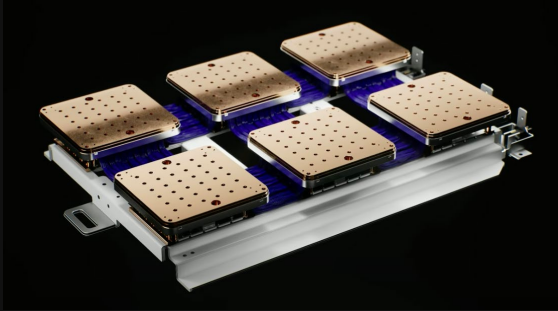




Improving AI Infrastructure Efficiency with Optical I/O

Miloš Popović | Senior Technical Architect & Co-Founder
milos@ayarlabs.com

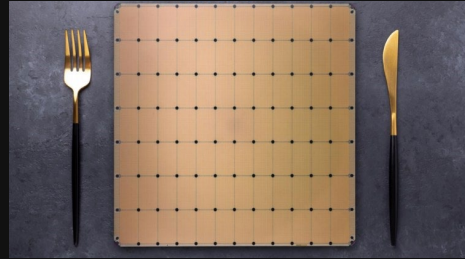
Optical I/O transformation of compute system scaling



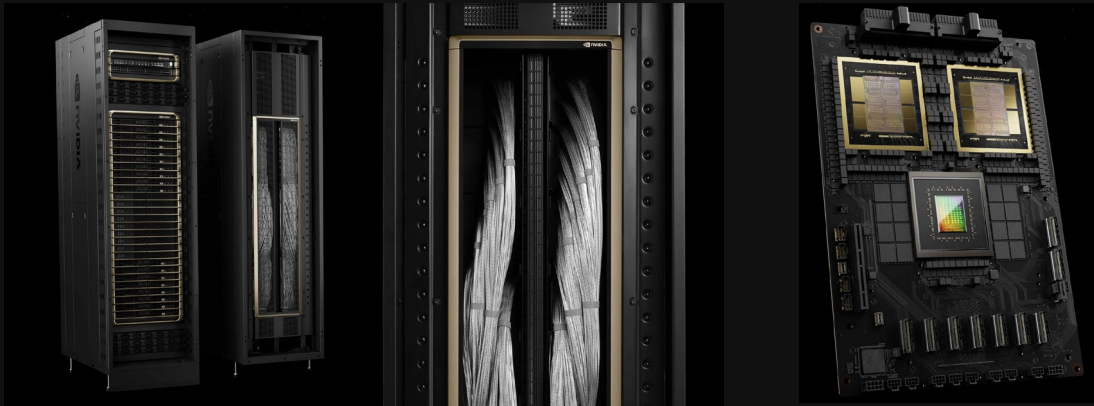
Tesla Dojo (Tesla AI Day 2022).



Cerebras CS-2 wafer size chip.

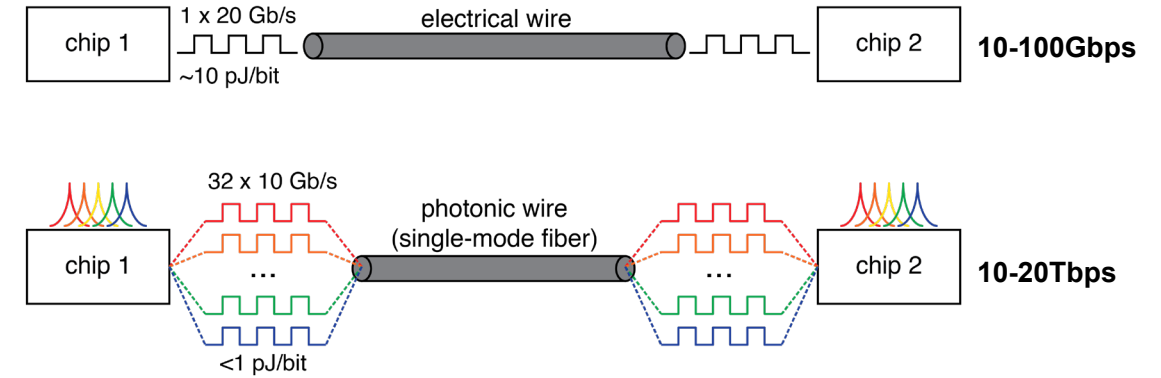


Nvidia GTC 2024 Keynote.

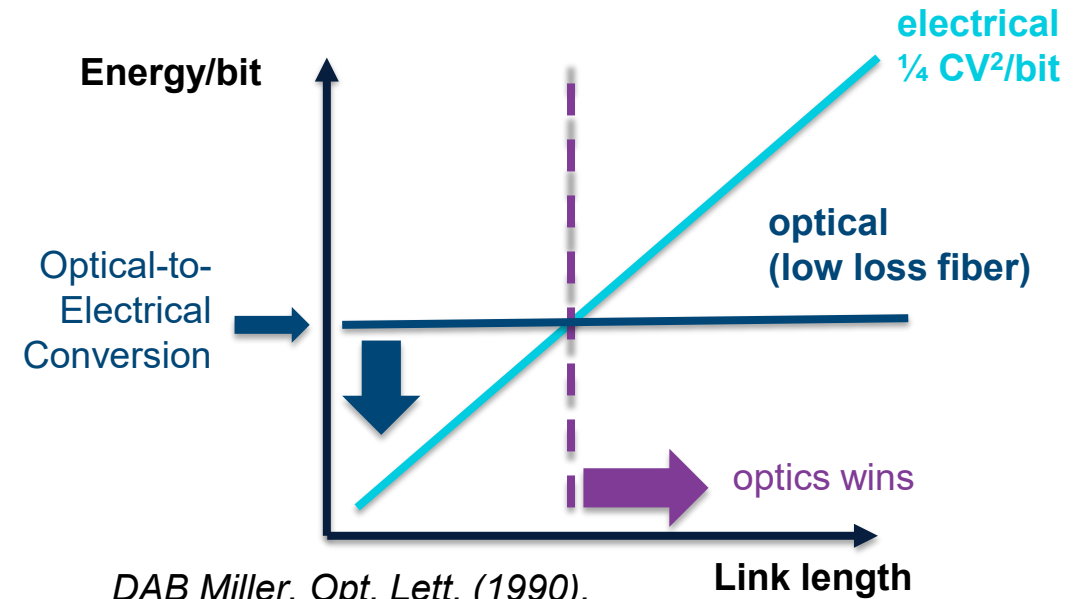


AI compute architectures limited by data movement:

- I/O communication energy and bandwidth density

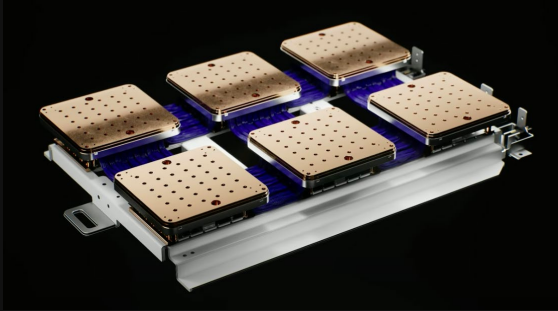


Optics provides bandwidth density: many wavelengths, smaller.



Optics provides low energy – requires deep integration.

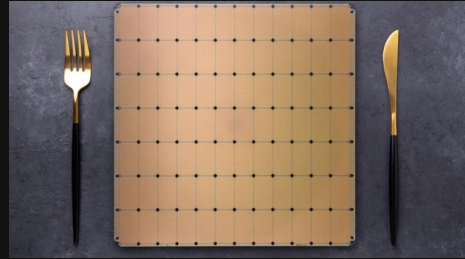
Optical I/O transformation of compute system scaling



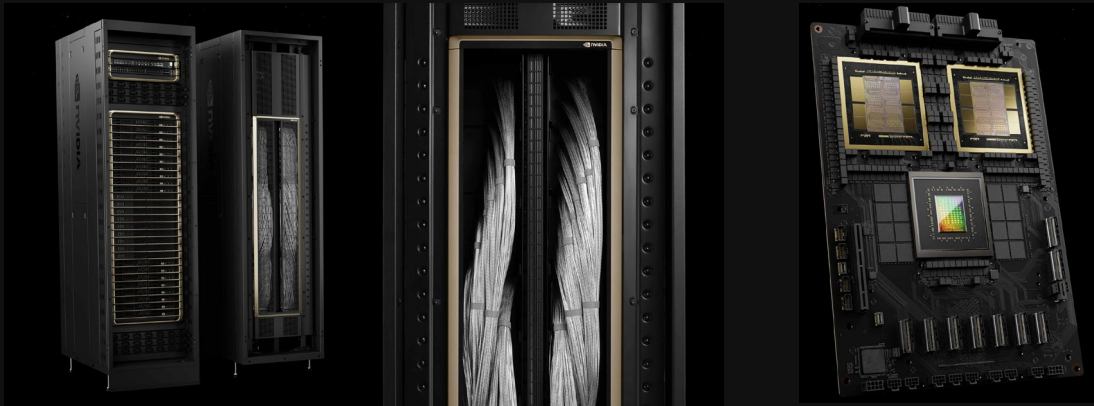
Tesla Dojo (Tesla AI Day 2022).



*Cerebras CS-2
wafer size chip.*

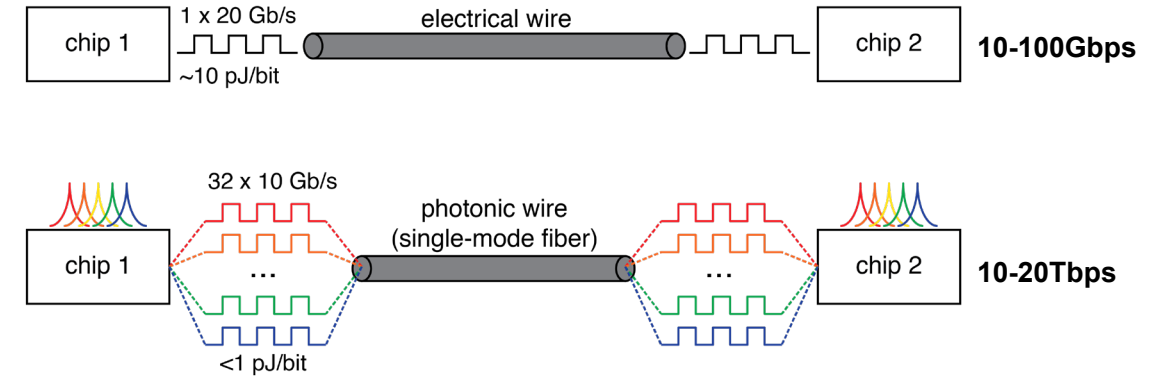


Nvidia GTC 2024 Keynote.

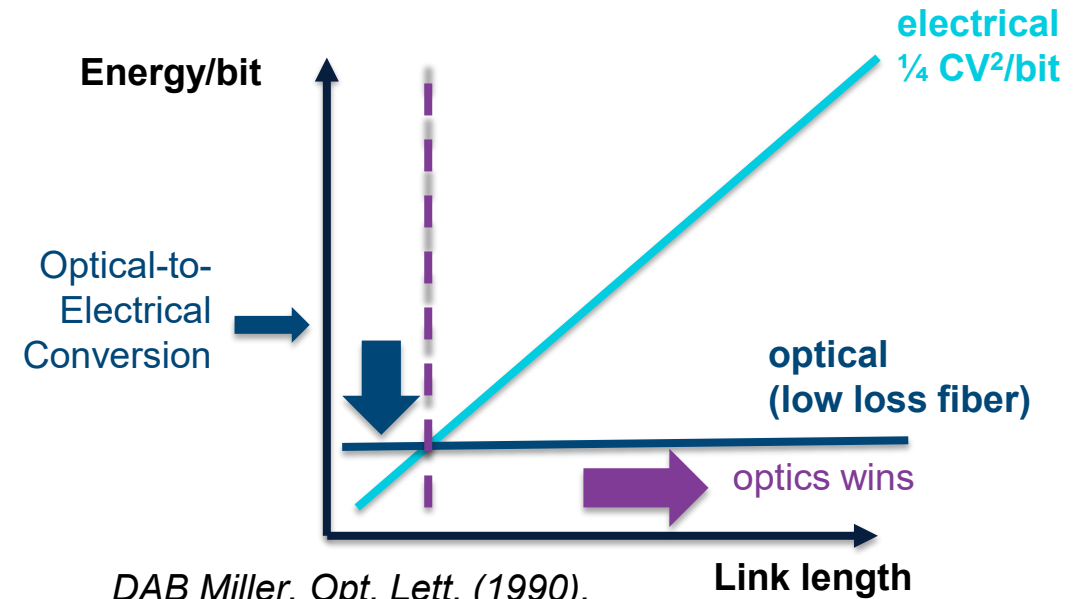


AI compute architectures limited by data movement:

- I/O communication energy and bandwidth density



Optics provides bandwidth density: many wavelengths, smaller.

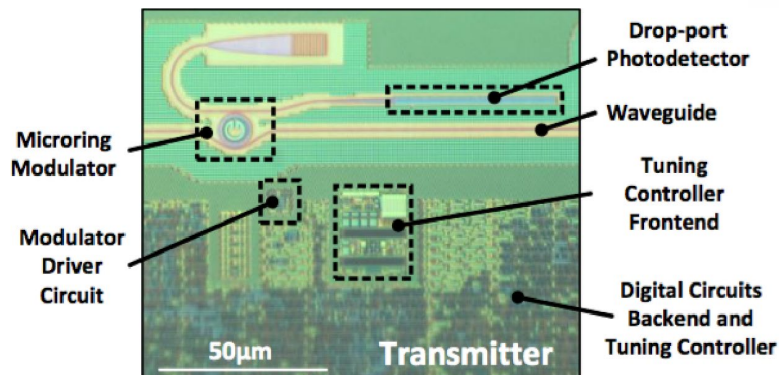
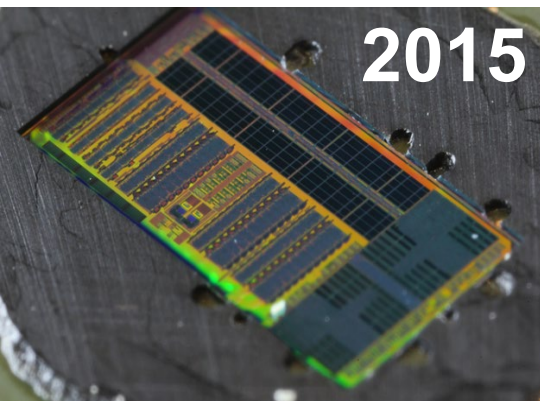


DAB Miller, Opt. Lett. (1990).

Optics provides low energy – requires deep integration.

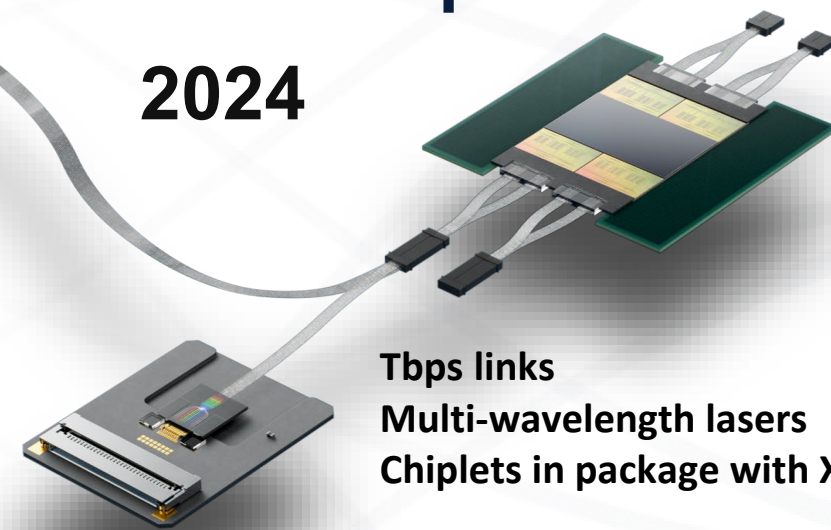
Optical I/O chiplet enabled interconnection of compute

1



2

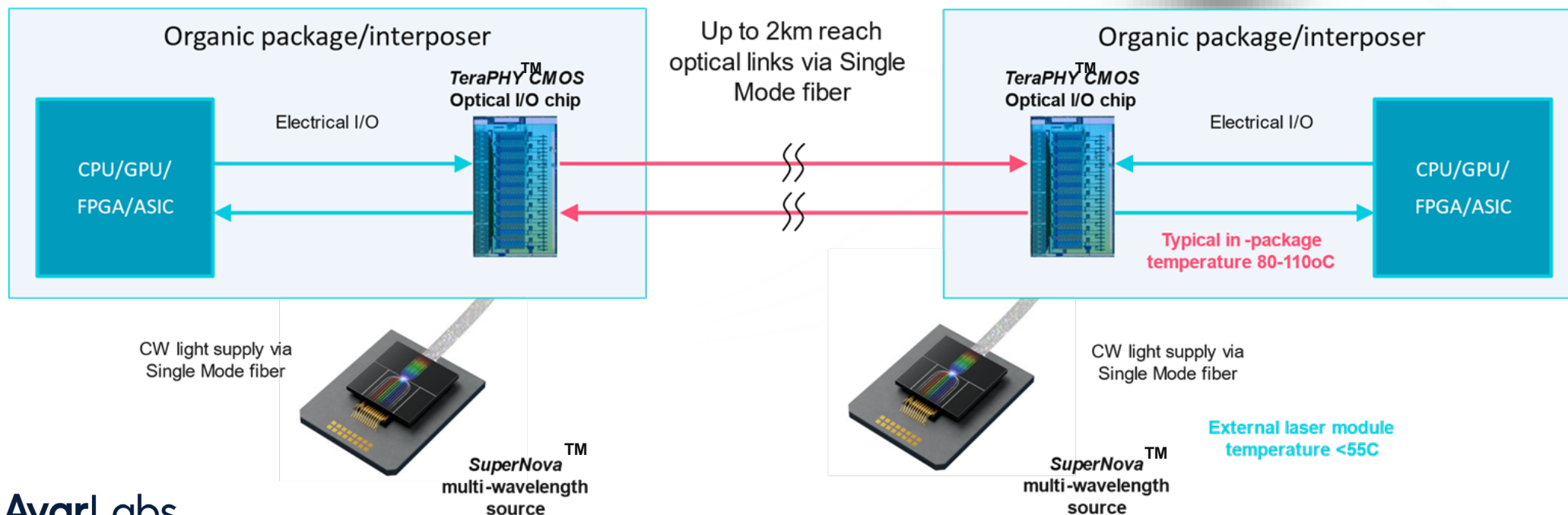
2024



- First microprocessor with light based interconnect (2015)
- University research -> commercial 300mm CMOS fab

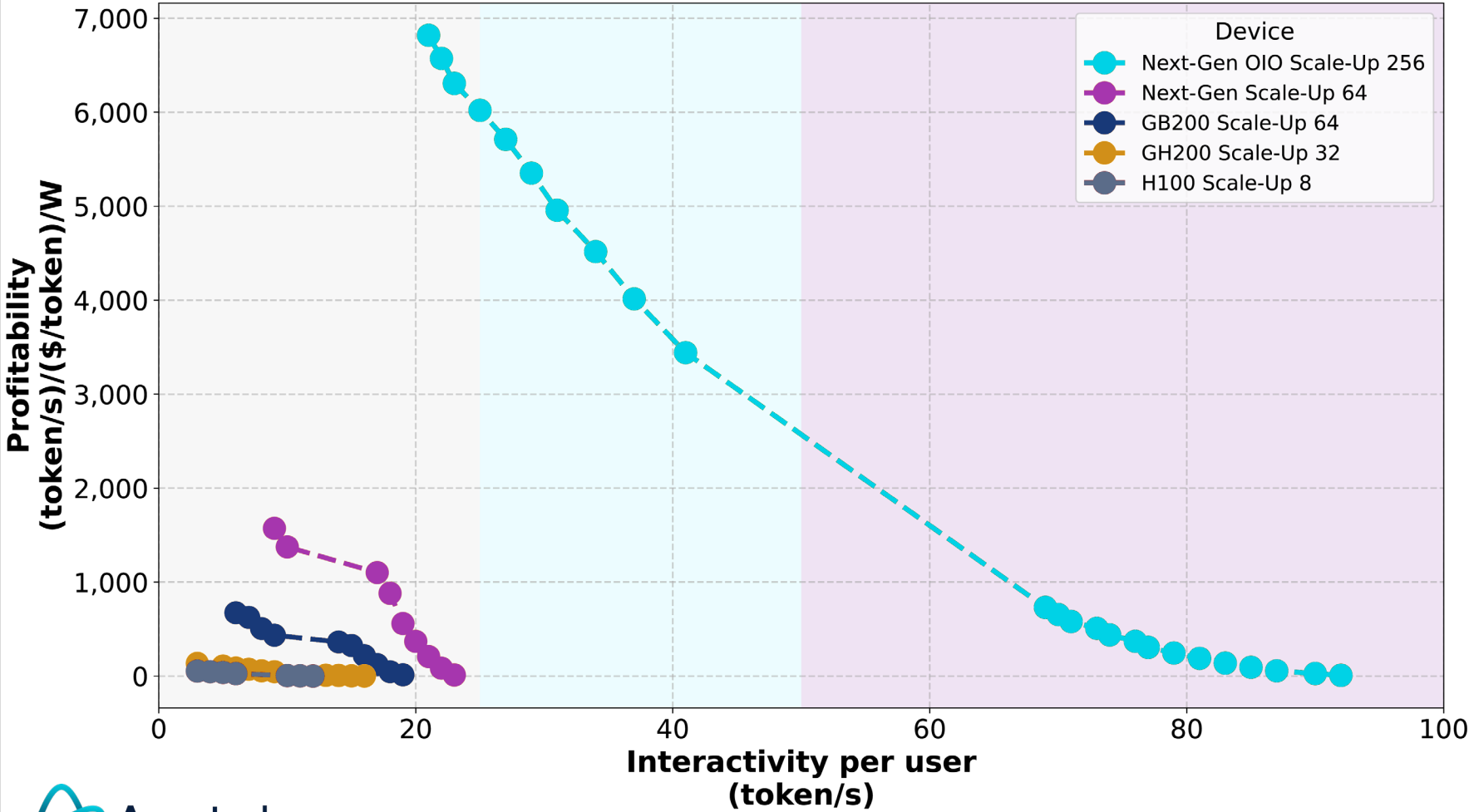


3



Optical I/O impact on future GPT-X performance

Figure of Merit	Description	Units
Throughput	Number of Users	Token
	Time Per Output Token	Second
Interactivity	1	Token
	Time Per Output Token	Second
Profitability (AI efficiency)	Throughput	Token / Second
	Cost X Power	\$ / Token
		X Watts



State of Art ~2026-27

14T parameters
Mixture of Experts
(MoE) w/ 32 experts

inference with:
32K/8K sequence length
5s Time to First Token

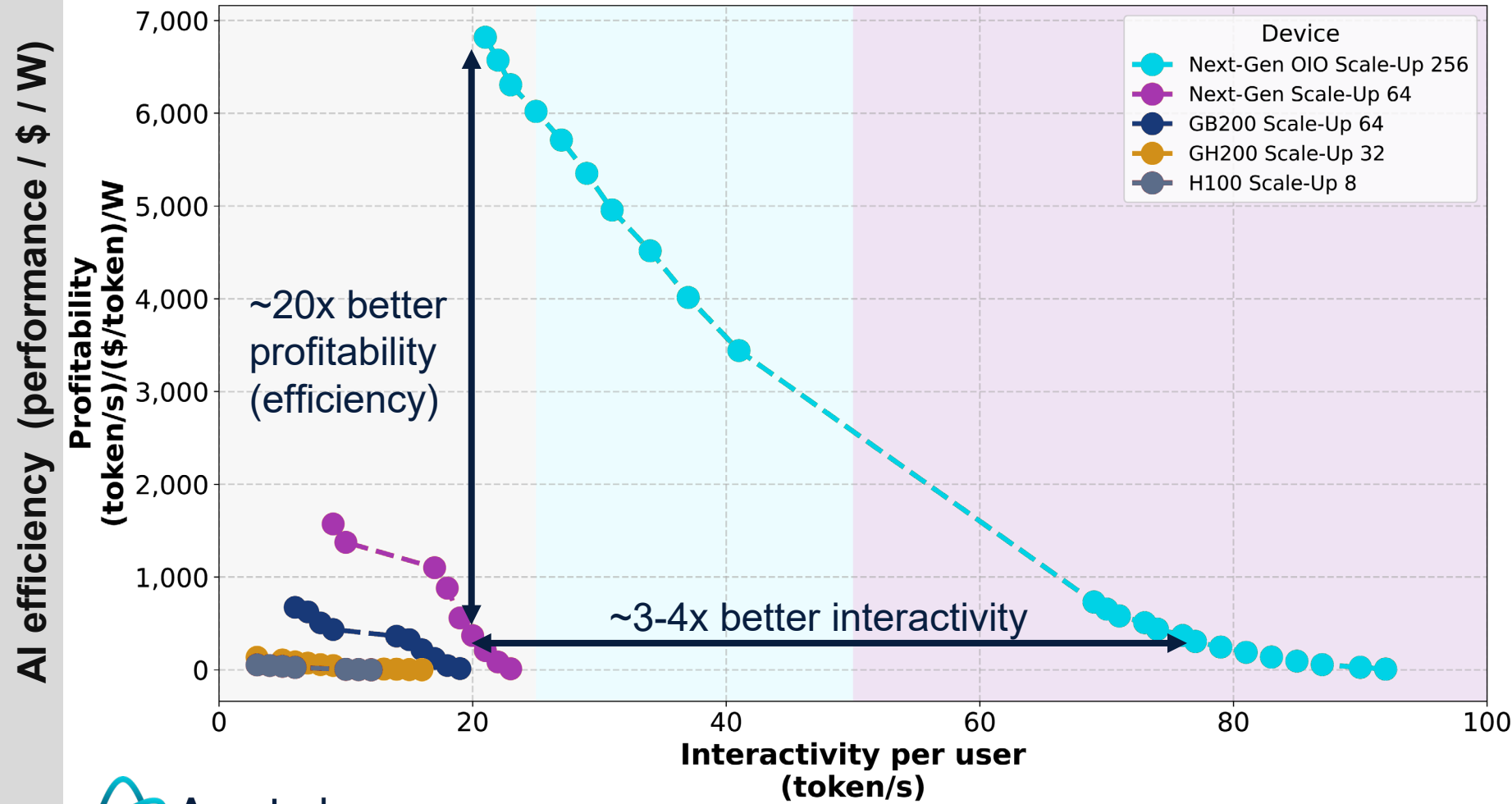
Optimizations:

Tensor/Sequence parallel
Pipeline parallel
Expert parallel

Optical I/O impact on future GPT-X performance

Figure of Merit	Description	Units
Throughput	Number of Users	Token
	Time Per Output Token	Second
Interactivity	1	Token
	Time Per Output Token	Second
Profitability (AI efficiency)	Throughput	Token / Second
	Cost X Power	\$ / Token
		X Watts

Batch	Human to Machine	Machine to Machine (Agentic)
-------	------------------	------------------------------



State of Art ~2026-27

14T parameters
Mixture of Experts
(MoE) w/ 32 experts

inference with:
32K/8K sequence length
5s Time to First Token

Optimizations:

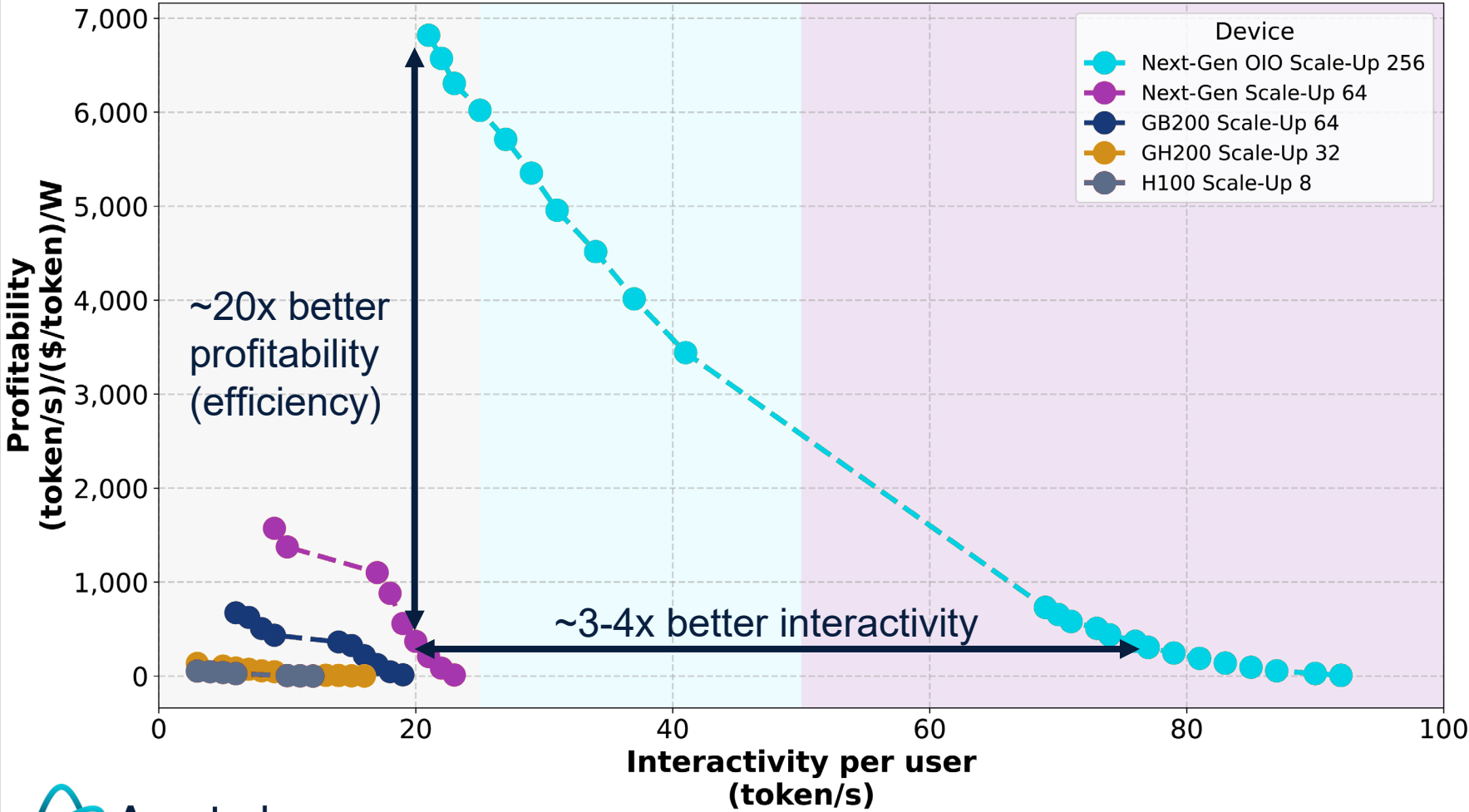
Tensor/Sequence parallel
Pipeline parallel
Expert parallel

Optical I/O impact on future GPT-X performance

Figure of Merit	Description	Units
Throughput	Number of Users	Token
	Time Per Output Token	Second
Interactivity	1	Token
	Time Per Output Token	Second
Profitability (AI efficiency)	Throughput	Token / Second
	Cost X Power	\$ / Token
		X Watts



AI efficiency (performance / \$ / W)



State of Art ~2026-27

14T parameters
Mixture of Experts
(MoE) w/ 32 experts

inference with:
32K/8K sequence length
5s Time to First Token

Optimizations:

Tensor/Sequence parallel
Pipeline parallel
Expert parallel



Optical I/O is an opportunity for major near-term AI hardware efficiency impact.

Source: Ayar Labs CTO
office team (Oct 2024)