



THE INTERMARKET UNIVERSE (IU) DIGEST

Conference Season

Field Dispatch — June 2026

Analytical coverage across three simultaneous technology conferences — Computex 2026 · Evercore Global TMT · BofA Global Technology

JUNE 2-5, 2026

TECHNOLOGY SECTOR · INSTITUTIONAL ANALYSIS

● LIVE · JUNE 2-5

Computex 2026

Taipei · “AI Together” · 1,500 exhibitors

● INAUGURAL · JUNE 2-4

Evercore Global TMT

San Francisco · Invite-only institutional

● ANNUAL · JUNE 2-3

BofA Global Technology

Fireside series · Semiconductor focus

Three simultaneous conference forums have produced a rare moment of analytical convergence: the agentic AI inflection is not a forward projection — it is the current operating condition. The hardware, software, and capital implications are being disclosed in real time by the executives managing them.

KEY THEMES

Agentic AI · Token Economy · Vera Rubin Ramp · Memory Tightness · Connectivity Infrastructure · x86 Turnaround

01 —

Computex 2026 • Taipei

June 2–5 • Record Scale

Computex 2026 convened 1,500 exhibitors from 33 countries across 6,000 booths at Taipei Nangang Exhibition Center. The event ran concurrently with NVIDIA's GTC Taipei (June 2–4), converting the opening week into a de facto global AI infrastructure summit. Keynotes from NVIDIA, Qualcomm, Intel, and Marvell all pointed to the same conclusion: the agentic AI inflection has arrived, and the hardware decisions being made right now will define the next decade of computing.

NVDA

NVIDIA Corporation

Jensen Huang, CEO • GTC Taipei Keynote • June 1, 2026

Jensen Huang delivered the de facto opening statement of the conference season. **Vera Rubin is confirmed in full production** — the supply chain built to support it is reportedly twice the scale of Grace Blackwell, with rack assembly time compressed from two hours to five minutes, a logistics efficiency signal with direct implications for data center deployment velocity globally.

Beyond Vera Rubin, NVIDIA made its formal entry into the PC market with the **RTX Spark superchip** (formerly N1X), built on TSMC's 3nm process. The chip pairs NVIDIA's 20-core Grace ARM processor with a Blackwell GPU, pushing 1 petaflop of AI compute throughput with 128GB unified memory at 300GB/s. OEM launch partners — Microsoft, Dell, and HP — are expected to ship systems in fall 2026.

\$81.6B

Q1 FY27 Revenue (+85% YoY)

\$75.2B

Data Center Revenue (+92% YoY)

75.0%

Non-GAAP Gross Margin

\$5.4T

Market Cap at Conference

SUPPLY CHAIN READ-THROUGH

RTX Spark, built on TSMC's 3nm process, creates a structurally durable demand vector for CoWoS advanced packaging and N3 capacity. The launch framing acknowledged the ongoing DRAM and NAND shortage — directly in line with what IU has been watching in the memory market across WDC and SNDK coverage.

QCOM **Qualcomm Technologies**

Cristiano Amon, CEO · Opening Keynote: “Year of the Agents” · June 1

Qualcomm CEO Cristiano Amon set the tone early, declaring 2026 “the year of agents” and framing the dominant theme across all three conferences: the **token economy**. Global AI token consumption currently runs at 31.7 billion tokens every ten seconds — a figure Amon projects will reach 1.27 trillion per ten-second window by 2030, a roughly 40-fold increase driven by agentic AI executing continuous multi-step background operations.

The engineering implication is stark: a simple conversational AI response requires roughly 10,000 tokens, while a multi-step agentic workflow demands up to one million — a 100x differential with direct consequences for memory bandwidth at scale. Amon closed by unveiling **Dragonfly**, Qualcomm's new data-center AI inference chip brand, with an Investor Day set for June 24 in New York.

PORTFOLIO IMPLICATION

Qualcomm's token math framework is the most analytically rigorous public framing of agentic compute demand this conference season. The 100x token-per-task differential has direct implications for HBM and DRAM demand modeling in IU's Micron and TSMC coverage.

INTC **Intel Corporation**

Lip-Bu Tan, CEO · Computex Keynote · June 2

Intel announced **Xeon 6+** — 288 e-cores, a 576MB L3 cache, built on Intel 18A — as its enterprise-grade inference answer to the agentic AI workload cycle. Management cited forecasts of AI inference accounting for approximately 40% of all data center power demand by 2030. Notably, **Intel 18A is now confirmed in volume production**.

Rackscale AI infrastructure integrates Xeon processors with SambaNova SN-50 RDUs, alongside strategic collaborations with Foxconn, Siemens, and Hitachi. A joint appearance with Perplexity CEO Aravind Srinivas highlighted agentic device applications as the next frontier for client computing.

MRVL Marvell Technology

Matt Murphy, CEO · Joint Appearance with Jensen Huang · June 2

The headline event of the week came on stage at Computex, where Jensen Huang declared Marvell **“the next trillion-dollar company”** in a joint appearance with CEO Matt Murphy. With Marvell's market cap at roughly \$192 billion at the time — implying more than 400% upside to reach the milestone — the endorsement nonetheless catalyzed a roughly 25% surge in premarket trading.

Huang's strategic rationale: as AI compute disaggregates across thousands of connected chips within massive clusters, Marvell's end-to-end connectivity infrastructure becomes structurally indispensable. The **\$2 billion NVIDIA investment in Marvell**, announced in March, formalized the silicon photonics and optical connectivity partnership. Marvell's custom ASIC business is forecast to surpass **\$10 billion in fiscal 2029**.

MARKET-MOVING EVENT

Huang's trillion-dollar framing reflects the structural reality that AI's disaggregated compute model creates durable demand for connectivity infrastructure at a scale the market had not fully priced. MRVL's one-year return of approximately 265% through conference date underscores the pace of market recalibration.

02 — **Evercore Global TMT Conference**Inaugural · San
Francisco

Evercore hosted its inaugural Global TMT Conference June 2–4 in San Francisco — an invite-only institutional forum. Senior Managing Director Amit Daryanani framed the analytical context: "AI is reshaping the underlying hardware and connectivity stack, creating a host of bottlenecks across the tech stack and resulting in significant investments flowing into compute, networking and data center infrastructure." Western Digital and Palo Alto Networks both presented across its two-day program.

WDC **Western Digital Corporation**

Kris Sennesael, CFO · Evercore TMT (June 3, 8:45 AM PT) + BofA (June 2)

WDC's 2026 production capacity is fully committed, with customer discussions underway on long-term agreements extending to 2028–2029, and early inquiries reaching out to 2030. As AI inference workloads shift from training to inference, data generation velocity increases sharply, and the economics of storing that data at scale favor hard drives for the bulk of hyperscale cloud storage.

On the NVIDIA KV Cache offload initiative, Sennesael sees it as a tailwind: more inference activity generates more data, and that data ultimately resides within the hyperscale storage tier. WDC's Innovation Day roadmap calls for 20%-plus annual revenue growth over three to five years, gross margins above 50%, and a HAMR transition expected to be margin-neutral to slightly accretive at scale.

COVERAGE NOTE

This lines up with what IU laid out in the May 2026 WDC deep-dive. As AI inference workloads keep growing, so does the volume of data that needs to be stored — and that storage largely lives on hard drives at the hyperscale level. Multi-year supply agreements give WDC a degree of revenue predictability the hard drive industry hasn't seen in a long time.

PANW Palo Alto Networks

Nikesh Arora, Chairman & CEO · Evercore Global TMT · June 2, 2026

Palo Alto Networks CEO Nikesh Arora arrived at the Evercore TMT Conference with fresh ammunition: Q3 fiscal 2026 results released that same day showed **revenue of \$3.0 billion, up 31% year-over-year**, with adjusted EPS of 85 cents beating expectations. The number that captured the room: Arora disclosed **1,200 inbound customer meeting requests** in the preceding weeks — organizations scrambling to assess their exposure as AI-powered attacks grow faster and harder to detect.

Arora's broader argument: AI doesn't threaten the cybersecurity business — it dramatically expands it. Every new AI agent deployed inside an enterprise is a new attack surface. Next Generation Security revenue grew 60% year-over-year in Q3, and management raised its full-year outlook.

THEMATIC READ-THROUGH

PANW's conference appearance adds an important dimension missing from the hardware-heavy conference narrative: every dollar of AI infrastructure being discussed at Computex and BofA creates a corresponding security obligation. Arora's 1,200 meeting request figure is the most concrete demand signal yet that the AI security cycle is running parallel to the same agentic AI inflection described from the hardware side.

AVGO

Broadcom Inc.

Hock Tan, President & CEO · BofA Global Technology Conference · June 3, 2026

Broadcom reported Q2 fiscal 2026 earnings the same day it presented at BofA — and the numbers were difficult to ignore. **AI semiconductor revenue reached \$10.8 billion, up 143% year-over-year.** Q3 AI chip revenue is guided to **\$16.0 billion — more than 200% year-over-year growth.** Total Q3 revenue guidance of \$29.4 billion, up 84% year-over-year, implies a quarterly run rate that would have been Broadcom's full-year revenue just a few years ago — a reflection of both organic AI acceleration and the structural step-up from VMware consolidation, which closed in late 2023.

The business behind those numbers is Broadcom's custom ASIC franchise. Google's TPU program is the primary customer relationship, with supply agreements running through 2031. Meta's MTIA accelerator and Anthropic's custom chip program are both ramping, with OpenAI deployments expected in fiscal 2027. Bookings for AI semiconductors now exceed three times quarterly shipments. Tan told investors the company has "line of sight" to AI chip revenue exceeding **\$100 billion annually in fiscal 2027.**

\$10.8B	\$16.0B	\$22.2B	\$100B+
Q2 AI Revenue (+143% YoY)	Q3 AI Revenue Guide (+200%+ YoY)	Q2 Total Revenue (+48% YoY)	AI Revenue Target, FY2027

COVERAGE READ-THROUGH

Broadcom's Q3 AI revenue guide — \$16 billion at 200%+ YoY growth — is the single most aggressive forward disclosure of the conference week. The custom ASIC ecosystem Tan is describing, with Google, Meta, Anthropic, and OpenAI all running separate multi-year programs simultaneously, is the demand story underneath the memory tightness, packaging constraints, and connectivity bottlenecks discussed across every other session.

AMD

Advanced Micro Devices

Jean Hu, CFO; Matt Ramsay, VP IR · June 2, 2:20 PM EDT

The **MI450 ramp begins in Q3**, scales to what management described as “billions of dollars” in Q4, and steps up further in Q1. The ZT Systems acquisition positions AMD to offer system-level design integration for AI infrastructure customers.

AMD expects additional multi-gigawatt customers beyond OpenAI and Meta, and both leading customers are now forecasting above AMD's original 2027 plan. On memory, CFO Jean Hu acknowledged unprecedented DRAM cost inflation, describing a strategy of long-term supplier agreements and disciplined cost-sharing with customers.

MI450

GPU Ramp: Q3 Launch, Q4 Scale

>\$800B

Market Cap at Conference

2027+

Demand Visibility Horizon

+136%

YTD Return at Conference

SUPPLY CHAIN READ-THROUGH

AMD's DRAM cost inflation acknowledgment — sitting alongside everything surfaced at Computex about agentic AI workload growth — adds another data point to the memory tightness picture. Leading customers revising demand above plan is the strongest institutional validation of multi-year AI infrastructure spending durability published this cycle.

DELL

Dell Technologies

Arthur Lewis, President — Infrastructure Solutions Group · June 2, 1:40 PM EDT

Dell's ISG President Arthur Lewis joined a session moderated by analyst Wamsi Mohan, who opened with a pointed observation: Dell's Infrastructure Solutions Group has been “nothing short of amazing lately.” That tracks — Dell sits at the center of enterprise AI deployment as both a system integrator and a confirmed OEM launch partner for NVIDIA's RTX Spark platform.

AMAT

Applied Materials, Inc.

Brice Hill, SVP & CFO · BofA Global Technology Conference · June 2, 2026 · 11:40 AM EDT

Applied Materials CFO Brice Hill opened with an unambiguous characterization: **"The growth in AI that Applied has been investing for is now in full force."** Applied now expects its semiconductor systems business to grow **more than 30% in calendar 2026** — upgraded from prior guidance of more than 20% — driven by new orders coming in above expectations across leading-edge foundry-logic, DRAM, and advanced packaging.

Hill put hyperscaler AI CapEx in context: approximately \$600 billion in 2026, rising to a projected \$700 billion in 2027. Against that backdrop, **advanced packaging revenue is on track to grow more than 50% in calendar 2026** — driven by surging demand for the chip-stacking and interconnect processes central to how NVIDIA, AMD, and custom AI chip designers build their most powerful products.

> 30%	> 50%	\$8.95B	50.0%
Semicon. Systems Growth, Cal. 2026	Adv. Packaging Revenue Growth 2026	Q3 FY26 Revenue Guidance	Q2 Non-GAAP Gross Margin (record)

WFE COVERAGE READ-THROUGH

Applied's guidance upgrade from 20%-plus to 30%-plus semiconductor systems growth in a single quarter — driven by incremental orders across logic, DRAM, and packaging — is the most direct confirmation that WFE spending is tracking ahead of prior consensus.

MU

Micron Technology

Sanjay Mehrotra, Chairman, President & CEO · BofA Global Technology Conference · June 2, 2026

Calendar 2026 DRAM and NAND supply is fully committed. HBM4 volume production is underway, aligned to NVIDIA's Vera Rubin platform, and Micron's fiscal Q2 2026 results showed record revenue across every business unit. Mehrotra's message was consistent: this isn't a cyclical blip. Cleanroom construction lead times are lengthening across geographies, and the HBM trade ratio — the DRAM wafer capacity consumed per HBM unit — rises with each successive generation.

Looking further out: the HBM total addressable market is projected to expand from roughly **\$35 billion in 2025 to approximately \$100 billion by 2028** — larger than the entire DRAM market in 2024. HBM4E is in development with volume production targeted for calendar 2027, and Micron has begun sampling a 48GB per-cube version of HBM4. Capex for fiscal 2026 is expected to exceed **\$25 billion**, with fiscal 2027 stepping up further.

\$100B	>\$25B	HBM4	2026+
HBM TAM Projection by 2028	FY2026 Capex Guidance	In Volume Production, Vera Rubin-aligned	Supply Tightness Persists

COVERAGE NOTE

Micron's conference commentary closes the loop on a read-through running across every section of this brief: from Qualcomm's token economy math, to AMD's DRAM cost inflation disclosure, to NVIDIA's RTX Spark supply constraint acknowledgment. The HBM TAM trajectory — from \$35B to \$100B in three years — is the single most consequential demand figure disclosed across all three conference forums this week.

04 — **Cross-Conference Synthesis**Analytical
Perspective**Five Structural Signals from Conference Week****01 • Agentic AI Is Not Forward Guidance — It Is Happening Now**

Three simultaneous conference forums produced a single structural consensus. Qualcomm quantified the transition (40x token demand by 2030). NVIDIA announced production infrastructure for it. AMD disclosed demand above plan from its leading hyperscaler customers. This is management disclosure from the operators of the relevant infrastructure — not a sell-side projection.

02 • Memory Tightness Is Structurally Confirmed Across Three Independent Disclosures

AMD's unprecedented DRAM cost inflation acknowledgment, WDC's fully-committed 2026 capacity with LTAs extending to 2030, and NVIDIA's RTX Spark launch constrained by the ongoing shortage together tell a consistent story: memory is tight, demand is real, and the shortage isn't going away soon.

03 • Connectivity Is the New Compute Bottleneck

Huang's Marvell endorsement was the most explicit public confirmation that as AI compute disaggregates across massive clusters, the connectivity layer — not the GPU alone — becomes the primary value-creation site. Qualcomm's Dragonfly entry into data-center inference reflects the same recognition.

04 • Intel's Turnaround Has Credible Operational Evidence for the First Time

Zinsner's BofA disclosures and Tan's Computex announcements together give the Intel recovery story its first real operational legs: a next-generation manufacturing process now in volume production, a management structure cut roughly in half, and a cultural shift toward accountability showing up in the numbers.

05 • Coverage Universe Alignment

Conference signals are broadly consistent with IU's published coverage across AVGO, Micron, TSMC, WDC, SNDK, and AMAT — with Broadcom's Q3 AI revenue guide (\$16B, +200% YoY) emerging as the single most aggressive forward disclosure of the week. Palo Alto Networks adds the security dimension: the obligation created by every new piece of AI infrastructure.

IU Editorial Position

The week of June 2–5, 2026 will be remembered as the moment the technology industry stopped debating whether agentic AI would arrive and started disclosing what it actually requires. Capital commitments, manufacturing constraints, and architectural decisions disclosed across these three forums all point in the same direction IU has been tracking. The story hasn't changed — it's just getting louder.

Access Full Deep-Dive Research

IU Digest subscribers receive proprietary DCF models, WACC sensitivity matrices, sector deep-dive reports, earnings analysis, and options strategy frameworks across all covered names — calibrated to depth of access.

■ Subscriptions Available Soon

Free Tier · Analyst Tier · Institutional Tier — Pricing & tier structure coming soon

This conference brief is produced by Intermarket Universe for informational and research purposes only. All data and company commentary are sourced from publicly available conference transcripts, press releases, and management disclosures. This is not investment advice. Forward-looking statements attributed to management are their own, subject to risks and uncertainties.

For analytical purposes only — not for redistribution without permission. © 2026 Intermarket Intelligence LLC.