

Single chip photonic deep neural network with accelerated training

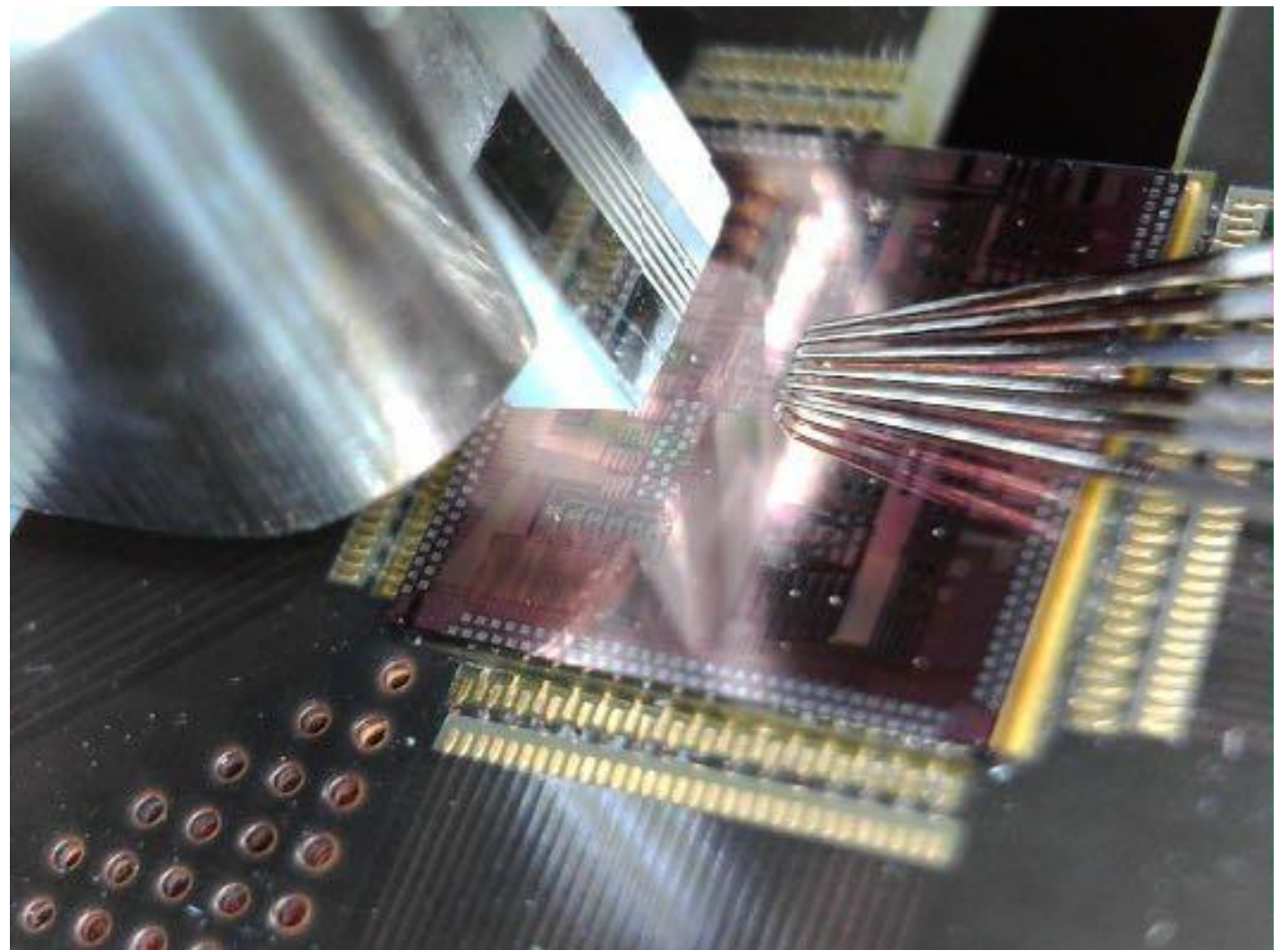
Saumil Bandyopadhyay¹, Alexander Sludds¹, Stefan Krastanov¹,
Ryan Hamerly^{1,2}, Nicholas Harris¹, Darius Bunandar¹, Matthew Streshinsky³,
Michael Hochberg⁴, Dirk Englund¹

¹Massachusetts Institute of Technology

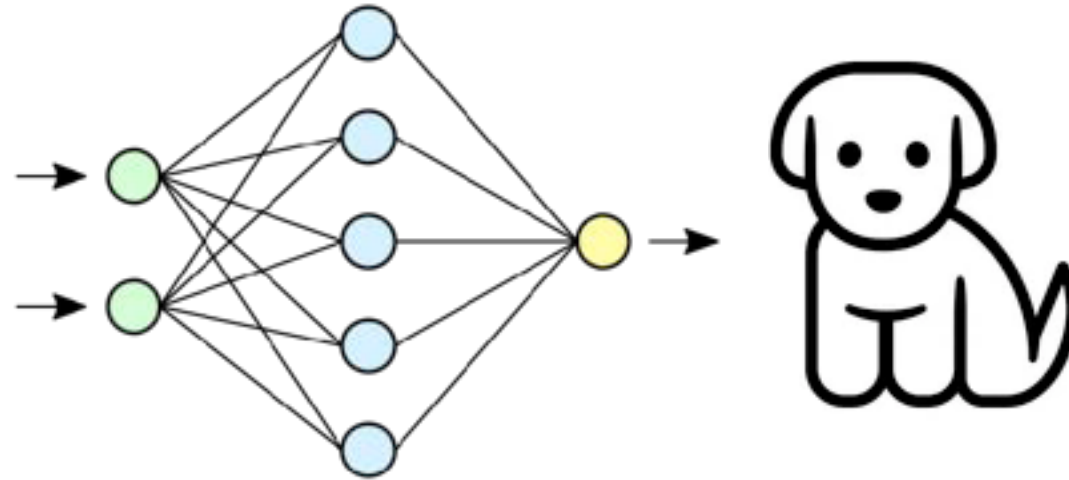
²PHI Labs, NTT Research

³Nokia Corporation

⁴Luminous Computing



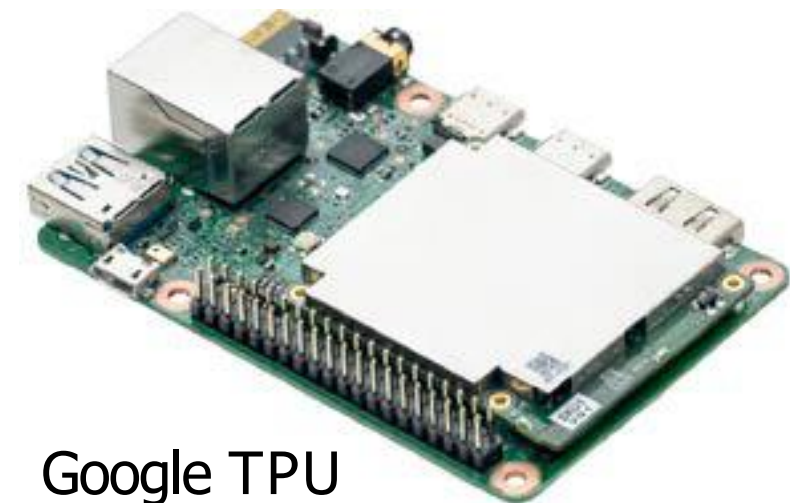
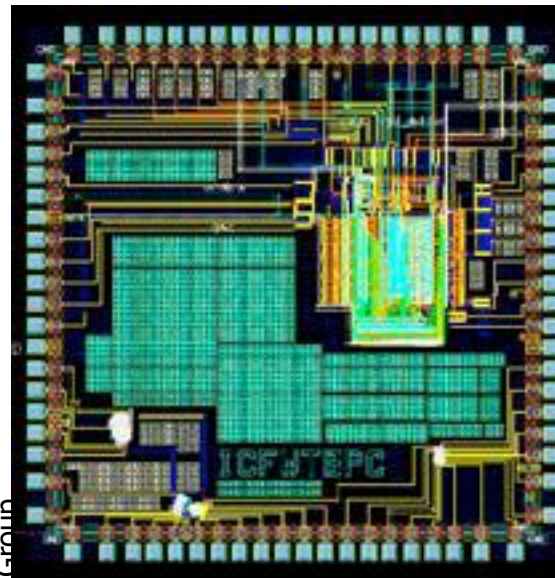
The “Cambrian explosion” of DNN hardware



Nvidia Tesla K80



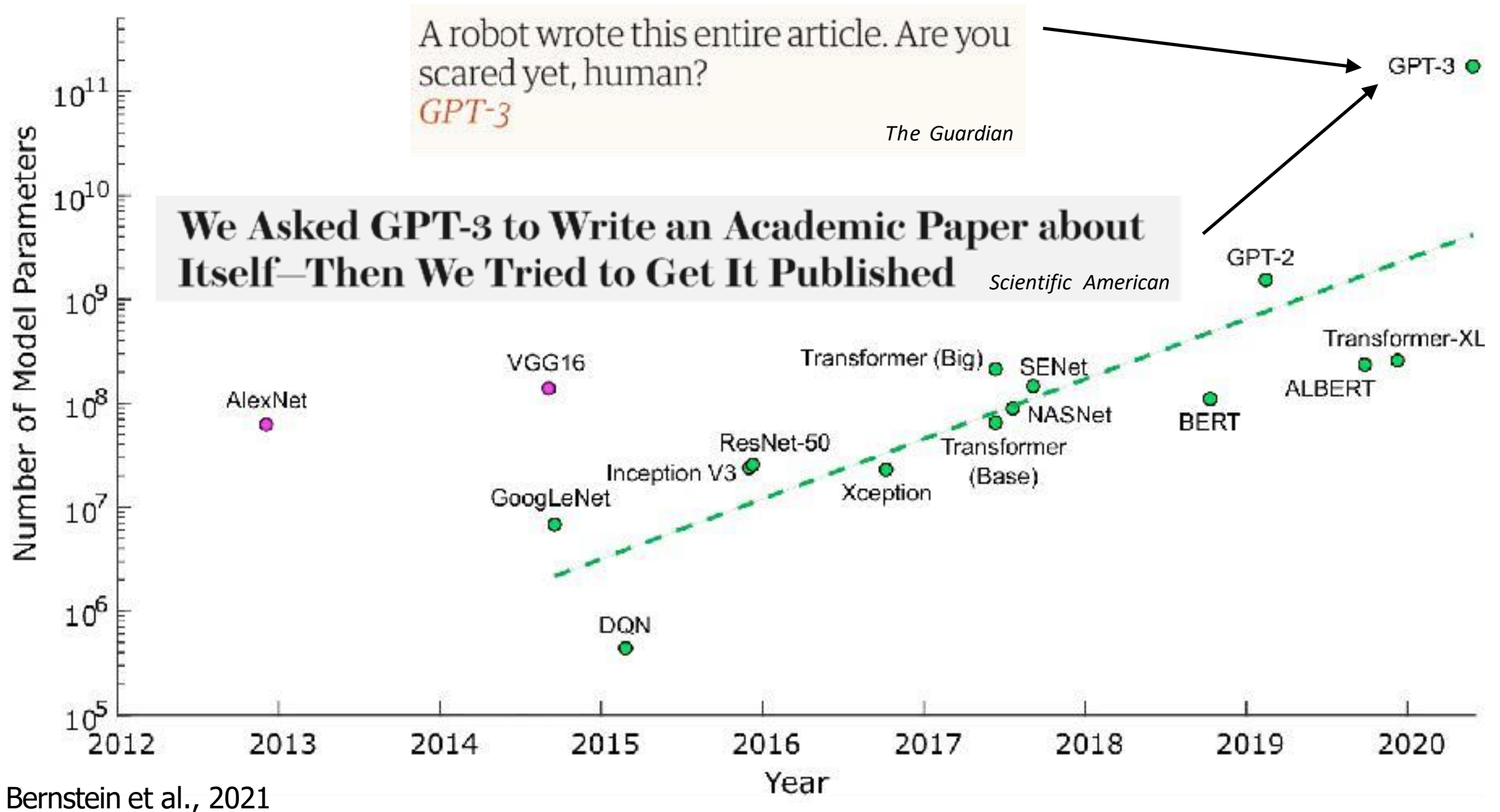
Waterloo CMOS Design and Reliability Group



Google TPU

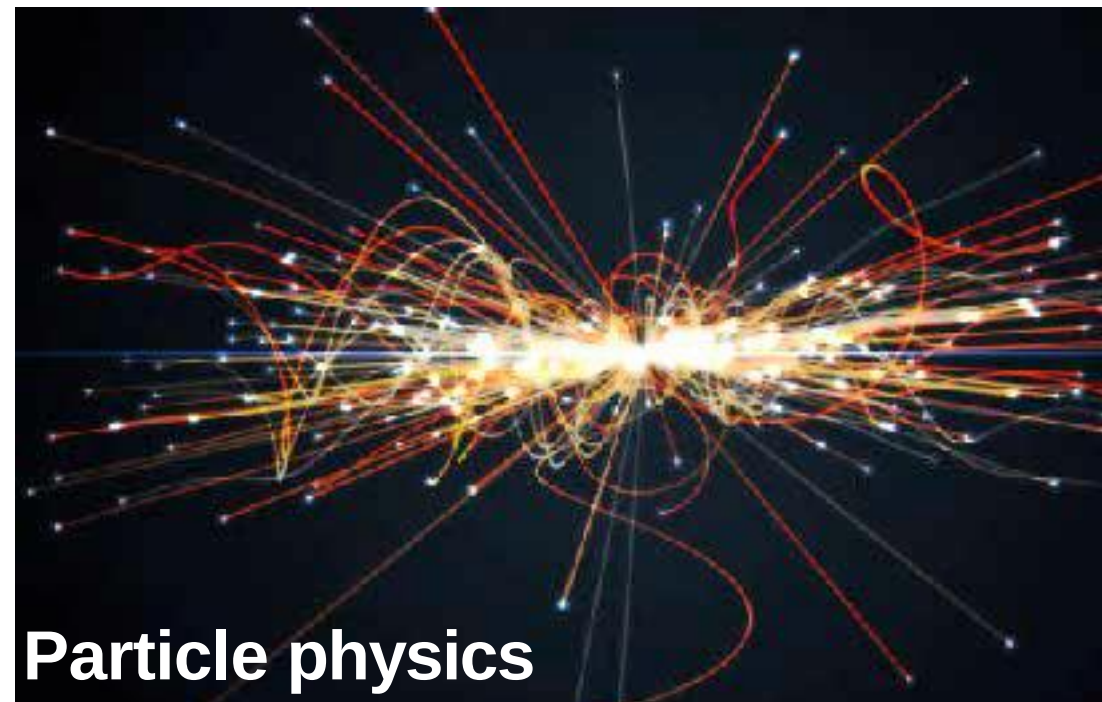
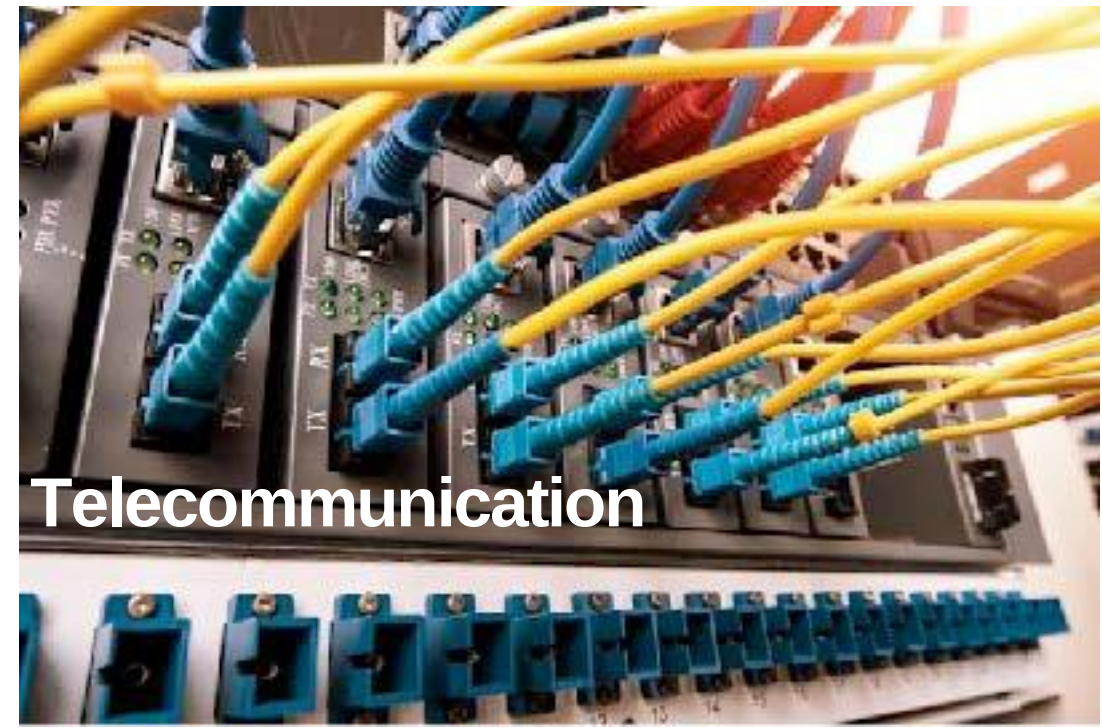


DNNs are getting better because they are getting *bigger*



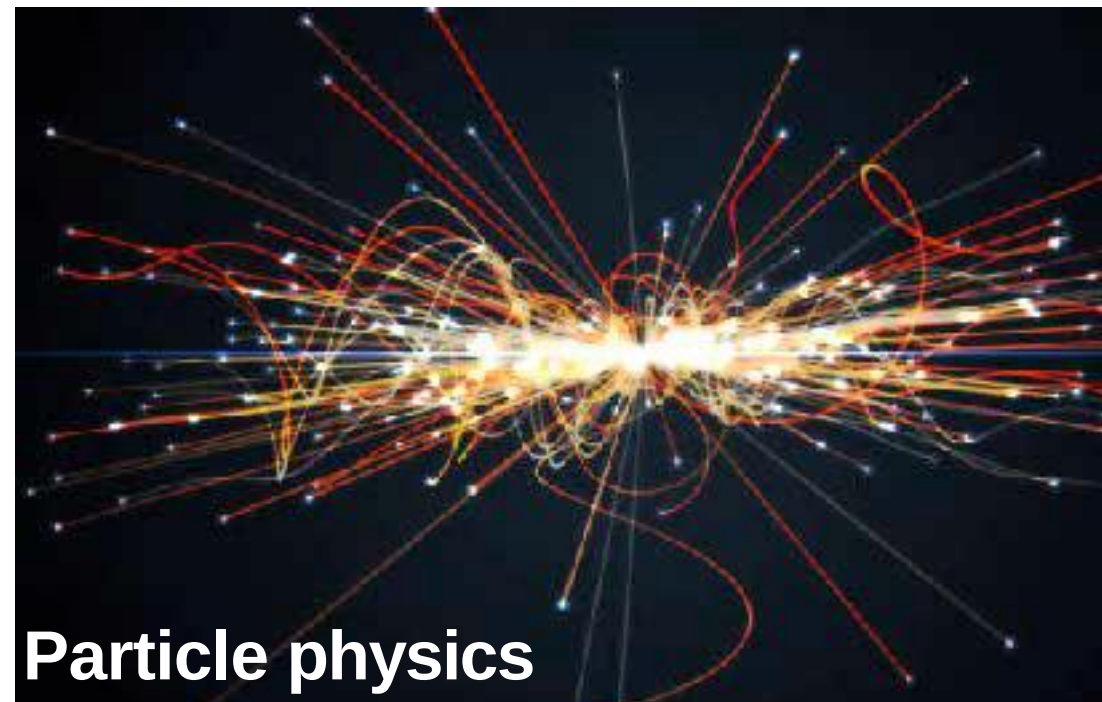
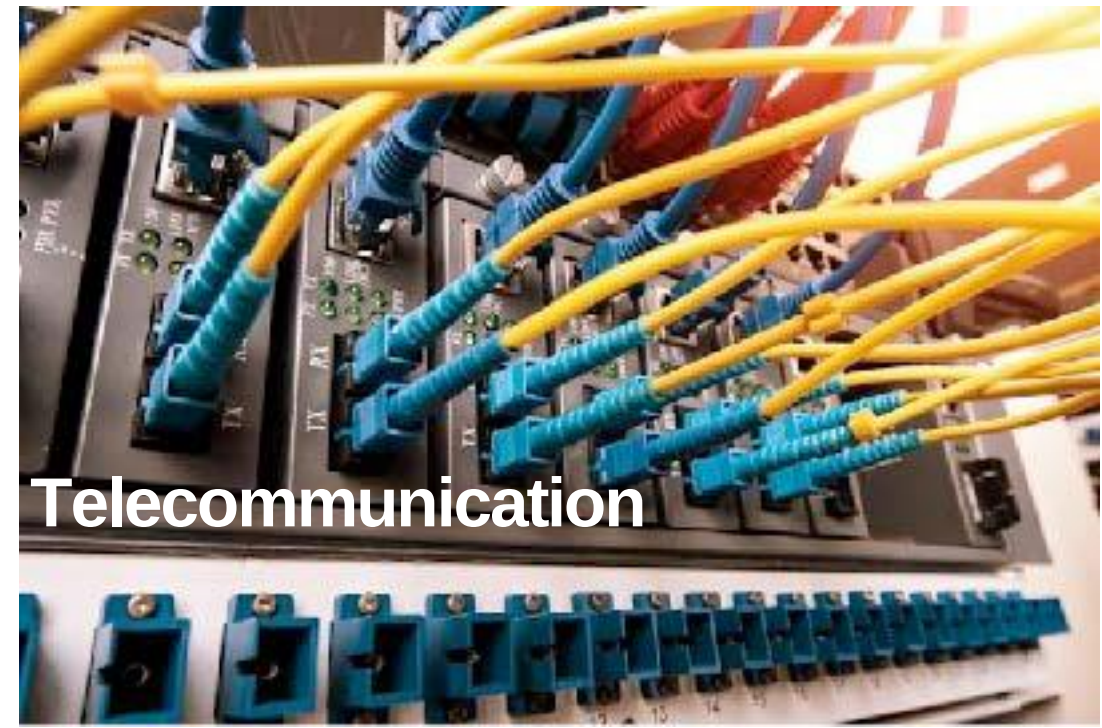
DNN scaling is **bottlenecked** by energy consumption, which is ~ 1 pJ/OP for CMOS.

New DNN hardware is needed for ultra-low latency processing



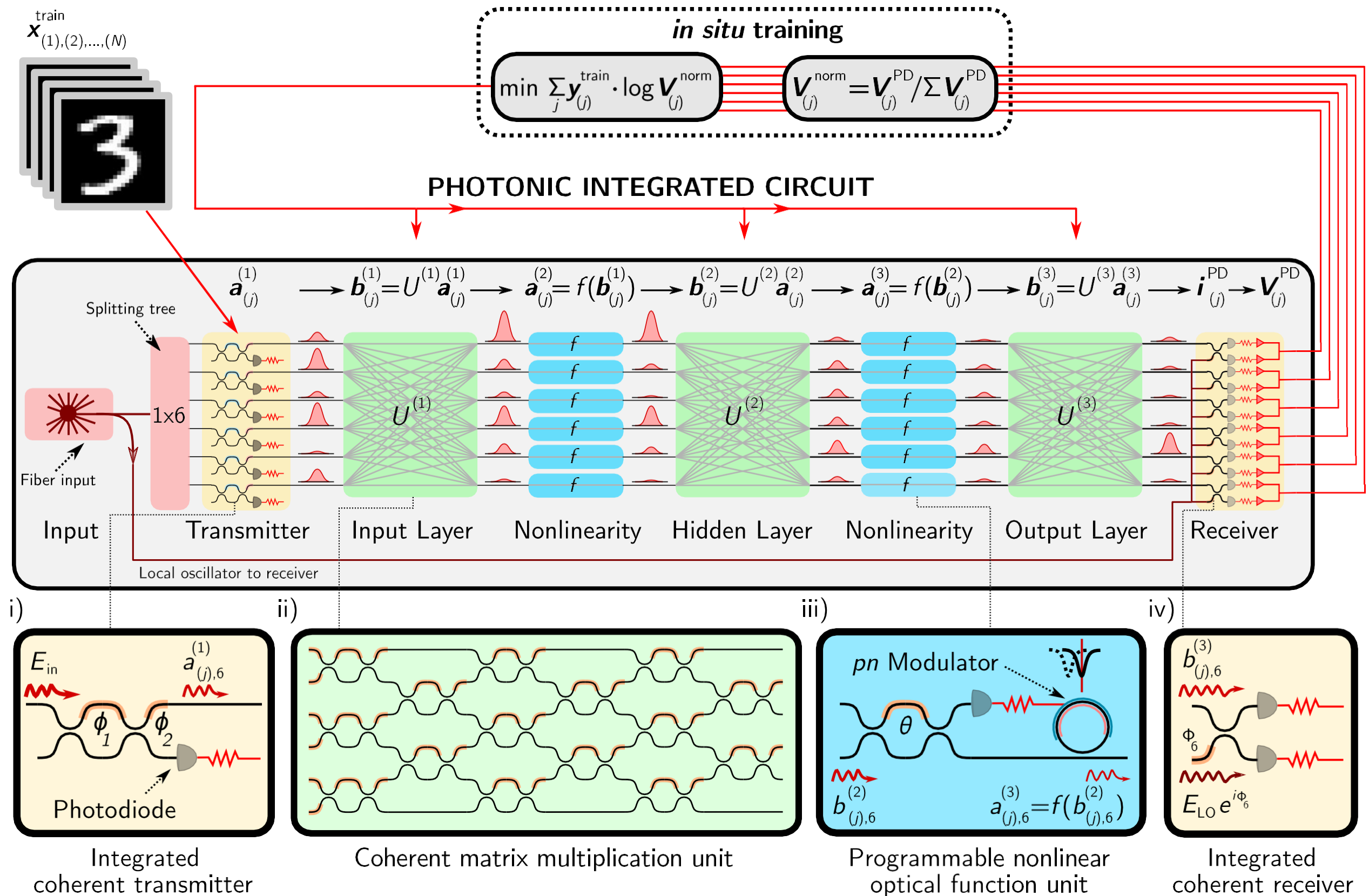
New applications require ultra-low latency (μs - ns) processing of (optical) data

New DNN hardware is needed for ultra-low latency processing

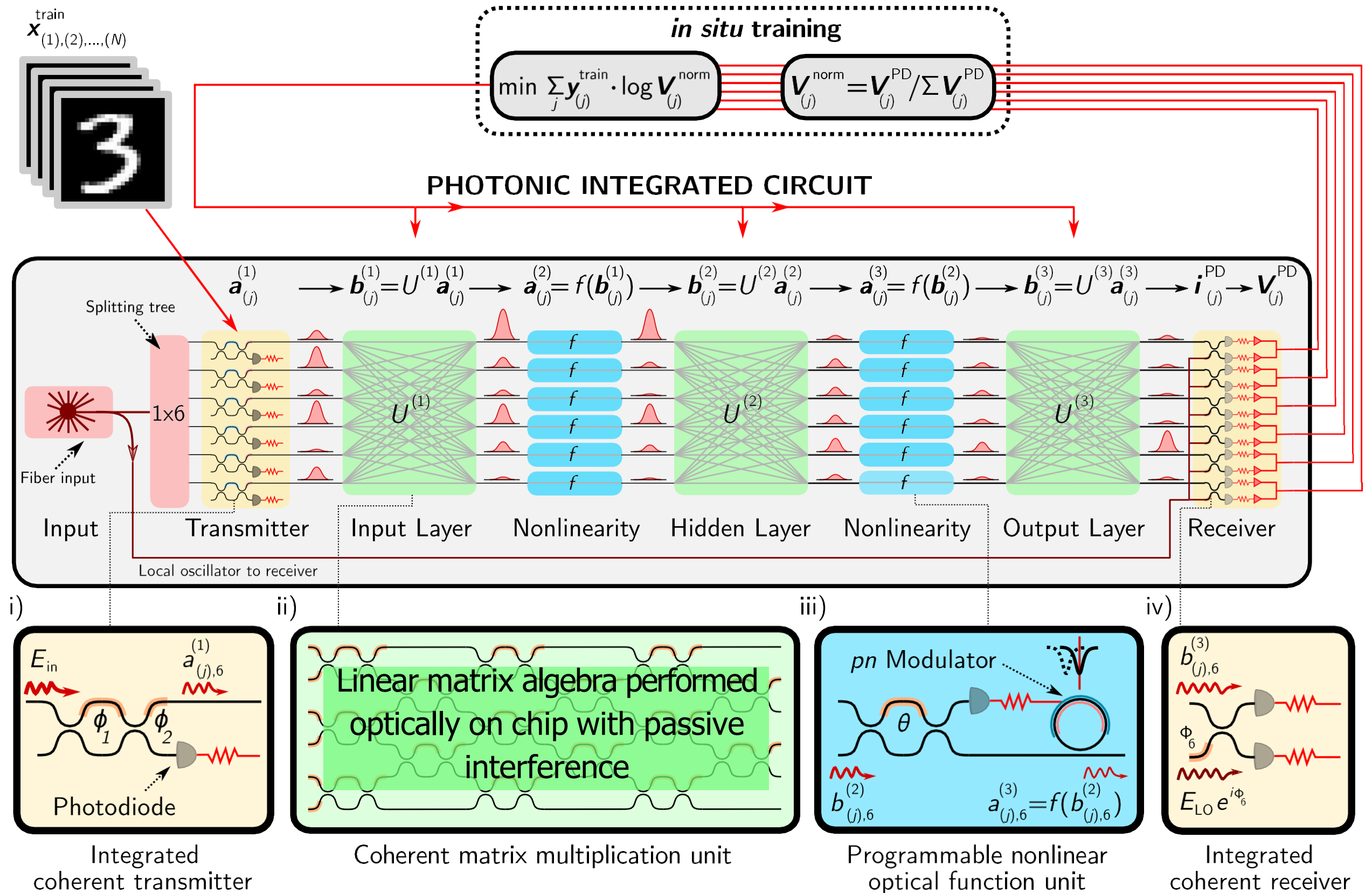


New applications require ultra-low latency (μs - ns) processing of (optical) data
 This requires time-of-flight (clockless) compute on a **single, integrated chip**

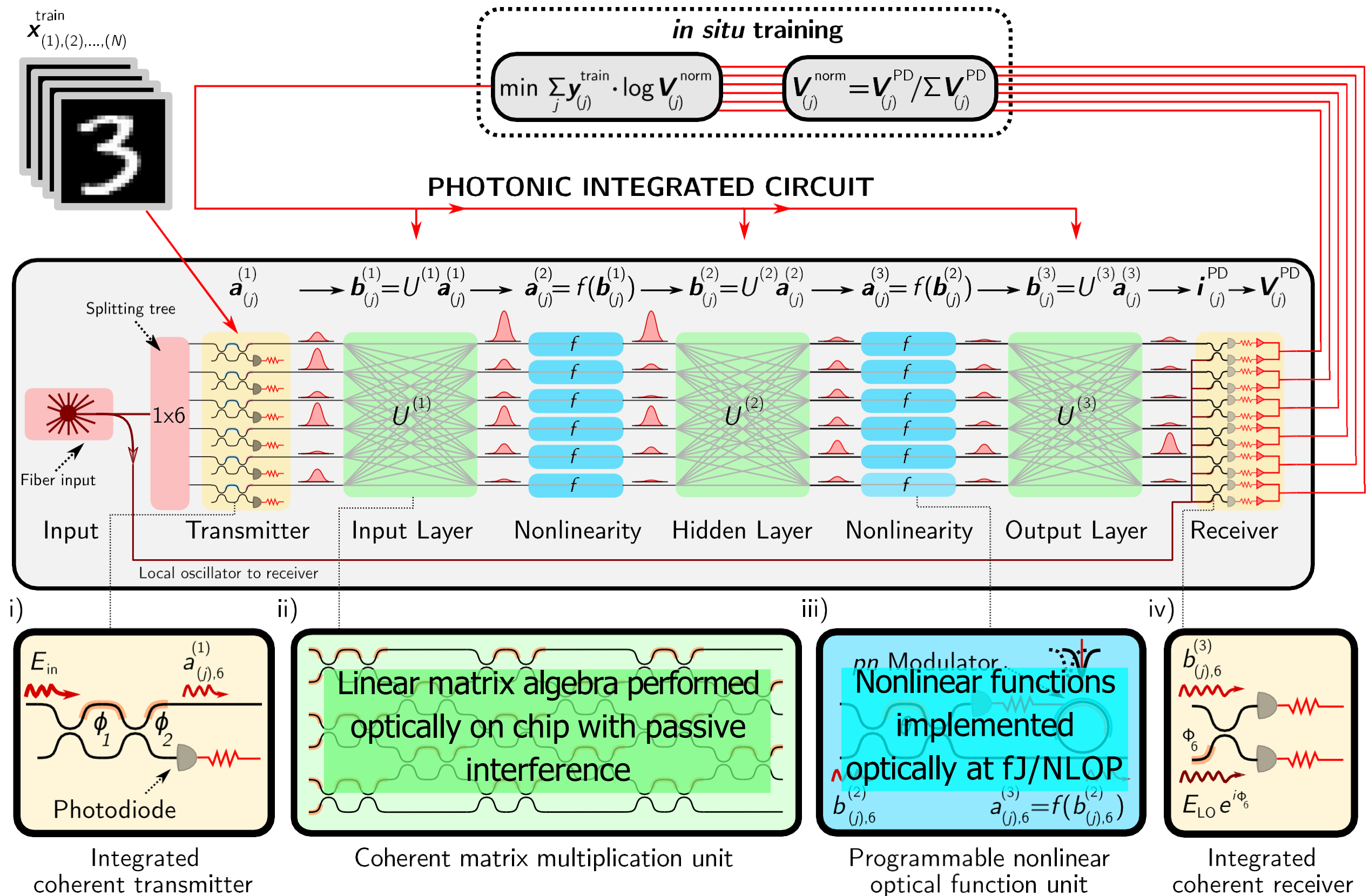
An end-to-end optical DNN processor on a single chip



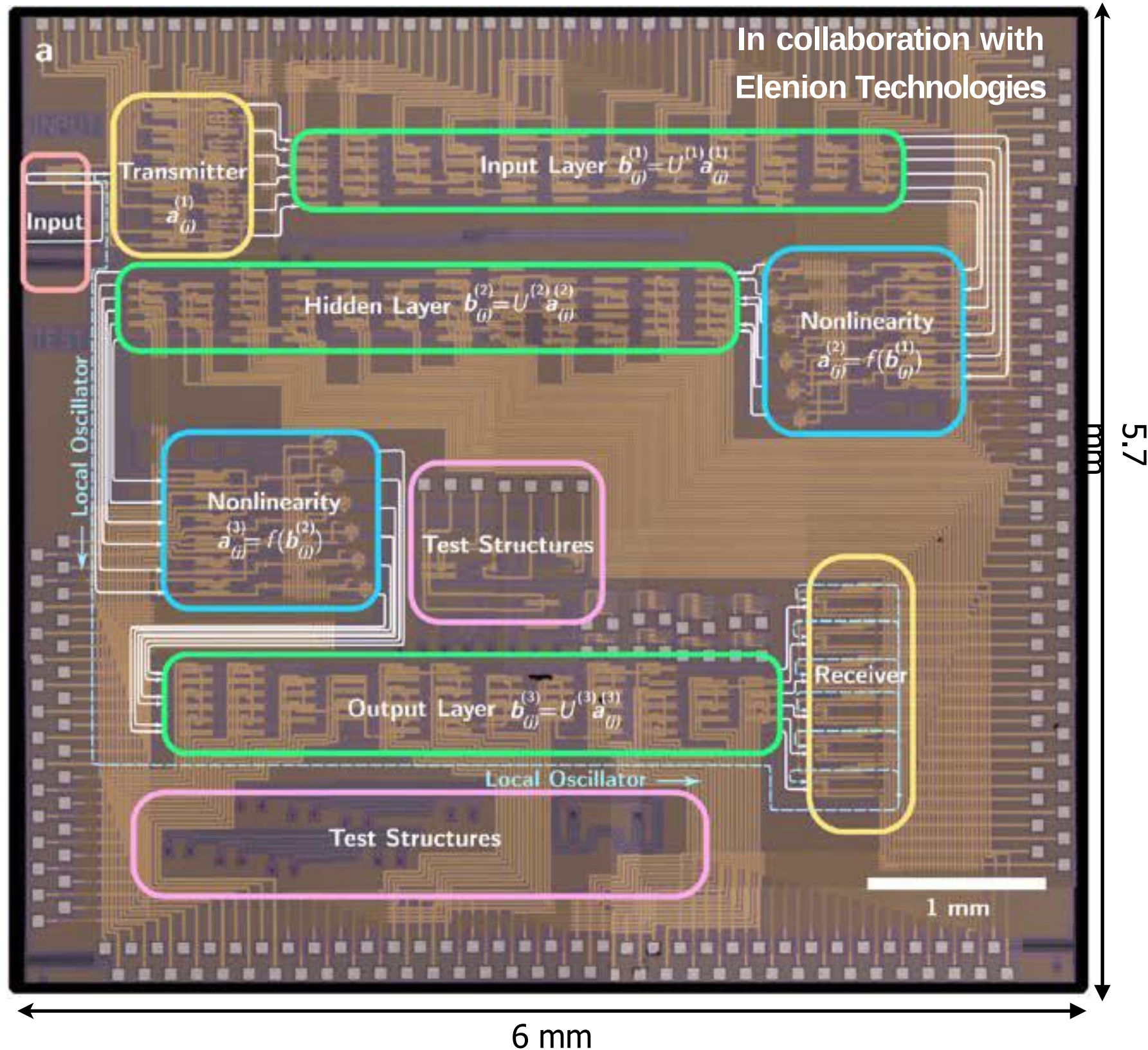
An end-to-end optical DNN processor on a single chip



An end-to-end optical DNN processor on a single chip



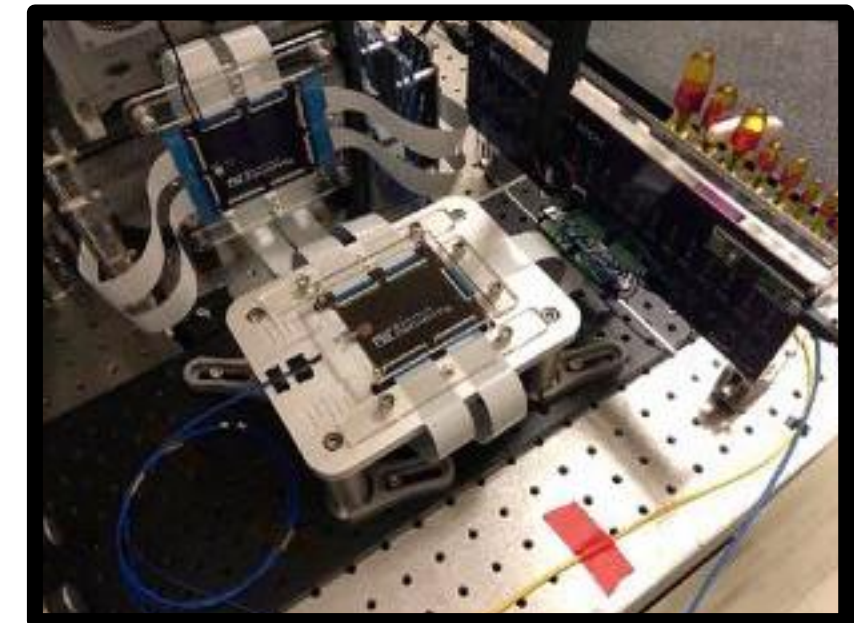
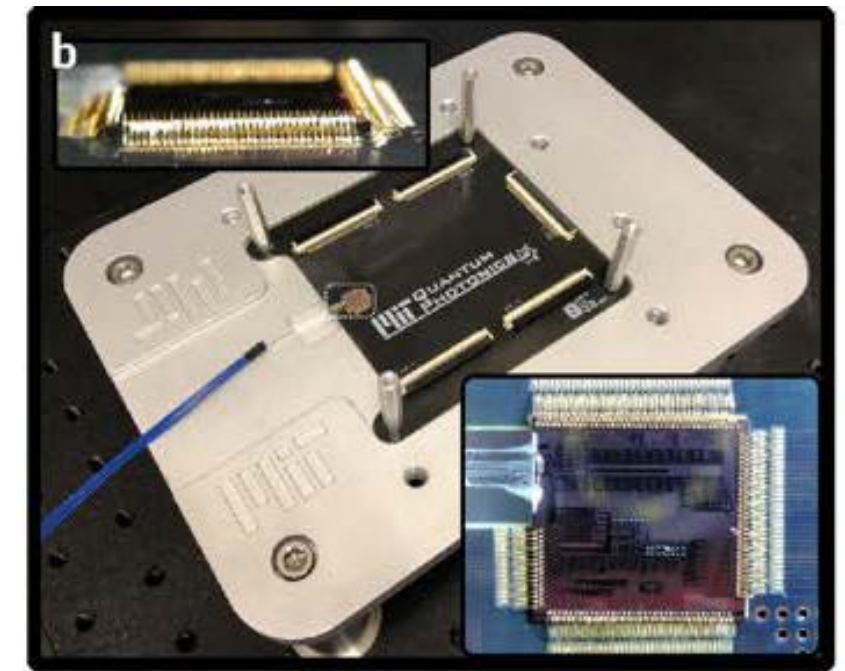
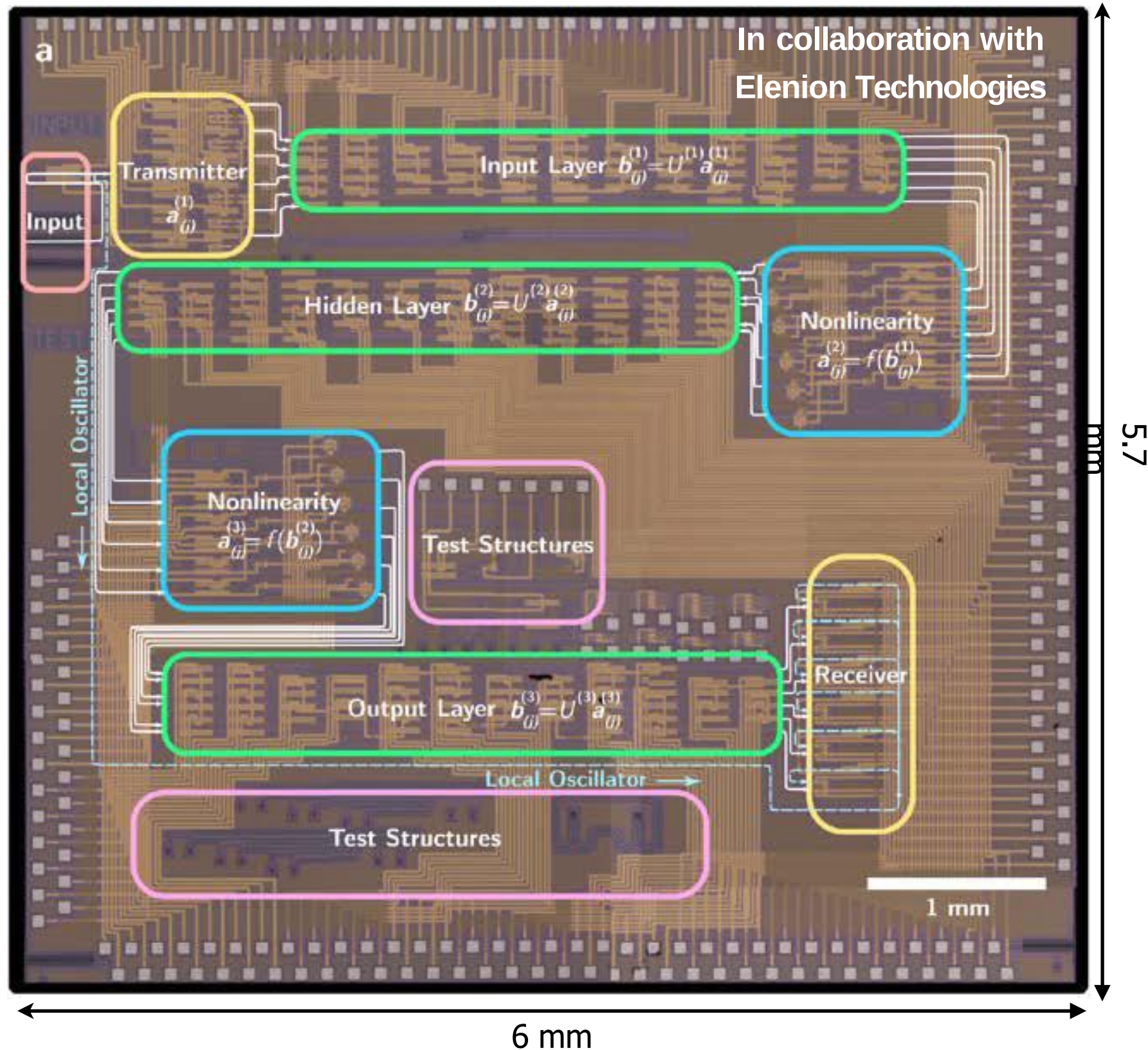
Application-specific photonic integrated circuit



132 model parameters on-chip, 90 layers of optical components

Preprint manuscript available at <https://arxiv.org/abs/2208.01623>

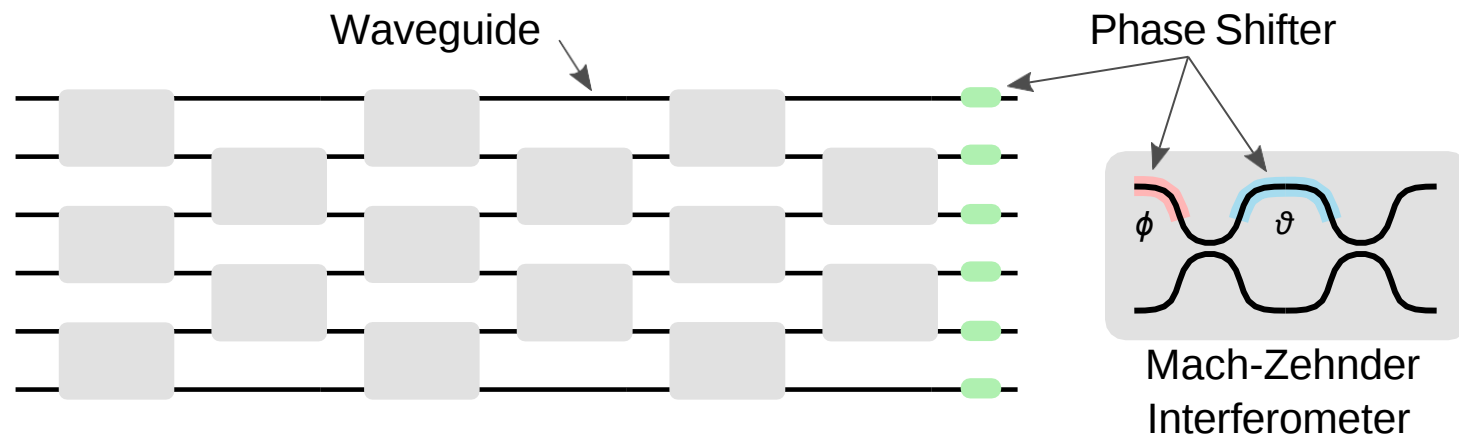
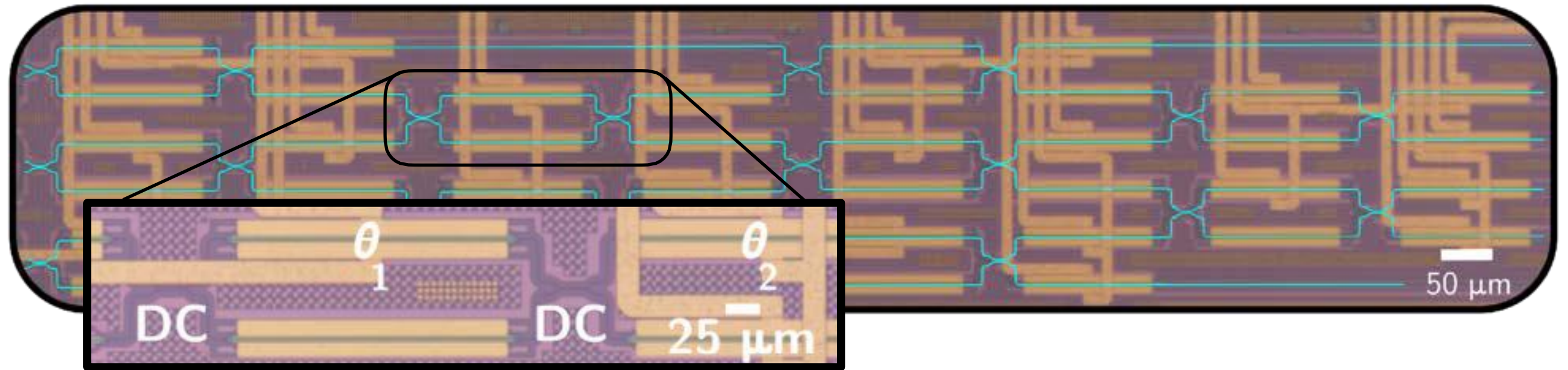
Application-specific photonic integrated circuit



132 model parameters on-chip, 90 layers of optical components

Preprint manuscript available at <https://arxiv.org/abs/2208.01623>

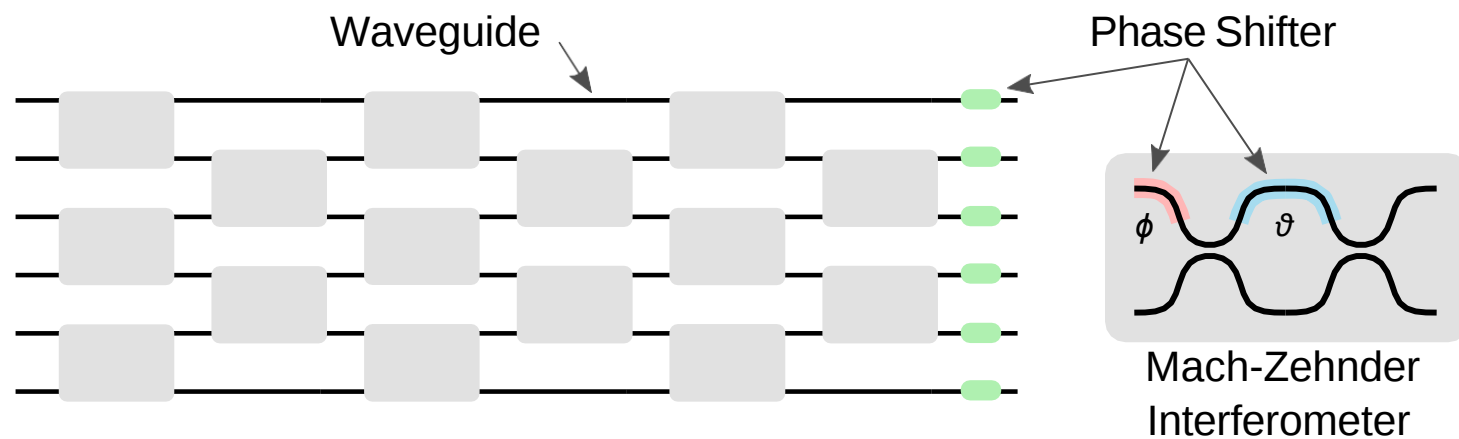
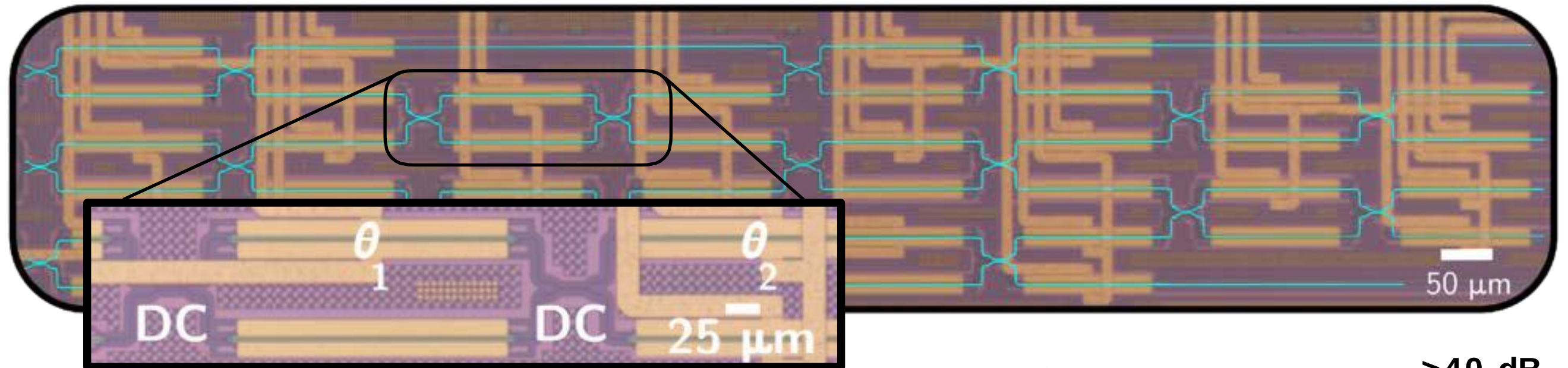
Coherent matrix multiplication unit



$$T_{ij}(\theta, \phi) = ie^{i\theta/2} \begin{bmatrix} e^{i\phi} \sin(\theta/2) & \cos(\theta/2) \\ e^{i\phi} \cos(\theta/2) & -\sin(\theta/2) \end{bmatrix} \quad U = D \prod T_{ij}(\theta, \phi)$$

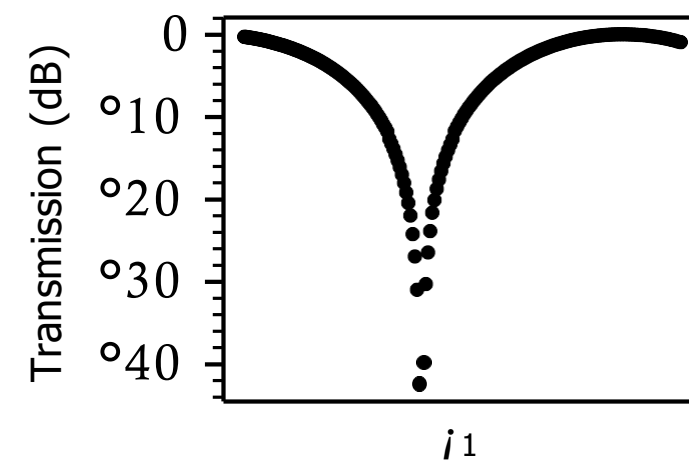
Programmable meshes of Mach-Zehnder interferometers compute linear algebra through passive optical interference

Coherent matrix multiplication unit

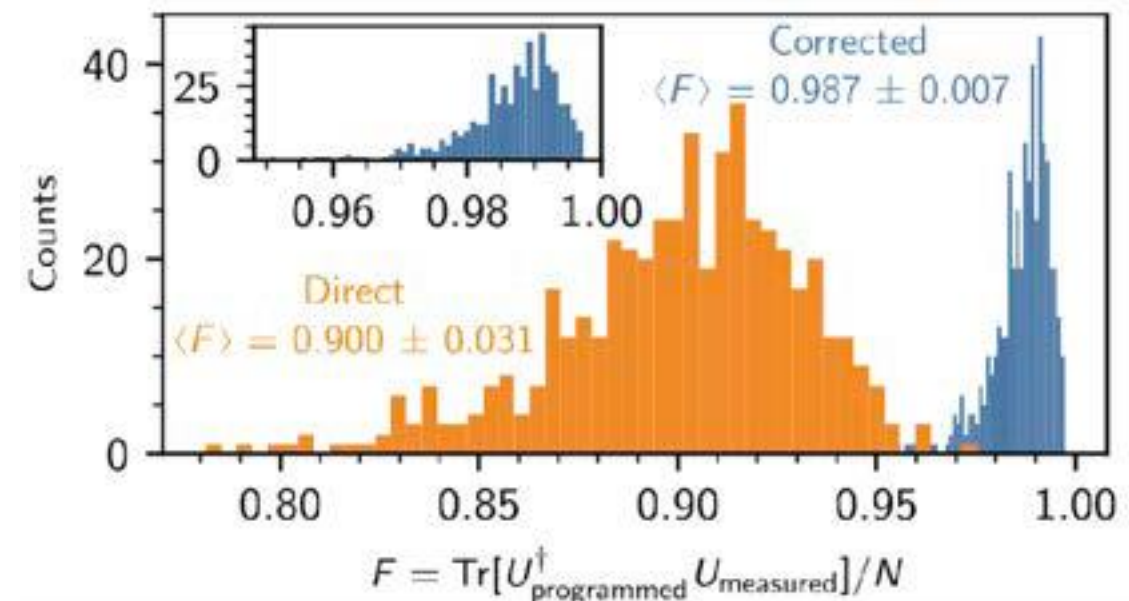


$$T_{ij}(\theta, \phi) = ie^{i\theta/2} \begin{bmatrix} e^{i\phi} \sin(\theta/2) & \cos(\theta/2) \\ e^{i\phi} \cos(\theta/2) & -\sin(\theta/2) \end{bmatrix} \quad U = D \prod T_{ij}(\theta, \phi)$$

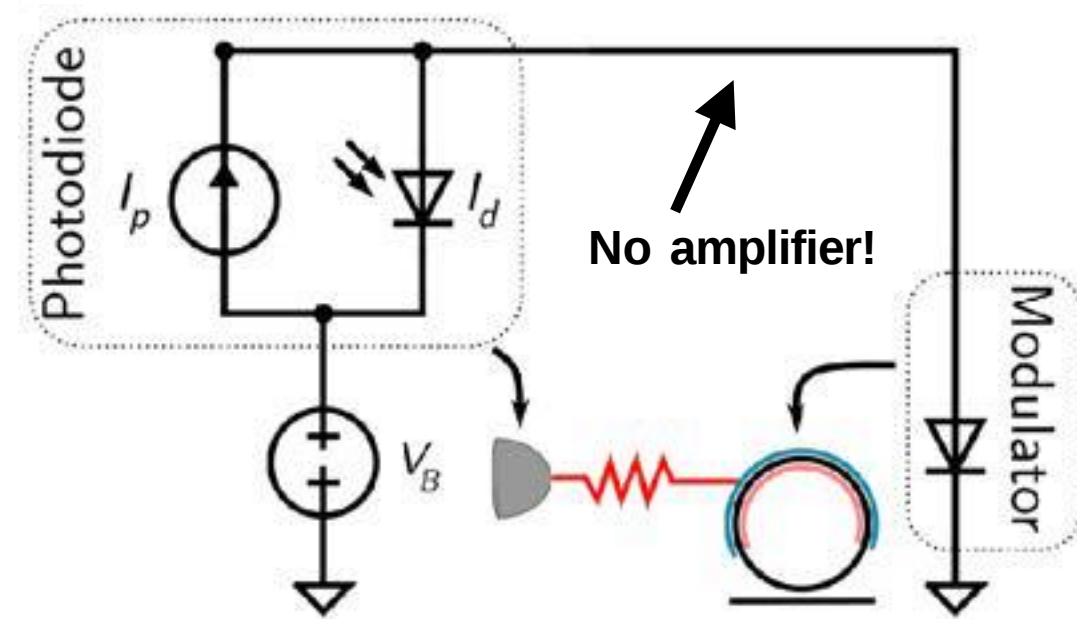
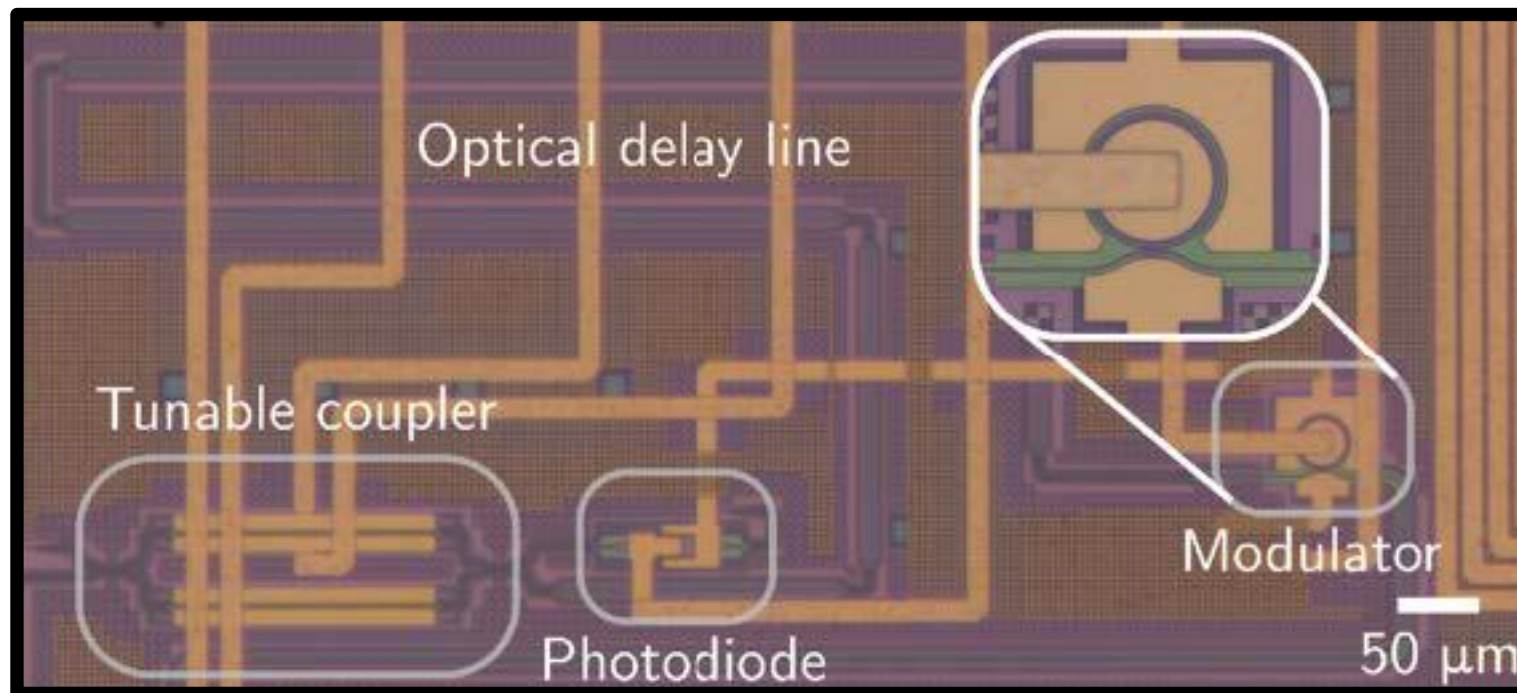
Programmable meshes of Mach-Zehnder interferometers compute linear algebra through passive optical interference



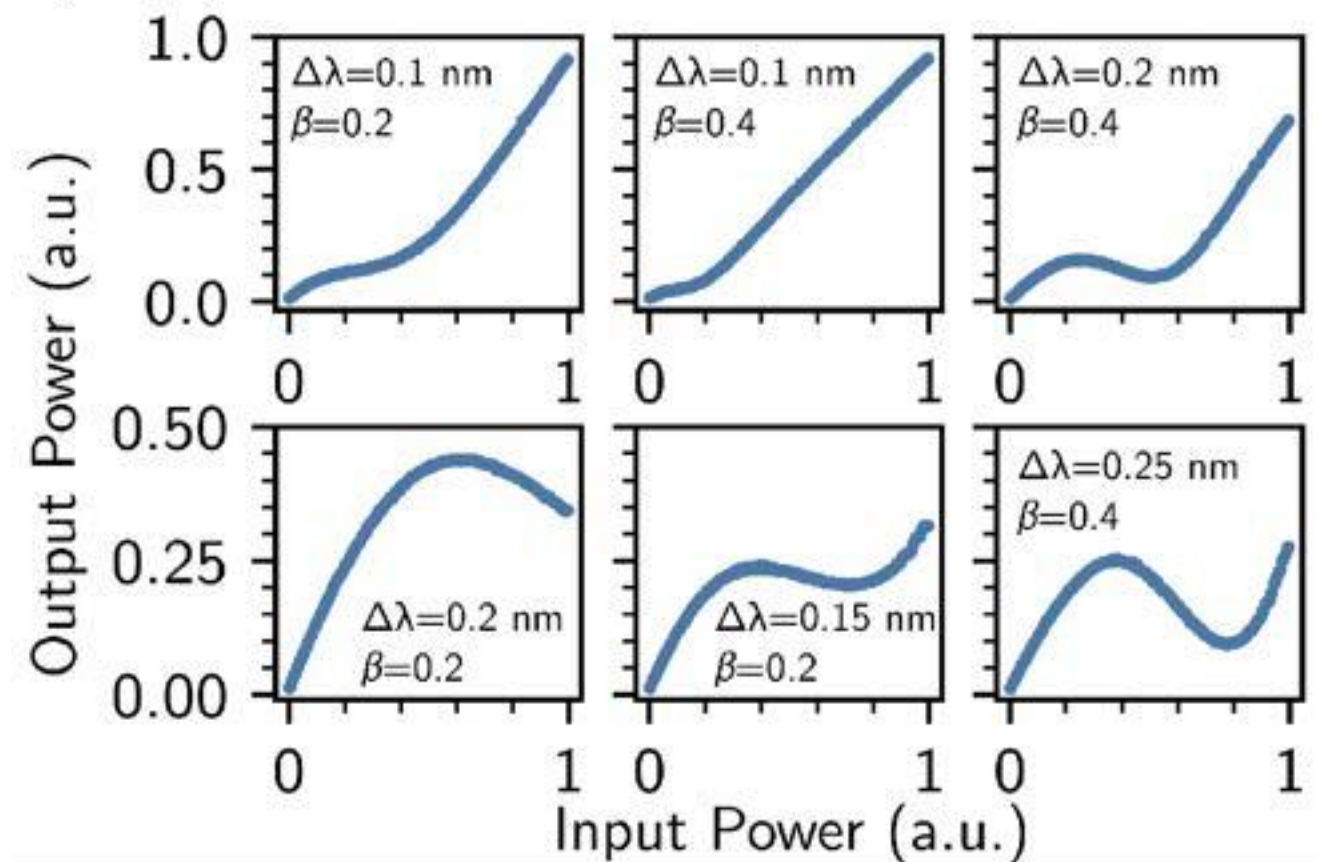
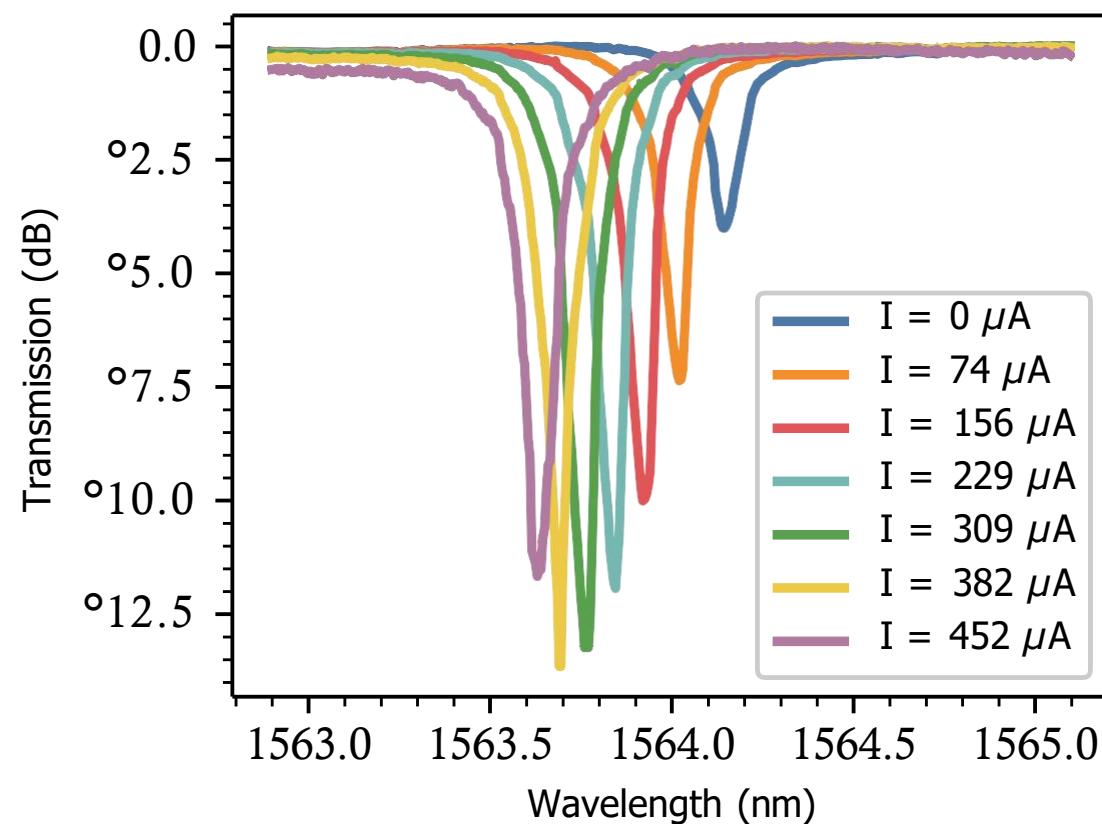
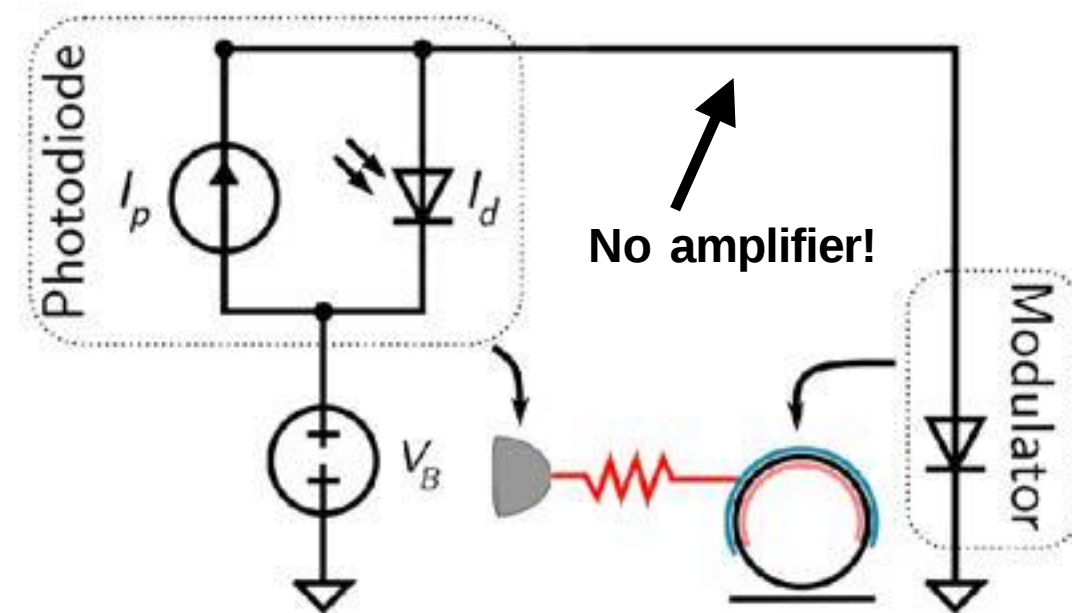
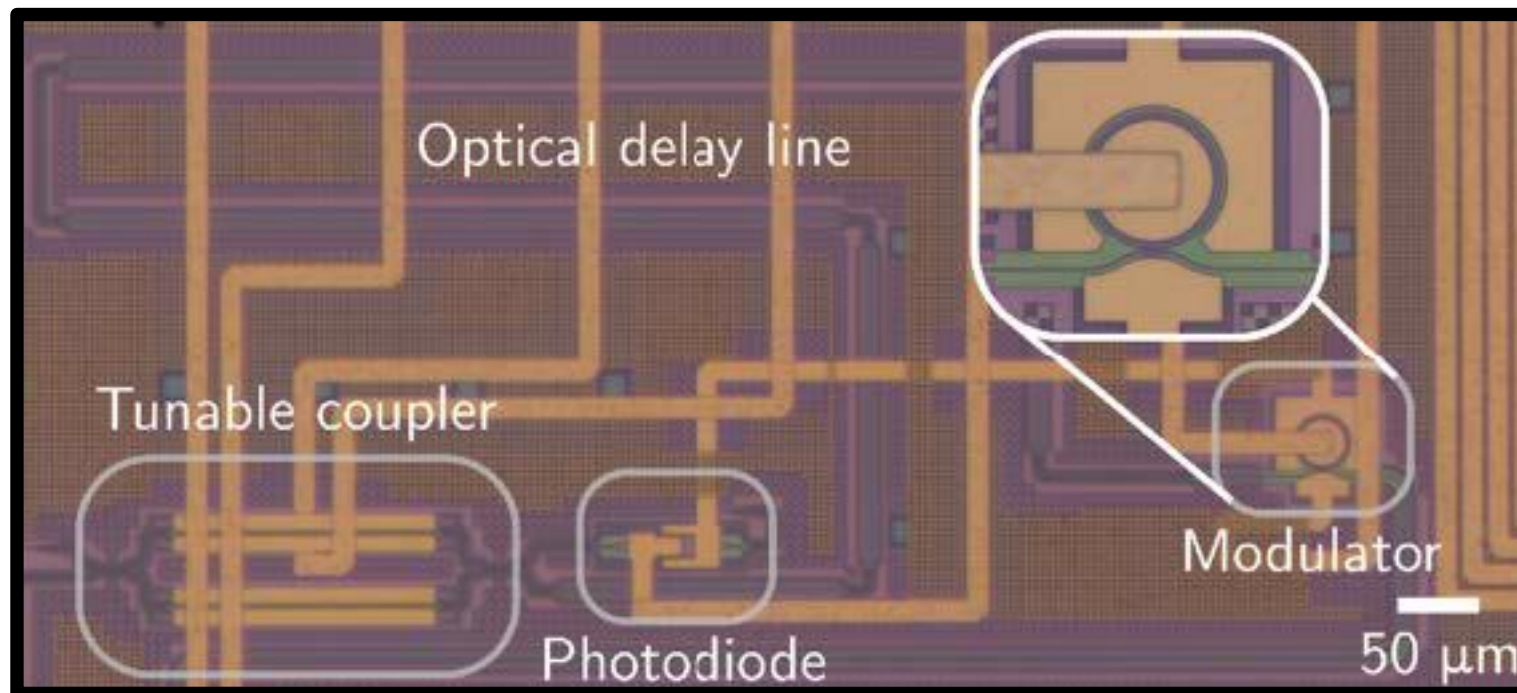
**>40 dB
optical
extinction,
>13 bits of
precision on
inputs**



Nonlinear optical function unit

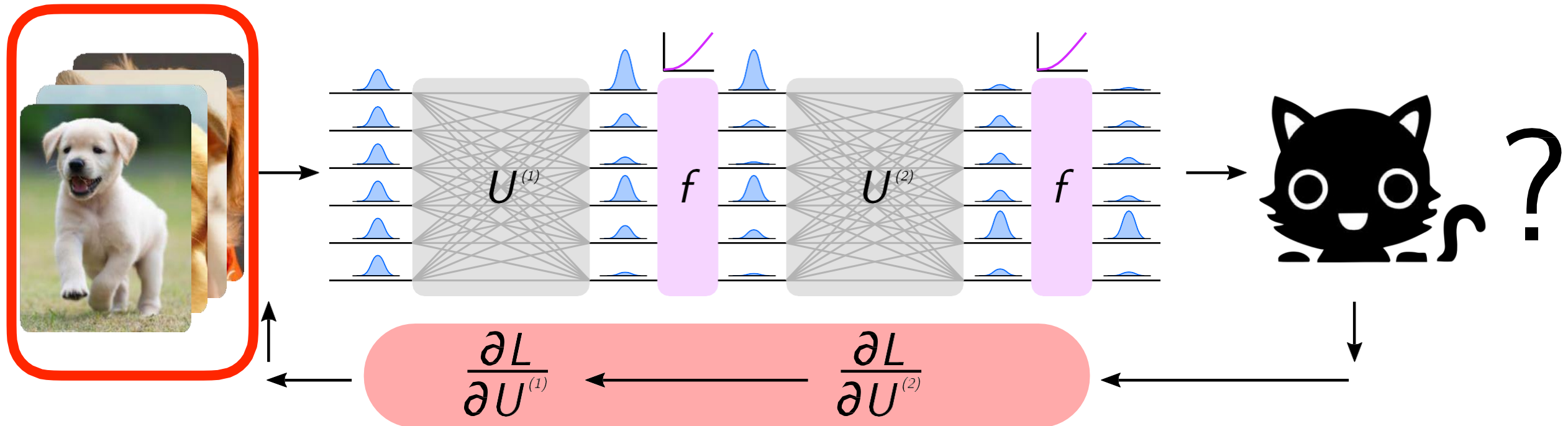


Nonlinear optical function unit



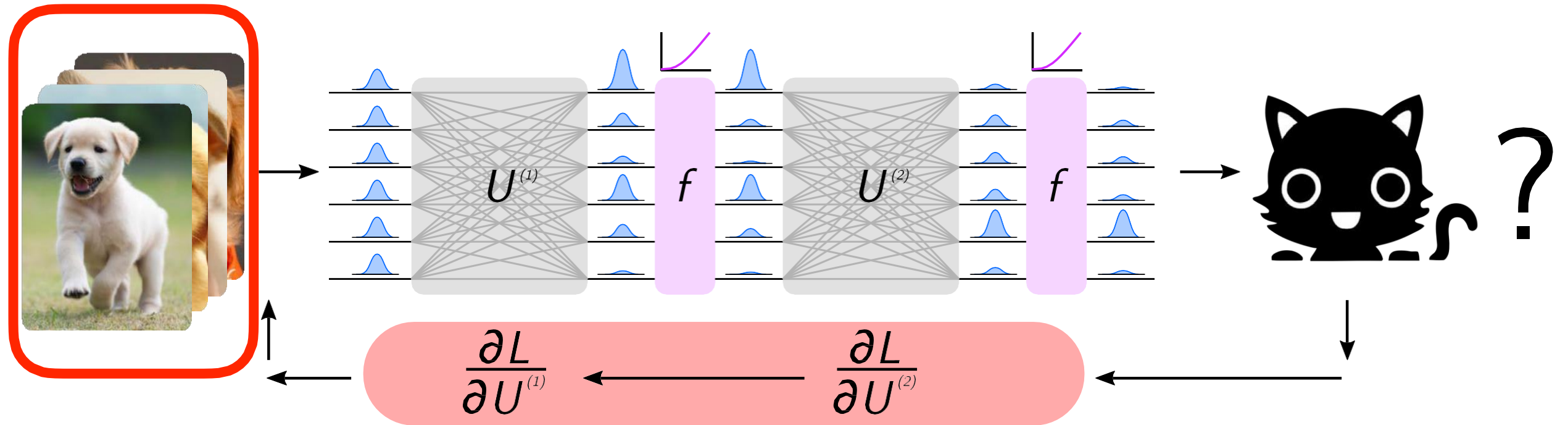
Coherent optical nonlinear functions without digitization or amplification

Model training benefits from optical acceleration



Model training also requires repeated forward inference.

Model training benefits from optical acceleration

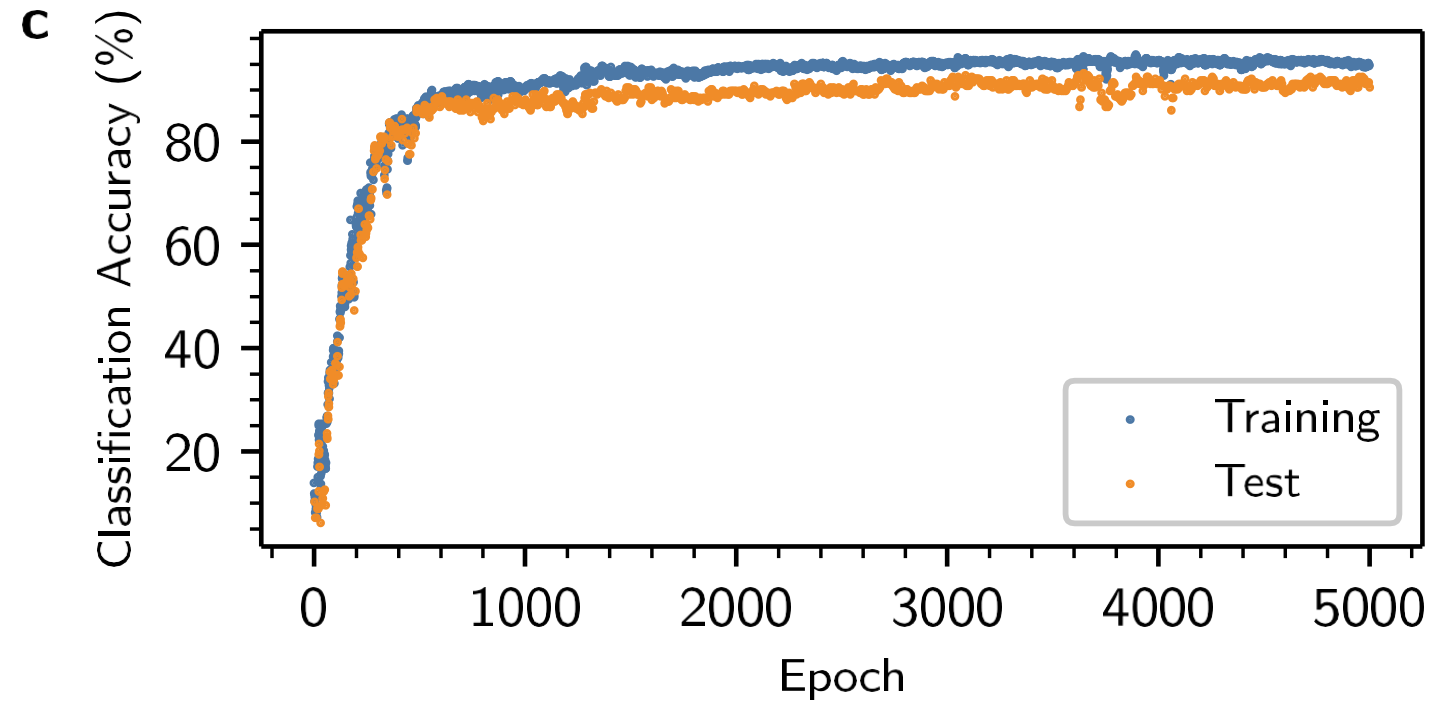
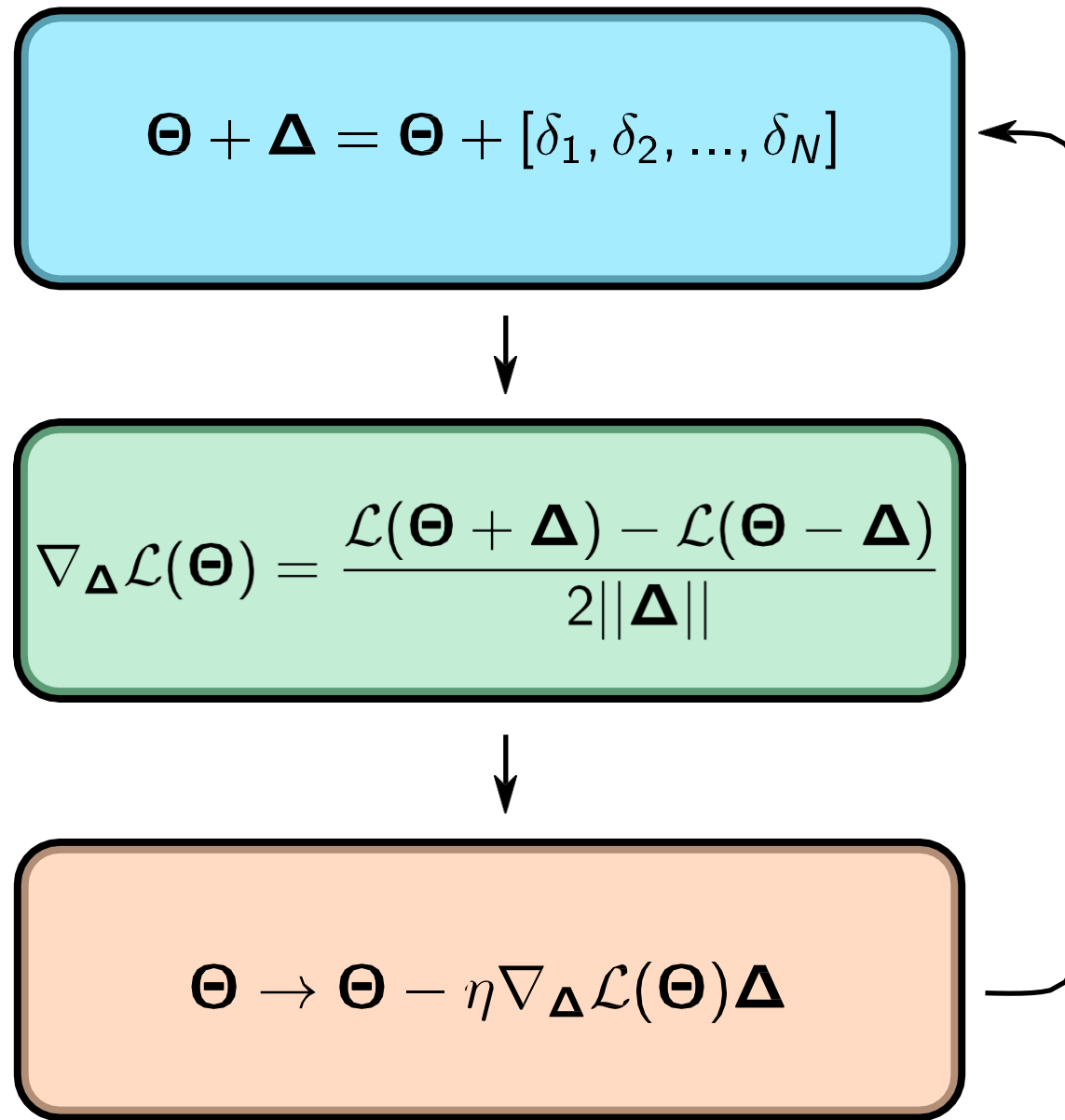


Model training also requires repeated forward inference.



Optically acceleration can enable low-latency model training.

In situ training on optical hardware



Training Accuracy: 96.5%							Test Accuracy: 92.7%								
Expected	ae	96	2	0	0	0	2	ae	94	0	0	0	0	0	6
	aw	0	100	0	0	0	0	aw	0	96	2	2	0	0	
	uw	0	0	99	1	0	0	uw	0	2	90	4	0	4	
	er	2	0	0	98	0	0	er	0	0	0	100	0	0	
	iy	6	0	0	0	92	2	iy	14	0	0	0	84	2	
	ih	0	0	0	0	6	94	ih	0	0	0	0	8	92	
		ae	aw	uw	er	iy	ih		ae	aw	uw	er	iy	ih	
Predicted							Predicted								

Forward inference is **optically accelerated** during *in situ* training, which achieves the **same accuracy** as a digitally trained system.

Summary

- Single-shot optical inference in a deep neural network on a photonic chip with time-of-flight (**500 ps**) limited latency

Summary

- Single-shot optical inference in a deep neural network on a photonic chip with time-of-flight (**500 ps**) limited latency
- All-optical processing of data without digitization between layers, eliminating the latency introduced by optical-to-electronic conversion
 - Coherent matrix multiplication in the optical domain
 - Coherent optical nonlinear functions without electronic amplification

Summary

- Single-shot optical inference in a deep neural network on a photonic chip with time-of-flight (**500 ps**) limited latency
- All-optical processing of data without digitization between layers, eliminating the latency introduced by optical-to-electronic conversion
 - Coherent matrix multiplication in the optical domain
 - Coherent optical nonlinear functions without electronic amplification
- Entirely fabricated in a commercial foundry

Summary

- Single-shot optical inference in a deep neural network on a photonic chip with time-of-flight (**500 ps**) limited latency
- All-optical processing of data without digitization between layers, eliminating the latency introduced by optical-to-electronic conversion
 - Coherent matrix multiplication in the optical domain
 - Coherent optical nonlinear functions without electronic amplification
- Entirely fabricated in a commercial foundry
- *In situ* training takes advantage of near-instantaneous inference, reducing latency and power consumption of model training.

Acknowledgments



Alexander Sludds



Prof. Stefan Krastanov



Dr. Ryan Hamerly



Dr. Nicholas Harris



Dr. Darius Bunandar



Dr. Matthew Streshinsky



Dr. Michael Hochberg



Prof. Dirk Englund

