

AI시대의 통신 기술

여인동, Ciena SE
2024. 6

목차

1. 통신기술의 발전
2. AI시대의 통신 기술
3. What's next?

통신기술의 발전

통신기술 발전 방향

전신, 전화등의 매체를 사용하여 정보나 의사를 전달하는 것.

from Wiki

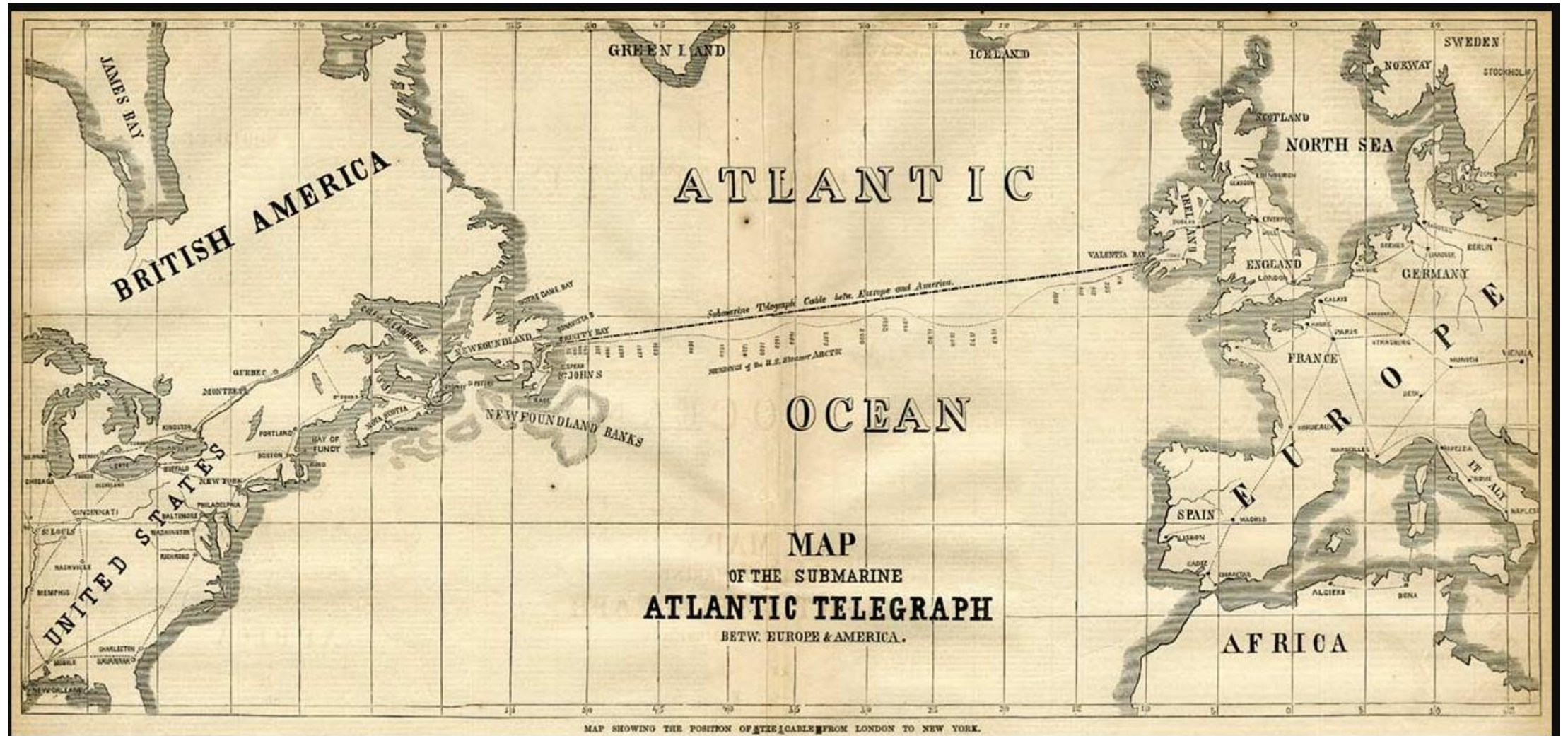


더 많이

더 멀리

신뢰성

최초의 해저케이블 통신 - 1851년 대서양 횡단





어떻게 더 효율적으로 목소리를 전달할까? - 압축/신장(양자화, Quantizing)

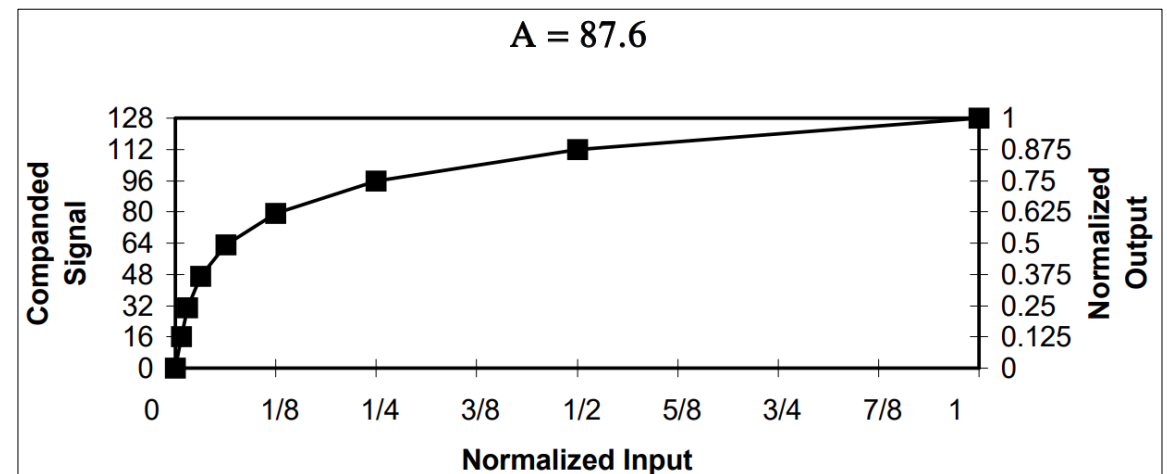
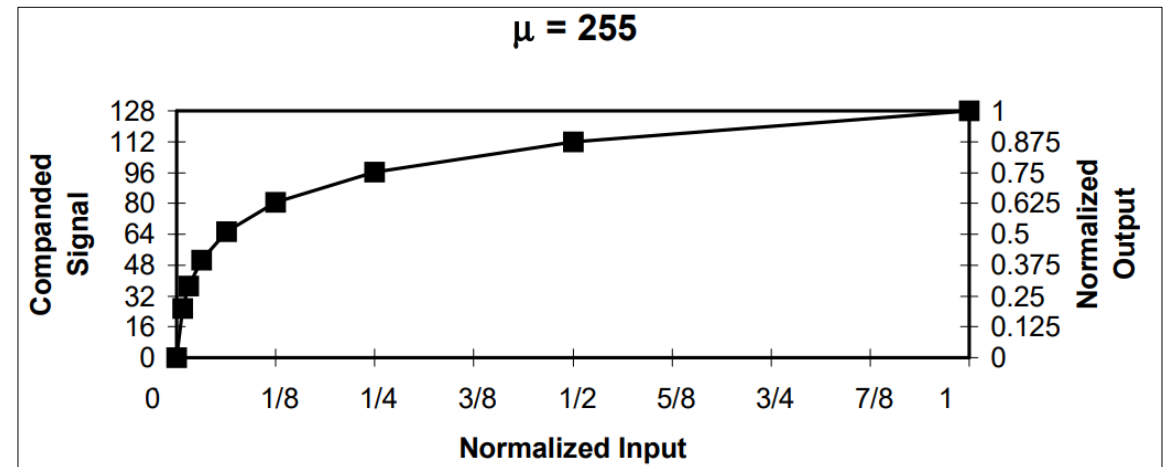
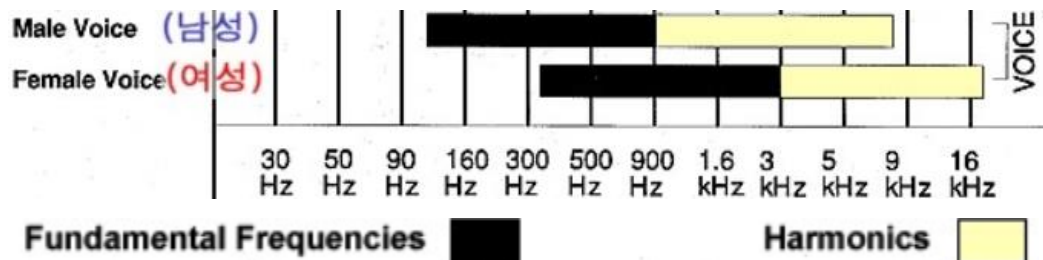
G.711 Companding algorithms

μ -law(used in T-Carrier)

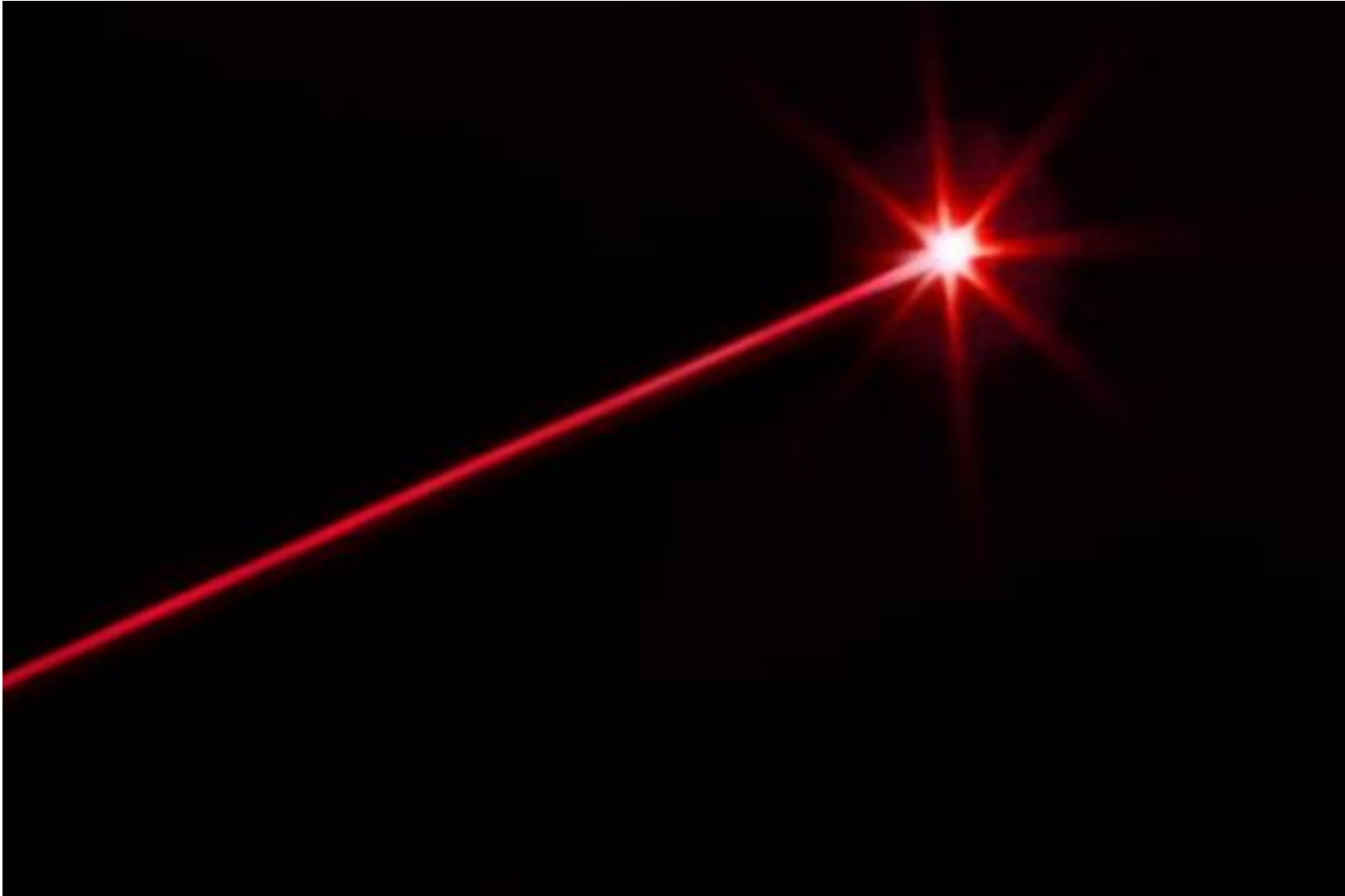
$$F(x) = \frac{\text{sgn}(x) * \ln(1 + \mu|x|)}{\ln(1 + \mu)} \quad 0 \leq |x| \leq 1$$

$$F(x) = \begin{cases} \frac{A * |x|}{1 + \ln(A)} & 0 \leq |x| < \frac{1}{A} \\ \frac{\text{sgn}(x) * (1 + \ln(A|x|))}{1 + \ln(A)} & \frac{1}{A} \leq |x| \leq 1 \end{cases}$$

A-Law Equation

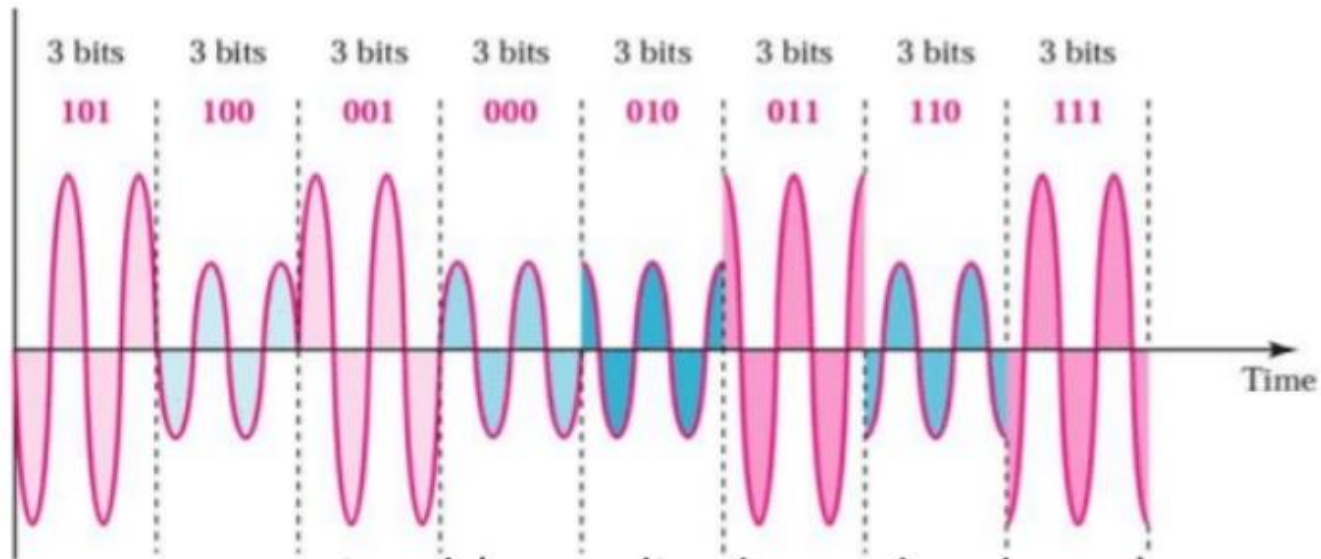
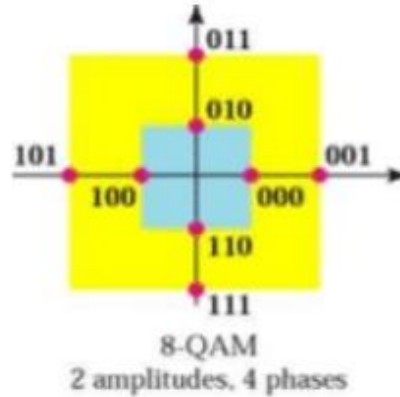


가장 효율적인 전송 방식 - 빛



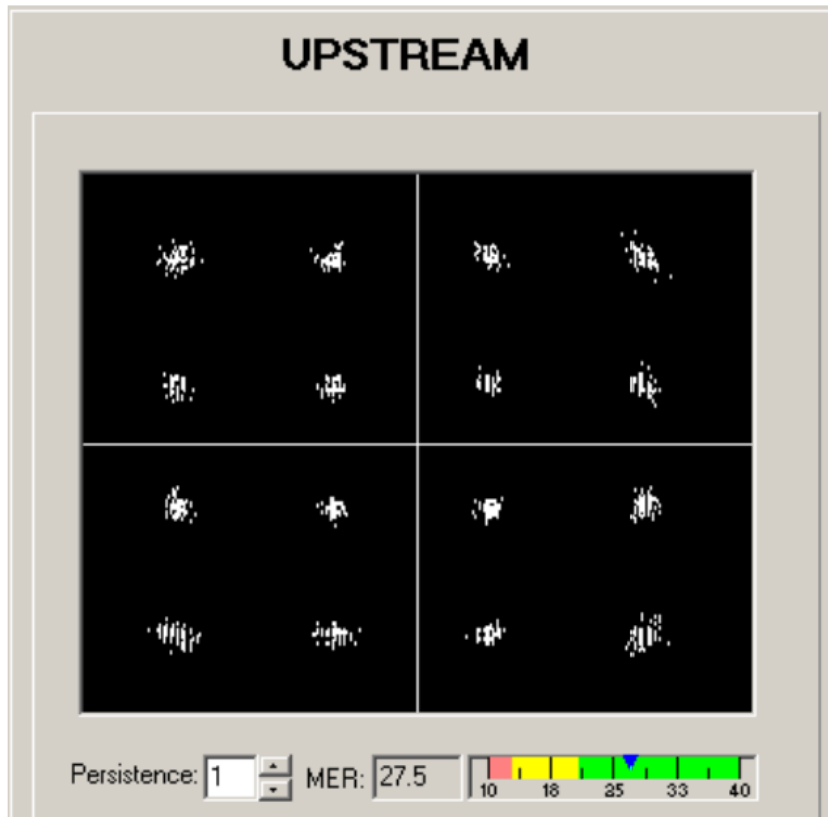
어떻게 효율을 높일까?

ASK/FSK/PSK -> QAM -> PCS

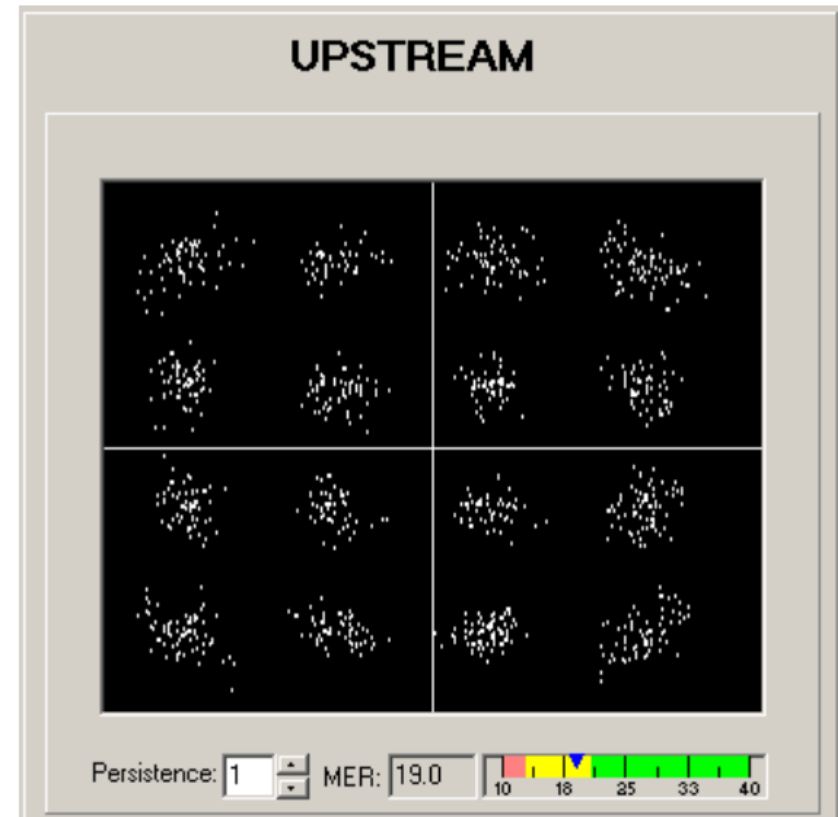


어떻게 효율을 높일까? ASK/FSK/PSK -> QAM -> PCS

Ideal



Noise



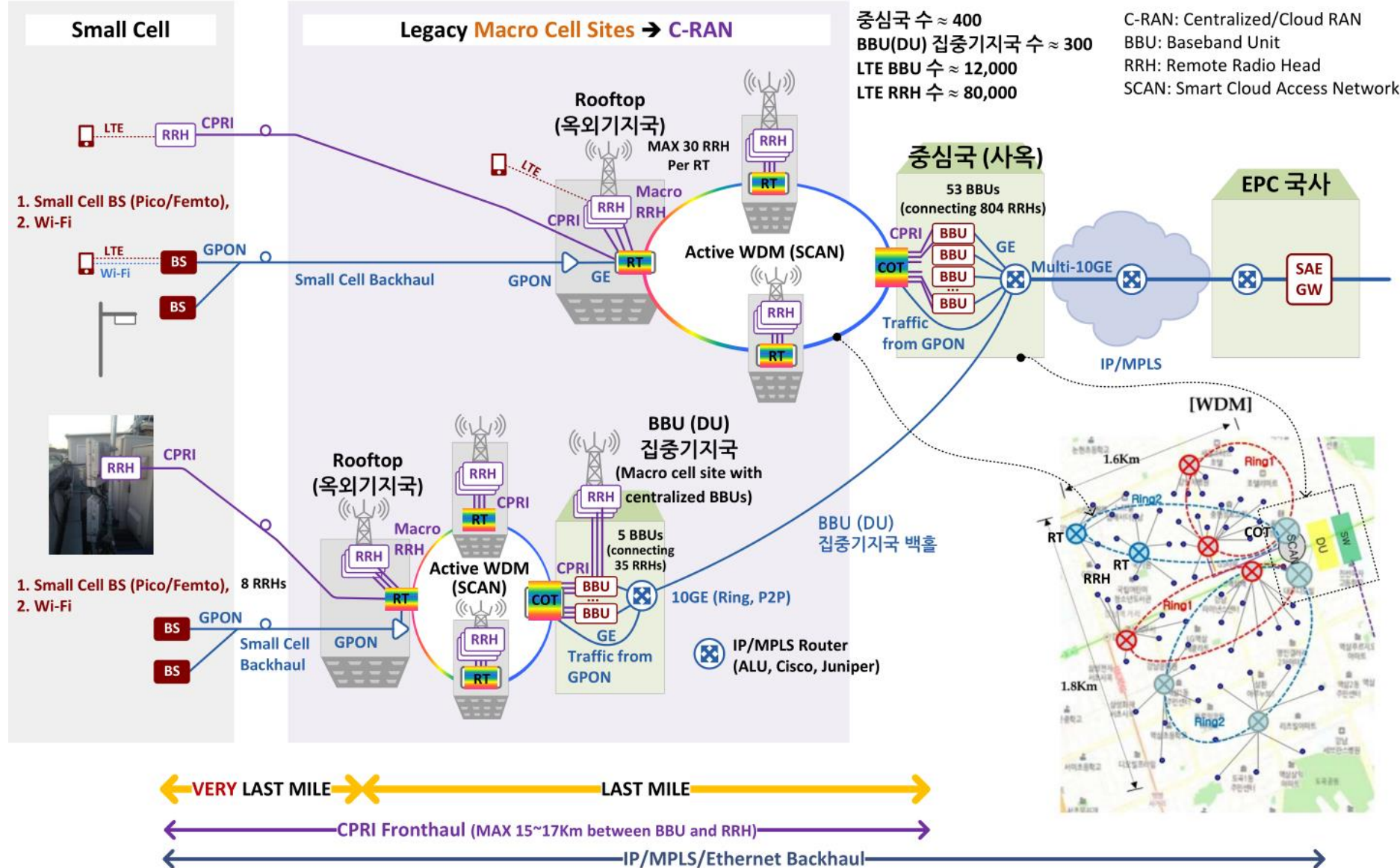
어떻게 더 많이 보내지? – 유선(WDM), 무선(MIMO)



어떤 통신기술들이 쓰이고 있을까? Access

5G/LTE
CPRI
PON
MSPP
WDM

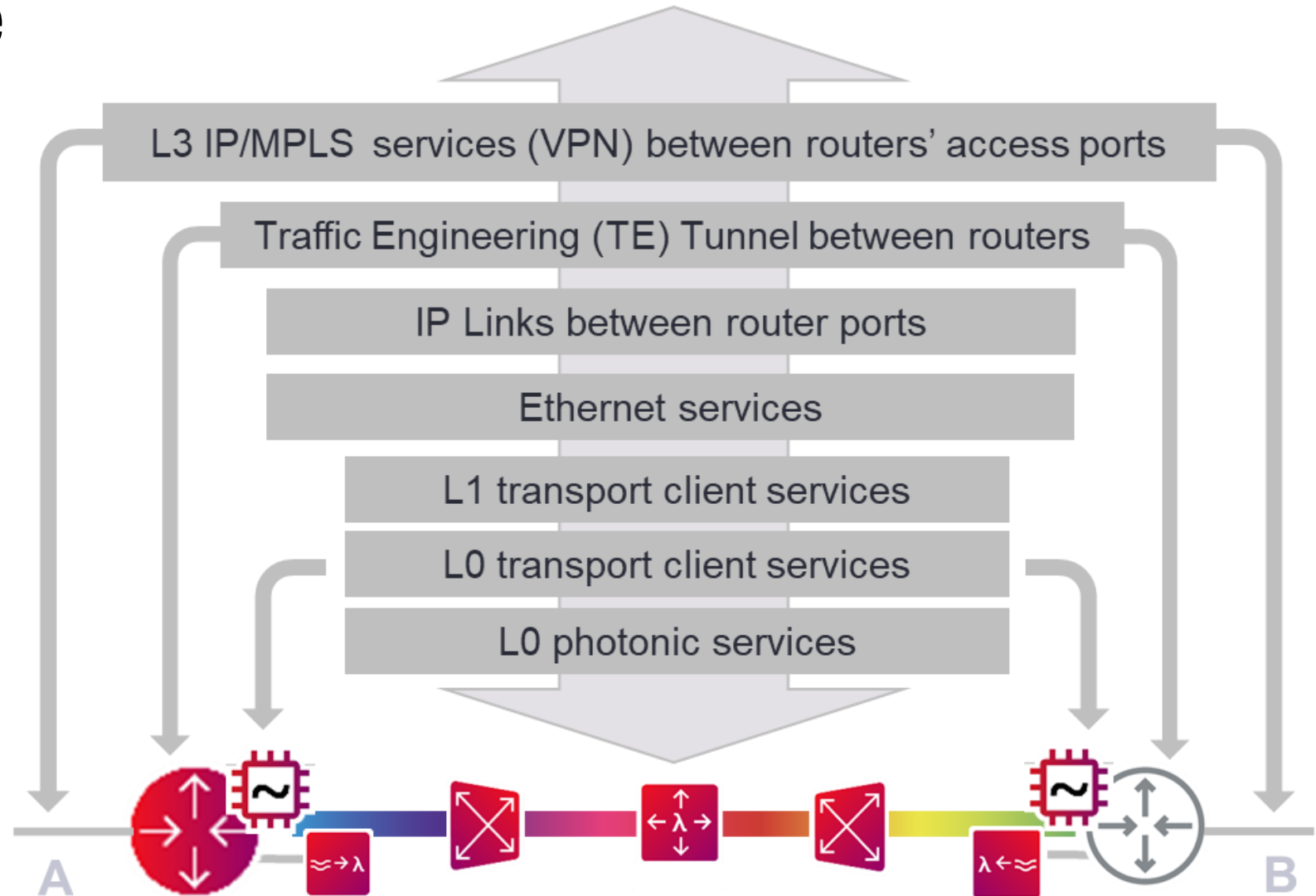
...



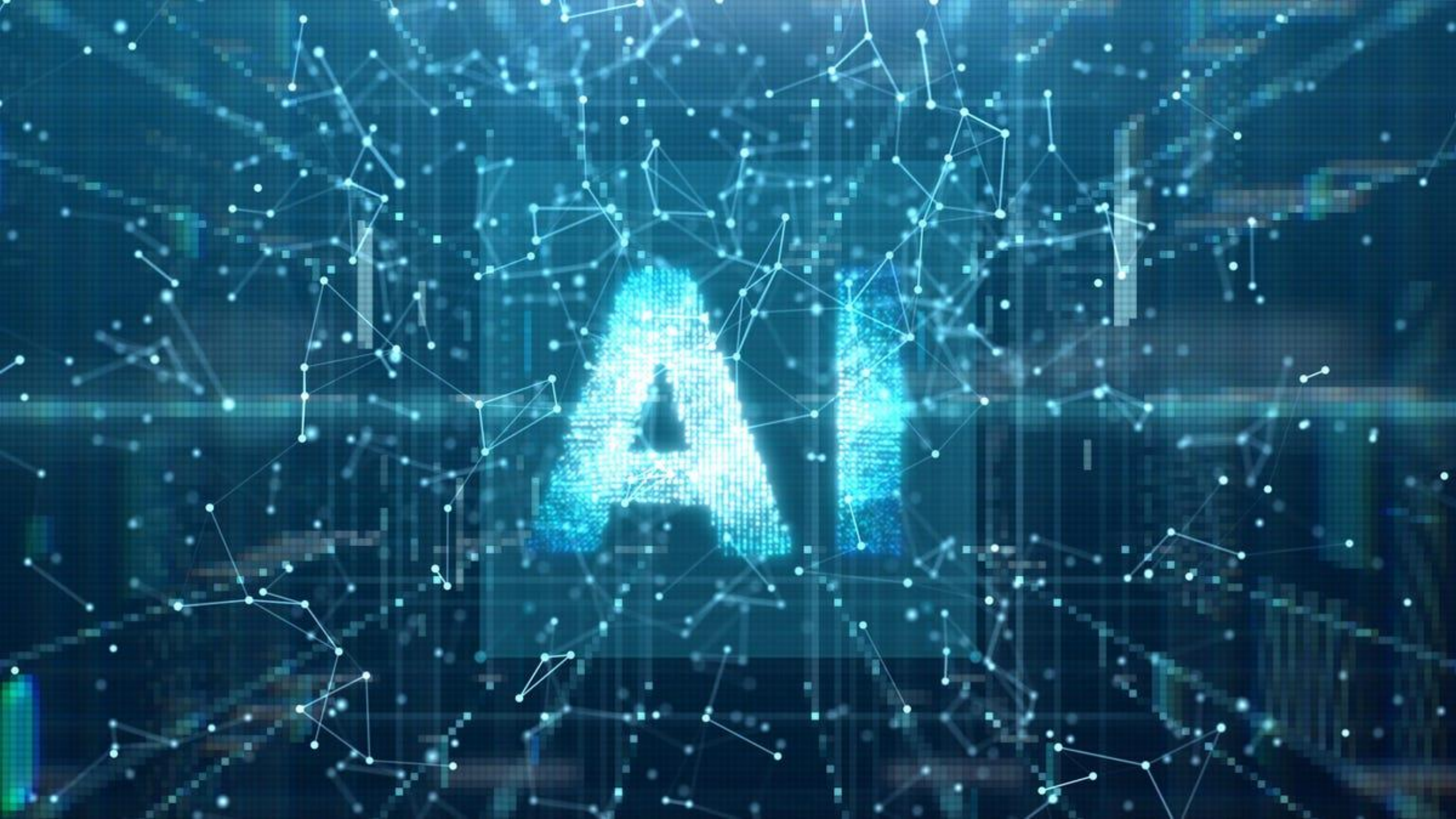
어떤 통신기술들이 쓰이고 있을까?

Metro/Core

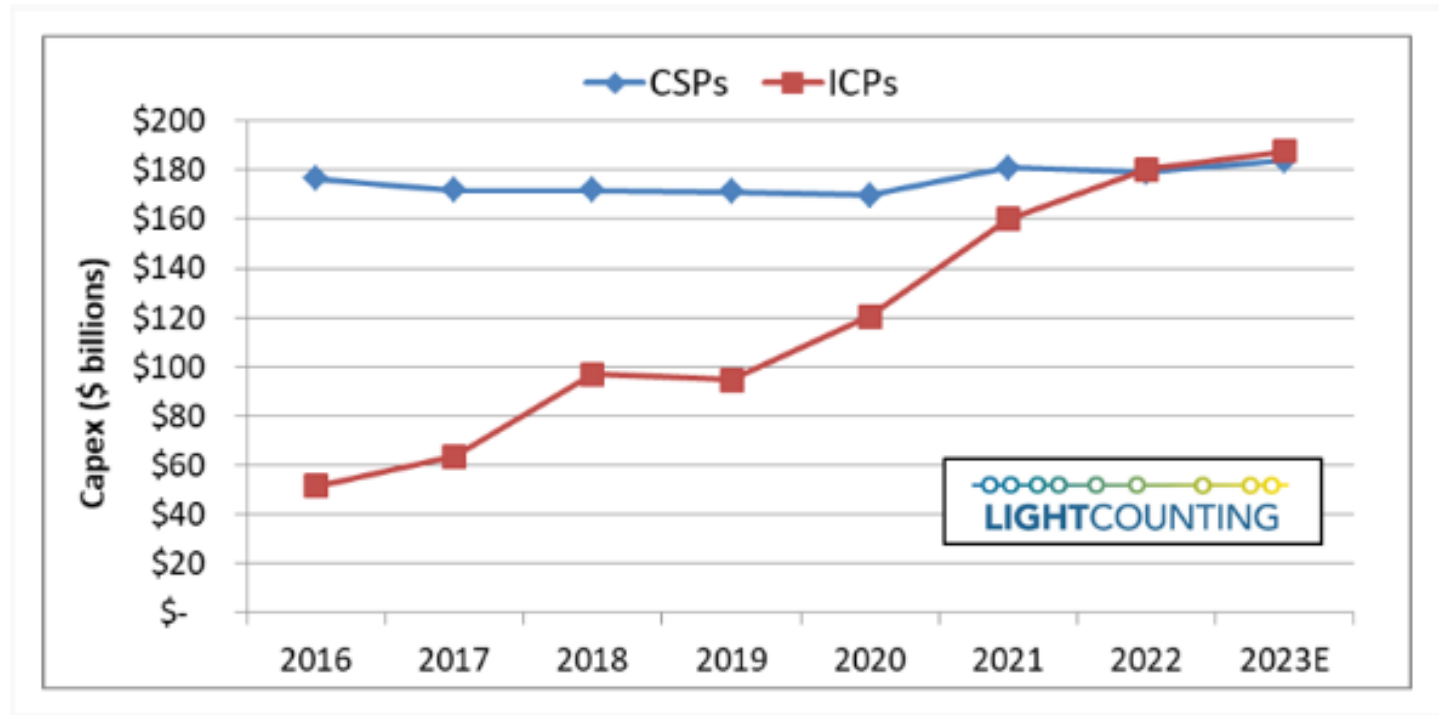
MPLS
Traffic Engineering
IP
Ethernet
OTN
ROADM
Coherent
Fiber



AI시대의 통신기술



CSP, ICP 투자예산 증감률 비교



AT&T, BT, DT, NTT,
Softbank, Verizon, etc.

Comm Service
Providers (CSPs)

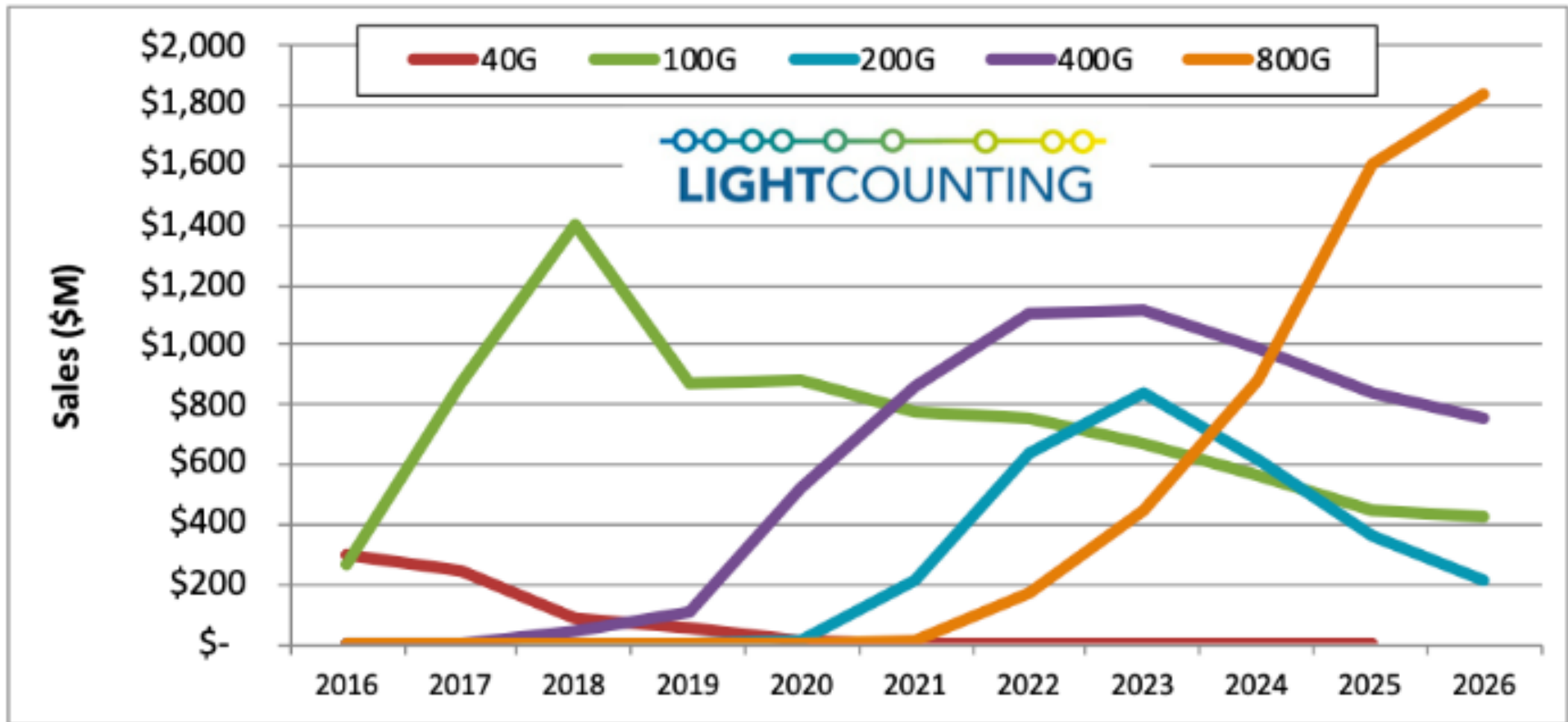
Internet Content
Providers (ICPs)

Alibaba, Alphabet, Apple,
Baidu, eBay, Equinix,
Meta, Tencent, etc.

Ethernet Transceiver 마켓 분석

Understanding the 400G+ market

Figure: Sales of Ethernet Transceivers to the Top 5 Cloud Companies



New supply chain shortages now
(100G VCSELs and 4x100G/8x100G DSP chips)

Gemini comes in three sizes



Ultra

Our most capable and largest model for highly-complex tasks.



Pro

Our best model for scaling across a wide range of tasks.



Nano

Our most efficient model for on-device tasks.

Chat GPT 4 이상 AI 요구사항

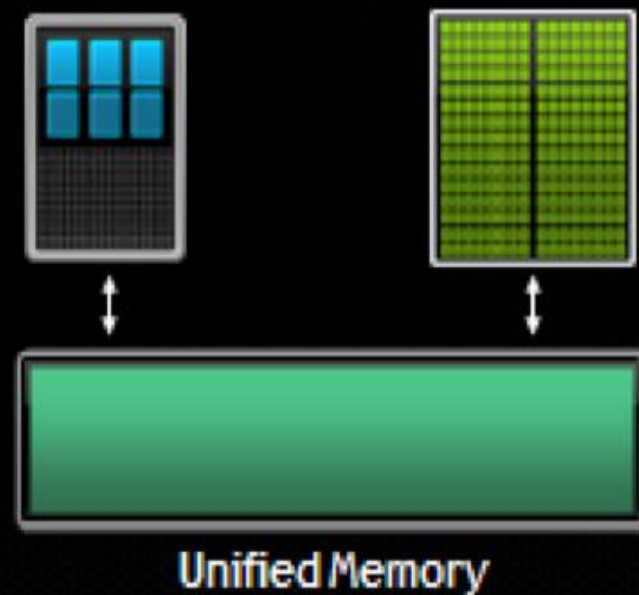
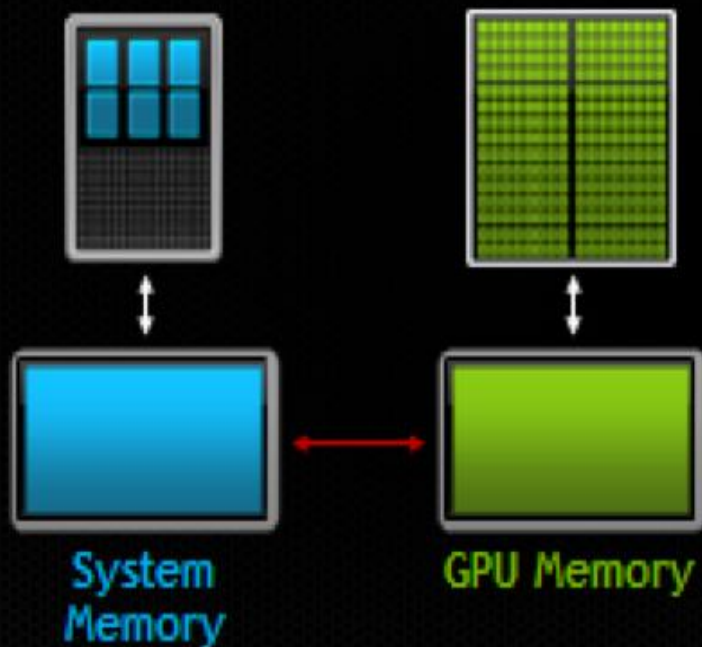


Mixtral 8×7B employs a **Mixture of Experts** (MoE) technique, incorporating eight expert models into one. This unique structure allows it to function at the speed of a **12B parameter-dense model** while containing four times the number of effective parameters. It's a breakthrough in efficiency but comes with a high VRAM requirement. The model needs a staggering **90GB of VRAM** in half-precision, posing a challenge for local setups and individual developers.



~145W

Unified Memory on Mac



COMPUTEX 2024 - NVIDIA Blackwell 소개 중..

GB200 GRACE BLACKWELL SUPERCHIP

40 PETAFLIPS OF AI PERFORMANCE

40 千萬億次的人工智慧效能



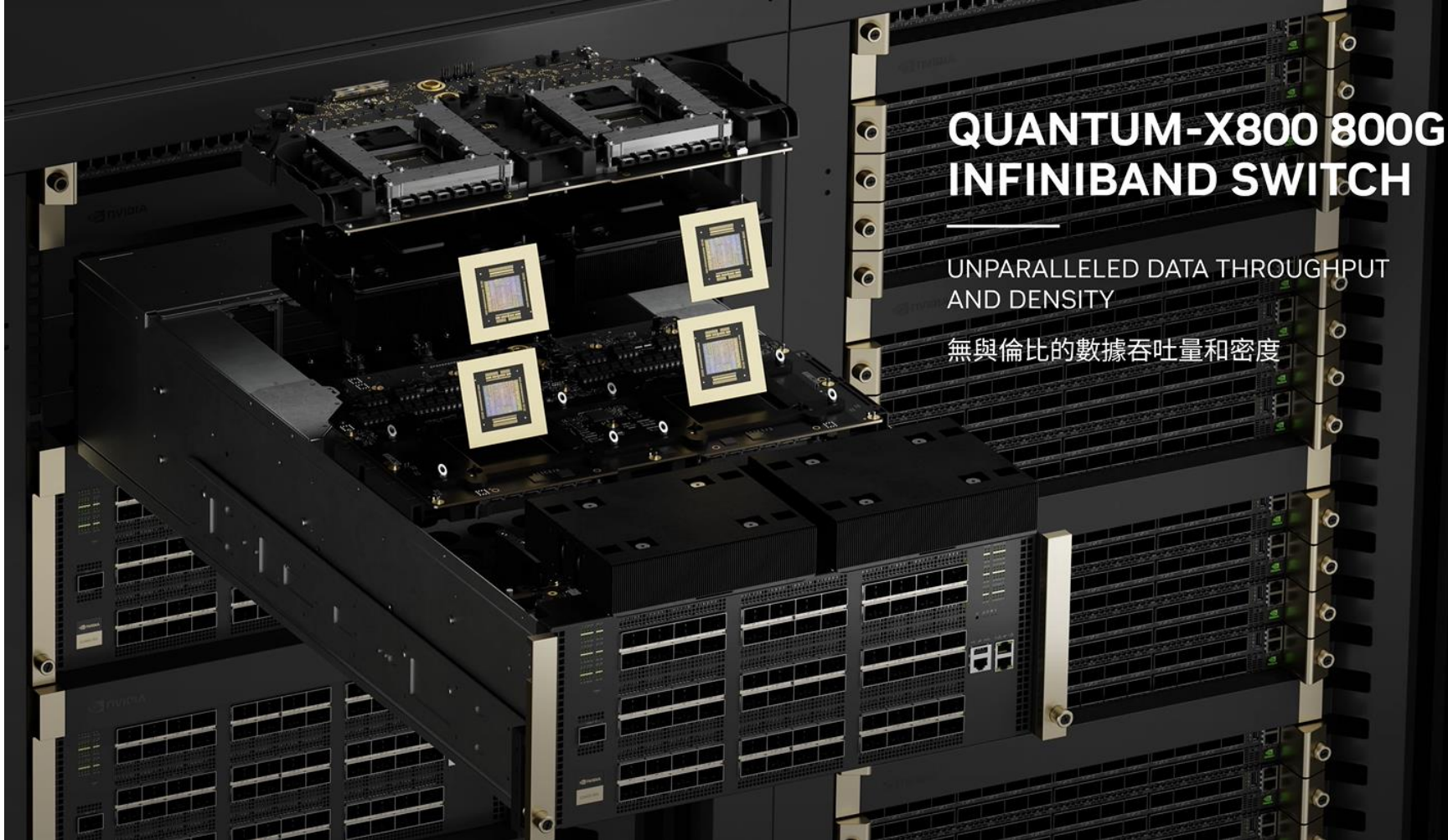
The AI Factory: How NVIDIA Builds the Engine for the New Industrial Revolution
https://www.youtube.com/watch?v=t_jiScDhf3M

COMPUTEX 2024 - NVIDIA Blackwell 소개 중..



The AI Factory: How NVIDIA Builds the Engine for the New Industrial Revolution
https://www.youtube.com/watch?v=t_jiScDhf3M

COMPUTEX 2024 - NVIDIA Blackwell 소개 중..



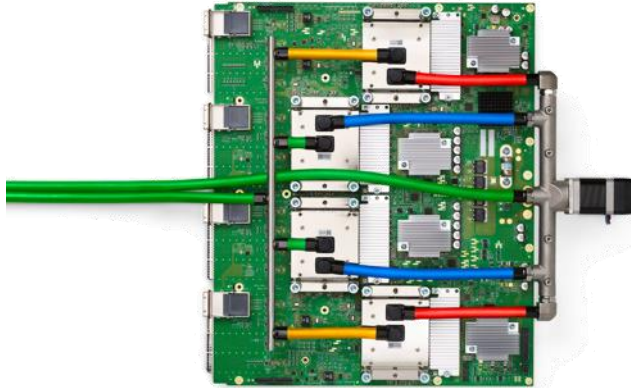
**QUANTUM-X800 800G
INFINIBAND SWITCH**

UNPARALLELED DATA THROUGHPUT
AND DENSITY

無與倫比的數據吞吐量和密度

OCS Theme : Already Deployed in ML/AI Cluster

Google TPUv4



TPU v4: An Optically Reconfigurable Supercomputer for Machine Learning with Hardware Support for Embeddings

Industrial Product*

Norman P. Jouppi, George Kurian, Sheng Li, Peter Ma, Rahul Nagarajan, Lifeng Nai, Nishant Patil, Suvinay Subramanian, Andy Swing, Brian Towles, Cliff Young, Xiang Zhou, Zongwei Zhou, and David Patterson

Google, Mountain View, CA

{jouppi,gkurian,lsheng,pcma,rahulnagarajan,lnai,nishantpatil,suvinay,aswing,btowles,cliffy,zhoux,zongweiz}@google.com and pattsrn@cs.berkeley.edu

ABSTRACT

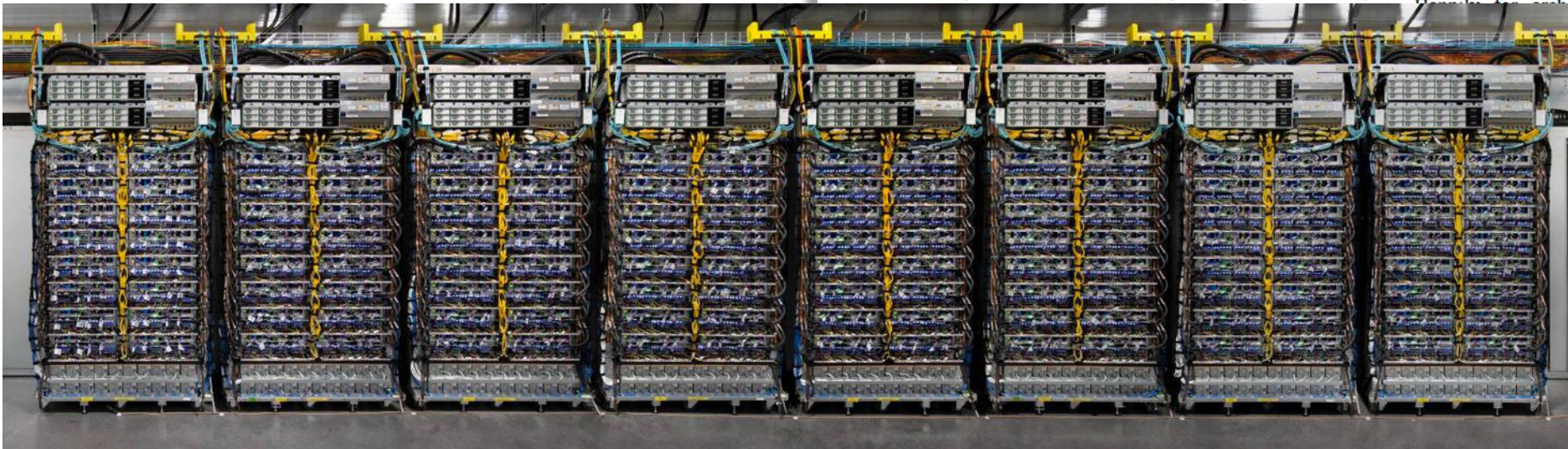
In response to innovations in machine learning (ML) models, production workloads changed radically and rapidly. TPU v4 is the fifth Google domain specific architecture (DSA) and its third supercomputer for such ML models. Optical circuit switches (OCSes) dynamically reconfigure its interconnect topology to improve scale, availability, utilization, modularity, deployment, security,

ACM Reference Format:

Norman P. Jouppi, George Kurian, Sheng Li, Peter Ma, Rahul Nagarajan, Lifeng Nai, Nishant Patil, Suvinay Subramanian, Andy Swing, Brian Towles, Cliff Young, Xiang Zhou, Zongwei Zhou, David Patterson. 2023. TPU v4: An Optically Reconfigurable Supercomputer for Machine Learning with Hardware Support for Embeddings: Industrial Product. In The 50th Annual International Symposium on Computer Architecture (ISCA '23), June 17–21, 2023, Orlando, FL, USA. ACM, New York, NY, USA, 15 pages. <https://doi.org/10.1145/3579371.3589350>.

1 INTRODUCTION

Translating machine learning (ML) models into challenging ways, both in scale and complexity (see Table 1 and Section 7.7). Examples of *language models (LLMs)* and examples of *embeddings* necessary for recommender systems or *DLRMs* (Deep Learning Recommendation Models) are the foundations of Transformers and BERT. The recent LLMs [Bro20,Tho22,Nay21] has

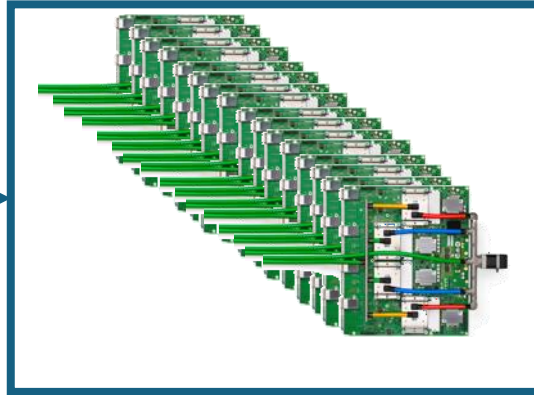


TPUv4 Architecture

1 rack = 64 chips (chip = TPUv4)

1 row = 8 rack = $64 \times 8 = 512$ chips

8 row = $512 \times 8 = 4096$ chips

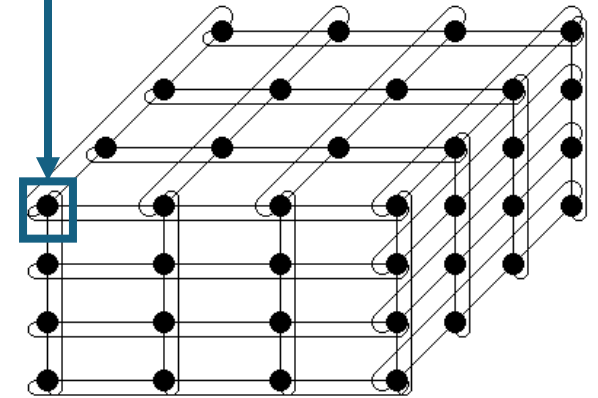
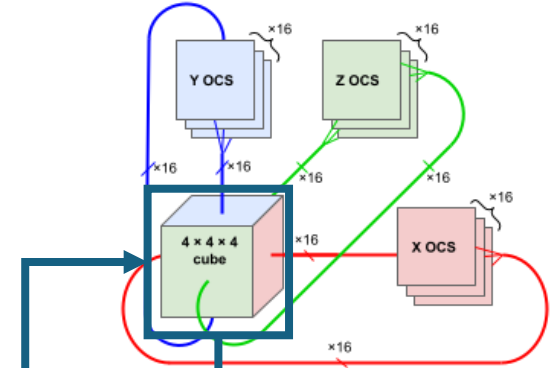
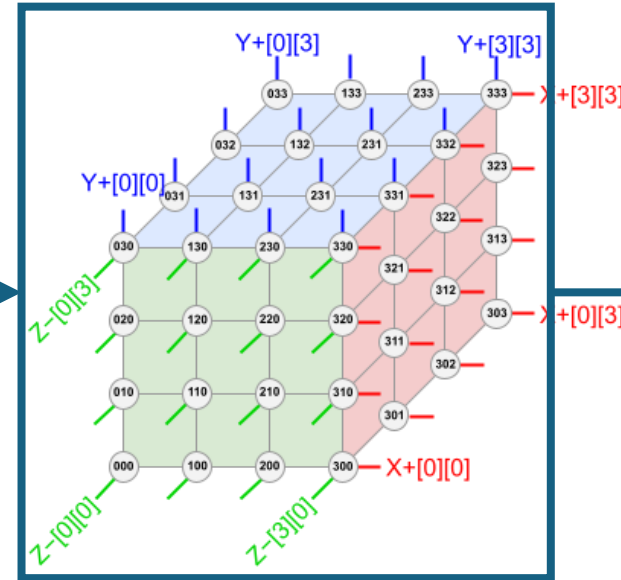


- Each rack contains 16 tray-host server pairs.
- 16 tray x 4 chips per tray = 64 chips
- Also 4 chip per CPU host.
- Each rack contains 64 chips and 16 CPU hosts (aka one 4^3 Cube)
- $4 \times 4 \times 4$ 3D electrical mesh per rack.

One 4^3 Cube

$4 \times 4 \times 4 = 64$ -chips / Cube

6×16 i/f = 96 i/f



3D Torus

64 cubes = 64×64 chips per
cube = 4096 chips

TPUv4 Architecture

Google Cloud

문서

기술 영역 ▾

크로스 프로젝트 도구 ▾

관련 사이트 ▾

Cloud TPU

가이드

참조

지원

리소스

필터

둘러보기

Cloud TPU 소개

시스템 아키텍처

TPU 런타임

TPU 버전

리전과 영역

지원되는 모델

내부 IP 주소 범위

홈 > Cloud TPU > 문서 > 가이드

TPU v5e

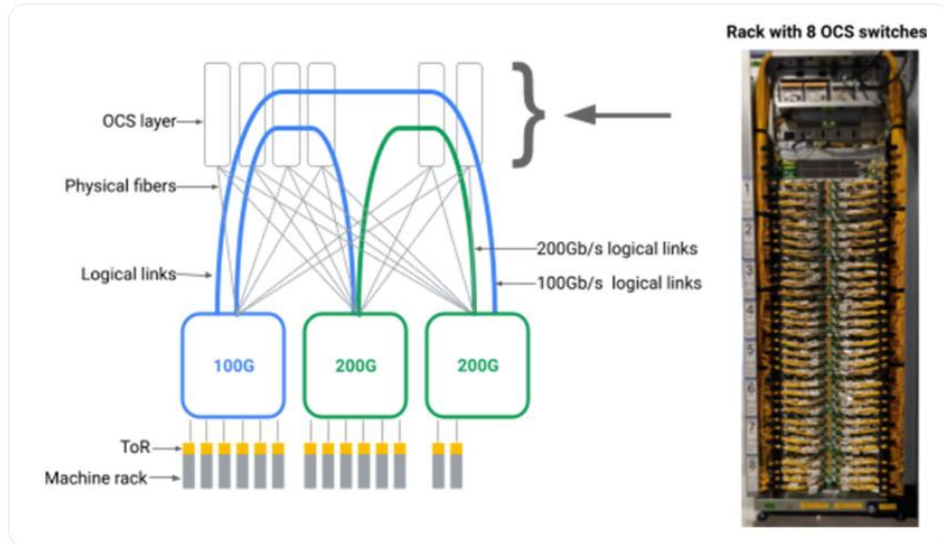
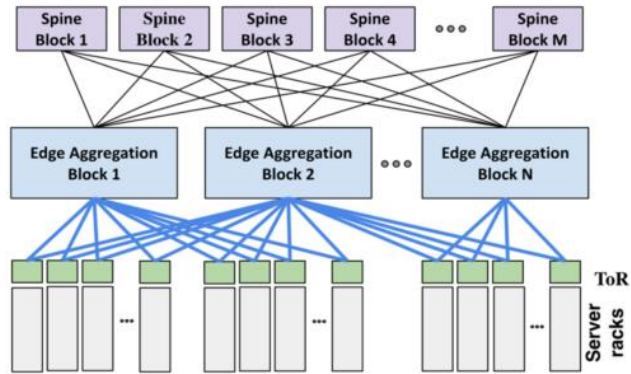
이 문서에서는 Cloud TPU v5e의 아키텍처와 지원되는 구성을 설명합니다.

의견 보내기

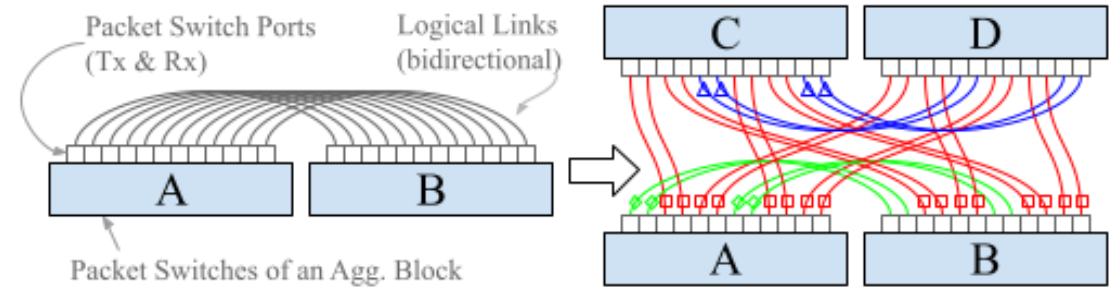
다음 표는 v5e의 포트 사양 및 값을 보여줍니다.

주요 포트 사양	v5e 값
TPU Pod 크기	칩 256개
상호 연결 토폴로지	2D 토러스
포트당 최고 컴퓨팅	100페타옵스(Int8)
포트당 올리듀스 대역폭	51.2TB/초
포트당 바이섹션 대역폭	1.6TB/초
포트당 데이터 센터 네트워크 대역폭	6.4Tbps

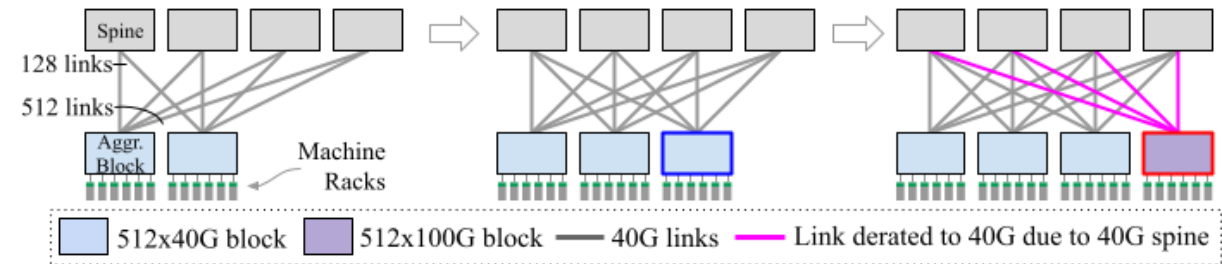
Google DC 내부 개선 사례 – “Mission Apollo”



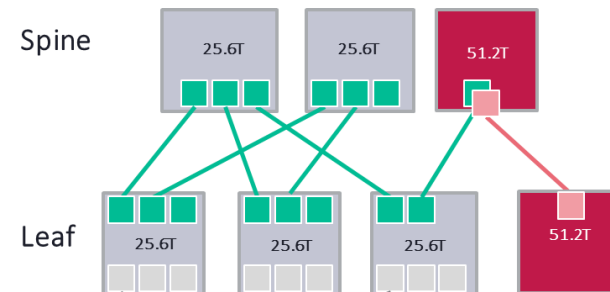
Pain #1: Rewiring



Pain #2: Spine Limit

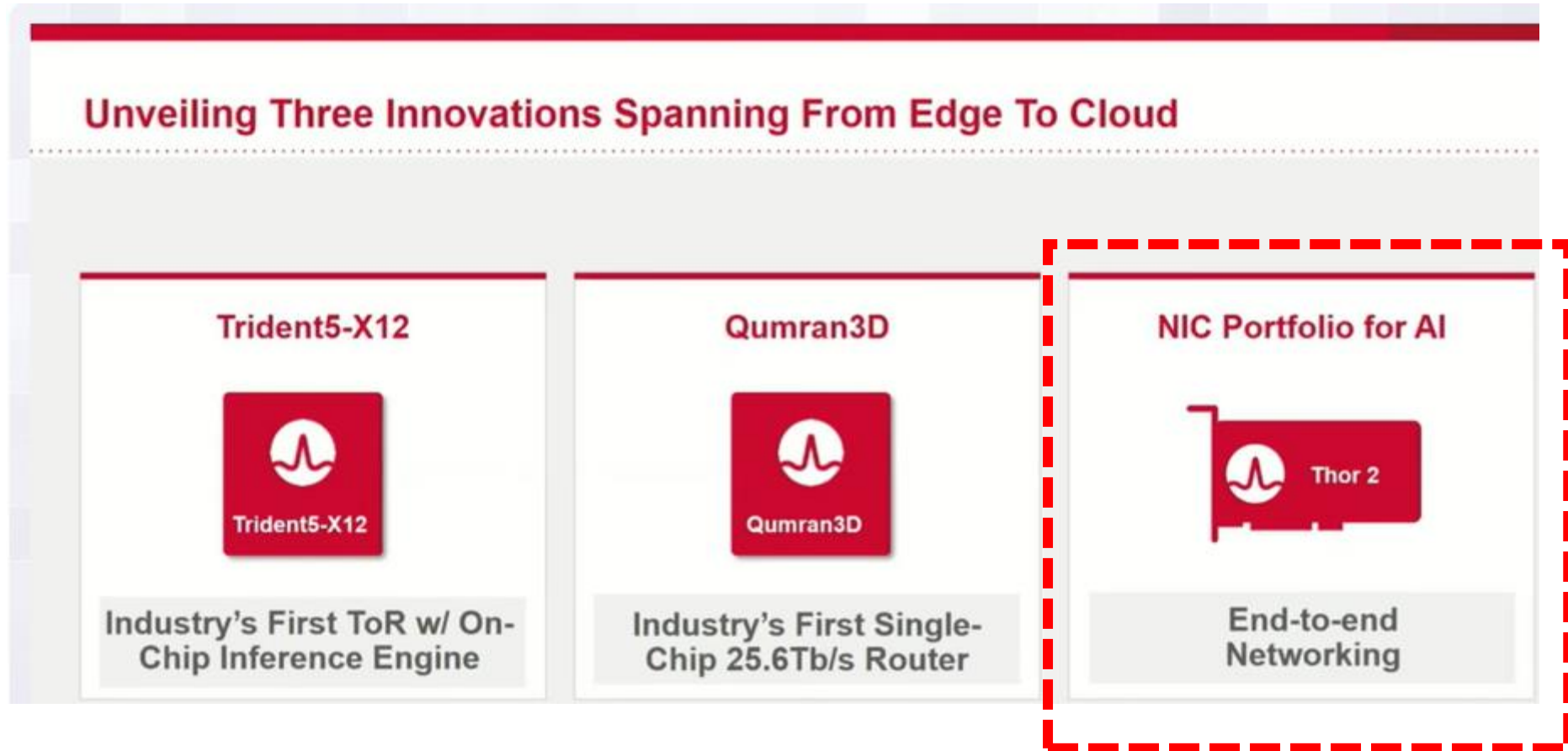


Pain #3: Retrofitting Optics

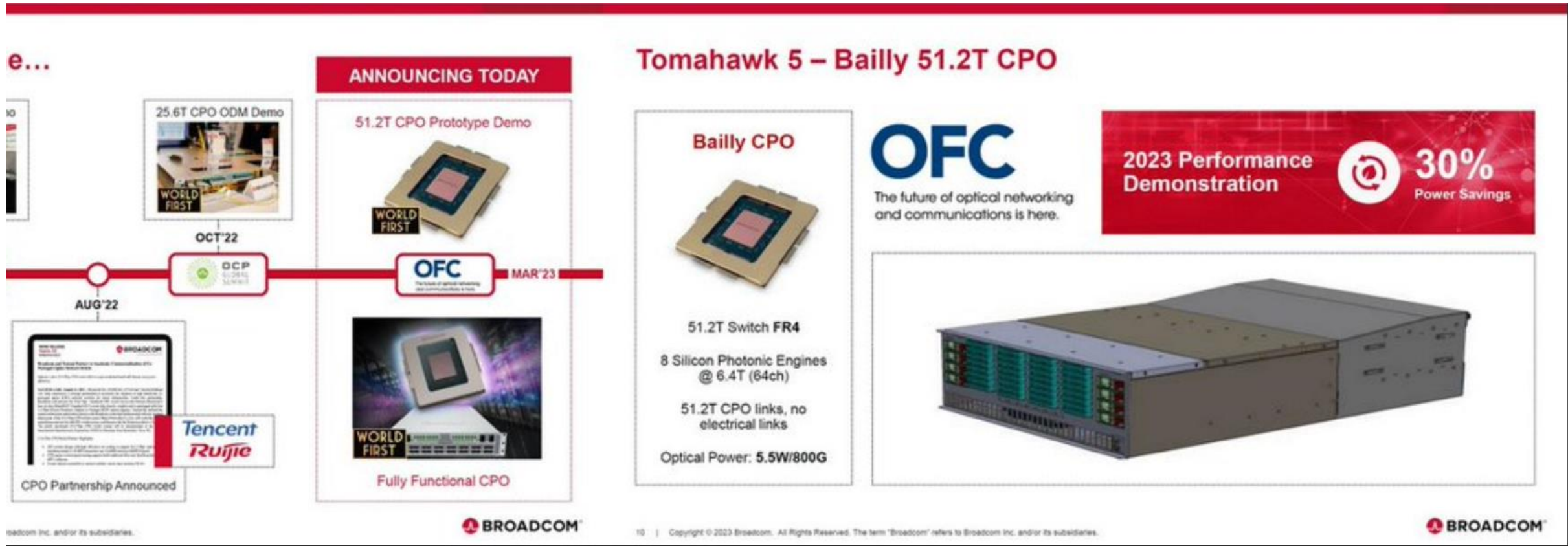


What's NEXT?

Broadcom Roadmap 소개



Broadcom Roadmap 소개



*https://x.com/ogawa_tter/status/1768455026991895031/photo/4

400/800G Client pluggable은 왜 비싸까?



Examples of optical front-end types and their subcomponents are outlined in the table below. There are many more varieties (e.g., LR1, SR4, etc.), but they follow the same patterns as shown.

Transceiver type	Laser Type	Number of Lasers	Fiber Type	Relative Cost
400G-SR8	50G VCSEL	8	8xMMF ribbon	\$
400G-DR4	100G wideband laser (DML)	Usually, 1, split 4 ways	4xSMF ribbon	\$\$
400G-FR4	100G CWDM laser (EML)	4	1xSMF	\$\$\$
400G-LR4	100G high-power CWDM laser (EML or SiPho)	4	1xSMF	\$\$\$\$
800G-SR8	100G VCSEL	8	8xMMF ribbon	\$\$
800G-DR8	100G wideband laser (DML or EML)	Usually, 1, split 4 ways	8xSMF ribbon	\$\$\$
800G-2xFR4	100G CWDM laser (EML or SiPho)	Usually, 4, split 2 ways	2xSMF	\$\$\$\$

400/800G Client pluggable은 왜 비쌀까?

OOK

On/Off
Keying

IMDD

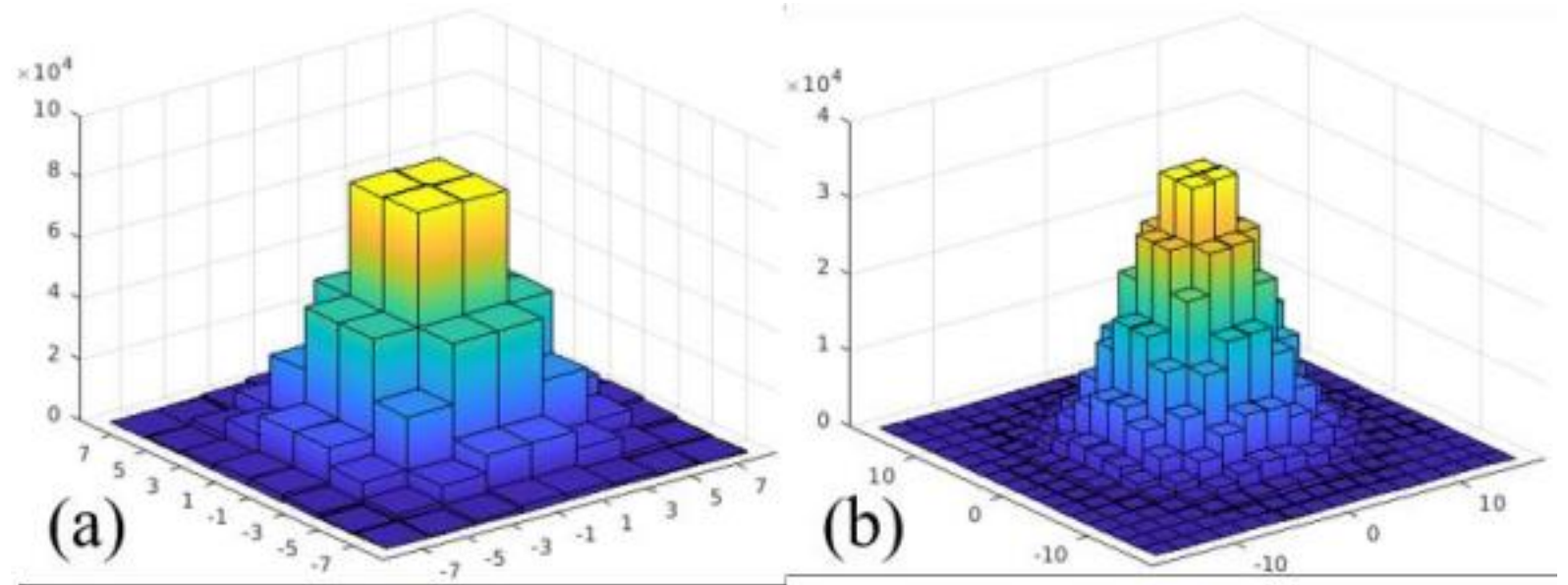
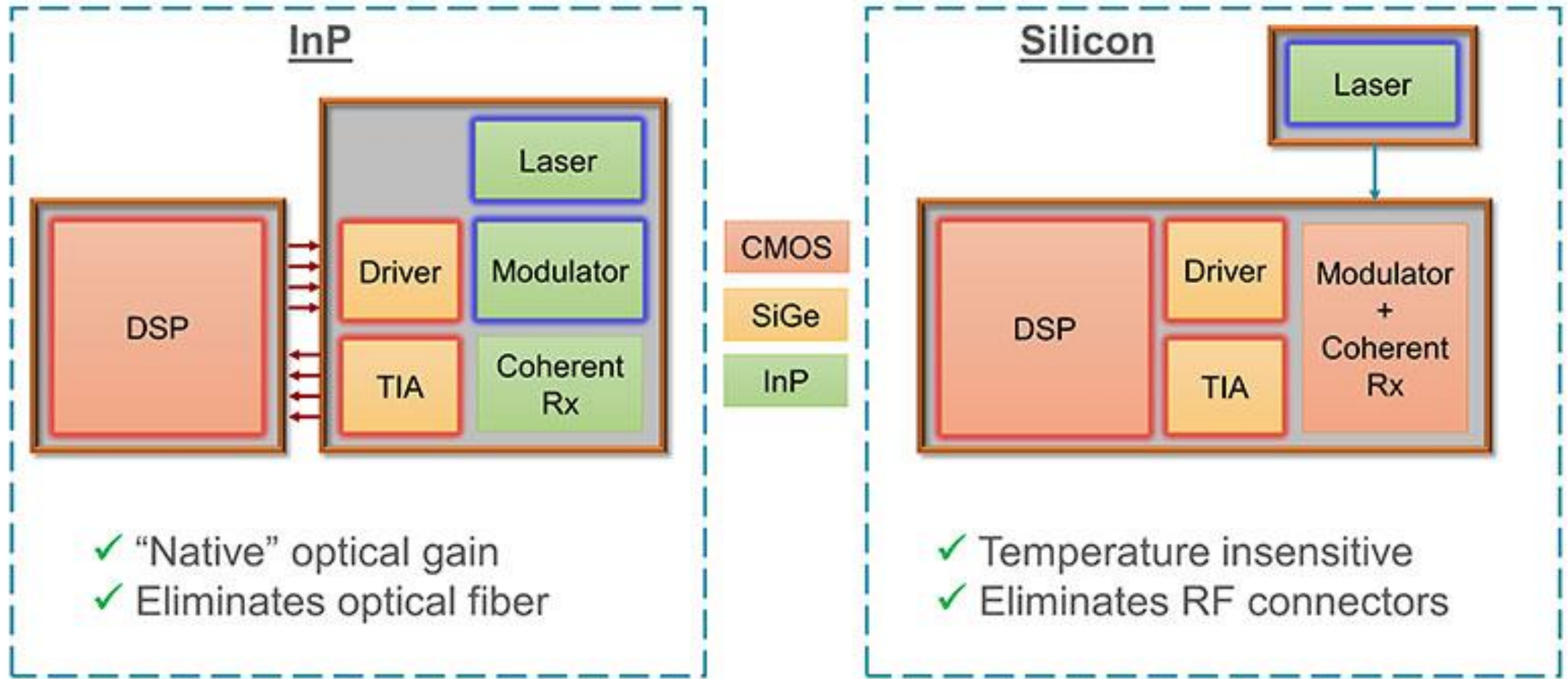


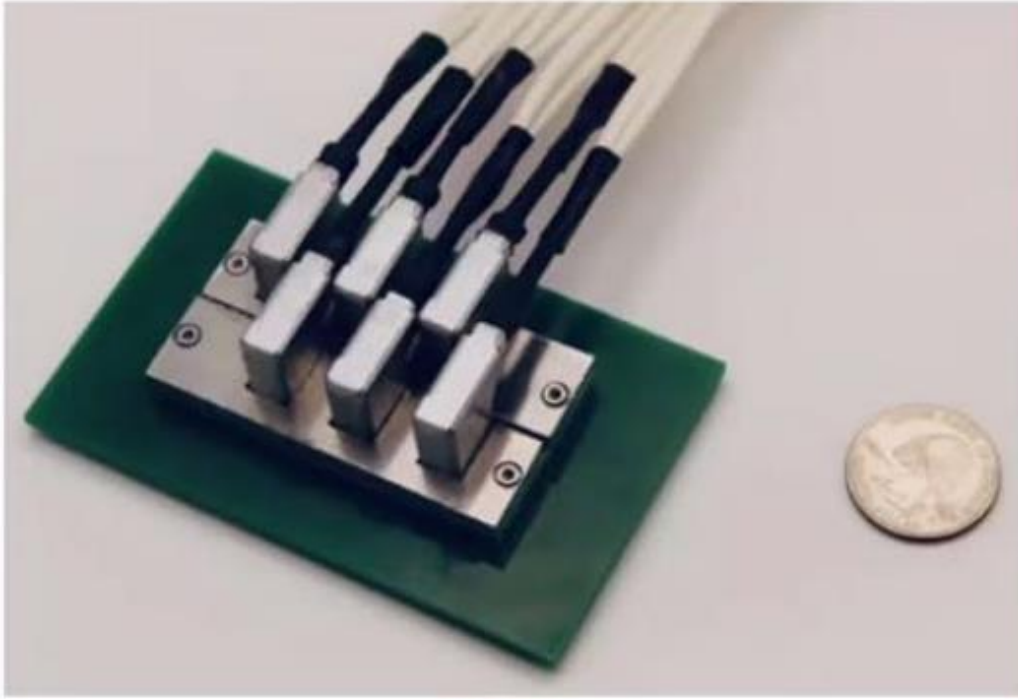
Fig. 1 Histogram of (a) PCS 64-QAM with $\nu = 0.0749$.

DSP

400/800G Client pluggable은 왜 비싸까?



400/800G Client pluggable은 왜 비쌀까?



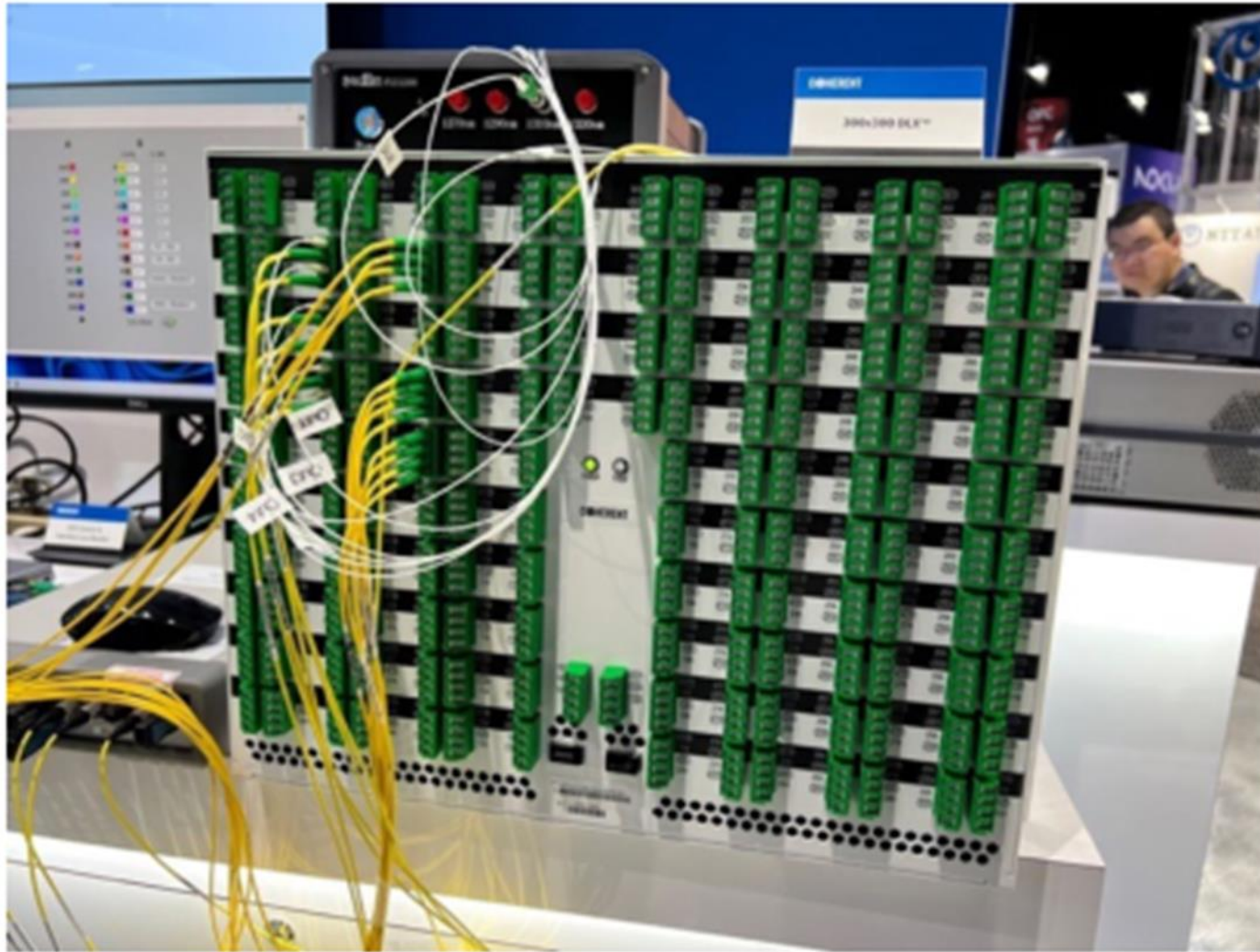
10 Tbps
(on half a business card)



8 Tbps
(ML/AI PCIe Card Interface)

Source: Nubis Communications

Coherent 社 광 스위치



Source: Coherent

노이즈 = 딜레이 = 투자비

감사합니다