过 4000 亿元，其中 AI 与云计算是重点投资领域，二者占比超过总投入的一半。

■ AI 驱动数据中心发生全景式革命

华为预测，到 2030 年全球 AI 计算算力将超过 105 ZFLOPS(FP16)：AI 计算算力成为数据中心发展的最大驱动力和决定性因素。未来 5 到 10 年通用大模型的发展有可能使 AI 对文字、音乐、绘画、语音、图像、视频等领域的理解力超过人类平均水平，并与互联网和智能设备深度融合，深度改变全社会的消费模式和行为。AI 技术与生产率之间显著的“扩散滞后”效应逐渐减弱，通用大模型能力将嵌入生产力和生产工具、行业

大模型和场景化 AI 等多路径融合，AI 技术创新对商业价值的影响将变得更加广泛和不可预测。通用大模型多模态泛化下的训练算力需求将保持远超摩尔定律的陡峭增长趋势，需要数据中心在算力规模、架构、算法优化、跨网协同等领域持续创新和快速迭代。展望未来，AI 的发展将加速平台型企业超级数据中心和国家级算力网络的建设。

■ 数据中心的产业标签从高耗能转变为绿色发展使能器

数据中心总耗电量在 ICT 行业占比超 80%，为保障数据中心行业的可持续发展，首先需要提升能源使用效率、实现绿色低碳。多个国家、国际组织发布数据中心相关政策，如美国政府通过 DCOI 数据中心优化倡议，要求新建数据中心 PUE 低于 1.4，老旧改造数据中心 PUE 低于 1.5。欧洲数据中心运营商和行业协会在《欧洲的气候中和数据中心公约》中宣布 2030 年实现数据中心碳中和。中国出台《全国一体化大数据中心协同创新体系算力枢纽实施方案》推动构建全国一体化大数据中心，启动“东数西算”工程，促进数据中心绿色可持续发展，加快节能低碳技术的研发应用，要求到 2025 年新建大型数据中心

PUE 低于 1.3。未来，随着各国相关政策的陆续出台和技术的持续发展，越来越多的先进节能技术将更广泛地应用到数据中心，推动 PUE 的进一步下降，预计到 2030 年，PUE 将进入 1.0x 时代。未来随着风光水等清洁能源占比的不断增加，通过数据中心微电网“源网荷储”的协同还可以进一步降低碳排放，实现数据中心的绿色零碳目标。其次除了自身降低碳排放之外，数据中心还可以为其他行业的智能化转型赋能，成为全社会降碳的使能器，据全球电子可持续性倡议组织（GeSI）预测，到 2030 年 ICT 技术通过使能其他行业，将帮助减少全球总碳排放的 20%，是自身排放量的 10 倍。



■ 超出物理数据中心边界，多流协同的数据中心普及化

一方面，规模化、中长期需求预测困难、技术迭代加速等成为所有骨干数据中心运营企业和领先数字化企业的共同挑战。数百万台服务器的云数据中心、数十万台服务器规模的行业数据中心将在 2030 年之前出现，ChatGPT 等突发的巨型超高密度任务涌现，土地、能耗获得的不确定性等因素使得基于超大单体、以 10 年为周期的数据中心规划模式难以为继。未来分阶段、模块化、集群化、服务化，逻辑上统一，物理上分布的数据中心新建设模式将逐渐普及。另一方面高性能计算的需求也随之不断提升，影视渲染、效果图

渲染等批量计算任务，基因测序、风机工况模拟等科学计算任务以及 AI 训练等可并行的计算任务，往往需要消耗大量的算力资源和运算时间。这类任务往往具有计算成本敏感、实时性不敏感、计算规模可变动的特点，针对这类需求可以通过实时传递价格信号，激励用户选择电力价格较低的时间段进行整体运算；也可以通过断点续训、可续渲染技术，在计算任务执行的过程中暂停乃至对并行规模进行改变，来平移和升降电力负荷。通过任务流、信息流、能量流的精准关联和多流协同，构建绿色低碳、算效领先的数据中心。

■ 系统级创新成为数据中心技术发展的主流

蚂蚁大脑一般只有 0.2 毫瓦的能耗，但是能够做很多复杂的事情：可以筑巢、寻找食物、养蚜虫等。相比之下，目前自动驾驶汽车还需要几十瓦甚至几百瓦来进行计算，在能效上与生物界相比还有很大的差距。应对十年百倍算力增长需求与能耗约束之间的矛盾，未来数据中心需要打破冯·诺依曼架构，基于新架构、新部件发展适应性与高效性的新计算模式。在信息计算领域，已经发展出了十几种广泛使用的计算模式，例如无线和光通信里大量使用基于快速傅里叶变换的蝶形计算模式；路由器里大量使用基于逻辑状态转移的有限状态机计算模式；在智能计算领域，除了统计计算之外，业界正在研究数理逻辑计算、几何流形计算、博弈计算等更高效的新计算模式，实现在特定场景下，计算能效的百倍提升。下一代数据中心还将构建“算存网安”多技术协同的全新系统，打破传统计算设备面临的功耗墙、I/O 墙、存算墙约束，从单设备到集群化、从单节点到

网络化，通过系统级创新、软硬协同实现数据中心效率的大幅提升。



■ 围绕算力供给和资源约束挑战持续创新突破

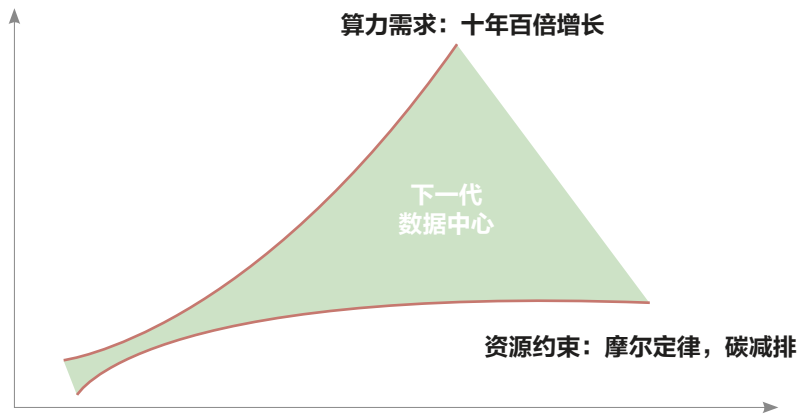


图 1-1 算力需求与资源约束挑战

面向 2030，数字经济加速发展对于算力的需求将呈现十年百倍的指数级增长；而与此同时，单芯片摩尔定律的失效，以及全球可持续发展对于碳减排的硬性要求，将成为制约数据中心未来发展的主要因素。可以预见，围绕算力需求和资源约束挑战的创新将成为未来数据中心发展的主旋

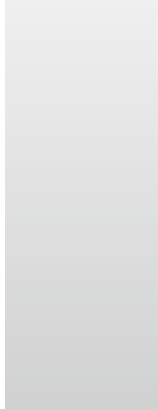
律。先进数字化企业和数字化国家，将在单个数据中心、数据中心集群、数据中心之间的“微中宏”观、多层次进行系统化创新，实现企业级或者国家级的“一台计算机”，通过整体效率的提升将算力供给和资源约束之间的剪刀差最大化，加速迈向智能世界。





未来场景 与创新方向

02



数据中心几乎涉及信息生活的所有方面，从科学研究的突破创新到生产生活的智能高效，都需要数据中心提供更强大的算力，处理更多的数据，算力需求将呈现远超摩尔定律的陡峭增长。与此同时，为了应对算力需求和资源约束的主矛盾，围绕效率提升持续创新，必将成为未来数据中心发展的核心方向。

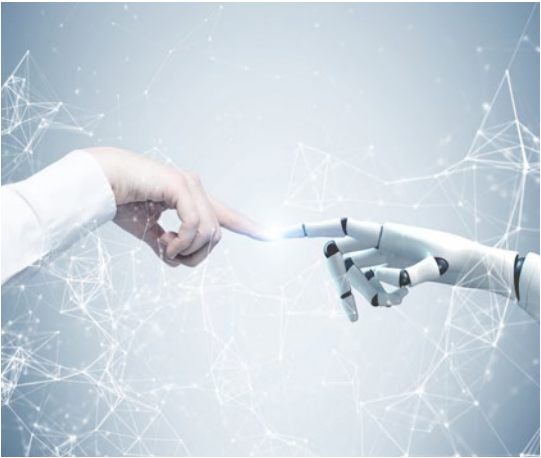
■ AI for All，创造新生产力

过去，人类在科学的边界之内，不断发现万物规律，并创造生产工具，推动社会从农耕文明、工业文明进入到数字文明的数字化阶段。未来，AI 以新的生产力形式出现，在人类定义的边界之内，以更高的效率和更快的速度进行分析和创造，将数字文明带入智能化阶段。

人类善于分析，但 AI 可能做的更好。“分析型 AI”已经得到广泛应用，可以分析一组数据，一组图片，并在其中找到模式，用于多种用途，无论是预防欺诈或是目标识别。

人类擅长创造，但 AI 可能做的更快。随着“生成式 AI”的快速发展，AI 已经开始创造有意义和美丽的东西，如写诗、绘图，并且效率更高。生成式 AI 在图像生成领域的进展来自扩散模型（Diffusion model）的应用，是一种从噪声中生成图像的深度学习技术。在自然语言处理（NLP）领域的进展来自于 ChatGPT，这是一种基于互联网可用数据训练的文本生成深度学习模型，用

于问答、文本摘要生成、机器翻译、分类、代码生成和对话的 AI。在代码生成领域的进展则来自代码生成系统 AlphaCode 和 Copilot。2022 年 2 月，DeepMind 推出了他们的最新研究成果 AlphaCode。它是一个可以自主编程的系统，在 Codeforces 举办的编程竞赛中，超过了 47% 的人类工程师，这标志着 AI 代码生成系统，首次在编程竞赛中，达到了具有竞争力的水平。



AI 技术正加速进入千行万业，如在气象行业，利用 AI 大模型能够在 10 秒内给出未来七天的天气预测结果，对比传统的 HPC 数值预报方法，在预测速度上提升了 10000 倍以上；在证券行业，某金融企业基于 AI 大模型实现了准确率达 90% 的企业财务智能预警，较传统机器学习模型准确率提升了 11%。AI 大模型正逐步从智能对话、短文创作、图片生成等消费应用场景，扩展到办公、编程、营销、设计、搜索等商业应用场景，并将进一步扩展到金融风控、智能客服、辅助诊

断、医疗咨询等企业应用场景，为千行万业注入新生产力。

人类正在从分析型 AI 理解世界迈向生成式 AI 创造世界。面向 2030，具备认知能力的 AI 像我们熟悉的土地、植物、空气、阳光一样无处不在：

“一辆会自己行驶的汽车、一个会自己做饭的机器人、一个会自己管理的通信网络、一个会自我优化的软件平台”将会成为人们日常生活的一部分，并支撑着人类文明的持续进化。

■ 科研第四范式，以数据密集型计算探索未知

几千年前科学以归纳为主，通过观测和实验来描述自然现象；过去数百年出现了理论研究分支，利用数学模型进行分析；过去数十年出现了计算分支，针对复杂问题，使用计算机进行仿真分析；21 世纪初期，新的信息技术已促使新范式的诞生，即基于数据密集型科学研究的“第四范式”，通过将理论、实验和计算仿真统一起来，由仪器收集或仿真计算产生数据、由软件处理数据、由计算机存储信息和知识、科学家通过数据管理和统计方法分析数据和文档。

元之间如何连接与工作，将带来每秒高达 100TB 的数据吞吐量；自动驾驶车辆每天将产生数十 TB 数据用于训练视觉识别算法；用电子显微镜重建大脑中的突触网络，1 立方毫米大脑的图像数据就超过 1PB；而天文专家需要从数十 PB 海量数据中分析发现新天体。PB 级数据使我们可以做到没有模型和假设就可以分析数据，将数据丢进巨大的计算机集群中，只要有相互关系的数据，统计分析算法可以发现过去的科学方法发现不了的新模式、新知识甚至新规律。

数据密集型科学研究，将产生海量数据需要分析处理，如模拟脑神经网络，探索人脑上亿个神经

科学数据已成为科学研究的关键成果和重要的战略性资源，面对喷薄而出的数据需求和数据量，

分类	时间	研究方法	模型
第一范式：经验科学	18 世纪以前	以归纳为主，带有较多盲目性的观测和实验	实验模型
第二范式：理论科学	19 世纪以前	以演绎法为主，不局限于经验事实	数学模型
第三范式：计算科学	20 世纪中期	对各个科学学科中的问题，进行计算机模拟和其他形式的计算	计算机仿真模型
第四范式：数据密集型科学	21 世纪初期	利用数据管理和统计工具分析数据	大数据挖掘模型

如何存储、管理、共享这些科学数据，成了全球科学家关注的热点，也是下一代数据中心的重要应用场景。当这些规模计算的数据量超过 1PB 时，传统的存储子系统已经难以满足海量数据处理的读写需要，数据传输 I/O 带宽的瓶颈愈发突出。而简单地将数据进行分块处理并不能满足数据密集型计算的需求，并与大数据分析的初衷相违背。

目前许多具体科学研究中所面临的最大问题，不是缺少数据，而是面对太多的数据，却不知道如何处理。

当前的超级计算机、计算集群、超级分布式数据库、基于互联网的云计算等并没有完全解决这些矛盾，计算科学期待一次全新的革命！

■ 空间互联网，带来多维虚实交互体验

虚实融合是下一代互联网发展的主要方向，一个具有高沉浸式交互体验的虚实融合的多维空间，将极大地提高人与信息的交互体验和经济活动效率。

虚实融合的发展包括两个方向：

一是由实向虚，基于虚拟世界对于现实世界的模仿，通过构建沉浸式数字体验，增强现实生活的数字体验，强调实现真实体验的数字化。在移动互联网时代，主要通过文字、图片、视频等 2D 形式建立虚拟世界，而未来在元宇宙时代，将真实物理世界在虚拟世界实现数字化重造，建立虚拟化，具备多维交互能力的虚拟世界。

二是由虚向实，超脱对于现实世界的模仿，基于

虚拟世界的自我创造，不但能够形成独立于现实世界的价值体系，还能够对现实世界产生影响，实现数字体验的真实化。如增强现实游戏通过设置与品牌联动特定地点发放限量购物券的方式，帮助品牌方吸引消费者关注，实现数字体验对真实消费的带动。

从技术层面来看，虚实融合的多维互动体验离不开计算机图形图像的多维空间计算能力支持和低延迟网络服务。同时，它还需要强大的人工智能认知能力的辅助，以及泛在通达的数据连接，计算和网络的能力将直接决定了虚实融合的深度和广度。

■ 行业数字孪生，推动智能升级

面向千行万业的数字孪生是数据中心的重要应用场景。根据第三方预测，全球数字孪生市场空间的年复合增长率将达到 40.1%，预计到 2030 年将达到 1310.9 亿美元。数字孪生涉及建模、感知、仿真、渲染、大数据、人工智能等新一代信息技术的综合集成应用，是数字经济发展的重点领域之一。

伴随各行业智能化的推进，城市、制造、交通、水利、能源的数字孪生应用需求快速增长，从端云两侧同时拉动数据中心算力需求。其中，基于 WebGL 的数字孪生应用快速发展，带来终端的升级需求；基于云渲染的数字孪生应用，带来云端算力的快速增长。应对算力快速发展的需求，应从加强算力供给、提升集约化利用水平、加强渲染算法研究等方面推动算力产业升级。

■ 普惠云原生，消除企业数字鸿沟

过去 10 年内，智能手机和移动互联网重塑了人类生活方式和企业生产模式；今天，智能化和电气化正在重构汽车行业的核心竞争力和生态。重塑和重构的背后是强大的算力、算法和数据构成的数据智能，是敏捷迭代、弹性伸缩、韧性自愈的云原生的 IT 系统。未来随着大模型 AI、万物互联、社会化数据协同和数字孪生的新技术推动，与现实世界结合更紧密的千行万业也将快步进入云原生为基础的智能世界。

各行业的领先者和现有分工的颠覆者正在凭借前瞻性思维实现更深层次的智能化，推动云原生特征明显的信息技术和运营技术的融合，赋予产品、

流程、组织精细化、敏捷化的全新竞争力。随着数字系统越来越复杂、发布变更频度越来越高、算力越来越密集、分布越来越广泛，企业将越来越依赖平台能力，越来越多的企业将全面拥抱云原生技术。

普惠化的云原生技术给传统的企业甚至个体带来将生产、经营活动现代化的机遇，消除数字化鸿沟，提供简单、经济而又专业、个性的智能化路径。当云端算力、数据服务 API、涂鸦化的 IOT 控制流程设计、商品化的行业 AI 算法组合时，每个拥抱变化的企业获得与领先者同步的智能化能力。



■ 系统化多流协同，提升能效

在全球积极应对气候变化目标下，绿色低碳成为数据中心的重要发展方向，大部分国家或地区均在单体数据中心领域相继发布了相应政策。中国在充分论证研究基础上，规划布局了 8 大算力网络国家枢纽节点，引导大规模数据中心适度集聚，通过实施“东数西算”工程，积极探索构建形成以数据流为导向的新型算力网络格局。

围绕绿色可持续发展，数据中心相关企业已经开

发了大量创新技术来实现基础设施建设与运营过程中的高效化和低碳化，并且已经在现有或新建的数据中心中实施。如苹果公司在数据中心范围内部署分布式太阳能、风能、沼气等可再生能源发电设施，以及与可再生能源电站签署长期采购协议，为自有数据中心供电，通过一系列措施实现 数据中心使用 100% 可再生能源。微软公司在智能云绿色数据中心建设时提出需要在选址、建设及运营的全流程将数据中心的“能源流”“数

据流”“业务流”有效协同起来，实现绿色高效。华为云贵安数据中心采用自然冷却技术，包括直通风制冷和部分高密度服务器就近利用湖水散热，并通过余热回收利用技术等将数据中心的热量进行采集，用于办公区取暖，在设计中既充分结合了贵州自然条件的优势，也融入了绿色低碳

的可持续发展理念。

实现“能源流”、“数据流”和“业务流”的多流协同，是面向 2030 年构建高能效数据中心的关键。

■ 多级化软硬协同，提升算效

算力的发展经历了单核、多核、网络化三大阶段。综合考虑技术和商业可行性，单核硅基芯片的计算能力将在 3 纳米达到极限。由于经济性原因，依靠增加核数换取算力提升的模式，也将在 128 核后迅速失效。这将推动算力架构从单设备多核向多设备网络化演进。此外，受网络技术发展及网络带宽成本约束，边缘的算力部署也将成为数据中心新的场景，最终形成云边泛在、多级化算力部署的新架构。

过去半个世纪，集成电路产业在摩尔定律的指引下飞速发展，算力一直保持着大跨度提升。在硬件主导算力快速提升的时代，计算过于依赖底层算力，对架构和代码优化重视不足，高级语言不断出现，程序执行效率越来越低，而这恰恰为未

来从“软硬协同”层面提升计算性能留下了优化空间。主流芯片和设备厂商已经纷纷开始通过软硬联合优化来提升整体计算性能。业界认为，硬件架构的每一个数量级的性能提升潜力，通过“软硬协同”能带来两个数量级的整体性能提升。华为云的异构计算服务，通过软件优化硬件直通能力，能够显著降低因计算资源虚拟化造成的性能损耗。图灵奖得主 David Patterson 曾提出，在计算领域，未来十年，我们将看到比过去 50 年更多的架构优化和性能提升的创新。

面向 2030，通过中心集群软硬协同优化、云边多级算力资源协作等提升算效是数据中心未来的重要发展方向。



■ 无损化网业协同，提升运效

未来数据中心的发展对网络将提出更高的要求。传统网络在业务配置和资源管理上不具备足够的灵活性，造成数据中心内和数据中心之间的算力资源利用率低，从而产生巨大的浪费。尤其是在 AI 大模型训练的场景下，需要用到大量的数据，模型参数也会变得非常大，为了让训练效率更高，往往需要上百张 GPU 卡来放置一个大模型作为一个数据并行组，训练大模型的时候往往需要很多个这样的数据并行组来缩短训练的时间。当 GPU 数量扩展到成千上万的时候，性能不仅取决于单一 GPU，也不仅取决于单一服务器，而是要取决于网络的性能。

构建高性能网络，提高数据在计算、存储之间的搬运效率（运效），除了带宽和时延之外，最重

要的是在数据包转发过程中实现无损化，即不允许出现数据包的丢失。业界实验发现，数据每丢失千分之一，计算性能就会下降 30%。

为了实现网络无损化，网络和计算、存储业务系统之间的协同越来越重要。在数据中心内，业界已有厂家在分布式存储、集中式存储、高性能计算等场景下实践了“网存协同”和“网算协同”的创新方案。在数据中心之间，领先的电信运营商也提出了算力网络的方案，在感知应用、感知算力需求的基础上，利用全光、端到端切片、弹性调度等技术，针对分布式存储、跨节点分布式计算等场景，提供零丢包的业务保障能力，为算力之间构建全程全网高效的无损网络。

■ 社会化数据协同，提升数效

生产要素反映了人类社会不同发展阶段的生产力水平。数据作为新型生产要素，是数字化、网络化、智能化的基础，已快速融入生产、分配、流通、消费和社会服务管理等各个环节，深刻改变着生产、生活和社会治理方式。数据的规模爆发式增长，不仅在数字经济发展中的地位和作用凸显，而且对传统生产方式变革具有重大影响：将催生新产业、新业态、新模式，成为驱动经济社会发展的关键生产要素。

在产业数字化方面，社会化将打破企业边界，使获取数据、运用数据的能力成为业务创新和提升用户体验的关键。销售平台可以根据买家的浏览记录做出精准推送以提高销量；制造企业可以通过分析生产流水线数据对生产情况及时做出调整

以提高生产效率；家居公司可以通过分析客户的生活习惯数据创造“智慧家庭”以提高生活质量，种种应用展示出数据在被有效的挖掘、整合后可能产生巨大的价值。业界有观点认为，数



据将逐渐成为与人、技术、流程同样重要的第四大核心竞争力。跨企业边界的数据共享和交换在当下已经比较流行，未来主要的变化是多领域数据汇聚、AI 集成、隐私保护和交易化。以普惠金融的农户贷款为例，风险分析数据包括家庭信息、政府征信信息、关联人信息、农田信息、农资信息，数据来源包括同行、政府、农资供应商、卫星遥感、互联网等，数据的交易将从点到点的原生数据交易向中介型的多阶数据交易转变。

在政府和公共事业数字化方面，社会化可以加速社会治理精准化和人性化。以中国政府的城市一网通管为例，一方面需要构建政企一体、多源整合的政务数据和社会化数据平台化对接机制，充分利用电信、供电、供水等公共数据；另外一方面需要丰富政府侧的数据创建和数据共享能力，不同部门的摄像头、传感器、等成为 24 小时的全场景“工作人员”。

与传统生产要素相比，数据要素表现出一系列新特征：首先具有非稀缺性，数据海量且能够重复

使用；其次具有较强流动性，数据要素的流动速度更快、程度更深、领域更广；第三具有非排他性，可以在一定范围按照一定权限重复使用。未来社会数据将通过“可用又可见”与“可用不可见”相结合的方式形成常态化跨企业、跨行业对接机制，为数字经济时代多元协同共治格局提供支撑。

数据社会化能够在流动、分享、加工和处理的过程中创造出新价值，但海量数据的汇集也将有可能带来严重的安全风险，基础设施一旦发生安全问题，将造成严重的后果，如 2021 年欧洲某云服务提供商的数据中心发生火灾，造成 360 万网站瘫痪，部分数据永久性丢失，社会损失巨大。如何有效利用和保护数据已经成为数字经济安全稳定运行的关注点。只有不断更新数据安全技术和方式，应对快速变化的安全需求，并加强数据中心及其相关的基础网络、云平台、数据和应用的一体化安全保障能力，才可以确保基础设施和数据的安全。





■ 智能化人机协同，提升人效

传统数据中心的运维模式以人为核心，人的能力将成为未来数据中心的运维瓶颈。根据中国信通院 2023 年最新研究显示，数据中心故障宕机场景中，人为操作的事故占比超过 60%。随着数据中心业务的增长，规模也越来越大，传统以人为中心的运维模式难以为继。

中国团体标准《数据中心基础设施智能化运行管理评估方法》中将数据中心自动化运行发展从全部人工运行的初级阶段到全自动运行的高级阶段分为五个等级，我们预测，到 2030 年，领先的

数据中心运维水平将达到 L4 高度运行自动化阶段。在这一级别将实现自动预测性排障和分析、全自动应急处置及 AI 能效管理，在运行态几乎可以达到“无人化”。

实现无人化的前提是数据中心全生命周期实现数字化、网络化和智能化。基于智能来支撑数据中心的规则设计，建设实施和运维运营。面向 2030 年，随着远程监控、数据分析、人机界面、机器人技术的快速发展，极简高效、人机协同的智能数据中心将成为产业发展的新方向。



图 2-1 数据中心自动化发展的五个阶段





愿景与 关键技术特征

03



愿景

人类社会正加速迈向智能世界，跨越式发展已经是全行业的共同诉求。数据中心是新型数字基础设施的算力底座，也是加速数字经济发展的“发动机”。未来十年，数据中心既要实现百倍算力提升，以满足快速增长的智能化业务需求，还要实现百倍能效提升，以满足绿色低碳可持续发展的长期目标。

我们认为，未来新型数据中心将具备多样泛在、安全智慧、零碳节能、柔性资源、系统摩尔、对等互联六大技术特征。

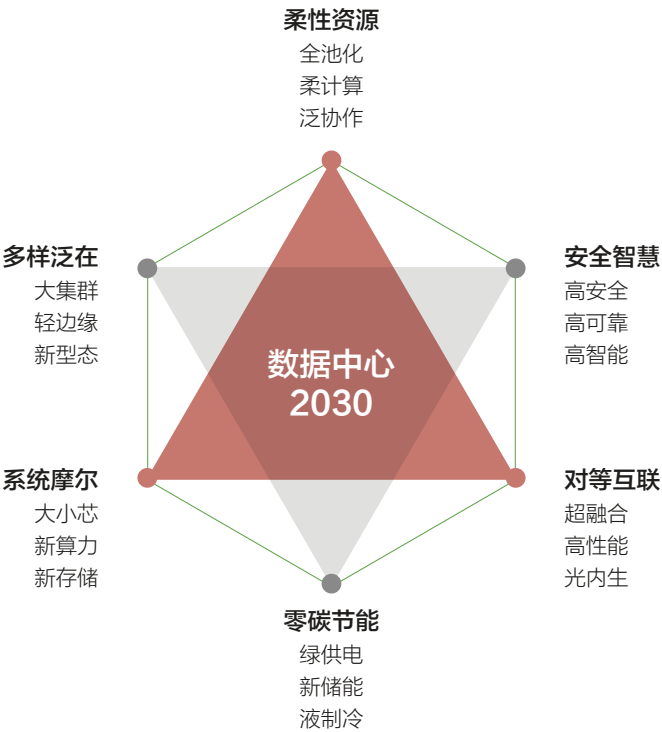


图 3-1 数据中心 2030 的关键特征



关键技术特征

■ 1. 多样泛在

未来数据中心的发展将出现两极分化，一方面超大型集约化数据中心的建设将持续增长；预计到 2030 年，单个集群提供的有效通用算力将达 70EFLOPS，有效的人工智能算力将达 750EFLOPS，配套的存储规模可达 EB 级；另一方面满足各行业低时延、数据安全需求的轻量级边缘计算节点将得到广泛部署；预计到 2030 年，通过轻边缘采集和处理的数据将超过 80%，企业生产设备通过物联化和数字化后，接入轻边缘的比例将超过 80%。同时面向新场景，多种创新型数据中心也将出现，如太空数据中心、海底数据中心等。多种形态满足不同场景部署需求的数据中心将为数字经济的发展提供源源不断的新动能。

1) 大集群

集约化枢纽数据中心部署的服务器规模达到万台甚至十万台，对服务器的部署运维效率提出了很高的要求。传统的数据中心按照服务器为单位进行部署，服务器需要拆包装，上机柜，接电源线，接网线 / 光模块 / 光纤，资产录入等一系列工序后才能部署上线。从运维来看，一个运维人员即使维护一千台服务器，考虑班次等因素数据中心

也需要配置近百人的运维团队。传统的部署和运维方式已经不能满足未来超大数据中心的要求。从单服务器走向计算集群，以机柜为单位进行包装、运输、部署，以机柜甚至整个数据中心为单位进行运维，可以大幅提升部署效率，并降低运维人力成本。

- **预制化交付：**

把服务器安装工作从数据中心前移到生产线可以全流程的提高效率、降低成本。在生产过程中就可以按照实际配置进行拷机测试，测试更完备，并可以增加温度应力等现场不具备的测试项，更有利于发现器件早期缺陷，当出现故障时，维修效率也更高，同时整柜运输比服务器单台运输，包装成本、仓储成本、运输成本能降低约 70%。

- **整机柜工程：**

机柜内采用集中供电，电源模块全局池化的技术，根据负载动态调节电源始终工作在最佳效率区间。通过动态调节供电和储能，应对算力峰值时的突发电力需求。如使用机柜内置水冷门或者使用液冷技术，将散热能力提升到 60KW/ 柜。

· **集群新背板：**

机柜采用线缆背板替代光模块 / 光纤，实现服务器和 TOR 交换机的连接。线缆背板是无源部件，没有功率消耗，可靠性更高。

通过预制化交付、整机柜工程、集群新背板等创新，实现服务器全盲插，杜绝手工接线错误，实现集群的自动化运维，满足未来大型数据中心算力规模增加，部署和运维复杂度不增加的目标。

2) 轻边缘

随着以云为底座的数字化、智能化从互联网行业广泛渗透到千行万业，从非实时 Web 交易、社交、搜索及后端 IT 支撑业务，拓展到实时互动媒体、元宇宙 AR/VR，工业生产系统、机器人，乃至万物智联场景。超大规模集约化的数据中心承载的应用和数据已无法保障遍布任意位置的消费者智能终端、工业 IOT 终端及机器人对低延迟接入与处理需要，亟待将云的弹性资源、应用服务及智能推理能力从超大规模中心延伸到距离各类接入终端更近的轻边缘系统。

轻边缘的形态包括“轻量级边缘集群”与“轻量级边缘服务与应用”两类。前者由云服务商提供小规模硬件算力集群，并分布式部署在合适的网络位置，将全栈云服务的部分核心能力如弹性虚机 / 容器、存储、网络，中间件、数据库、媒体处理、流数据处理及 AI 推理等时延敏感类服务及应用软件通过物理或逻辑专线从中心云服务区扩展到边缘集群站点；后者则以更为轻量化的容器、函数形式，将中心服务区的中间件、数据库、媒体处理、流数据处理及 AI 推理等时延敏感类服务及应用软件部署在由云服务商、运营商、企业客户、家庭客户、个人消费者及任意第三方提供的硬件及 OS 环境上，并可通过开放互联网

及 HTTP/HTTPS 协议穿越防火墙建立与中心云服务区的连接。后者不与边缘算力硬件及中心到边缘的物理专线绑定，因此更为轻量 and 灵活，而前者从全栈云中心服务区下沉，其云服务与应用能力则相对更丰富一些。



· **轻量级边缘集群**

按照是否具备 Internet 公网接入，轻量级边缘集群可以分为两类：

第一类是具备公网 Internet 就近接入能力的开放式公共轻量边缘。其特点是支持将公有云资源池、云服务及网络接入能力下沉至城市 IDC、CDN 边缘站点、5G 接入 MEC（Multi-access Edge Computing）等相关位置，提供小规模（数台服务器）起步，并可扩展（数千服务器）的大带宽、低时延、高性能边缘云能力。其核心技术特征在于：

（1）低时延接入，具备 ISP 本地入口，可以将多家运营商网络接入，为城市区域提供 <10ms 的接入能力；（2）多样化边缘算力，除 CDN 热点 Web 及视频内容缓存加速之外，特别面向视频渲染、边缘 AI 推理、云手机 / 游戏等场景，将 ARM、GPU、NPU 等多种异构算力下沉到边缘，大幅度提升边缘数据

处理的效率；（3）云边协同，边缘计算和中心 Region 通过高速骨干网 / 专线互联，在满足边缘侧高频低时延大带宽的热数据处理能力后，将低频温冷数据传输至中心云进行处理归档，实现分级处理；使能中心云的基础服务和高阶服务扩展到边缘基础设施上，实现中心 - 边缘协同、全网算力调度、全网统一管控的能力。

第二类是特定企业云租户独占，不对外呈现 Internet 公网出口的轻量边缘。除低时延保障之外，更强调本地合规和多地区分支部署与云中心统一管理诉求。通过将公有云端基础设施与云服务高度集成后，部署到用户机房，在用户本地提供标准化的全栈公有云服务能力。面向各类企业用户业务和场景需求，通过高度集成的硬件和灵活适配的云服务软件，为用户在距离业务更近的位置，提供综合性、一致性的云服务体验。其核心技术体现在：（1）高集成度：包含适用于多种环境、多种场景，复用公有云标准的计算和存储服务器，为用户提供标准化弹性算力云服务，如云主机、云容器、云存储等；（2）弹性部署：针对边缘场景进行独立设计，搭

配轻量边缘专用的融合节点机型，提供单站点 1-500 高弹性节点规模；（3）定制硬件：针对不同机房环境（无机房、野外、标准机房、微模块），以及安全可信计算、HPC、AI 计算、Serverless 等场景，提供的定制化的轻量级硬件支持。此类轻量边缘需要企业用户与公有云服务商共同维护基础设施可靠性，保障机房供电和网络的正常。

· **轻量级边缘服务与应用**

云以服务和应用形态从大规模中心服务区扩展延伸到最终用户近端的轻量级边缘，由于与硬件解耦，可充分利用全球数千异构边缘节点，百万服务器资源，不依赖于硬件、基础设施、部署形态（裸机、虚拟机、容器），从而驱动云上应用可更广泛地服务于千行万业。可以为实时新媒体应用提供就近接入，低于 1ms 的端云交互时延，可以提供极致轻量的高级边缘函数能力，使得对实时交互有要求的媒体业务、机器人业务、WEB3 等业务逻辑天然运行在边缘，构建起算子在边缘内、边缘与边缘、边缘与云之间跨地域、跨业务的实时交互式、交叉式新计算模式。



3) 新型态

随着大数据、人工智能、元宇宙为代表的数字经济高速发展，数据中心一方面要匹配用户对低时延、极致体验的业务诉求，又要应对土地紧缺和能源供应的挑战。除了主流数据中心走向大集群、轻边缘的两级化发展方向之外，业界也正在探索一些新型态的数据中心，以满足特定场景的需求，如运营商接入网络边缘数据中心、海底数据中心、太空数据中心等。

· 运营商接入网络边缘数据中心

运营商接入网络边缘数据中心是一种近几年新出现的创新型形态，它通过在接入网络边缘节点提供用户所需的服务和计算功能，使应用服务和内容更靠近用户，并实现与网络协同，为用户提供可靠、极致的业务体验。根据华为预测，2030 年全球运营商部署的 MEC 边缘计算节点数量将超过 1 万个，其具有以下优势：

低延迟：MEC 可以提供高效、高性能的异构计算能力，并且部署在离用户接入网络更近的地方，因此可以大大降低内容数据分发的延迟，为用户提供更快速、更流畅的业务体验。

数据属地化：MEC 将算力直接部署到地市和区县，能够在边缘节点对数据进行本地化处理，避免数据跨越网络传输，实现了属地化管理，确保企业核心数据资产的安全性，做到了可视、可管、合法合规，从而充分保护了企业数据安全性和隐私性。

高可靠：MEC 能够在边缘节点提供最优路径的网络连接，根据业务需求和网络状态，动态调整网络东西向连接和南北向路径，避免

单点故障，支持就近接入，实现最优化的连接效果，为用户提供更可靠、更稳定的数据传输能力。

云网一体：MEC 提供云网原生一体化的底座和北向接口，满足 ETSI 和 3GPP 规范，实现了虚拟机、容器、裸金属、网络、存储等多种资源服务化，支持业务自助集成，支持云、网、业的集成发放自动化，最小化边缘数据中心的维护成本，实现边缘业务的快速开发、快速部署。支持统一的北向接口，方便端、网、边、云一体化的性能统计和故障管理，实现边缘业务快速故障定位与恢复。

MEC 数据中心具有独特的应用场景和技术实践：在内容分发领域，MEC 能够为裸眼 3D、XR 等业务提供更流畅、更清晰的观看体验；在 5G 融合应用领域，MEC 能够提供实时、高效的数据推理能力，为企业提供更高效、更便捷的行业数字化体验。截至 2023 年，中国运营商已经部署了超过 1200 多个 MEC 节点，覆盖了中国 90% 以上的地市。MEC 数据中心通过构建高性能、高集成度的一体化边缘硬件，优化边缘业务体验，是未来数据中心发展中一种全新型态的探索。



· 海底数据中心

数据中心运行过程会产生巨大的热量，散热



制冷所带来的耗电量极大，约占总能耗的三分之一。为了节约能耗支出，把数据中心放在水下，利用海水冷却给“服务器”降温，成为降低数据中心能耗的解决方案之一。海底数据中心在提供数据信息存储、计算、传输服务的同时，实现了绿色、节能、高效的目标。

和陆地数据中心相比，海底数据中心具有很大的优势。首先，海底数据中心具有节省资源的优势，海水作为数据中心的自然冷源，将其产生的热量利用周围的流动水带走。由于海水的比热容较高，数据中心对水温的影响变化可以忽略不计。这一做法，不仅对环境的影响较小，也极大的降低了数据中心用于散热制冷的能量消耗：中国首个海底数据舱的 PUE 值为 1.076，远低于传统数据中心。海底数据中心无需蒸发散热，减少了冷却塔和冷水系统，水资源消耗为零，且大部分设施位于海底，土地资源消耗极少，仅为传统数据中心的十分之一。

其次，海底数据中心具有低时延的优势，全球主要的互联网公司、高科技企业云大多部署在沿海发达地区，建立靠近用户的水下数据中心能够缩短数据的传输距离，降低传输时延，有效满足工业互联网、远程医疗等对时延敏感的业务需求。

最后，海底数据中心的部署和运营成本较低。沿海发达地区的地价昂贵，数据中心建在海底能够大大降低土地成本。在运营过程中，能源消耗降低，电费大幅减少，运营成本急剧下降。

海底数据中心优势明显，正在逐步进行商业化探索：2023 年 3 月全球首个商用数据中心海底舱在中国海南正式运行，是全球最大的海底数据舱，有望成为数据中心绿色发展的新范式。

• **太空数据中心**

随着无人值守和自我管理型数据中心日益受到大家关注，太空作为网络的极端边缘和无人值守计算的理想运行地点，逐渐成为数据中心建设的目标地点之一。随着太空商业化的不断发展，卫星提供电路、广播、导航等服务逐渐普及，使部署太空数据中心成为可能。未来十年，多个商业空间站和数千颗卫星将发射进入近地轨道。业界正在研究将数据中心送入太空轨道运行的可能性，打造一个由数据中心、边缘计算组成的为太空轨道经济服务的数字基础设施。

数据中心部署到太空存在着一定的优势，其成长空间有望被打开。第一，数据中心将更

加绿色高效，太空中太阳能储量丰富，能源效率更高，可以更加稳定持续的为数据中心进行供电，从而减轻地球的能源压力，降低二氧化碳的排放。太空低温环境会使数据中心的温控方式产生较大变化，减少能源消耗。第二，太空数据中心建设会提高太空数据的利用效率和传输速率，减少卫星与地面间传输的数据量。地球的近地轨道资源逐步被占满，低轨卫星的数目越来越多，产生的数据量越来越大，卫星数据传输的窗口期逐渐被压缩，未来卫星数据很难在有限时间内被完全传输。在太空中部署数据中心，使数据靠近计算端和应用端，数据在太空中直接完成采集和应用。数据中心和卫星通信网络结合，有利于边缘计算，系统效率不断提升，进一步降低服务延迟。第三，太空数据中心运营成本低，数据安全性高。数据中心的主要运营成本包括维护和能源，太空的固有环境优

势会大大降低数据中心的运营成本。干扰和拦截卫星发送的数据更加具有挑战性，卫星数据的边缘计算也进一步保障了数据安全。

太空数据中心的建设目前也存在一些问题。比如太空数据中心的建设成本过高，从研发到发射将耗费至少 10 亿元人民币，对应的企业估值应在百亿以上。太空环境存在辐射，这对服务器的整体可靠性和抗辐射性提出了更高的要求：专用计算芯片的创新，磁随机存储器（MRAM）等将为太空数据中心提供新的选择。

到 2030 年，随着有效载荷送上太空的成本进一步降低，在太空建设数据中心的可能性进一步增大。

■ 2. 安全智慧

1) 高安全

数字经济的核心技术涉及数据、算法与算力三个要素，确保上述三个要素的安全与合规是数字经济发展的基础。数据中心作为数字经济的基础设施，不仅是承载上述三个要素的平台，同时也是数据流通交易的平台，因此，数据中心必须在系统设计之初就内置安全能力，并贯穿计算、存储、网络等数据处理的全过程，覆盖从芯片到应用全栈，以抵御全方位的安全威胁，预计到 2030 年，数据中心在安全方面的投资占比将达到 20%。

数据中心的数据基础设施需要在设计上充分考虑对不同类别和级别的数据提供保护，数据在使用，

存储，传输过程中，能够根据不同的数据价值和合规要求配套不同的安全策略，保证数据在数据中心内的全生命周期的安全合规，可管可控。同时，系统应具备根据新的分级和安全防护规则，自动调整适配的能力。

数据中心需要提供基于零信任理念的安全方案，使用包括 RBAC, ABAC 等细粒度的访问控制模型，对数据的使用进行严格限制。

数据在传输和落盘过程中需要提供加密能力，整个密码学体系需要充分考虑到抗量子计算攻击的风险。对于高价值数据提供基于机密计算，联邦学习，同态加密等使用态的保护方案，实现可用

不可见。

高价值数据的防护要具备更强的抗攻击能力，高价值数据与普通数据隔离存储，高价值数据的完整性，机密性保护方案，比如 WORM，密钥管理和分发的机制，需要考虑提供软硬结合的安全设计，应对恶意攻击。

面向 2030 年，除了传统的设备安全、网络安全之外，可信计算、机密计算和 AI 大模型时代的新安全生态等成为主要研究方向。

· 新计算范式下的可信计算

可信计算技术 TPM 起源于上世纪 90 年代末期，其巧妙的应用了密码学原理解决了软件完整性保护与加密密钥保存问题。但其已经越来越难以有效支撑针对云化虚拟机、异构计算软件的完整性度量。为了解决这些问题，

主流开源社区已经开始支持软件 TPM 虚拟化功能 swTPM，但在 VM 中提供的 TPM 服务需要依靠软件 libTPM 模拟，缺乏硬件可信根，导致无法有效证明其安全性。

为彻底解决针对虚拟机、异构计算环境的软件完整性度量问题，需要扩展当前可信计算 TPM 技术与标准，从硬件开始支持可信计算多实例，在此基础上不仅可以为多个 VM 提供软件完整性度量服务，也可以为可信计算环境 TEE、异构计算环境中 XPU 提供软件完整性度量服务，最终实现以硬件可信根为基础的，支持云化、异构计算环境的可信计算技术与标准体系。

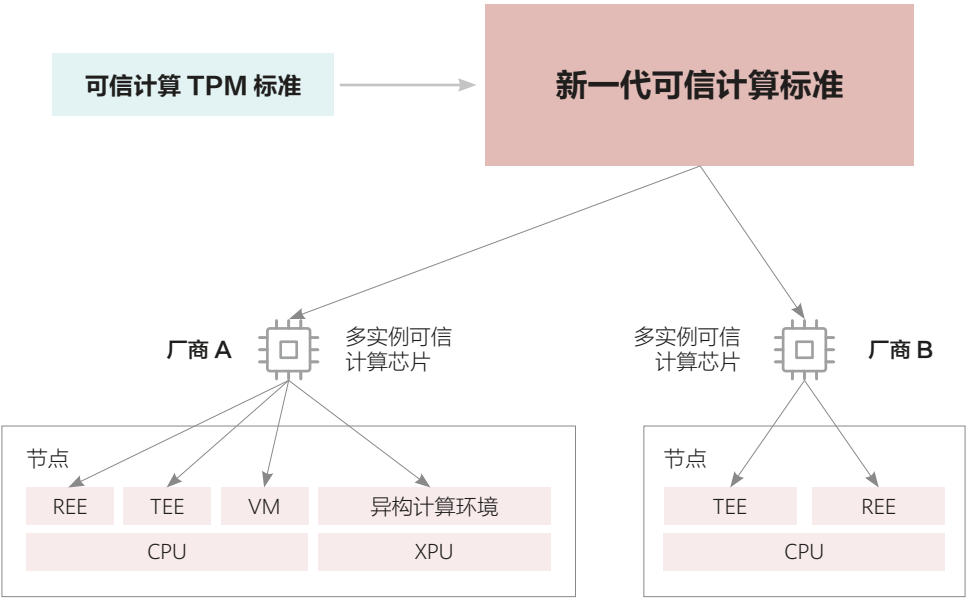


图 3-2 新计算范式下的可信计算

· **数字经济时代的异构机密计算**

机密计算是在基于硬件的可信执行环境（TEE）内运行代码，以保证环境内数据的机密性、完整性以及运算过程机密性的新计算模式。数字经济时代对机密计算的需求来自多个方面，其一，是来自企业与用户的互信需求，需要在用户不依赖企业计算环境可信的前提下保护用户数据安全及隐私，典型场景如云上的内鬼管理员作案；其二，来自企业自身的防御需求，在不可信的环境中通过有效手段来保护企业数据安全，典型场景如部署在公共场所的边缘计算设备中的密钥管理保护；其三，来自机构之间数据共享的需求，在互相不暴露各自数据的前提下，仍然能进行数据合作，典型场景如多方计算与建模。

当前，随着全球范围内数据与隐私保护法律法规的陆续出台，机密计算技术已经逐渐被业界认可和接受，越来越多的企业青睐于软硬件结合的数据安全隐私保护解决方案，为数据共享和交换构建安全可信的计算环境。然而，当前的机密计算技术落地主要覆盖的还是较小数据规模、安全性要求较高的应用，如在手机、平板等端侧设备上的支付、人脸识别应用，或在云上采用机密虚拟机或者机密容器实现的区块链、密钥管理等应用。面向数字经济时代的大规模数据的通用计算和AI计算场景，机密计算仍然处在试点和探索阶段，距离全面落地仍有一定差距。

到2030年，随着大数据应用的进一步成熟和AI大模型的迅速发展，多机构之间的通过数据共享进而挖掘数据价值，训练更精准的大模型等场景会变得更加普遍。机密计算技术可以在大规模的数据分享与计算过程中既

保证计算的高性能，又保证数据的安全与“可用不可见”，因此有望成为未来数据安全的主流技术。为了能够高效率的适配大模型、大数据等对算力要求极高的场景，以CPU为中心的机密计算技术将会逐步演变至以数据为中心的异构机密计算技术，同时兼容已有的大模型软件框架，支持多样化算力设备如GPU、DPU、NPU等对机密计算安全算力的加速和协同。

具体来说，面向2030的异构机密计算架构有以下几个特点：

（1）安全算力与普通算力生态完全兼容且可灵活配置，同样的计算任务和数据，可以由用户灵活选择是否部署安全算力，不会因为采用了机密计算技术而带来应用和软件生态的改变，进而带来额外的工作量。

（2）安全算力从CPU内的可信执行环境扩展到多样化的算力设备，如GPU、DPU、NPU等，既能高效的利用异构硬件进行计算加速和卸载，并兼容已有的软件生态，又能通过访问控制、通信加密等手段保障异构计算的安全可信。

（3）安全算力从单节点扩展至多节点，在整个数据中心实现安全算力的灵活调度，完成算力资源的统一管理，从而有效满足多机构、大数据量的数据共享与联合需求。

· **大模型时代的 AI 安全新生态**

随着以 ChatGPT 为代表的大模型时代的到来，AI 将在更多的领域发挥关键性的作用，从根本上改变人类的生活和工作方式。当各种各样的 AI 应用落地到越来越多的关键基础设施行业，AI 承载了越来越多的商业价值，就会诱生出各种新型的安全威胁与攻击手段：既有针对 AI 模型自身的脆弱性的新型攻击，如数据投毒（药饵）、模型后门、对抗样本、模型萃取，以及专门针对大语言模型的提示注入（prompt injection）等；又有对 AI 技术有意识或无意识的滥用，如将 AI 用于欺诈、造谣、身份伪造（deepfake），造成数据与隐私的大规模泄露等。

AI 系统与应用的安全可信越来越成为国家、社会、行业以及广大用户普遍关心的问题。全球主要国家、地区、国际标准组织也在着手探索对新兴的 AI 系统、应用与服务进行有效监管的方式与方法。2021 年，欧盟率先发布了《欧盟 AI 法》草案，对高风险 AI 系统明确提出了数据治理，责任追溯，准确性、健壮性和网络安全方面的要求。针对大模型与 AIGC 带来的 AI 滥用的迫切威胁，中国政府也出台了《互联网信息服务深度合成管理规定》、《生成式人工智能服务管理办法（征求意见稿）》。

在此大背景下，业界迫切需要针对 AI 业务发展中遇到的安全问题、安全威胁，提出创新的安全技术，构建安全解决方案以应对新的威胁：

（1）AI 全生命周期安全：将安全融入到 AI 的整个生命周期中，将安全治理的理念贯彻到整个 AI 研发、使用的流程中，以数据安全与治理为起点，通过持续的模型安全测评与改进，不断提升 AI 模型的安全性，并对使用中的 AI 应用进行持续监测与及时响应。

（2）用 AI 保护 AI：由于传统安全手段无法感知和防御对抗样本、提示注入等新型 AI 安全威胁，在新的一轮“矛与盾”的对抗中，需要基于深度神经网络的强大的端到端学习与泛化能力，利用创新的 AI 安全模型，感知与防御各种安全攻击，从而实现“用魔法打败魔法”。

（3）面向监管的透明可回溯的技术：针对 AI 加强监管，以避免或减弱其对人类社会的负面冲击，实现 AI 向善，已经成为各国的共识。借助于创新的透明责任可回溯技术，在整个 AI 生命周期中实现对方权责的透明、可信界定，才能真正实现 AI 向善的目标。



2) 高可靠

高可靠的数据中心是支撑数字经济发展的重要基石。面向 2030，数据中心将从设备级、节点级、同城级的高可靠，走向更大范围的跨域高可靠；从主要关注数据级高可靠，走向关注业务连续性的业务级高可靠，系统级和业务级可用性均将达到 5 个 9 的水平。

为实现数据中心跨域、业务级高可靠，需要研究多 DC 数据一致性保障、多 DC 异地多活、基于 AI 的高可靠等关键技术。

· 多 DC 数据一致性保障

当前通过双活、同步复制等技术基本能解决单 DC 集群内、同城双 DC 集群间的数据一致性问题，但无法解决一致性与远距离时延的矛盾。

未来，面对远距离跨域数据中心间的数据一致性保障，将需要进一步探索超低时延光网络传输技术、多 DC 分布式数据库技术、以及结合时序化、精准时钟同步等技术，并考虑时延、数据一致性等 SLA 策略，最终实现远距离、大规模、灵活的多 DC 数据一致性保护能力。

· 多 DC 异地多活

数据中心多地多活技术的实现是一个系统工程，它需要从网络接入层到数据层、再到存、算资源层端到端的实现业务多活分担、精准调度及流量自治。

未来，随着云计算、低时延大带宽网络互联技术的不断发展，跨多个数据中心的资源池将整合成一个“虚拟数据中心”，上层业务的部署将对地域无感知，即业务的



Regionless 化，从而，数据的高可靠、业务的连续性也将实现去地域化，具备多 DC 异地多活的能力。

· 基于 AI 的高可靠

当前，通过预设操作、人工决策、手动触发等进行数据中心故障切换和应急管理的方式，难以支撑业务的真正连续。

未来，数据中心将利用 AI 技术提前预防发现隐患，不仅与数据中心 IT 健康、供电等内部环境相结合，同时与供电网络、地震感知系统等外部环境以及安全态势等要素结合，再利用 AI 预防算法深度自学习、大数据分析算法，进行灾难关联智能预测，并做到自动化预防响应。在发生故障和灾难时自主开展全链条自愈恢复，进而在隐患影响到业务前，执行有效预判并开展计划性、应急性响应，提前解决影响业务运行的问题。做到要素全量感知、过程全监控、决策智能化、切换自动化、指挥可视化，实现全面合规的灾备运营管理体系。

通过多 DC 数据一致性技术、端到端多地多活技术以及基于 AI 的高可靠技术，可大幅度提升跨 DC 业务连续性保障能力，充分调度数据中心资源，提升资源利用率。

3) 高智能

数据中心投资飞速增长，规模日益增大、密度不断提升，数据中心的复杂程度越来越高，数据中心的大型化增加了数据中心的建设运营的复杂度与约束，传统建设运营模式面临挑战。AI 和数据使能数据中心规划、建设、运营全生命周期，促进数据中心朝着高效、节能、智能的方向发展。

· AI 使能

AI 技术用于数据中心规划、建设、运营全生命周期，能够大大提高数据中心的建设运营效率，实现降本增效。基于 AI 的智能管理系统与数据中心供电、制冷系统结合能显著降低数据中心的能耗、降低运行故障率、提升运营效率：例如 AI 技术和 UPS 结合大幅提高数据中心的用电质量，UPS 对输入电网和输出负载质量相关主要参数进行实时跟踪，引入 AI 智能算法，对历史数据进行主动学习和分析，能够大幅提高数据中心的供电质量。AI 技术和数据中心内各个系统和部件的进一步融合，推进数据中心运营更加高效。应用智能运维，有效提高数据中心运营的流程化和标准化，提升运维管理效率。网络智能运维是数据中心智能运维最重要的应用场景之一，可持续推动网络可视、可管、可控能力的提升。网络运维技术通过融合 AI 技术实现对数据中心网络的管控、自动运维及优化，对各类网络故障进行恢复自愈，对阻塞进行管理，同时支持网络自我调优和自我演进，为用户提供更加优质的网络服务。网络自身形成面向任意场景具备执行、监视分析和决策能力，完全实现闭环管理和自动化能力，基于 AI 的智能运维系统能够大幅提升数据中心的运行效率。

· 数据中心数字孪生

数字孪生技术借助历史数据、实时数据、算法模型等，实现对物理实体全生命周期的模拟、验证、预测、优化、控制。数字孪生技术在数据中心的应用将大幅提升数据中心自动化、智能化水平，给数据中心所面临的安全运营、节能减排等挑战带来极具竞争力的解决方案。数字孪生技术在数据中心设计阶段的应用主要包括仿真评估和 3D 可视等，随着技术的进一步发展，数字孪生技术还将实现设计方案的自动优化。在数据中心建设阶段，运用数字孪生技术对施工进度、质量、安全进行管理，实现进度过程可视，人力物力协同，提升这一阶段的智能化水平。在数据中心运维阶段，数字孪生可视化利用 3D 技术，将孪生体对象，包括园区建筑、机房布局、基础配套、冷通道和机柜、IT 设备、强弱电链路等，通过数据处理和建模仿真实现数字化映射。这样可以实现 IT 设施、动环、容量、链路、告警等信息的可视化，并对数据进行模拟、分析、预测和验证，提供决策依据，为数据中心减员增效提供助力。

随着数据中心大型化、集中化的发展，传统数据中心由僵化的结构、低效的管理与运营，向具有数字化、网络化、智能化特征的数据中心转变，运用 AI 技术、数字孪生技术等使能于数据中心全生命周期，以实现投资效益与运营效率的最大化，无疑是大势所趋。智能运维机器人等新技术的使用进一步实现了独立诊断、自动排障、防御升级的监控模式，解放了运维人力，预计在 2030 年，业界领先的数据中心自动化运行能力将发展到 L4 级，几乎达到“无人化”。数据中心的高智能发展实现了由“人力密集”到“技术密集”，为数字经济提供了更可靠的保障。

■ 3. 零碳节能

1) 绿供电

数据中心整体的碳排放量较大，对于数据中心服务商，高能耗也意味着更多经营成本。随着“碳中和”逐渐成为全球的共识，数据中心将加速向绿色低碳的方向迈进：绿电的发展为数据中心提供了更加丰富、优惠的电能供给，全面助力数据中心零碳目标的实现。随着全球绿色低碳发展政策的不断强化，风能、太阳能等清洁能源在数据中心能源结构中所占比例将会提升，预计到 2030 年，大型数据中心绿电使用率将达到 100%。

- **提高绿电采用率**

- (1) 风电**

风电产业是可循环新能源产业，作为一种可再生能源，其分布广、无污染的特性有效控制了不断增加的能源供给对环境所带来的影响。在全球能源过度消耗的生态环境下，对新能源的研究和利用已成为世界热门话题。风力发电是新能源发电技术中最具规模开发和商业化发展前景的发电方式，目前各国都在加大对风力发电及其相关的技术研究。全球风电行业年度市场增长率达 40%，已有一百多个国家涉足到风电行业，该行业已经成为世界能源市场的重要组成部分。

数据中心产业的绿色低碳需求，给风电企业提供了稳定的消纳空间，使数据中心和新能源企业能够协同发展。例如，中国乌兰察布地区周边的风电资源丰富、装机充裕、电价低，可实现清洁能源供电，华为早在 2013 年就在乌兰察布建设了数据中心，利用风电为数据中心降本增效。

- (2) 光伏**

随着光伏技术的不断发展，光伏的经济性（电价在规模化扩张中基本与煤价持平）、技术性和公众认可度方面取得了巨大的进步，光伏发电能够为解决气候问题、降低能源使用成本、保障能源安全等方面带来的巨大作用。分布式光伏可以在数据中心就近建设、就近供电，进而减少供电成本。由于分布式光伏主要建设在建筑屋顶，不会额外占用土地资源，使其可以在数据中心建筑屋顶建设。目前，光伏发电越来越多的应用于数据中心辅助设施或次要负荷的供电，如照明、电梯、监控系统等。而“光伏 + 储能”和“光伏 + 电网”等多种互补模式可以不间断的为数据中心提供清洁电能，满足数据中心昼夜不停运行的用电需求。

- (3) 水电**

水电作为一种清洁、可再生的能源，在优化电力结构、保障电力安全运行、降低电力资源消耗、提高电力经济效益等方面具有重大优势。把数据中心建在水电资源富集的地区，不仅能够用清洁能源对数据中心进行供电，还能就地取材利用水资源帮助数据中心冷却。为了降低能耗及成本，中国很多数据中心位于水电资源丰富的地区：位于中国三峡地区的东岳庙数据中心，运行时所需要的能源完全来自于附近的三峡水电站，并且它的冷却也直接应用三峡水库中的水，解决了数据中心的绿色能源供应和冷却问题。

· **动态微电网**

利用风能、水能、太阳能等绿色能源对数据中心供电、能源可持续发展以及绿色低碳目标的达成具有重要意义。然而，绿色能源的利用具有不可预测性，不能长时间持续稳定的供电，可能会导致电力系统的波动甚至崩溃。

微电网是指由分布式电源、储能装置、能量转换装置、相关负荷和监控、保护装置汇集而成的小型发电配电系统。数据中心微电网在绿色能源充足之时就近消纳，省去在电网中传输的损耗，提高能源的使用效率，并和电网单点连接，在能源供应不稳定时，利用电网供电。微电网采用先进的控制方式以及大量电力电子装置，将分布式电源、储能装置、可控负荷连接在一起，使得它对于电网系统成为一个可控负荷，并且可以施行并网和独立两种运行方式，充分维护了微电网和大电网的安全稳定运行。

目前数据中心微电网已经有很多实践，例如中国张北云计算基地绿色能源中心新能源微电网项目，总装机规模达到 220 兆瓦，项目建成后年发电量约 4.5 亿千瓦时。

面向 2030，随着微电网关键技术的进一步研发以及数据中心绿色低碳发展的进一步推进，微电网将进入快速发展期。

2) 新储能

储能技术通过“削峰填谷”，成为降低数据中心电力成本的重要方式。数据中心能耗高，电力成本占运营总成本的 60%-70%，供电公司通常会提供波峰或者波谷电价，数据中心可利用储能系统在波谷时存储电力，并在高峰期利用，以降低数据中心用电成本。国际能源署发布的《2022 年世界能源展望》报告显示，随着越来越多国家开始加速能源转型，许多国家和地区制定了可再生能源发展目标及规划，加速能源转型。根据法国政府计划，到 2030 年，法国可再生能源发电占比将提升至 40%；到 2050 年，法国太阳能发电装机容量将增加 10 倍，并建成 50 个海上风电场。日本最新版能源基本计划提出，2030 年可再生能源发电占比将提高到 36% 至 38%。全球可再生能源产业进入一个快速发展期，可再生能源的渗透率越高，平衡电力系统负荷需求越大，所需储能时常越长。

· **锂进铅退**

随着数据中心对内部空间容量管理及高效运营的需求逐渐加强，以提升数据中心功率密度为目标的数据中心改造方式正在成为推动数据中心升级的重要路径。在数据中心储能方面，锂离子电池以其高能量密度、高输出电压及高安全性特点迅速成为替代当前数据中心铅蓄电池的下一代储能设施：目前越来越



越多的数据中心开始指定锂离子电池作为供电单元。相较于传统铅蓄电池，锂离子电池具有诸多优势，在体积方面，锂离子电池体积小、重量轻，数据中心运营人员可直接将锂离子电池放在更高位置而无需采用加固地板；在占地面积方面，锂离子电池占用空间仅有铅蓄电池的三分之一，能够更好地适应模块化数据中心机房环境，轻便灵活的特点使其更加易于运输和安装；在使用寿命方面，铅蓄电池寿命较短，通常在三到六年内报废，而锂离子电池工作寿命为 10-15 年，更长的使用寿命意味着数据中心电池更换及维护成本将得到极大地降低；在安全可靠方面，当前主流的数据中心锂离子采用长寿命的磷酸铁锂电芯及四级架构系统保护方法，电池充放电性能能够得到有效保障；在电池管理方面，锂离子电池可与更为先进的电池监控系统（BMS）结合，为数据中心运维人员提供电池运维时间、健康状态等信息。随着市电供电环境的日益完善，数据中心储能电池使用场景将不断减少，锂离子电池的应用能够有效压缩数据中心电池运维管理成本，助力数据中心安全高效管理，未来必将得到更为广泛的应用。

氢储能

氢能由于燃烧时不排放温室气体和细粉尘，与全球碳减排战略相契合：和光伏和风电等可再生能源相比，氢能可以克服其波动性和间歇性的短板，成为世界能源转型的关键补充。氢储能在全球范围内正得到越来越多的关注和应用：全球已经有 20 多个国家和地区已发布氢能战略，相关技术不断取得突破。

电力储能方式目前主要有抽水蓄能、锂电子电池等等，与其他储能方式比，氢储能具有

放电时间长、规模化储氢性价比高、储运方式灵活、不会破坏生态环境等优势。另外，氢储能应用场景丰富，在电源侧，氢储能可以减少弃电、平抑波动；在电网侧，氢储能可以为电网运行调峰容量和缓解输变线路阻塞等。数据中心作为数字经济的“发动机”，承载着支撑各行业实现智能化转型的重要使命，因此保障数据中心的安全具有重要的战略意义。在使用氢能供电时，数据中心面临的主要风险是氢气泄漏导致燃烧、爆炸，造成对数据中心的物理性伤害。因此，如何有效管理氢安全、保证零事故是核心前提。

随着氢能的不断发展，已经有一些企业在数据中心中应用，比如业界某公司的数据中心项目采用功率达 4MW 的燃料电池，进而替代柴油发电机作为备用电源。氢储能目前仍处于起步阶段，随着氢能的不断发展完善，将与数据中心深度耦合。



3) 液制冷

制冷系统是数据中心除 IT 设备之外的第二大耗能部分：数据中心的 IT 设备在运行时持续产生热量，在超出额定功率和范围时会导致服务器宕机引起业务中断、设备寿命缩短等一系列问题，所以需要通过制冷系统用来维持 IT 设备的正常运行。为了降低数据中心的 PUE，先进制冷技术的研究尤为重要，逐渐成为数据中心的关键技术，其发展以绿色节能为导向，朝着融合技术创新、模块化和集成化方向发展，并通过智能化的手段与 IT 设备的运行状况相结合，进行动态的适配和调整。

液冷技术多适用于高功率、高密度数据中心，数据中心的液冷技术目前处于探索阶段，总体发展趋势良好。随着数据中心机柜平均功率密度的逐年上升，对液冷技术的需求也就不断增加，液冷数据中心的市场规模持续扩大。液冷技术的应用不仅有助于数据中心的节能和降噪，也有助于提升单位空间的服务器密度，从而提升数据中心的计算效率以及稳定性。

• 全液冷技术

全液冷目前主要有冷板式液冷、浸没式液冷和喷淋式液冷三种技术路线。冷板式液冷技术是通过冷板将发热元器件的热量间接传递给封闭在循环管路中的冷却液体，通过冷却液体将热量带走的一种实现方式。在该项技术中，工作液体与被冷却对象分离，工作液体与电子器件不直接接触，而是通过液冷板等高效导热部件将被冷却对象的热量传递到冷却液中，因此冷板式液冷技术又称为间接液冷技术。

浸没式液冷是最近几年备受业界关注的新型散热技术：它主要是采用特定的冷却液作为散热介质，将 IT 设备直接浸没在冷却液中，通过冷却循环带走 IT 设备运行过程中产生的热量。同时，冷却液通过循环过程与外部冷源进行热交换，将热量释放到环境中去。由于架构特殊，浸没式液冷具有独特的优势。首先，浸没式液冷的冷却液与发热设备直接接触，散热效率较高；第二，冷却液具有较高的导热率和比热容，运行温度变化较小；第三，支持更高功率密度的 IT 设备部署，能极大地提升能源的使用效率。这种散热方式同风冷相比，密度更高、更节能、防噪音效果更好。

喷淋式液冷通过在服务器内部部署喷淋模块，使用一种对人、IT 设备和环境无害无腐蚀的绝缘液体，有针对性的对发热器件进行喷淋降温的解决方案。喷淋液冷具有器件集成度高、散热效率高、高效节能和静音的特点，是解决大功耗机柜在 IDC 机房部署以及降低 IT 系统制冷费用、提升能效的有效手段之一。

液冷的发展仍存在很多挑战，在液冷方式下，由于制冷方式改变，部署环境也会随之变化。在传统机房部署液冷会带来部署成本和部署难度方面的问题。浸没式和喷淋式液冷还需要考虑 IT 设备和液体的兼容性以及液体对人的友好程度等。而冷板式液冷不需要昂贵的水冷机组，在减少总体拥有成本的同时，可以有效提升数据中心的能源利用效率。

· 风液混合

风液混合数据中心是由于目前液冷成本还较高，为了平衡初投资和 PUE 而采用的折中方式。因此这两种方案混合部署，能够带来成本降低和 PUE 目标达成。在风液混合的数据中心中，风冷和液冷分别处于不同的子机房，互相之间独立无干扰。

风液混合成为数据中心制冷技术发展新趋势，未来数据中心市场将出现“风冷 + 液冷”混合发展的新格局：风冷技术不会被液冷技术完全取代，客户会根据不同需求，选择不同的数据中心制冷方案。对于部署了较低单机柜功率的数据中心，客户依然会选择风冷的制冷方式。而针对超算、能源勘测等高密度大规模计算需求，综合考虑成本因素，灵活地选择冷板和浸没的混合液冷，部分功耗大的部件和设备改用液冷的解决方案也将成为选择。总而言之，液冷技术将与风冷技术配合，助力行业发展。

· 极致 PUE

“能耗指标”及“碳排放指标”成为数据中心的核心竞争力所在，挖掘数据中心的节能减排潜力，提升数据中心建设的能效标准至关重要。

数据中心极致 PUE 是多个因素平衡的综合结果。其中“极致”是针对单个数据中心而言。数据中心有自己的特质，不仅分地域特点，也要分行业特点。数据中心从规划设计、部署实施到管理运维全生命周期采用关键系统低碳节能技术，从而达到极致 PUE 目标。

· 最优 WUE

水资源是人类赖以生存和发展的基本元素，是维系生态系统和支撑社会经济发展不可替代的战略资源，数据中心是耗水大户，寻求 PUE 和 WUE 的平衡，实现最优 WUE 极其重要。

数据中心 PUE 和 WUE 是互相紧密关联的两个指标，水蒸发是最高效的换热方式之一，多耗水能帮助数据中心实现更优的 PUE，因此，必须基于实际的业务诉求、地理环境，在 PUE 和 WUE 之间寻求平衡，实现最优 WUE。如在水资源丰富的区域，可以通过用水来实现极致 PUE，而在水资源匮乏的区域，则减少用水，并通过收集雨水、废水处理循环利用、减少机组湿工况运行的时间等开源节流的方式尽可能降低用水量，预计到 2030 年实现 $WUE < 0.5L/KW \cdot h$ 。



4. 柔性资源

随着公有云、行业云、私有云作为数字化、智能化平台底座在全球各行业的广泛深度普及，加之 AI 大模型、元宇宙、数字孪生等大颗粒应用的方兴未艾与爆炸式增长，云化架构将成为未来数据中心基础设施的“标配”之一，即通过在全球分布式数据中心硬件之上叠加云操作系统软件，为千行万业、政企客户以及多样化业务应用提供多租安全隔离、性能 SLA 保障前提下的大规模集约化算力、存力、运力共享，以及动态随需而变的算力、存力、运力供给。

下一代云数据中心架构，将沿着“全池化”、“柔计算”、“泛协作”的方向持续演进，从而实现数据中心硬件资产投入产出比的极致最优。

1) 全池化

通过超大规模资源池化，实现多租、多应用对数据中心算力、存力、运力资源的最大化共享，是云最本质的特征之一，然而，由于当前云数据中心架构与技术的制约，计算、存储、网络资源虽已具备池化能力，仍存在大量资源碎片化孤岛，限制了规模化、集约化共享效率的进一步提升。因此，未来 5-10 年云数据中心的“全池化”是主要趋势，具体体现在：跨硬件代次 CPU 统一

池化、跨节点内存统一池化、存储统一池化（存算分离）、异构算力统一池化，及 DCN 网络统一池化。

· 跨硬件代次 CPU 池化

当前云数据中心的算力供给仍采用“以资源为中心”的模式，将 CPU 代次信息透传给面向最终云租户的弹性虚拟机及容器，不同 Flavor 规格对应 Intel Sandy Bridge, Ivy Bridge, 鲲鹏 920/930 等不同 CPU 代次，使得每个 CPU 代次的算力资源事实上形成了彼此间无法动态共享的“独立资源池”，而云租户在消费算力时总是倾向于选择最新 CPU 代次的虚拟机、容器，必然导致不同硬件 CPU 代次的“独立资源池”之间存在分配率不均衡问题，特别不利于仍处于可靠性浴盆曲线有效期范围内存量老 CPU 代次算力资源的有效利旧及价值转换。为应对上述挑战，下一代数据中心需将算力供给转变为“以应用为中心”模式，将一定 CPU 代差范围内的算力资源进行统一整合，通过云算力服务及资源调度层屏蔽底层 CPU 代次的差异，并基于实时的黑盒式性能 QoS 检测机制，在满足云租户应用性能 SLA 的前提下实现最优的 CPU 资源自使用动态超分复用比。

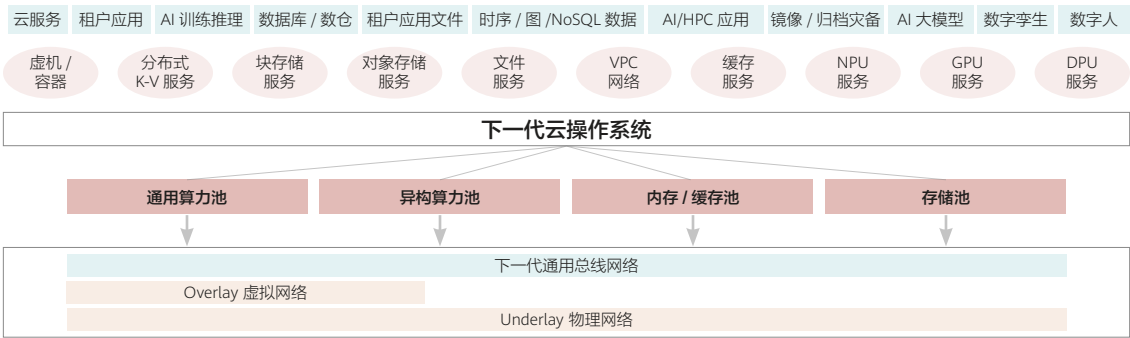


图 3-3 云数据中心全池化框架

· **跨节点内存池化**

云数据中心算力基于预定义的系列化 Flavor 规格，将特定大小的内存资源与特定数量的 CPU 核资源绑定，在匹配硬件规格的服务器节点内进行分配从而实现最小化的资源碎片，而内存不同于 CPU，在资源分配过程中无法超分，原因是 CPU 超分下的资源不足至多影响应用性能，而内存不足则可能引发应用运行异常甚至失败，也因此导致弹性虚拟机、容器中大量过配置的内存资源只能保持闲置而无法得到充分利用。若考虑引入内存超分，在部分服务器面临动态内存资源不足时，如果网络时延和带宽能力可支撑，则完全可通过网络向算力集群内其他服务器“借用”空闲内存资源。而随着百纳秒级超低时延、百 GB 级超大带宽远端内存异步访问网络技术（UB/RDMA）的不断成熟与发展，使得打破服务器物理边界实现面向虚拟机、容器无感的统一内存资源分配与读写访问成为可能。考虑到跨服务器远端内存访问依然相比 CPU 通过 DDR 内存通道访问本地内存慢 1-2 个数量级，并会带来应用性能下降 20%-30%，因此跨节点统一内存池化共享仅普适于对应用性能劣化有一定忍耐度（如性能劣化最大 30%）的云上应用类型。而对于性能敏感的多租户虚拟机 / 容器复用场景，跨节点内存池化并非最佳选择，需考虑通过极速毫秒级迁移，及时进行内存资源腾挪，从而保障性能敏感应用。

· **跨异构海量数据的存储&缓存池化**

基于统一的去中心化跨可用区分布式 K-V 存储引擎，当前云数据中心实现了块存储、对象存储、文件存储等面向基础非结构化数据的存储资源统一池化，然而云上 SQL 行存、列存数据库、NoSQL 数据库、多维数据仓



库、图数据、时序数据库等半结构化、结构化数据服务，软件架构仍多采用存算合一模式，也即同时覆盖了计算侧的数据查询、变更、分析处理，以及存储侧的数据持久化可靠性、可用性保障、并行 IO 读写，以及业务无损的弹性容量管理等功能，由此带来了一系列问题与挑战：（1）不同数据处理与分析场景下对计算及存储资源的弹性扩缩容器需求不同步；（2）计算侧与存储侧跨节点数据冗余机制存在冲突导致部分数据库、数仓及大数据集群只能使用存算一体的服务器硬件类型，无法与其他计算类云服务及租户应用共享通用算力资源，也无法充分享受软件定义弹性存储服务带来便利；（3）计算侧与存储侧的数据 IO 路径存在多跳迂回，数据访问性能存在瓶颈；（4）云租户数据全生命周期流水线数据处理与分析的不同阶段之间的数据资产难以共享，跨阶段数据拷贝代价大，多份数据冗余存储成本高；（5）跨云可用区多活部署架构计算层数据冗余处理逻辑复杂度高。应对上述挑战，未来云数据中心通过一份数据拷贝跨不同数据计算引擎共享、近计算统一池化缓存、近数据分布式算子卸载、跨异构计算引擎的统一元数据管理、智能化数据分层存储等关键措施实现跨结构化、半结构化及非结构化海量数据的统一存储池化。

· **异构算力池化**

随着 AI 大模型及元宇宙数字孪生时代的到来，云上 GPU/NPU 异构算力将逐步取代通用 CPU 成为 AI 大模型训练推理，以及数字人、数字孪生城市渲染仿真的关键生产资料，其需求量必将迎来指数型爆炸式增长。然而当前主从计算架构下 GPU/NPU 作为服务器 CPU 的 PCI 从属设备，仅能以独占 GPU/NPU 卡形式与特定数量 CPU 核绑定作为云主机或云容器实例提供给云租户和开发者。对于单块 GPU/NPU 卡及多卡 GPU 服务器无法胜任的大模型训练场景，必须依赖服务器内 PCI 总线及跨服务器的 TCP/IP 或 RDMA 网络来实现更大范围 GPU 卡集群之间的紧密协同，严重限制了 GPU/NPU 集群的线性加速比及大模型训练的性价比。

基于数据中心网络异步 RDMA/UB 的百微秒级超低时延、10 倍降低尾时延、以及百 T 级超大带宽网络技术方面的突破，实现了 NPU/GPU 集群内的卡间 Full-Mesh 互联，大幅提升了大模型训练性价比。而通过软件定义的 GPU/NPU 池化算力，一方面可将一颗物理加速芯片（GPU 或 ASIC）切分为几个到几十个互相隔离的小的计算单元，也可将分布在不同物理服务器上的 GPU/NPU 芯片聚合给一个操作系统（物理机 / 虚拟机）或容器完成分布式任务，而没有加速芯片（GPU/NPU）的 CPU 服务器也可调用远程服务器上的加速卡（GPU 或 ASIC）完成 AI 运算，实现 CPU 与 GPU 设备的解耦，通过异构算力的统一池化，可提供更有弹性的 GPU/NPU 资源。

· **DCN 网络池化**

云数据中心采用叠加在物理交换网络之上的

分布式软件定义 Overlay 网络，实现了云数据中心 DCN 网络的全面池化，支撑了百万级服务器物理网络互联规模下千万级虚拟机、百万级租户 VPC 虚拟网络，在统一物理网络之上的，弹性按需动态复用，突破了传统硬件路由器及网关节点的水平扩展规模瓶颈，并实现了云业务负载逻辑网络地址与物理网络地址的解耦，以及高度灵活的网络互通 ACL 安全策略定义和实施。然而，正因为云数据中心普遍采用了 Overlay 软件定义网络与 Underlay 物理网络的分层式架构，从而增加了云数据中心内多租户、多应用网络技术栈的复杂度，当云租户应用的网络连接出现故障后，端到端故障定位也面临更大挑战。此外，当前云数据中心内承载多租户、多应用之间业务流量的网络传输及路由层，仍沿袭了互联网上存在了数十年并广泛应用的 TCP/IP 协议栈，尽管云数据中心已将以太 /IP 的服务器网络端口升级到百 G 级，物理交换机容量升级到 T 级，并引入了用户态 DPDK 及 DPU 硬件卸载机制，持续优化了多租户 Overlay 虚拟网络的吞吐及时延性能，但从云租户与应用端到端网络 QoS 视角来看，由于链路层以太协议缺乏高效的 E2E 流控能力，传输层 TCP 协议存在较高的用户侧网络包收发处理延迟与资源开销，以及丢包场景下低效的重传机制，已越来越无法适应 AI 大模型训练推理、元宇宙仿真渲染、搜索推广等日益普及的紧耦合、大颗粒云应用内部的分布式并发处理单元 / 微服务之间对极大带宽、极低时延、可预测网络延迟与丢包性能日益苛刻的要求，也制约了存储、内存及异构计算资源全池化效率的进一步提升。虽然业内已有 RDMA 网络技术可以部分解决了上述挑战，但在网络扩展规模及端到端精确流控方面，则仍远未达到云数据中心对

十万、百万节点网络扩展规模的要求。为解决上述问题，未来数据中心将引入全新架构的下一代云数据中心极简网络架构，支持基于 CPU/NPU 侧协议处理的无中断、DPU 全卸载，将云数据中心 DC 网络由 Overlay 虚拟网络 + Underlay 物理网络的复杂双层池化架构，简化为具备百万级节点平滑可扩展能力的轻量化、异步远端内存访问协议的极简、单层池化架构，从而最大限度减少不必要的跨多层协议栈的软件握手及开销，更好地支撑计算池化、存储池化，及紧耦合、大颗粒云应用的性价比倍增目标的达成。

2) 柔计算

“弹性计算”是当前云数据中心主流的算力分配与供给模式，即先由云服务商预定义固定规格的虚拟机或容器，再由云租户从中选取适合自己的资源规格，以及应用业务是否性能敏感选择 CPU 资源超分比，并基于该固定资源规格以及超分比进行资源消费的计量计费。该模式有利于最大限度减少算力分配的碎片化，但却由于为云租户分配的固定规格资源往往显著大于应用业务的实际资源需求，导致计算池的平均利用率仅 20%、远低于 80% 的平均分配率水平，全球云数据中心内数千万台服务器近一半以上物理算力实际上仍

处于空转状态，其能耗及碳排放并未产生实际价值。针对上述问题和挑战，未来云数据中心将引入相比“弹性计算”模式更为灵活、更智能、具备动态适配应用算力资源需求变化的下一代算力分配与供给模式“柔性计算”。预计 2030 年，领先的云数据中心资源分配粒度将达到“函数级”。在大幅提升云服务商的算力资源池有效利用率、减少算力资源空置率的同时，也能让最终云租户与开发者像用水和用电计费那样，为自己的动态算力消耗而非固定资源规格支付费用，降低不必要的算力支出与浪费。

“柔性计算”一方面强调极致的弹性，即算力不仅具备平滑的水平伸缩能力，也支持无缝垂直伸缩能力，同时也与业界“柔性制造”理念所强调的高度适应市场需求变化的生产模式相对应，强调基于普适化的应用精细化资源需求测量，以及应用 QoS 感知，具备更强定制化及快速需求应变能力算力分配与供给模式。实现“柔性计算”需要支持“应用驱动”精细化画像及初始资源分配、基于 AI 大模型的应用性能 QoS 劣化感知、柔性实例二次调度和柔性内存动态复用四个关键技术。

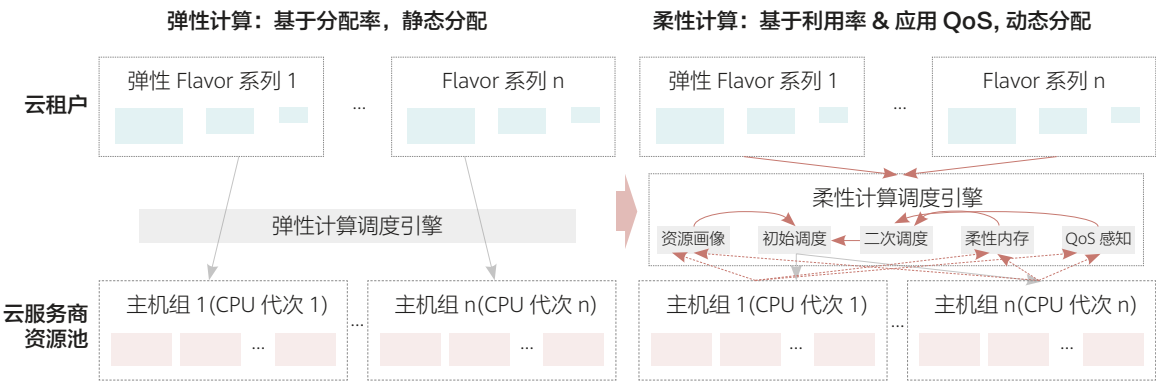


图 3-4 从弹性计算走向柔性计算

· **“应用驱动”精细化画像及初始资源分配**

柔性计算的算力资源分配，是以“应用驱动”为根本出发点，对租户业务应用负载的资源需求进行精细化洞察，可分为实例级和集群级两个层面：

在实例级的资源需求层面，无论该算力资源的最小颗粒度大小，以及是否云原生，柔性计算都打破了弹性计算固定 CPU 内存配比的桎梏，在实例发放前基于业务历史的资源用量采样进行画像，为实例设置最匹配其业务需求的精细化规格。实例发放后，主机还会持续监控实例的资源用量，对实例进行动态画像，调整主机的资源分配策略，持续确保主机资源的供需平衡。在实例级的“应用驱动”资源需求画像层面，考虑到不同云租户对性能 QoS 和成本的优先级诉求不同，以及应用性能 QoS 超限后第二次调度优先级差异，可设计不同优先等级的柔性实例，打破弹性计算超分与非超分集群无法统一共池，以及跨硬件 CPU 代次无法统一共池的制约，实现了跨不同硬件 CPU 代次、多优先级共主机混合部署与抢占式调度的“真正资源池化”。

在集群级的“应用驱动”资源需求画像方面，柔性计算打破了传统弹性计算模式下主要依赖人工介入及历史经验，甚至多轮试错来进行物理算力资源的规划置备，以及云主机、云容器的弹性按需伸缩容量规划与策略设定的模式。对于物理算力资源的规划，柔性计算以应用业务集群的历史任务起止时间序列及多任务并发特征的画像输入，利用排队论工具，计算出在一定量的任务排队概率下，集群所需要维持的物理算力资源水位线，省去了人工维护物理资源水位线的成本和试错

代价。对于虚机及容器弹性资源规划，柔性计算则利用 AI 时序预测或时频域机理式建模工具，对集群级的资源需求给出分钟级、甚至秒级的实时动态资源建模预测。

· **基于 AI 大模型的应用性能 QoS 劣化感知**

柔性计算区别于弹性计算模式的关键特征，除了业务负载动态资源需求感知之外，还体现在其对于业务负载性能 QoS 的量化感知能力。虽然柔性计算调度系统可依据对柔性实例的 CPU95 峰值、CPU 用量均值 + 方差 CPU 用量，以及柔性主机的累计 CPU 用量，分别对高、中、低不同优先级柔性实例进行首次资源分配，从而将业务性能冲突的概率控制在可预期的量化阈值范围内，但仍不能排除共宿主机的多个柔性实例之间因资源竞争导致租户应用负载的性能 QoS 劣化达到甚至超出可接受阈值边界的小概率情况，并在此情况下触发必要的二次资源调度，或租户应用负载与主机资源实例之间进行重均衡调度。因此，柔性计算调度系统云上租户业务负载的性能 QoS 量化感知能力必不可少。

柔性调度系统对于租户业务负载性能 QoS 劣化的感知，业务负载的应用层性能白盒式测量是最直接的途径，然而由于云上应用多数为非单体实例，普遍由多实例构成，导致仅基于业务负载的应用层性能指标变化并不能直接定位多个资源实例中每个资源实例对最终应用性能劣化的贡献占比，因此仍需对每个资源实例进行应用层性能 QoS 劣化感知。柔性计算可通过应用层非入侵的黑盒模式从底层主机 OS 收集的所有资源实例级的多维度性能指标，包括 CPU、内存、存储 IO、网络 IO，NUMA、L3 Cache Miss、前后端等待微指令周期等应用性能相关 Perf 系统指标

序列，并通过可持续在线迭代优化的上游预训练自监督学习任务与下游有监督任务微调及反馈式强化学习的 AI 建模方式，从采集观测到的海量多租户业务负载实例的性能流数据中，经特征提取、模型挖掘，提炼出能充分刻画租户业务负载性能特征，如 CPU 密集型、内存密集型、存储密集型、网络密集型，以及上述类型的特定组合等关键特征参数，并在下游有监督学习任务中通过少量已知典型业务负载训练样本，最终建立可黑盒式观测的柔性实例资源层性能 QoS 流数据，及其对应的业务负载实例的应用性能 QoS 劣化百分比之间的拟合关系。与 GPT 大模型可持续在线训练类似，柔性计算的黑盒式业务负载性能 QoS 劣化模型，同样可依据观测到的多样化业务负载性能特征，持续不断丰富其性能 QoS 特征库，从而支撑该模型具备对更多未知业务负载类型的性能 QoS 劣化程度的泛化普适能力。

• 柔性实例二次调度

当柔性调度系统通过多实例叠加 CPU 利用率，以及上述应用性能 QoS 劣化感知机制检测到共宿同一主机的多个虚机、容器实例的应用性能 QoS 劣化达到可接受临界阈值的情况下，则需触发二次调度以消除或缓解 QoS 劣化超限的情况，其二次调度策略分为 2 类：

资源层的业务无感热迁移或冷迁移，以及业务感知协同的二次调度机制。

资源层的业务无感热迁移 / 冷迁移的优势在于租户业务应用层软件无需为二次调度引入任何修改适配，缺点是面向虚拟机运行时的热迁移计算和网络资源开销较大，特别是热迁移持续时长与柔性实例内存大小、CPU 忙闲度，以及网络带宽 / 延迟等因素强相关；而面向容器运行时的跨主机 / 虚机 CRIU 冷迁移则因需要依赖共享 SSD 云存储进行进程级 CPU 运行时上下文的快照式持久化 IO 写入，及 IO 读出与内存加载，而面临更长时间的业务中断（热迁移为毫秒级，冷迁移则可能为百毫秒级），因此建议仅在物理服务器支持百 G 以上带宽、微秒级时延及 RDMA 类通讯协议，虚机 / 容器实例 CPU 利用率低于 60%，内存规格小于等于 16G 的前提下，优先选择性能 QoS 劣化百分比最大的柔性实例进行热迁移或冷迁移。

与业务无感的资源层二次迁移对比，业务感知协同的二次调度机制的优点，在于二次调度本身的计算与网络资源开销更小，通过业务层与资源层的配合，如对过载的资源实例进行业务分发隔离，可实现未完成任务或未完成 Web 会话的“优雅关闭”，从而保障



二次调度过程中的业务体验更平滑，缺点则在于需要应用层任务调度或 Web 负载均衡层软件的少量适配修改，当然，最佳选择无疑是将资源层 QoS 劣化感知驱动的二次调度机制，沉淀到云原生部署与治理框架内，所有基于该框架进行部署与流量治理的云原生应用，可在无需软件适配修改的前提下，实现资源代价最小的业务层二次调度。

· **柔性内存动态复用**

与 CPU、存储 IO 及网络带宽资源不同，云上租户虚拟机 / 容器实例的内存资源缺省采用独享模式，仍属 “非超分刚性资源”。未来伴随着云数据中心内存取代 CPU 成为资源池成本支出占比（约 2/3）最大的部件，如何实现内存资源跨租户、跨应用的安全高效动态服务，成为云算力利用率提升的关键瓶颈。

为解决客户机 Guest 与宿主机 Host 内存管理机制各自独立，导致各客户机实例内空闲内存资源无法被及时回收到主机侧，并按需分配给有需要其他客户机实例的问题，业界提出了内存气泡 / 空闲页报告的解决方案，然而该方案依赖客户机与主机之间异步非实时的内存资源空闲通知机制，仍存在主机内存资源临界条件下的突发性能瓶颈，以及内存资源释放不及时的风险。未来云数据中心通过引入柔性内存架构，将主机侧与客户机侧独立分层的内存页管理架构，演进到主机与客户机侧拉通的“扁平式内存管理”架构。在不影响虚拟机 / 安全容器实例间的内存页资源 Hypervisor 级安全隔离强度的前提下，建立主机侧与客户机侧实时同步共享的内存页空闲列表元数据信息，使得主机侧能实时感知所有客户机侧的物理内存页申请与释放状态，从而彻底解决内存气泡机制的性能突



发瓶颈与内存资源释放不及时的问题。未来数据中心基于 “扁平式内存管理” 架构，可以无侵入方式获取所有虚拟机、容器实例的物理内存页申请与释放历史数据，对其进行内存用量峰值、均值及标准差的统一画像，并与相关画像特征值进行高、中、低不同柔性实例优先级的动态内存复用分配，在面向中优先级柔性实例的物理内存页缺页中断处理中，透过百纳秒级超低时延、百 GB 级超大带宽的远端内存网络（UB/RDMA），触发其他服务器的远程空闲物理内存页资源的分配及远程访问，或直接基于 SCM/SSD 云存储介质进行内存交换区读写（相比 CPU 本地内存可带来额外 20%-30% 性能开销），而在面向高优先级柔性实例的缺页中断中，则同样基于百纳秒级超低时延、百 GB 级超大带宽的远端内存网络（UB/RDMA）触发中低优先级柔性实例的热迁移或冷迁移，或基于柔性资源侧与业务层协同的二次调度机制，从而为高优先级柔性实例腾挪必要的内存资源。

3) 泛协作

云数据中心计算、存储、网络资源的统一”全池化“，仅局限于单个地理服务区（Region）范围内，池化典型容量百万台服务器规模，由 50-100 公里范围内采用 10Tb 级光纤互联的多个可用区（Availability Zone）物理资源池构成。考虑到不同地理服务区在机房单位面积建设 / 租赁成本、单位耗电成本、PUE 水平、电网碳排放因子、一次性规划可扩容算力空间存在差异，因此可能导致不同地理服务区单位算力成本存在差异。以中国西部乌兰察布 / 贵阳及东部北上广一线城市服务区对比为例，二者相差超过 10%。如何突破地理服务区的物理边界制约，实现跨区域的“横向泛协作”，成为下一代云数据中心架构势在必行的重要发展趋势。

与此同时，随着移动终端及物联网的技术发展和普及应用，以及企业业务的跨区域部署，从互联网到政企客户均产生了从中心云服务区向距离云租户接入点更近的分布式边缘计算延伸的强烈诉求，进一步催生了云数据中心基础设施部署将从集中式向分布式演进，出现了跨大规模中心服务区与分布式边缘站点之间的“纵向泛协作”新发展趋势。

• 横向协作

为解决横向协作诉求，未来云数据中心将从云服务区（Region）感知的架构，向 Regionless 全域计算架构演进。

从云服务商的视角，将一定地理区域范围内的多个物理 Region 及分布式边缘节点（Edge）整合为一个“全域统一逻辑资源池”，通过全域资源调度引擎，拉通该逻辑资源池范围内所有物理 Region 及 Edge 节点，为租户的所有资源请求分配提供支撑，从而实现

极致优化数据中心投入产出比，以及更低的能耗与碳排放，解决跨物理 Region 资源分配不均衡及地域性算力资源供需不均衡的问题。

从云租户与开发者视角，为全域应用的开发、部署及业务路由提供唯一的“逻辑入口”，使得云租户与开发者无需再感知多个独立 Region 云服务 API 及开发框架入口，也无需再关注跨多 Region 及 Edge 节点的业务部署、资源发放，业务路由、数据同步，及相关广域网络连接管理与网络成本问题，从而不仅实现物理 Region 内的无服务器化、全自动弹性伸缩的能力，更进一步可为其提供跨物理 Region，跨云边的“单机式”的极简开发体验。

考虑到跨 Region、跨云边的广域网络与 Region 内的无阻塞物理网络相比，时延及带宽成本巨大差异，因此全域资源编排与调度层，需要定义应用资源实例及数据实例的接入时延（冷热温分层）及其相互间的亲和性，以及进一步拉通感知各数据中心 L1/L2/L3 综合 TCO 成本、广域网动态带宽成本，资源池总体分配率以及能耗、碳排放等基础信息形成云资源全域编排和调度的统一模型。



· **纵向协作**

轻量级边缘 / 分布式云的核心是跨地理位置以及用户可以将全部或者部分边缘基础设施交由云服务提供商进行运营和运维管理。在实施上具有全域协同的特征。

(1) 服务协同，也即统一租户业务应用，分布在主地理服务区 Region 及多个不同分布式边缘地理位置，需要跨不同站点进行微服务之间的协同：包括微服务跨边云发现和通信，通信全链路进行路由、限流、熔断等治理能力，灰度发布、金丝雀、蓝绿发布等典型发布流程；

(2) 业务管理协同，即用户可在云端对边缘侧的复杂业务模型进行管理，以集中式的管理方式降低管理运维成本；

(3) 数据协同，支持分布式站点与中心 Region 间的数据同步和共享，完成应用在不同位置分布式云之间的无缝衔接，满足应用在分布式云上使用数据的一致性；

(4) 资源协同，也即构建统一的分布式资源调度机制，能够根据不同位置的分布式的能力、位置、业务运行状态、资源使用情况，以及用户的习惯和意图，选择合适的站点进行资源分配、容灾编排；

此外在 IOT 物联网场景的边缘计算中，还需要提供云边端的资源协同管理，在云端统一管理边端的节点和设备，需要对节点、设备进行功能抽象，在云边端之间通过各种协议完成数据接入，在云端统一管理和运维。

· **全域弹性网络服务**

面向未来 5-10 年云数据中心的目标架构，我们预期跨云内不同地理服务区、跨不同云厂家异构云之间的横向协同，及跨云内中心地理服务区与分布式边缘站点、从云租户最终用户端到分布式边缘站点或就近地理服务区之间的纵向协同所带来的广域网络连接资源需求必将呈现百倍的增长，然而考虑到长距传输及路由设备的造价成本，广域网络不可能像数据中心内 DCN 网络那样作为不向云租户收取费用的“公摊成本”，广域网络属于单位带宽成本昂贵的稀缺资源，如何引入传输 QoS/SLA 保障、性价比最优的“全域弹性网络服务”以支撑上述横向协同及纵向协同带来的广域带宽互联爆炸式增长的需求就成为解决上述矛盾的关键。

“全域弹性网络服务”首先需突破传统光纤 / MPLS 专线连接的物理独占式及预定义最大带宽的静态式广域链路容量分配与路由模式，构筑具备云租户应用负载及数据资产的跨地理服务区，跨云边缘，乃至跨异构云的广域网业务互通及数据同步 / 异步复制流量的时域、空域特征感知能力的动态式广域链路容量分配与路由能力，并形成与开放互联网平面之间的“多活冗余”及“弹性按需”的互助关系。“全域弹性网络服务”需突破开放互联网 Internet 仅提供“尽力而为”的广域传输 QoS 保障，面向用户体验敏感的 Web 会话（远程 API 调用、Web 页面访问）及实时媒体类业务提供端到端实时时延优化、拥塞控制能力，为云租户构建“一站式就近入云”的极致性价比、极致体验保障。

5. 对等互联

1) 超融合

1978 年自 intel 开创 x86 体系以来，计算机体系经过 40 多年的高速发展，衍生出各种物理特性、传输特性和功能特性各不相同的互联协议。如图 3-5 所示，处理器间有 UPI、NVLink、CXL 等，处理器与外设及存储之间有 PCIe、CXL、NVLink、SATA 等，节点间有 Ethernet 和 IB 等。芯片需要为每种接口设计物理层和控制器，实现不同外设功能语义。当通信流量需要跨越协议接口时，协议间的转换会产生桥接的硬件代价、软件开销和功耗代价。

针对多样化连接诉求，需要建立新的超融合互联架构，打通所有芯片物理边界，减少协议转换开销，消除通信软件栈开销，实现更低通信延迟、更大通信带宽和更大互连利用率。

- 纵向打通 die 内协议
统一互连协议，减少转换，避免片上总线、PCIe 总线、网络端口带宽逐级收敛，使得端到端互连带宽是处理器端口直出带宽。

- 横向统一链路接口
提供统一的内存管理机制，让内存语义直达软件，各组件之间可以直接通信，互相调用，实现节点间数据高效流转，提升访存效率，减少通信开销。
- 以数据为中心，构建存算网融合架构
单一类型的计算资源，单一节点的计算能力、存储能力，以及配比固定的扩展模式已经难以满足日益复杂且快速变化的应用部署的需求。同时，海量数据的交互计算，也对数据中心算力效率和互联性能提出挑战，为了提升数据处理效率和存储资源利用率，未来数据中心需要走向“以数据为中心”，满足多样性计算，融合计算、存储、网络的超融合架构。

将计算、通信以及存储承载在统一协议栈上，打破传统分散架构限制，实现从通用计算、高性能计算和存储网络的三张网到一张网的融合部署，统一网络架构，推动无损网络向超融合网络架构演进，预计到 2030 年，超

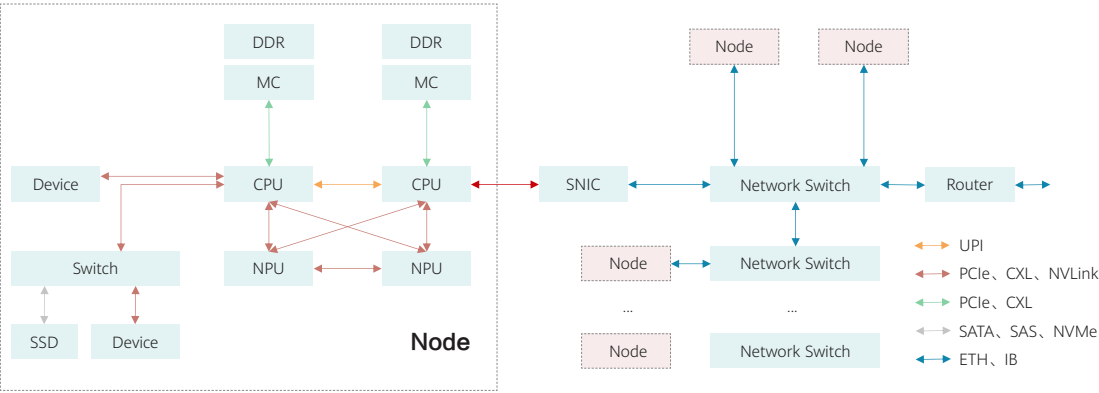


图 3-5 计算体系多种互联协议

融合以太网在大型数据中心的渗透率将达到 80%。

架构演进包括两个方面：1）在宏观上存算分离，计算、存储资源独立部署，通过高通量数据总线互联，统一内存语义访问数据，实现 CPU、GPU 等异构计算、存储资源解耦灵活调度，资源利用率最大化。2）在微观上存算一体，围绕数据，近数据处理，减少数据非必要移动，在数据产生的边缘、数据流动的网络中、数据存储系统中布置专用数据处理算力，通过网存算融合提升数据处理效率。

2) 高性能

2030 年，人类活动产生的数据量进入 YB 数据时代，从医、食、住、行、城市、企业和能源等多个场景展望，未来数据的存储、计算需求将高速增长。全球每年新增数据只有 5% 被使用，严重阻碍了数据价值的发挥，创新设计下一代高性能计算数据中心越来越重要。我们认为下一代高性能计算数据中心会有以下几个特征：

芯片到 DC 的统一可扩展大并行技术

当前计算芯片的算力增长已经落后于数据增长，根据第三方的研究显示，在人工智能方向，计算需求 2 年增长了 750x，摩尔定律驱动芯片算力只增长了 2 倍。因此下一代数据中心需要支持大并行来缓解芯片算力和应用计算需求之间的矛盾。依托大并行系统，大规模的数据集被分割成无数小块，数据中心的每块计算芯片只需处理其中一个小块，最终满足大数据的计算需求。当前数据中心的并行计算软件中间件分为两个层次：跨节点的 Spark、flink 和 hadoop，节点内芯片级有 CUDA、OpenCL 和 SysCL。未来的数据中心将会发展出从芯片到数据中心的统一可扩展的大并行计算软件中间件。

高速对等互联架构

在大并行系统中，新应用、新场景的时效性需求，需要计算芯片实时互相高速通信，交换计算中间结果，然而当前摩尔定律驱动的算力，互联带宽和内存带宽增长严重不平衡，根据第三方的研究显示，在过去 20 年，算力增长了 9 万倍，互联带宽和内存带宽只增

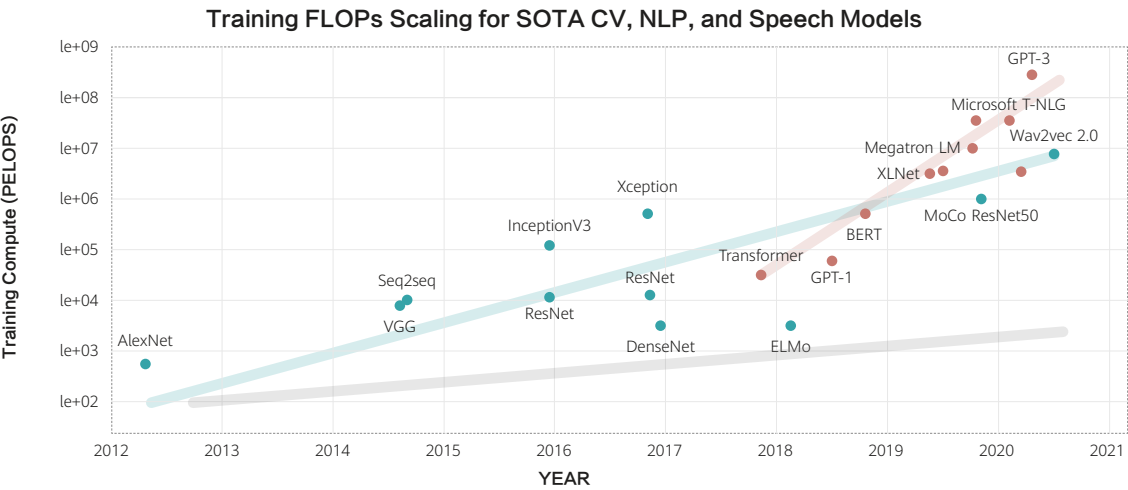


图 3-6 不同模型对算力的需求变化

长了 30 倍。下一代数据中心需要新的高速互联架构去缓解算力和互联带宽的不平衡，以芯片出光，光交换，动态 Torus 和光互联技术为基础，设计下一代高速互联架构，满足下一代数据中心高带宽和低时延的互联需求。数据中心将会发展出全新一代统一互联协议，消除数据通信协议转换代价，实现对等高速互联。

无损数据中心网络

ChatGPT 推动 AI 大模型进入万亿参数时代，远超单 GPU 芯片能力，不同 GPU 执行部分任务，并在不同 GPU 间共享结果，需要将大量芯片通过稳定低时延、零丢包的无损网络连接到一起，打造大规模算力集群。业界实践表明时延、丢包将严重制约了 AI 大模型的 GPU 利用率。无损数据中心网络已经成为研究热点，业界已经推出了专为 AI 设计的高性能以太网产品和芯片。

为了实现无损网络，在数据中心内部将引入超融合交换技术，实现零丢包、10us 级的低

时延的转发能力。为了保障超算等时延敏感类应用，数据中心网络设备可以参与计算信息汇聚和同步，通过算网协同降低通信时延，提升计算效率。

伴随着数据中心从一个节点走向网络化，跨数据中心也需要具备无损网络的能力，目前运营商正在探索算力与网络相互感知技术，网络可以参与到算力的调度过程之中，为时延敏感类应用提供零丢包、确定时延的通信保障。

芯片级长流水技术

为了缓解内存带宽不足的问题，新一代计算芯片需要减少访存频率，长流水技术通过将计算流程分成多个阶段，每个阶段并行处理不同的数据，实现芯片级的大并行，长流水的中间阶段数据不写回内存，减少内存带宽需求，缓解芯片算力和内存带宽之间的不平衡。

分布式多级缓存系统

面向数据中心 2030，需要发展支持分布式多

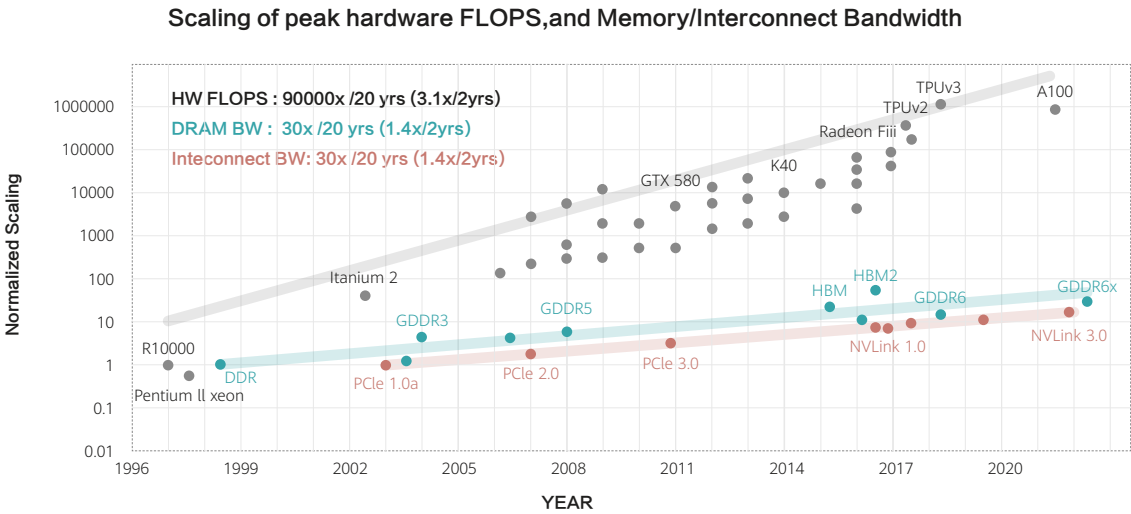


图 3-7 摩尔驱动的算力和互联、内存带宽的增长对比

级缓存的新一代缓存系统，持续挖掘数据局部性，减少数据中心长距离的通信需求。

新一代数据中心多级缓存系统由多个层级的缓存组成，每个缓存层级都有不同的容量和速度，通过分布式处理，可以提高计算芯片数据访问速度，减少计算芯片等待时间，并依据数据的使用频率和重要性进行自动管理，充分发掘数据中心的存储资源，提高数据中心的整体吞吐率。

3) 光内生

高算力芯片的 IO 带宽将越来越高，预计 2030 年，端口速率达 T 级以上。根据第三方的预测，2028 年数据中心内将实现 100% 的全光化连接。



随着单路速度提升，100/200Gbps 以上的高速串行通信带来功耗、串扰和散热挑战，传统光电转换接口将无法满算力增长需要，芯片出光在数据中心连接中的占比将持续提升，相比传统方案芯片出光端到端能耗有望降低至 1/3，成为未来突破带宽瓶颈，实现数据中心绿色发展的关键技术。

同时数据中心的网络架构也将发生变化，业界已

经开始研究新型的光交叉连接技术，通过利用光交换在带宽、端口、低功耗和时延等方面的优势，解决数据中心网络规模和流量带宽两个关键系统需求。

· 高速光接口 (1.6T/3.2T)

数据中心设备之间连接由高速光接口提供，并且根据连接距离不同分为 SR、FR、LR 等规格，不同传输距离采用的技术方案也会有所不同。高速光接口的速率发展与数据中心的交换机容量以及 SerDes 技术的发展息息相关。交换机容量每 2 年增长翻一番，预计 2030 年会出现 200Tbps/400Tbps 交换容量，单端口速率需要增长到 1.6Tbps/3.2Tbps。

光连接技术根据接收技术不同可以分为直检检测技术和相干检测技术。直检检测技术由于成本低、功耗低，在 800GE 之前，为数据中心高速光接口的主要技术。随着速率的提升，直检检测技术受到色散，四波混频等问题的影响，传输距离下降，使得相干技术下沉到数据中心存在了可能。在 800G 时代，IEEE 802.3dj 针对 10km 场景将会定义相干和直检两条技术路径。但相干技术面临功耗高以及成本高的挑战。未来 1.6T/3.2T 时代，直检技术和相干技术将同时存在。

直检检测技术在 1.6T /3.2T 时代仍是主力技术路径之一，并沿着 Scale Up 和 Scale Out 同时发展，在单 lane 速率持续提升的同时，通过增加光纤或者波分复用技术增加并行路数也将持续发展。800GE 时代延续单 lane 100G 技术并发展了单 lane 200G 的技术。1.6T/3.2T 时代将会依托单 lane 100G，单 lane 200G 技术进行多路复用，或会发展单 lane 400G 的技术。如 IEEE 802.3dj 已立项

16*100G 的 1.6TSR 的技术方案。也有公司陆续表达了对 8x200G 构建 1.6T 技术方案的预期，由于其采用 8 波复用的技术方案，将会面临着色散，四波混频等挑战，需要研究新的波长分配方案，色散管理技术，低功耗均衡技术等。对于单 lane 400G 技术，可采用高带宽器件，高阶调制格式，以及偏振复用等技术。

相干技术传统是长距光传输采用的技术方案。由于直检技术面临色散，四波混频等挑战，传输距离不断缩小。业界出现了相干技术下沉到数据中心互联的发展趋势。相干技术传输性能好，且可以灵活的采用 oDSP 进行色散的补偿，但是成本和功耗较高。为了降低成本和功耗，许多高校及企业提出了 Coherent-Lite 的概念。例如，利用 DFB 灰光源，量子点光源等低成本光源代替长途相干使用的 DBR 光源，进一步通过光源池共享光源来降低光源的成本及功耗。通过光域偏振跟踪方案来降低数字信号处理的复杂度，使用分段硅光调制器避免发端 DAC 等技术。

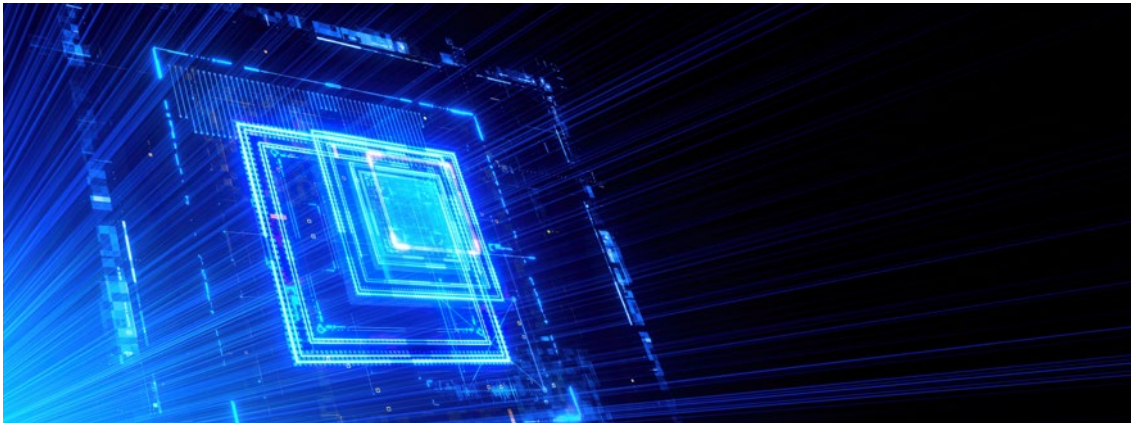
· **芯片出光**

单比特成本和功耗的降低是高速光接口技术发展的持续追求。过去十几年，交换机的容量提升了 80 倍，整体功耗下降 4 倍，其中 ASIC 功耗下降 10 倍，光接口的功耗降低了 3 倍。虽然光接口的单比特成本和功耗在不断下降，但是下降的速率远落后于交换机 ASIC 部分的功耗降低。究其根本原因，光接口依赖于 SerDes 技术，SerDes 为数模混合的技术，其能效演进低于 ASIC 部分。为了进一步降低功耗，必须要通过缩短 SerDes 的距离或者减少 SerDes 的数量来降低功耗，因此在光接口的系统结构上出现了很多新型技术如 OBO、CPO 等，芯片直接出光的 CPO 技术已经成为业界热点。

(1) 面向数据中心交换机的芯片出光技术——CPO 技术

目前主要的技术路径有两条，分别是基于硅光的技术路线和基于 VCSEL 的技术路线。

硅光技术因其集成度高、CMOS 工艺兼容有望实现低成本的特性成为多通道集成收发机的主要路径。针对硅光平台 CPO 技术光源部分主要有两种思路，一种是可插拔光源池模块技术，考虑到光源部分失效率较高，方





便后续更换，将多通道、大功率的激光器芯片封装后组装成可插拔模块置于面板侧，通过保偏光纤与交换芯片的四周的光引擎芯片连接，提供连续的激光源，这也是业界普遍认可的一种光源形态。另一方面，少数厂家具备较强的 III-V/Si 异质集成能力，能够直接在硅光引擎上实现光源的集成，通过采用 2:1 备份的方式改善光源的良率，该方式成为第二条光源技术路径。针对硅光平台的高速调制器部分，当前主要有三种技术路径：第一种是相对成熟的 MZ 调制器技术，由于 MZ 尺寸较大（百 μm 量级），多通道集成后，光引擎尺寸较大，功耗相对偏高；第二种是微环调制器技术，微环具备小尺寸（几十 μm 量级）、低功耗（驱压小）的特点，但是微环调制器需要非常稳定的工作波长跟踪系统；第三种是基于 Ge 材料的 EA 调制器，调制器尺寸也在几十 μm 左右，通过法兰兹-卡尔迪西 (Franz-Keldysh) 效应实现对光的吸收。

业界部分厂商也在推动基于 VCSEL 的 CPO 技术，主要原因在于 VCSEL 具有优异的功耗特性（ $< 5\text{pJ/bit}$ ），基本可满足 100m 以内的互联需求，后续通过器件进一步升级为少模或单模的 VCSEL，也有望能够实现 km 级互联长度。当前，VCSEL 较为成熟的器件为 25GBd 量级，后续 50GBd 有望在近几年成

熟商用，虽然带宽发展趋势上略慢于硅光技术，但 VCSEL 技术可以通过外置合分波器实现波分复用以提高单纤容量，也可以通过阵列化的 VCSEL 器件 /PD 器件配合多芯光纤（ $\sim 40\mu\text{m}$ 芯间距）实现大容量传输。

（2）面向高性能计算的芯片出光技术——光 I/O 技术

高性能计算集群是由高速通信网络连接的强算力平台，高速互联网络的通信能力已经成为 xPU 集群的重要支撑，如何进一步提升互联带宽成为业界关注的重点，光 I/O 技术开始步入大家的视野，该技术通过将光学收发芯片放进计算芯片封装内，因此也被称为封装内光学连接技术（In-packaged Optics）。通过采用该技术可以大幅改善芯片扇出带宽，降低光互联功耗，实现可媲美板内 / 框内电互联的带宽密度 / 功耗水平，同时，又能提供电互联无法达到的互联距离（ $\sim \text{km}$ 级），为集群系统互联提供了一种低功耗，大容量的新技术路线。光 I/O 技术的具体实现技术路径以硅光技术为主，具体为采用低调制速率（30-60Gbps）的微环总线型波分技术。一方面，在该调制速率区间，具有相对最优的端到端功耗水平（ $\sim 5\text{pJ/bit}$ ），另一方面，利用微环本身的窄带工作特性，实现多路合一的波分型总线，可以大幅扩展边缘互联带宽密度，很容易达到百

Gbps/mm，甚至 Tbps/mm 互联密度。当前，该领域主要研究热点聚焦于密集波分微环调制器的实现，多通道微环调制器的控制，多波长外置光源技术以及先进封装技术等多个技术方向。

· **光交叉连接**

近年来，业界和学术界广泛研究新型的光交叉（OXC）连接技术，通过利用光交换在带宽、端口、低功耗和时延等方面的优势，解决数据中心网络规模和流量带宽两个关键系统需求。OXC 主要技术方向分为波长级交叉连接和光纤级端口交叉连接，面向 2030，在数据中心场景的重点研究方向是 MEMS OXC 以及亚 μs 快速光交叉技术。

（1）MEMS OXC

MEMS OXC（Micro-Electro-Mechanical Systems Optical Cross-Connect）是一种基于微机电系统技术的光交叉系统设备，由一对光学准直器组成阵列作为输入和输出（I/O）端口和一对 MEMS 微镜阵列芯片来控制光束，以便任何输入端口都可以连接到任意输出端口，具有高集成度、高速率、低功耗的特点。

（2）片上集成光开关

基于片上集成快速光开关使用的关键技术可将片上集成光开关分为五类：热光效应（thermo-optic effect）；硅基载流子效应（free carrier effect）；泡克耳斯效应（Pockels effect）；克尔效应（Kerr effect）；波导型 MEMS 技术（Si-MEMS）。

热光效应，即利用材料晶格结构对温度敏感的特性，实现对材料折射率的调控。制成的光开关可实现 $100\mu\text{m}$ 量级的超紧凑尺寸、

10mW 级开关功耗、亚微秒（sub- μs ）量级切换时延；载流子效应是基于硅材料的一种特殊效应，光开关长度在 $300\mu\text{m}$ -mm 量级，开关时延为 ns 级；泡克耳斯效应和克尔效应均属于非线性光学效应，非线性光学效应光开关电光响应时间在 ps-fs 量级，且不产生额外损耗，但是需要更高的驱动电压或更长的器件尺寸；硅波导微型 MEMS 系统（Si-MEMS）依靠静电力对悬空波导结构的吸引/排斥行为直接改变波导物理间距从而改变光路径，光开关切换速度为亚 μs 量级。相比其他技术，Si-MEMS 可提供更高的隔离度和更低的损耗，并且允许结构尺寸更加紧凑，但 Si-MEMS 依靠移动波导或金属电极结构，限制了开关可靠性和耐久性。

· **新光纤介质**

下一代数据中心互联的发展趋势是高速率、高密度、低时延、低成本和易运维，新型光纤的应用将对数据中心光互联产生革命性的影响。其中空芯光纤和多芯光纤，由于其特殊和优异的光纤特性，将进一步推动数据中心实现更低时延、更高密度、更低成本的光互联。

（1）空芯光纤

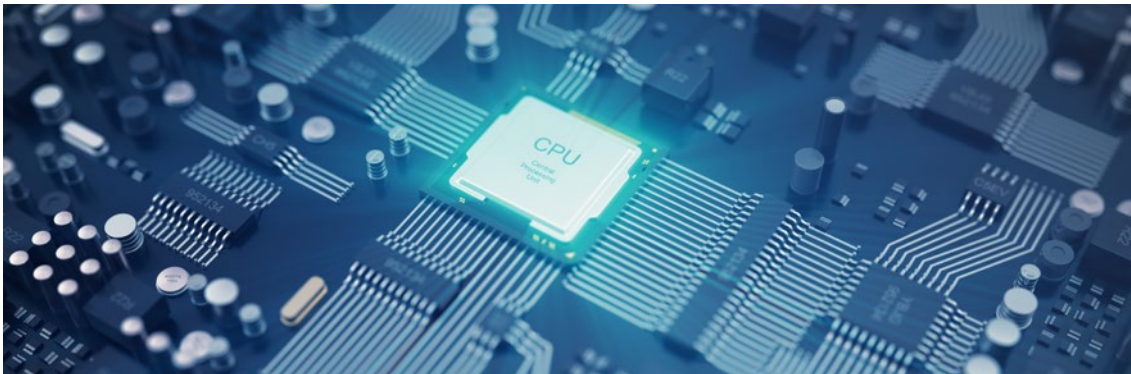
空芯光纤突破了传统石英光纤的局限，其基于反谐振机理，通过特定的包层结构设计，可将光限制在空气纤芯中进行传输，改变了光在光纤中的传输介质，从根本上避免了因材料本征限制而带来的问题。相比实芯光纤，空芯光纤具有低时延、低色散和低非线性等优点。首先，光在空气纤芯中的传输速度是光速，是在玻璃介质中传输速度的 1.5 倍，可大幅缩短 AI 数据中心内各个服务器以及 GPU 之间的通讯时延；其次，由于空芯光纤

的传输介质是空气，材料色散低，有助于扩展数据中心内高速光模块的传输距离，降低光互联成本；第三，与低材料色散类似，空气相对于二氧化硅等玻璃材料，其非线性折射率系数小，具有更低的非线性效应，极大程度地抑制了数据中心内光互联产生的信号畸变，保证更好的通信和网络质量。

(2) 多芯光纤

多芯光纤是多个纤芯共享一个包层，其中每个纤芯都是单模，且纤芯之间的串扰很小，这将使密度比传统单模光纤提高数倍。在多

芯光纤中，多路光信号可分别在多个纤芯中同时传输，信号之间串扰小，极大地提高了通信容量，其应用将对数据中心光互联产生革命性的影响。单模代替多模、多芯代替单芯、COBO/CPO 代替热可插拔将是数据中心未来的布线趋势。多芯光纤有潜力成为未来超 800G 的互联方案，可大幅提升光传输容量和频谱效率，节约布线成本和管道资源、降低能耗，且具有多个平行的物理通道，在下一代数据中心布线中更具应用潜力。



■ 6. 系统摩尔

1) 大小芯

半导体集成电路是现代信息产业的基石，但主导其发展的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，亟需采用新型芯片技术推动未来信息产业持续蓬勃发展。

· Chiplet 小芯片封装技术

提升芯片算力和产品竞争力传统芯片受 wafer（硅片）曝光尺寸限制（1 Reticle：

25mm*32mm），芯片 Die 的尺寸及 Die 良率提升受到严重技术瓶颈，直接制约芯片整体性能提升及芯片成本降低。2.5D Silicon/FO Interposer+Chiplet 技术可以有效提升 Die 良率、降低芯片成本，通过小芯片堆叠集成实现更大规模芯片性能，且对于不同产品规格应用更加灵活。同时 2.5D 封装性对于传统封装板级互连方案单 bit 能耗降低至约 1/2。

基于行业发展与超大规模芯片需求，预计 2025 年 2.5D silicon/FO interposer 尺寸将超过 4xReticle，未来封装 substrate（基板）预计会超过 110mm*110mm。更大尺寸的 2.5D 与 substrate 应用直接面临良率、交期、可靠性等一些列工程难题，融合创新基板架构需求迫切。

• **3D 大芯片制造**

与传统 2D/2.5D 先进封装及异质集成芯片技术相比，3D 芯片技术在互连密度及带宽、芯片尺寸、功耗性能、芯片综合性能方面优势显著，是解决未来高性能计算、AI 等关键场景芯片与系统集成的核心技术。3D 芯片技术未来会从 D2W（Die-to-Wafer，芯片到晶圆）->W2W（Wafer to Wafer，晶圆片对晶圆片），uBump->Hybrid Bonding->Monolithic 3D 技术逐渐演进，应用场景将会广泛覆盖 3D Memory on Logic、Logic on Logic 及 Optical on logic 等，并且未来会逐步走向更多层异质堆叠。3D 芯片在堆叠工艺方面需要采用小于 10μm 甚至更小 pitch 超高密 Bonding 技术，3D 芯片相对于传统 2.5D 封装在带宽及功耗性能优势显著，

单 bit 功耗降低有望降低至 1/10。更小尺寸 TSV（Through Silicon Via，硅通孔）技术需要从材料、工艺基础技术深入持续探索；同时 3D 堆叠带来局部功耗密度和电流密度倍增，直接影响系统整体供电与散热路径。

从发展节奏看，基于小芯片集成的 Chiplet 技术将最先成熟应用，未来伴随工艺和技术的成熟，3D 大芯片将逐渐崭露头角。

2) 新算力

登纳德缩放定律在硅基半导体上已经失效，如何延续或超越摩尔定律成为计算领域的重大挑战，学术界、工业界都在寻找新的计算范式，通过探索模拟计算、非硅基计算等来提升计算能效。

• **量子计算加速工程化**

量子计算硬件目前处于高速工程化的阶段，量子比特数快速增长，预计未来五年将出现超过 10000 物理比特的量子芯片。在含噪声的中等规模量子（NISQ）时代，构建经典计算机与量子计算机混合的计算系统是最具可行性的技术方向。其中量子模拟、量子组合优化算法及量子机器学习三大方向是业界



主流的应用场景。量子模拟能为药物研发与新型材料研发提供全新的计算范式；量子组合优化算法充分利用量子计算的并行计算能力，能更快更好的解决物流调度、行程规划及网络流量分配等问题，量子机器学习将成为人工智能计算加速的新路线。未来十年在硬件上需要不断提升单量子芯片的物理比特规模，增强量子比特的相干时间和量子操作的保真度，并通过量子芯片互联提升系统的扩展能力。在软件和算法上一方面要完善量子软件栈，另一方面需要结合应用场景优化量子算法，降低线路深度和复杂度，逐步推动 NISQ 量子计算走向商用。此外，还要逐步增强量子计算的容错设计，提升量子系统可靠性。但要实现一台通用量子计算机，道路更加漫长、更加充满挑战。

• **模拟光计算构建光电混合加速器**

光的传播速度快、能耗低，其干涉、散射、反射等物理现象背后，都有对应的数学模型，通过对光信号的调制、控制、探测，可完成某些特定的计算任务。同时光作为玻色子天然具有波分复用、模分复用、OAM（Orbital Angular Momentum，轨道角动量）复用等特性，通过模拟光计算实现多维度并行，是未来光计算发展的重要方向，有望在光信号处理、组合优化、AI 加速等场景的提供计算加速能力。光计算要实现规模应用，首先需要解决有源 / 无源器件在芯片上的异质集成问题，提升光信号耦合效率、控制插损和噪声、满足特定应用场景的计算精度要求，并基于此构建光电混合系统，实现特定计算任务的加速。

• **非硅基计算逐步走向规模应用**



二维材料晶体管具备沟道短、迁移率高、可 2D/3D 异质集成的优势，有望作为晶体管沟道材料延续摩尔定律至 1nm 节点。此外具有超低介电常数的二维材料：也可以用作集成电路的互连隔离材料，二维材料可能首先在光电、传感等领域应用。当前二维材料及其器件仍处于基础研究阶段，未来五年首先需要解决工业级二维材料晶圆制备的良率问题，其次要不断改善电极和器件结构，提升二维晶体管器件综合性能；碳纳米管具有超高的载流子迁移率、原子级的厚度，具有高性能、低功耗的优势，在尺寸极端缩减的情况下，碳管晶体管能效比硅基晶体管提升约 10 倍，5 年内有望在生物传感、射频电路实现小规模商用。未来还要继续改进碳管材料的制备工艺，降低表面污染和杂质，提升材料纯度和碳管排列的一致性；优化器件接触电阻和界面态，提升注入效率。当碳基半导体器件的尺寸能够微缩到与硅基先进工艺相当水平时，在高性能、高集成度的应用场景中，将迎来规模应用的机会。

3) 新存储

随着 Big Data 和 AI 的广泛应用，数据驱动的计算受到高度重视，数据的价值得到广泛认可。但数据存储系统面临两大挑战：一是如何快速满足计算单元的数据处理需求；二是如何低成本长周期的保存数据。为了应对这些挑战，新的数据存储有望通过多样化的存储介质和以数据为中心的体系架构，进一步发挥数据价值。

· 多样化的存储介质

预计到 2030 年，全球每年新增 1YB 的数据，其中有接近 50ZB 的价值数据需要存储，相比 2020 年增长 23 倍，要求存储介质必须具备大容量、高性价比、低能耗，要求存储系统具备高可靠、高扩展、长寿耐用和高安全性，同时具有数据计算和分析能力，以便更快的获取数据。

围绕着数据全生命周期的热温冷差异，未来介质将向高速高性能、中速大容量、低速低成本三个方向演进。

(1) DRAM 仍是高速高性能介质主流选择

当前性能最好的存储器依然是 DRAM，随着制造工艺演进到 1 α 制程，单位面积存储容量达到 0.315Gb / mm²。由于 DRAM 结构中电容尺寸过大，平面微缩基本接近极限，业界开始研究 3D DRAM 工艺、晶圆减薄和混合键合技术，以期进一步提升存储密度、降低功耗；与此同时，业界在新型非易失存储器上的研究从未止步，FeRAM、MRAM、ReRAM、PCM、氧化物半导体存储器等都取得了不错的进展。FeRAM 已有 Mb 级产品，以及采用 1x nm DRAM 工艺的 8 Gb 阵列原型展示；MRAM 已有 Gb 级独立式、Mb 级嵌入式产品，当前面向 SRAM/DRAM 缓存等

场景进行研发；PCM 已有 512GB 3D Xpoint 产品，用于持久化内存或 SCM(存储级内存)；ReRAM 已有 Mb 级独立式产品及嵌入式量产准备，同时面向存算一体正在研究；氧化物半导体如 IGZO 可构建 2T0C DRAM，有望通过 3D 堆叠实现 <4F² 单元密度。总体来看，这些新型存储介质具备非易失、低功耗的特点，能大幅降低存储器功耗，但同时存储密度或擦写次数上，与 DRAM 还有较大差距，到 2030 年 DRAM 仍然是主流的高性能存储介质。

(2) 中速大容量介质 SSD 具备明显优势

基于磁存储技术的 HDD 一直是大容量存储介质的主流，随着铁铂合金薄膜介质以及热辅助磁记录 (HAMR) 和微波辅助磁记录 (MAMR) 技术的逐步成熟，3.5 吋机械硬盘的存储容量将从当前的 30TB 演进到 80TB，由于增加了激光和微波等辅助技术，HDD 未来的每 TB 成本将不会有明显的降幅；受益于半导体制造工艺的不断进步，以及 3D 堆叠技术的创新，3D NAND 在中速大容量介质上后劲十足。当前 QLC 已经规模发货，PLC (Penta Level Cell) 技术也提上日程，在单个 CELL 中存储更多 Level 的数值是不断突破的方向；与此同时，3D NAND 目前已经堆叠超过 230 层，未来十年可能达到 1000 层，因此 3D NAND 的存储密度还能持续提升，预计到 2030 年，SSD 的每 TB 成本将与 HDD 相当，且读写时延和带宽有明显优势，数据中心采用 SSD 全闪存存储的占比将进一步达到 80%。

(3) 低速低成本介质以磁带和光盘为主

随着数字化进程的推进，数据中心汇聚的数据量成指数增长，人工智能的兴起，使得数据的价值再一次被发掘。同时由于政策法规

的需要，很多数据都被要求保存 30 年以上甚至更长时间，因此“古老”的磁带和光盘存储，再次受到数据中心运营者的重视。磁带因为制备工艺简单、可用存储面积大，在每 TB 成本上有绝对优势，同时磁带可以借鉴硬盘的磁头和磁粉制备技术，不断提升容量密度，确保技术可持续演进，当前商用的 LTO9 单盒磁带容量 18TB，未来可演进到 576TB。由于格式兼容和介质寿命等原因，磁带需要每隔七年将数据拷贝到新磁带上；光盘存储在蓝光 BD (Blu-ray Disc) 技术基础上发展出 AD (Archival Disc)，实现单盘容量 500GB/1TB，可保存 50 年以上，12 张盘容量密度与当前的磁带接近。随着介质材料和伺服技术的进步，未来单盘可演进到 2TB/4TB。磁带和光盘将是数据中心冷数据保存的重要介质，磁带的优势是成本低，光盘的优势是保存时间长。

在大数据、人工智能、HPC、IOT 等新型数据密集型应用的推动下，数据量爆炸增长，年复合增长率近 40%，其中热数据占比将超过 30%；另一方面，摩尔定律、Dennard 缩放定律的放缓，CPU 性能年化增长降低至 3.5%。高速增长的数据与缓慢增长的数据处理能力成数据产业的基本矛盾，数据存力与数据发展严重失衡。

在传统的以 CPU 为中心的体系架构中，业务在空间、时间的不均匀性导致本地存储资源利用率低，本地内存、存储闲置率超过 50%，数据的移动、数据格式的反复转换消耗了大量 CPU 时间，使得数据处理效率低下。

为了提升数据处理效率和存储资源利用率，未来数据中心体系架构需要从“以 CPU 为中心”走向“以数据为中心”，包括四个方面：

以数据为中心的体系架构

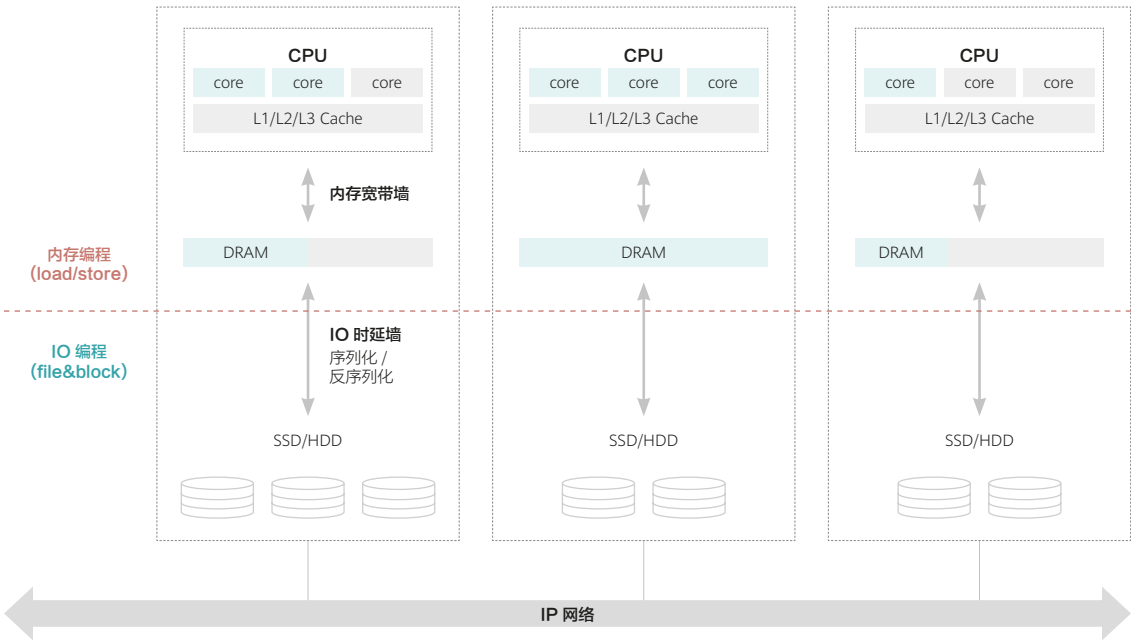


图 3-8 传统以 CPU 为中心的体系架构

(1) 数据中心内存算分离

计算、存储资源独立部署，通过高通量数据总线互联，统一内存语义访问数据，实现计算、存储资源解耦灵活调度，资源利用率最大化。存算分离不再局限于 CPU 与 SSD、HDD 外部存储解耦，而是彻底打破各类计算

存储硬件资源的边界，将其组建为彼此独立的硬件资源池(例如 CPU 池、DPU 池、内存池、闪存池等)，实现各类硬件的弹性扩展及灵活共享。存算分离架构具备三个特征：存储资源池化、全内存语义访问、高通量对等互联总线。

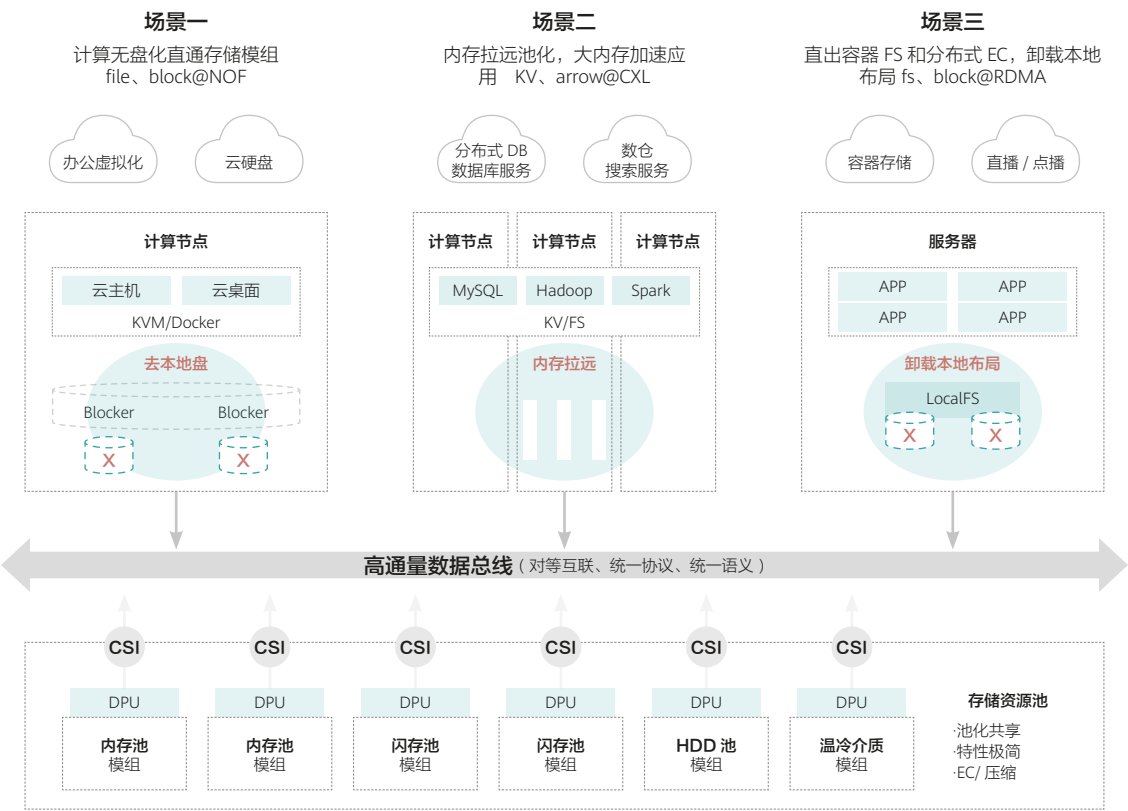


图 3-9 存算分离架构

(2) 存储系统内数控分离

数据面和控制面分离，CPU 只处理控制面的任务，数据面的由 DPU 等异构算力处理，避免反复的上下文切换，提高数据读写性能。

传统存储以 CPU 为中心设计，数据读取、写入都要经过 CPU，使 CPU 成为系统性能的瓶颈，无法满足新兴应用越来越高的性能诉求。

存储 IO 处理可基于 IO 直通等技术，数据处理路径可从智能网卡、DPU 直通到盘，实现前端卡到后端介质的快速直通，构建极简的快速数据访问路径，从而减少 IO 路径 CPU 的参与，时延和吞吐挑战理论极限。

(3) 跨数据中心智能数据编织

数字技术的不断发展催生了大量的跨域数据

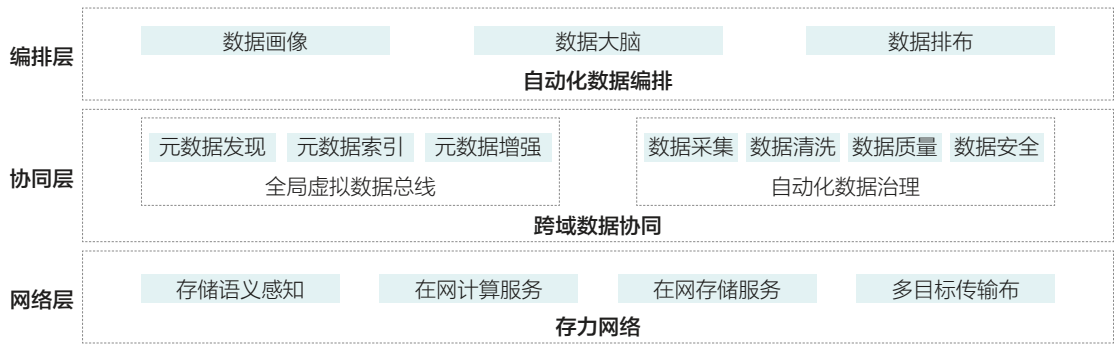


图 3-10 智能数据编织框架

流动的需求，对数据的可用性和质量提出了更高的要求。但地域的阻隔和数据治理的困难限制了数据的自由流动，最终形成了数据重力。数据编织是以一种自动化的方式，动态地协调分布式的数据源，跨数据平台地提供集成和可信赖的数据，支持广泛的不同应用的使用。

智能数据编织可基于人工智能和知识图谱等技术，不断识别和连接来自不同应用的数据，以发现可用数据点之间独特的业务相关关系。在数据网络中，边缘、数据中心、云端频繁的数据交换，智能数据编织可通过对现有的、可发现和可推断的元数据资产进行持续分析，完成跨平台的数据整合，为应用提供高效数据流动和处理。为了更好地实现智

能数据编织，需要在跨域数据协同、自动化数据编排和高效快速存力网络等技术方向上持续突破，以解决数据重力问题。

（4）面向应用的数据加速

在以数据为中心处理范式中，数据处理由通用计算走向数据处理专业化，由数据搬移到处理器走向近数据布置算力，在靠近数据的地方，以最合适的算力来处理数据，在数据产生的边缘、在数据移动中、在数据存储中就近完成数据处理，预计到 2030 年，采用近存和存内计算的数据将超过 80%。数据存储作为数据载体，不仅提供数据存取服务，还提供近数据处理加速服务，数据就近处理有三种主要方式：多样化存算融合、数据存储与网络融合、数据处理与网络融合。

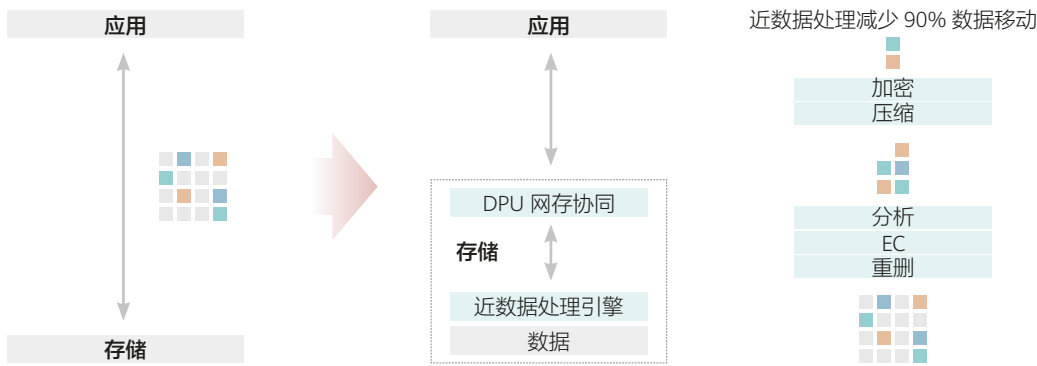


图 3-11 面向应用的数据加速





新型数据中心 参考架构

04



面向 2030，随着千行万业智能化对算力需求的急剧增长，以及算存网云、制冷及绿色供储能等数据中心相关技术创新的持续加速，驱动数据中心功能定位发生着深刻变化，如数据中心由封闭孤立到需要更广泛的参与社会化算网协同、由粗放式资源管理到需要为各类应用提供更精细更高效的算力供给、由数据孤岛到需要保障大规模数据的跨域安全可信流通、由以 CPU 为中心到以数据为中心的的计算架构等，这些变化在影响数据中心现有架构的同时，为企业数据中心的建设带来巨大挑战，为此，我们提出了新型数据中心“6 新”参考架构，具体如下：



图 4-1 新型数据中心“6 新”参考架构

■ 新基础设施，供电制冷走向全天候绿色零碳

随着数据中心规模的不断增大，数据中心耗电量将持续攀升，并带来供电、制冷等诸多挑战。首先，在供电方面，绿电比例低、为保证可靠性带来的电网利用效率低、供电损耗环节多以及使用大量的柴发备用电源，最终实际用于 IT 的有效电力普遍不足 80%。其次，在制冷方面，当前数据中心大部分时间依靠压缩机制冷，制冷效率低下，并且静态的制冷架构难以匹配算力快速的变化需求。此外，数据中心产生的余热利用难，大量废弃。应对这些问题和挑战，面向 2030，我们提出了具备全天候零碳、极低的能源损耗、更灵活弹性制冷等核心功能的新型数据中心参考架构。

面向 2030，新型供电系统通过采用长时储能、氢燃料发电机、本地光伏等，与虚拟电厂形成“源

网荷储”互动，使电网能够充分利用数据中心多余的电力储备来满足不断变化的负荷需求，辅助解决风电、光伏随机性和间歇性问题，提升大比例清洁能源电网稳定性和利用效率，实现数据中心近 100% 绿色供电。供电系统也将进一步融合，减少损耗，数据中心中实际应用于计算的电力占比将提升到 95% 以上。

面向 2030，新型制冷系统通过风冷、液冷兼容性架构设计，支持风液动态灵活调配，能够更好的匹配急剧增长的算力需求。通过降低传热温差，因地制宜充分利用干空气、湖水等自然冷源，将实现近 100% 自然冷却，制冷能效提升 2~3 倍。由于余热品味提升及余热发电、余热配套产业合理规划，100% 余热利用将成为可能。



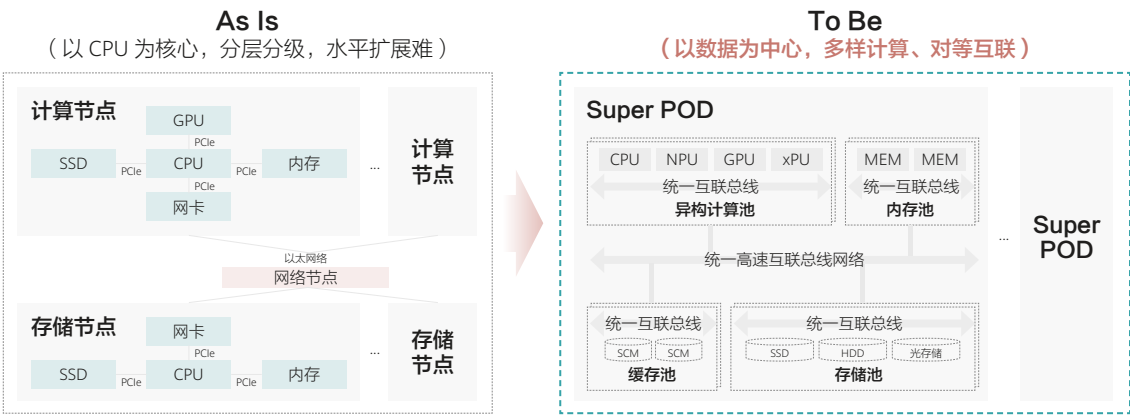
图 4-2：新型数据中心全绿色供电及动态制冷参考架构

■ 新算力底座，构建以数据为中心多样算力系统

当前，数据中心大多还是采用传统的计算、存储、网络分层多级架构，每一层都独立围绕CPU、内存、总线、硬盘等组件构成的完整计算机系统。该架构普遍存在三堵“墙”问题，即：内存墙、I/O 墙和算力墙。这些“墙”也造成了数据访问和搬移速度瓶颈，且限制了大规模分布式水平扩展。

面向 2030，下一代新型数据中心计算架构，将从

以 CPU 为核心的分层分级架构走向以数据为中心的对等互联多样计算架构。该架构基于“内存语义”构建统一的高性能、可编程、规模可扩展的互连网络 / 总线（uBus Fabric），聚焦数据的搬移、转换和分发，突破“内存墙”和“I/O 墙”，释放 CPU 和异构加速器算力，做到计算和网络深度融合，共同构成一个高效的超级计算机系统。



■ 新资源调度，应用为中心实现柔性调度

正如每台计算机都有操作系统（Operating System，简称 OS）来调度 CPU、内存、硬盘等硬件资源一样，数据中心也有“操作系统”来为整个数据中心提供分布式调度与协调功能，实现数据中心级弹性伸缩能力，它将数据中心的资源当作一台计算机来调度。数据中心操作系统的发展历经最早的物理机时代，到了当前的虚拟化 / 云化时代，未来将走向以应用为中心的时代。

物理机时代的数据中心操作系统独立于每个物理服务器设备上，每台服务器上运行一个应用程序，

这时单台服务器的性能限制了应用程序的部署规模。虚拟化时代的数据中心操作系统以虚拟机为单位，将资源提供给用户。虚拟化的操作系统可以将一台高性能的服务器虚拟成多个虚拟机，在物理上虚拟机共享宿主服务器的硬件资源，而逻辑上各自独立，可在各虚拟出的服务器上运行不同的应用，各司其职，互不干扰。如此一来，大大提升了服务器的使用率，降低了数据中心的运营成本。如今的数据中心操作系统里到处都是虚拟化的身影，核心技术有 SDS、SDN、

OpenStack 等等。但是，虚拟化构成的集群难以运维，尤其是出了故障后，很难分析出故障原因和位置。

面向 2030 越来越多元化的应用场景，用户希望能够直接获取资源、快速启动、服务可以无限扩展、应用易于迁移。这时一切以应用为中心，将数据中心甚至需要将多个数据中心的算存网资源进行整合，CPU、NPU、GPU、内存和 I/O 这些基本资源都进行池化，根据各个应用，按需分配。并且，人工智能、科学研究以及元宇宙等新兴领域快速崛起都对算力提出了更高要求。据预测，未来 3~5 年，十万亿、百万亿参数的 AI 大模型将出现，单中心算力将无法满足 AI 训练的需求，需要通过集群方式突破单点算力的性能极限。如何打破数据中心物理上的“四面墙”，实现跨

DC 集群算力资源灵活调度和快速应用部署，成为了产业界共同关注的热点。

要突破单个 DC 的资源和平台限制，使能大规模多中心分布式应用。首先，要构建跨多中心间的高速互联网络，实现域内多个数据中心间的超低时延、超大带宽、超高可靠的互联，充分保障数据和任务的快速调度流转。其次，要构建以应用为中心的下一代数据中心操作系统，一方面能够提供跨 DC 硬件资源的抽象与协同，充分释放硬件能力；另一方面能够提供诸如实例精细画像、负载动态监测、AI 性能 QoS 感知以及柔性资源调度等精细化、智能化的资源统筹管理功能，提升全局能效；最后要能够提供大规模分布式应用部署工具和运行框架，使能分布式应用的高效运行和快速部署。



图 4-4 新型数据中心资源调度系统参考架构

■ 新数据管理，数据全局可视助力高效流通

面向 2030，数据跨域流通的诉求越来越强烈，但存在效率、安全、协同、管理等挑战：一是存在大量数据孤岛，缺乏数据全局视图，导致数据利用率低，价值难以挖掘；二是缺乏分级的热温冷数据流动技术，数据中心间数据流动困难；三是跨域数据协同困难，缺乏跨地域统一元数据管

理，无法支撑数据并行分析；四是数据存储效能不高，数据存储成本高，数据处理性能不足以支撑跨域查询和分析。由此，需要一个跨域、跨 DC、跨存储形态的逻辑统一的数据湖，结合数网大脑，实现数据全局可视、跨域安全高效流通和自动分级最优放置。

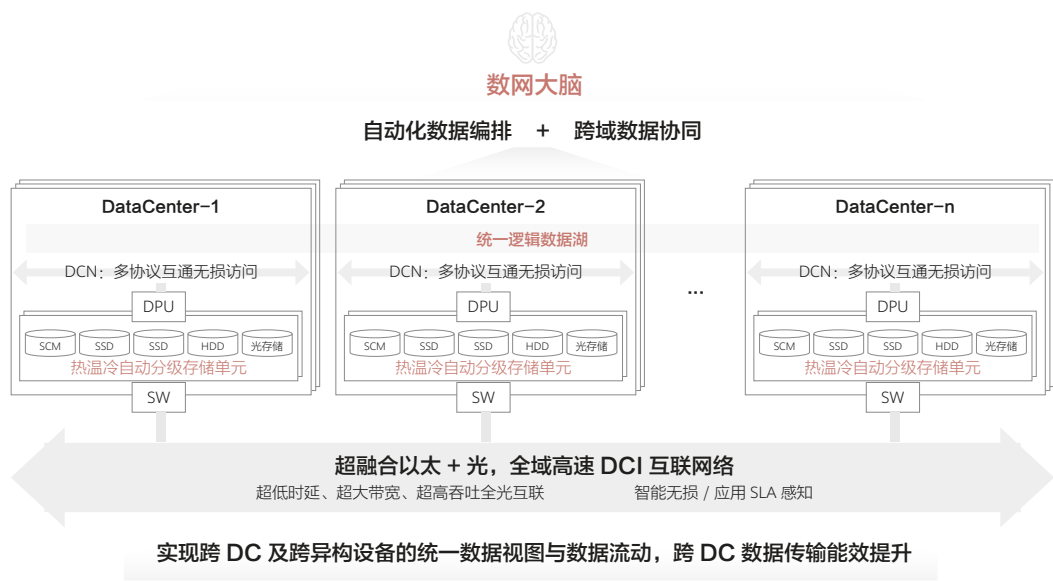


图 4-5：新型数据中心全局数据管理系统参考架构

■ 新协同服务，开放架构融入社会化算力

未来，随着应用场景的多元化，对数据中心功能定位也提出了新要求：数据中心正在由通用数据中心占主导，演变为通算中心、智算中心、超算中心，甚至由光计算、量子计算等共同构成的多类型数据中心共同发展的新局面，数据中心间协同，以及云边协同的体系将不断完善。以应用为驱动，多种类型的数据中心协同一体，共同提供算力服务的模式，将成为数据中心算力供给重要形态，持续支撑数字经济的发展。

新型数据中心不再是一个单独孤立的数据汇聚处理中心，而是泛在普惠算力服务基础设施的一部分，是整个社会化算力网络的有机构成单元，需要和外部进行更广泛的协同，广泛参与到社会生产生活的各个领域并实现全面赋能。

新型数据中心需要一个更加开放的协同架构，让所有数据中心在遵循统一标准的基础上，开放算力协同、数据协同、作业协同等接口，并能够与

外部的算力共享交易平台、数据共享交易平台、作业需求分配平台等可信共享交易平台进行快速对接，无缝参与到社会化普惠算力网络的分工协

作当中，共同实现“人人”为“人人”的开放共享经济模式，支撑千行万业智能化急剧增长的算力需求。



图 4-6 基于统一标准的开放共享新型数据中心协同架构

■ 新智能管理，AI 驱动实现 DC 自动运维

数据中心安全、高效、稳定的运行，是每一个运营者的核心目标。随着行业智能化程度越来越高、数据中心规模不断增大（到 2030 年，预计超大规模数据中心的 IT 硬件设备数量将超百万，实时运行的应用实例数量也将达数亿规模）、组件监控粒度越来越细、监控数据量越来越大以及新技术、新组件不断引入。传统数据中心运维管理面临着越来越严峻的挑战：1）竖井式管理，运维工具多、协同差，问题定界定位难。当前数据

中心管理往往太分散，单一视角解决问题，基础设施机房、计算、存储、网络等不同设备均有不同的监控告警系统、日志记录系统等，且缺乏系统间联动，日志、数据、告警等整合分析能力弱，造成故障难以定位、修复。2）运维数据分散量大，价值难发挥。当前数据中心一方面各类运维数据分散，缺乏集中沉淀，并且缺乏统一的数据指标体系，不利于数据整合分析；另一方面，数据获取难，多依赖手工收集，质量参差不齐，并

且分析手段、方法单一，数据价值未深入挖掘。
3) 自动化智能化水平低，根据中国信通院在《数据中心智能化运维发展研究报告（2023 年）》中的调研发现，当前中国大部分数据中心还处于人工运维为主的水平。

域数据汇聚、全局统一可视，并借助专家知识库及 AI 运维大模型训练，实现更早的发现风险、更快的解决问题、更精准的运营决策、更智能的运维管控，最终实现更深层次的高度“自动驾驶”数据中心。

面向 2030，数据中心管理需要由“人力密集型”向“技术密集型”演进，必须借助大数据、人工智能、知识图谱、数字孪生等关键技术来打造全局一体化智能运维体系。实现全栈数据采集、全



图 4-7：新型数据中心全局一体化智能运维体系参考架构







发展与 倡议

05



我们所处的既是充满挑战的时代，更是充满希望的时代。

以 AI、5G、云等为代表的数字化技术正深刻影响着我们的生活，并向千行万业加速渗透。人们一方面在畅想科技对生活的改变，同时也在忧虑新技术对环境和生态的影响。

数据中心是数字经济的“发动机”，只有不断地提高数据中心的效率，才可以为数字经济的发展提供澎湃的动力。我们认为只有从能效、算效、数效、运效、人效五个方面综合考虑，才能更全面的评价数据中心整体的先进性。通过“五效合一”，才能最终解决高速增长的算力需求与数据中心绿色低碳可持续发展之间的结构性矛盾，为智能社会创造更大的价值。

未来十年，构建“五效合一”，并且具备多样泛在、安全智慧、零碳节能、柔性资源、系统摩尔、对等互联六大技术特征的新型数据中心将是产业发展的愿景。

“长风破浪会有时，直挂云帆济沧海”，让我们携美好愿景一起前行，通过架构创新、系统创新、理论创新、工程创新，以及产、学、研、用全产业的共同努力，加速迈向波澜壮阔的智能世界。

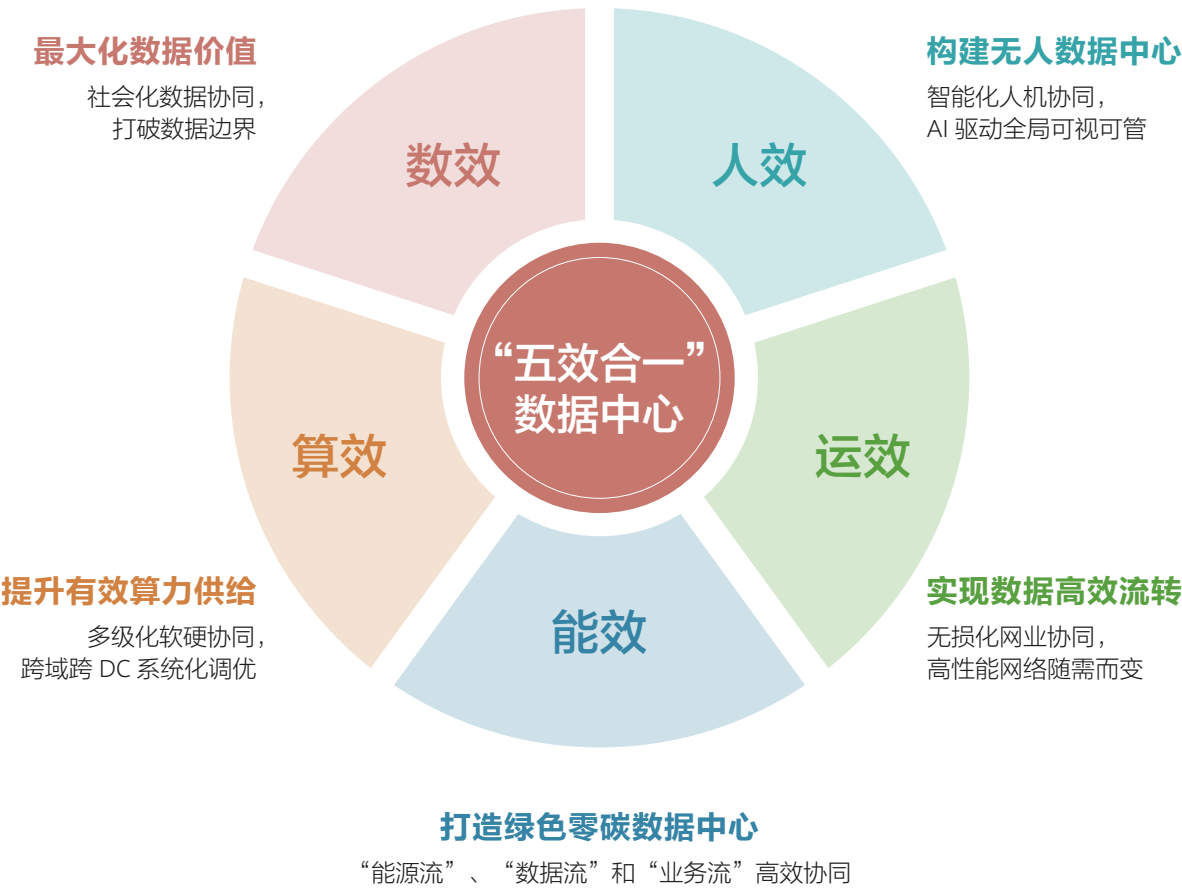


图 5-1 五效合一数据中心

附：关键预测数据指标体系

技术特征	指标	指标定义	2030年预测值
多样泛在	集群通用算力规模	单个集群通过软硬件联合调优，所获得的有效算力规模	>70EFplops
	集群 AI 算力规模	单个 AI 集群能够提供的有效算力规模	750EFlops
	集群存力规模	单个集群所能提供的有效存力规模	EB 级
	云边协同处理数据占比	需要边缘和中心分级处理的数据占总数据量的比例	80%
	企业生产设备数字化接入率	企业生产设备通过物联化和智能化后，通过边缘数据中心接入的占企业总生产设备的比例	80%
安全智慧	数据安全投资占比	数据安全投资占整个数据中心投资的比例	20%
	系统级可用性	系统级可用性 = 系统年平均故障间隔时间 MTBF / （系统年平均故障间隔时间 MTBF + 平均故障修复时间 MTTR）	99.999%
	重要数据灾备覆盖率	重要数据及关联应用系统具备容灾备份能力的比例	100%
	自动化运行能力等级	L1：运行辅助、L2：部分运行自动化，L3：有条件运行自动化，L4：高度运行自动化，L5：完全运行自动化，自动化包括自动预测性排障和分析、自动应急处置级 AI 能效管理等，L4 基本实现无人化，L5 完全实现无人化	L4 级
零碳节能	电能利用效率（PUE）	数据中心总消耗电量 /IT 设备消耗电量	1.0x
	可再生能源占比(REF)	可再生能源电量 / 数据中心总消耗电量	80%
	水使用效率（WUE）	水资源消耗量 / IT 设备消耗电量	0.5 L/kWh
柔性资源	DC 级资源池化率	在单个 DC 内，可进行全局调度的计算、存储和网络资源占全部资源的比例	80%
	企业新增应用云原生占比	新增云原生应用在全部新增应用中的比例	90%
	资源分配粒度	计算、存储、网络资源分配、调度和计费粒度	函数级
对等互联	超融合互联总线技术渗透率	统一的超融合互联总线技术的渗透率	60%
	超融合以太网网络渗透率	通用计算、高性能计算和存储网络的三张网融合部署占全部数据中心组网的比例	80%
	光算协同渗透率	采用全光直连 AI 参数面 Spine 层实现算力协同的集群算力力量，占集群总算力力量的比例	50%

技术特征	指标	指标定义	2030年预测值
对等互联	光存协同渗透率	采用全光直连 SSD 实现跨广域高速传输数据，占总传输数据的比例	50%
	全闪存存储占比	全闪存存储占数据中心总存储容量的比例	80%
系统摩尔	RDMA 存储网络渗透率	基于 RDMA 技术的存储网络使用比例	80%
	近存 / 存内计算数据处理量占比	采用近存 / 存内计算技术的数据处理量占全部数据处理量的比例	30%

附： 缩略语

缩略语	英文全称	中文全称
3GPP	3rd Generation Partnership Project	第三代合作伙伴计划
5G	5th Generation of Mobile Communication	第五代移动通信
ABAC	Attribute-Based Access Control	基于属性的权限控制
AI	Artificial Intelligence	人工智能
AIGC	AI-Generated Content	AI 生成内容
API	Application Programming Interface	应用编程接口
AR	Augmented Reality	增强现实
ARM	Advanced RISC Machine	高级精简指令集计算机
ASIC	Application-Specific Integrated Circuit	专用集成电路
BMS	Battery Management System	电池管理系统
CDN	Content Delivery Network	内容分发网络
CMDB	Configuration Management Database	配置管理数据库
CMOS	Complementary Metal-Oxide-Semiconductor	互补型金属氧化物半导体
CPO	Co-Packaged Optics	光电合封
CPU	Central Processing Unit	中央处理单元
CRIU	Checkpoint/Restore In Userspace	用户态实现 Checkpoint/Restore 功能的软件工具
CUDA	Compute Unified Device Architecture	通用并行计算架构
CXL	Compute Express Link	处理器、内存扩展和加速器的高速缓存一致性互连协议
DAC	Digital-to-Analog Conversion	数模转换
DBR	Distributed Bragg Reflector	分布式布拉格反射器
DC	Data Center	数据中心
DCN	Data Center Network	数据中心网络
DCOI	Data Center Optimization Initiative	数据中心优化倡议

缩略语	英文全称	中文全称
DDR	Double Data Rate	双倍数据速率
DFB	Distributed Feedback Bragg grating	分布式布拉格光栅
DPDK	Data Plane Development Kit	数据平面开发套件
DPU	Data Processing Unit	数据处理单元
DRAM	Dynamic Random Access Memory	动态随机存取存储器
E2E	End-to-End	端到端
EA	Electronic Absorption	电子吸收
EB	Exabyte	艾字节
EC	Erasure Code	纠删码
EFLOPS	ExaFLOPS	每秒一百京（=10 ¹⁸ ）次的浮点运算
ETH	Ethernet	以太网
ETSI	European Telecommunications Standards Institute	欧洲电信标准协会
FeRAM	Ferroelectric Random Access Memory	铁电式随机存取内存
FLOPS	Floating-point Operations per Second	每秒浮点运算次数
FPGA	Field Programmable Gate Array	现场可编程门阵列
FS	FusionSphere OpenStack	华为云操作系统
GeSI	Global e-Sustainability Initiative	全球电子可持续性倡议组织
GPT	Generative Pre-trained Transformer	生成式预训练 Transformer 模型
GPU	Graphical Processing Unit	图形处理单元
HAMR	Heat Assisted Magnetic Recording	热辅助磁记录
HDD	Hard Disk Drive	机械硬盘
HPC	High-Performance Computing	高性能计算
HTTP	Hypertext Transfer Protocol	超文本传输协议

附：缩略语

缩略语	英文全称	中文全称
HTTPS	Hypertext Transfer Protocol over Secure Sockets Layer	安全套接字层的超文本传输协议
I/O	Input/Output	计算机系统输入输出
IB	InfiniBand	无限带宽
ICT	Information and Communications Technology	信息与通信技术
IDC	Internet Data Center	互联网数据中心
IGZO	Indium Gallium Zinc Oxide	氧化铟镓锌
IO	Input/Output	输入输出
IoT	Internet of Things	物联网
ISP	Internet Service Provider	互联网服务提供方
IT	Information Technology	信息技术
K-V	Key-Value	键 - 值
KVM	Kernel-based Virtual Machine	内核虚拟机
KW	Kilowatt	千瓦
LR	LongRange	长距离光接口
MAMR	Microwave Assisted Magnetic Recording	微波辅助磁记录
MC	Main Control	主控
MEC	Multi-access Edge Computing	多址边缘计算
MPLS	Multi-Protocol Label Switching	多协议标签交换技术
MRAM	Magnetic Random Access Memory	磁阻式随机存取存储器
ms	Millisecond	毫秒
MW	Megawatt	兆瓦
MZ	Mach-Zehnder modulator	马赫 - 曾德调制器
NAND	Non-volatile Memory Device	非易失性存储设备

缩略语	英文全称	中文全称
NG DCOS	Next Generation Data Center Operating System	下一代数据中心操作系统
NISQ	Noisy Intermediate-Scale Quantum	嘈杂中型量子
NOF	NVMe over Fabrics	NVMe-oF 协议
NoSQL	Not only SQL	非关系型数据库
NPU	Neural Processing Unit	神经处理单元
NUMA	Non-Uniform Memory Access	非一致性内存访问
OBO	On Board Optics	板载光学系统
oDSP	optical Digital Signal Processor	光数字信号处理芯片
OS	Operating System	操作系统
OXC	Optical Cross-Connect	光交叉连接
PB	Petabyte	拍字节
PCI	Peripheral Component Interconnect	外围部件互连标准
PCIe	Peripheral Component Interconnect express	快捷外围部件互连标准
PCM	Phase Change Memory	相变存储器
PUE	Power Usage Effectiveness	能源利用效率
QLC	Quad-Level Cell	四层式存储单元
QoS	Quality of Service	服务质量
RBAC	Role-Based Access Control	基于角色的访问控制
RDMA	Remote Direct Memory Access	远程直接存储器访问
ReRAM	Resistive Random Access Memory	电阻型随机存储器
SATA	Serial Advanced Technology Attachment	串行高级技术附件
SCM	Storage Class Memory	存储级内存
SDN	Software-Defined Networking	软件定义网络

附：缩略语

缩略语	英文全称	中文全称
SDS	Software-Defined Storage	软件定义存储
SLA	Service Level Agreement	服务水平协议
SNIC	Standard Network Interface Card	标准网络接口卡
SQL	Structured Query Language	结构化查询语言
SR	ShortRange	短距离光接口
SSD	Solid-State Drive	固态硬盘
swTPM	Software Trusted Platform Module	软件 TPM
TB	Terabyte	太字节
TCO	Total Cost of Operation	总运营成本
TCP/IP	Transmission Control Protocol/Internet Protocol	传输控制协议 / 互联网协议
TEE	Trusted Execution Environment	可信执行环境
TOR	Top of Rack	机架交换机
TPM	Trusted Platform Module	可信平台模块
UB	Unified Bus	灵衢总线
UPI	UltraPath Interconnect	超级通道互连
UPS	Uninterruptible Power Supply	不间断电源
VCSEL	Vertical Cavity Surface Emitting Laser	垂直共振腔表面放射型激光部件
VM	Virtual Machine	虚拟机
VPC	Virtual Private Cloud	虚拟私有云
VR	Virtual Reality	虚拟现实
WebGL	Web Graphics Library	Web 图形库
WORM	Write Once Read Many	一次性写入多次读出
WUE	Water Usage Effectiveness	水使用效率

缩略语	英文全称	中文全称
xPU	A portfolio of architectures (CPU, GPU, FPGA and other accelerators)	CPU、GPU 等各种处理器的统称
XR	eXtended Reality	扩展现实
YB	Yottabyte	尧字节
ZB	Zettabyte	泽字节
ZFLOPS	ZettaFLOPS	每秒十万京（=10 ²¹ ）次的浮点运算

致谢

人类社会正加速进入智能时代，未来最大的需求是算力，最关键的基础设施是数据中心，数据中心也被称为“数字经济发动机”。

2023 年 9 月 20 日，华为在 HC 2023 期间面向全球成功发布《数据中心 2030》报告，获得客户伙伴及全球数据中心产业相关方的高度关注和热烈反响。

《数据中心 2030》报告从算力需求与资源约束的核心矛盾出发，描绘了影响数据中心发展的五大未来场景，围绕效率提升的五大创新方向，在业界首次定义未来数据中心六大技术特征，系统性阐述了数据中心所包含的云服务、计算、存储、网络、能源等全栈技术的发展挑战与突破方向。最后提出了一个面向未来的新型数据中心参考架构，给出了 22 个指标及预测数据，对数据中心产业未来发展前景进行了定量预测。

《数据中心 2030》是华为《智能世界 2030》思想领导力产品系列的延续，报告由数据中心产品组合 SDT、ICT 战略与业务发展部和 EBG CTO 办公室三方共同推动，聚合公司各领域专家及业界智库的智慧，共同编写而成。面向全球客户伙伴、产业组织、行业智库等，《数据中心 2030》报告有效传递了华为对未来数据中心的探索与思考，也必将给全球数据中心产业发展及建设带来积极影响。

探索未来数据中心，引领智能时代！

谨以此对项目组所有成员的贡献表示最诚挚的感谢！

一、报告统稿小组

项目组	主要贡献者
赞助人	马海旭、盖刚
统稿小组	刘树清、张蕾、杜伟、窦雪峰、蓝飞翔、鞠德刚
主编	杜伟、窦雪峰、蓝飞翔

项目组	主要贡献者	
编委	华为云	顾炯炯、黄瑾
	计算 / 2012 实验室	姜涛、惠涛、邹斌、秦佩峰、高俊恩、张瑞、邓春梅、LIAO HENG、程柏、刘华伟
	存储	庞鑫 、方卫峰、张国彬、何苗、袁燕龙、张大成
	光 / 2012 实验室	肖新华、王景燕、马会肖、宋小鹿、黄科超、李洋
	数通	钱骁、李军
	数字能源	张峰、张宗望
	云核心网	马亮、尹东明
	战略研究院	黄华、李岑
	数据中心产品组合 SDT	赵云凯、宋春尧、周胡根
其他 贡献人	党文栓、朱照生、张浩、林春光、朱国军、赵俊峰、伏兴平、姜险峰、王志新、夏卓新、唐晓军、赵祎、毛杉乡、陈斌、邓彬林、孟万红、李社明、高文、齐立炜、孙春、刘军、秦云鹭、林帅、熊浩宇、张敏威	

二、报告发布与传播小组

组别	主要贡献人员
总体组	刘树清、张蕾、张宛琪、任竟慧、高政军、魏彤彤
传播组	邱斯娴、陈璐
翻译组	徐寿娟、陈夏欢、王锦娴、刘鹏、刘丽敏、邱智胜、王爱香、潘媛、曾懋琳、雷亚琳、李明霞、刘璇、巫梦妮、严慧玲、张维瑜、钟美玲、高木子、白友员、李盼望、余珊珊、王攀鹏、冯文超、Ekaterina Christova、Zachary Overline、Megan Young、Scott Winnen、Omar Belove、George Fahy、Gavin Wills

Thank you

华为技术有限公司
深圳龙岗区坂田华为基地
电话: +86 755 28780808
邮编: 518129
www.huawei.com



扫码下载报告

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。