

US Semiconductors and Semi Equipment

SemiBytes: Ultra-Fast Inference, QCOM ADay, WFE Upside, AMD-MEXT

Equities

Americas

Semiconductors

Timothy Arcuri

Analyst

timothy.arcuri@ubs.com

+1-415-352 5676

Natalia Winkler, CFA

Analyst

natalia.winkler@ubs.com

+1-415-352 4626

Alex Kivali

Analyst

alex.kivali@ubs.com

+1-212-713 3945

Gianmarco Vella

Associate Analyst

gianmarco.vella@ubs.com

+1-415-352 4555

Aaryan Wadhwa

Associate Analyst

aaryan.wadhwa@ubs.com

+1-212-821 6481

UBS Chip Chats - A Discussion on Ultra-Fast Inference with a Hardware Expert

We will be hosting an expert call on fast inference and disaggregated architectures on Tuesday, June 23 from 11:00am to 12:00pm ET. Please register [here](#).

QCOM Analyst Day Preview - \$20+ EPS by FY31

We expect QCOM Analyst Day this Wednesday to lean more heavily into a long-term data center narrative—which was largely absent back in 2024 and not part of QCOM's standing financial model. Qualcomm has already begun to showcase this buildout, including a rack-scale AI inference system around AI200 with a 43TB GDDR memory rack (a 350B-parameter model fitting on a single card) optimized for decode, with initial deployments this year. Beyond announced accelerators for decode applications, we expect incremental updates on its server CPU efforts and potential inorganic options. On the latter, the Tenstorrent angle (per [press reports](#)) makes sense to us as a stack extension, moving QCOM from a largely decode-optimized system provider toward a more balanced inference architecture (CPU + accelerator + system), adding exposure to higher-value prefill / latency-sensitive workloads with Tenstorrent hardware (also based on GDDR memory, as is QCOM's AI line). This aligns with our work (see [here](#)) pointing to disaggregated inference (prefill vs. decode on different hardware stacks) as both very attractive for performance / economics, but still presenting hardware and software challenges in the near term, according to engineers in the field. From a messaging standpoint, we think QCOM offers a new model that incorporates the data center and CPU/agent opportunity - these markets could ultimately add ~\$20B to QCOM's existing financial model thus maybe extending the model to revs/EPS of \$65B/\$25 in a 5yr horizon. Outside of data center, we continue to see solid execution in Automotive, where the business remains on track toward QCOM's \$8B FY29 target. With increasing industry focus on "physical AI," we would not be surprised to see QCOM take a more constructive stance on incremental opportunities across Auto and IoT, potentially implying a combined non-handset revenue opportunity of ~\$25B by FY31 (vs. the current ~\$22B FY29 target). This might translate into QCOM's thinking of ~\$70B of revenues for FY31 (with handset dropping to ~50% of revenues) and potential EPS of \$20+.

SPE look cheap on a \$300B WFE upside case in C2029

We continue to field numerous questions around our bullish stance on the SPE vertical into C2027/2028 (and beyond). The recent rally in SPE names (1mos % change: LRCX +29%; AMAT +44%; KLAC +37%) reflects, in part, a growing investor appreciation for the strength of WFE growth over the next 2-3 years ([UBS WFE of \\$198B/\\$247.5B in C2027/2028](#), and C2029 up strongly again, we think). However, we are increasingly asked by investors how much higher stocks can go and consequently estimate what EPS could be in a scenario where WFE spend reaches \$300B in C2029. To this end, we calculate EPS of ~\$17 for LRCX, ~\$36 for AMAT and ~\$10 for KLAC ([Figure 1 - Figure 3](#)) - effectively, ~70-100% above where C2027 consensus stands today. From a valuation standpoint, our upside scenario implies SPE names to be effectively trading at ~22x (LRCX), ~17x (AMAT), and ~26x (KLAC) - levels that screen cheap vs. current multiples. Looked at another way, if we use current Street estimates for the Top 5 SPE to back into an implied WFE, we show in [Figure 4](#) that current Street estimates for each company imply a consensus WFE for C2029 of ~\$232B. We think this will prove far too low. Net, we continue to see a path for SPE names to be one of the most compelling opportunities across our coverage over the next few years.

Sifting through AMD's MEXT deal - Boosting system NAND performance (but NOT a DRAM replacement)

Earlier this week, we received a host of investor questions around the significance of AMD's acquisition of MEXT, particularly as a read-through to DRAM consumption. Importantly, as was the case w/ TurboQuant, we push back on the misconception among some investors that MEXT could reduce overall DRAM consumption (in favor of NAND) or, even more improbable, replace DRAM. Based on the company's website, MEXT describes its 'Predictive Memory' as a (software-only) offering that presents NAND flash to the operating system as if it were DRAM. In doing so, the company claims 2-4x greater effective memory capacity and, subsequently, up to 50% cost reductions - noting that in [its disclosed MoonRay \(application rendering application\) white paper](#), MEXT used a system with 256GB DRAM plus 256GB MEXT memory backed by NVMe, compared with a 512GB DRAM baseline, and reported comparable system performance (with 97.9% prediction accuracy - more below on this). From a technical standpoint, we understand MEXT's Predictive Memory technology as consisting of predictive prefetching capabilities that allow a system to constantly decide which data is either "hot" or "cold". The software demotes data that is used less often to NAND and leverages its Predictive Memory Engine (which leverages AI and is really the core of MEXT's IP, as we understand it) to forecast which data in NAND the application is likely to request next and proactively promotes it back into DRAM - with the goal being that the application will effectively experience the working set as if it were still resident in DRAM. However, the critical variable here is prediction accuracy - as MEXT's efficacy almost entirely depends on the Predictive Memory Engine's ability to correctly identify which data needs to be pulled back into DRAM before the application requests it. To this end, a wrong prediction risks enforcing the system to retrieve data from flash at far higher latency. Net, we think of MEXT's value proposition to be ultimately tied to the robustness of its prediction engine, making it best suited for workloads with predictable access patterns (such as LLMs and analytics) but, in contrast, less suited for highly random or latency-sensitive workloads (e.g., high-frequency trading). Strategically, we view the deal as primarily a memory-efficiency and TCO-related for AMD, as the company could package MEXT with EPYC/Instinct systems to improve effective memory capacity, reduce system cost, and improve AI inference economics - but noting that we do not view this as a threat to DRAM consumption. In fact, similar to what we have observed with TurboQuant, memory optimizations often expand the economically viable use-cases for AI and, in turn, spark incremental overall memory usage rather than reducing it.

SPE SENSITIVITY ANALYSIS

Figure 1: LRCX - EPS upside in a C2029 \$300B WFE scenario

		C2025	C2029 Upside Model	
		\$116,000	\$300,000	
WFE Share		11.5%	14.5%	
	LT Model	Reported	Model	CAGR
Revenue	\$30B+	\$20,561	\$57,903	30%
Systems		\$13,378	\$43,500	34%
CSGB		\$7,183	\$14,403	19%
Gross Profit		\$10,265	\$31,847	
Gross Margin	50%+	49.9%	55.0%	
Opex		\$3,245	\$7,817	
Opex % Revs		15.8%	13.5%	
Operating Income		\$7,020	\$24,030	
Op Margin	35%+	34.1%	41.5%	
Net Income		\$6,236	\$20,725	
Diluted Shares		1,274	1,195	
EPS	\$8.00+	\$4.89	\$17.34	

Source: Company reports, UBS estimates

Figure 2: AMAT - EPS upside in a C2029 \$300B WFE scenario

		C2025	C2029 Upside Model	
		\$116,000	\$300,000	
WFE Share		18.1%	21.8%	
	LT Model	Reported	Model	CAGR
Revenue	-	\$26,933	\$76,396	30%
Semi Systems		\$20,985	\$65,250	33%
AGS		\$5,948	\$11,146	17%
Gross Profit		\$13,149	\$42,018	
Gross Margin	-	48.8%	55.0%	
Opex		\$5,295	\$11,077	
Opex % Revs		19.7%	14.5%	
Operating Income		\$7,854	\$30,940	
Op Margin	-	29.2%	40.5%	
Net Income		\$7,560	\$27,929	
Diluted Shares		803	785	
EPS	-	\$9.42	\$35.58	

Source: Company reports, UBS estimates

Figure 3: KLAC - EPS upside in a C2029 \$300B WFE scenario

WFE Share	LT Model	C2025	C2029 Upside Model	
		\$116,000	\$300,000	
		8.5%	8.5%	
		Reported	Model	CAGR
Revenue	\$14B +/- \$500MM	\$12,745	\$30,577	24%
Systems		\$9,842	\$25,500	27%
Services		\$2,903	\$5,077	15%
Gross Profit		\$8,006	\$19,875	
Gross Margin	~63%	62.8%	65.0%	
Opex		\$2,211	\$5,045	
Opex % Revs	13% R&D ; 8% SG&A	17.3%	16.5%	
Operating Income		\$5,510	\$14,830	
Op Margin	41-43%	43.2%	48.5%	
Net Income		\$4,700	\$12,652	
Diluted Shares		1,326	1,265	
EPS	\$38.00 +/- \$1.50	\$3.54	\$10.00	

Source: Company reports, UBS estimates

Figure 4: Street-implied WFE estimates

	2025	2026E	Y/Y	2027E	Y/Y	2028E	Y/Y	2029E	Y/Y
ASML	\$28,227	\$33,869	20%	\$42,301	25%	\$47,227	12%	\$51,120	8%
AMAT	\$20,887	\$28,212	35%	\$36,683	30%	\$43,304	18%	\$45,901	6%
LRCX	\$13,378	\$18,366	37%	\$23,748	29%	\$28,283	19%	\$29,797	5%
KLAC	\$9,842	\$11,979	22%	\$15,011	25%	\$17,126	14%	\$18,089	6%
TEL	\$11,639	\$14,670	26%	\$17,710	21%	\$20,619	16%	\$20,112	-2%
Top 5 SPE Revenue	\$83,973	\$107,095	28%	\$135,453	26%	\$156,559	16%	\$165,020	5%
Top 5 WFE Share (10YR avg)	72.2%	72.2%		71.7%		71.7%		71.2%	
Implied Consensus WFE	\$116,323	\$148,354	28%	\$188,946	27%	\$218,387	16%	\$231,805	6%
UBSe WFE	\$116,000	\$147,000	27%	\$198,000	35%	\$247,500	25%		

Source: UBS estimates, VisibleAlpha

Valuation Method and Risk Statement

We use various valuation techniques such as P/E, EV/FCF for valuing the companies in this report. Risk factors include but are not limited to macroeconomic factors such as a downturn in the economy, a disruption of international trade, technological disruption due to new inventions, or business model innovation whereby structural changes in the industry alter the future course of unit sales, ASPs, and revenues.

