

硅谷坐标 x TeraHop 于让尘：AI 光互连的超级周期 – 图文解说

一张关键画面对应一段完整解说，按原视频时间推进。

文章摘要

本期围绕《硅谷坐标 x TeraHop 于让尘：AI 光互连的超级周期》展开，按时间顺序讨论AI 的下一堵墙：不只是“算”，而是“连”、从脑灰质到脑白质：互联本身就是算力、Scale Up、Scale Out、Scale Across：从机架到数据中心园区、“Use copper when you can, use optics when you must”的经济账。以下内容保留完整问答、例子、数字、画面信息和限定条件，适合先读这段摘要建立框架，再结合每个关键帧逐段阅读。

01. AI 的下一堵墙：不只是“算”，而是“连”（00:04–01:26）



场景 1

本节主旨 AI 的下一堵墙：不只是“算”，而是“连”

完整内容

曹卿云：过去三年多，AI 算力增长了大约 300 倍，但光连接的总带宽只增长了大约 30 倍。这意味着，今天限制 AI 集群继续做大的，越来越不是单点算力，而是芯片之间的数据搬运效率。换句话说，真正卡住 AI 系统上限的，已经不只是“算”，而是“连”。而要打穿这堵通信墙，最关键的一类技术就是光互联。

今天这期节目，我们请到的嘉宾是全球光互联龙头、硅谷 AI 云计算顶级巨头的核心供应商 TeraHop 营销副总裁于让尘博士。

于让尘：当算力发展了 **300 倍**，光互联只发展了 **30 倍**，中间就有约 **10 倍的 gap**。要填补这个间隙，必须整个产业一起努力。所以，公开场合可能会有很多文章在说这家公司选什么、那家公司选什么，但所有互联网公司都很实际，不会做“宗教式选择”。可插拔模块、NPO、CPO，甚至同一家互联网公司三项都会选。

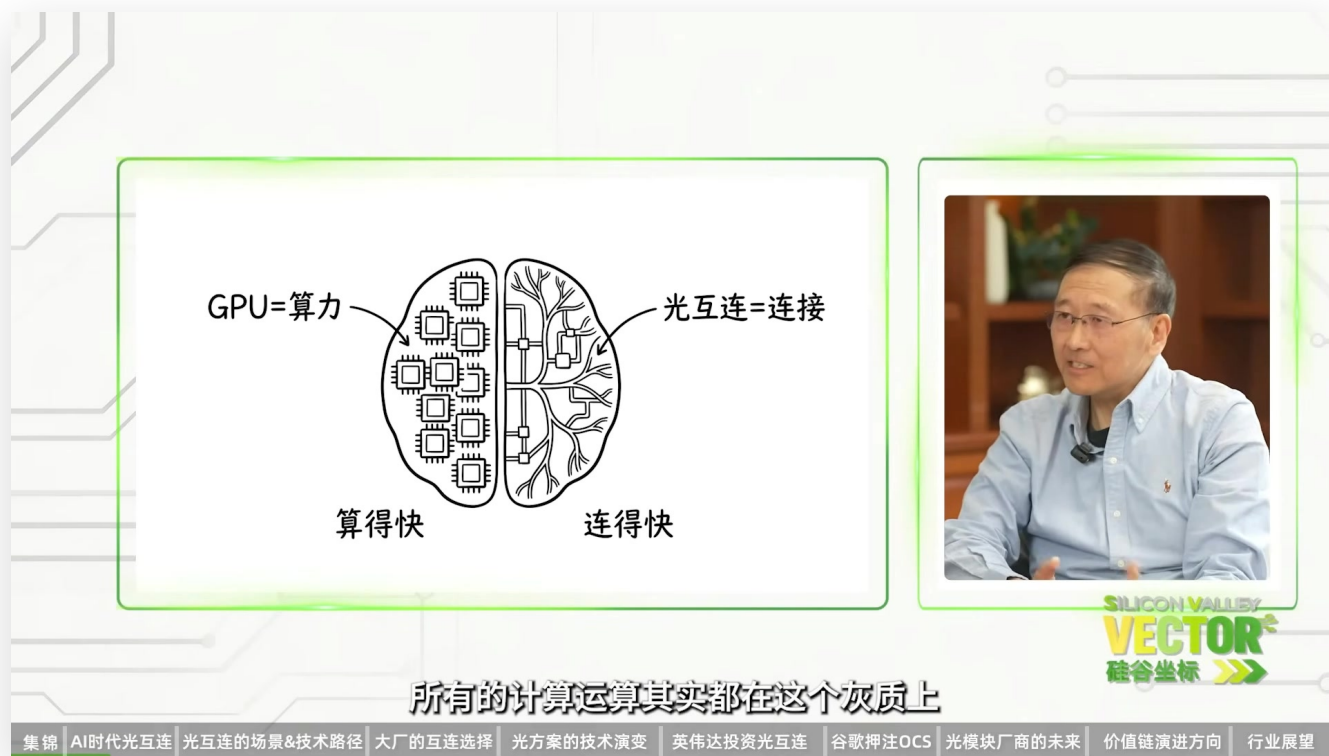
曹卿云：就是“成人不做选择题”。那么，整个产业链中，哪一边的蛋糕在变大，利润又会怎样重新分配？

于让尘：这很快就会变成一个万亿美元级别的问题。

画面说明

画面是数据中心机柜内密集的高速连接，发光端口将“芯片之间的数据搬运”这个抽象瓶颈变成了可见的硬件连接。

02. 从脑灰质到脑白质：互联本身就是算力（01:42–05:43）



场景 2

本节主旨 从脑灰质到脑白质：互联本身就是算力

完整内容

曹卿云：欢迎您来到《硅谷坐标》。您是光互联领域的资深专家，也一直在硅谷一线感受科技大厂对光互联需求的变化。请您先介绍一下，在整个 AI 的广义格局下，光互联为什么重要？

于让尘：光互联已经成为 AI 算力非常重要的组成部分。它不只是一个配件，而是已经成为算力中心不可替代的连接；这种连接本身，也已经变成算力的一部分。当单颗算力芯片无法继续满足需求时，需要靠光互联把它们组成一个更大的“大脑”。这个过程不是一步完成，而是循序渐进的。

从 2022 年到现在，算力发展了大约 300 倍，AI 连接虽然也发展得很快，但数量级上只有大约 30 倍。全行业已经认识到，算力与光互联之间还有约 10 倍的间隙需要填补。

可以用人脑的演化来做比喻。人脑的计算主要发生在表层的脑灰质，而灰质之间的高度连接主要由脑白质承担：一个是算力，一个是连接。在人脑中，灰质与白质的体量已经大致接近 1:1，也有研究给出约 40:60。但在更低等动物中，白质的比例可能只有 **10% 到 20%**，因为它们并不需要处理人类这么复杂的问题。过去几万年里，在大脑总体资源有限的条件下，人脑演化出了大量白质，因为增加连接是更有利的方向。

今天的 AI 集群，可能还处在“小鼠、小猫”的阶段：灰质式的算力比例很大，白质式的连接相对较小。但大脑越做越大，连接的发展速度就应该高于单纯算力的发展速度，总体才会更优。**因此，光互联在整体中的比例会持续增加，而且增速应该大于算力增速。**

从资本开支看，光互联目前只占 AI 数据中心总资本开支的个位数百分比，但增长很快。如果把算力投入视为 10，现在连接可能只在其中的 **10% 到 20%** 左右。今后每增加一份算力投入，连接投入的增长应该超过一份，否则新增算力就可能被浪费。

画面说明

画面将 GPU 算力画成脑灰质，将光互联画成负责连接的脑白质，直观呈现“算力”与“连接”需要共同扩展的比喻。

03. Scale Up、Scale Out、Scale Across：从机架到数据中心园区（05:44–11:10）

The diagram illustrates three scaling strategies in data center architecture:

- Scale Up:** Represented by a 3D cube containing multiple GPU units. Arrows indicate high-bandwidth connections between units within the same rack. Below the diagram is the text: 机柜内/近距离/超高带宽 (In-rack/Short distance/Very high bandwidth).
- Scale Out:** Represented by three server racks connected in a row. Below the diagram is the text: 楼内/机柜间/中距离 (In-building/Between racks/Medium distance).
- Scale Across:** Represented by multiple server racks connected to a central hub, which then connects to other racks. Below the diagram is the text: 园区/数据中心间/更长距离 (Across campus/Data center/Longer distance).

On the right side of the diagram is a video frame showing a man speaking. Below the video frame is the logo for **SILICON VALLEY VECTOR 硅谷坐标**.

Below the diagram is the title: **大语言模型的训练的话**

At the bottom of the slide is a navigation bar with the following items: 集锦 | AI时代光互连 | 光互连的场景&技术路径 | 大厂的互连选择 | 光方案的技术演变 | 英伟达投资光互连 | 谷歌押注OCS | 光模块厂商的未来 | 价值链演进方向 | 行业展望

场景 3

本节主旨 Scale Up、Scale Out、Scale Across：从机架到数据中心园区

完整内容

曹卿云：如果从数据中心架构俯瞰，业界经常提到 Scale Up、Scale Out 和 Scale Across。能不能用普通人能理解的方式，说明它们的区别？

于让尘：过去几年发展得最多的是 Scale Out。它解决的是：单个算力不够时，怎么在不同算力之间建立连接。用大脑做比喻，就是一个人的大脑不够，那就连接更多个大脑。今天的大型训练可能需要万卡级集群，一个万亿参数模型的训练，基本就需要一座大型数据中心，内部要靠多层网络把算力互联起来。

Scale Up 是把“单个大脑”本身做得更大。它对带宽的需求远高于 Scale Out，大致可以高一个数量级。例如英伟达可以先用 NVLink 等技术在一个机架内把 72 颗芯片组成一个大脑。机架内目前主要使用铜缆，但铜缆的带宽与距离有物理极限，已经在限制大脑继续变大。**下一阶段必须用光互联跨越多个机架，把更多芯片连在一起。**英伟达已经开始从 72 颗走向 576 颗，有的互联网公司还在规划 1000 颗；谷歌 TPU 的大脑甚至已经可以连接到约 9000 颗，采用机架内用铜、机架外用 OCS 与光互联的混合方式。

Scale Across 是当一座数据中心的物理大楼都不够时，因为地理、供电和功耗限制，需要把多座数据中心连成一个训练系统。

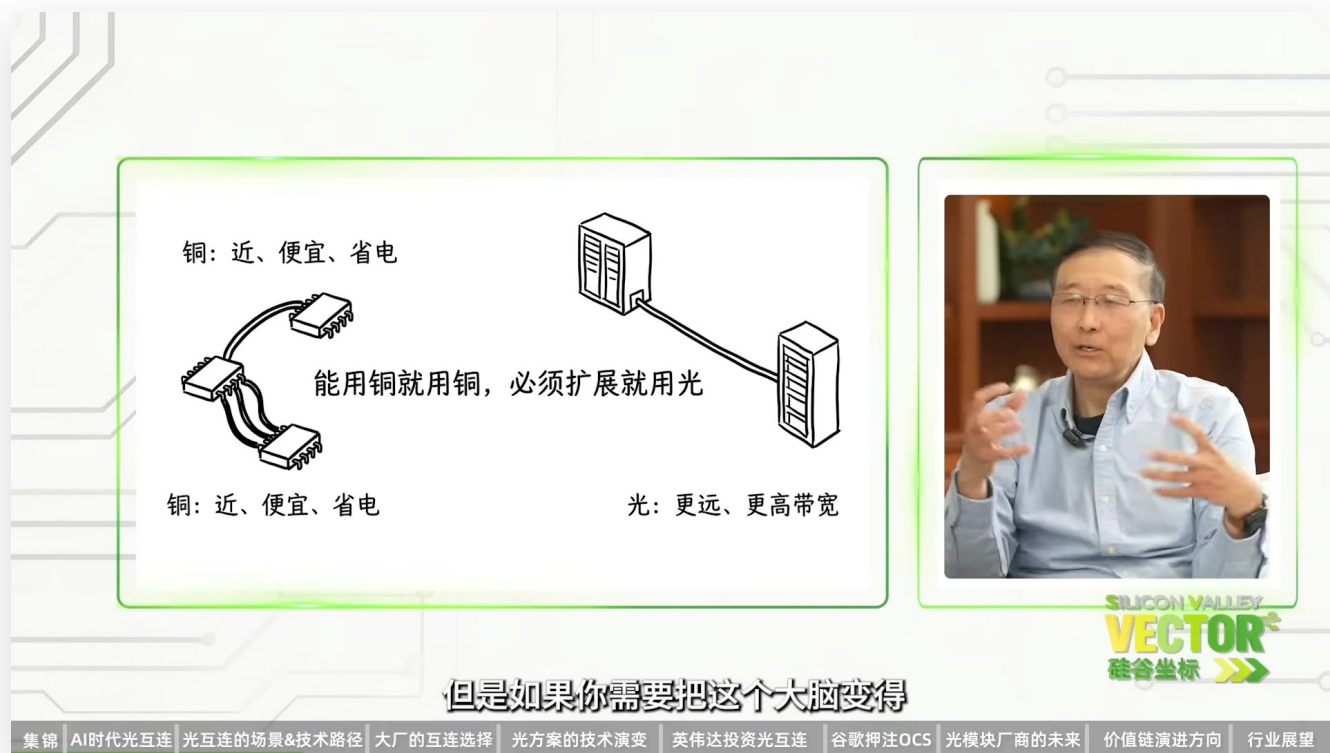
从增长空间看，Scale Up 起点最低，但可能是十倍级机会，因此是潜力最大的节点。从物理距离看，Scale Up 要从机架内的一两米，扩展到八个、十个乃至几十个机架之间的几十米，但它的带宽又最高，因此对方案选择和功耗效率的要求极高。

Scale Out 可能覆盖一座大楼内的一到两公里；带宽相对小一些，但距离长，对光互联又是另一类要求。Scale Across 可以从几公里延伸到几十公里。Meta 近期公布的一些数据中心园区，规模可以大到类似曼哈顿。**从机架、多机架、大楼到整个园区，每一层都需要不同的光互联方案支撑。**

画面说明

图表将 Scale Up、Scale Out 和 Scale Across 并列：分别对应机架内/多机架的算力扩展、数据中心内的集群扩展，以及数据中心之间的跨地域扩展。

04. “Use copper when you can, use optics when you must”的经济账（11:11–14:06）



场景 4

本节主旨 “Use copper when you can, use optics when you must”的经济账

完整内容

曹卿云：您刚才说，Scale Up 对带宽和时延的要求最苛刻，需要让成千上万颗 GPU 像一个整体运行。现在机架内主要使用铜，行业却在说“光进铜退”。普通人会直觉地觉得光的成本比铜高好几倍，这笔经济账应该怎么算？

于让尘：有一句话叫“Use copper when you can, use optics when you must”，能用铜时就用铜，必须用光时再用光。铜互联仍然会在相当长时间内继续存在，但它的距离与带宽局限是物理局限，无法通过意愿改变。机架内能继续用铜，就继续用铜；但当你需要把“大脑”做得更强、更大，而铜已经做不到时，就必须用光。

成本不能只把一个光模块与一根铜缆横向比较，而要从终端客户构建算力的角度，看单位算力输出或单位 Token 输出的总成本。**即使在光互联上多花了钱，只要它让 Token 的单位能耗和单位成本更优，让集群产生更多 Token、更多有效算力，客户就会愿意为它买单。**最终要比的是用户价值、整体收益与总支出。

曹卿云：可以这样理解吗？一块 GPU 可能价值 3 万到 5 万美元，一块光模块可能是几千美元。如果更好的光互联方案能提高成千上万颗 GPU 的集群利用率，那么客户对光互联的 willingness to pay 实际上可以更高。

于让尘：对。如果你在发展算力，但互联没有跟上，就会造成算力闲置或利用率不高。你花了 3 万美元买 GPU，却为了节省互联成本，让这块 GPU 的潜力没有发挥，失去的就是算力和 Token。所以必须把算力和互联当作一个整体来优化。不能只想着在互联上省钱，因为省下的互联成本，可能牺牲的是算力和最终产生 Token 的效率。

画面说明

图表左侧是铜：便宜、低速、近距离；右侧是光：更高带宽、更远距离。**中间文字强调“能用铜就用铜，必须用光就用光”。**

05. Pluggable、NPO、CPO：光接口应该离主芯片多远？ (14:08–16:45)

Pluggable vs NPO vs CPO

可插拔	NPO	CPO
主芯片 光模块 PCB 较远距离	主芯片 光引擎 PCB 基板 更近距离	主芯片 封装包 PCB 距离最近
✓ 最灵活 ✓ 易维护 ✗ 功耗较高	✓ 更近一步 ✗ 折中方案 ✗ 过渡路线	✓ 距离最近 ✓ 带宽密度高 ✗ 散热维修难

本质问题：光接口到底离主芯片有多近

这就会影响后面在不同的应用场景

SILICON VALLEY
VECTOR
硅谷坐标

集锦 | AI时代光互连 | 光互连的场景&技术路径 | 大厂的互连选择 | 光方案的技术演变 | 英伟达投资光互连 | 谷歌押注OCS | 光模块厂商的未来 | 价值链演进方向 | 行业展望

场景 5

本节主旨 Pluggable、NPO、CPO：光接口应该离主芯片多远？

完整内容

曹卿云：目前硅谷科技大厂在光互联上走到了技术路径的十字路口：可插拔、NPO 和 CPO。它们本质上都在回答同一个问题：光接口到底应该离主芯片多远？

可插拔方案是外挂式、可以插拔和更换的独立部件。它离主芯片最远，中间的电信号传输距离最长，功耗也相对较高；但它很灵活，坏了可以直接拔下来换新，不会影响主芯片，供应链也最开放。

CPO (Co-Packaged Optics) 的思路相反：尽量把光与主芯片做成深度集成的整体。好处是电路径被压到最短，带宽密度最高、延迟最小；代价是散热更集中、系统耦合更深，某个部件损坏可能影响整块板卡甚至整个系统。

NPO (Near-Packaged Optics) 介于可插拔与 CPO 之间。它比可插拔更靠近芯片，但没有激进到与主芯片封装在一起，相当于在性能和可维护性之间找平衡。

于让尘：这不是一道单纯的选择题，而是平衡题。一个算力中心要同时优化功耗、成本、带宽和密度，还要加上可维护性与可靠性。所有维度对客户都重要，但在具体设计中必须确定优先级。例如设计算力卡时优先什么，与构建交换机时优先什么，得出的方案可能不同。所以最难的不是“选哪个名字”，而是怎样在具体场景中把这些维度做到总体最优。

画面说明

对比图将 Pluggable、NPO 和 CPO 并列：光学部件逐步靠近 ASIC，性能和密度随之提高，但可更换性、维护难度与系统耦合也随之改变。

06. 大厂不做“宗教式选择”：同一家公司也会多路线并行 (16:45–20:53)

大厂不是选最炫的，而是选最合适的组合

而且不局限于只选一种

集锦 | AI时代光互连 | 光互连的场景&技术路径 | 大厂的互连选择 | 光方案的技术演变 | 英伟达投资光互连 | 谷歌押注OCS | 光模块厂商的未来 | 价值链演进方向 | 行业展望

场景 6

本节主旨 大厂不做“宗教式选择”：同一家公司也会多路线并行

完整内容

曹卿云：硅谷几家科技大厂是怎样选择可插拔、NPO 和 CPO 的？背后有什么样的考量？

于让尘：对互联网大厂来说，它们一方面买英伟达的算力和垂直整合系统，另一方面也在自建芯片，基本没有例外。谷歌有 TPU，Meta 有 MTIA，亚马逊有 Trainium，微软有 Maia，OpenAI 在做自己的芯片，xAI 也在做。这些公司都是两条腿甚至多条腿走路：同时采购外部系统，并构建自己的算力大脑。

在这种格局下，客户一定欢迎开放的光互联生态。大家也高度认可硅光的方向，因为算力持续增长，不可能只靠传统分立器件的组装，必须走向半导体化、高度集成。但不同公司的优先级会不一样：有的公司希望延续 Pluggable，会更喜欢 12.8T XPO；有的公司更强调高密度算力卡与交换机互联，会更喜欢 NPO、开放式 CPO，或者能够把 CPO、NPO 都接入的开放 Socket 方案。

曹卿云：外界常说谷歌继续偏向可插拔，Meta 和亚马逊倾向 NPO，英伟达希望走向 CPO。这些不同路线，是否来自各自公司的不同条件？

于让尘：只能根据公开信息来说。谷歌在 OCP 等场合公开表达过对可插拔的偏好，还有一句话是“CPO 总是 two years later，而且每年都是 two years later”。但这不等于它们私下不会评估其他工具，也不等于同一家公司在不同算力和不同 AI 组网中只有一个选项。

互联网公司不会做宗教式选择。所谓宗教式选择，就是不管代价如何，都坚持只用一种方案，这很不合理。更合理的做法是：不同架构有不同偏好和优先级，在特定的排列组合下选最合适的方案。即使在同一家大型互联网公司内部，也会有多种架构，并在不同场景中同时选用可插拔、NPO 和 CPO。外界看到的常常只是表面现象；更核心的原则是，为自己的算力中心选择总体最优的组合，不局限于一种。“成人不做选择题”。

画面说明

图中是互联网大厂的复杂算力网络示意，没有只通往单一技术的路径；画面下方字幕也点出“宗教式选择”并不符合大厂实际决策。

07. 光互联如何演进：更高速、更多 Lane，也可能“Slow and Wide”（20:55–26:01）

传统方案
大而笨重
单通道大载量
速度受限

Slow and Wide
慢但是宽
多通道小载量
总吞吐量更高
并行高效

是有可能在这条路上也会继续发展

SILICON VALLEY
VECTOR
硅谷坐标

集锦 | AI时代光互联 | 光互联的场景&技术路径 | 大厂的互连选择 | 光方案的技术演变 | 英伟达投资光互联 | 谷歌押注OCS | 光模块厂商的未来 | 价值链演进方向 | 行业展望

场景 7

本节主旨 光互联如何演进：更高速、更多 Lane，也可能“Slow and Wide”

完整内容

曹卿云：光模块或光互联的总带宽，可以粗略理解为两件事相乘：第一，每条“车道”有多快，也就是每通道速率；第二，一共有多少条“车道”，也就是通道数量。带宽可以通过增加通道数或提高单通道速率来提升。在新增的十倍市场中，未来的近期、中期和远期技术节奏会怎样演变？

于让尘：这个行业是多个维度同时发展。Scale Out 和 Scale Across 会继续增长，同时 Scale Up 是目前行业最大的增长机会，因为它本身就有十倍带宽的扩张空间。近期最重要的事，是帮助客户在 Scale Up 里构建更大的大脑，让成百上千颗 GPU、XPU 组成更大的 Scale Up domain。为此，行业已经推出 12.8T XPO，也在推进 6.4T 级 NPO、CPO，给客户更多选项。

第一个增长维度是单通道速率。200G per lane 和 1.6T 开始量产后，未来两三年会成为重要增长点。行业已经在讨论 400G per lane，这不只是概念，而是可见的下一步。

第二个维度是通道数。传统可能是八条 Lane，现在可以一次跳到 32 或 64 条，以后还可能向 128 条发展。速率与通道数都可以增长，而且这条路不会在 32 或 64 条停下。

铜的物理极限终究会出现。现在铜在某些速率下可以走一米半到两米，到 400G 时可能连一米都很勉强，以后可能只剩半米。当铜实在走不下去，整个系统就会向全光发展，这才是“光进铜退”最终的含义。

一旦把铜缆连接完全拿掉，接口的底层逻辑也可能重构。这时，速率不一定只能往高走，也可以往更低的单通道速率走，用大量通道来换总带宽，同时简化接口、降低功耗。这条探索路线叫“Slow and Wide”：每辆车载量小一些，但同时有非常多辆小车并行，总效率可能高于一辆很大但很笨重的车。这在物理上是可实现的，已经有早期尝试。

另一条路是 Coherent Light，也就是相干光与更高阶的调制技术。它不一定立刻进入 Scale Up，但在距离更长的 Scale Out 和 Scale Across 中，会越来越需要。所以技术演进不是单线的：可以提高单通道速率，可以增加通道数，可以 Slow and Wide，也可以在长距离上引入相干技术。

画面说明

图表对比传统“大车、少车道”与 Slow and Wide 的“小车、多车道”，说明总带宽可以由更多低速 Lane 并行实现，不必只靠单 Lane 不断提速。

08. 12.8T 为什么会突然跳出来？（26:01-27:22）



场景 8

本节主旨 12.8T 为什么会突然跳出来？

完整内容

曹卿云：过去常听到 400G、800G、1.6T、3.2T，现在 TeraHop 突然提出 12.8T，对外行人来说是一个很大的飞跃。为什么带宽能一次跳高这么多？

于让尘：要从需求和可行性两个维度回答，缺一不可。只有需求、技术做不到，没有用；只有技术可行、没有需求，也不会发生。

需求端最重要的瓶颈，就是算力增长 300 倍、光互联只增长 30 倍造成的十倍 gap。同时，Scale Out 正在向 Scale Up 扩展，机架内的铜缆已经遇到瓶颈，必须有新的光解决方案。12.8T 要解决的正是这个需求。

可行性来自硅光。对硅光技术来说，只要制程开发完成，良率和可靠性发展到一定程度，向更多通道扩展就是必然方向。行业已经生产过大量八通道产品，下一步自然可以走向 16、32 乃至 64 通道。**从底层技术看，成熟度已经能支撑 12.8T 方案。**当需求与可行性同时到位，这个跃迁就会必然发生。

画面说明

于让尘在画面中直接说明 12.8T solution 的可行性；本段没有更完整的数据图，因此保留嘉宾论证时的近景。

09. 英伟达 40 亿美元投资：用真金白银锁定光互联供应链 (27:22-30:30)



场景 9

本节主旨 英伟达 40 亿美元投资：用真金白银锁定光互联供应链

完整内容

曹卿云：今年 3 月初，英伟达向 Lumentum 和 Coherent 合计投资了 40 亿美元，并签署战略合作协议。这释放了什么样的产业信号？

于让尘：这是一个公开证明，证明英伟达对光互联的重视程度。它担心供应链能不能跟上增长步伐，所以用真金白银来锁定供应。这个具体交易主要是锁定激光器供应。黄仁勋已经在两次 GTC 上强调，光互联是整个算力构建非常重要的组成部分。下一步的 CPO 等方案需要大量激光器，所以英伟达需要提前保障供应。放到更大的产业格局里，这只是一个插曲，但是非常明确的公开信号：构建 AI 大脑时，光互联已经是核心环节。

曹卿云：这与芯片供应链的逻辑是否类似？英伟达芯片很强，但客户不希望被一家锁死，因此一边用英伟达，一边从自研芯片、供应链和战略上做多元化。

于让尘：这些是相辅相成的。终端客户不愿意被一家锁死，它们已经在用钱投票（vote with the wallet），而且是花大钱自建算力芯片。它们认为，今后构建算力本身会成为核心竞争力，不只提供服务，还要进入基础设施业务。所以做自己的芯片已经不是选项，而是必须做的事情。

与自研芯片配套的连接方案也必须建起来。互联已经不是配件，而是与算力一起优化的必选项。客户需要解决方案的多样性、多元协议和充足供应，也需要供应商既满足特定需求，又保留整个 AI 生态的丰富性。一家供应商选中一种方案，然后强推给所有客户，可能在特定时间和条件下推得动，但很难走向大规模商用。只有提供足够多的选项和排列组合，让客户自己优化，生态才会健康。

画面说明

画面上方的信息条直接标出，英伟达投资美国 Lumentum 等光互联企业，作为供应链锁定的产业证据。

10. OFC 为什么突然出现这么多开放标准？（30:31–34:12）

开放标准 vs 封闭生态

开放标准

多源
可替代
不易被锁死
生态更开放

封闭生态

垂直整合
效率高
绑定强
议价强

客户追求高密度效率，也不想被单一供应商锁死



SILICON VALLEY
VECTOR
硅谷坐标

集锦 | AI时代光互连 | 光互连的场景&技术路径 | 大厂的互连选择 | 光方案的技术演变 | 英伟达投资光互连 | 谷歌押注OCS | 光模块厂商的未来 | 价值链演进方向 | 行业展望

场景 10

本节主旨 OFC 为什么突然出现这么多开放标准？

完整内容

曹卿云：今年 3 月，与 GTC 几乎同期举行的 OFC 是光互联行业的重要大会。与去年相比，今年有很多开放标准组织和多元协议出现。为什么在这个时候，会有这么多不同声音一起维护开放生态？

于让尘：今年与去年相比，有多项重大的多元协议和标准推出。TeraHop 也是这一轮新标准和新协议的主要推动者之一。我们发布了非常重要的 12.8T XPO 多元协议，也推动了 OpenCTX 等开放方案。这些新协议在带宽上可以带来四到八倍的增长，现在已经有超过百家公司愿意参与，包括互联网公司、芯片公司和模块公司，很快就会形成大家乐见其成的生态。特别是终端客户，对共建下一代 AI 数据中心硬件非常有兴趣。

市场上现在很热闹的一些方案，往往是少数大客户与大供应商采用的 proprietary、垂直整合方案。这种形态可以有一定市场，但终端客户未必喜欢，因为它会变成一个垂直整合的打包销售，供应商的溢价和议价能力变得很强，生态不一定健康。这也是 CPO 谈了很久、却一直没有大规模商用的瓶颈之一：客户不一定喜欢被绑定，同时垂直整合的技术难度也非常高。

对于密度要求更高的网络和算力卡，也可以走开放的 CPO/NPO 与多芯光纤路线。这样既能在板上或基板上提供小尺寸、高密度连接，又希望继续保留可维护、可替换的特点。

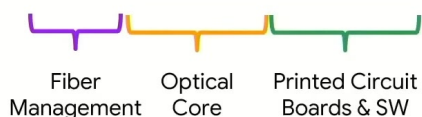
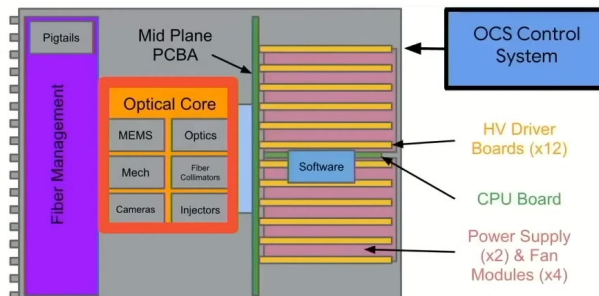
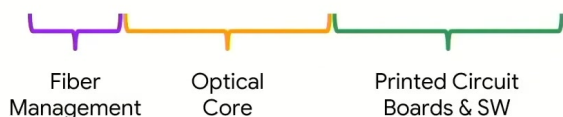
为什么连光纤本身也需要新的多元协议？因为当 AI 大脑容量持续增加，从一个机架里引出大量光纤，再铺到整个数据中心，光纤的物理密度也会变成瓶颈。即使单根光纤只有头发丝粗，数量太大时也很难布线。多芯光纤协议希望让一根光纤跑多个通道，并让不同供应商都能兼容，从而给终端客户更多选项和排列组合，去实现十倍增长。

画面说明

图表对比“开放标准”与“封闭生态”：开放标准有多家供应商、更多组合与更高可替换性；封闭生态则由单一供应链垂直整合。

11. OCS 如何把 9000 颗 TPU 连成一个“超级大脑”（34:12–37:41）

Google OCS: Palomar 136x136



它不需要像传统电交换机那样

集锦 | AI时代光互连 | 光互连的场景&技术路径 | 大厂的互连选择 | 光方案的技术演变 | 英伟达投资光互连 | 谷歌押注OCS | 光模块厂商的未来 | 价值链演进方向 | 行业展望

场景 11

本节主旨 OCS 如何把 9000 颗 TPU 连成一个“超级大脑”

完整内容

曹卿云：OCS 的全称是 Optical Circuit Switch，中文叫光交换机。传统电交换机需要在中间站点处理、判断和转发数据包；OCS 则直接建立端到端的光路径，让数据以光的形式沿专用路径通过。因此，它很适合 AI 训练这种大吞吐、数据流相对可预测的场景。

刚才我们一直在谈单个光模块的带宽，但谷歌把问题提升到了网络重组的层面：不只看单点带宽，还要看整个 AI 网络怎样组织。谷歌已经在 TPU 系统中大规模使用 OCS。您怎么看 OCS 对产业的意义？

于让尘：OCS 是当前的热点之一。它首先在不同网络层级上解决算力连接问题。对谷歌 TPU 方案来说，OCS 是构建 Scale Up 的重要组成部分。根据谷歌的公开信息，它可以用约 9000 颗 TPU 再加上 OCS，组成一个 super brain，也就是它的最强大脑。

用 OCS 部分取代电交换，有几个明显好处。第一是节省功耗；第二是降低时延，因为数据进入 OCS 后，可以直接转到另一条光通道，不需要像电交换那样再处理数据包；第三是有可能降低总成本。

但 OCS 不是没有代价，也还不能完全取代电交换。OCS 本身的路径重配速度比较慢，它更像一个重新规划数据路径的工具。当工作流相对稳定、可预测时，它很好用。另一个适用场景是容错：如果某条路径出现可靠性问题，可以通过重新调度数据流，快速绕开故障。

各大互联网公司都在评估怎样把 OCS 用到更多地方，但谷歌走在最前面，也走得最远。它大约十年前就开始布局，这不只是硬件，还包括大量软件和调度算法的开发。其他互联网公司也需要开发自己的算法，决定什么时候怎样调度数据。这个过程需要时间，但当整个行业向全光互联发展时，OCS 会是其中非常重要的组成部分。

画面说明

画面是 Google Palomar 136×136 OCS 系统的实物与内部结构示意，标出光纤管理、光学核心、电路板与 OCS 控制系统。

12. OCS 会取代电交换吗？更现实的未来是长期共存（37:42–40:54）



场景 12

本节主旨 OCS 会取代电交换吗？更现实的未来是长期共存

完整内容

曹卿云：未来除谷歌以外的玩家，是不是也都会选择 OCS，让它成为一个产业趋势？

于让尘：谷歌已经开始混合部署，它并没有不用电交换，电交换仍然是大规模使用的。光交换与电交换是互补的，可以做混合部署；只是谷歌的 OCS 比例已经开始上升，而其他互联网公司的比例可能仍然非常低，甚至接近于零。随着这些公司开始布署，OCS 的占比会增加。对供应商来说，这意味着大量 greenfield 机会，因为很多地方现在还没有 OCS。

曹卿云：终极情况下，电交换机会不会完全消失？

于让尘：我相对保守。我不认为电交换会完全消失，因为它本身非常快，可以在纳秒级重新转换和导流数据。OCS 受到物理限制，在切换速度上还有非常长的路要走。未来五年，完全取代基本不可能。十年、二十年后会不会发生，可以说 never say never，但我仍然认为比例会越来越高，却不会完全取代。电交换芯片厂商不需要担心生意完全消失，更可能是长期共存。

曹卿云：OCS 除了切换速度较慢，还有哪些局限？

于让尘：做光交换必然会带来光损耗。损耗一增加，对光模块和整条链路预算的要求就会更高。不管前后采用 CPO、NPO 还是 Pluggable，都必须把 OCS 带来的额外损耗算进链路预算。不同形态会因此发生取舍，因为链路里可能多出两到三 dB 的损耗。

在某些场景下，Coherent Light 可能成为重要选项。它可能成本更高、功耗更高、技术更复杂，但优势是接收和数字处理能力更强，链路预算更高。因此，相干光与 OCS 可能是一个很好的配对，特别是当速率再向下一代、下两代提升时，这个趋势会越来越明显。

画面说明

画面中的**于让尘**正在说明 OCS 与电交换的共存关系；本段候选画面中没有完整损耗图，因此保留主要发言人的对谈镜头。

13. 从 Celero 到供应商布局：相干光的难点不只是一颗芯片 (40:55-43:41)



场景 13

本节主旨 从 Celero 到供应商布局：相干光的难点不只是一颗芯片

完整内容

曹卿云：谷歌的战略投资部门今年投资了一家相干光初创公司 Celero。从供应链实现的角度看，未来推进还会遇到哪些困难，节奏会如何变化？

于让尘：这会推动几个维度。第一是 DSP，也就是数字信号处理，这非常重要。但不能直接把常规长距离相干传输的 DSP 搬过来，因为 OCS 场景要求低功耗和低时延，必须针对这个应用重新设计。谷歌战投 Celero，有一个考量就是让供应链更丰富，并确保自己获得足够供应。

第二是底层光学与硅光。相干光需要更复杂的调制、偏振处理和平衡接收机，对硅光平台的要求比常规方案高很多。所以，这不是投一家 DSP 公司就能解决的问题，而是整条供应链上多个环节要同时成熟。

曹卿云：现在有 Pluggable XPO、NPO、CPO 等多种形态。对光连接解决方案供应商来说，未来会怎样布局？

于让尘：不管是什么 form factor，都只是一个载体。作为主流的头部分解决方案供应商，我们不替客户排除选项，而是尽量把各种 solution 都提供给客户，由客户根据自己的系统构建做选择。

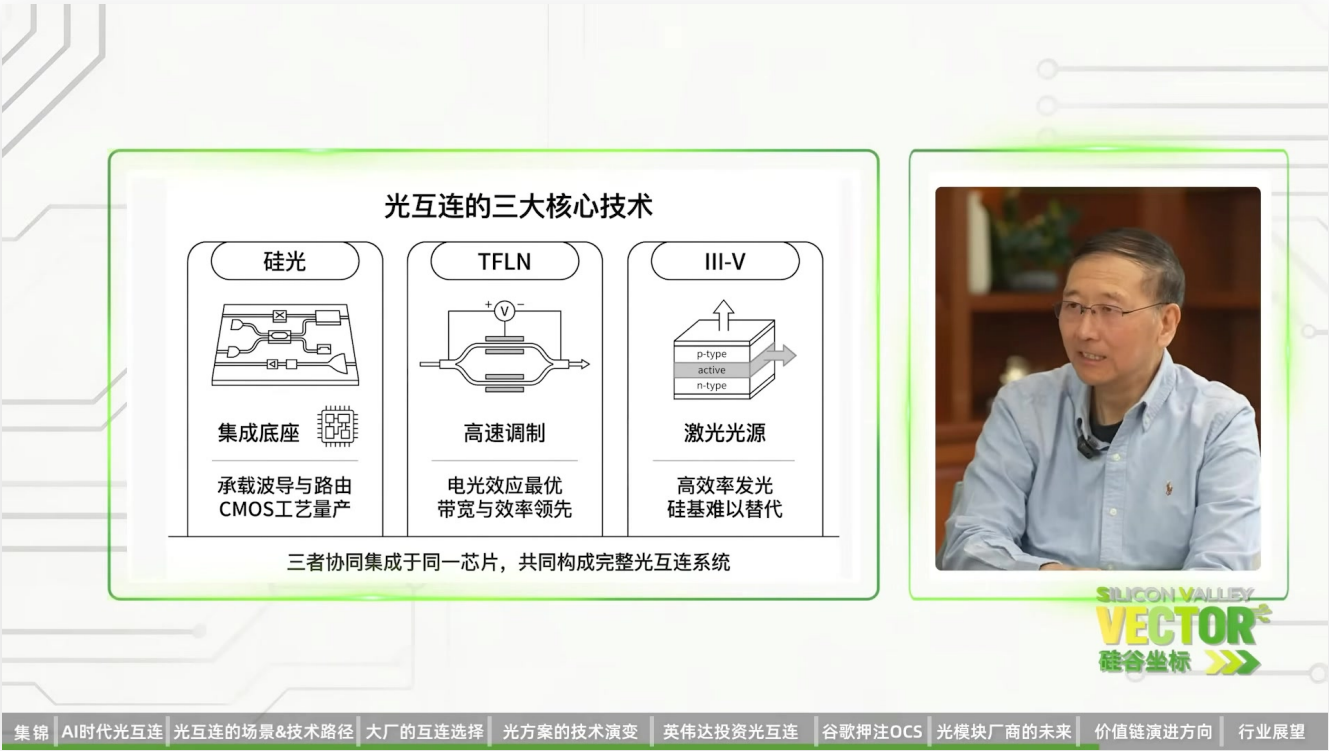
曹卿云：TeraHop 手上也有很多现金，未来会怎么花？会不会向上下游发展？

于让尘：作为上市公司，具体资本规划不便在公开场合讨论。可以确认的是行业大趋势：光互联会成为 AI 算力构建的重要部分。过去光互联可能只是相对较小的比例，以后要与整个“大脑”共生增长，而且光互联的比重增长会更快。在这个过程中，有多个可发展方向，我们会同时推进。对于上游、下游、同行和客户，也不应做排他式选择，而应该共同发展与合作。

画面说明

画面信息条标出“光模块厂商的未来布局”，于让尘说明头部供应商不会只绑定一种 form factor。

14. 400G per lane 的三马竞跑：硅光、薄膜铌酸锂与 EML (43:42-45:10)



场景 14

本节主旨 400G per lane 的三马竞跑：硅光、薄膜铌酸锂与 EML

完整内容

曹卿云：您刚才说，未来每通道带宽会提升到 400G，这也是今年 OIF 的热门话题。现在市场上有三类主要技术：硅光、薄膜铌酸锂和 EML 激光器。怎样评价它们各自的长处与短板？

于让尘：现在还不到说谁是主流的阶段，只能说大家都在 PK。今年 OFC 非常热闹，不同厂商都进入了早期验证和演示阶段。硅光有一些已公开的方案，薄膜铌酸锂也有演示，三五族和 EML 路线同样有新发布。现在的关键不是立刻宣布胜者，而是三条路线都在快速发展。

对我们来说，可集成的方向更值得长期关注。而且不一定只有纯粹的单一技术，也可以做混合集成，例如把薄膜铌酸锂嫁接到硅光平台上，组成混合集成平台。这在我们看来是很有前景的方向。

但作为解决方案提供商，还是那句话：成年人不做选择。**EML/三五族、薄膜铌酸锂和硅光都不放弃，让三匹马同时跑，再由实际的性能、功耗、良率、可集成性和客户需求来决定不同场景下用什么。**

画面说明

对比图并列三类光源/调制路线：**硅光、TFLN（薄膜铌酸锂）与 III-V/EML**，图下列出了它们在集成、器件和制造方面的差异。

15. 万亿美元连接市场：形态会变，半导体化的方向不会变 (45:11-46:46)



场景 15

本节主旨 万亿美元连接市场：形态会变，半导体化的方向不会变

完整内容

曹卿云：从 Pluggable 走向 NPO、CPO，整个产业价值链会怎样重新改变？哪些环节的蛋糕在变大，利润又会怎样重新分配？

于让尘：这很快就会变成一个万亿美元级别的问题。整个“连接”市场，不论是做交换机还是做光互连，总体市场价值已经接近 one trillion dollars。现在各家都在开发自己的方向，有人做 CPO，有人做 NPO，有人继续做 Pluggable。我个人认为，产业会以混合方式向前走。

但更重要的底层逻辑是，大家都会朝半导体集成方向发展，光与电会越来越深地集成。“半导体化的光互联”是确定性很高的方向。但到底以哪种形态落地，不同环节最后分到多少价值，可能需要在行业继续向前走的过程中边走边看。

对光解决方案供应商来说，不能把自己只定义成“可插拔模块供应商”。我们是光互联供应商，并不锁定只做某一种 form factor。在技术演进中，我们乐见各种形态往前发展，并用底层的光半导体和电半导体技术，同时支撑 Pluggable、NPO 与 CPO。最后能成为大玩家的公司，应该都会以这种方式参与。

画面说明

画面是主持人与嘉宾的双人镜头，字幕列出 CPO、NPO 与 Pluggable，对应“形态多样，底层技术共通”的讨论。

16. 光互联的真正瓶颈：激光器还在小晶圆时代（46:47-48:47）



场景 16

本节主旨 光互联的真正瓶颈：激光器还在小晶圆时代

完整内容

曹卿云：您目前看到的整个产业链，最大的瓶颈在哪里？

于让尘：最大瓶颈是，作为一个产业，我们还处在半导体集成的早期阶段。以激光器为例，它现在仍然是一个相对小众的生产生态。过去两年，AI 需求逼着大家快速扩张，但整个生态还没有完全发展到位。

今天主流电子半导体已经在 12 英寸、300 毫米晶圆上大规模制造。但如果问一家激光器厂商的晶圆尺寸，它可能还只有 3 英寸，现在正在努力走向 6 英寸。这与成熟半导体生态相比，仍然落后很远。未来几年，光源如何扩大晶圆、提升产能、良率和可靠性，是很重要的问题。

现在也有初创公司在做量子点激光器。如果量子点激光器能够直接在硅线上制造，就可能完成一次跳跃：从现在的 3 英寸、4 英寸，直接跳到 12 英寸的量子点光源制造。如果做到，光源供应瓶颈就会得到大幅突破。但这仍然是对未来的展望，不是已经完成的量产事实。

另一个重要方向是光电集成，特别是封装。光芯片与电芯片会出现更深的共生发展，从 2D、2.5D 到 3D 的多层次集成都在开发。这条路径与先进半导体封装生态非常接近。例如 TSMC 在推进 COUPE（Compact Universal Photonic Engine）等集成平台，把硅光芯片、电控芯片与主 ASIC 向更紧密的共同封装推进。其他晶圆厂也会沿这个方向努力，让光与电进入更深度的集成。

画面说明

画面中字幕提到 3D、2D 集成是各大厂正在开发的重点，对应本段从小晶圆光源生态走向先进光电封装的讨论。

17. 五到十年展望：不要让算力与互联争夺资源（48:47-49:46）



场景 17

本节主旨 五到十年展望：不要让算力与互联争夺资源

完整内容

曹卿云：如果把时间拉长到五年、十年，您觉得光互联行业会怎样发展？

于让尘：技术层面的方向刚才已经谈了很多。如果从算力支出的角度看，过去很多人会说，光互联应该尽量省钱，把省下来的钱放到算力上。我认为这个概念需要纠正。

大家应该把光互联看作算力的一部分，最后追求的是整个算力系统的最佳平衡。不应该把算力和互联看作两块互相争夺资源的支出，而应该把它们看作一个共生、共同成长、共同优化的生态。

可以预期，光互联在 AI 算力构建中所占的比重会越来越大。它会更像现代人脑中的脑白质与脑灰质：计算能力与连接能力缺一不可，而且随着大脑变得更大更复杂，连接所占的份量还有很大成长空间。未来的光互联不会只是一个附属部件，而会成为算力中非常不可分割的一部分。

画面说明

画面信息条标出“光互联未来 5—10 年展望”，双人对谈镜头保留了节目最后的长期结论。