

COMMScope®

数据中心的发展趋势

2023 年发展趋势展望

目录

简介	3
关于作者	4
第 1 章： 适应数据中心中越来越多的光纤数量	6
第 2 章： OM5 多模光缆应用的成本/效益分析	14
第 3 章： 数据中心中的 400G 以太网：光收发器选择	19
第 4 章： 数据中心中的 400G 以太网：高密度化和园区架构	22
第 5 章： 请勿只着眼于现在 - 800G 即将来临！	26
第 6 章： 网络边缘的 MTDC	30
第 7 章： 数据中心在 5G 世界中的角色演变	33
第 8 章： 遍布园区，进入云端：MTDC 连接的推动因素有哪些？	38
第 9 章： 通往 1.6T 的道路正在开启	44
结论	49



来年展望：是什么正在影响着数据中心？

在数据center里，没有什么是一成不变的，展望 2023 年，也不例外。随着数据center的数据流量不断攀升，受更大的连接需求驱动，网络规划者正在重新考虑如何才能在这些挑战中保持领先。

回顾 2014 年，25G 以太网联合会提出单通道 25 Gbps 以太网和双通道 50 Gbps 以太网时，在行业路线图中创建了一个大的分支，降低了每比特成本，并可轻松地过渡到 50 G、100 G 及以上。

2020 年，100G 以太网应用大量涌入市场，导致光纤使用量不断增长，超大规模和云计算数据center将不可避免地走向 400G 以太网跃进。随着交换机和服务器逐步采用 400G 和 800G 连接，物理层也必须具有更高性能，以持续优化网络容量。

发展数据center物理层基础设施是满足用户对低延迟、高带宽和可靠连接需求的关键。我们通过观察数据center管理者对于 800G 以太网的规划，以及 5G 技术将带来的蘑菇效应（数据流量爆发式增长），来看看这些重要的趋势。

关于作者



Matt Baldassano

Matt Baldassano 是康普美国东北地区的企业解决方案部门技术总监，专门负责数据中心布线产品。他曾担任过康普数据中心业务部门的业务开发经理和技术营销工程师。

此外，他还曾在纽约市和德克萨斯州达拉斯的 EMC2 公司担任过客户工程师一职，负责提供数据中心有线网络和室内无线系统方面的服务，并写过无线安全主题方面的文章。Matt 持有圣约翰大学的计算机科学学士学位和先进技术大学的技术硕士学位。



Jason Bautista

作为超大规模和多租户数据中心 (MTDC) 的解决方案架构师，Jason 负责康普的数据中心市场开发。他负责关注数据中心市场的趋势，以帮助推动针对超大规模和多租户数据中心客户的产品战略、解决方案和产品计划。

Jason 在网络行业拥有 19 年的丰富经验，曾经担任多个面向客户的职位，致力于各种不同网络的产品开发、营销和支持，服务于全球各地的客户。

关于作者



Ken Hall

Ken Hall 是康普北美洲数据中心架构师，负责技术和思想领导，以及全球范围和相关数据中心的光纤基础设施规划。他参与开发和发布了高速、超低损耗光纤解决方案，以便高效地实现数据中心运营商的网络迁移。

此前，Ken 曾在 TE Connectivity/Tyco Electronics/AMP 担任过多种职务。他在全球网络 OEM 和数据中心的计划管理和战略，以及项目管理、市场营销、行业标准和销售管理方面拥有丰富的经验。Ken 还曾负责面向网络电子设备 OEM 的小型铜缆和光纤接头以及高密度接口的行业标准化和普及。

迄今为止，Ken 在光纤接头和基础设施管理系统方面拥有 9 项专利。

Ken 获得了宾州西盆斯贝格大学的理学学士学位。Ken 既是一位注册通讯设计师 (RCDD)，也是一位网络技术系统设计师 (NTS)。



Hans-Jürgen Niethammer

Hans-Jürgen 于 1994 年 7 月加入康普布线部门，曾出任产品管理、技术服务和市场营销部门的多种关键职务，包括欧洲、中东和非洲项目管理总监，欧洲、中东和非洲市场营销总监，以及欧洲、中东和非洲技术服务和销售运营总监。

自 2013 年 1 月起，Hans-Jürgen 就一直负责康普在欧洲、中东和非洲的数据中心市场开发，确保康普解决方案能够实现客户数据中心基础设施的敏捷性、灵活性和可扩展性，以满足这一动态细分市场的当前和未来需求。

Hans-Jürgen 是数据中心、光纤和 AIM 系统领域的国际专家，多个 ISO/IEC 和 CENELEC 标准化委员会的成员，以及多个国际标准的编辑。

Hans-Jürgen 持有电子工程学的认证工程师学位，并且获得了商业经济学家认证。

关于作者



Alastair Waite

Alastair Waite 于 2003 年 9 月加入康普，担任公司企业光纤部门产品经理一职，从那以后，他在公司担任了许多关键职位，包括欧洲、中东和非洲的企业产品管理主管，欧洲、中东和非洲市场管理主管和数据中心业务部门领导者。

自 2016 年 1 月，Alastair 就一直负责康普数据中心解决方案的架构，在这一动态细分市场确保客户的基础设施能够随运营需求而扩展。

在加入康普之前，Alastair 曾是 Conexant Semiconductor 公司光学芯片部门高级产品线经理，全面负责该公司的所有光学接口产品。

Alastair 威尔士大学电子工程理学学士学位。



James Young

James 是康普企业数据中心部门总监，负责战略监督并领导全球产品和现场团队。

此前，James 曾在加拿大的 Tyco Electronics/AMP、Anixter、Canadian Pacific 和 TTS 担任过多种职务，包括通信解决方案的销售、营销和运营。James 在 OEM 产品、网络解决方案和增值服务销售方面，拥有丰富的直接和间接渠道销售经验。

James 获得了西安大略大学的理学学士学位。他既是一位注册通讯设计师 (RCDD)，也是一位获认证的数据中心设计专家 (CDCP)。



1/

适应数据 centers 中越来越多的光纤数量

涌入数据中心的数据流量持续攀升；与此同时，5G、人工智能和机器与机器 (M2M) 通信等先进技术推动的新一代应用，正将延迟要求提升至一毫秒范围内。这些趋势和其他趋势正在一起影响数据中心的基础设施，迫使网络管理者重新思考如何在变化中保持领先。

传统上，网络使用四个主要方法来满足不断增长的更低延迟和更大流量需求：

- 减少链路中的信号损耗
- 缩短链路距离
- 加快信号速率
- 加大管道尺寸

虽然数据中心在某种程度上使用了所有这四种方法（尤其是超大规模数据中心），但目前主要集中在增加光纤的数量。过去，核心网络布线主要采用 24、72、144 或 288 芯光缆。在这一层级上，可以对数据中心内交换机或着服务器之间的主干网络进行分段式管理，然后使用光纤耦合器与跳线对光缆进行分配和跳接，以提高安装效率。如今，光缆的芯数增加了多达 20 倍——每根光缆达到 1,728、3,456 或 6,912 芯。

更大的光纤芯数结合更紧凑的光缆结构，对于数据中心互联更有帮助。连接两个超大规模设施时，通常会使用超过 3,000 芯光缆作为数据中心互连 (DCI) 的主干布线，在不久的将来，运营商的设计需求还将在此基础上翻倍。在数据中心内部，相同的问题会出现在从核心交换机或会接间到机柜列内 Spine 交换机之间的主干布线线缆上。

无论数据中心配置是需要点对点还是交换机至交换机连接，不断增加的光纤芯数，都给数据中心在提供更高带宽和容量方面带来了巨大的挑战。

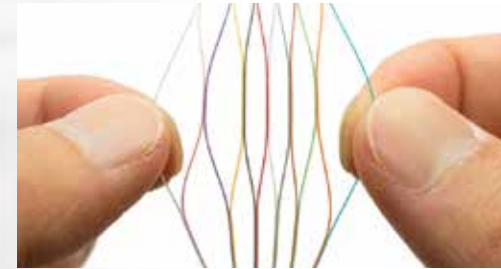
第一个挑战是：如何以最快速、最有效的方式部署光纤？如何将光纤放到线轴上？如何将光纤从线轴上拉出？如何在两点之间进行布线和穿过线槽？

一旦完成安装，第二个挑战是：如何在交换机和服务器机架上对光纤进行配线和管理？



可卷曲带状光纤布线

随着市场对更快更大的数据管道需求的增加，光纤和光网络也在不断发展。随着这类需求的加剧，光缆中光纤的设计和封装方式也发生了变化，这使得数据中心可以在无需增加布线空间的情况下增加光缆中的光纤数量。可卷曲带状光纤布线就是这个创新链中的一个重要环节。



可卷曲带状光缆以不连续的方式粘连在一起。

资料来源：ISE 杂志

从某种程度上说，可卷曲带状光缆是以早期开发的中心束管带状光缆为基础。自 20 世纪 90 年代中期推出以来，中心束管带状光缆主要用于 OSP 网络，其特点是在一个中心缓冲管中堆叠了多达 864 芯带状光纤。这些光纤被分组，并沿着光缆的延伸方向完全粘合在一起，这种方式增加了光缆的刚性。虽然这对于在数据中心 OSP 应用中部署这类光缆影响不大，但是要在有限的空间中进行敷设，最好不要使用又粗又硬的光缆。

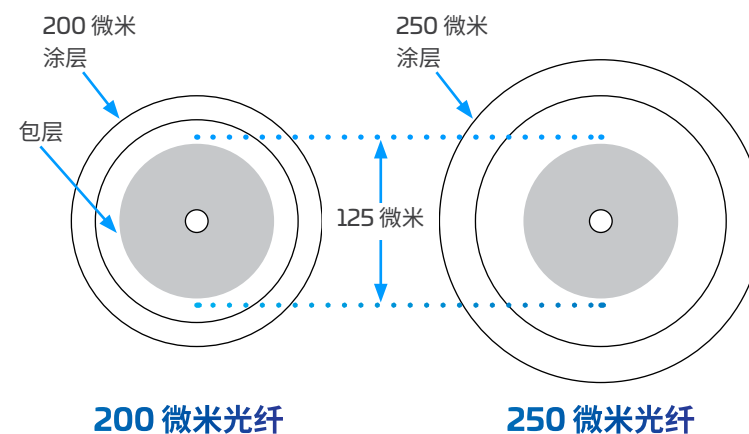
在可卷曲带状光缆中，光纤以不连续的方式粘结在一起，形成了一个松散的网状结构。这种结构使得带状光缆更加柔软，这让制造商能够在一个两英寸的套管内放置多达 3,456 芯光纤（传统封装光纤密度的两倍）。这种结构减小了弯曲半径，因而这些光缆更易于用在空间有限的数据中心。

以不连续的方式粘连在一起的光纤具有松散光纤的物理特性，易于收缩和弯曲，因此更易于敷设在狭小空间内。此外，可卷曲带状光纤布线方式与干式光缆完全相同，这有助于减少准备熔接所需的时间，从而降低人工成本。不连续的粘连技术仍保持了典型带状熔接所需的光纤对准方式。

减少线缆直径

几十年来，几乎所有通信光纤的标称涂层直径都是 250 微米。随着市场上对更细的线缆的需求越来越多，这种情况已经开始有所改变。许多光缆设计都已经达到了标准光纤直径缩减的实际极限。但需要更细的光纤促进光纤物理直径进一步缩减。200 微米涂层的光纤现已用于可卷曲带状光缆和微管道光缆。

必须强调的是，缓冲涂层是光纤中唯一被改变的部分。200 微米光纤保留了传统光纤 125 微米的纤芯/包层直径，以确保熔接操作中的兼容性。200 微米光纤与 250 微米光纤在剥除缓冲涂层后，熔接流程完全相同。



在光学性能和熔接兼容性方面，200 微米光纤具有与 250 微米光纤相同的 125 微米纤芯/包层。资料来源：ISE 杂志

全新芯片导致问题进一步复杂化

一列机柜中的所有服务器都被配置为特定的连接速度。但在当今高度融合的矩阵网络中，一列机柜内的所有服务器都同时以最大速率运行的情况极其少见。所需的服务器上行带宽与部署的下行带宽之差称为收敛比。在网络的某些区域中，如交换机间链路 (ISL)，收敛比可高达 7:1 或 10:1。选择更高的收敛比可以尽可能地降低交换机成本，但大多数现代云和超大规模数据中心网络设计的目标收敛比为 3:1 或更低，以实现更好的网络性能。

收敛比在构建大型服务器网络时更为重要。随着交换机至交换机的带宽容量增加，交换机之间的链接数量减少。为了满足所需的服务器上联需求，需要将多层叶脊网络结合起来，同时通过交换机与交换之间的连接来实现整个网络的收敛比目标。不过，每个交换机层都会增加成本、功耗和延迟。交换技术一直致力于解决这一问题，这推动了商用 ASIC 硅交换芯片的快速发展。2019 年 12 月 9 日，Broadcom Inc. 开始交付新的 StrataXGS Tomahawk 4 (TH4) 交换机，从而在单个 ASIC 中实现了 25.6 Tbps 的以太网交换能力。这距离 Tomahawk 3 问世还不到两年，当时 Tomahawk 3 (TH3) 的最高速率为 12.8 Tbps。



这些 ASIC 芯片 不仅提高了通道速度，还增加了其所含端口的数量。数据中心可以更好的控制收敛比。使用单个 TH3 ASIC 芯片构建的交换机支持 32 个 400G 端口。每个端口可以分解成八个 50GE 端口，以实现服务器连接。端口可以组合在一起，支持 100G、200G 或 400G 连接。在采用与 QSFP 相同的规格时，每个交换机端口可在 1 对、2 对、4 对或 8 对光纤之间迁移。

虽然这看似复杂，但却在帮助控制收敛比方面非常有用。如今，这些新型交换机可连接最多 192 台服务器，同时仍保持 3:1 收敛比，同时提供 8 个 400G 端口用来进行叶脊交换机之间的互联。现在，一台这种新型交换机可取代 6 台上一代交换机。

新型 TH4 交换机将配有 32 个 800Gb 端口。ASIC 芯片的通道速率提升至 100G。目前正在制定新的电气和光学规范以支持 100G 通道。全新的 100G 生态系统将进一步优化网络基础设施，更适用于新的应用，如机器学习 (ML) 或人工智能 (AI)。

布线系统供应商的角色演变

在这个更为复杂的动态环境中，布线系统供应商的作用变得重要起来。虽然光纤布线曾一度被认为仅仅是一种商品，而不是

工程解决方案，但现在不再是这样了。由于有如此多的信息需要了解，同时又面临如此多的风险，供应商的角色已转变为技术合作伙伴，在确保数据中心成功运转方面，他们与系统集成商或设计人员同等重要。

数据中心所有者和运营商越来越依赖其布线合作伙伴在光纤成端、收发器性能、熔接和测试设备等方面的专业知识。这一新增角色的要求布线合作伙伴与基础设施生态系统相关人员以及标准机构建立更密切的工作关系。

随着行业标准和多源协议 (MSA) 数量的增加以及通道速度的加速，布线合作伙伴在实现数据中心技术路线图方面发挥着越来越重要的作用。目前，有关 100G/400G 和正在演进的 800G 的标准包含一系列令人眼花缭乱的选择。每个选项都包含多种传输数据的方法，包括双工、并行和波分复用——每一种方法都考虑了一个特定的应用场景。布线基础设施的设计应该在其使用寿命范围内支持尽可能多的传输应用。



一切都要归于平衡

光纤数量的增长使得数据中心可利用空间持续缩减，导致空间不足。我们也期待着采用更小规格的服务器与机柜，以便提供更多空间。

空间并不是唯一需要最大化的变量。通过将新的光纤结构（如可卷曲带状光缆）与更小的光缆尺寸和先进的编码技术相结合，网络管理者及其布线合作伙伴就会有許多方法可供使用。他们将用到所有这些方法。

如果技术发展速度可以预示未来，那么数据中心，尤其是超大规模云数据中心，应做好充分准备。随着带宽需求和服务等级

不断提高，以及延迟对最终用户/机器来说变得越来越重要，光纤将会在网络中得到进一步普及。

随着用户、设备 and 应用数量越来越多，为了提供超可靠的连接，超大规模和云端设施面临着越来越大的压力。而要满足这些需求，就必须具备部署和管理更高光纤容量的能力。

我们的目标是通过在适当设备中提供适当数量的光纤来实现平衡，同时实现良好的维护和可管理性并支持未来增长。所以，请明确您的发展方向，并在您的团队中选择一个像康普这样的可靠合作伙伴。





2 /

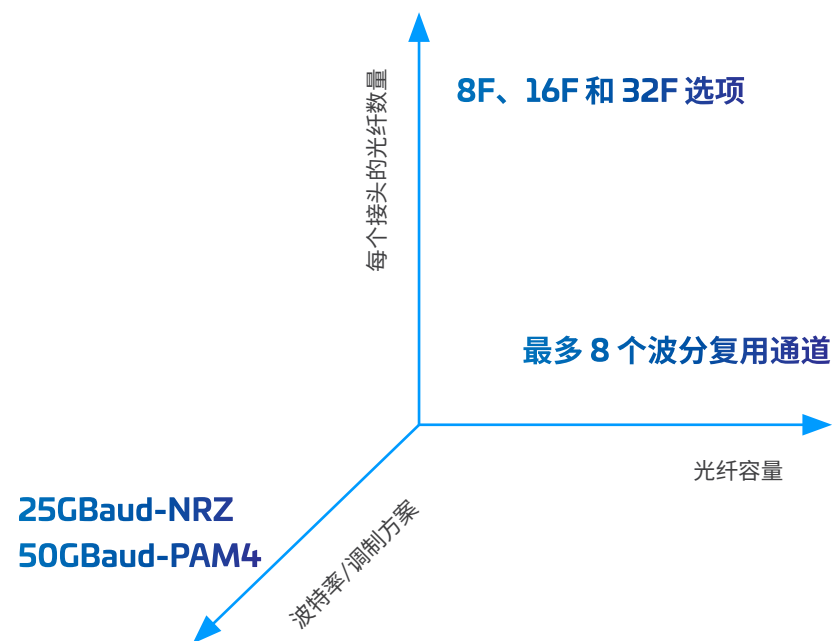
OM5 多模光缆应用的成本/效益分析

为了满足日益增长的网速的需求，IEEE（以太网标准化委员会）开始使用 3 项关键技术来提高以太网带宽：

- 通过增加用于传输的光纤数量来增加数据流（通道）的数量。虽然传统上，每个数据通道使用 2 芯光纤，但我们发现当今有些以太网应用使用 8 芯、16 芯甚至 32 芯光纤。从布线角度来看，我们使用多芯光纤接头（MPO）来满足每个应用中增加的光纤数量。
- 增加波特率调制。更具体地说，这包括从 25 Baud-NRZ 方案过渡到 50 Baud-PAM4 方案。当然，由于 PAM4 方案的传输速度增加了一倍，所以需要在信号质量和收发器成本方面进行权衡。
- 升级单条光纤容量。通过在每个光纤纤芯上使用不同波长的光信号，WDM 波分复用技术通过使用不同波长的光信号可以在光纤上传输多个数据流，每芯光纤最多可以支持 8 个 WDM 通道。

许多应用采用上述的其中一种技术来提高速度，而某些应用则采用多种技术。比如说，400GBase-SR4.2 集成了并行光纤（8 条）和使用短波分复用（SWDM；主要用作为双向技术）的优势。

数据中心网络管理者面临的挑战是，在不知道未来会有什么变化、400G/800G 及更高速率发展路径如何相交或在何处相交的情况下进行规划，例如 400GBase-SR4.2。这就是 OM5 多模光纤的价值所在，OM5 是一种经过标准化设计，以便在单个光纤纤芯中支持多种波长的新型多模光纤。



提高以太网速度的三条路径

OM5 多模光纤

OM5 于 2016 年推出，属于 WBMMF（宽带多模光纤）。

OM5 的特性经过优化，在同一芯光纤中传输多个波长光信号的技术 (Bi-Di) 来支持高速数据中心应用。OM5 技术的技术细节和操作优势是众所周知的：

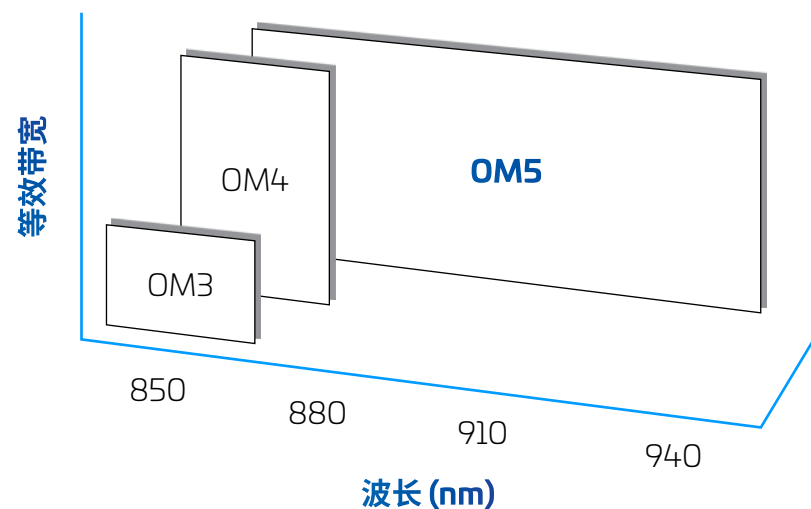
- 减少了并行光纤数
- 减少了已光纤成缆的衰减
- 实现更宽的有效模式带宽 (EMB)
- 应用支持距离比 OM4 长 50%

因为 OM5 采用与 OM3 和 OM4 相同的物理外形（50 μm 纤芯，125 μm 包层），所以它完全可以向后兼容这些光纤类型。

OM5 与 OM4：成本/效益分析

与 OM4 相比，OM5 具有一些明显的技术和性能优势。

尽管 OM5 有许多优势，但其采用和普及却受到一些数据中心运营商的阻力（就像推出 OM4 时，数据中心使用 OM4 替换 OM3 的速度也很慢）。不愿意转而使用 OM5 的一个潜在原因就是其价格更高。然而，仔细分析 OM5 与 OM4 的成本/效益就会发现并非如此。



有效模式带宽差异

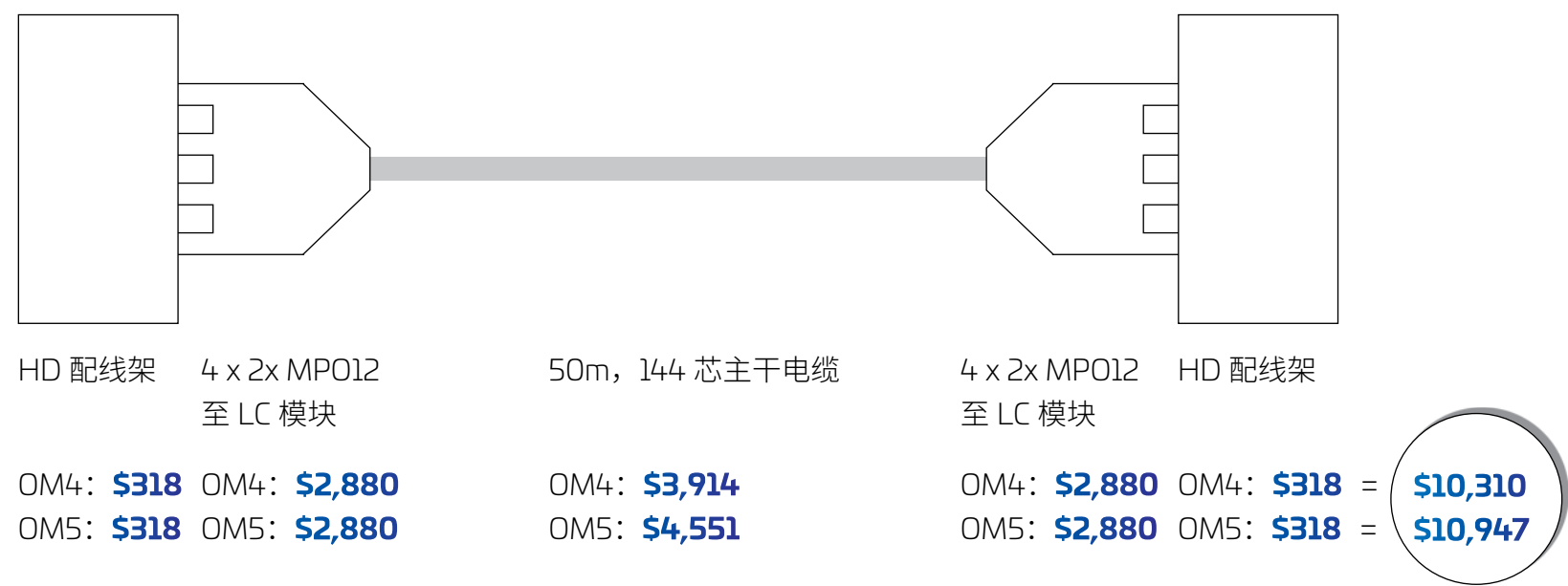
成本

OM5 光纤的反对者通常会指出 OM5 的购买价格比 OM4 高 50-60%。但如果只关注光纤价格，就会忽略数据中心管理者的总体运营需求。首先，在将 OM5 光纤制成主干光缆中后，OM5 光缆与 OM4 光缆的价格差异就会缩小至大约 16%。其次，当您将主干布线两端的配线架和光纤盒的成本加起来时，最初的 50-60% 价格差异就会显著减小。事实上，当您比较链路配置相同的 OM4 和 OM5 的总成本时，您会发现 OM5 只比 OM4 贵大约 6.2%。

示例场景

试想一个实际用例：我们将一根 50 米长的 144 纤芯主干电缆的两端连接至安装在 1U 高密度配线架中的 4 个 2xMP012-LC 模块上。每套组件的大致成本已给定。请注意，OM4 和 OM5 的配线架和光纤盒的总成本相同，区别只在于两者的光纤主干电缆成本相差大约 16%。

当您计算每个端对端链路的总成本时（OM4 为 10,310 美元，而 OM5 为 10,947 美元），您会发现成本差为 637 美元，意味着成本提高 6.2%。



此外，请记住综合布线只占数据中心总资本支出（包括结构、支持的基础设施（如电力设备、散热设备、UPS）以及所有 IT 设备（如交换机、存储器和服务器）的 4% 左右。因此，改用 OM5 将使数据中心的总资本支出提高 0.24%（不到 1% 的四分之一）。以绝对美元计算，这意味着每 1,000,000 美元的数据中心资本支出就会增加 2,400 美元。

优势

对于数据中心管理者来说，问题在于 OM5 的成本增量是否超过其效益。这里只介绍其中一部分直接和间接的优势。

OM5 可提供更高的单芯容量，因此可以在双向应用中减少光纤数，延长距离。在 100G 和 400G 应用中，采用 Bi-Di 技术时，使用 OM5 的支持距离比使用 OM4 远 50%，且使用的光纤数比 OM4 少 50%。OM5 只需两条光纤就能够支持 100G（甚至更多，参考 800G 和 1.6T Bi-Di 应用）。此外，OM5 能够支持 150 米，而 OM4 只能支持 100 米，所以随着布线架构的发展，OM5 可实现更大的设计灵活性。

使用 OM5 光纤可减少纤芯数量，从而能够更好的利用现有的光纤线槽，更少的光缆意味着更少的空间占用。

OM5：避免不确定性的一种手段？

也许最重要的是，OM5 可以给您提供更多的自由度利用未来的技术。无论您升级到 400G/800G 甚至更高数据速率的路径涉及增加每个接头的纤芯数、每条光纤的波长种类，还是需要采用更高阶调制方案，OM5 都可根据您的需要提供应用支持，扩展带宽以及延长传输距离。

当涉及应对快速发展环境下更高速网络升级的持续挑战时，对各种选择持开放态度就显得非常重要。您或许不需要 OM5 提供的所有优势，或者不了解 OM5 可能会在未来发挥关键作用。但您无法提前知晓，而这就是问题的关键。OM5 使您能够以尽可能低的风险应对每次挑战。以上就是康普的观点；我们也不想听听您的看法。



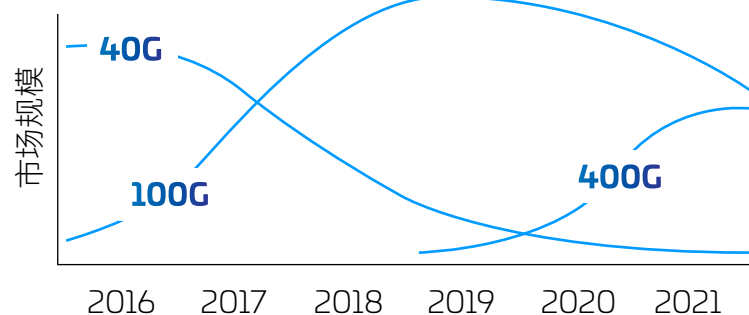
3 /

数据中心中的 400G 以太网： 光收发器选择

衡量组织是否成功的首要标准就是看其是否能够适应其所在环境的变化。我们称之为生存能力。如果不能适应新的变化，就会丢失客户。

对于云计算数据中心来说，由于对带宽、容量和低延迟需求的不断上升推动了网络速度的提高，其适应能力和生存能力每年都在经受考验。在过去几年中，整个数据中心行业的网络速度已经从 25G/100G 提升至 100G/400G。每次飞跃至更高速度后，就会进入一个短暂的平稳期，之后数据中心管理者需要为下一次升级做准备。

目前，数据中心正在寻求升级到 400G。主要考量因素就是哪种光纤技术最优。在此，我们来分析一些考量因素和选项。



400G 端口数量包括 8x50G 和 4x100G 两种实现方式。

资料来源：NextPlatform 2018

400G 光收发器

400G 的光学市场正在受到成本和性能的驱动，设备制造商也正在试图占据数据中心市场的最佳位置。

2017 年，CFP8 成为第一代 400G 模块收发器类型，用于核心路由器和 DWDM 传输客户端接口。该模块的尺寸略小于 CFP2，但其光学器件支持 CDAUI-16 (16x25G NRZ) 或 CDAUI-8 (8x50G PAM4) 电子 I/O 接口。最近，关注焦点已从模块技术转向尺寸缩小的第二代 400G 收发器类型模块：QSFP-DD 和 OSFP。

这些拇指大小的模块是专为结合使用高密度数据中心交换机而开发，可通过 32 x 400G 端口在 1RU 中实现 12.8 Tbps 的速度。请注意，这些模块只支持 CDAUI-8 (8x50G PAM4) 电子 I/O 接口。

尽管 CFP8、QSFP-DD 和 OSFP 均为热插拔模块，但并非所有 400G 收发器模块都是如此。某些模块被直接安装在了主机印刷电路板上。由于这些嵌入式收发器采用非常短的 PCB 印制线，所以可实现低功耗和高端口密度。

尽管嵌入式光模块具有更高的带宽密度和更快的通道速率，但以太网行业仍更青睐于 400G 可插拔光模块；因为可插拔光模块更易于维护，且可实现按需付费的成本效益。

从目标开始

对于业内资深人士来说，升级到 400G 只是数据中心发展之路上的又一个中转站。如今，800G MSA 行业联盟和标准委员会已开始致力于使用 8 x 100G 收发器实现 800G 的速度。作为 800G MSA 行业联盟的一员，康普也在与其他 IEEE 成员合作，寻找使用多模光纤支持每波长 100G 服务器连接速度的解决方案。这些发展的目标是在 2021 年进入市场，随后可能在 2024 年实现 1.6T 方案。

尽管过渡到越来越高速度所涉及的细节充满挑战，但它有助于我们正确认识这个过程。随着数据中心服务能力的发展，存储速度和服务器速度也必须提高。为支持越来越高的速度，需要使用合适的传输介质。

在选择最能满足网络需求的光纤模块时，首先要考虑项目的建设目标。您对所需服务以及交付这些服务所需连接拓扑结构的预测越准确，网络就越能更好地支持新应用和未来应用。

4 /

数据中心中的 400G 以太网： 高密度化和园区架构

400G 创造了对布线设备的新需求

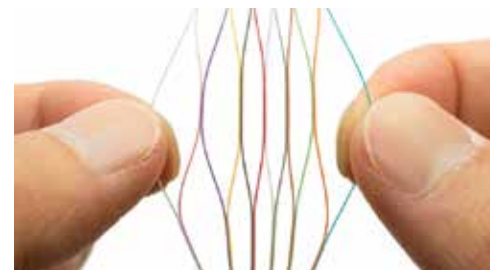
更高的带宽和容量需求推动着光纤数量的增加。十五年前，数据中心的大部分光纤主干均使用不超过 96 芯的光缆，包括多样化和冗余路由的需求。

目前，144、288 和 864 光纤芯数已开始成为常态，而主干光缆以及超大型数据中心和云计算数据中心使用的光缆正过渡到使用 3,456 芯光缆。如今，一些光缆制造商可提供 6,912 芯光缆，并考虑在未来提供更高纤芯数的光缆。

新的光纤封装和设计增加了密度

大芯数光缆的布线占用了管道宝贵的空间，且其较大的光缆外径提高了弯曲半径方面的施工难度。为解决这些问题，光缆制造商开始采用具有 250 和/或 200 微米包覆层的可卷曲带状光缆结构。

传统带状光纤中的 12 芯光纤沿着光缆方向完全粘合在一起，而可卷曲带状光纤则采用不连续地粘合方式，因此可将光纤卷起来，而无需将其平放。通常，这类型的设计允许在一个两英寸的管槽内放置 2 根 3,456 芯光缆，而扁平设计只能在同一空间内放置单根 1,728 芯光缆（使用管道最大占用率为 70% 的光缆）。



可卷曲带状光缆以不连续的方式粘连在一起。

资料来源：ISE 杂志

200 微米光纤保留了标准的 125 微米包层，可完全向后兼容当前光纤和新兴光纤；区别在于典型的 250 微米涂层减少为 200 微米。当与可卷曲带状光纤连接时，更小的光纤直径使布线制造商能够保持光缆尺寸不变，同时使光纤数量比传统的 250 微米扁平带状光缆增加一倍。

为满足数据中心之间日益增长的连接需求，超大规模数据中心开始部署可卷曲带状光纤和 200 微米光纤等技术。在数据中心内部，叶交换机到服务器的连接距离更短，密度更大，主要考虑因素是光纤模块的成本和运营成本。

为此，许多数据中心仍在使用由多模光纤提供支持的低成本垂直腔面发射激光（VCSEL）收发器。还有许多数据中心则选择混合方法，即在核心层网络中使用单模光纤，而服务器与一级叶交换机之间采用多模光纤连接。随着越来越多的设施采用 400G，网络管理者将需要这些选项来平衡成本与性能，因为 50G 和 100G 光纤连接至服务器已成为常态。

80 km DCI 空间：相干检测与直接检测光通讯技术

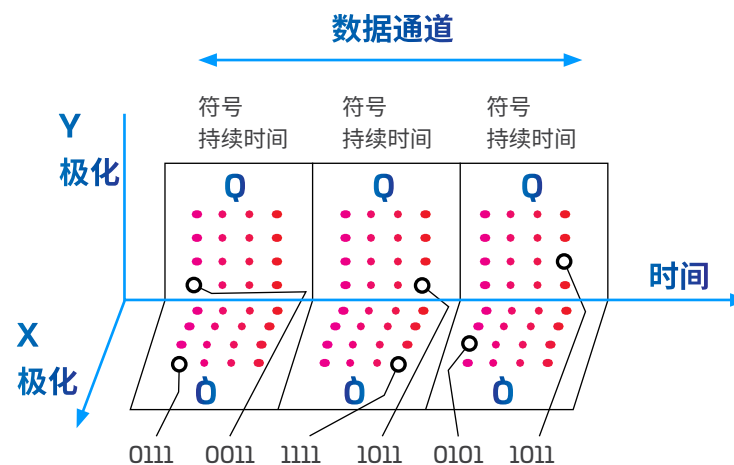
随着地区数据中心集群趋势的继续发展，对高容量、低成本 DCI 链路的需求变得越来越迫切。开始出现一些新的 IEEE 标准，以提供各种支持即插即用、点对点部署的低成本选项。

基于用于直接检测的传统四电平脉幅调制（PAM4）的收发器将可实现高达 40 km 的链路，同时直接兼容近期推出的 400G 数据中心交换机。还有其他一些发展也针对传统 DWDM 传输链路的类似功能。

随着链路距离增加至 40 km 至 80 km 及以上，为远程传输提供增强支持的相干系统有可能占领大部分高速市场。

相干光纤克服了模式色散和极化色散等限制，成为更长链路的理想技术之选。一直以来，相干光纤都是高度定制的解决方案（且成本非常高），同时需要使用定制“编码模式”，而非即插即用光纤模块。

随着技术进步，相干解决方案可能变得更小巧，并有可能降低部署成本。最终，相对成本差异可能会缩小，直至将该技术用于更短链路。



资料来源：www.cablelabs.com/point-to-point-coherent-optics-specifications

对持续的高速网络升级进行整体分析

数据中心不断提高速度的过程是一个渐进的过程；随着应用和服务的发展，存储速度和服务器速度也必须提高。采用一种系统化方法来处理周期性重复升级有助于减少计划和实施更改所需的时间和成本。康普建议对交换机、光纤和光缆布线作为三个独立的协作方式进行整体分析。

最终，这些组件的协同工作方式将决定网络能否可靠且有效地支持未来的新应用。目前的挑战是 400G；未来，将变成 800G 或 1.6T。即使网络技术不断变化，对于高质量光纤基础设施的基本要求仍将保持不变。





5 /

请勿只着眼于现在 - 800G 即将来临!

100G 光纤正大量涌入市场，预计明年将实现 400G。与此同时，数据流量继续增加，数据中心的压力只会越来越大。

实现三方平衡

在数据中心，容量就是服务器、交换机和连接三方之间相互制衡的问题。每一方的发展都会推动其他方提升速度，并降低成本。多年以来，交换机技术一直都是主要驱动力。由于 Broadcom StrataXGS Tomahawk 3 的推出，数据中心管理者如今可以将交换和路由速度提升至 12.8 Tbps，并将每个端口的成本降低 75%。那么，如今的限制因素就是 CPU，对吗？事实并非如此。今年早些时候，NVIDIA 推出了适用于服务器的新型 Ampere 芯片。事实证明，用于游戏的处理器非常适合处理人工智能和机器学习所需的训练和基于推理的处理操作。

网络变成瓶颈

随着交换机和服务器将如期支持 400G 和 800G，保持网络平衡的压力转到了物理层。2017 年批准通过的 IEEE 802.3bs 为 200G 和 400G 以太网铺平了道路。然而，IEEE 最近才完成了 800G 及以上的带宽评估。鉴于制定和采用新标准所需的时间，我们可能已经落后了。

因此，布线和光模块制造商正努力保持发展势头，因为业界希望支持从 400G 到 800G、1.6Tb 及以上的持续过渡。以下是我们了解到的一些趋势和发展。

交换机在演进

首先，服务器排列和布线架构在不断变化。汇聚交换机开始从机架顶部 (TOR) 移至列中 (MOR)，并通过结构化布线配线架连接至交换机矩阵。如今，向更高速度迁移只需更换到服务器的跳线，而无需更换较长的交换机至交换机的主干线缆。这种设计还无需安装和管理交换机与服务器（每一条都是针对某一应用的，因此速度也都是特定的）之间的 192 条有源光缆 (AOC)。

不断变化的收发器类型

可插拔光纤模块的新设计为网络设计人员提供了额外的工具，主要是支持 400G 的 QSFP-DD 和 OSFP。这两种收发器类型都具有 8 个通道，并采用可提供 8 x 50G PAM4 的光模块。当部署在 32 端口配置中时，QSFP-DD 和 OSFP 模块可在 1RU 设备中实现 12.8 Tbps 的速度。OSFP 和 QSFP-DD 收发器类型支持目前的 400G 光纤模块和下一代 800G 光纤模块。通过使用 800G 光模块，交换机将实现每 U 25.6Tbps 的连接速度。

新的 400GBASE 标准

此外，还可以使用更多接头选项来支持 400G 短程多模光纤模块。400GBASE-SR8 标准支持 24 芯 MPO 接头（适用于使用 16 条光纤的传统应用）或单排 16 芯 MPO 接头。云级服务器连接的早期首选是单排 MPO16 接头。另一个选项 400GBASE-SR4.2 采用具有双向信号传输功能的单排 MPO12 接头，因此适用于交换机至交换机连接。IEEE802.3 400GbaseSR4.2 是在多模光纤上利用双向信号传输的首个 IEEE 标准，该标准采用 OM5 多模布线。OM5 光纤可为 BiDi 等应用提供多波长支持，使网络设计人员可以将应用距离延长至 OM4 的 1.5 倍。

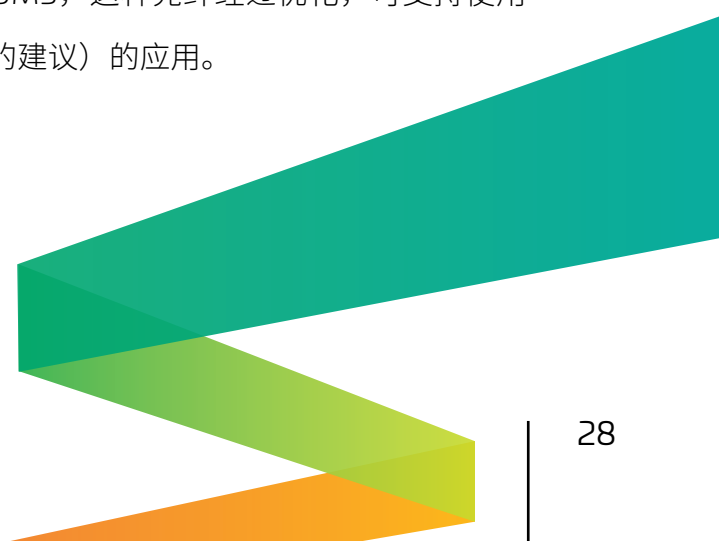
但我们的速度足够快吗？

据行业预测，未来两年内将需要使用 800G 光模块。因此，2019 年 9 月成立了 800G 可插拔光模块 MSA 行业联盟，以开发新应用，包括覆盖范围在 60 至 100 米范围内的低成本 8x100G SR 多模模块。我们的目标是提供初期低成本 800G SR8 解决方案，使数据中心能够支持低成本服务器应用。800G 可插拔光模块解决方案支持增加交换机的收敛比和提升每机架服务器数量。

与此同时，IEEE 802.3db 工作小组还在研究面向 100G 单波长的低成本 VCSEL 光模块解决方案，并证明了通过 OM4 MMF 实现 100 米覆盖范围的可行性。如果成功，这项工作可以将服务器连接从机架内 DAC 转为 MOR/EOR 提高交换机收敛比。这样就可以实现低成本光纤连接，并可扩展针对传统 MMF 布线的长期应用支持。

企业数据中心对更大容量的需求持续上升，并且需要采用新策略来提高较大安装量的多模光纤 (MMF) 布线基础设施的安装速度。在过去，在 MMF 中增加更多波长能够非常成功地提高网络速度。

TB 级双向 (BiDi) 多源协议 (MSA) 联盟建立在 40G BiDi 成功的基础之上，旨在开发面向并行 MMF 的可互操作 800 Gbps 和 1.6 Tbps 光学接口规范。作为该 BiDi MSA 联盟的创始成员，康普推出了多模光纤 OM5，这种光纤经过优化，可支持使用多种波长（如该 MSA 的建议）的应用。



MMF 备受数据中心运营商的欢迎，因为它支持旨在降低网络硬件资本支出以及运营支出（由于供电需求降低）的短距离高速链路。通过扩大面向 BiDi 应用的距离支持，OM5 还可以进一步提高 MMF 的价值。就 IEEE 802.3 400G BASE4.2 来说，OM5 提供的距离比 OM4 布线长 50%。未来，当我们介绍向 800G 和 1.6T BiDi 发展的后续步骤时，OM5 的优势会越发显著。

使用在 IEEE802.3.db 和 IEEE802.3.cm 中开发的技术，这种新的 BiDi MSA 将提供基于标准的网络，还将支持单光纤 100G BiDi，支持 200G BiDi 的双工光纤，以及基于不断发展的 QSFP-DD MSA（8 通道）和 OSFP-XD MSA（16 通道）而添加的额外光纤，以分别实现 800G 和 1.6T 数据传输速率。

2022 年 2 月 28 日，MSA 表示：

“通过利用安装量较大的 4 对并行 MMF 链路，该 MSA 可将基于 400 Gb/s BiDi 的并行 MMF 升级到 800 Gb/s 和 1.6 Tb/s。实践证明，BiDi 技术可成功地将已安装的双工 MMF 链路从 40 Gb/s 升级到 100 Gb/s。TB 级 BiDi MSA 规范将满足现代数据中心中交换机之间以及服务器-交换机互连之间的关键型大容量链路应用需求。”

“由于这一 MSA，同样的并行光纤基础设施将能够支持从 40 Gb/s 到 1.6 Tb/s 的不同数据速率。该 MSA 成员正在响应行业对低成本、低功耗、基于 BiDi 多模技术解决方案的需求。有关 TB 级 BiDi MSA 的更多信息，请访问：terabit-bidi-msa.com。”

资料来源：terabit-bidi-msa.com

那么，我们的进度如何？

万物发展瞬息万变，而且提醒一下它将变得更快。好消息是，标准机构和行业之间正在进行重要且有潜力的发展，有助于数据中心实现 400G 和 800G。然而，扫除技术障碍只是挑战的一半。另一半是时机。由于更新周期为每两到三年，且新技术也在加速投入使用，因此对运营商来说，适当地把握转型时机变得更加困难，如果未能做到这一点，其成本就会变得更高。

在变革中存在许多变数。像康普这样的技术合作伙伴可以帮助您应对不断变化的形势，并做出符合您最佳长期利益的决策。

A man and a woman in business attire are walking on a modern glass-enclosed walkway. The man is holding a folder and pointing at it, while the woman is holding a coffee cup and a tablet. They appear to be in a professional setting, possibly a corporate office or a modern building. The background shows a large glass building with multiple floors.

6 / 网络边缘的 MTDC

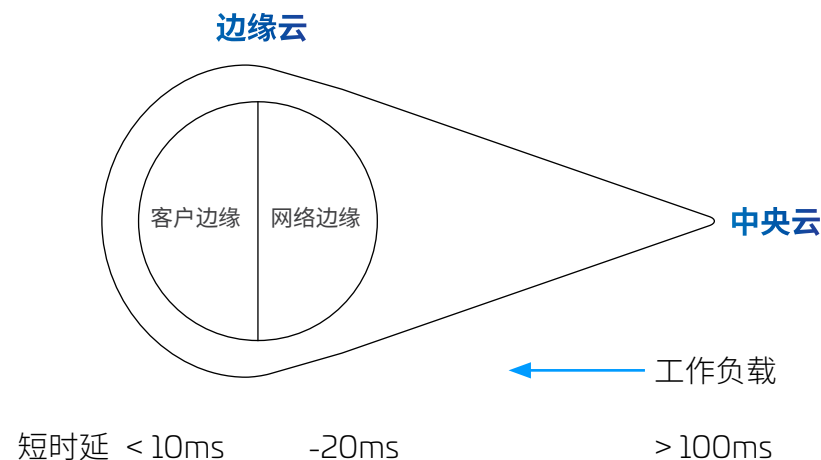
“边缘计算”和“边缘数据中心”是近来 IT 行业中越来越常见的术语。多租户数据中心 (MTDC) 现在正处于利用其网络位置的边缘。要想了解方式和原因，我们首先需要定义“边缘”。

“边缘”是什么？位于何处？

“边缘”这个词有点让人误解，因为它可能比名称所暗示的更靠近网络的核心；并且不是一个具体的对于边缘的定义，而是两个。

第一个定义是客户边缘，位于客户营业场所附近，用于支持超低延迟应用。例如，一家制造厂需要一个网络来支持 5G 全自动机器人。

另一个定义是网络边缘，位于网络核心附近。这种模式有助于支持云辅助驾驶和高清游戏等应用所需的低延迟。MTDC 就是在网络边缘得以蓬勃发展。



灵活性和适应性

MTDC 既灵活又可适应各种客户配置，因此能够充分利用其在网络边缘的位置以及靠近人口密集区域的优势。一部分 MTDC 客户将能够清楚他们的要求，并提供其自己的设备。另一部分客户将其应用从自己的原来的场所中转移到 MTDC，这将需要专家指导以支持他们的应用。成功的 MTDC 应能够适应这两种场景。

不仅在初始设置中需要操作灵活性；MTDC 内部连接在初次建设和升级改造中也必须具有灵活性。为实现这种灵活性，您需要考虑您的基础设施，即综合布线。在用户围笼内的建议架构是基于叶脊网络的架构。使用大芯数主干光缆（如 24 芯或 16 芯 MPO）可使叶交换机和脊交换机之间的布线保持稳定，因为它们有足够数量的光纤来支持未来几代的网络速度。

比如说，由于以太网光模块可以从双工端口转变为并行端口，然后又变回双工端口，所以您只需要简单的更换叶交换机和脊交换机所在机架中的预端接模块和光纤跳线的类型就能够满足光模块的变化。这样就无需破坏和更换主干布线。

当叶 (Leaf)-脊 (Spine) 架构就位后，就需要考虑其他因素，以确保 MTDC 能够轻松地满足未来围笼内外链接的速度和带宽需求。为实现这一点，我们必须了解服务器机柜及其组件，并确定至这些机架的布线路径是否有足够的空间来支持未来的移动、增添和变更操作，尤其是引入新服务和客户时。请记住，增添和修改操作必须简单迅速，并且可从远程地点进行。在这种情况下，自动化基础设施管理系统可监控、映射和记录整个网络中的无源连接。随着更多应用和服务进入市场，手动监控和管理布线网络很快变得不切实际。



要想更深入地了解 MTDC 如何才能优化边缘的资本化，请查阅康普的白皮书：“**网络边缘的 MTDC 面临的新挑战和机遇。**”

A blue-tinted photograph of a server room aisle. The aisle is lined with rows of server racks on both sides, receding into the distance. A person is standing in the middle of the aisle, looking down at a device or document they are holding. The ceiling has circular light fixtures. The overall image has a strong blue color cast.

7 /

数据中心在 5G 世界中的角色演变

几十年来，数据中心一直位于或靠近网络中心。对于企业、电信运营商、有线电视运营商以及近来的 Google 和 Facebook 等服务提供商而言，数据中心是 IT 的心脏和肌肉。

云技术的出现更是彰显了现代数据中心的重要性。但仔细聆听，你会听到变革来临时的隆隆声。

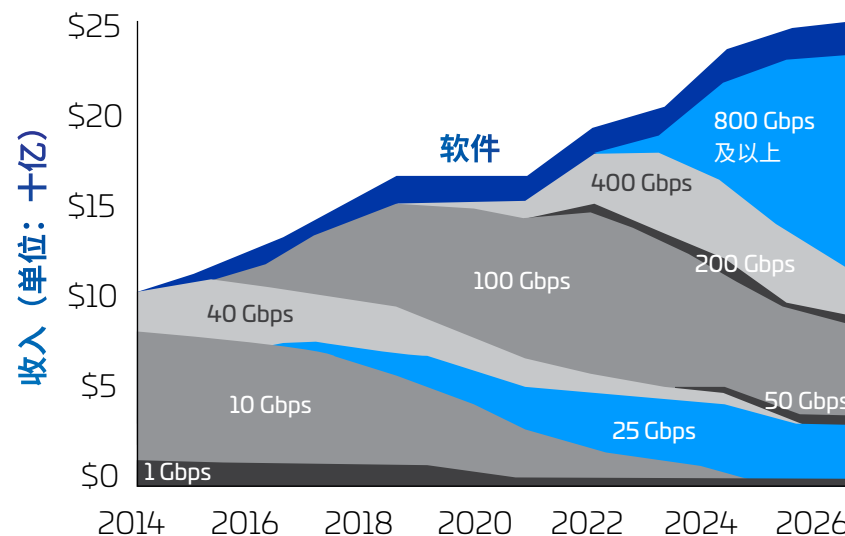
随着网络计划迁移至 5G 和物联网，IT 管理者开始关注边缘，并且需要将更多容量和处理能力部署在更靠近最终用户的位置。在此过程中，他们将会重新评估数据中心的作用。

Gartner¹ 认为，到 2025 年，将有 75% 的企业生成数据在边缘创建和处理——而 2018 年这一数据仅为 10%。

与此同时，数据量方面也要准备迎接新的挑战。一辆自动驾驶汽车预计平均每小时可产生 4T 数据。

现在网络都在争相确定如何才能更好地应对边缘流量的激增以及对低延迟性能的需求——而且不会破坏其现有数据中心的设施投资。

答案之一是大量投资于东西向网络链路和对等冗余节点，在产生数据的地方建设更多的处理能力。但数据中心呢？它们将会发挥什么样的作用？



资料来源：650 Group，2020 年 12 月市场情报报告

¹ 边缘计算对于基础设施和运营领导者意味着什么；Smarter with Gartner；2018 年 10 月 3 日

AI/ML 反馈循环

未来超大规模和云计算数据中心的商业应用主要在于其庞大的处理能力和存储容量。随着边缘活动日益活跃，需要借助数据中心的力量创建数据处理算法。在支持物联网的世界，人工智能 (AI) 和机器学习 (ML) 的重要性不能低估。将其付诸实施的数据中心的角色也同样不容忽视。

为了生成驱动 AI 和 ML 所需的算法，需要处理大量数据。核心数据中心已经开始部署更强大的 CPU，以配合 TPU 或其他专用硬件。此外，这通常需要超高速、大容量网络，以及先进的交换层，为处理同一问题的服务器库提供支持。AI 和 ML 模型正是这些艰辛努力的成果。

另一方面，需要将 AI 和 ML 模型放在能够对业务产生最大影响力的位置。例如，对于面部识别等企业 AI 应用，超低延迟要求需要它们在本地部署，而不是在核心部署。但是，还必须定期调整模型，以便将边缘收集的数据反馈到数据中心，从而更新和优化算法。

利用沙盒还是拥有沙盒？

AI/ML 的反馈循环是一个例子，说明数据中心需要能够支持更广泛和多样化的网络生态系统，而不是控制它。对于超大规模数据中心领域的重要参与者来说，适应分布更广、协作性更强的环境并不是件易事。他们希望确保，如果您在部署人工智能或机器学习或访问边缘时，您会使用他们的平台，但不一定要使用他们的设施。

像 AWS、Microsoft 和 Google 之类的提供商现在正在向客户所在的位置（包括私人数据中心、中心机房和企业内部的本地存储）配置具有一定容量的机架。这使客户能够使用提供商的平台，在其设施中构建和运行基于云的应用。由于这些平台也嵌入在许多运营商的系统中，因此客户也可以在运营商存在任何位置运行其应用。该模式仍处于起步阶段，为客户提供了更大的灵活性，同时使提供商能够在边缘进行控制并声明所有权。



与此同时，其他模型展现出了一种更开放、更具包容性的方法。边缘数据中心制造商开始设计可以提供标准化计算、存储和网络资源的托管数据中心。较小的客户（如游戏公司）可租一台虚拟机来服务其客户，而数据中心运营商则会以收益分享模式向您收费。对于一个正在竞争进入边缘市场的小企业来说，这是一个极具吸引力的模式（或许是它们赢得竞争的唯一方式）。

基本挑战

随着下一代网络前景逐渐明朗，行业必须应对实施方面的挑战。就数据中心内部而言，我们了解到：服务器连接将从每通道 50G 增加到 100G；交换机背板带宽将增加到 25.6T；而迁移到 100G 技术将让我们获得 800G 可插拔模块。



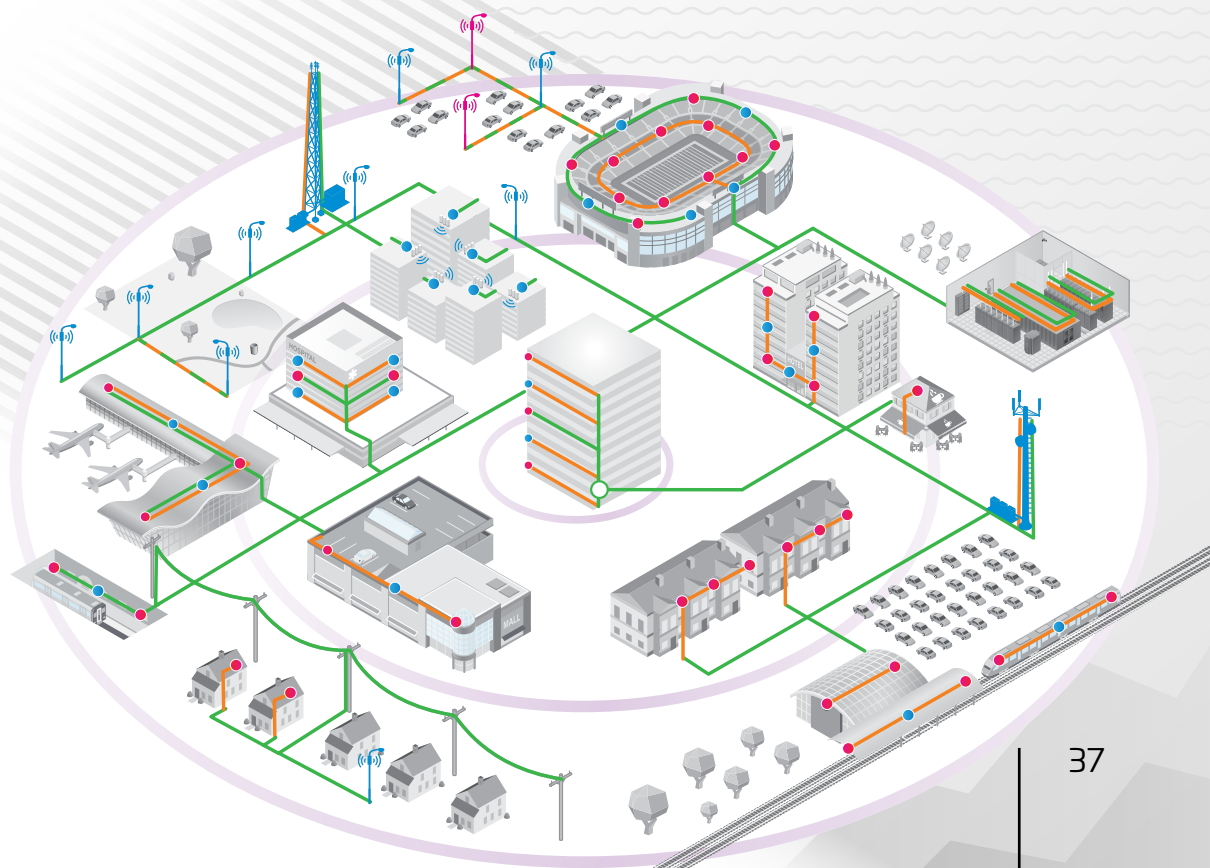
我们尚不明确如何设计从核心到边缘的基础设施——具体来讲，就是我们如何实施 DCI 架构和城域网/远距离链路并支持高冗余对等边缘节点。另一个挑战是开发实现海量流量管理和路由所需的协调和自动化功能。随着行业迈向支持 5G/物联网的网络，这些问题都亟需解决。

携手共进

我们深知构建和实施下一代网络需要协调一致共同努力。

云数据中心的低成本、大容量计算和存储能力都无法在边缘数据中心复制，因此云数据中心依然会发挥一定作用。但是，随着网络内部的责任变得越来越分散，数据中心的工作将从属于更大的生态系统。

将所有这些组合在一起，将得到一个更快、更可靠的物理层，从核心开始一直延伸到网络的最边缘。这个布线和连接平台——由传统的以太网光学和相干处理技术提供支持——将扩大容量。使用共封装光模块和硅光子技术的新型交换机将进一步提高网络效率。当然，到处都需要更多的光纤——采用超高密度、紧凑的光纤布线，这将支撑网络性能的发展。





8 /

遍布园区，进入云端：
MTDC 连接的推动因素有哪些？

在数据中心，尤其是在多租户数据中心 (MTDC) 工作将会是一段令人难以置信的经历。最近，机械、电气和冷却设计领域取得了较大进展。如今，人们的关注点开始转移至物理层连接，以使租户能够快速、轻松地扩展至云平台，并从云平台向外扩展。

在 MTDC 内部，客户网络正在迅速扁平化，并向东西方向扩散，以满足日益增长的数据驱动型需求。曾经不相干的围笼、楼层和园区现如今均可实现互连，以跟上物联网管理、增强现实集群和人工智能处理器等应用的发展步伐。然而，进入数据中心的连接以及数据中心内部的连接却滞后了。

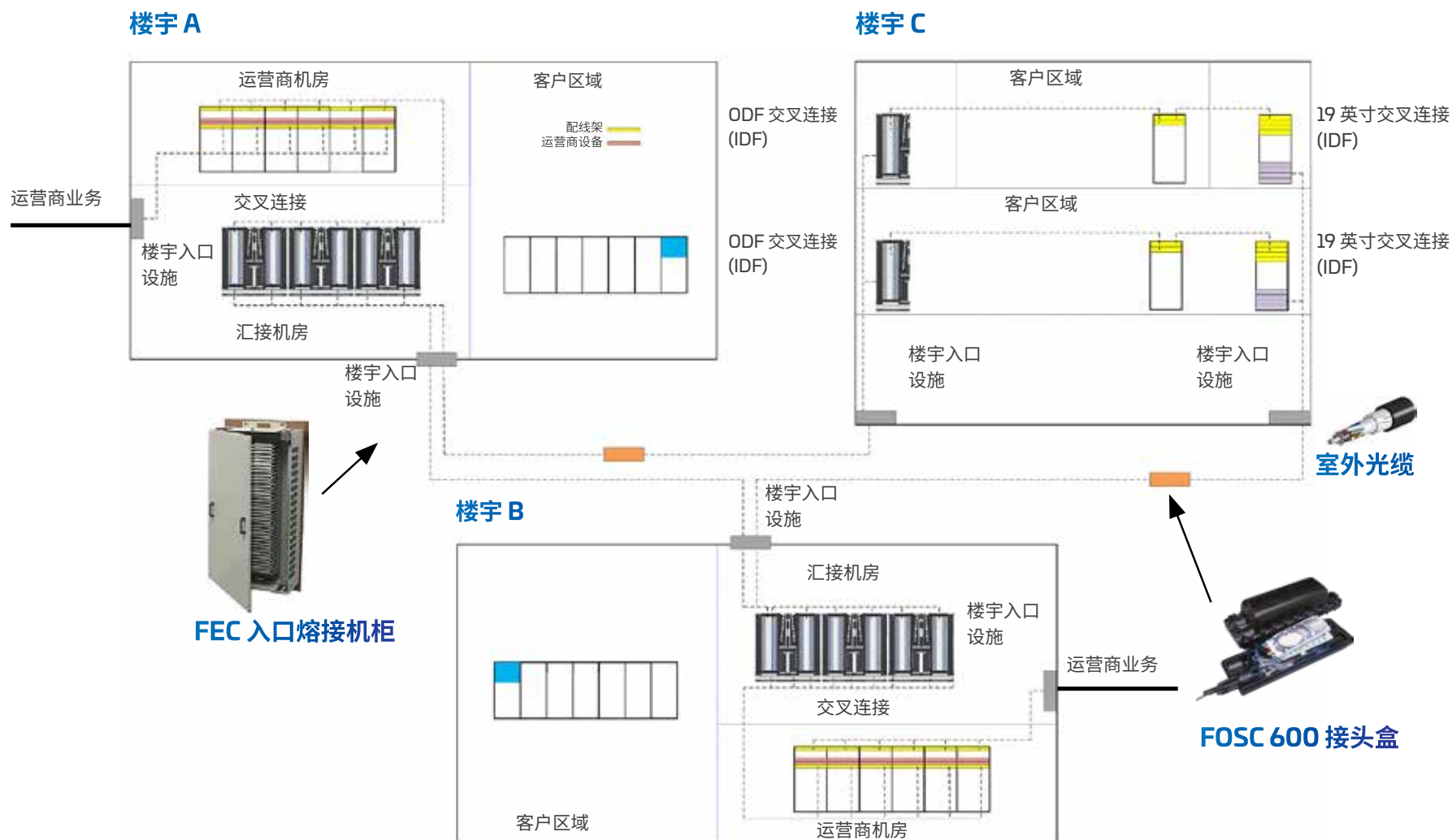
为了解决这些连通性方面的差距，MTDC 提供商开始使用虚拟网络作为云接入平台。设计公有云、私有云和混合型云网络内部及相互之间的布线架构极具挑战。以下重点介绍了 MTDC 用于创建可扩展的云互连方法的许多趋势和策略中的一部分。



连接 MTDC 园区

云连接的挑战始于室外。大芯数光纤布线和多样化路由可在现有和未来建筑之间形成网格。在进入设施之前，可使用光纤

接头盒将这些室外 (OSP) 光缆与室内光缆熔接在一起，以便在进线间内部进行配线。这适用于进线间内的配线架和机架已预先端接了光缆的情况。或者，可以使用大芯数的光纤进线机柜 (FEC) 将 OSP 室外光缆直接熔接在每个建筑的进线间 (EF) 内。



当园区内建造了其他楼宇 C 时，就会由数据中心 A 或 B 提供服务。最终，A, B, C 建筑内任何网络流量都可以轻松地在整个园区重新定向，从而提高了可用性，并降低了网络中断风险

这些建筑互连越来越多地采用高密度可卷曲带状光纤进行互联。独特的网状配置使整个光缆结构更小巧、更灵活，同时允许制造商将 3,456 根光纤或更多光纤敷设在现有管井中，或使得新建的更大的管井组得到最大化的利用。可卷曲带状光纤可实现传统封装光纤的两倍密度。其他优势包括：

- 更轻巧的线缆可简化操作、安装和子单元分支
- 无过度弯曲可降低出现安装错误的风险
- 易于分辨的标识便于备线/熔接和预制
- 更细的线缆具有更小的弯曲半径，适用于接头盒、配线架和手孔

改进的入口设施连接性

数以千计的 OSP 室外光纤汇聚在 EF 进线间 内并连接至 ISP 室内 光纤，因此对 EF 进线间 内部可管理性的关注促进了光纤进线机柜 (FEC) 和光纤配线架 (ODF) 的显著改进。

作为光纤布线管理的战略要点，ODF 通常会被忽略。然而，能否精确识别、保护和重复使用闲置容量可能是数天实现和数月实现全园连接的效率差异。

康普 FEC 产品包括地板式安装、壁挂式安装和机架式安装等设计，可管理超过 10,000 芯光纤。其它优点包括：

- 具有更大的托盘密度，实现大规模熔接
- 可从 OSP 室外光缆有序地转换到 ISP 室内光缆
- 能够将大芯数光纤分支为小芯数光纤

ODF 对于现代会接间 (MMR) 的顺畅运行至关重要，自最初为电信和广电网络而开发以来，已经实现了长足进步。例如，自带内置直观路由的 ODF 可组合成一排，以支持采用同一长度的跳线，并在机架内管理超过 50,000 芯光纤。从机械方面看，ODF 还可提供出色的前端跳线管理，从而简化了库存管理和安装程序。

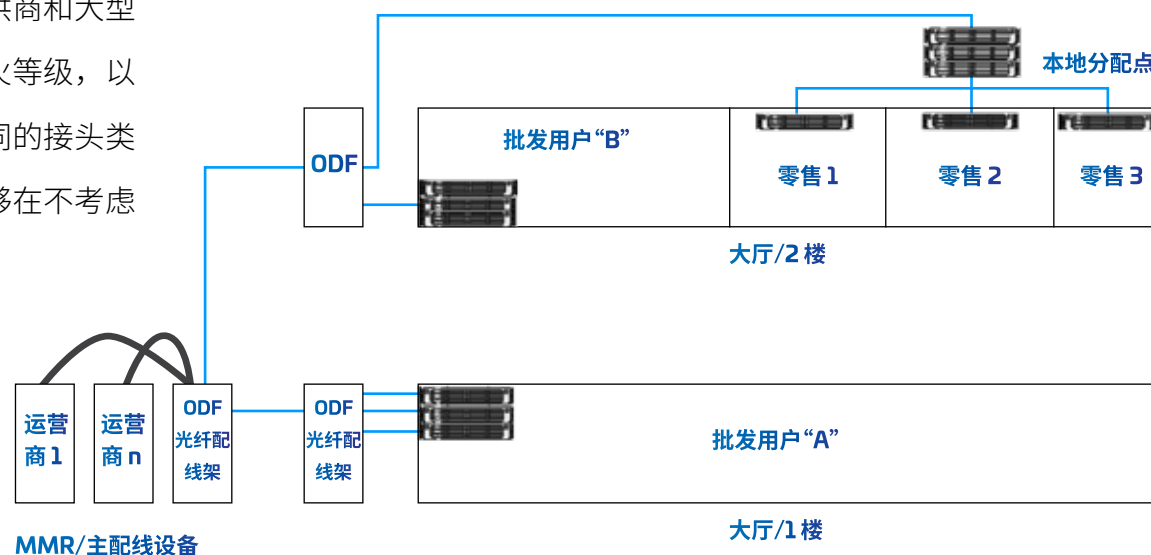
随着对单端连接器布线的需求持续增长，熔接大芯数预端接光缆的能力被设计到组件中。

支持 MTDC 内的云连接

随着 IT 应用将备用设备转移至公共和私有云领域，访问 MTDC 园区内的云服务提供商变得越来越重要。云服务提供商和大型企业由于业务遍布全球，需要不同的光缆结构和防火等级，以满足各个国家的本地法规要求。它们还需要使用不同的接头类型和光纤数，以匹配其网络基础设施架构，从而能够在不考虑安装者技能的情况下进行快速扩展，并保持一致。

当然，云连接要求根据租户的类型而有所不同。例如，使用私有和混合云的传统企业可能需要在围笼（或机柜）区域内建立更低密度的连接。

为连接从 MMR 到围笼/机架区域，MTDC 正以 12 和 24 SMF 的初始标配增量部署光纤。一旦租户搬走，就无需拆卸大量线缆。MTDC 只需将“最后一米”光纤卷起来，并重新将其部署至另一个围笼或界线位置，即可重复使用“最后一米”光纤，将其敷设至重新配置的可用空间。这些围笼（通常少于但不限于 100 个机架）内的结构化布线可实现至私有和公共云提供商的可扩展连接。



另一方面，云服务提供商具有各种极其多变的连接要求。连接至这些围笼的光纤数量通常要比连接至企业的高得多，且有时候可以将整个园区范围内的多个围笼直接整合在一起出租。这些提供商每年都要部署几次新的物理基础设施布线，并根据施工成本不断评估和改进设计。

这涉及到审查从光收发器和 AOC 到光纤类型和预端接组件的所有产品的成本效益。

通常，至 MTDC 的云服务提供商布线链路采用大芯数光缆和多样化路由，以减少故障点。其最终目标就是，为不同的密度和空间交付可预测的模块。一致性可能很难保证，因为随着光收发器变得更加专业，与直觉相反，找到正确的光模块和连接器往往变得更加困难，而不是更容易。

例如，目前的收发器具有不同的接头类型和损耗预算要求。双工 SC 和 LC 接头不再支持所有光收发器选项。如今，云级网络中开始部署密度更高的新型特定应用接头，如 SN 接头。因此，在接头尺寸和纤芯数量之间选择具有高互操作性的收发器至关重要。

保持连网，了解动态

在整个 MTDC 园区内，要实现各种建筑的互连，并提供对于零售和批发客户来说至关重要的云连接，这一需求推动着网络架构的内外变革。该前瞻性观点只触及了这个日益复杂的话题的一小部分。

有关发展趋势的更多信息或想了解日新月异的技术发展，请联系康普。了解未来趋势是我们的工作。

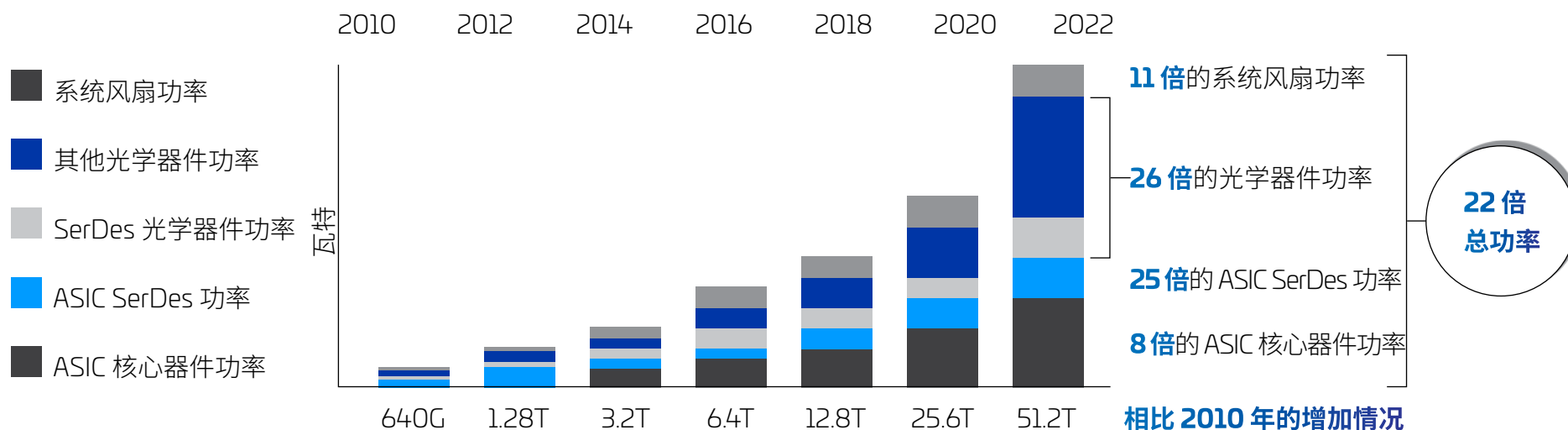


9 /

通往 1.6T 的道路正在开启

规划指数级增长的挑战在于，变化往往比我们预期的更频繁，也更具破坏性。超大规模和多租户数据中心管理者正在亲身经历这一挑战。他们才刚开始升级到 400G 和 800G 数据速率，但标准已经上升到 1.6T。竞争已经拉开帷幕，如果数据中心运营商能够成功地增加应用容量，并降低服务成本，那么每个人都能从中获益。这样它们就能够降低最终用户的成本，同时有助于提高互联网的能效。就任何飞跃来说，每次成功都孕育着另一项挑战。更高的容量会产生更多需要大量数据且耗电量高的新型应用，而这些应用又需要更多容量。如此循环往复。

在 Google、Amazon 和 Meta（原 Facebook）等企业的推动下，云服务、分布式云架构、人工智能、视频和移动应用工作负载的爆炸性增长将很快超越 400G/800G 网络功能的发展。问题不仅是带宽容量，还有运营效率。数据网络开销在整体传输成本中所占的比例越来越大。这些成本反过来又受功耗的驱动，从而产生下一代设计目标。最终目标就是降低每比特功耗，并使这种不可思议的爆炸性增长成为一种持续的可能性。



资料来源: www.ethernetalliance.org/wp-content/uploads/2021/02/TEF21.Day1_Keynote.RChopra.pdf, 2021 年 1 月 25 日; Rakesh Chopra、Mark Nowell、Cisco Systems

思考——功率与网络的关系

如果使用当前的网络交换机来扩展网络容量，很可能无法满足功率需求。（请注意，在环境可持续性背景下审查公司的每个决策时，也会遇到功率问题。）

在降低每比特功率比（常见的能效指标）方面，网络面临的压力越来越大，其最终目标是降至 5pJ/bit。事实表明，增加网络交换机的密度（收敛比）是解决这一问题的有效途径。这可大大提高交换机的容量和效率。

简单来说，交换机的整体功耗越来越令人担忧。2021 年技术探索论坛上发表的主题演讲表明，2010 年至 2022 年，交换机功耗增长了 22 倍。详细地说，功率增加的主要组件与 ASIC 芯片和光发射器/接收器之间的电子信号有关。因为电气效率会随着交换速度的增加而降低，而交换速度受电气速度的限制。目前，实际限制为 100G。

降低功耗的途径在于使用更大且更高效的交换组件、更高的信号传输速度和更高的密度。理论上说，这种途径最终可实现 102.4T，但鉴于当前交换机的设计，这个目标似乎极具挑战性。因此，有人主张基于单点解决方案的策略。这可以解决电气信号传输挑战 (flyover cables vs PWB)，并支持继续

使用可插拔光模块。此外，还可以选择将信号传输速度提高至 200G，而其他一些人则建议将通道数翻倍 (OSFP-SD)。但是，另一方则主张使用运营平台方法，以推动行业使用更长远的解决方案。从根本上提高密度和降低每比特功率的系统方法涉及使用共封装光纤 (CPO)。

共封装光模块和近封装光模块 (CPO/NPO) 所发挥的作用

CPO 和近封装光模块 (NPO) 的拥护者认为，要想实现 1.6T 和 3.2T 交换机所需的每比特功率目标，就需要使用新的架构，而 CPO/NPO 刚好符合要求。它们是理想之选，因为 CPO 技术可将电气信号传输限制在非常短的距离内，无需使用重定时器，并可优化 FEC 方案。将这些新技术大规模推向市场需要整个行业致力于改良网络生态系统。新标准将大力推进这一产业转型。

使用 CPO 面临的一项挑战就是，它们含有不可现场维修的光模块，且要求实现极低故障率 (FIT)——与可现场维修的可插拔光模块相比，CPO 必须实现这一点。使用 CPO 时要注意的一个重要事实就是，它需要时间来变得成熟。

行业需要采用新的互操作性标准，而供应链也必须发展，以便支持 CPO。许多人认为，考虑到 CPO 相关的风险因素，可插拔模块似乎更适用于实现 1.6T。

交换机设计人员和制造商也建议使用可插拔光学器件来实现 1.6T（基于 100G 或 200G 电气 SERDES 速率）。这种方法不需要进行全面改革，因此可降低风险，并缩短这种备选方案的上市时间。这并不是毫无风险的方法，但支持者认为，这种方法所面临的挑战远比使用 CPO 少。

200G 电信号传输

通过使用更多 I/O 端口或更高信号传输速度，可进入下一个交换节点（容量翻倍）。对于每种替代方案的应用级和系统级驱动程序，选择时需要基于带宽的使用方式。更多的 I/O 可用于增加交换机支持的设备数量，而更高的聚合带宽组合则可用于距离更长的应用，以减少支持更高带宽所需的光纤数。

2021 年 12 月，4x400G MSA 联盟 建议采用 1.6T 模块，可选择使用 16x100 或 8x200G 电气通道，以及通过 16 通道

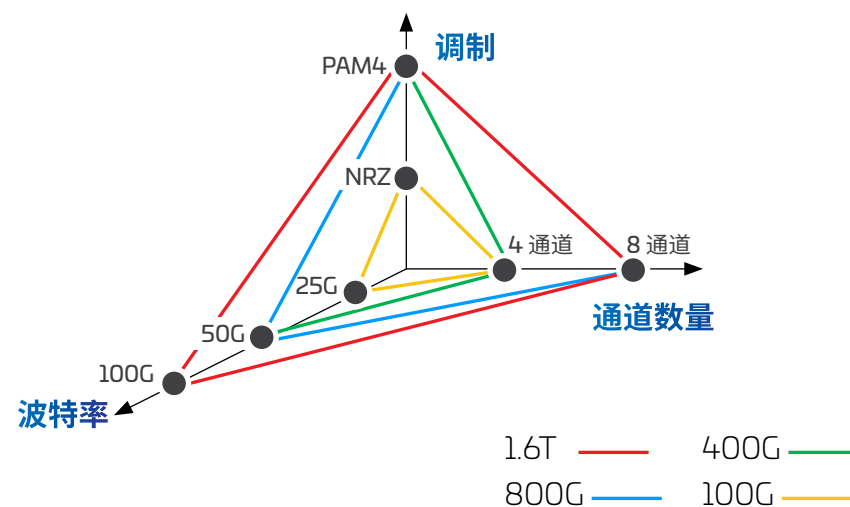
1 OSFP-XD MSA 包含两个 MPO16 接头（总共支持 32 条 SMF 或 MMF 光纤）的选项。

2 实现 1.6T PAM4 光学器件的正确途径：8x200G 或 16x100G；Light Counting；

2021 年 12 月

OSFP-XD 收发器类型映射的各种光学选项。高收敛比应用在数据速率为 100G 时，需要 16 个 100G 的双工连接（32 条光纤）¹（可能为 SR/DR 32），而更长距离选项在数据速率为 200G/400G 时则符合之前几代的标准。

尽管供应商已证明了 200G 通道的可行性，但客户仍担心业界是否能够制造出足够的 200G 光学器件，以便降低成本。再现 100G 的可靠性以及验证芯片质量所需的时长也是潜在问题。²



升级至 200G 通道的潜在路径。资料来源：Marvell

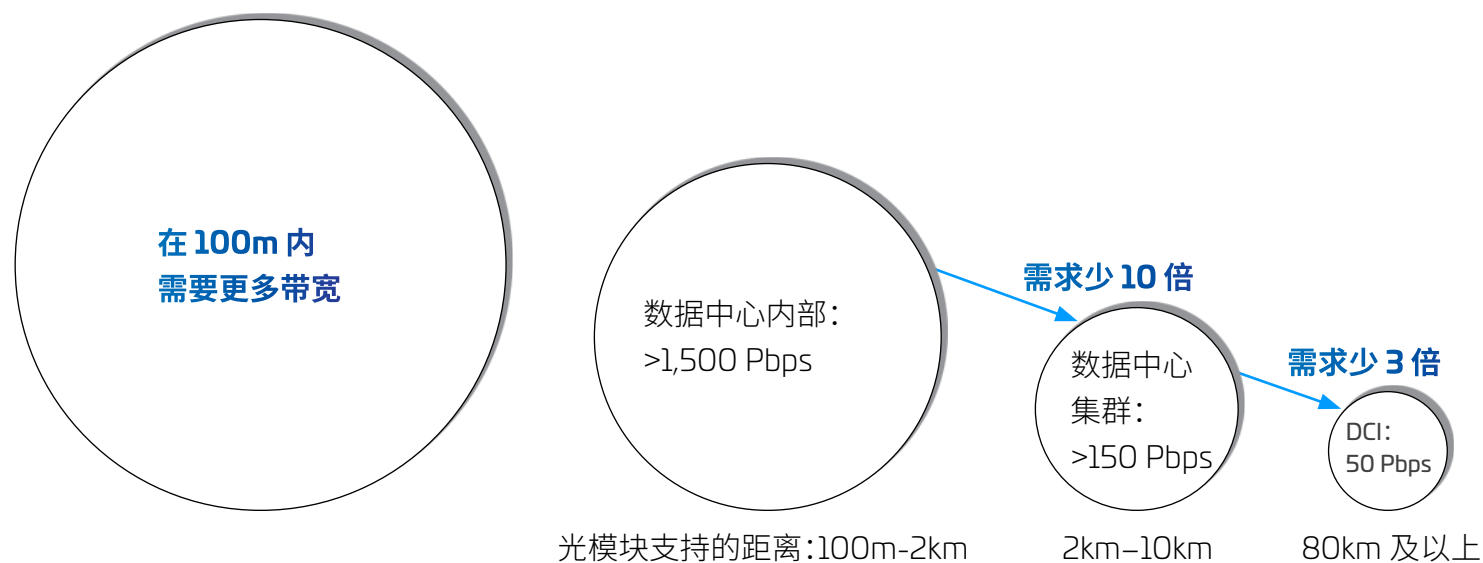
最终，任何 1.6T 迁移路线都需要使用更多的光纤。MPO16 可能发挥至关重要的作用，因为它可提供具有极低损耗和较高可靠性的更宽通道。此外，它提供的容量和灵活性可支持更高收敛比的应用。与此同时，随着数据中心内部的链路变得越来越短，会逐渐倾向于使用多模光纤，因为其光模块的成本更低，而且延迟、功耗以及每比特功率性能都得到了改进。

那么，如何看待长久以来关于铜缆会被逐渐取代的这种观点？速度越高，使用铜缆 I/O 的限制就越大，因为根本不可能实现每比特功率与距离之间的合理平衡。即使是最终由光纤系统占主导的短距离应用亦是如此。

我们所知道的

以上这一切都表明，尽管实现 1.6T 的最佳途径尚难以确定，但有些方面已经开始成为人们关注的焦点。在未来几年内扩大容量、提高速度并显著提高效率是大势所趋。为了准备好日后扩展这些新技术，我们需要从现在开始进行设计和规划。

欢迎访问 zh.commscope.com，进一步了解您当前可采取的步骤，以确保您的光纤基础设施能够满足未来所需。



当数据速率为 100G、400G 及以上时，AOC 和 SR 光学器件将占用 ToR、EoR、MoR。资料来源：《LightCounting 超级数据中心光学报告》

未来又怎样？

世界瞬息万变，而且将变得更快！对所有人来说，2021 和 2022 年充满了不可预测性，但面临不可预知挑战时，数据中心已实现了新的扩展和增长，以满足不断增长的连接需求。展望 2023 年及以后，这种增长只会继续。

5G 和人工智能等技术的出现是数据中心发展轨迹上的关键，并将为 800G、1.6T 等方案奠定基础！随着网络不断增强对 5G 和物联网的支持，IT 经理开始关注边缘计算以及不断增长的容量需求。从可卷曲带状光纤到 400G 光收发器，网络提供商一直都在开发面向未来的解决方案，为在每个节点实现端对端无缝连接铺路。

随着行业的持续发展，无论您是专注于边缘计算的企业、超大规模数据中心、多租户提供商，还是系统集成商，各方都会有一席之地。康普一直在关注数据中心的发展趋势，以及不断变化的数据中心领域的前沿。如果您想在准备过渡到更高速度时与我们探讨您的选择，请与我们联系。

COMMScope®

zh.commscope.com

如需了解更多信息，请访问我司网站或联系您的康普销售代表。

© 2022 CommScope, Inc. 保留所有权利。带有™或®标识的所有商标均为在美国的商标或注册商标，并且可能已在其他国家/地区注册。所有产品名称、商标和注册商标均为其各自所有者的财产。本文件仅供规划设计之用，不涉及对任何康普产品或服务相关规格要求或保证的修改或补充。康普致力于最高标准的商业诚信和环境可持续发展，其全球诸多分支机构已获得 ISO 9001、TL 9000、ISO 14001 等国际标准认证。

更多有关康普公司承诺的信息，请访问
www.commscope.com/About-Us/Corporate-Responsibility-and-Sustainability。

EB-115375.1-ZH.CN (08/22)