

US Semiconductors

CPU Reborn: agentic drives growth, but watch for competition

Price Objective Change

Raising server CPU forecasts, ests/POs for AMD, INTC, ARM

We raise server CY26-30E server CPU industry sales forecasts to +15% CAGR towards \$71bn from +13%/ \$59bn driven by: 1) accelerating spend from hyperscalers to address emerging agentic AI workloads, and 2) stronger CPU pricing on supply tightness and improved mix. We further raise our estimates and POs on Buy-rated AMD (to \$310 from \$280), Neutral-rated ARM (to \$180 from \$155) and Underperform-rated INTC (to \$48 from \$40). Near-term, we see AMD as best-positioned with its leading pipeline (current-gen Turin, upcoming Venice). INTC benefits also, but we see the recent ~\$110bn market cap addition over 3 months as excessive. Longer-term, we believe CPU market gets more competitive with several ARM-based variants including NVDA Vera, ARM AGI, and many internal cloud options (Amazon Graviton, Google Axion, Meta/ARM, MSFT Cobalt).

AMD: reiterate Buy, raise estimates/PO to \$310

We continue to like AMD's strong CPU product pipeline and incumbency in cloud server CPUs, as well as the upcoming first 2.0-2.5 GW of compute capacity (GPUs, CPUs, NICs) buildout beginning 2H26 at Meta, OpenAI, and others (potentially Anthropic). Each GW installment represents \$15-20bn of net revenue for AMD, resulting in >60% YoY Data Center growth in both CY26/27E. Our new \$310 PO is based on 29x CY27E PE (vs. 27x prior) on larger addressable market.

INTC: raise estimates/PO to \$48, remain Underperform

We like INTC's improving opportunity in server CPU, EPS accretion from recent fab buy-out and scarcity value of its manufacturing & packaging assets for external customers. However, the ~\$110bn market cap increase over the last 3 months seems excessive to us. On a conceptual sum-of-parts basis: 1) Any external manufacturing business, wafers or packaging should not be valued more than the ~10x EV/S of leader TSMC. We model \$3bn/\$6bn in CY27/28E (MediaTek TPU packaging, Apple M processors, other ASICs), and anything more will require incremental capex, shifting out foundry breakeven & wiping cash accretion, 2) Core Products can generate \$2-\$2.50 but should also be valued no more than the ~16x CY28E PE of profitable fabless peers AMD/AVGO/NVDA. Last, we are skeptical of durable benefits from Tesla Terafab beyond some IP sharing & packaging, since Tesla is creating a rival fab, and since it already has foundry relationships with Samsung & TSMC. Our new \$48 PO is now based on 3.7x CY28E EV/S (vs. 3.5x CY27E) as we look out into further foundry opportunity, toward high-end of 1.7x-4x historical range.

ARM: reiterate Neutral, CPU chip TAM limited

We like the upside in ARM royalties from rising NVDA/cloud CPU IP adoption that could offset headwinds in mid/low-end Android sales. However, on the chip side, we are yet to see any third-party benchmarking of ARM AGI CPU, which enters a very crowded field of incumbents and larger rivals. ARM's recent CPU event featured Softbank/Stargate and Meta, but did not feature other hyperscale customers. Successes of NVDA Vera and x86 is also competition for any new merchant incumbent in our view, limiting ARM's actual opportunity. Medium/longer-term we expect ARM to leverage Softbank Graphcore assets to enter AI accelerators, another growing but crowded field. Our new \$180 PO is based on 69x CY27E EPS (vs. 60x prior) on larger addressable market, but still within 2x-3x PEG framework for IP peers.

BofA Securities does and seeks to do business with issuers covered in its research reports. As a result, investors should be aware that the firm may have a conflict of interest that could affect the objectivity of this report. Investors should consider this report as only a single factor in making their investment decision. Refer to important disclosures on page 17 to 21. Analyst Certification on page 16. Price Objective Basis/Risk on page 15.

12960058

Timestamp: 17 April 2026 12:53AM EDT

17 April 2026

Equity
United States
Semiconductors

Vivek Arya
Research Analyst
BofAS
vivek.arya@bofa.com

Duksan Jang
Research Analyst
BofAS
duksan.jang@bofa.com

Michael Mani
Research Analyst
BofAS
michael.mani@bofa.com

Liam Pharr
Research Analyst
BofAS
liam.pharr@bofa.com

See inside for Glossary



Contents

The Rise of CPUs in AI Inference	3
CPUs in Training	4
CPUs in Inference	4
CPUs in Prefill Inference	5
CPUs in Decode Inference	6
Server CPU Update	8
Server CPU growth towards Siphon+ TAM	8
Server Unit Growth vs. Server CPU Unit Growth	9
Server CPU Competition	9
Server CPU Roadmap	10
End Market Poster	11



The Rise of CPUs in AI Inference

AI training and AI inference perform fundamentally different workloads, and thus require fundamentally different hardware/semis with varying degrees of roles for CPUs, GPUs, memory, and other compute/networking components.

Particularly, GPUs or accelerators typically excel in parallel compute-intensive workloads, such as in AI training.

Meanwhile, CPUs excel in sequential and latency-sensitive workloads, as well as in disaggregated architectures where scheduling, controlling, batching, and routing functionalities are critical to the overall system performance. This is particularly the case in AI inference which requires constant back and forth between the head node, compute (accelerators), and memory (KV cache).

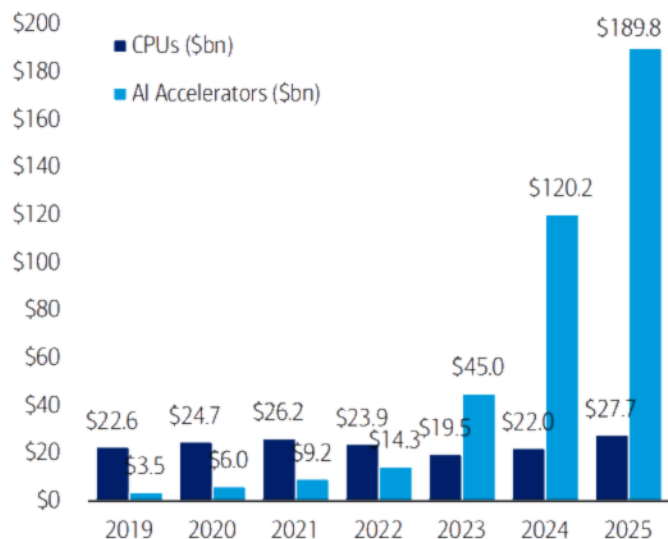
CPUs vs. Accelerators During AI Training Era

As such, during the era of AI training, we saw a significant decline in the role/value of CPUs, being overtaken by GPUs and other accelerators as seen in the charts below. Specifically, between 2022 and 2025, the AI accelerator market grew at a +137% CAGR due to a rise of tremendous training demand, vs. server CPUs at only +5% CAGR.

By 2025, AI accelerators have become 87% of total compute spend and CPUs just 13%, versus accelerators at just 26% (CPUs 74%) in 2021 when the training market had just begun to emerge.

Exhibit 1: AI accelerator market grew at a +137% CAGR between 2022-2025 from a tremendous rise of AI training demand, vs. CPUs at just +5% CAGR

CPUs vs. AI Accelerators TAM (\$bn)

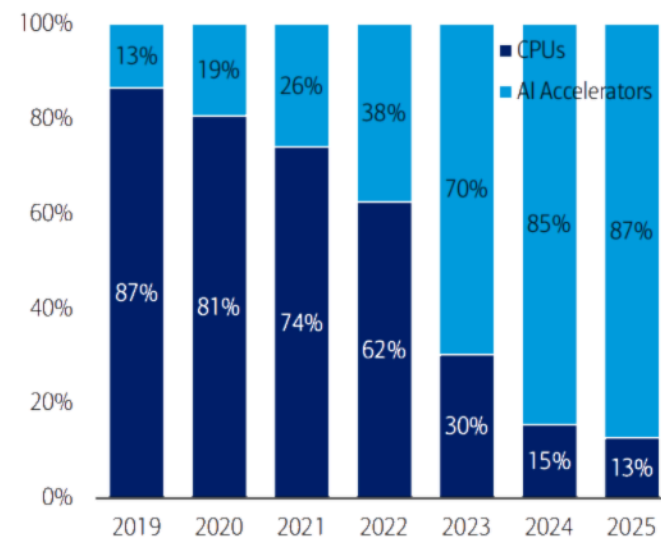


Source: BofA Global Research estimates, Mercury Research

BofA GLOBAL RESEARCH

Exhibit 2: AI accelerators as a % of total data center compute spend has risen to 87% by 2025 (CPUs 13%), from just 26% in 2021 (CPUs 74%) before the rise of AI compute demand

CPUs vs. AI Accelerators TAM (%)



Source: BofA Global Research estimates, Mercury Research

BofA GLOBAL RESEARCH

The shift in value from traditional compute to accelerated compute has brought forth this massive transition, with the CPU market almost not growing at all between 2021 (\$26.2bn) and 2025 (\$27.7bn) given limited growth in cyclical traditional/enterprise server demand as well.

However, as the market slowly transitions from AI training to AI inference, we expect to see increasing roles of CPUs again (vs. training era), with the CPU market growing together with other critical hardware components such as GPUs, accelerators, memory, and networking.

CPUs in Training

In training, CPUs are responsible for decompression, tokenization, batching, and ultimately feeding the GPUs. CPUs' main job is to make sure the GPUs/accelerators are utilized as much as possible by constantly preparing and queuing batches, making sure the expensive accelerators are fully in work without bottlenecks.

CPUs control the orchestration and preprocessing of data, i.e. they are typically not involved in the actual heavy processing that AI training requires (mostly done in parallel processing by GPUs/accelerators).

CPUs essentially perform a secondary or supporting role for the GPUs/accelerators in AI training, hence explaining the decline in the % of value in the overall system (vs. GPUs).

CPUs in Inference

In AI inference, the role of CPUs becomes more broad and extensive.

The AI inference process can be broken into three phases: load phase, prefill phase, and decode phase.

In **load or initialization**, models are loaded from the disks into the memory, with performance largely dependent on overall disk (I/O) and CPU speed. Weight loading, quantization, and compilation all occur during the load phase, with the latter two (quantization and compilation) much better to be done offline vs. online to avoid runtime overhead.

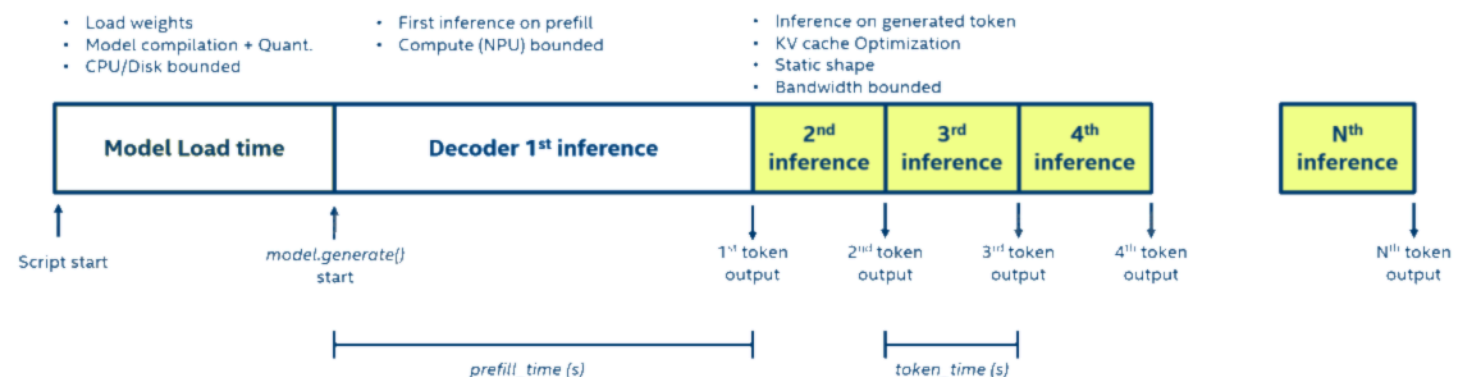
Overall, the role/speed of CPU is key during the load phase of AI inference.

In **prefill**, time to first token (TTFT) and KV cache creation are the two main roles, with both prefill time and first token latency heavily dependent on overall compute/large matrix multiplication abilities, i.e. GPU/accelerators and networking rather than CPUs.

In **decode**, performance is heavily impacted by KV cache reuse, so optimizing KV cache and/or increasing the bandwidth/capacity of system memory is critical, as such with high-bandwidth SRAM or HBM memory, inference storage context memory (i.e. low-latency storage), or even CXL memory (pooling of system memory) methodology.

Exhibit 3: LLM inference process can also be broken into three phases: load phase, prefill phase, and decode phase

LLM Inference Process in Three Phases



Source: Intel

BofA GLOBAL RESEARCH



CPUs in Prefill Inference

In prefill inference, the system runs (i.e. performs compute on) the entire model, builds essential KV cache that could be used during the decode process.

The CPU's role in prefill is quite similar to their role in training. CPUs compute tokenization and format the prompt before GPUs come in for processing. Note prefill is an immensely compute-heavy and compute-bound process, with the primary goal of the CPUs being setting up the system so that GPUs/accelerators are not bottlenecked and are at full utilization.

CPUs allocate and manage memory structures (KV cache setup), handle the arrival of user prompts, batch/route those prompts to the GPUs. As such, the CPUs determine the overall throughput by allocating compute responsibilities to the GPUs.

Specifically, CPUs may be involved in compression (i.e. quantization/tuning, pruning, knowledge distillation, and mix precision) as seen in INTC's Neural Compressor Architecture example below.

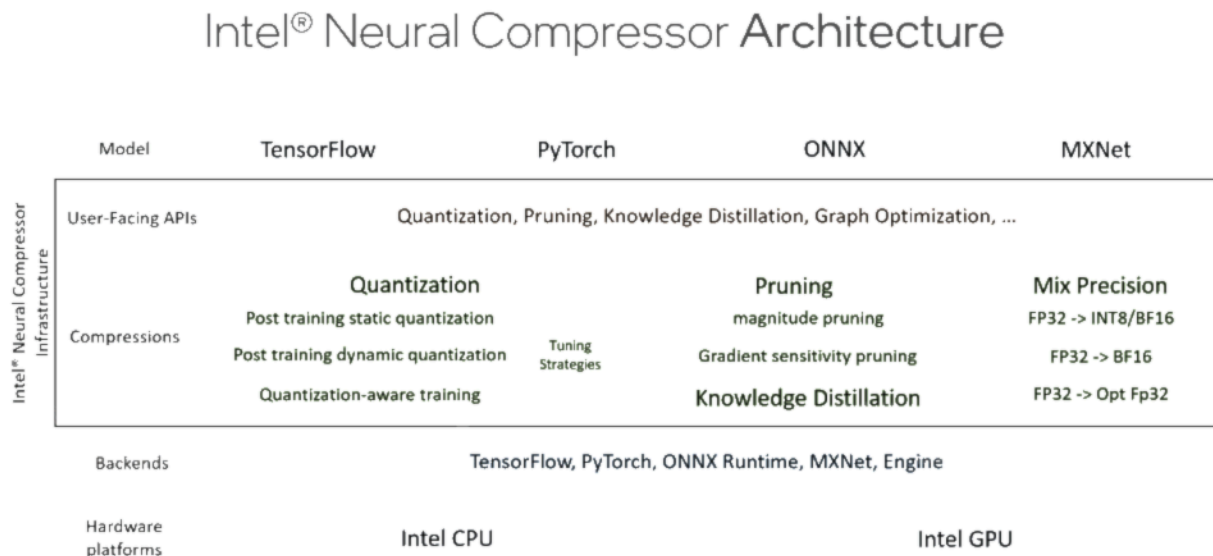
They also help optimize and export between different models in TensorFlow, PyTorch, or ONNX (Open Neural Network Exchange), essentially preparing the overall system ahead of massive parallel compute done by GPUs/accelerators.

While GPUs/accelerators can also help/perform some of these pre-/post-compute techniques, they remain majority CPU workloads as CPUs are best suited for the control flow, orchestration, graph manipulation, and data-handling components.

Overall, CPUs are an important part of the prefill process, though its significance generally remains secondary and a supporting role to the main compute processors of the system – AI accelerators.

Exhibit 4: CPUs may be involved in data compression – quantization, tuning, pruning, knowledge distillation, mix precision, among others – to help improve overall system speed/efficiency in both AI training and AI inference

Intel Neural Compressor Architecture



Source: Intel

CPUs in Decode Inference

In decode inference, the role of CPUs becomes more important.

Recall the decode process is highly sequential and arithmetic, and each new token generated depends on: 1) the entire prompt (from prefill), and 2) all previously generated tokens. In other words, token #2 refers to token #1, and token #3 refers to both tokens #1 and #2.

Role of CPUs in decode

Here, CPUs are responsible for issuing (scheduling) thousands of many small GPU kernels per second with minimal overhead. CPUs also compute the sampling, safety/guardrails checks, logits processing, as well as memory management.

Memory management is particularly critical because KV cache must be allocated, evicted, and distributed effectively for an overall efficient decode process, i.e. KV cache routing.

Lastly, CPUs handle the token-by-token output that the GPUs process, with minimal latency.

Overall, the role of CPUs in decode becomes far more reaching than in AI training or in prefill. Furthermore, CPUs with larger RAM (i.e. DDR/LPDDR) have advantages in providing repeated weight access, and they can also provide for higher KV cache capacity via CXL memory pooling.

Below, we highlight some of key AI workloads that cater to CPUs vs. GPUs vs. CPU+GPU clusters.

Exhibit 5: CPUs excel in document processing/classification, indexing, and captioning, among other AI workloads

Role of CPUs in AI Inference

AI Inference Workload	Good fit for...		
	CPUs	CPUs + PCIe-Based GPU	GPU Clusters
Document processing and classification	✓		
Data mining and analytics	✓		✓
Scientific simulations	✓		
Translation	✓		
Indexing	✓		
Content moderation	✓		
Predictive maintenance	✓		✓
Virtual assistants	✓	✓	
Chatbots	✓	✓	
Expert agents	✓	✓	
Video captioning	✓	✓	
Fraud detector		✓	✓
Decision-making		✓	✓
Dynamic pricing		✓	✓
Audio and video filtering		✓	✓
Financial trading			✓
Telecommunications and networking			✓
Autonomous systems			✓

Source: AMD

BofA G. R. ARCH

CXL memory pooling

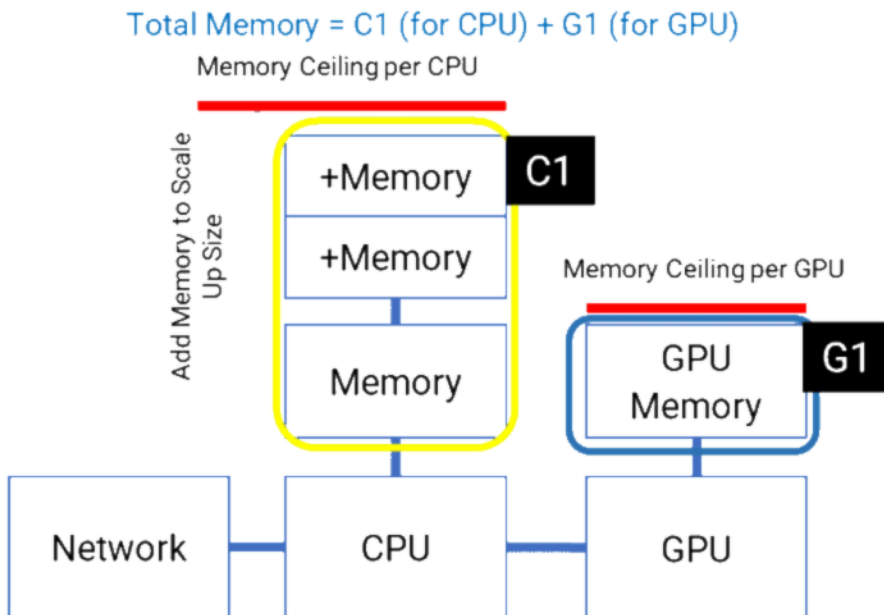
Besides CPU's compute capabilities, its onboard SRAM and attached DDR/LPDDR memory can also be critical to improving overall decode performance via memory expanding or memory sharing/pooling of CXL.

Importantly, CXL enables the pooling of multiple memory types in a system together, such as GPU memory (HBM), host/CPU-attached memory (DDR/DIMMs), non-AI GPU memory (GDDR), and any other memory.

With expanded system memory, the model can utilize: 1) much larger KV cache capacity beyond just expensive and limited AI GPU HBM, 2) generally low latency (200-500ns) vs. slower NVMe and distributed storage options, and 3) memory-semantic access with near-DRAM semantics.

As a result, CPUs with high memory capacity can provide additional benefits to improving overall system beyond its traditional batching, routing, scheduling, and post-processing tasks.

Exhibit 6: Total system memory could equal CPU memory + GPU memory + any other memory
CXL Memory Expansion/Pooling



Source: Astera Labs

BofA GLOBAL RESEARCH

Server CPU Update

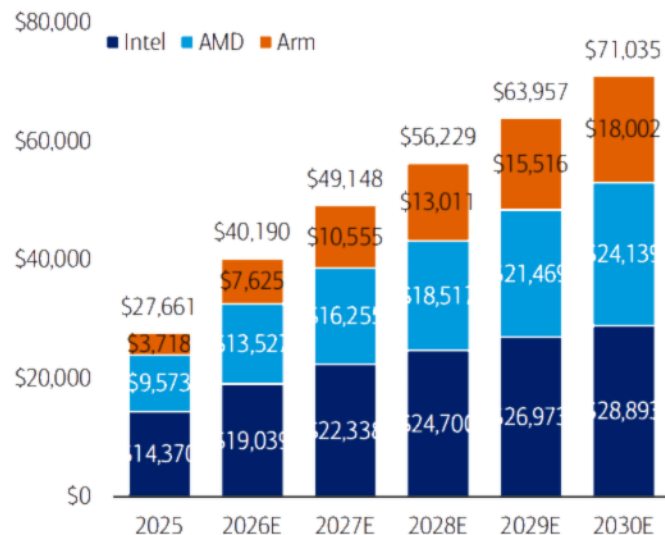
Server CPU growth towards \$70bn+ TAM

As we noted above and in [our 'CPU in AI inference' primer report](#), we believe CPUs are increasingly becoming essential in sequential and latency-sensitive workloads, and are an integral part of overall AI infrastructure. Particularly, CPUs excel in disaggregated architectures where scheduling, controlling, batching, and routing functionalities are critical to the overall system performance. This is particularly the case in agentic AI inference which requires constant back and forth between the head node, compute (accelerators), and memory (KV cache).

Overall, we expect CPUs to make up a consistent ~5% share of overall AI data center TAM (~\$1.4Tn), growing at +21% CAGR toward \$70bn+ TAM by CY30E from just ~\$28bn in CY25. This share is modestly faster than our previous outlook of 4-5% of TAM, given the expanding role of CPUs in agentic AI.

Exhibit 7: We expect all three major server CPU vendors to see expanding CPU sales from rising CPU demand (agentic AI, etc.) and overall TAM increase between 2024-2030

Server CPU TAM (\$mn) and Sales by Vendor through 2030

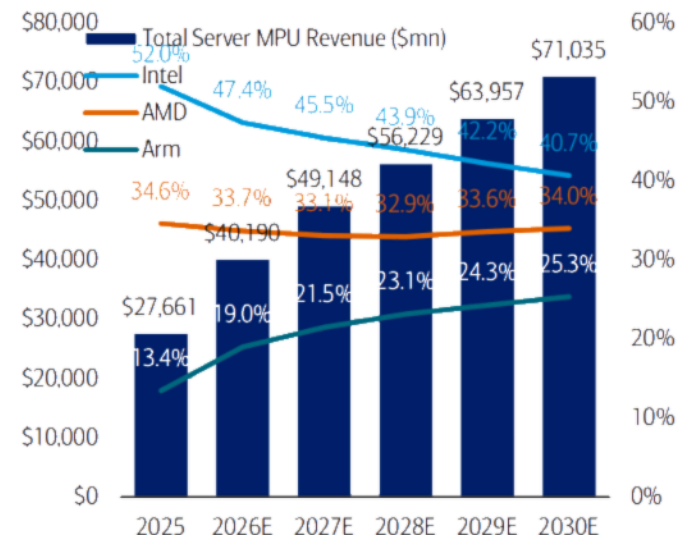


Source: BofA Global Research estimates, Mercury Research

BoFA GLOBAL RESEARCH

Exhibit 8: We expect INTC to generally lose share over time given strong product pipeline and increasing traction (AMD in enterprises, ARM in cloud/AI) at competitors AMD/ARM

Server CPU TAM (\$mn) Market Share Outlook through 2030



Source: BofA Global Research estimates, Mercury Research

BoFA GLOBAL RESEARCH

Given **INTC's** incumbency in sticky enterprise customers (35-40% of total workload today), we believe INTC server sales can also generally increase over time, albeit slower than overall market (+15% CAGR vs. +21% market) given increasing competitive pressures. However, improving INTC's product roadmap (on new 18A/14A nodes) remains a priority given recent concerns around the mainstream 8-channel Diamond Rapids (18A) product cancellation.

We expect **AMD** to continue to lead in cloud (particularly public cloud instances) while expanding its exposure in traditional enterprise market, with sales growing to ~\$24bn by CY30E at ~34% share (+20% CY25-30E CAGR).

Lastly, we think **ARM** could see its share expand from ~13% in CY25 to ~25% in CY30E, on the back of multiple new product ramps from multiple new vendors (ARM, QCOM, NVDA, Ampere, etc.). We expect ARM server CPU total sales to reach ~\$18bn by CY30E, growing at +37% CY25-30E CAGR.

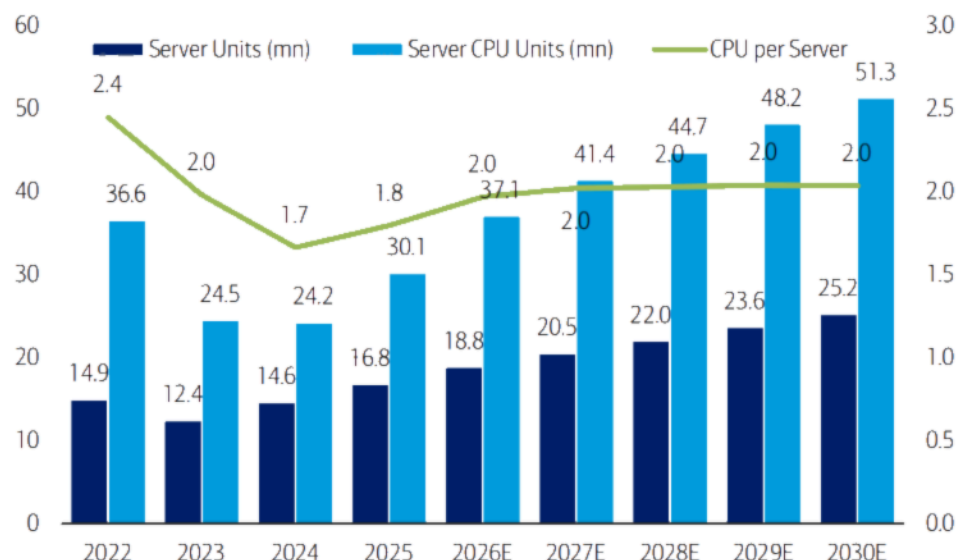


Server Unit Growth vs. Server CPU Unit Growth

The key driver of our server CPU TAM growth is growing server unit TAM to ~25mn units by CY30E from just ~17mn units in CY25, +8% CAGR. We particularly forecast strong server unit growth in CY25-27E from rising agentic AI demand, as well as shifting of capacity from PC/smartphone-oriented chips toward server-oriented chips.

Moreover, we expect the number of CPUs per server to generally remain at 2.0 CPUs through CY30E, despite a modest rise in 1-socket CPU server demand recently.

Exhibit 9: We expect Server units to grow at +8% CAGR CY25-30E vs. Server CPU units +21% CAGR
Server Unit Outlook vs. Server CPU Outlook



Source: BofA Global Research estimates, IDC, Mercury Research

BofA G. DBA, RESEARCH

Server CPU Competition

Meanwhile, we highlight the server CPU market is getting very crowded.

There are incumbents in both x86 and ARM camps who have a very wide breadth of portfolio and established software/ecosystem. INTC specifically caters well to enterprise/telco customers. AMD is the leader for cloud offerings. There are also new ARM-based merchant providers ARM and QCOM who aim to further compete in the cloud space. Finally, there are also hyperscalers who have their own customized CPUs for internal/AI workloads.

As such, we flag it could especially difficult for new merchant solution vendors such as **ARM** and **QCOM** who lack the years-worth of established software/ecosystem, as well as the customization capability of internal solution vendors (hyperscalers who develop their own Arm-based chips).

Exhibit 10: New merchant-based server CPU vendors such as QCOM and ARM could face tough competition and limited (merchant) TAM once their products launch in 2027+

Key Server CPU Providers and Products

Vendor	INTC	AMD	Google	AWS	Microsoft	Alibaba	NVDA	QCOM	ARM	SoftBank/Ampere
Architecture	x86	x86	ARM	ARM	ARM	ARM	ARM	ARM	ARM	ARM
Product	Xeon	EPYC	Axion	Graviton	Cobalt	Yitian	Vera	TBD	AGI	TBD
First Launch	5yr+	5yr+	4Q24	5yr+	4Q24	5yr+	2H26	2027?	2027	2027?
Merchant/Custom	Merchant	Merchant	Custom	Custom	Custom	Custom	Merchant	Merchant	Merchant	Merchant
Key Customers	Enterprise	Cloud	Internal	Internal	Internal	Internal	CoreWeave	HUMAIN	Meta, Stargate?	Neoclouds?

Source: BofA Global Research estimates

BofA GLOBAL RESEARCH



Server CPU Roadmap

Below, we present our expected product roadmap of major merchant server CPU vendors through CY28E.

Overall, we expect INTC to ramp its 18A-based Diamond Rapids CPU throughout 2026, AMD to launch its 2nm-based Venice in 2H26, NVDA to ramp its 3nm-based Vera CPU alongside its Rubin GPU platform, and ARM to launch its new 3nm-based AGI processor in 2027.

Exhibit 11: We expect INTC to ramp its 18A-based Diamond Rapids CPU throughout 2026, and AMD to launch its 2nm-based Venice in 2H26

Major CPU Vendor Roadmaps

BofA Securities		1Q24	2Q24	3Q24	4Q24	1Q25	2Q25	3Q25	4Q25	1Q26	2Q26	3Q26	4Q26	1Q27	2Q27	3Q27	4Q27	1Q28	2Q28	3Q28	4Q28
Server CPUs	AMD (EPYC)	Zen 4 "Genoa-X" (T-5nm)			Zen 5 "Turin" (T-4nm)						Zen 6 "Venice" (T-2nm)				Zen 7 "Verano" (T-A16)						
	AMD (EPYC C)	Zen 4c "Bergamo" (T-5nm)			Zen 5c "Turin Dense" (T-3nm)						Zen 6c (T-2nm)										
	INTC (Xeon P)	"Emerald Rapids" (I-7)		"Granite Rapids" (I-3) (HVM - 1H25)						"Diamond Rapids" (I-18A)				"Coral Rapids" (I-18A-P)							
	INTC (Xeon E)		"Sierra Forest" (I-3)						"Clearwater Forest" (I-18A)												
	NVDA	"Grace" (T-4nm)						"Vera" (T-3nm)													
	ARM													"AGI" (T-3nm)							

Source: BofA Global Research estimates

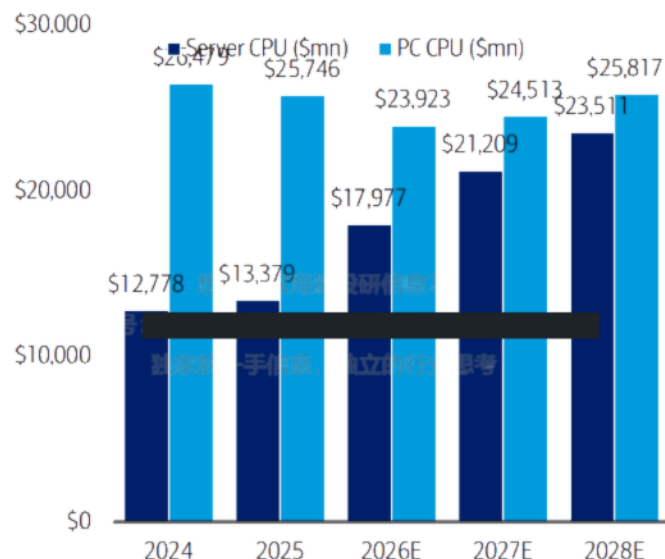
BofA GLOBAL RESEARCH

INTC Server Opportunity

Alongside the increased overall TAM, we see similarly increased INTC server CPU opportunity over time given its sticky enterprise customer base (35-40% of total workload today), as well as its upcoming partnership with NVDA to produce traditional DGX form-factor Rubin GPU platform with INTC server CPUs. Meanwhile, we also see INTC ASPs to benefit over time given the ongoing CPU supply tightness (albeit at the expense of PC capacity) and improved mix (higher core-count every new generation).

Exhibit 12: INTC Server CPU sales could reach \$22bn by CY28 from just \$13bn in CY25

Server CPU Sales Outlook (\$mn) and PC CPU Sales Outlook (\$mn)

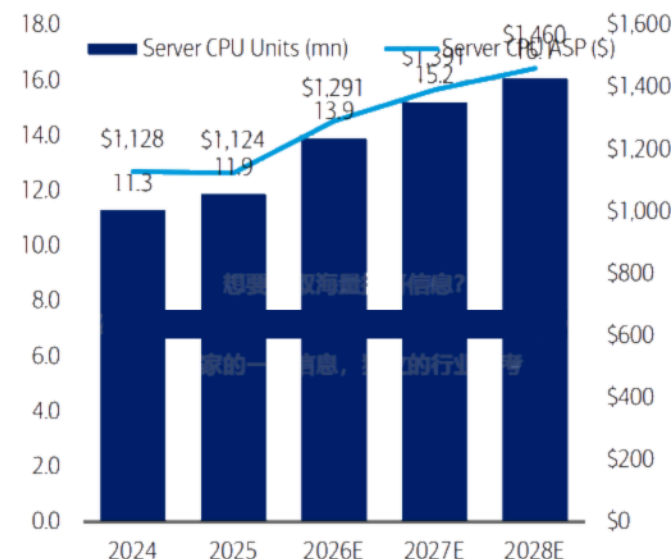


Source: BofA Global Research estimates, Mercury Research

BofA GLOBAL RESEARCH

Exhibit 13: We expect INTC Server units/ASPs are expected to rise over time

Server CPU Units Outlook (mn) and Server ASP Outlook (\$)



Source: BofA Global Research estimates, Mercury Research

BofA GLOBAL RESEARCH



End Market Poster

The next exhibit shows relative exposure of major semiconductor companies by end market, including for companies exposed to PC and data center markets.

Exhibit 14: Relative exposure of major semiconductor companies by end market, including for companies exposed to PC and data center markets

End Market Poster

Company	Ticker	Data Ctr. /Telco								Key Customers	Key Products
		Auto	Indus- trial	Cons- umer	PC	Smart- phone	Data Ctr. Compute	Network- ing	Mobile Infra		
Advanced Micro Devices	AMD	2%	4%	6%	25%	2%	59%	1%	2%	Meta, OpenAI, Oracle, HP, Sony	PC/Server CPU, AI GPU, FPGA
Allegro Microsystems	ALGM	71%	25%	0%	0%	0%	4%	0%	0%	-	Magnetic Sensors, Power Mgmt ICs
Ambarella	AMBA	22%	39%	39%	0%	0%	0%	0%	0%	Motorola, Bosch, Amazon	Video and Image Processors
ams OSRAM	AMS	38%	34%	4%	1%	20%	3%	0%	0%	Apple	LEDs, Optical
Analog Devices	ADI	27%	46%	7%	2%	5%	4%	3%	6%	Apple	High-End Mixed Signal Analog
Arm Holdings (Royalties)	ARM	6%	1%	16%	3%	51%	23%	0%	0%	Qualcomm, MediaTek, Microsoft	CPU, CPU IP Licensing
Broadcom (Semis)	AVGO	1%	1%	3%	0%	9%	56%	30%	0%	Google, Anthropic, Apple, Meta	AI ASIC, AI Networking, RF, Storage
Cirrus Logic	CRUS	1%	0%	8%	13%	77%	0%	0%	0%	Apple, Samsung, Xiaomi	Audio Codecs, Amplifiers, Power Mgmt
Coherent	COHR	0%	17%	0%	1%	7%	0%	75%	0%	HPE, Cisco, Dell, Apple	Specialty lasers, Optical Switches, Tx
ChangXin Memory (CXMT)	-	0%	1%	9%	2%	64%	24%	0%	0%	Chinese OEMs	DRAM
Diodes	DIOD	17%	28%	20%	5%	8%	12%	5%	5%	Apple	Analog, Discretes
GlobalFoundries	GFS	26%	20%	0%	3%	39%	5%	8%	0%	Qualcomm, AMD, MediaTek	Foundry Services
HiSilicon	-	4%	12%	7%	0%	41%	22%	8%	6%	Huawei, Chinese OEMs	CPU/SoC
Infineon Technologies	IFX	58%	20%	7%	1%	3%	11%	0%	0%	Siemens, BYD, Bosch, Conti, ZF	MCU, Discretes, Sensors
InnoLight	300308 CH	3%	0%	0%	0%	0%	90%	7%	0%	Google, Nvidia, AMZN, META, MSFT	Optical Transceivers
Intel	INTC	0%	2%	0%	55%	0%	32%	7%	3%	Dell, Lenovo, HP, Amazon	PC/Server CPU, Foundry Services
Kioxia	285A JT	0%	0%	4%	15%	33%	48%	0%	0%	Apple, Dell, Lenovo	NAND, SSD
Lattice Semiconductor	LSCC	2%	31%	6%	2%	0%	44%	12%	3%	Arrow, Weikeng, Macnica	Low-power, mid-range FPGA
Lumentum	LITE	0%	9%	0%	0%	4%	0%	87%	0%	Google, Ciena, Apple	Lasers, Optical Switches, Tx
M/A-COM Technology Solutions	MTSI	4%	39%	0%	0%	0%	0%	36%	21%	Defense, SpaceX (Starlink)	Power RF, Optical, Discretes
Marvell Technology	MRVL	1%	0%	1%	0%	0%	14%	65%	19%	Amazon, Microsoft, Google, Dell	AI Connectivity, AI ASIC, Storage, DPU
MaxLinear	MXL	0%	24%	21%	2%	9%	0%	32%	12%	CommScope	Wi-Fi, Optical, Transceivers
MediaTek	2454 TT	2%	1%	10%	5%	65%	5%	0%	12%	Xiaomi, Samsung, Google	Modem, CPU/SoC, AI ASIC
Melexis	MELE	88%	7%	4%	0%	0%	0%	0%	1%	Bosch, Sensata	Sensors, Drivers
Microchip Technology	MCHP	14%	48%	9%	9%	0%	10%	0%	10%	-	MCU, Analog, FPGA
Micron Technology	MU	12%	3%	3%	14%	18%	46%	3%	1%	Nvidia, AMD, Apple, Dell	DRAM/HBM, NAND
Monolithic Power Systems	MPWR	21%	7%	9%	6%	0%	46%	8%	3%	-	Power Mgmt ICs, Power Modules
Murata Manufacturing	6981 JT	26%	13%	8%	7%	37%	-	7%	2%	-	Passives, MLCCs, RF
Nanya Technology	2408 TT	3%	1%	50%	25%	13%	6%	1%	1%	OEMs, ADATA, Kingston	DRAM (Commodity, Low-Power)
Novatek	3034 TT	5%	4%	50%	0%	34%	0%	0%	8%	-	Display Driver ICs, Display SoC
Nvidia	NVDA	1%	0%	1%	4%	0%	79%	14%	0%	MSFT, Amazon, Meta, Dell, SMCI	GPU, CPU, DPU, AI Networking
NXP Semiconductors	NXPI	59%	10%	9%	0%	12%	0%	0%	10%	Continental, Apple	MCU, Power RF, Connectivity, Discretes
Onsemi	ON	56%	25%	4%	2%	2%	4%	2%	5%	-	Analog, Power Discretes, Image Sensors
Qorvo	QRVO	2%	14%	7%	0%	70%	1%	3%	3%	Apple, Samsung	RF, GaN Amplifiers, UWB modules
Qualcomm	QCOM	11%	1%	15%	4%	68%	0%	0%	1%	Apple, Samsung, Xiaomi, Meta	Mobile SoC, RF, Modem, PC/DC CPU
Renesas Electronics	6723 JT	48%	16%	3%	0%	15%	12%	4%	2%	Denso, Apple	MCU, Analog, Power Mgmt
ROHM	6963 JT	49%	12%	23%	6%	4%	5%	1%	1%	-	Power Discretes, Power Mgmt ICs
Samsung Electronics	005930 KS	2%	1%	3%	13%	27%	52%	1%	1%	Nvidia, AMD, Google, AWS, Apple	DRAM, NAND, CIS, DDI
Sandisk	SNDK	3%	3%	30%	32%	18%	15%	0%	0%	-	NAND, SSD
Semtech	SMTC	3%	41%	8%	5%	3%	8%	26%	6%	Trend-tek, Frontek, Samsung	Analog, Power Mgmt, Discretes
Silicon Motion	SIMO	2%	1%	1%	63%	23%	8%	1%	1%	OEMs, Nvidia, SK Hynix, Micron	Controller IC
SK Hynix	000660 KS	1%	1%	1%	13%	18%	64%	1%	1%	NVDA, AMD, Google, AWS, Apple	DRAM, NAND
Skyworks Solutions	SWKS	7%	4%	11%	4%	66%	1%	5%	2%	Apple, Samsung	RF, BAW Filters, Power Mgmt, Tuners
Semi Manufacturing Int'l (SMIC)	981 HK	8%	3%	43%	15%	23%	2%	6%	0%	Huawei, GigaDevice, China fabless	Foundry Services
Sony (Electronics Related)	-	5%	2%	13%	0%	80%	0%	0%	0%	Apple, Samsung	Image Sensors
STMicroelectronics	STM	36%	17%	8%	3%	18%	6%	0%	12%	Apple, Samsung, Tesla	MCU, Analog, Discretes, Sensors
Synaptics	SYNA	9%	10%	39%	21%	15%	0%	6%	0%	Samsung, OEMs	Display Driver ICs, Touch Sensors
Taiwan Semi Manufacturing	TSM	6%	6%	6%	7%	30%	28%	10%	7%	Apple, Nvidia, AMD, Qualcomm	Foundry Services
Texas Instruments	TXN	27%	37%	4%	5%	7%	11%	2%	6%	Apple	Analog ICs, Power Mgmt ICs, MCU
Vishay	VSH	35%	38%	8%	3%	1%	6%	8%	0%	-	Discretes, Passives, Optical
Yangtze Memory (YMTC)	-	0%	0%	1%	16%	55%	22%	3%	2%	Chinese OEMs	NAND, SSD

Source: BofA Global Research

BoFA GLOBAL RESEARCH



Glossary:

AVGO: Broadcom

NVDA: Nvidia

AMD: Advanced Micro Devices

ARM: Arm Holdings

INTC: Intel

MRVL: Marvell

GOOGL: Google

GCP: Google Cloud Platform

META: Meta Platforms

MSFT: Microsoft

AMZN: Amazon

AWS: Amazon Web Services

GCP: Google Cloud Platform

ORCL: Oracle

CSP: Cloud Service Provider

AGI: Artificial General Intelligence

TSMC: Taiwan Semiconductor Manufacturing Company

GPU: Graphics Processing Unit

CPU/MPU: Central/Micro Processing Unit

TPU: Tensor Processing Unit

ASIC: Application-Specific Integrated Circuit

MTIA: Meta Training and Inference Accelerator

ASP: Average Selling Price

DT: desktop

EDA: electronic design automation

GR: Granite Rapids

IDC: International Data Corporation

IDM: integrated device manufacturer

IP: intellectual property

NB: notebook

Nm: nanometer

NVDA: Nvidia

PC: personal computer

SKU: stock keeping unit

TOPS: trillions of operations per second

WS: workstation

TFLOP: Teraflop

PFLOP: Petaflop

EFLOP: Exaflop

HBM: High-Bandwidth Memory

NVL: NVLink

GB200/300: Grace Blackwell 200/300



SKU: Stock Keeping Unit
ESUN: Ethernet for Scale-Up Networking
IF: InfiniBand
SUE: Scale-Up Ethernet
UEC: Ultra Ethernet Consortium
AI: Artificial Intelligence
TAM: Total Addressable Market
GW: Gigawatt
MW: Megawatt
W: Watt
kW: Kilowatt
perf: Performance
TCO: Total Cost of Ownership
FP: Floating Point
NVFP: Nvidia Floating Point
CUDA: Compute Unified Device Architecture
JAX: JAX library (Google/Nvidia)
XLA: Accelerated Linear Algebra
OEM: Original Equipment Manufacturer
LLM: Large Language Model
DC: Data Center
CoWoS: Chip-on-Wafer-on-Substrate
Trn: Trainium
NIC: Network Interface Card
SAM: Serviceable Addressable Market
KV: Key Value
SRAM: Static Random Access Memory
CXL: Compute Express Link
RAM: Random Access Memory
DDR: Dual-Data Rate
LPDDR: Low-Power Dual-Data Rate
DIMM: Dual In-Line Memory Module
GDDR: Graphics Double Data Rate
NVMe: Non-Volatile Memory Express
DRAM: Dynamic Random Access Memory
P+E: Performance+Efficiency
HPC: High-Performance Computing
IC: Integrated Circuit
Tx: Transceiver
LED: Light-Emitting Diode
RF: Radio Frequency
SSD: Solid State Drive
SoC: System-on-Chip

MCU: Microcontroller Unit

MLCC: Multi-Layer Ceramic Capacitor

UWB: Ultrawide Band

DPU: Data Processing Unit

CIS: CMOS Image Sensor

CMOS: Complementary Metal Oxide Semiconductor

DDI: Device Driver Interface

BAW: Bulk Acoustic Wave



Exhibit 15: Stocks mentioned

Prices and ratings for stocks mentioned in this report

BofA Ticker	Bloomberg ticker	Company name	Price	Rating
AMD	AMD US	Advanced Micro	US\$ 278.26	C-1-9
ARM	ARM US	Arm Holdings	US\$ 162.33	C-2-9
INTC	INTC US	Intel	US\$ 68.5	C-3-9

Source: BofA Global Research

BofA GLOBAL RESEARCH

Price objective basis & risk**Advanced Micro Devices, Inc (AMD)**

Our \$310 PO is based on 29x our 2027E non-GAAP EPS. Our PO basis is towards the middle of AMD's historical 13x-58x range, justified by AI growth and CPU share gains offset by slower growth in cyclical embedded/console markets.

Downside risks: 1) Execution on first rack-scale product (MI400 Series), 2) Timing/Magnitude of Middle East AI Projects, 3) Lumpy nature of consumer and enterprise spending that could create delays in acceptance and success of new products, 4) High reliance on one outsourced manufacturing partner, 5) Maturity of current game console cycle.

Upside risks are greater share gain potential in the PC and server processor market against competitors.

Arm Holdings (ARM)

We assign a \$180 PO, which is based on 69x our CY27E non-GAAP EPS. Our PO basis is towards the middle of the 35x-92x historical trading range given increasing SoftBank dependence and opex ramp, though we flag longer-term data center content and silicon/chiplet opportunities. This basis is also in-line with 2x-3x PEG framework for IP/EDA peers.

Upside risks: 1) better than expected smartphone/PC/server unit shipment, 2) faster adoption of higher-royalty rate v9 and CSS IP at customers, 3) share gains in emerging ARM-based PC/server CPUs, 4) further proliferation of AI data centers and hyperscale-specific custom products, 5) Qualcomm/Nuvia content expansion post-settlement, 6) small floating volume.

Downside risks: 1) historically cyclical nature of semiconductor units, 2) high exposure to mature smartphone market, 3) competition against established x86 in the data center, 4) emerging competition from RISC-V in low-end consumer markets, 5) rising geopolitical tensions and deterioration of Arm China relationship, 6) ongoing Qualcomm/Nuvia litigation, 7) small trading float.

Intel (INTC)

Our \$48 price objective is based on 3.7x our 2028E EV/S. Our PO basis is within the 1.7x-4x historical range which we view as appropriate given increasing server CPU opportunity, offset by near-/medium-term manufacturing and external foundry uncertainties.

Upside risks to our price objective are 1) key external foundry packaging/wafer deals that could significantly boost sales/utilization, 2) greater than expected yields/ramps at 18A and upcoming 14A nodes, resulting in greater GM/utilization profile, 3) stronger than expected PC market from Windows 10 refresh or AI uplift, 4) geopolitical tensions boosting sentiment for domestic manufacturing asset.



Downside risks to our price objective are 1) lower than yield/ramp at Intel Foundry, particularly for its new 18A and upcoming 14A nodes, 2) lack of material external foundry customer in wafer processing, 3) weaker-than-expected trends in a mature PC market, which is largest revenue generator for Intel, 4) accelerated share loss to major CPU competitors.

Analyst Certification

I, Vivek Arya, hereby certify that the views expressed in this research report accurately reflect my personal views about the subject securities and issuers. I also certify that no part of my compensation was, is, or will be, directly or indirectly, related to the specific recommendations or view expressed in this research report.

Special Disclosures

BofA Securities is currently acting as joint financial advisor to ABB Ltd in connection with the sale of its Robotics Division to SoftBank Group Corp, which was announced on October 8, 2025.



US - Semiconductors and Semiconductor Capital Equipment Coverage Cluster

Investment rating	Company	BofA Ticker	Bloomberg symbol	Analyst
BUY	Advanced Energy Industries	AEIS	AEIS US	Duksan Jang
	Advanced Micro Devices, Inc	AMD	AMD US	Vivek Arya
	Allegro Microsystems	ALGM	ALGM US	Vivek Arya
	Analog Devices Inc.	ADI	ADI US	Vivek Arya
	Applied Materials, Inc.	AMAT	AMAT US	Vivek Arya
	Broadcom Inc	AVGO	AVGO US	Vivek Arya
	Cadence	CDNS	CDNS US	Vivek Arya
	Camtek	CAMT	CAMT US	Michael Mani
	Credo Technology	CRDO	CRDO US	Vivek Arya
	KLA Corporation	KLAC	KLAC US	Vivek Arya
	Lam Research Corp.	LRCX	LRCX US	Vivek Arya
	M/A-Com	MTSI	MTSI US	Vivek Arya
	Marvell Technology, Inc.	MRVL	MRVL US	Vivek Arya
	Microchip	MCHP	MCHP US	Vivek Arya
	Micron Technology, Inc	MU	MU US	Vivek Arya
	MKS Instruments	MKSI	MKSI US	Michael Mani
	Nova	NVMI	NVMI US	Michael Mani
	NVIDIA Corporation	NVDA	NVDA US	Vivek Arya
	onsemi	ON	ON US	Vivek Arya
	Synopsys	SNPS	SNPS US	Vivek Arya
	Teradyne	TER	TER US	Vivek Arya
NEUTRAL	Ambarella	AMBA	AMBA US	Vivek Arya
	Ambiq Micro, Inc.	AMBQ	AMBQ US	Vivek Arya
	Arm Holdings	ARM	ARM US	Vivek Arya
	Astera Labs Inc	ALAB	ALAB US	Vivek Arya
	Coherent Corp	COHR	COHR US	Vivek Arya
	Lumentum Holdings	LITE	LITE US	Vivek Arya
	NXP Semiconductors NV	NXPI	NXPI US	Vivek Arya
UNDERPERFORM	Texas Instruments Inc.	TXN	TXN US	Vivek Arya
	Axcelis Technologies	ACLS	ACLS US	Duksan Jang
	GlobalFoundries	GFS	GFS US	Vivek Arya
	Intel	INTC	INTC US	Vivek Arya
	Lattice Semiconductor	LSCC	LSCC US	Duksan Jang
	Qualcomm	QCOM	QCOM US	Vivek Arya
	Skyworks Solutions, Inc.	SWKS	SWKS US	Vivek Arya
RVW	Wolfspeed Inc	WOLF	WOLF US	Vivek Arya

Disclosures

Important Disclosures

