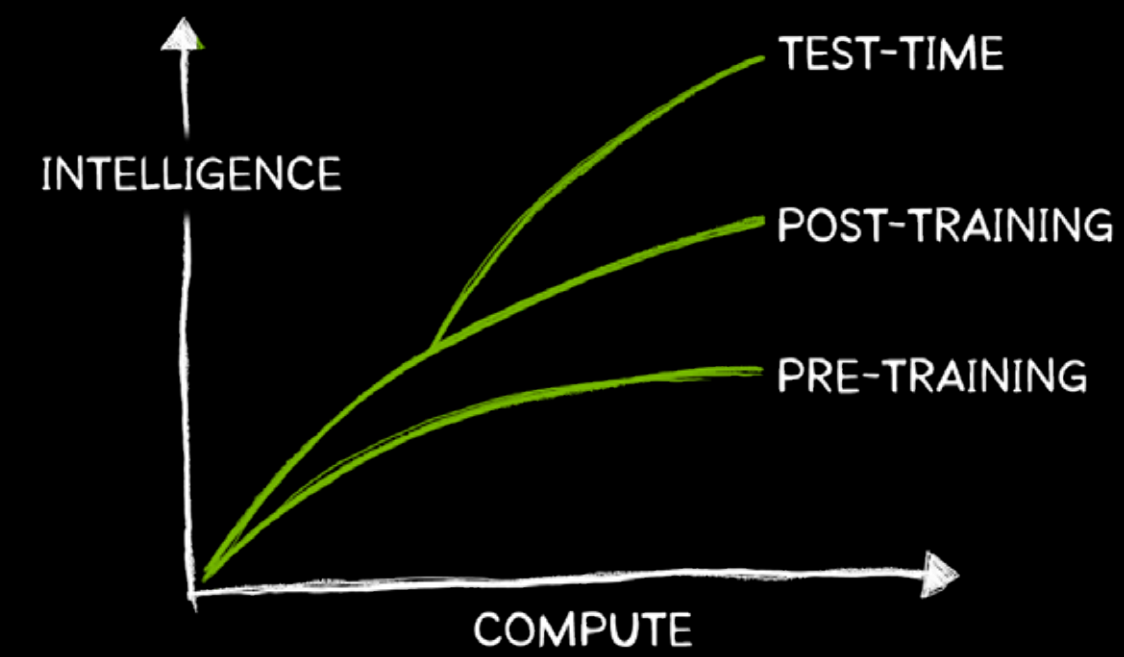




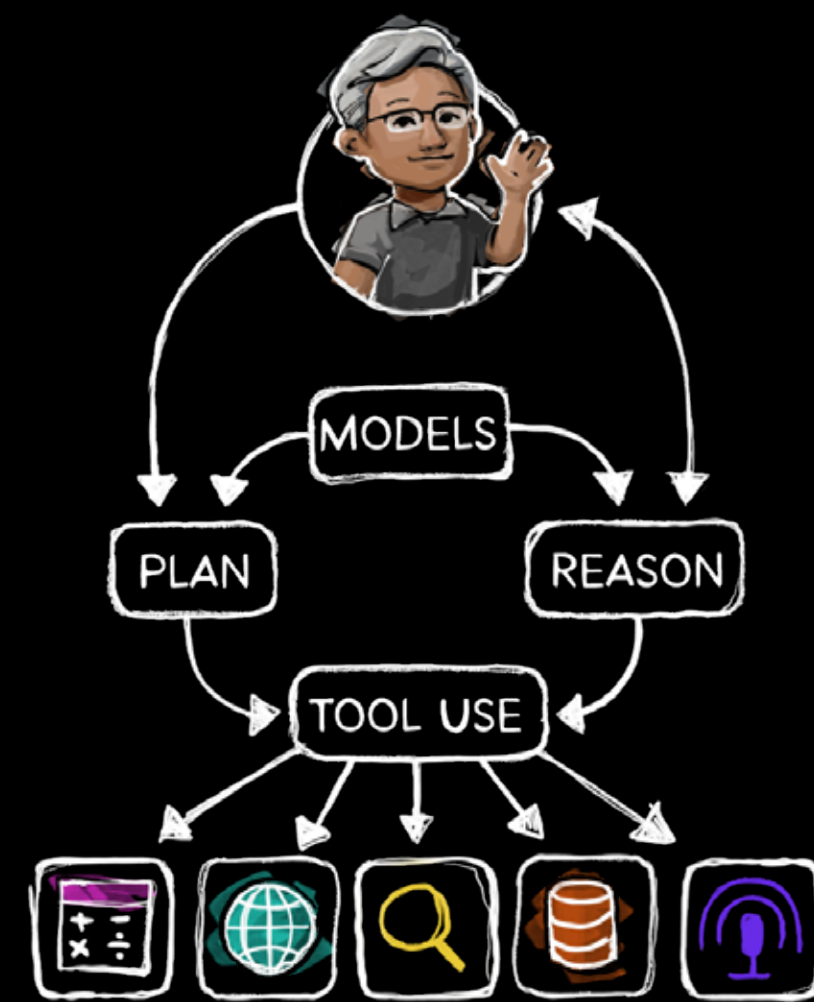
AI Scales Beyond LLMs

“Everything Everywhere All at Once”



3 SCALING LAWS

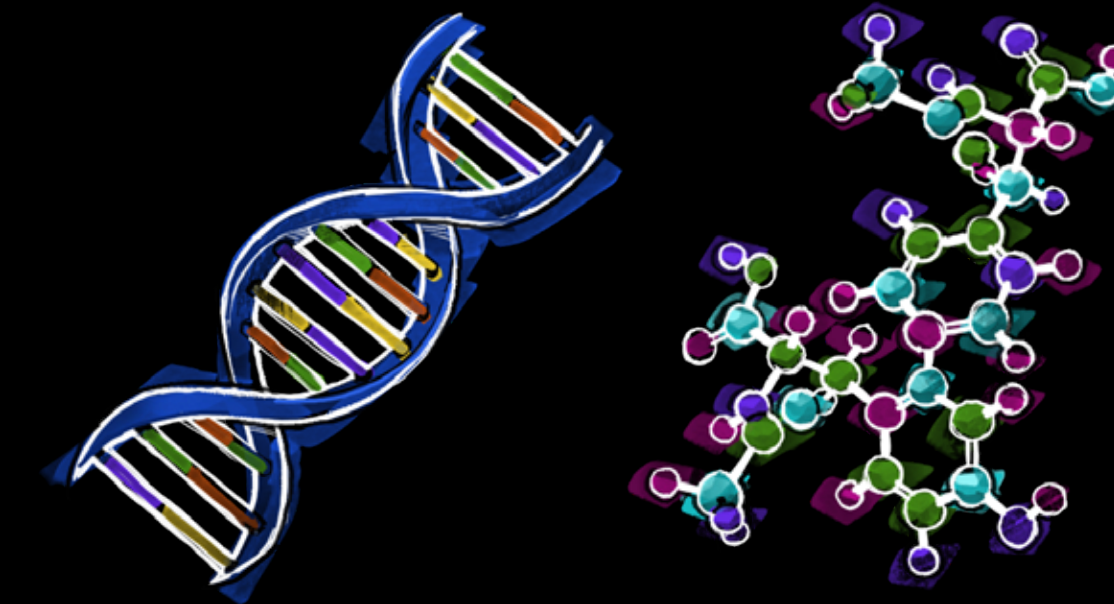
COMPUTE IS
DATA



AI BECOMES
AGENTIC



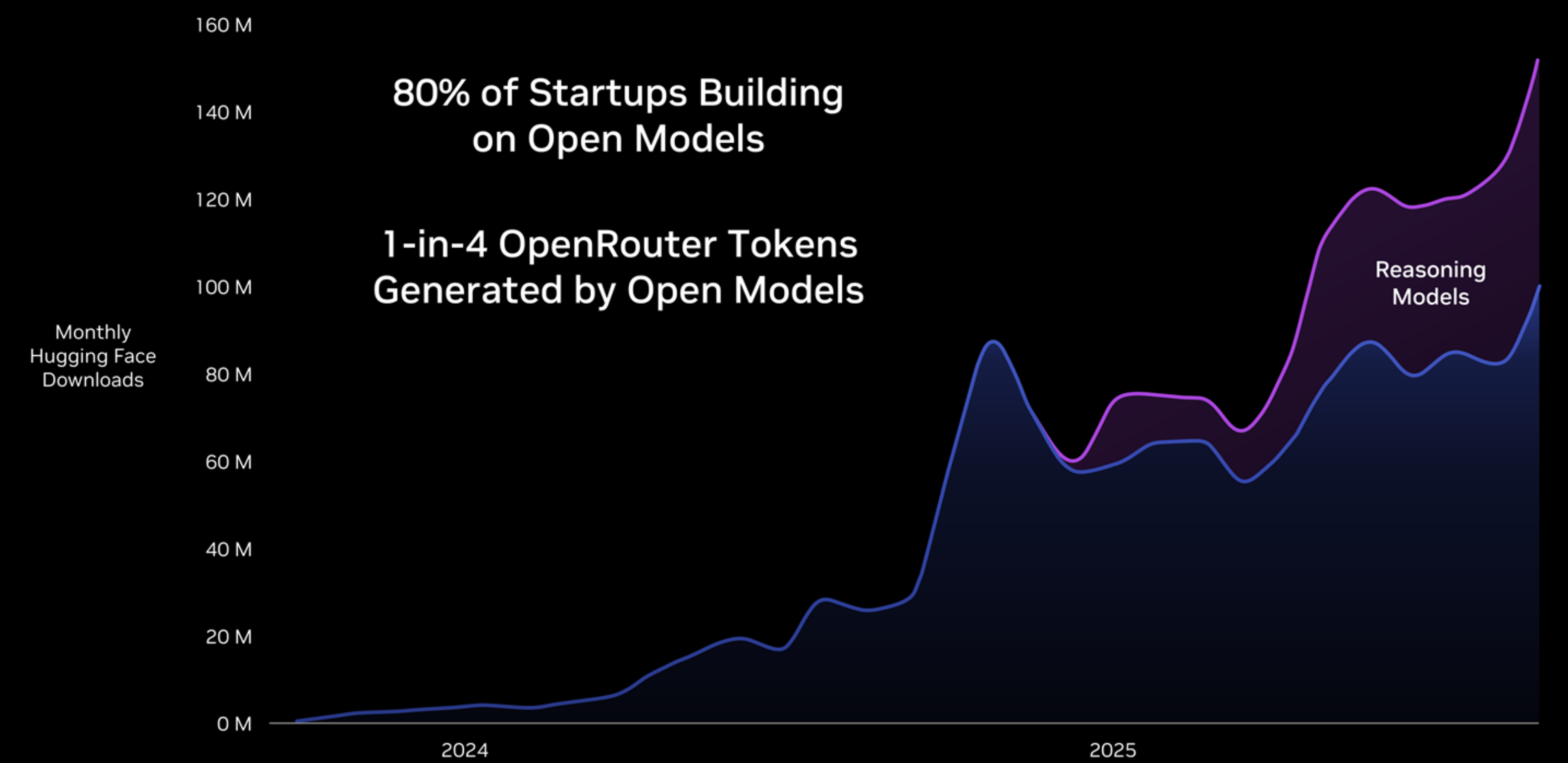
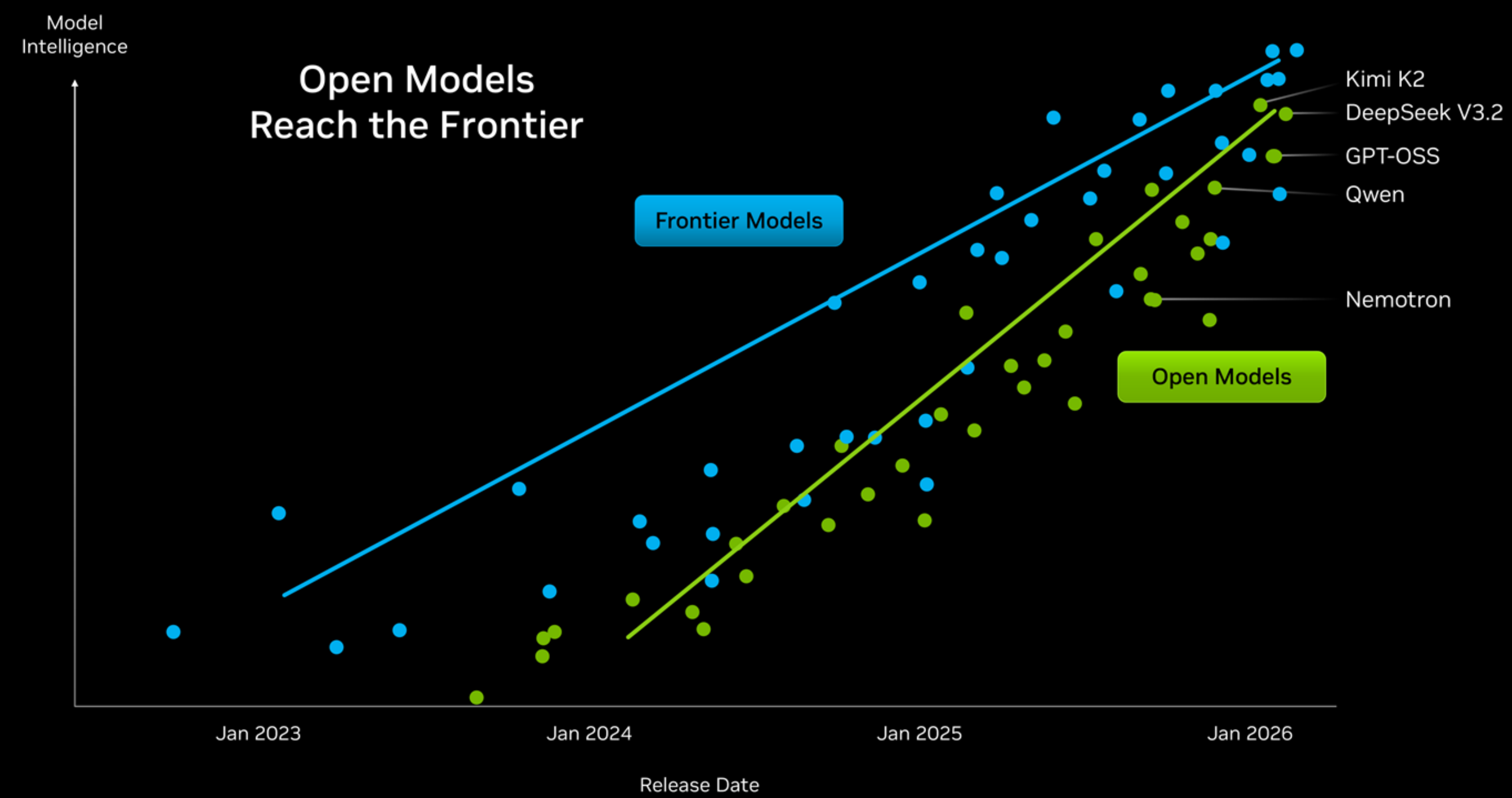
PHYSICAL AI
TAKES LEAP



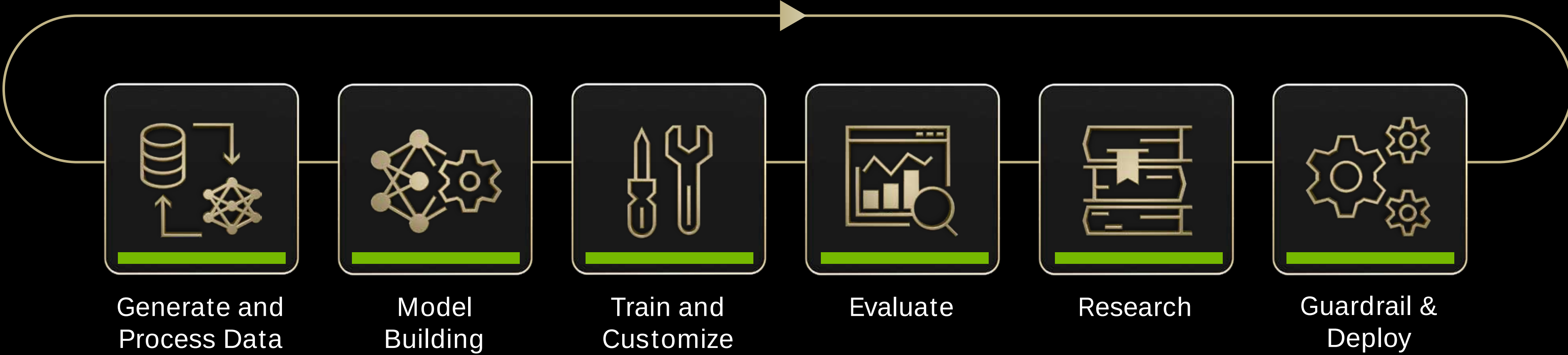
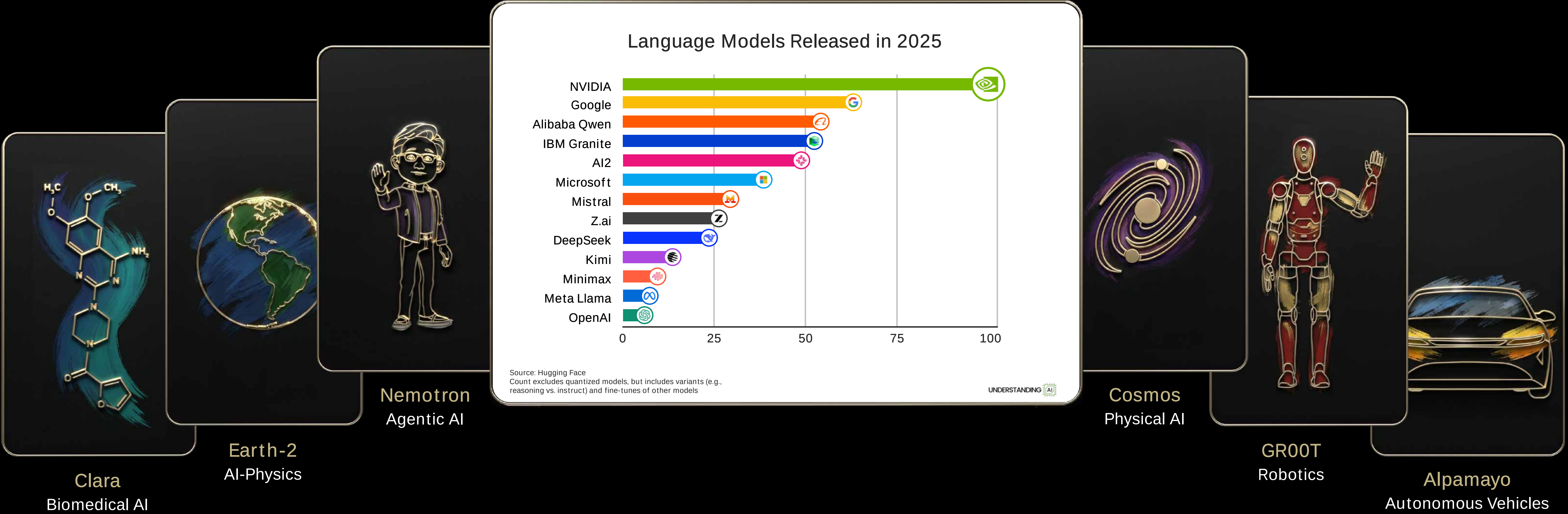
AI LEARNS LAWS
OF NATURE



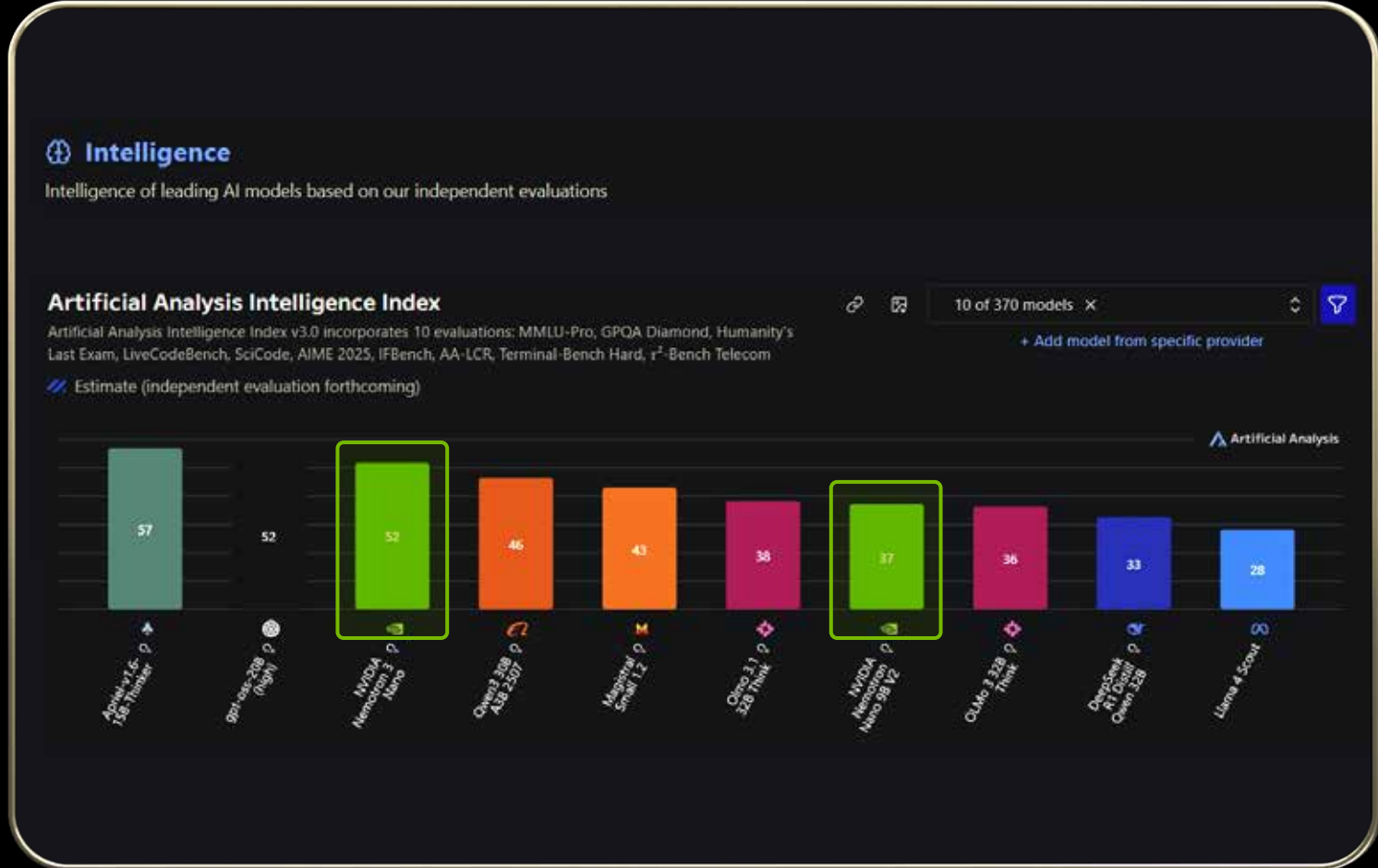
OPEN MODELS
REACH FRONTIER



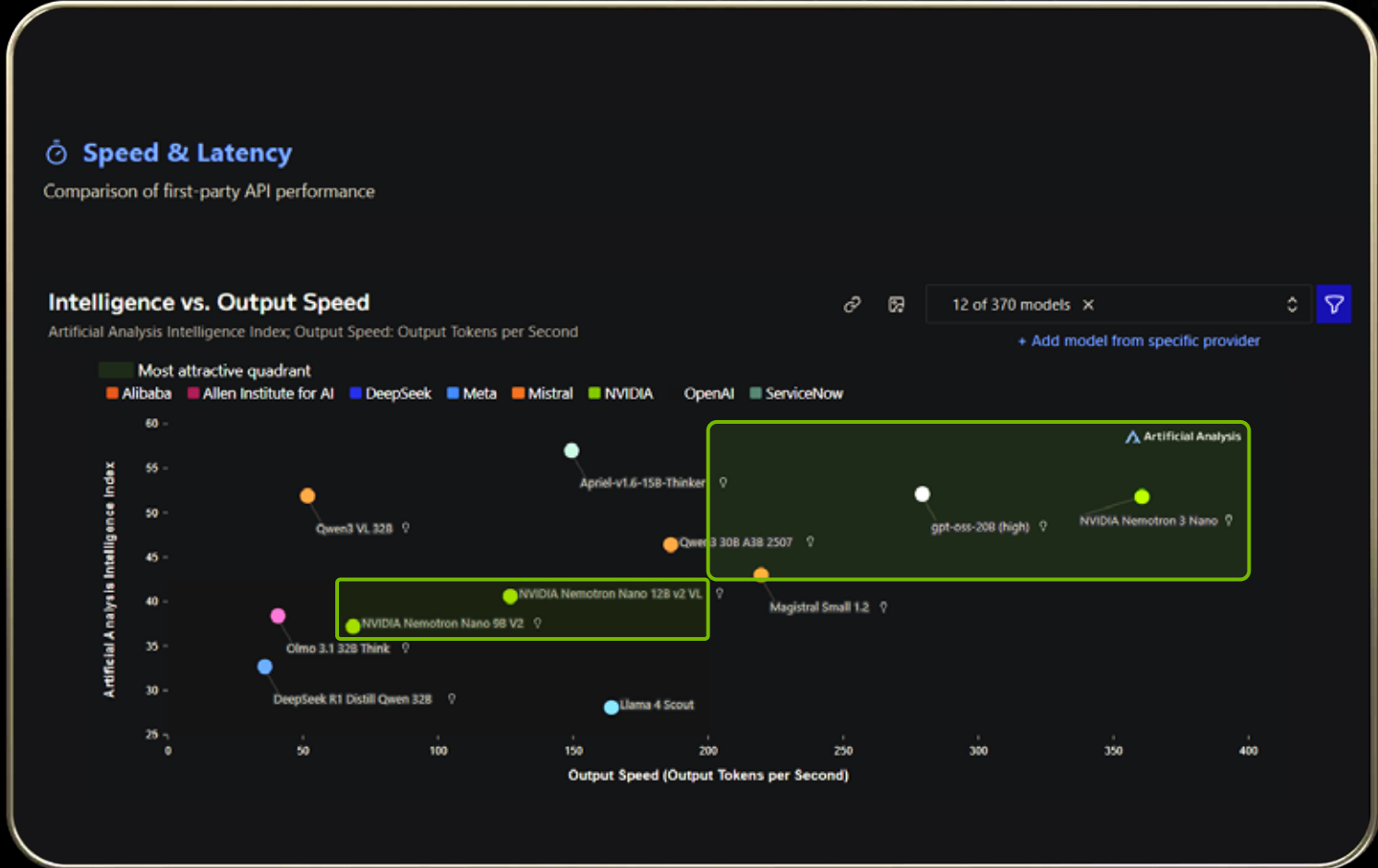
NVIDIA Leads Open Model Ecosystem



NVIDIA Tops Leaderboards



Artificial Analysis Intelligence Index



Artificial Analysis Intelligence per Token

#	User (Team)	Algorithm	Created	Mean Position
1st	Alina Amirajab (UJM_PM)	Report2CT	11 Aug. 2025	1.7
1st	Zohaib Alahuddin (UJM_PM)	CT Gen with Classifier Guidance	29 Aug. 2025	1.7
3rd	Danielle Molino (ArCo Lab)	Rep2CT No Rag	15 July 2025	3
4th	Ray W (CTFlow)	CTFlow	30 Aug. 2025	4.7
5th	Delweicheng4536 (always_wrong)	MedSyn v2.448	31 Aug. 2025	5.7
6th	pmath2012	LoRAD: Low-Resource Adaptive Diffusion via Pretrained Latent Representations	11 Aug. 2025	6.7
7th	Vim3challenge (CT-RATE Team)	GenerateCT (Scores will be updated)	6 July 2025	7.7
8th	Bagazrev	LimVAE decoder	2 Sept. 2025	8
8th	Uhuho (Potato Rebellion)	baseline	25 July 2025	8
10th	Sufan (MIC-DKFZ)	ctgen2	22 Aug. 2025	8.3

VLM3D

model	Average MER	RFX	License	AMI	Exam
nvidia/canary-gemini-2.5b	5.63	418.28	Open	10.19	10.45
ibm-granite/granite-speech-3.3-8b	5.74	146.42	Open	8.98	9.42
ibm-granite/granite-speech-3.3-2b	6	270.57	Open	8.9	10.25
microsoft/Phi-4-multimodal-instruct	6.82	151.1	Open	11.89	10.16
nvidia/sawheel-t4i-9.6b-v2	6.85	3386.82	Open	11.16	11.15
awsvoice/avalon-v1-m0	6.24	NA	Proprietary	11.58	11.37
nvidia/sawheel-t4i-9.6b-v3	6.32	3332.74	Open	11.39	11.19
nvidia/canary-ib-flash	6.35	1845.75	Open	13.11	12.77
kyral/s1t-2.6b-m0	6.4	88.37	Open	12.17	10.99
nvidia/canary-ib	6.5	235.34	Open	13.9	12.19
nyxhealth/crisper-whisper	6.67	84.85	Open	8.71	12.89
elevenlabs/scribe-v1	6.88	NA	Proprietary	14.43	12.14

OpenASR

Model Name	IntPhys 2 (K)	MVPBench (K)	CasualVQA (K)	Model Type
Human	92.44	92.9	84.78	Human
Cosmos-Reasoner-7B	59.88	41.31	48.17	Open
Y-JEPA-2	56.4	44.5	44.89	Open
o1-1.5	53.19	32.5	58.95	Closed
Gemini-2.5-Flash	49.12	36.7	49.85	Open
Gemini-2.5-Flash	56.1	-	61.66	Closed
Perception Language Model (PLM)	-	39.7	58.86	Open
InternVL2.5	-	39.9	47.54	Open
Gemini-1.5-Pro	52.1	29.6	-	Closed
LLaVA-Owl2-72B	-	28.7	45.27	Open
Test1	0	-	25.47	Open
Open-Mistral-7B-Instruct	-	-	67.59	Closed

Physical Reasoning MVPBench, IntPhys, CasualVQA

Name	Open Source	Text Recognition	Text Referring	Text Spotting	Relation Extraction	Element Parsing	Mathematical Calculation	Visual Understanding
LLaMA	-	-	-	-	-	-	-	-
Nemotron-Nano-VL-8B	Yes	78.2	69.1	61.8	81.4	39.2	31.9	73.1
InternVL2.5-14B	Yes	67.3	36.9	11.2	89	38.4	38.4	79.2
Gemini-Pro	No	61.2	39.5	13.5	79.3	39.2	47.7	75.5
Open2-VL-7B	Yes	72.1	47.9	17.5	82.5	25.5	25.4	78.4
InternVL2.5-28B	Yes	65.6	26.1	1.6	86.9	36.2	37.4	78.3
InternVL2.5-8B	Yes	68.6	30.4	8.8	85.3	34	27.1	77.5
o1-1.5	No	69.7	26.9	0.3	75.6	36.7	42.9	71.5
Owl2-72B	Yes	73.2	24.6	0.7	62.4	44.8	40.6	72.7
InternVL2-28B	Yes	63.4	26.1	0	76.8	37.8	32.3	79.4
Stem-VL	No	67.8	31.3	7.2	73.6	37.2	27.8	69.8

OCRBench

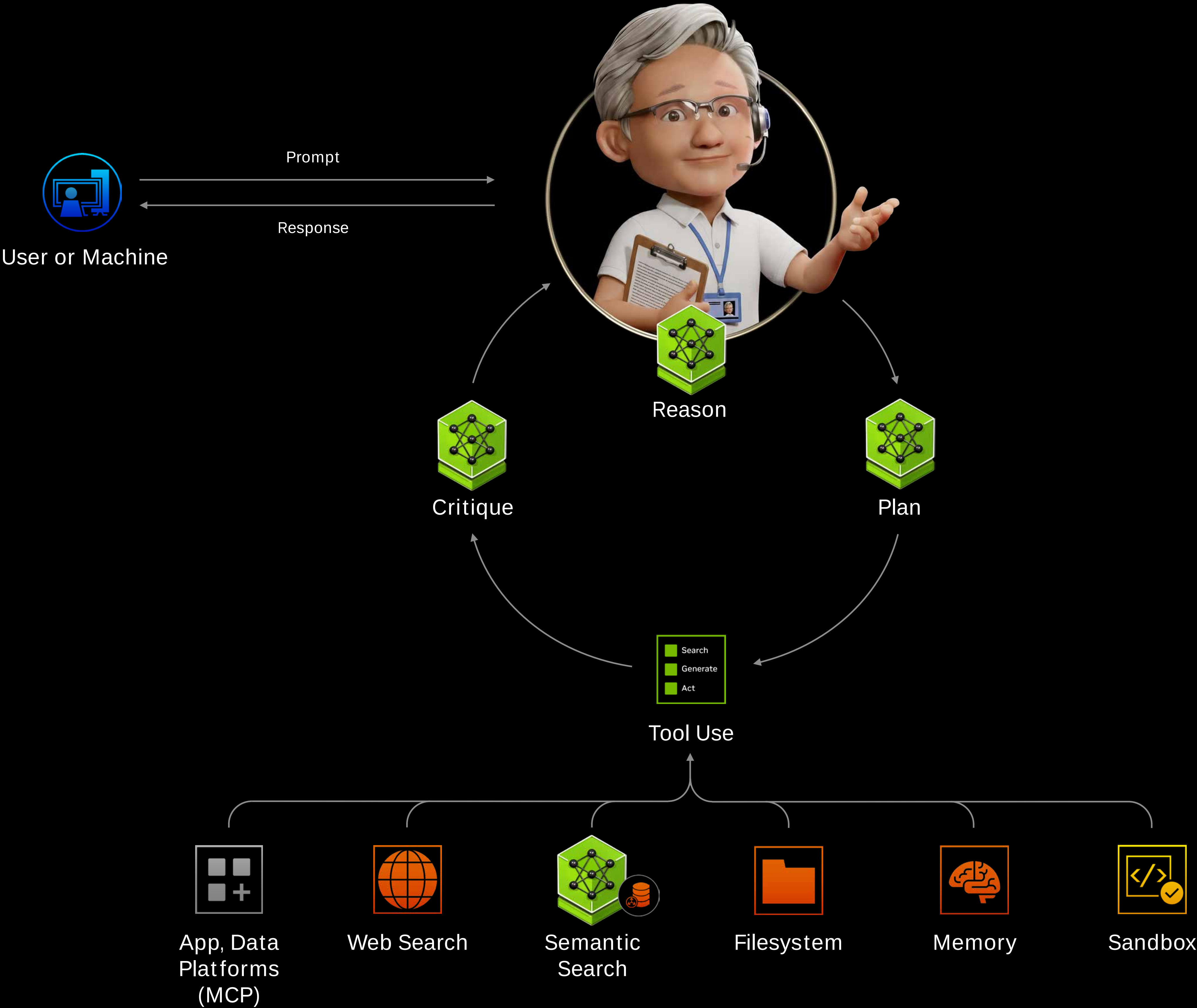
Model	Zero-shot	Memory Usage	Number of Para.	Embedding Dim.	Max Tokens	Mean F1
WLN-Embedding-Gemini-1.5B-2511	73%	44884	11.8	3840	32768	72.32
llava-embed-nemotron-8b	99%	28629	7.5	4896	32768	69.46
OpenAI-Embedding-8B	99%	14433	7.6	4896	32768	70.58
gemini-embedding-8B	99%	-	3072	2840	32768	68.37
OpenAI-Embedding-8B	99%	7671	4.8	2560	32768	69.45
OpenAI-Embedding-8B	99%	14433	7.6	4896	32768	67.85
Seed-4-embedding-1215	89%	-	2848	32768	70.26	-
OpenAI-Embedding-8.0B	99%	1136	8.596	1824	32768	64.34
gte-Open2.7B-Instruct	NA	29948	7.6	3584	32768	62.51
Lim-Embed-R100	99%	13563	7.1	4896	32768	61.47
multilingual-v5-large-instruct	99%	1068	8.568	1024	514	63.22
embedding-gemini-300b	99%	1155	8.388	768	2840	61.15

Retrieval ViDoRe, MTEB, MMTEB

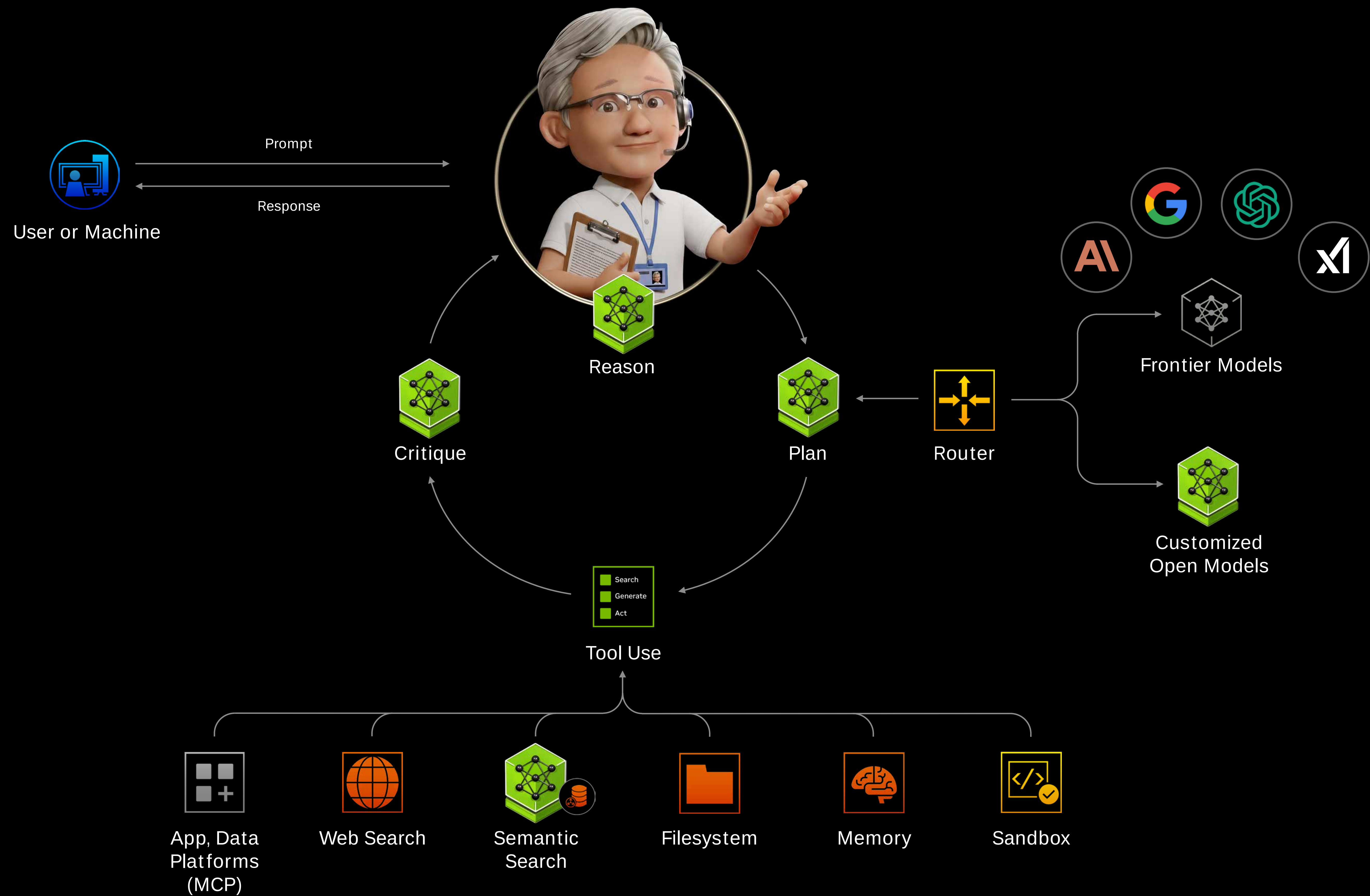
Model	Overall	Domain
Source	82.6	87.1
Veo-3	82.1	86.7
Cosmos-Predict2.5-14B	81.0	83.8
Cosmos-Predict2.5-14B	81.0	84.0
Man2.2-T2V-ALB	80.6	84.1
Man2.2-T2V-5B	80.4	83.4
Cosmos-Predict2.5-14B-Video2World	80.0	84.3
Man2.1-T2V-14B-720P	79.8	82.7
Cosmos-Predict2.5-14B-Video2World	79.6	83.9
Man2.1-28B	78.5	80.5
CopVideoX1.5-5B-T2V	78.3	80.1

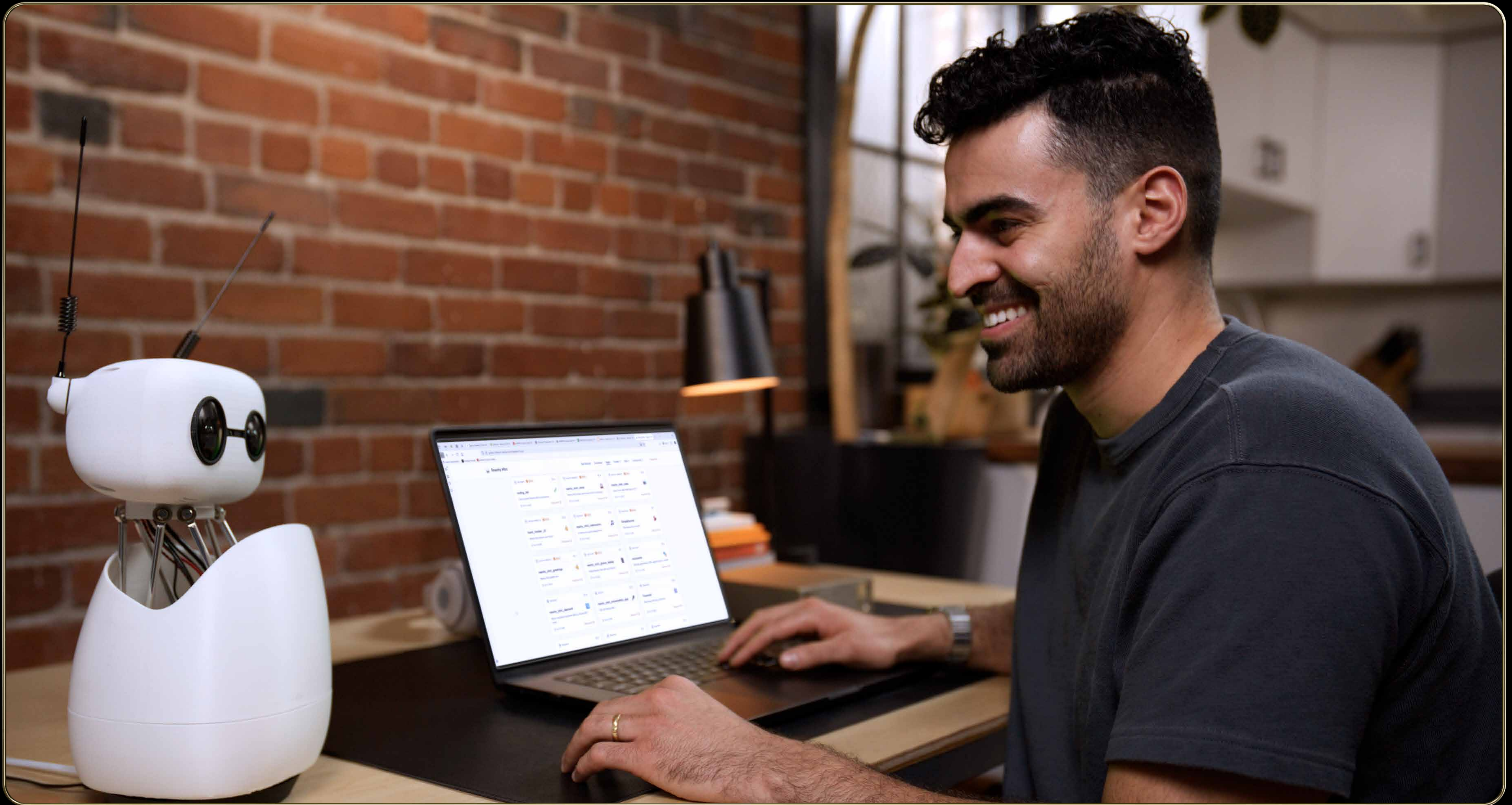
Physical AI PAIBench

Agents Are Multi-Model, Multi-Cloud, and Hybrid-Cloud



Agents Are Multi-Model, Multi-Cloud, and Hybrid-Cloud

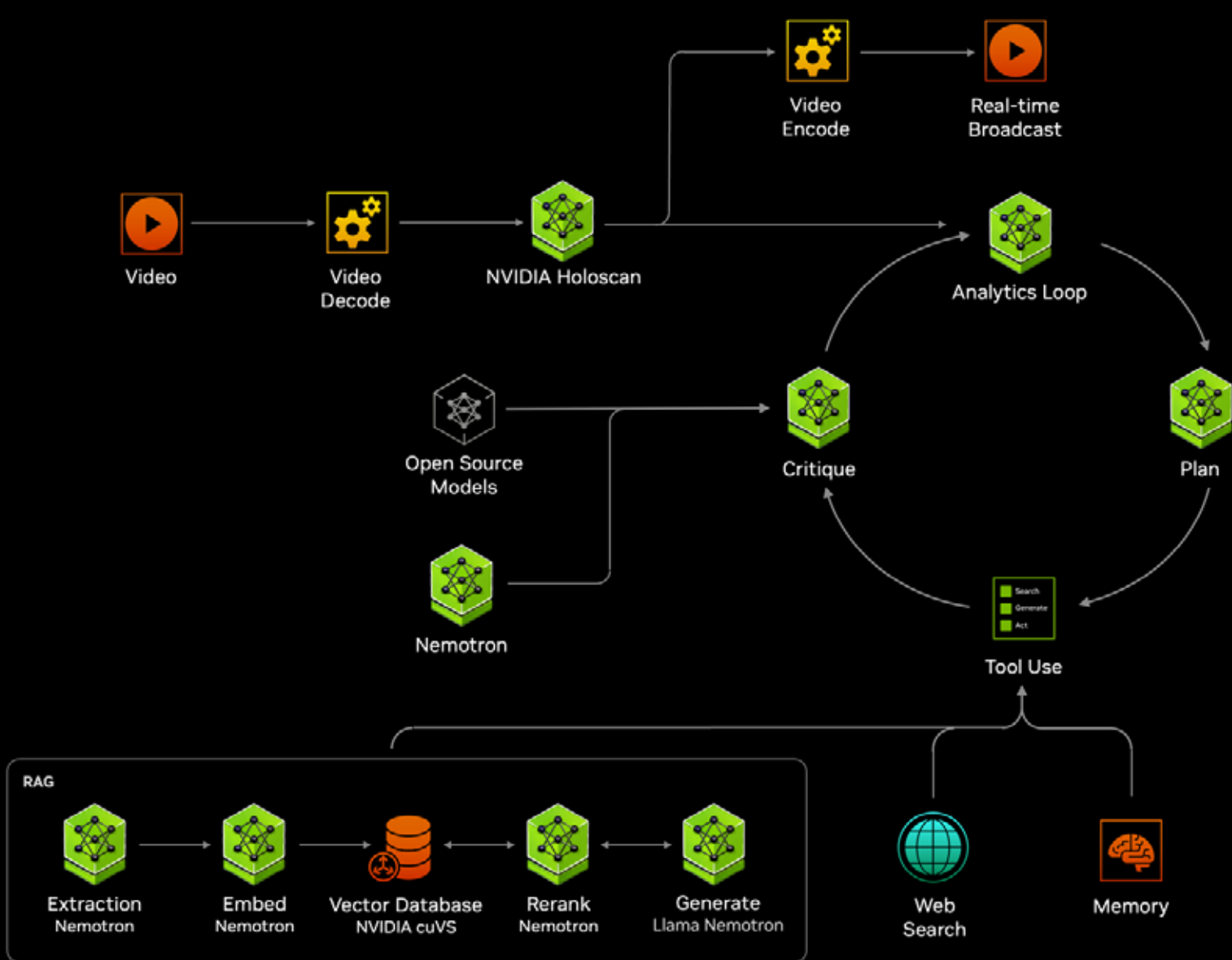




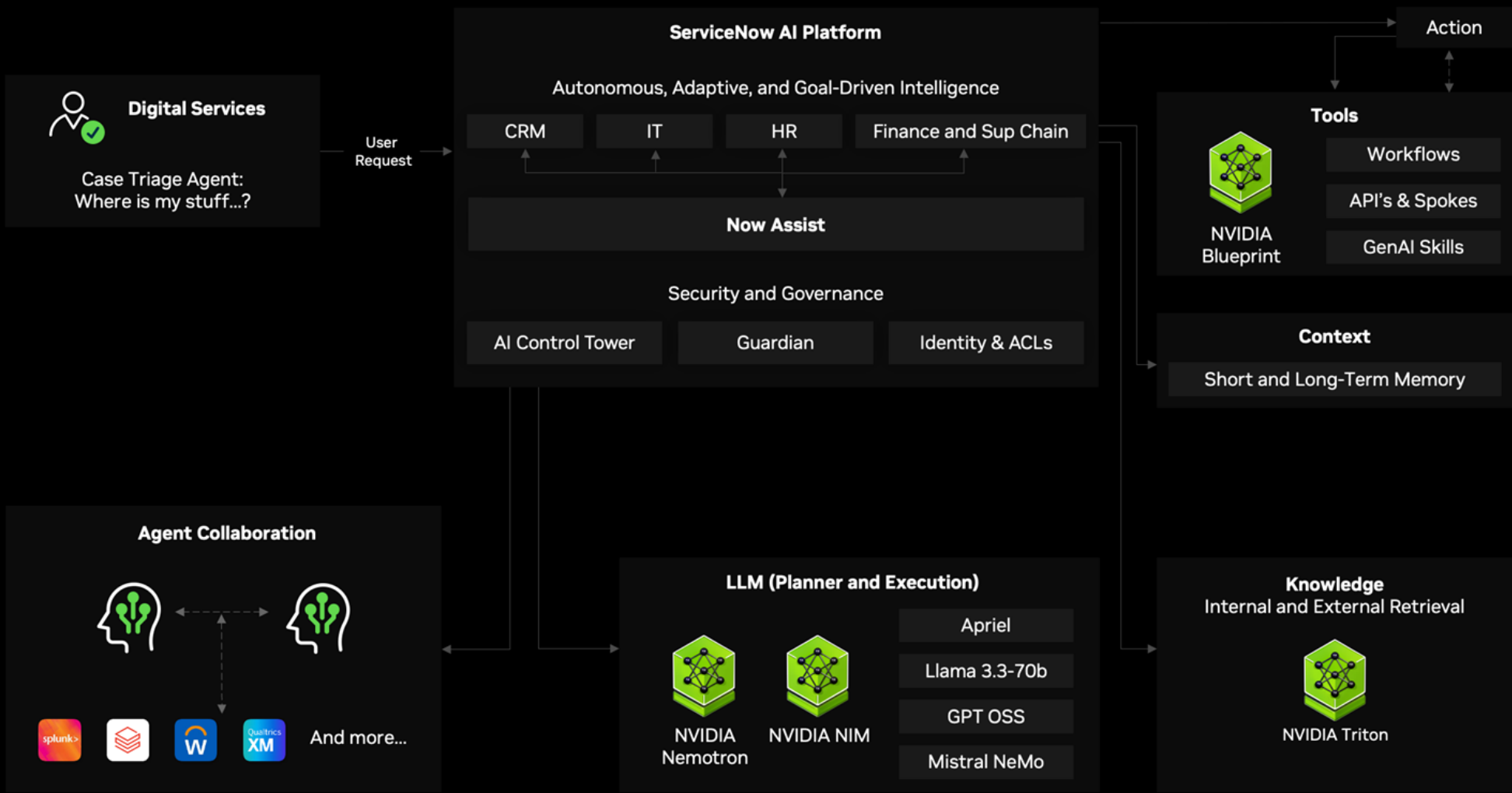
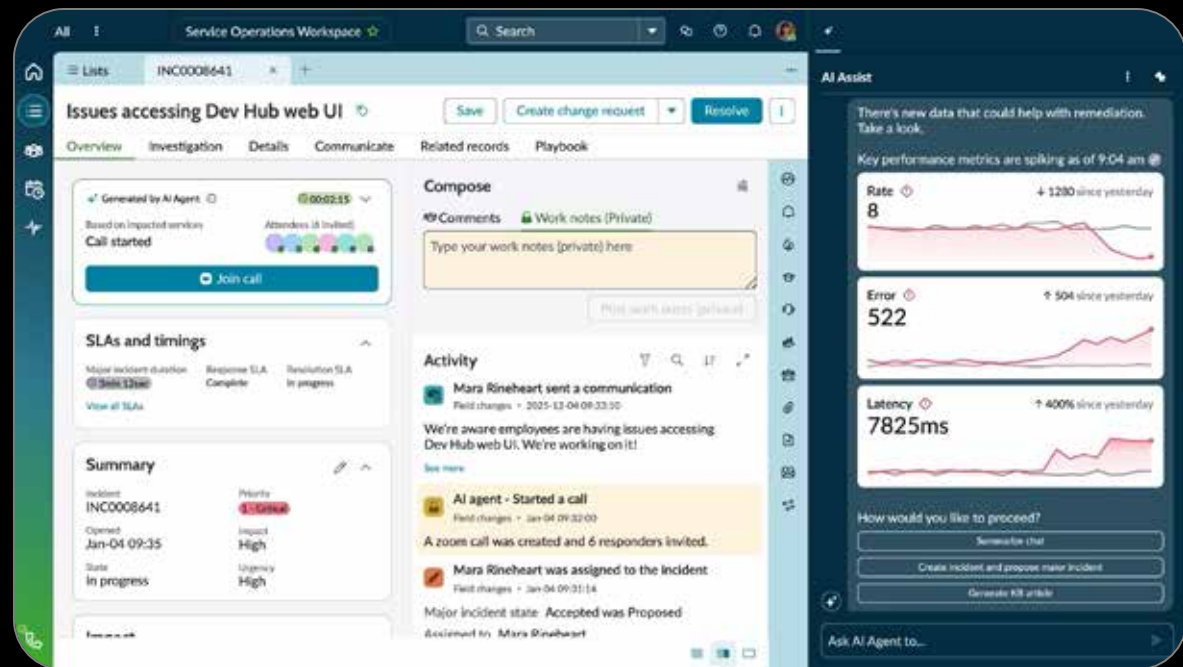
Enterprise Platforms Adopt NVIDIA AI



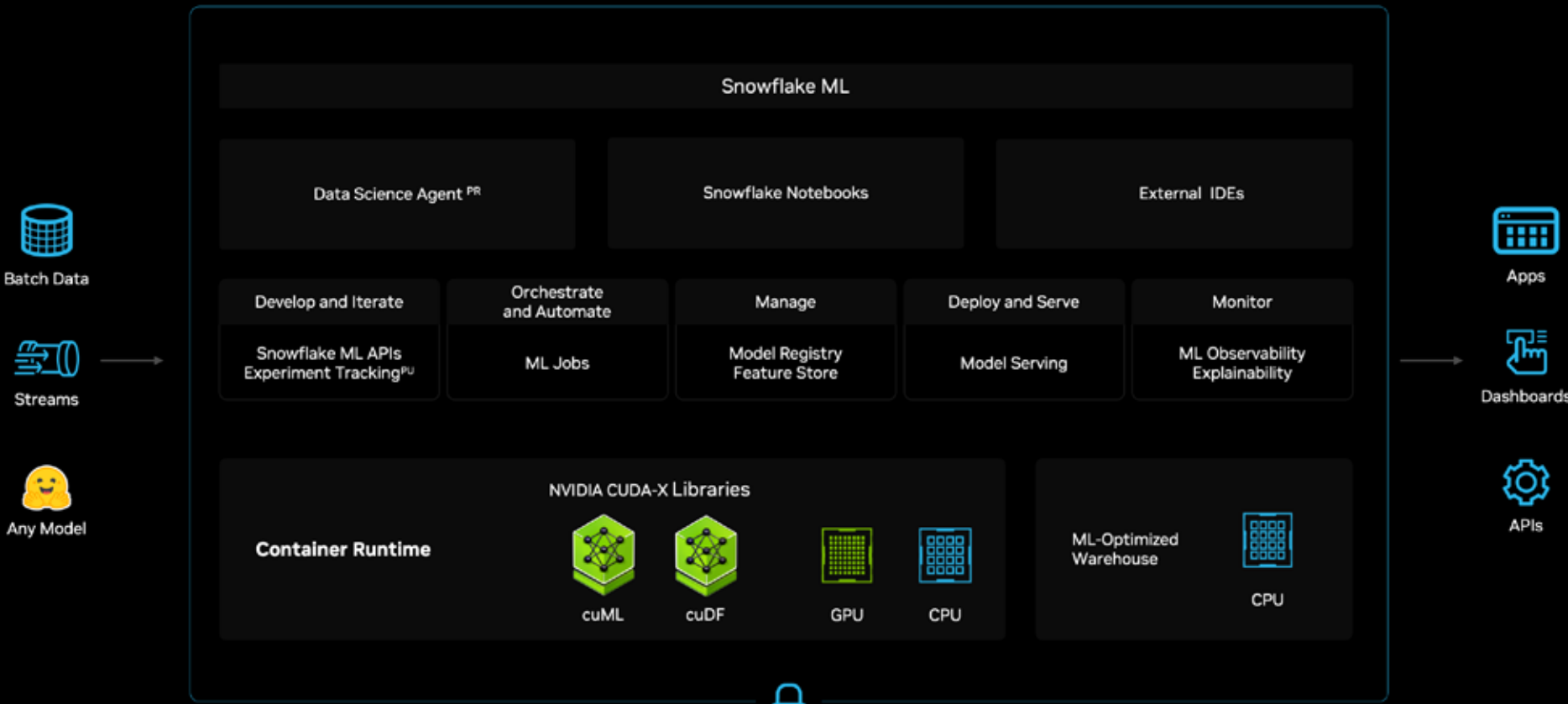
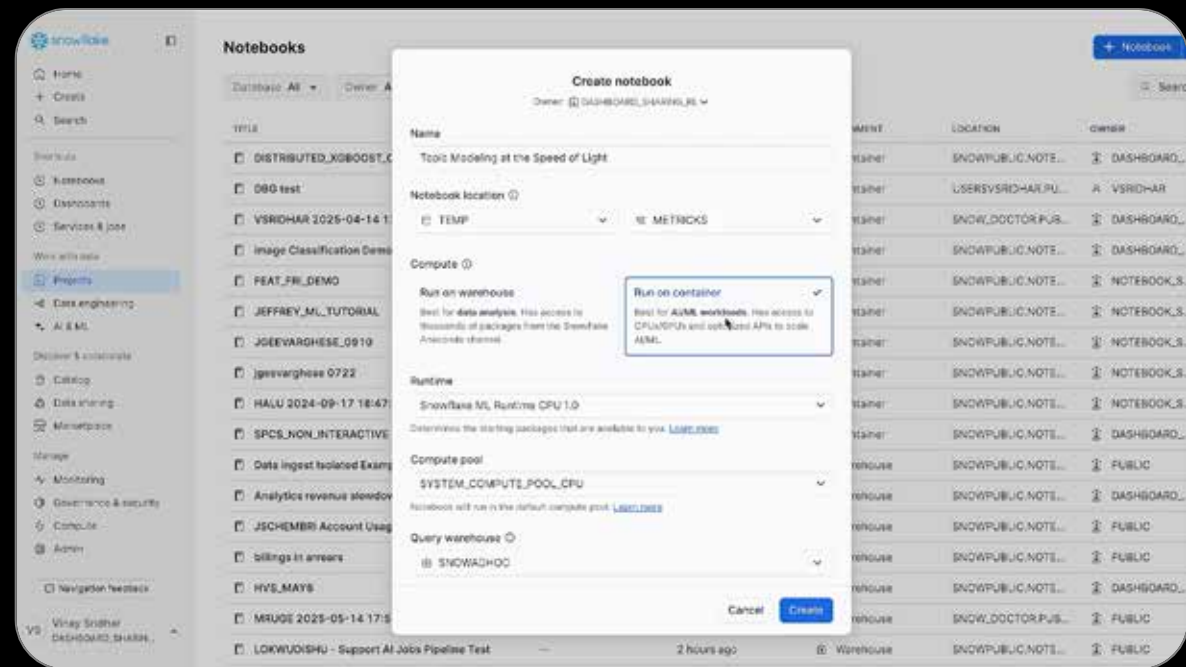
Edge Video Analytics



IT and HR Workflow Automation



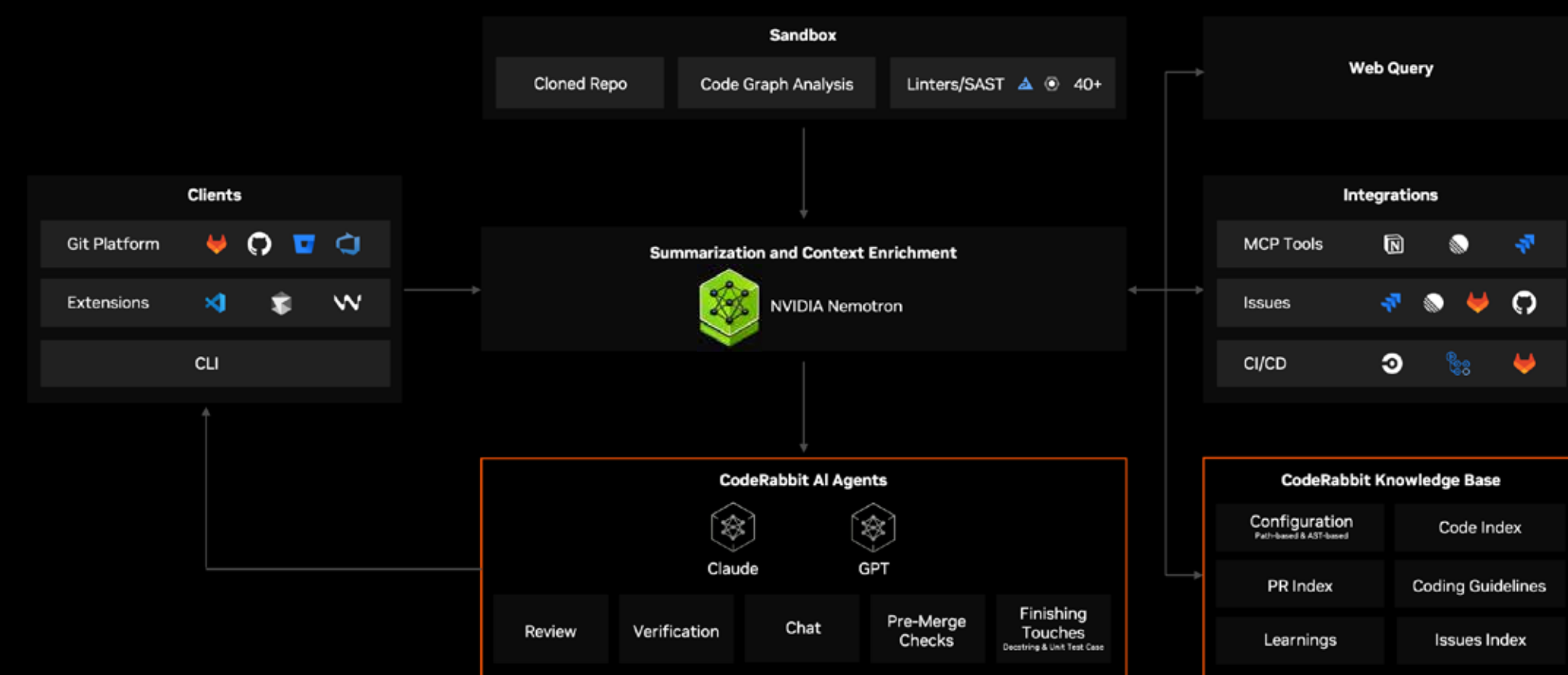
Accelerated Data Science and ML



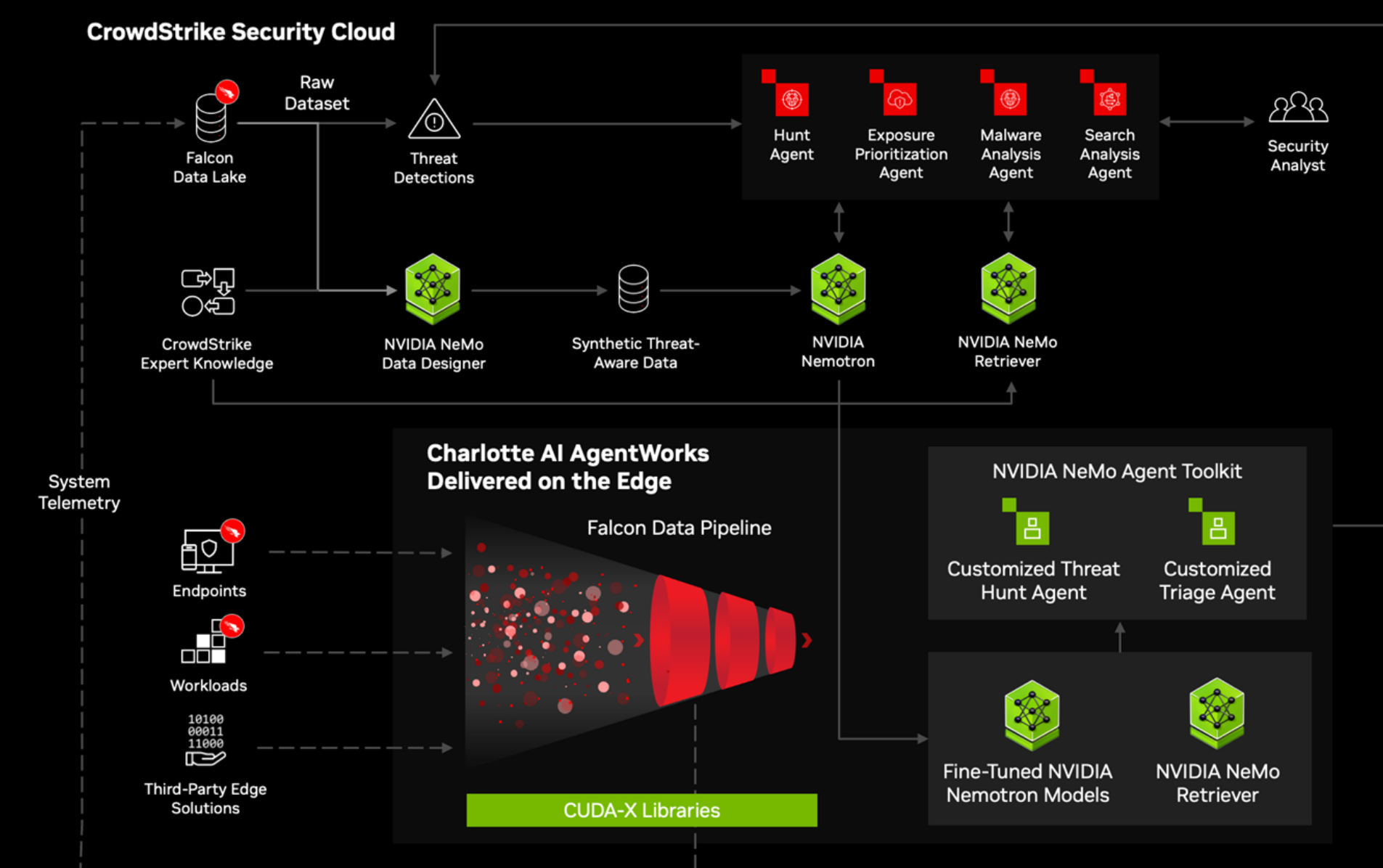
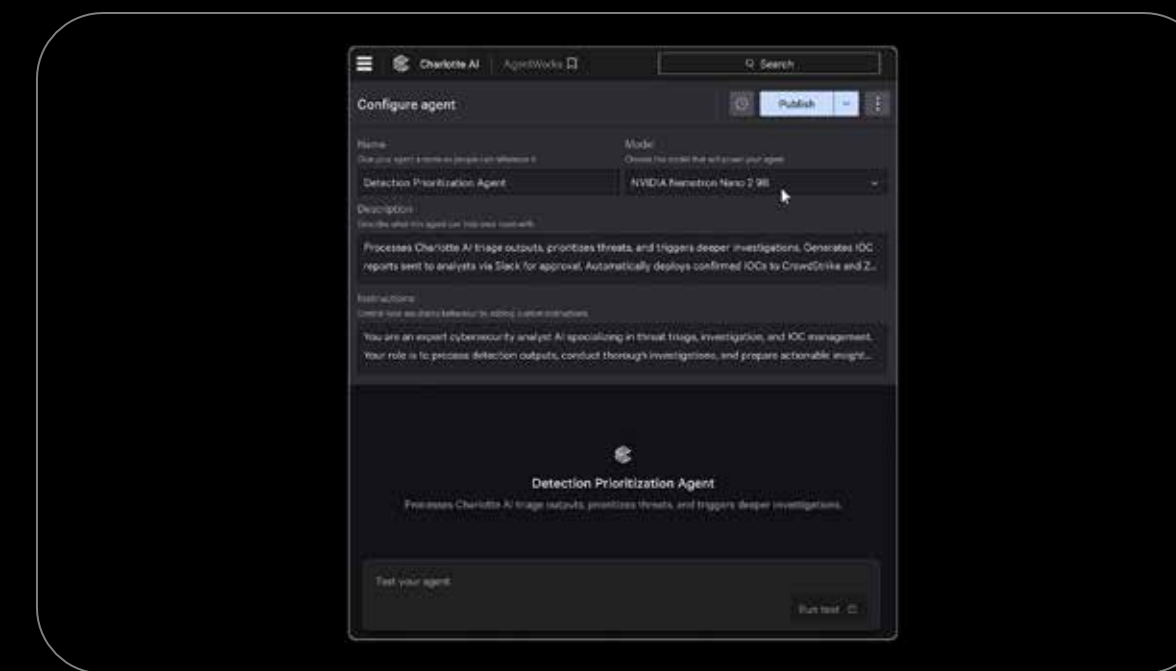
Enterprise Platforms Adopt NVIDIA AI



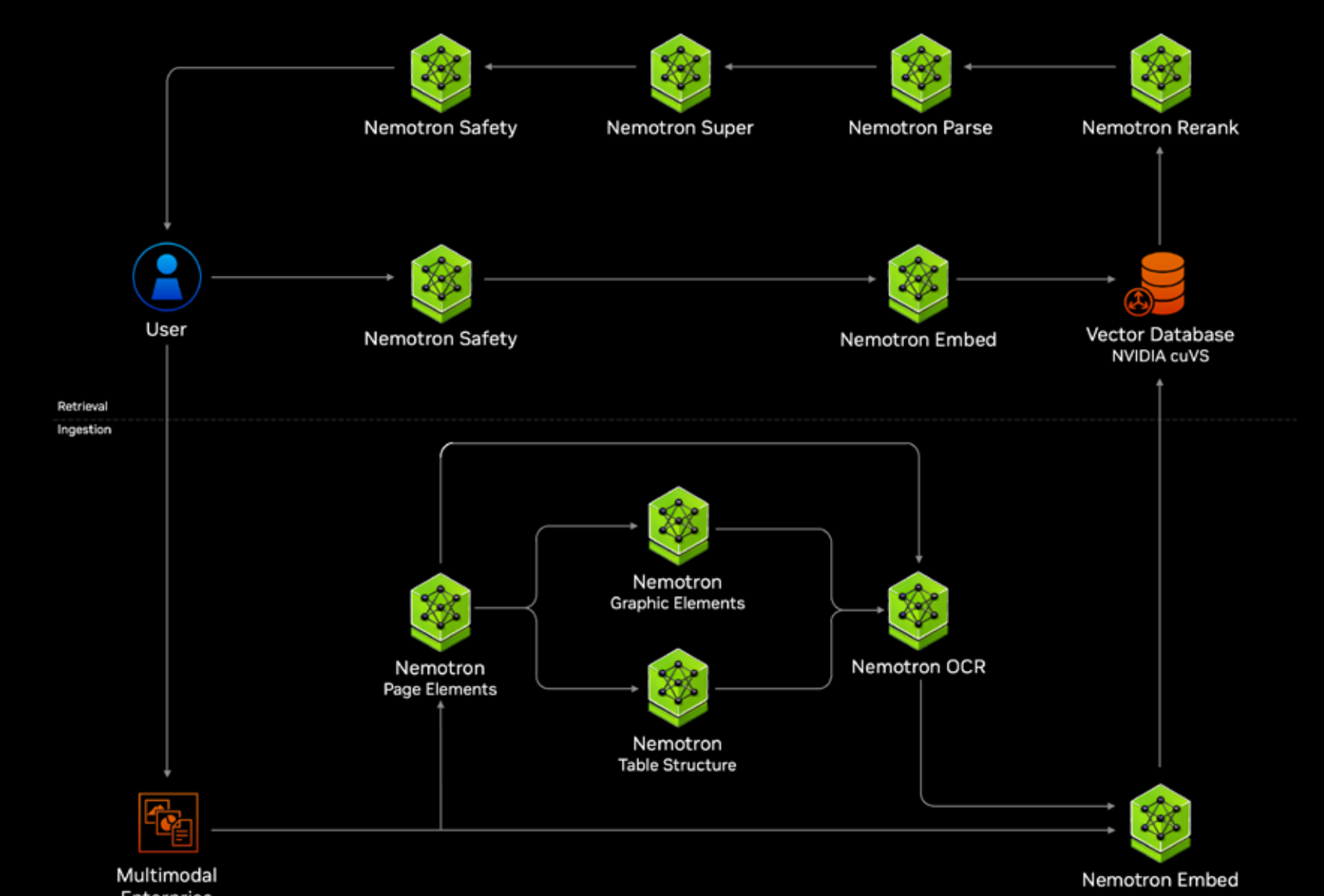
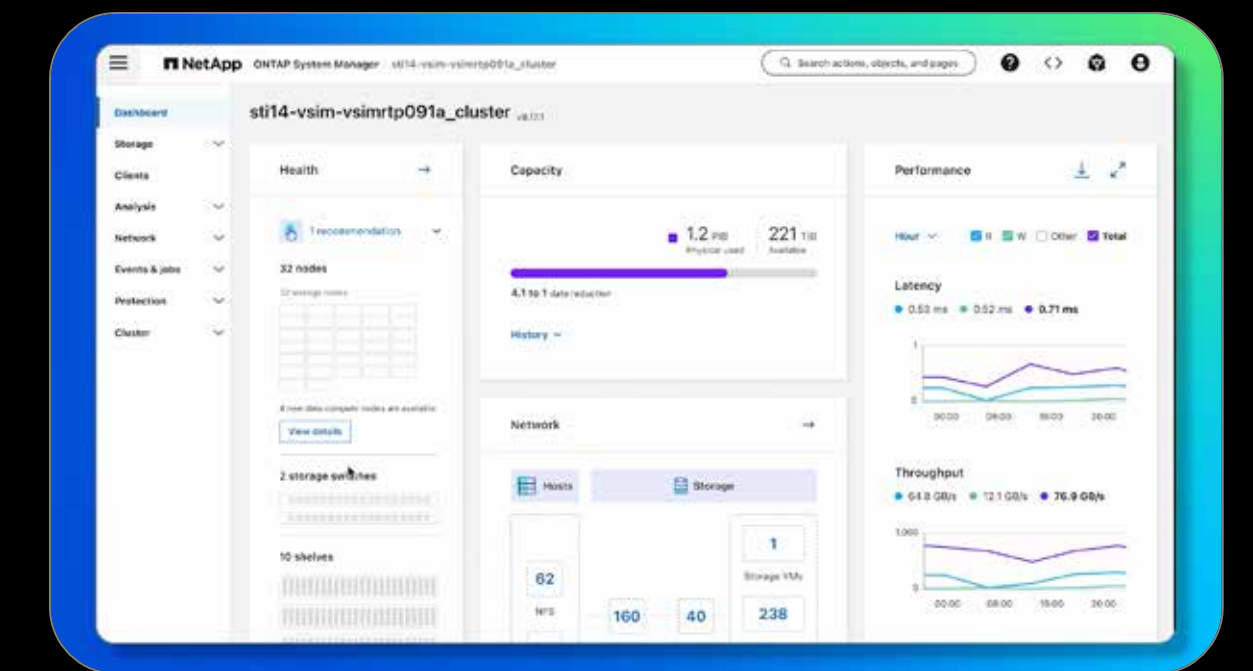
Code Review



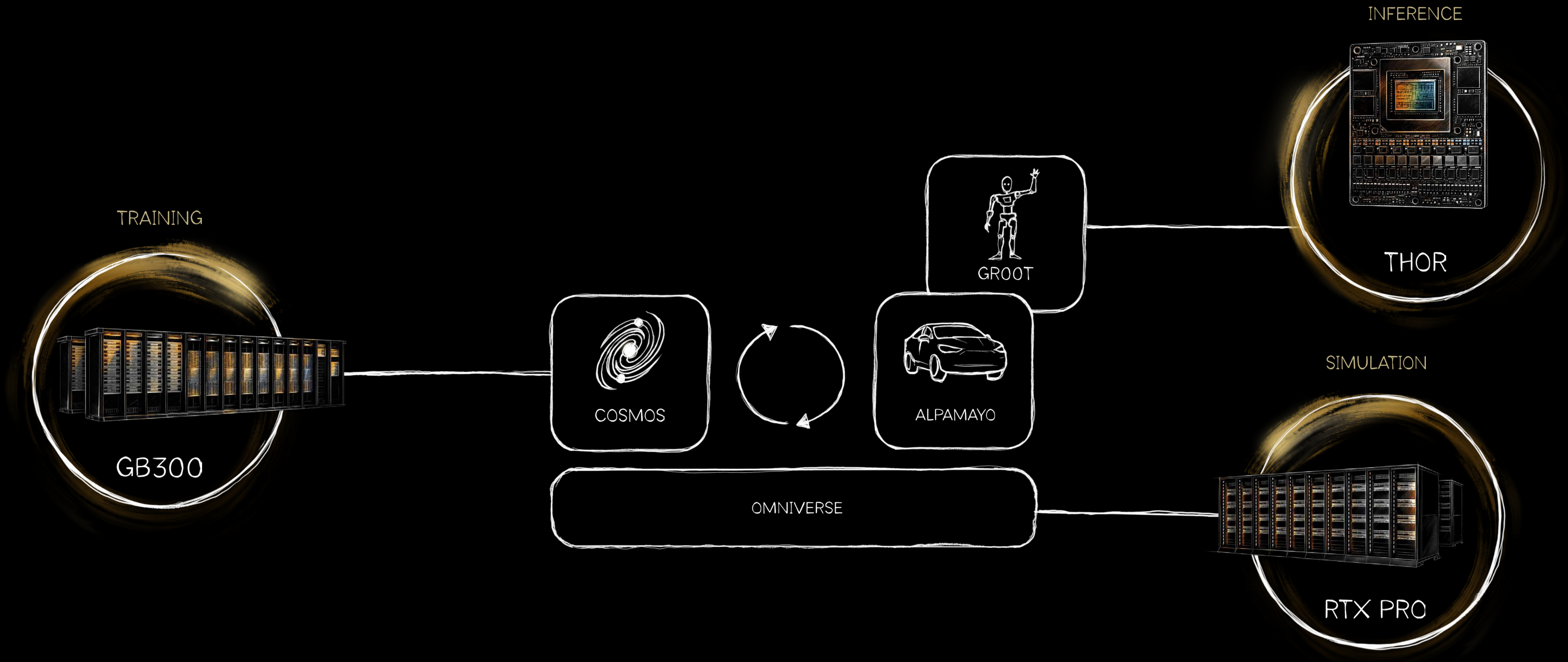
Real-Time Security Analyst



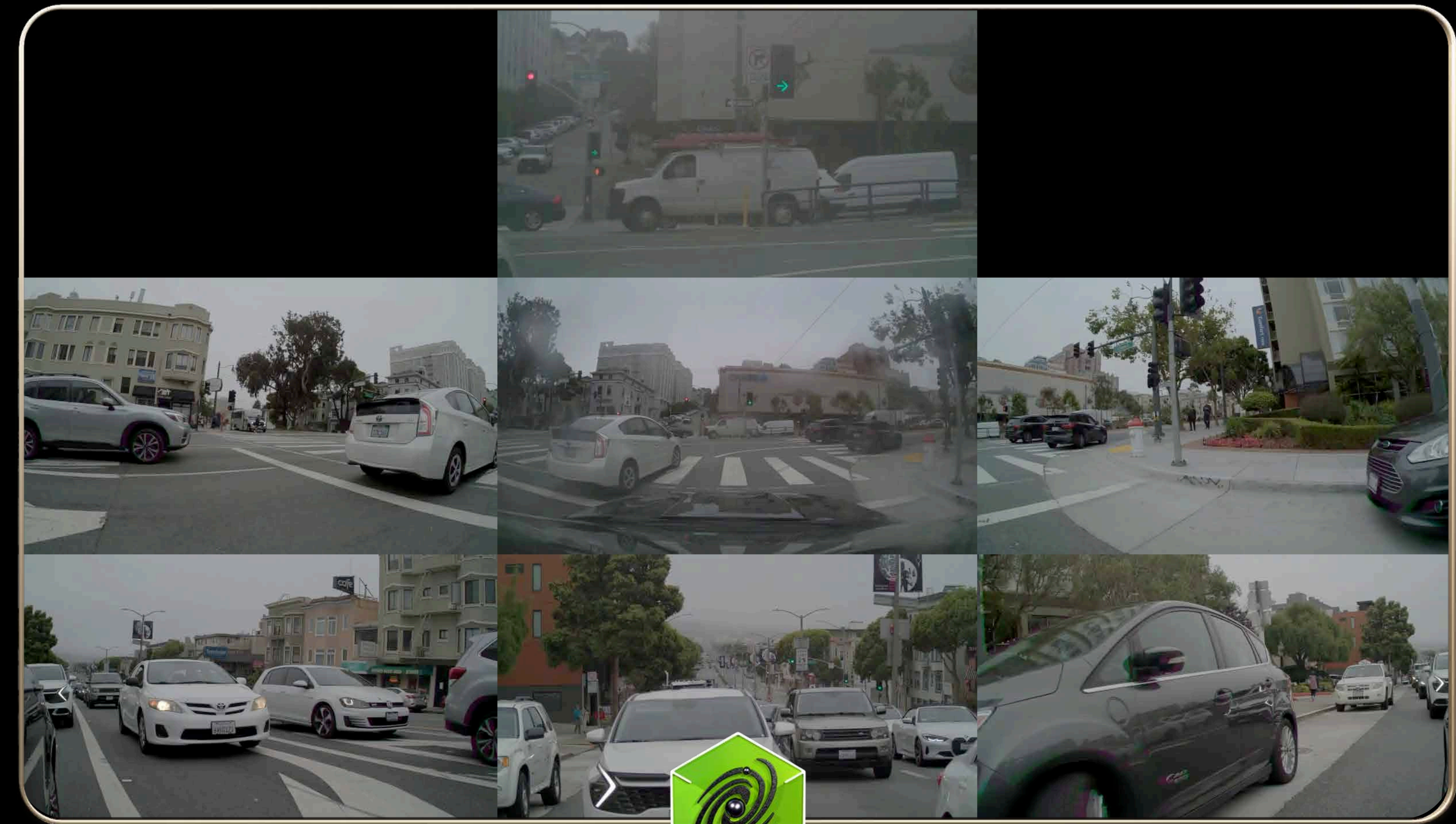
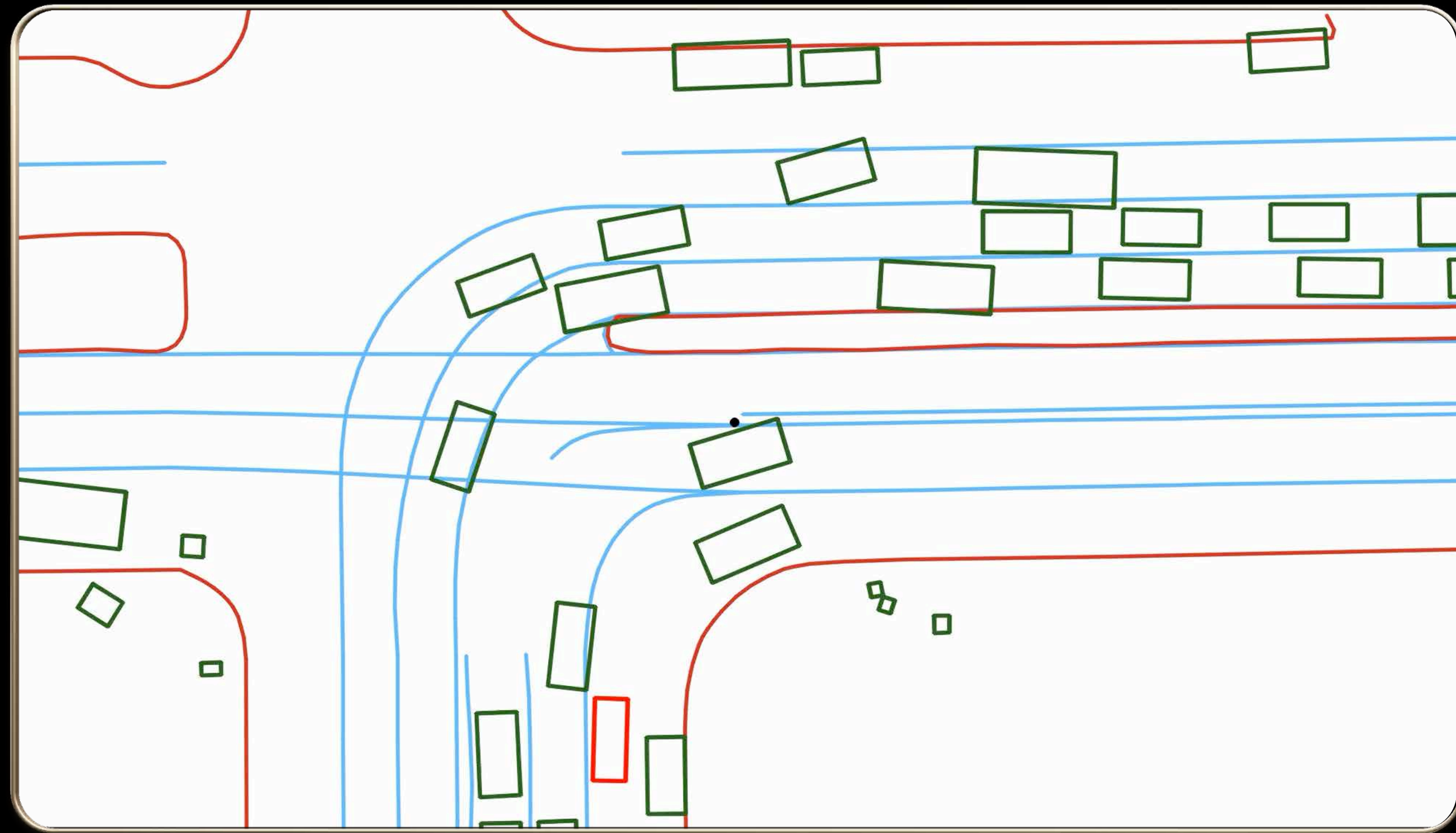
Retrieval Augmented Generation



NVIDIA Full-Stack Physical AI Platform



Compute is Data



Cosmos



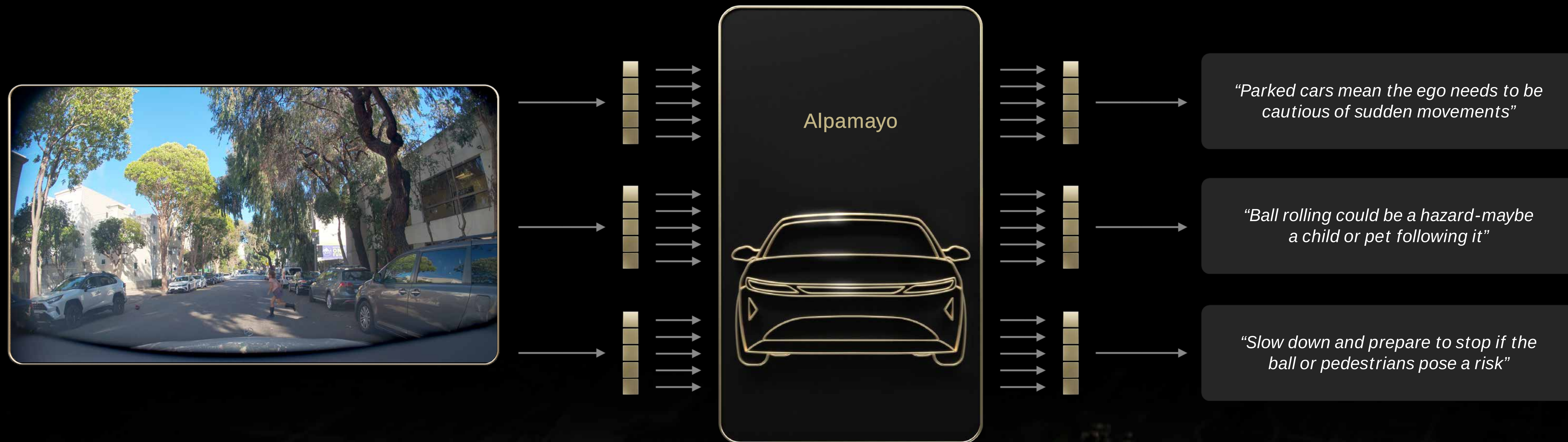
Announcing NVIDIA Alpamayo

Open Reasoning VLA for Autonomous Vehicles



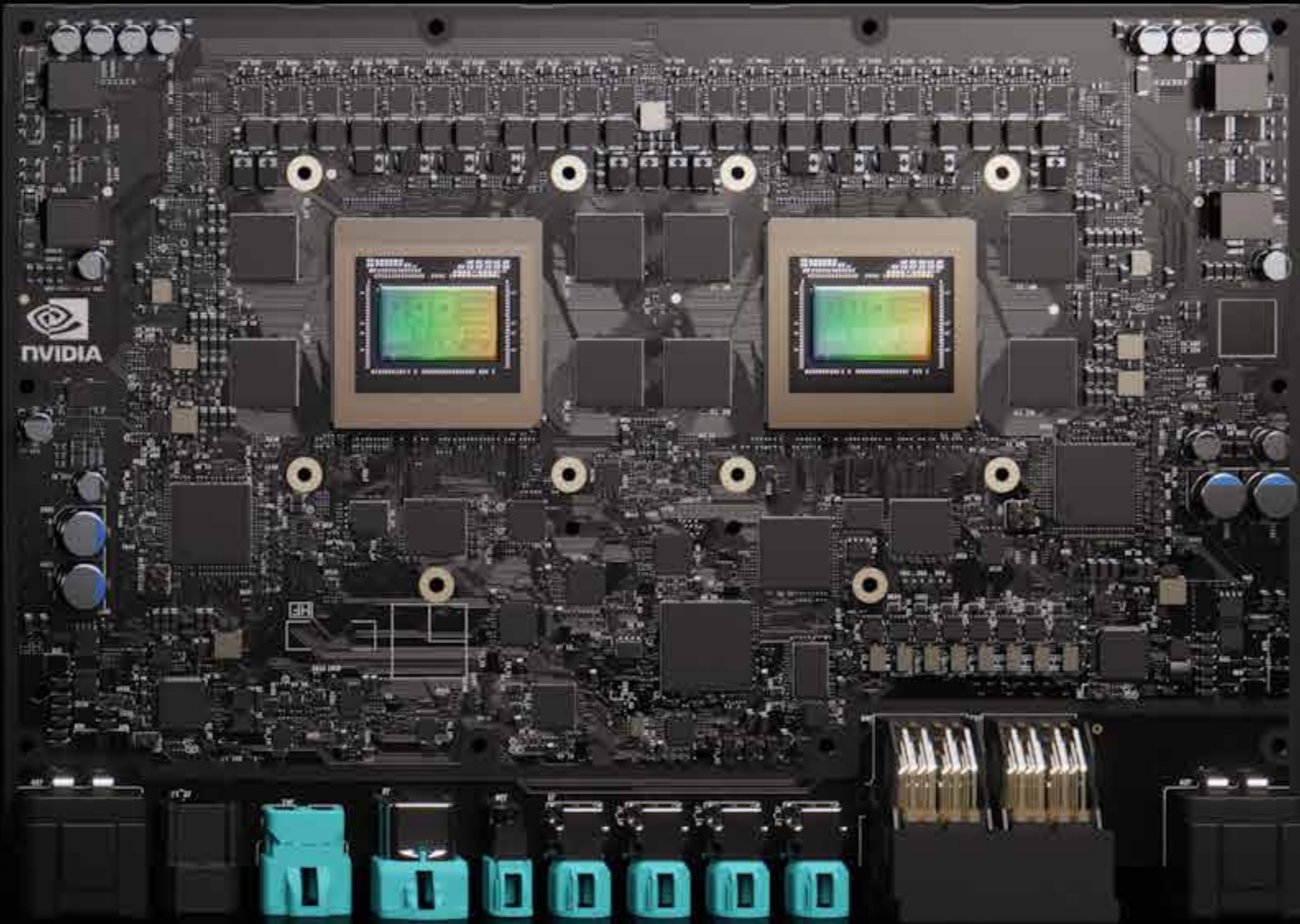
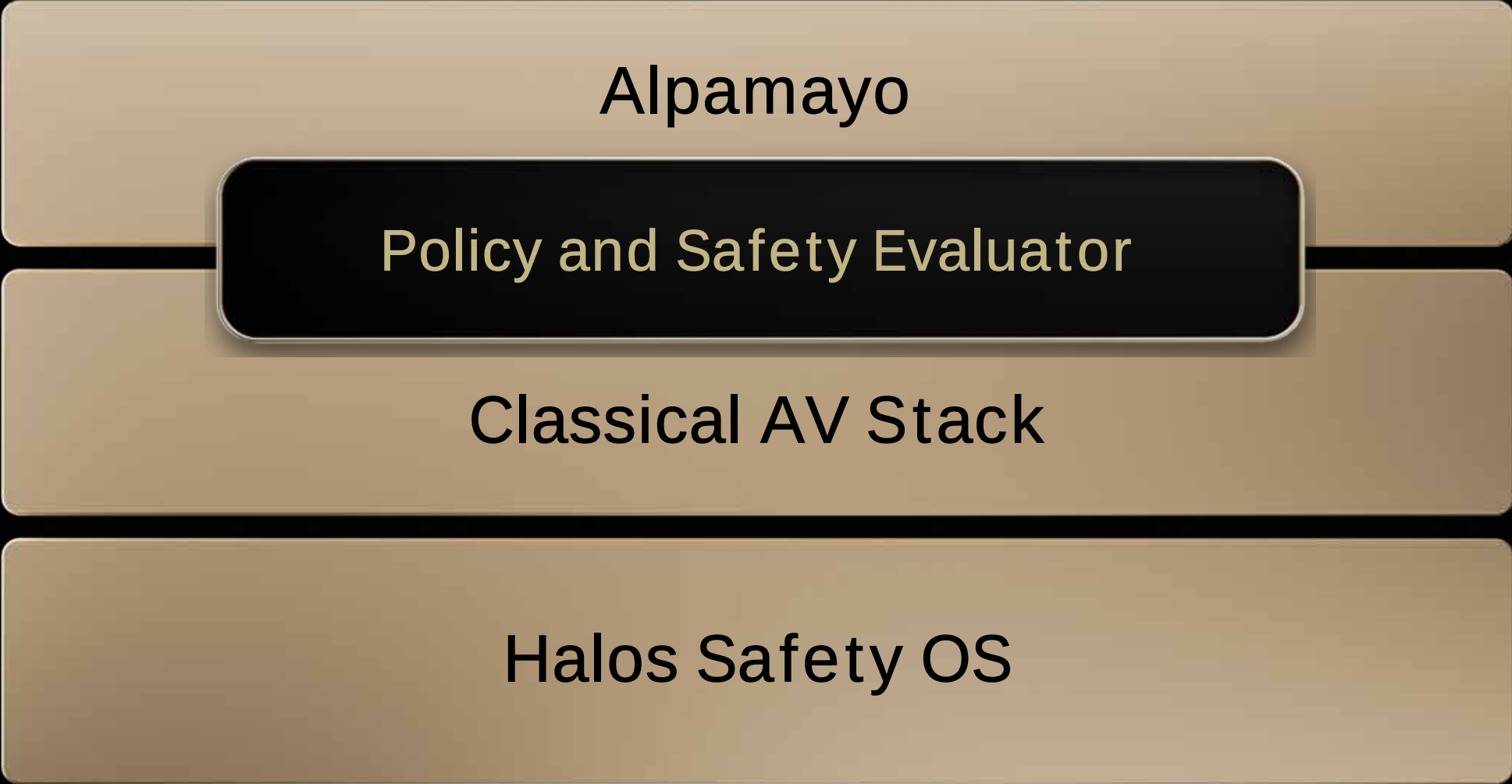
Announcing NVIDIA Alpamayo

Open Reasoning VLA for Autonomous Vehicles





NVIDIA Ships Full-Stack AV on 2025 Mercedes Benz CLA



Global L4 and Robotaxi Ecosystem Building on NVIDIA

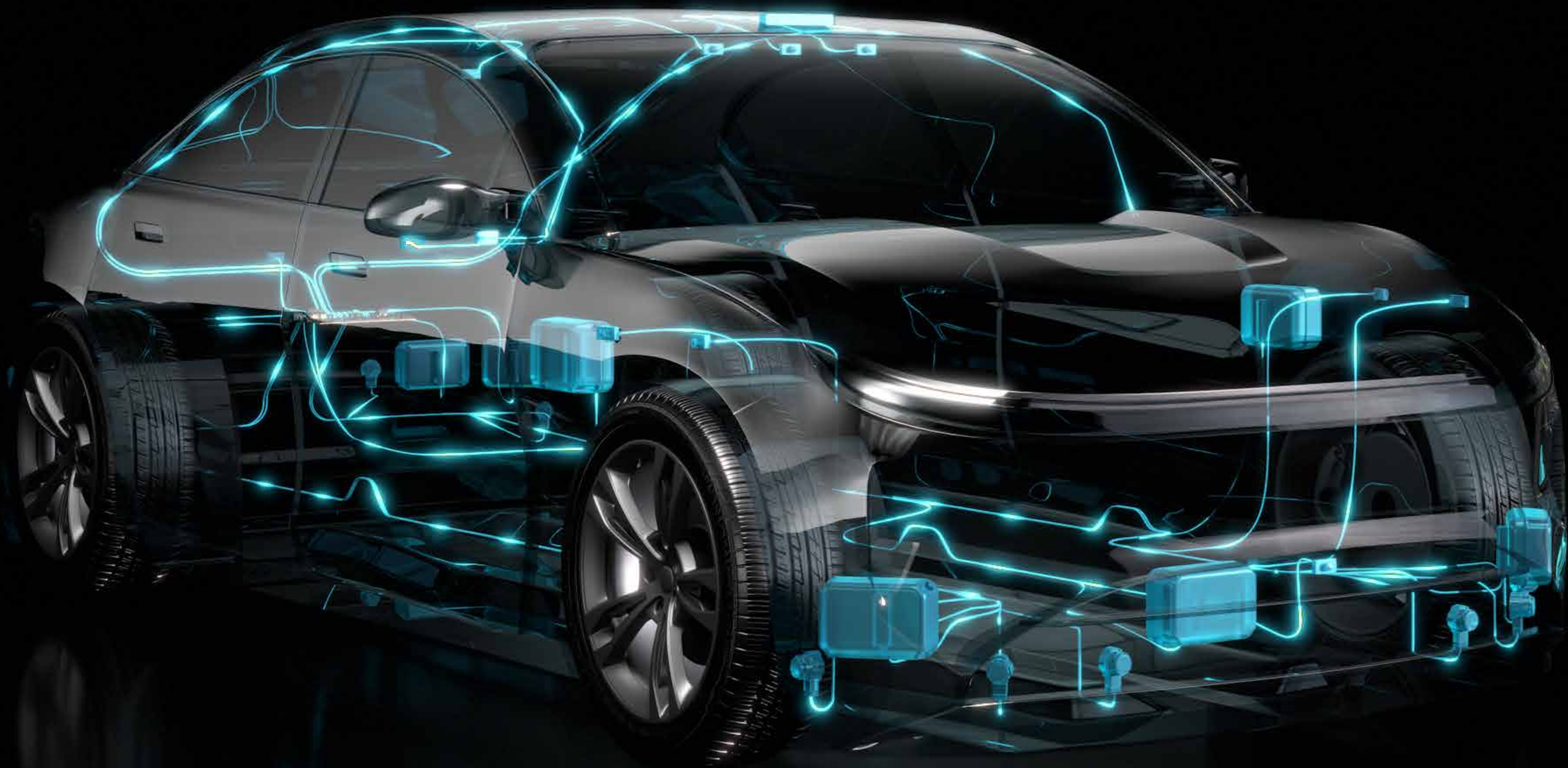
AV Software



OEMs and Mobility Platforms



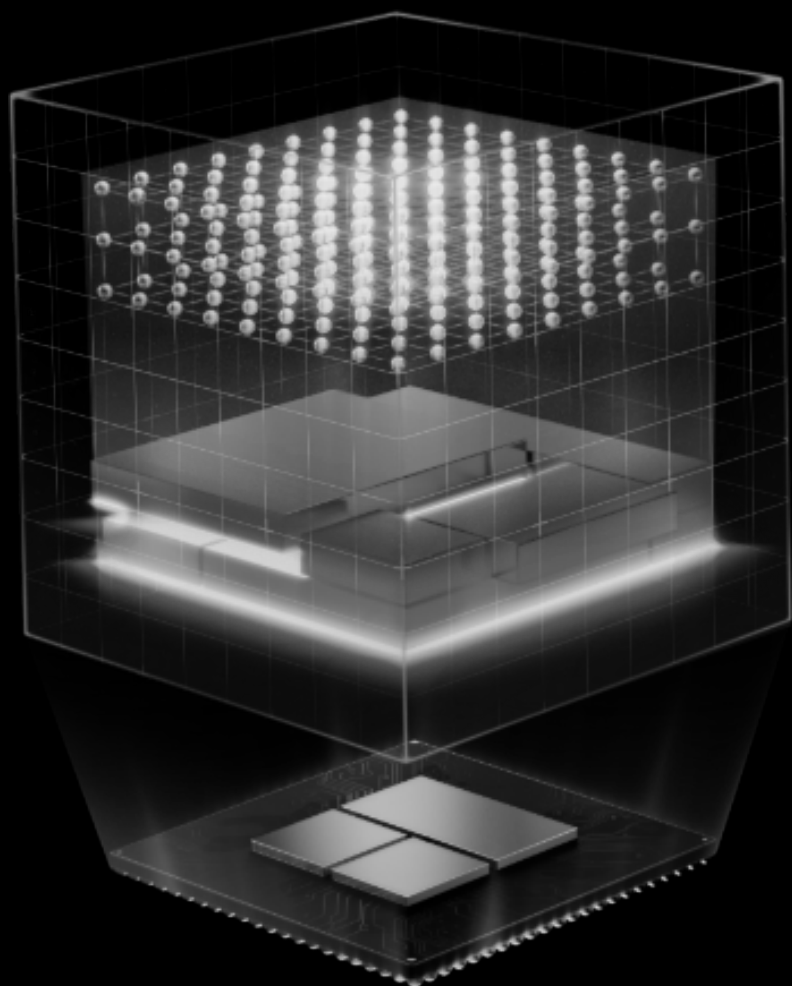
Tier 1 and Hardware Partners



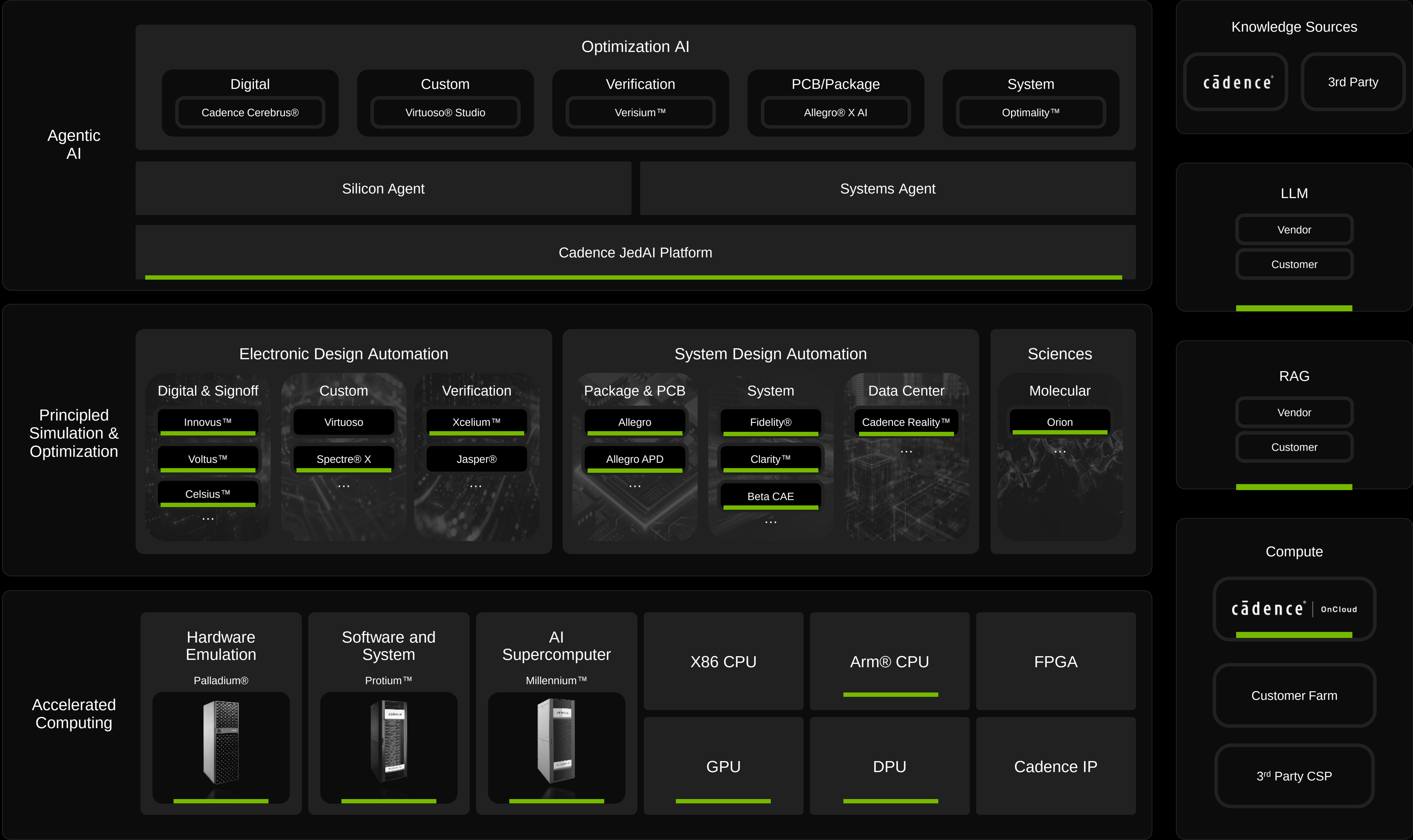








This slide contains forward-looking statements regarding Cadence's business or products. Actual results may differ materially from the information presented here.





Electronic Design Automation

Design

PrimeLib

Redhawk-SC

Totem SC

Manufacturing

Proteus

S-Litho

S-Device

Quantum ATK

Verification

PrimeSim

VCS

Computer Aided Engineering

Structures

Mechanical

LS-Dyna

Discovery

Fluids

Fluent

FreeFlow

Rocky

Electronics,
Electromagnetics

HFSS

Icepak

Maxwell

Optics

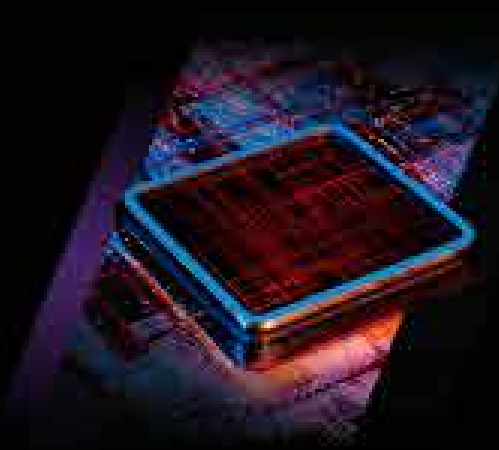
SPEOS

Zemax

Lumerical FDTD

Synopsys AgentEngineer™ Technology

Physical AI



Semiconductor DT



Industrial DT



Medical DT



Aerospace DT



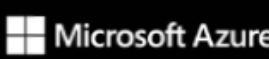
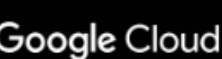
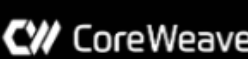
Robotics DT



Automotive DT

NVIDIA Omniverse

Targeting Availability in Every Cloud and OEM



Legend

In production



Industries

Chip Design

Fuse

Solido

Calibre

Tessent

Questa One

Veloce

Aprisa

Catapult

Xpedition

HyperLynx

Design / Engineering

Teamcenter
Digital Reality Viewer

Designcenter

Simcenter

Digital Twin

Operations

Industrial AI Suite

Industrial Edge

Totally Integrated
Automation (TIA)

Industrial Cybersecurity
Services

Opcenter

Infrastructure

BuildingX

Gridscale X

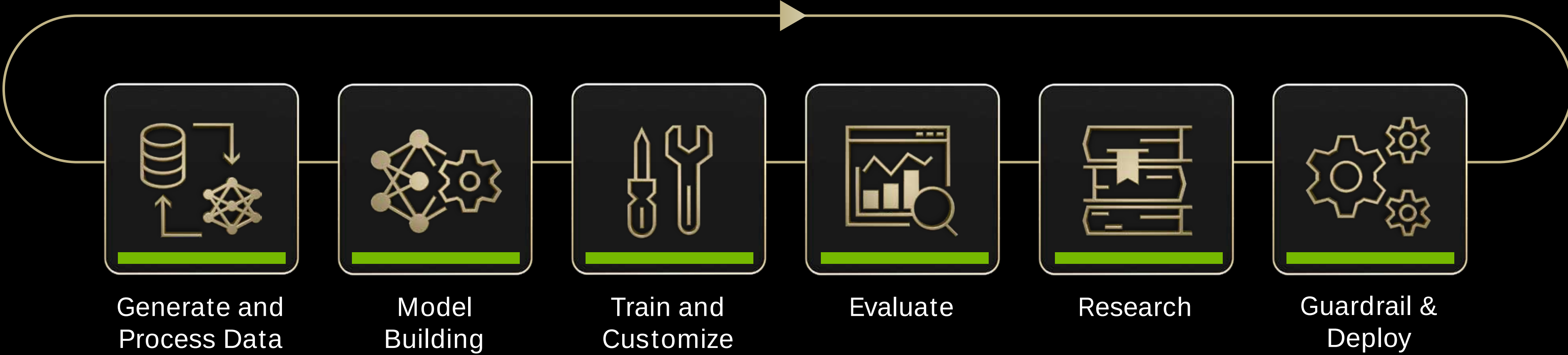
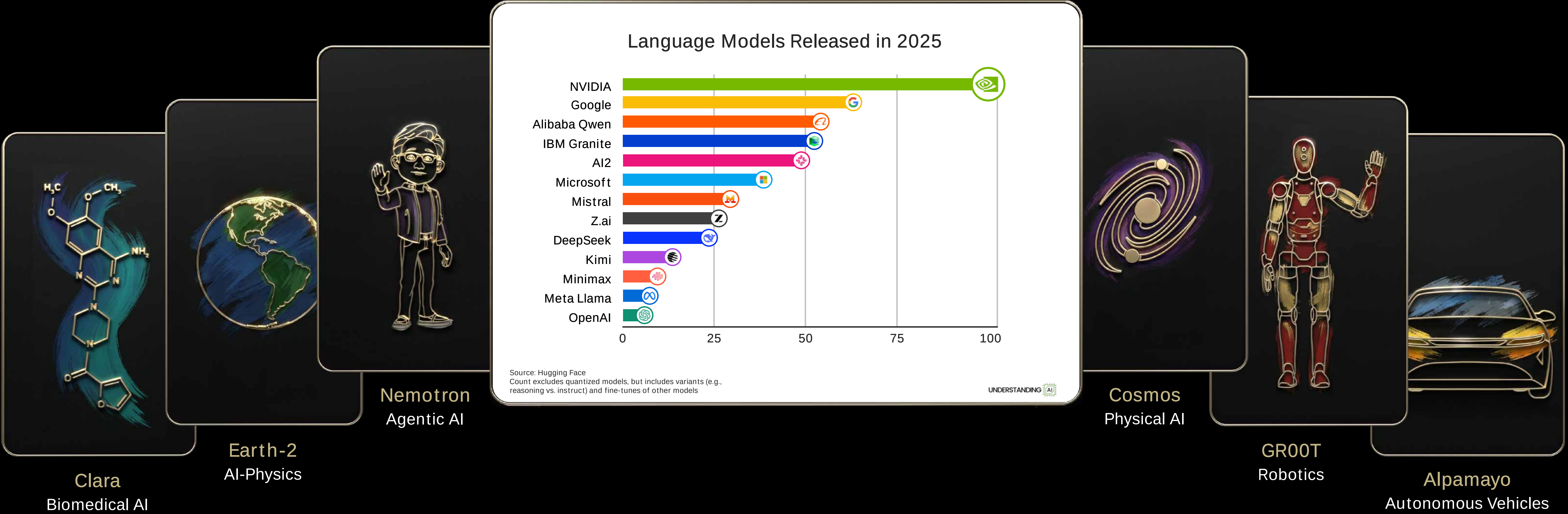
Electrification X

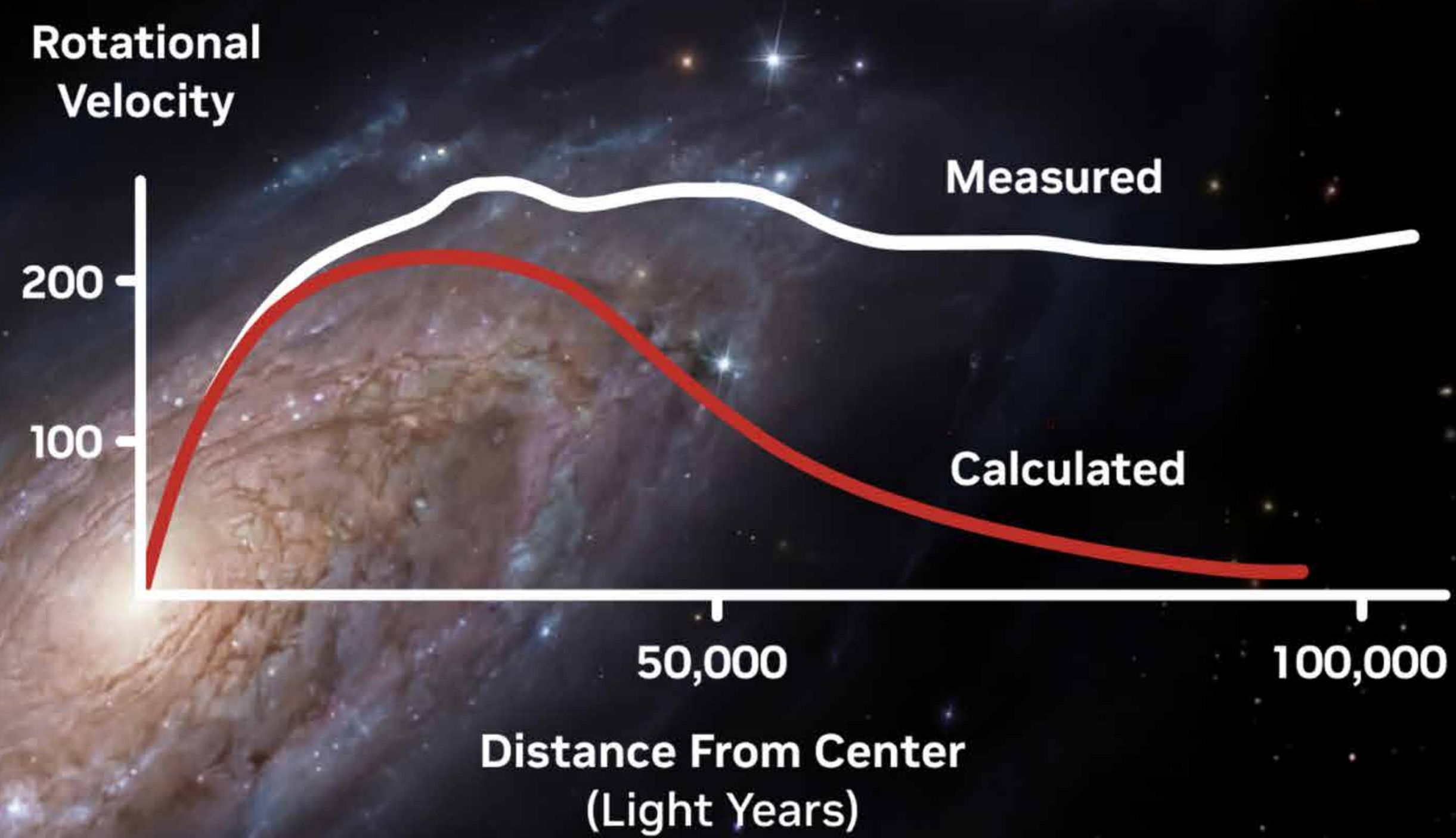
Powered by Siemens Xcelerator





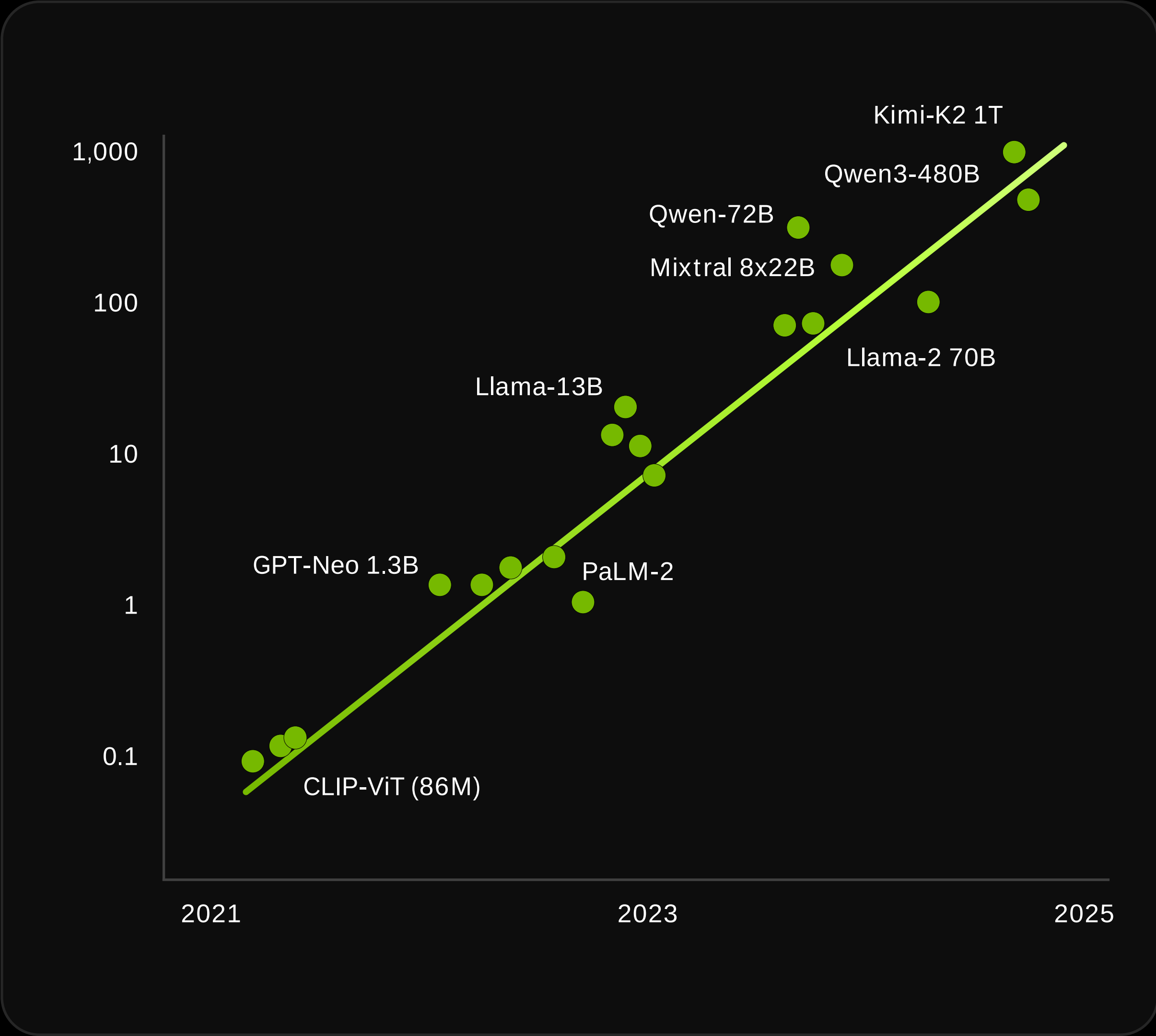
NVIDIA Leads Open Model Ecosystem





Insane Demand for AI Computing

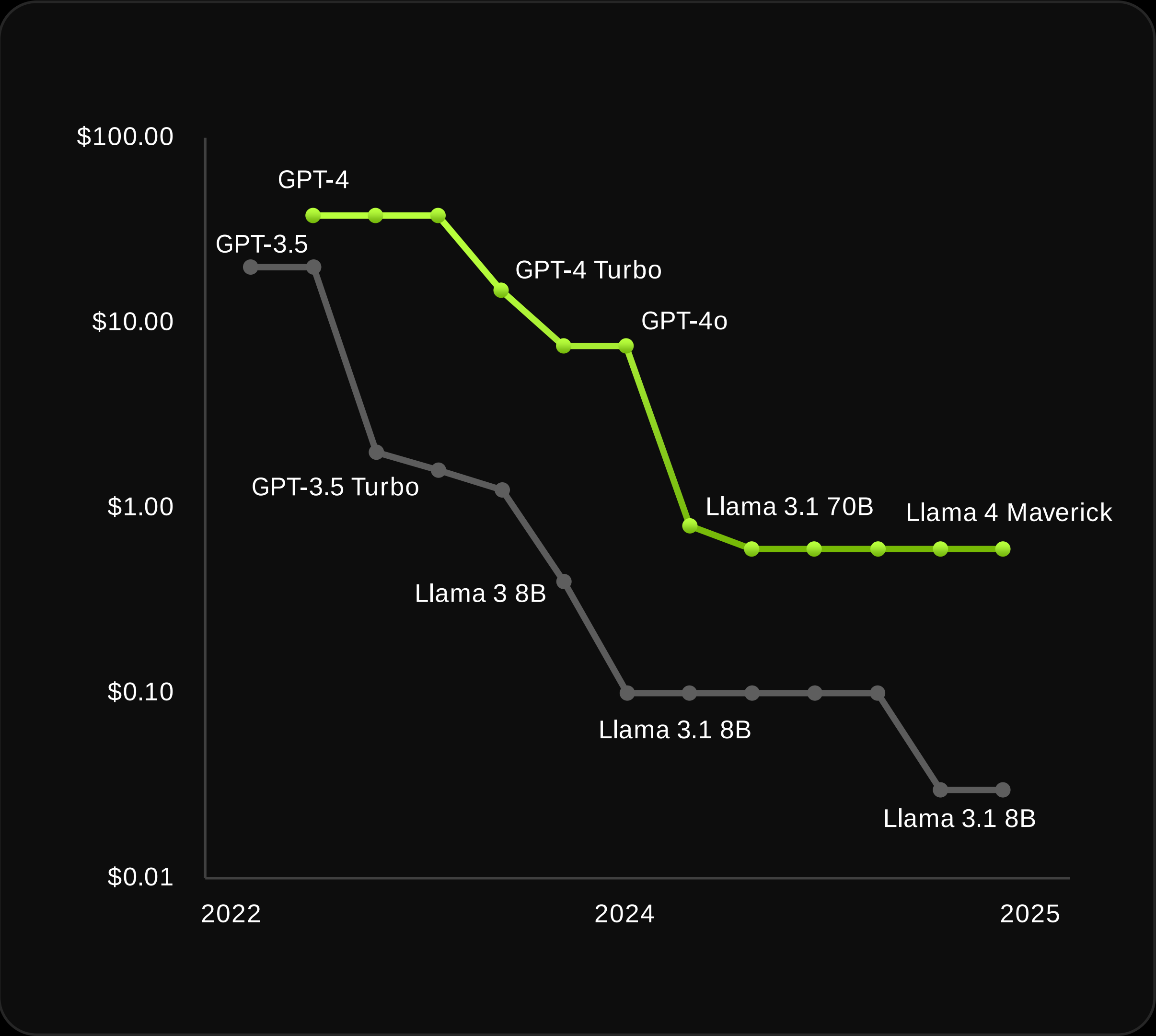
Model Size Growing
10X Parameters Per Year



Test-Time Scaling “Thinking”
5X Tokens Per Year

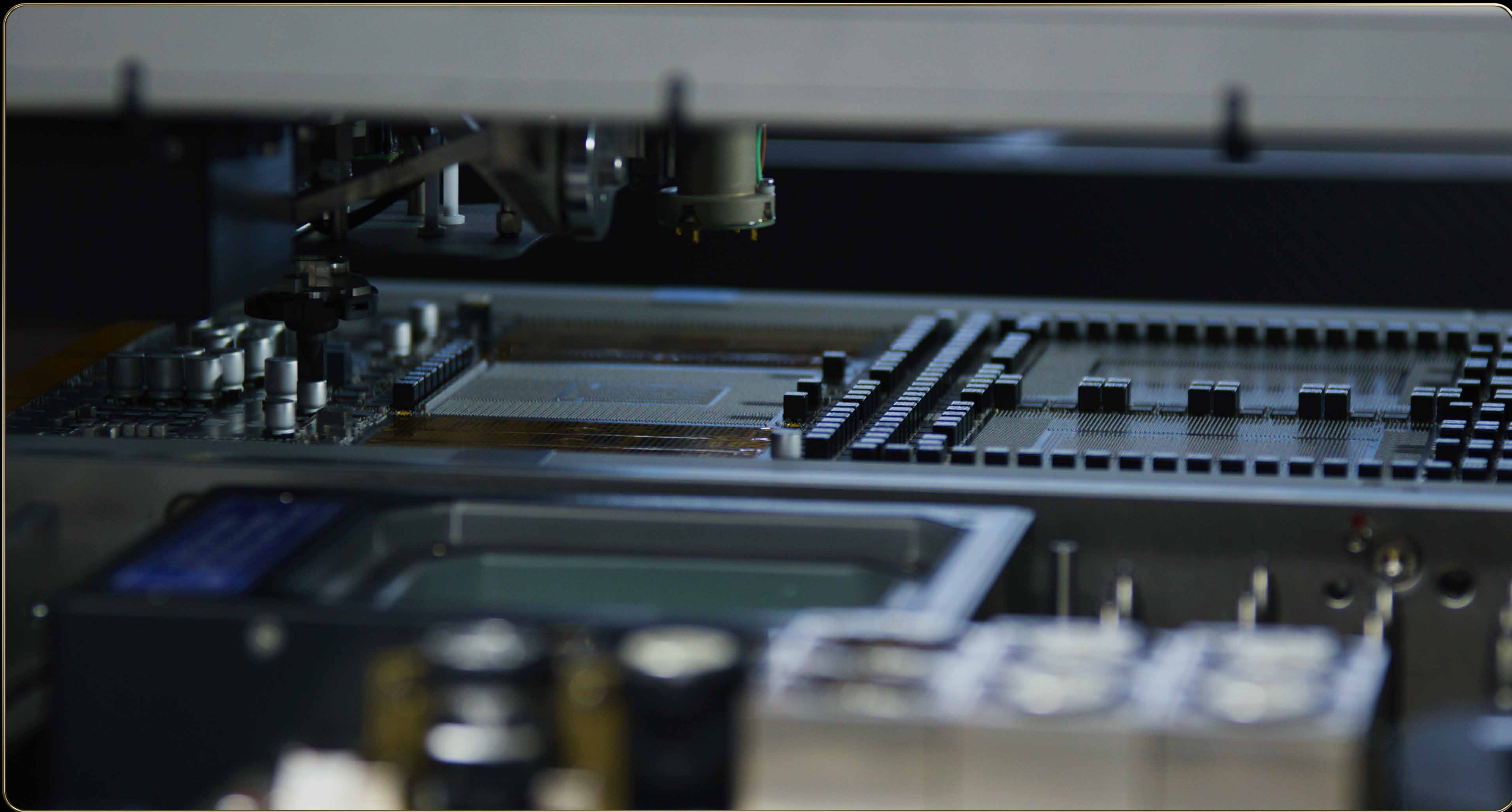


Token Cost
10X Cheaper Per Year



Source : EPOCH AI

Source : Artificial Analysis



NVIDIA Vera CPU

88 NVIDIA Custom Olympus Cores

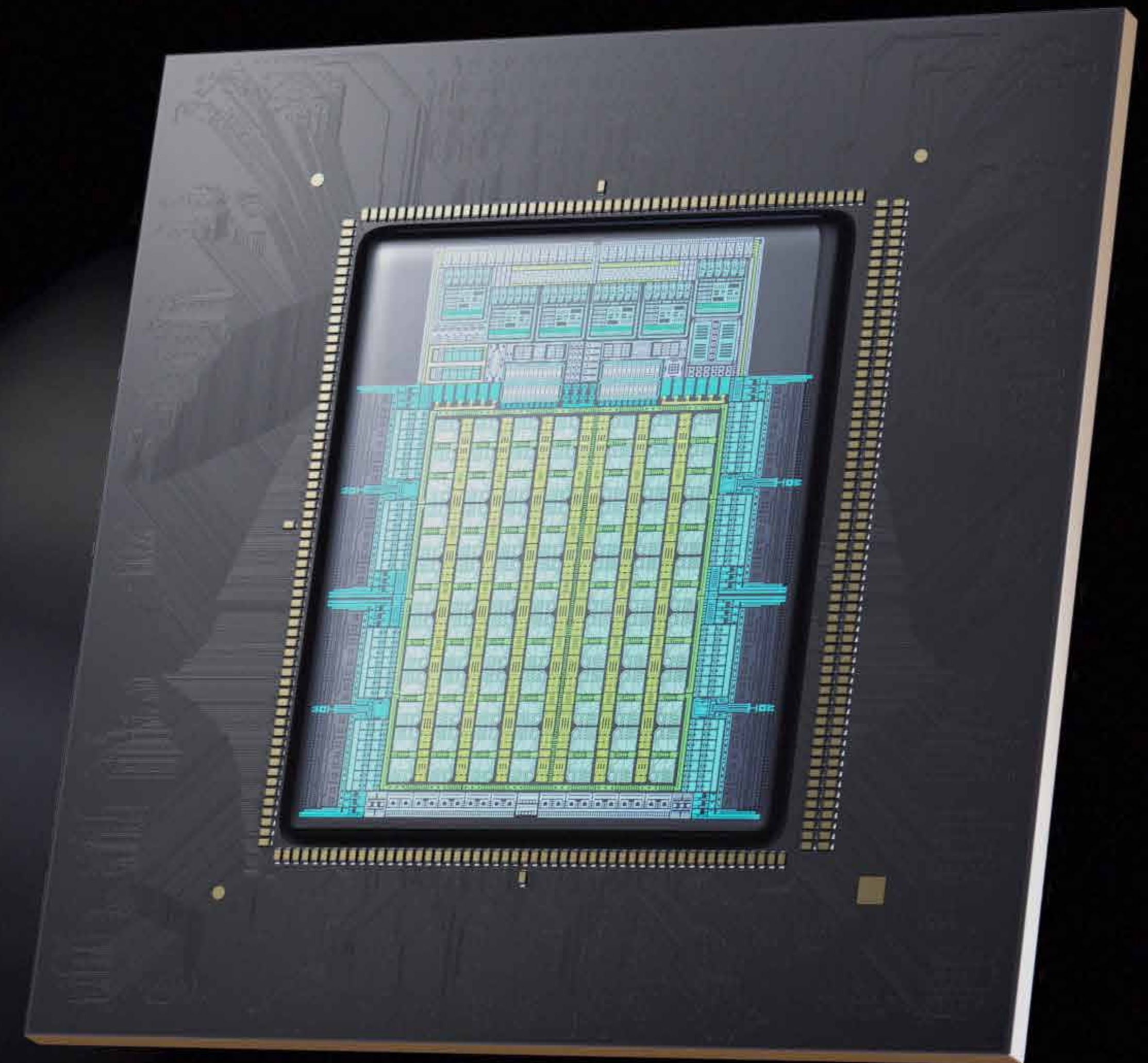
176 Threads with NVIDIA Spatial Multi-Threading

1.8 TB/s NVLink-C2C

1.5 TB System Memory (3X Grace)

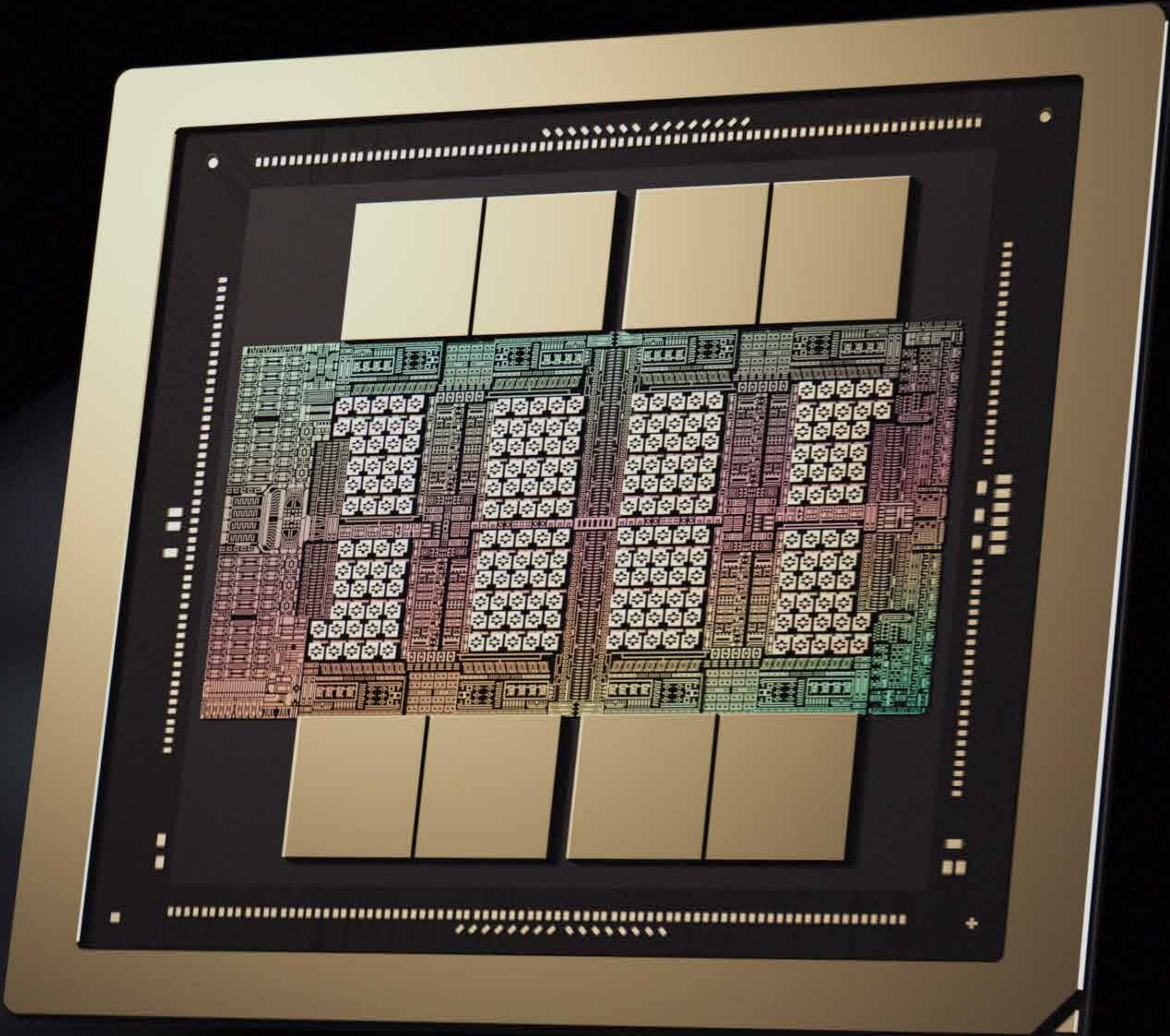
1.2 TB/s LPDDR5X

227 Billion Transistors



NVIDIA Rubin GPU

NVFP4 Inference	50 PFLOPS	5X Blackwell
NVFP4 Training	35 PFLOPS	3.5X
HBM4 Bandwidth	22 TB/s	2.8X
NVLink Bandwidth per GPU	3.6 TB/s	2X
Transistors	336 Billion	1.6X



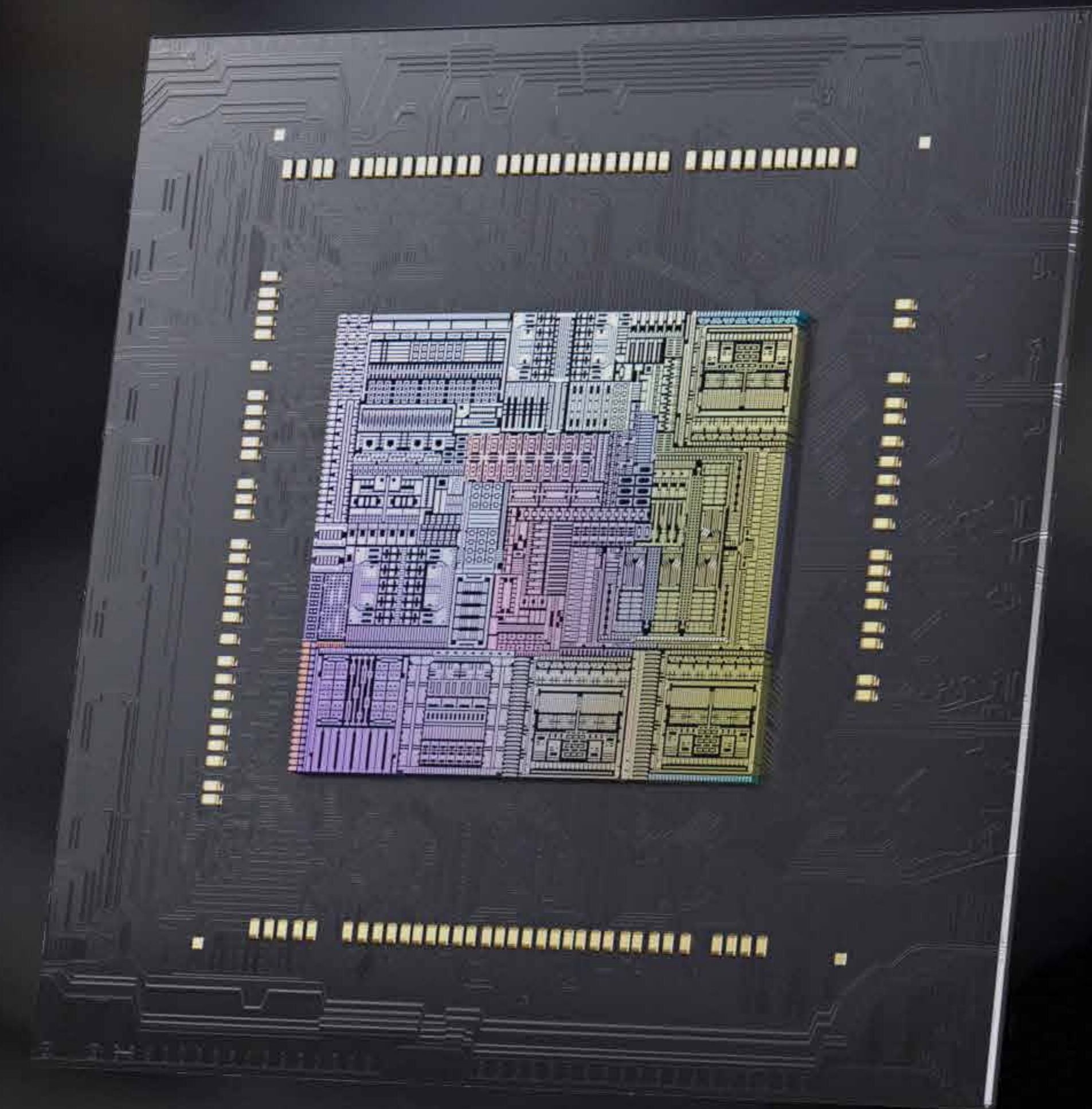
NVIDIA ConnectX-9 Spectrum-X SuperNIC

800 Gb/s Ethernet with 200G PAM4 SerDes

Programmable RDMA and Data Path Accelerator

State-of-the-Art Security
Line-Speed Encryption | Security Enclave and Attestation
CNSA and FIPS Certified

23 Billion Transistors



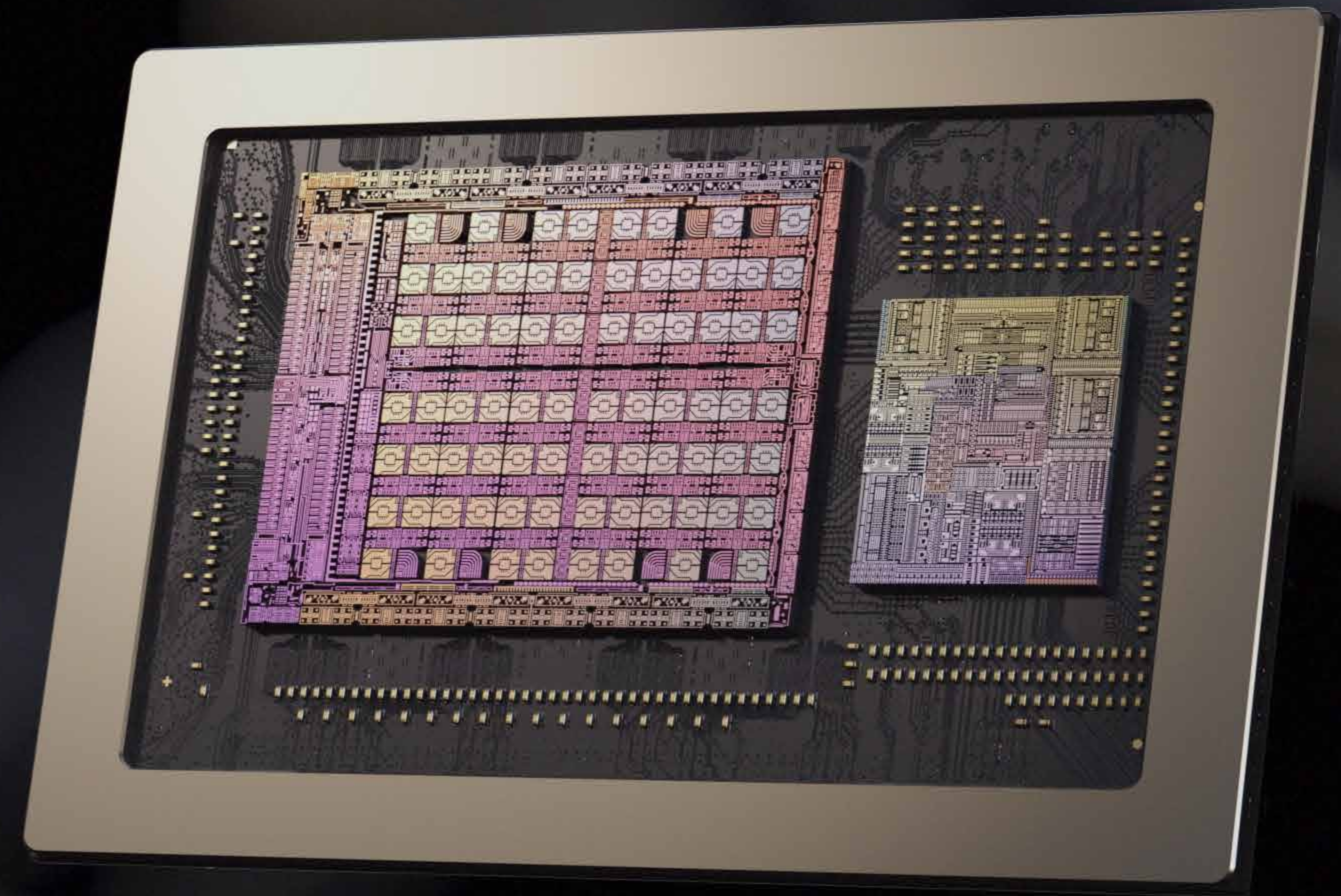
NVIDIA BlueField-4

800G Gb/s DPU for SmartNIC and Storage Processor

64 Core Grace CPU with ConnectX-9

2X Networking		vs BF3
6X Compute		
3X Memory BW		

126 Billion Transistors



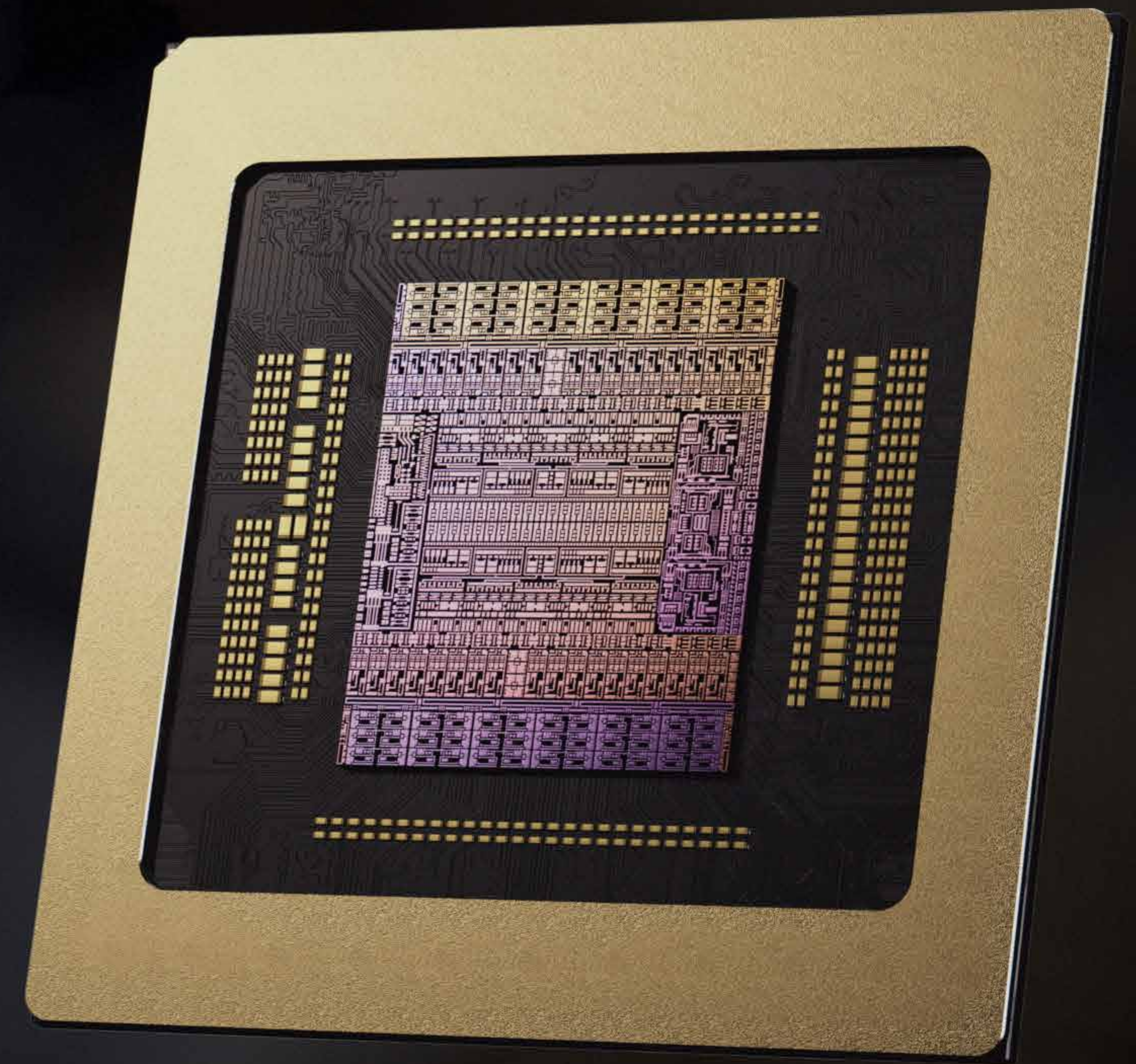
NVIDIA NVLink 6 Switch

Scale-Up Fabric with 3.6 TB/s Per-GPU All-to-All BW

400G SerDes

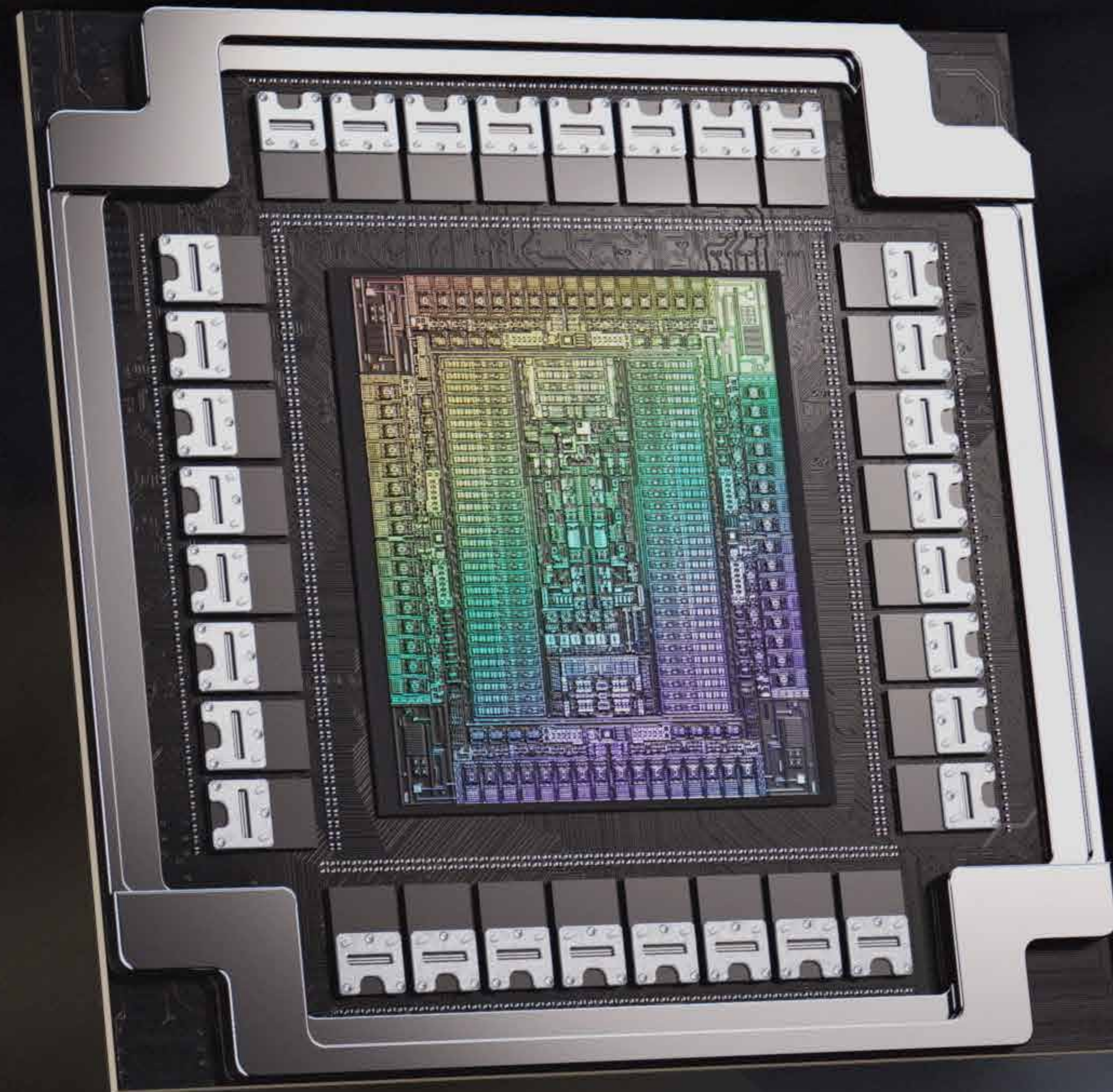
In-Network SHARP Collectives

108 Billion Transistors



NVIDIA Vera Rubin NVL72

NVFP4 Inference	3.6 EFLOPS	5X Blackwell
NVFP4 Training	2.5 EFLOPS	3.5X
LPDDR5X Capacity	54 TB	3X
HBM Capacity	20.7 TB	1.5X
HBM4 Bandwidth	1.6 PB/s	2.8X
Scale-Up Bandwidth	260 TB/s	2X
Transistors	220 Trillion	1.7X



NVIDIA Spectrum-X Ethernet Co-Packaged Optics

102.4 Tb/s Scale-Out Switch Infrastructure

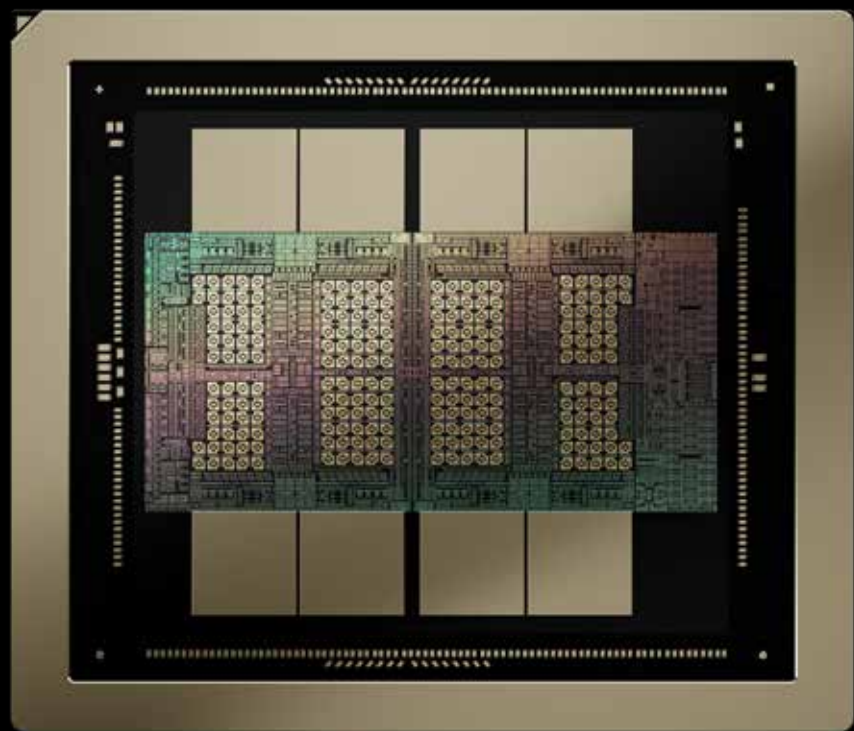
Co-Packaged 200G Silicon Photonics

128 Ports of 800 Gb/s

512 Ports of 200 Gb/s

352 Billion Transistors

Context is the New Bottleneck – Storage Must be Rearchitected



HBM

Memory

Rack SSD

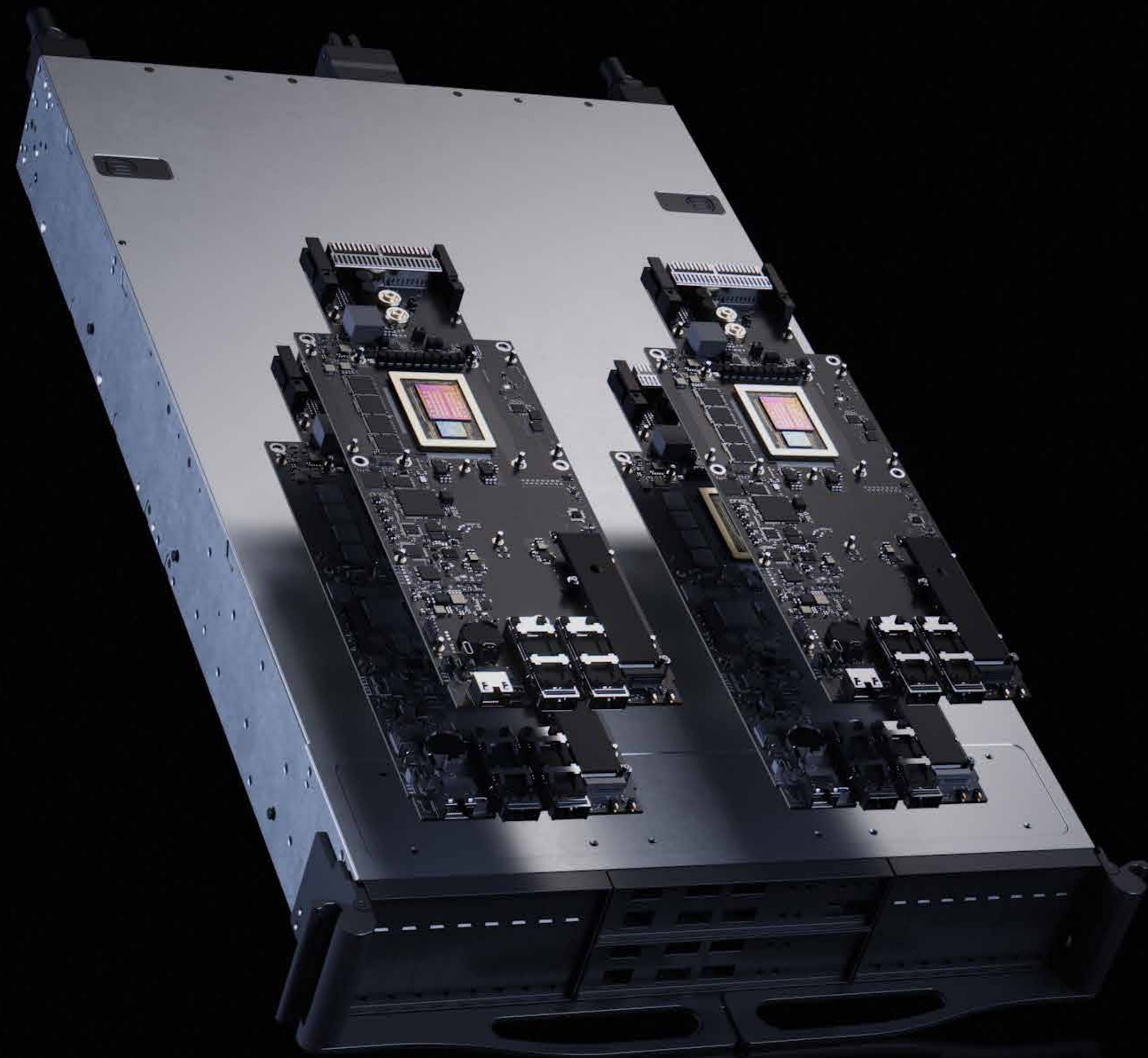
Network SSD

Model Size Growing

Context Length Growing

Multi-Turn Conversations - Context Accumulates

More Concurrent Users and Sessions



NVIDIA Context Memory Storage Platform

A New POD-Level Context Memory Platform

BlueField-4 Data and Storage Processor

Spectrum-X Ethernet RDMA Infrastructure

5X Higher Tokens per Second

5X Higher Power Efficiency

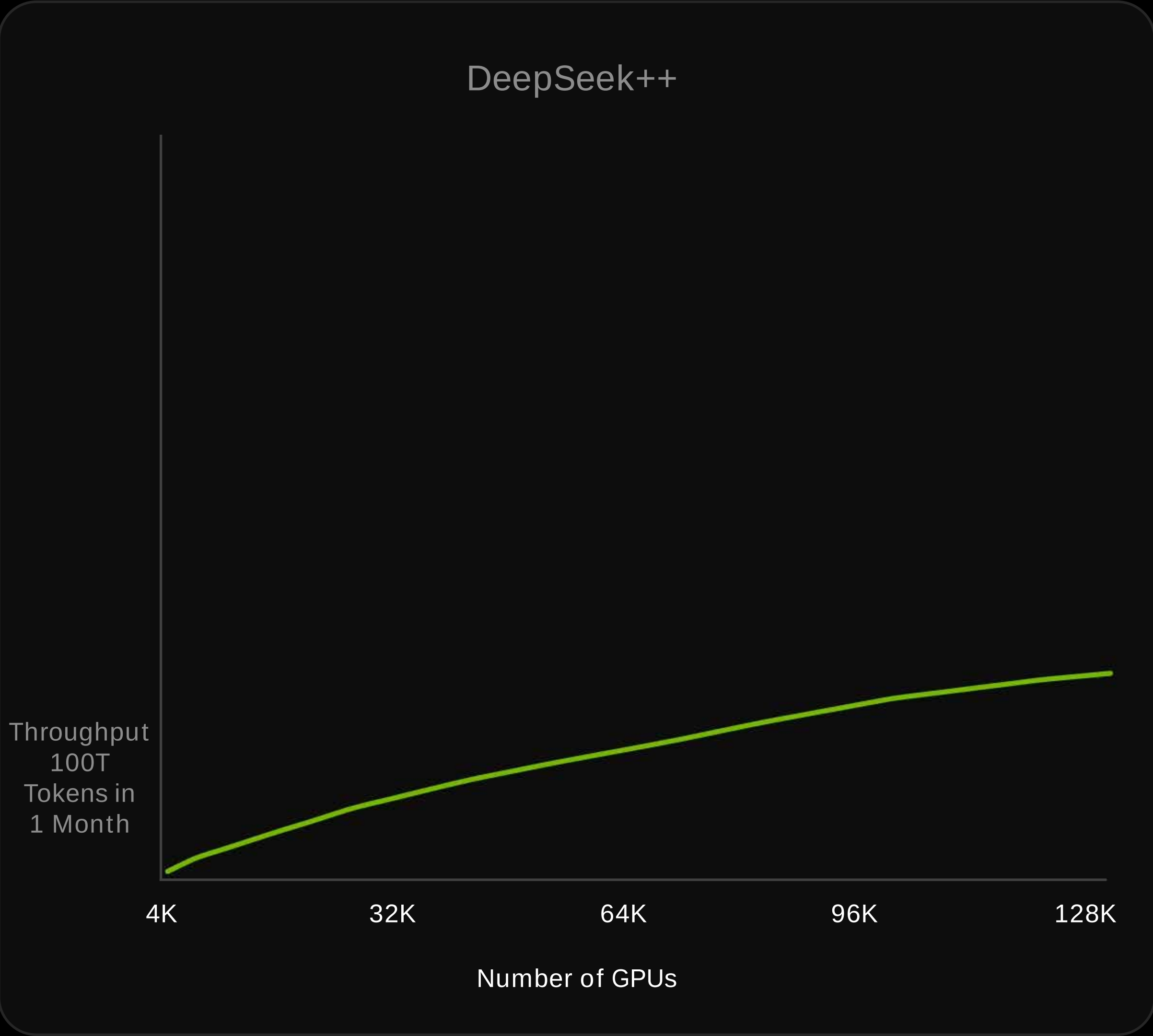


Vera Rubin
Six New Chips – One Giant Leap to the Next Frontier

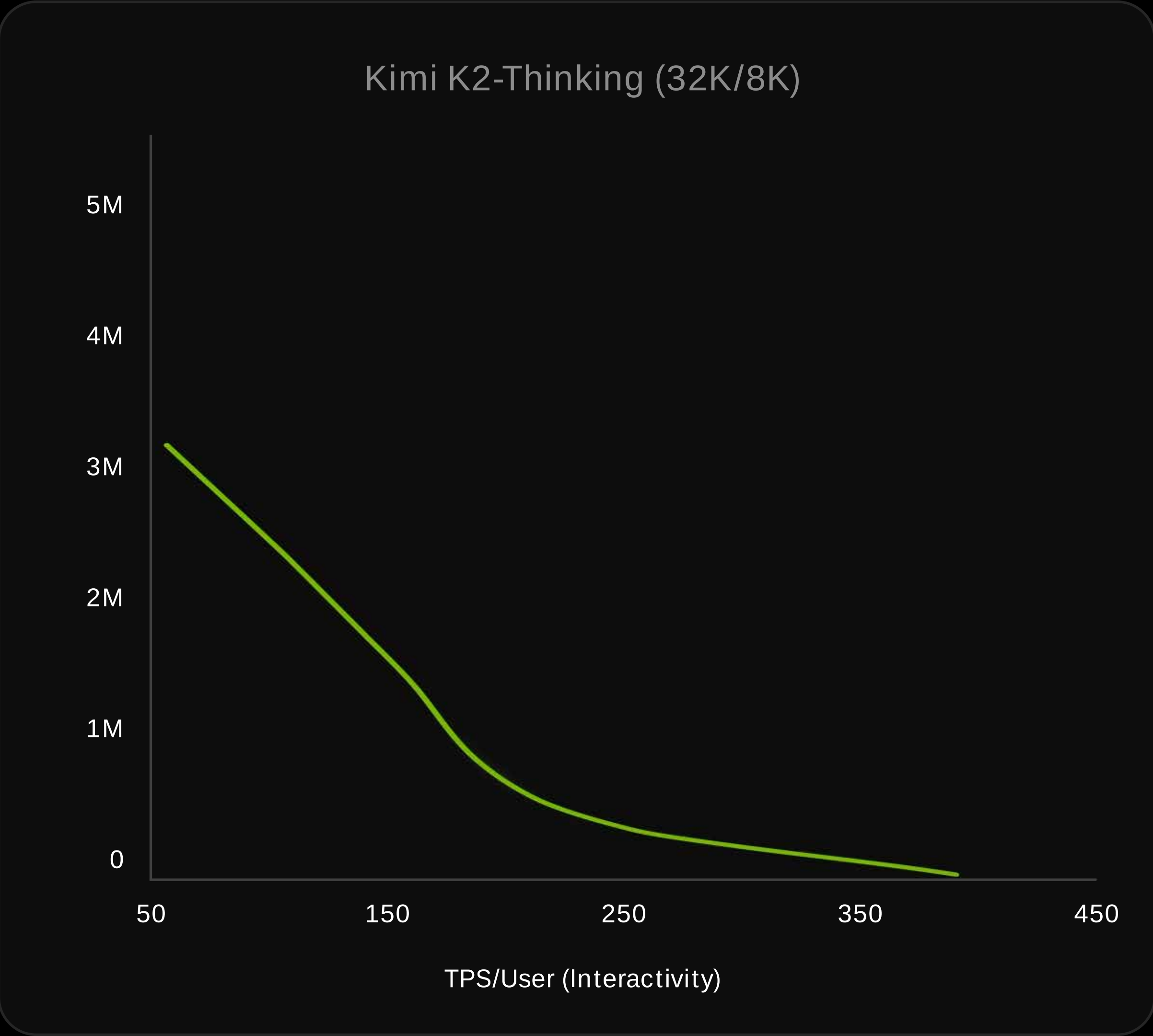


Six New Chips – One Giant Leap to the Next Frontier

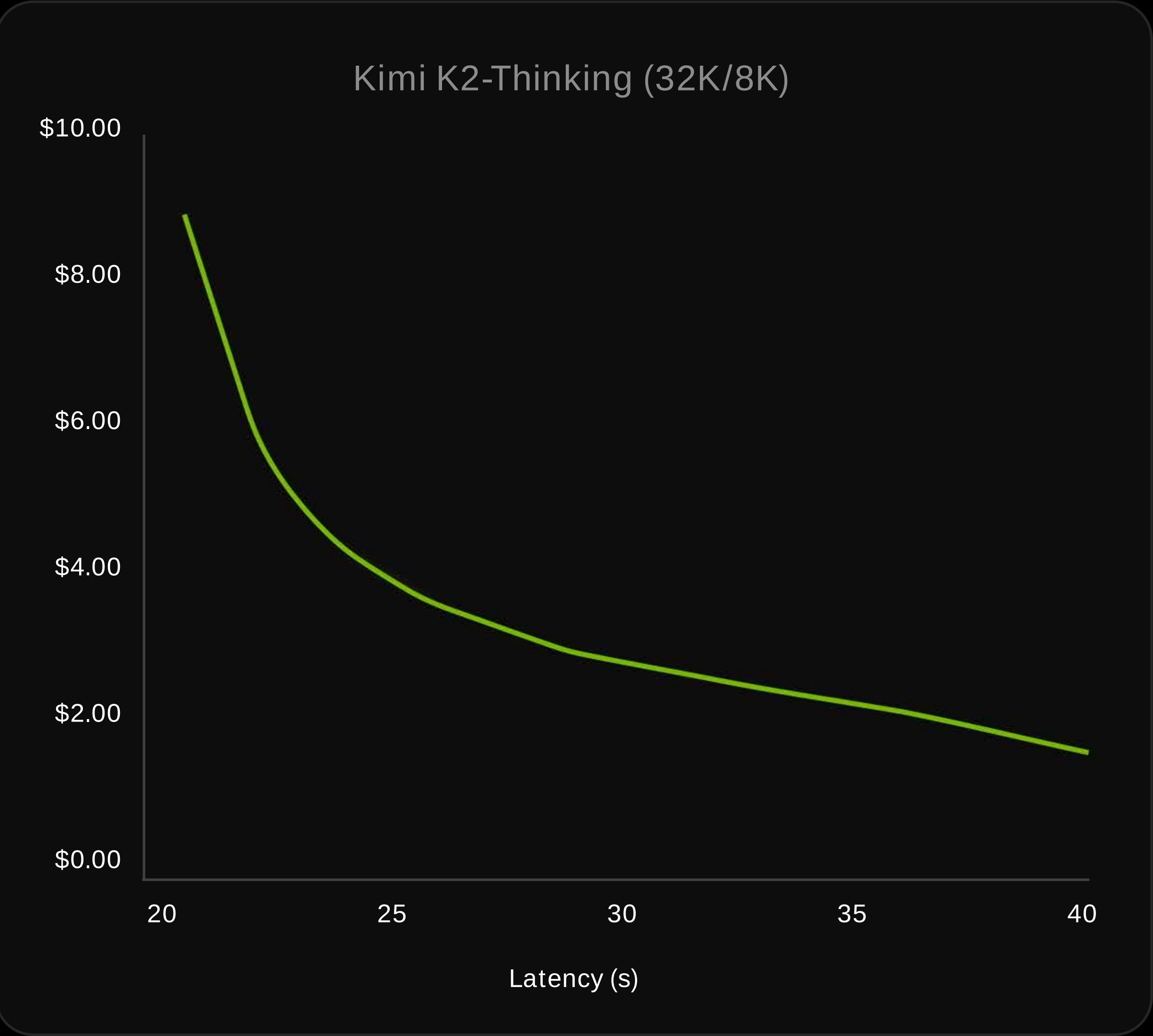
Time to Train
1/4th Fewer GPUs



Factory Throughput
Up to 10X More Tokens



Token Cost
1/10th Lower Cost

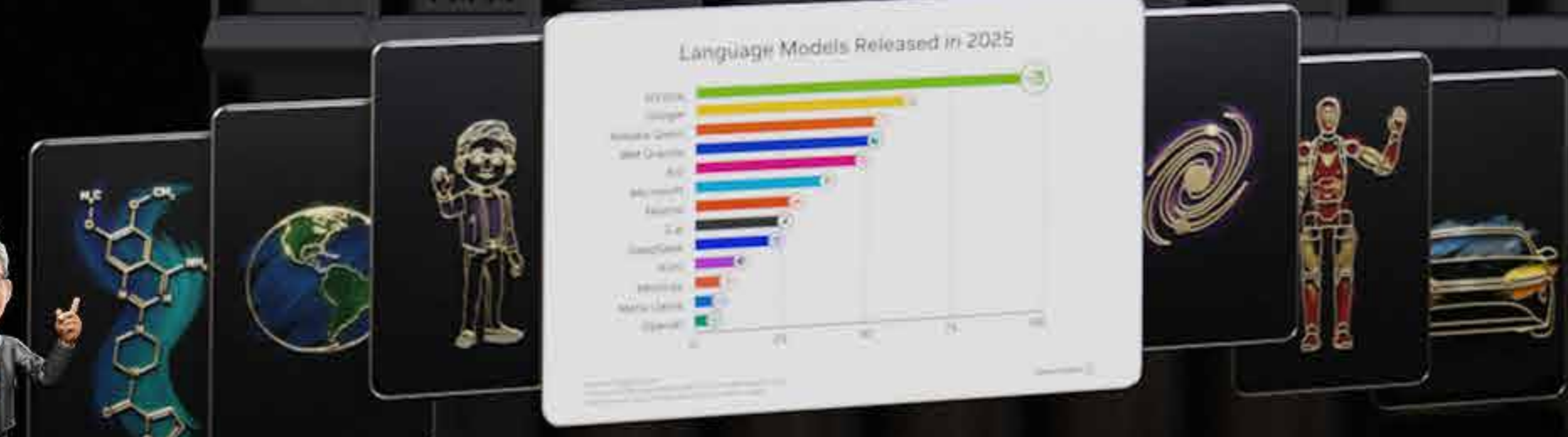


Blackwell NVL72 Rubin NVL72

NVIDIA

One Platform for Every AI

Vera Rubin



Open Models



Isaac

Alpamayo

Agentic AI

