

AMD 2Q26 Earnings Call Transcript

JOHN: Greetings, and welcome to the AMD second quarter 2026 conference call. At this time, all participants are in a listen-only mode. A question and answer session will follow the formal presentation.

If anyone should require operator assistance during the conference, please press star 0 on your telephone keypad. And please note that this conference is being recorded. I will now turn the conference over to Matt Ramsay, VP, Financial Strategy, and IR. Thank you, Matt. You may begin.

MATT RAMSAY: Thank you. And welcome to AMD's second quarter 2020 financial results conference call. By now, you should have had the opportunity to review a copy of our earnings press release and the accompanying slides. If you have not had the chance to review these materials, they can be found on the Investor Relations page of AMD.com.

Today, we will refer primarily to non-GAAP financial measures during the call. The full non-GAAP to GAAP reconciliations are available in today's press release and slides posted on our website. As a reminder, our second quarter 2025 results included approximately \$800 million of inventory and related charges associated with US export control restrictions on MI308 shipments to China.

Unless otherwise noted, comments making year-over-year comparisons exclude the impact of those charges to provide a more comparable and meaningful view of our underlying business performance. Participants on today's conference call are Dr. Lisa Su, our chair and CEO, and Jean Hu, our executive vice president, CFO, and treasurer. This is a live call, and will be replayed via webcast on our website.

Before we begin, I would like to note that AMD will participate in the following events for the financial community. Key bank's Technology Leadership Forum on Tuesday, August 11, Citi's 2026 Global TMT Conference on Tuesday, September 8th, and the Goldman Sachs Communacopia and Technology Conference on Friday, September 11th.

Today's discussions contain forward-looking statements based on the current beliefs, assumptions, and expectations, including forward-looking statements regarding financial projections, business and industry trends that speak only as of today.

And as such, involve risks and uncertainties that could cause actual results to differ materially from our current expectations. Please refer to the cautionary statement in our press release for more information on these factors that could cause actual results to differ materially. With that, I will hand the call over to Lisa.

LISA SU: Thank you, Matt, and good afternoon to all those listening today. We delivered another outstanding quarter with record revenue and profitability as adoption of our leadership products continued to expand.

Revenue increased 50% year-over-year to 11.5 billion, driven by significantly higher sales of EPYC, Instinct, Ryzen, and embedded processors. Data center revenue more than doubled year-over-year and now represents 58% of total revenue, up from 42% a year ago, reflecting the rapidly expanding scale of our server and data center AI businesses.

Our record results mark another clear step up in AMD's financial performance, and demonstrate the strength of our product portfolio and execution. We are still in the early stages of a multiyear AI adoption cycle, as deployments grow across a broad set of markets and workloads, driving demand for more compute and creating a clear path to significant revenue growth and earnings power in the years ahead.

Turning to our segments, data center revenue grew 107% year-over-year to a record 6.7 billion, driven by strong demand for EPYC processors and Instinct accelerators. In server, we delivered our fifth consecutive quarter of record server CPU revenue with and enterprise sales each growing more than 70% year-over-year, exceeding the outlook we provided last quarter.

We gained x86 server revenue share year-over-year as customers expanded deployments of both 5th Gen EPYC Turin and 4th Gen EPYC Genoa families. In , hyperscalers continued expanding EPYC across their internal infrastructure and public offerings, including AWS, Microsoft, Google, Oracle, and others.

Fifth Gen EPYC Turin now powers nearly 1/3 of the more than 1,600 EPYC public instance types available globally, as providers broaden their offerings with new database storage and AI workloads. That expanding footprint is translating into growing adoption of EPYC in the with health care, financial services, media, and technology companies adding tens of millions of instances in the last quarter.

In enterprise, we delivered record sales in our fourth consecutive quarter of record sell through as on-prem adoption accelerated, driven by the leadership performance and TCO advantages of our EPYC portfolio.

Growth was broad-based as we won large deployments with leading financial services, manufacturing, telecom, retail, and technology companies. More than 230 5th Gen EPYC platforms are now in market from HPE, Dell, Lenovo, Supermicro, and others, our broadest enterprise portfolio to date.

Looking ahead, agentic AI is creating a new growth vector for server CPUs, spanning high frequency AI host nodes, high density agentic servers, and general purpose and enterprise workloads. Our 6th Gen EPYC Venice family is purpose built for this expanding range of workloads, and delivers one of the largest generational performance gains in EPYC history.

Built on our all new Zen 6th core and 2 nanometer technology, Venice extends EPYC leadership in performance and efficiency, delivering more than twice the performance per watt of leading x86 CPUs and up to 3.3 times the performance per watt of leading ARM-based CPUs.

The Venice family includes more than 30 processors that combine leadership per core and per socket performance with a broad range of memory and I/O configurations, giving customers greater flexibility to optimize performance, efficiency, and TCO across the most widely used , enterprise, and HPC workloads.

Venice is in production now, with every major OEM on track to launch platforms and the leading providers planning deployments beginning later this year. Customer demand for Venice is stronger than for any prior EPYC generation, and we expect to continue growing market share across and enterprise in the coming quarters.

Turning to our data center AI business, revenue more than doubled year-over-year, driven by strong demand for Instinct accelerators. MI355X adoption continued to broaden as leading AI companies scale deployments across a growing range of inferencing and training workloads, and providers expanded MI350 series availability.

At our advancing AI event, we launched Helios, our rack scale AI platform combining EPYC Venice CPUs, MI450 series GPUs, Pensando networking, and ROCm software. Across a broad range of inferencing workloads, Helios delivers up to 15% more throughput at the same rack power and up to 30% more tokens per dollar than the competition.

Customer pull for Helios is very strong and tracking ahead of our initial forecasts. In addition to our multi-generation gigawatt scale deployments with OpenAI and Meta, we announced a new strategic partnership with Anthropic.

Anthropic will deploy up to 2 gigawatts of MI450 series GPUs in Helios, with deployment of the first gigawatt beginning in the first half of 2027. The partnership includes a multi-year joint engineering collaboration using Claude to optimize workloads for Instinct GPUs and accelerate ROCm software development.

We also expanded our long-standing partnership with Microsoft. Microsoft will deploy Helios at scale on Azure for frontier model inferencing across Microsoft, its AI customers, and Azure AI services. Together, these commitments broaden the group of leading AI companies and cloud providers, building their next generation infrastructure on AMD.

Helios is now in production, with initial shipments on track to begin later this quarter and ramp through the fourth quarter and into 2027 to meet very strong customer demand. Looking beyond Helios, we plan to launch a new rack scale AI platform every year, with each generation delivering significant performance, efficiency, and TCO gains.

In 2027, our next generation platform combines MI500 series GPUs, Verano CPUs, and Pensando networking with expanded scale up domains and both copper and optical-based interconnects. Customer engagement on MI500 is very strong, with multiple customers working closely with us as they plan their next generation AI infrastructure.

We expect MI500 to deliver the largest generational leap in Instinct history, putting us on track to increase inferencing performance more than 2,000 times in just four years. Turning to our AI software stack, ROCm has reached an important inflection point.

With the performance capabilities and developer experience, customers need to deploy AI in production at scale. The breadth of the ecosystem also continues to expand. More than three million models now run out-of-the-box on AMD, the leading open models launch with day 0 support Instinct, and open-source contributions to ROCm have increased more than tenfold over the past year.

We introduced ROCm.ai, our new AI assisted development platform for AMD GPUs last month. ROCm.ai lets developers use today's leading coding agents, including Claude, Codex, and Cursor, to create port and optimize code for Instinct, making it significantly faster and easier to bring new models and workloads to AMD.

ROCm.ai delivers more than twice the training performance and more than three times the inferencing performance of ROCm 7 across a broad range of models. We are also working closely with the leading AI labs, including OpenAI, Anthropic, Meta, and others to co-optimize ROCm for their models, with the improvements benefiting the entire AMD ecosystem.

Taking a step back, the overall data center market opportunity is expanding far more rapidly than we projected just six months ago. As AI moves into production across a broader range of applications and workloads, demand for both accelerators and CPUs is growing well above our prior expectations.

We now expect the data center AI accelerator market to grow more than 45% annually to approximately \$1.4 trillion by 2030, and we expect the server CPU market to grow more than 50% annually to approximately \$220 billion by 2030.

For AMD, this larger opportunity, combined with the strength of our portfolio and growing customer visibility, is creating a steeper growth trajectory for our data center business. In data center AI, the growing number and scale of Helios and MI450 series Instinct deployments positioned the business for significant growth in the second half of the year, with growth accelerating in 2027.

In server CPUs with very strong customer demand and improved supply, we now expect server revenue to grow more than 80% year-over-year in the second half of 2026, and more than 70% for the full year 2027 off a much higher base. Taken together, we now expect data center segment revenue to more than double year-over-year in 2027.

Turning to client and gaming, segment revenue grew 6% year-over-year to 3.8 billion. In client, revenue increased 23% year-over-year to 3.1 billion, driven by record mobile processor revenue and continued share gains.

Commercial adoption continued to expand in the quarter, with Ryzen PRO sales growing more than 50% year-over-year as we closed new wins with large health care, technology, automotive, and financial services companies.

To build on this momentum, Dell, HP, Lenovo, ASUS, and others launched a broad portfolio of new commercial PCs powered by our latest generation Ryzen AI PRO 400 series processors. Demand was also strong for Ryzen AI Halo developer systems, which went on sale in the quarter.

In July, we introduced our next generation Ryzen AI Halo platform, powered by our new Gorgon Halo processor, featuring an industry leading 192 gigabytes of unified memory and can run models with up to 300 billion parameters.

And to make it even easier for developers to build and test large AI models locally, we are partnering with Hugging Face to include one year of Hugging Face Pro with every Ryzen AI Halo system beginning later this year.

Looking to the second half of the year, we're planning for softer PC market as higher memory and component costs weigh on demand. Against this backdrop, we expect our client business to perform better than the market, driven by the strength of our Ryzen portfolio and growing commercial adoption.

In gaming, revenue declined 31% year-over-year to \$779 million, primarily due to lower semi-custom sales at this stage of the console cycle. Gaming graphics revenue also declined year-over-year, as higher industrywide component costs contributed to higher graphics card prices and weighed on overall demand.

Turning to our embedded segment, revenue increased 19% year-over-year to \$977 million, our strongest growth in more than three years. Demand was broad-based, with strength across networking, aerospace and defense, test, measurement, and emulation, and communications customers.

Our embedded x86 business grew significantly in the quarter, as hyperscalers and networking customers increasingly adopted our CPUs to power critical networking and control plane functions in the data center.

We also continued to expand our portfolio introducing Ryzen AI embedded x100 processors for demanding real-time edge AI workloads and the Kria AI robotics platform for physical AI. Looking more broadly, the strategy we have been executing over the last few years is now delivering strong results.

Embedded x86 is becoming a significant growth driver for the segment. Our overall embedded portfolio is outgrowing the market and gaining share, and our embedded semi-custom engagements are expanding.

Design win momentum also remains very strong. We are tracking towards another record year with more than 18 billion of new design wins, led by major wins with networking, data center, communications, test, and aerospace and defense customers.

In summary, we delivered record revenue and profitability in the second quarter, reflecting our strong execution and the growing adoption of our leadership products. We enter the second half with strong momentum across our businesses, with Venice and MI455X now in production, initial Helios shipments set to begin this quarter, Ryzen PRO CPUs driving continued commercial share gains, and our embedded segment returning to strong year-over-year growth.

More than a decade of focused investment has given us the strongest and broadest product portfolio in the industry, deep strategic relationships with the companies driving the future of computing, and a proven ability to deliver multi generation roadmaps and ramp complex products at scale.

At the same time, AI is driving demand for dramatically more compute across all of our markets. We now see the overall market for high performance and AI computing growing approximately 40% annually over the next several years, approaching 2 trillion by 2030. And we expect to grow well above the market.

As a result, we are tracking materially ahead of the long-term financial model we shared at our financial analyst day last November. We now expect revenue to grow substantially above our prior target of greater than 35%, and we expect to significantly exceed our \$20 annual EPS target within our strategic time frame.

We are still in the early innings of a multiyear AI adoption cycle, and the opportunity ahead is enormous. We are exceptionally well positioned to capitalize on this opportunity and deliver significant growth in the coming years. Now I will turn the call over to Jean to provide additional color on our second quarter results. Jean.

JEAN HU:

Thank you, Lisa. And good afternoon, everyone. I'll start with a review of our second quarter financial results and then provide our current outlook for the third quarter of fiscal 2026. We had an outstanding second quarter, marking our sixth consecutive quarter of greater than 30% year-over-year revenue growth.

Revenue increased 50% year-over-year and 13% sequentially to a record \$11.5 billion, driven by continued momentum across our businesses. Importantly, data center accounted for approximately 58% of total revenue, underscoring the continued shift in our business mix and its increasing contribution to AMD's growth.

On a comparable basis, diluted earnings per share increased approximately 82% year-over-year, significantly outpacing our revenue growth and demonstrating our earnings power as we scale our business.

Gross margin for the quarter expanded to 56% up over 200 basis points year-over-year and 80 basis points sequentially, reflecting a favorable product mix and a growing contribution from our data center business.

Operating expenses were \$3.4 billion, an increase of 40% year-over-year as we continue to invest in R&D to expand our AI silicon systems and software capabilities to support our long-term growth opportunities. Operating income was \$3.1 billion, representing a 27% operating margin.

Now turning to our reportable segment, starting with the data center segment. Revenue was a record 6.7 billion, more than doubling year-over-year and up 16% sequentially. Growth was driven by greater than 70% year-over-year increase in EPYC sales, with record enterprise sell through and robust demand across our cloud customer base.

Both unit shipments and ASP increased significantly year-over-year, reflecting the continued mix shift toward to our latest Zen 5 generation processors. Instinct sales more than doubled year-over-year, with a continued ramp of our MI350 series products, as adoption broadened across the largest AI labs, cloud providers, leading AI startups, national labs, and the sovereign AI deployment.

Data center segment operating income was \$2.1 billion or 31% of revenue. Client and the gaming segment revenue was \$3.8 billion, up 6% year-over-year and 7% sequentially. The client business revenue was \$3.1 billion, up 23% year-over-year and 6% sequentially, led by record mobile revenue.

The gaming business revenue was \$779 million, down 31% year-over-year, primarily due to lower semi-custom revenue. Sequentially, gaming revenue increased 8% driven by a higher semi-custom sales, partially offset by lower Radeon shipments.

Client gaming segment operating income was \$582 million or 15% of revenue, compared to \$767 million or 21% a year ago, reflecting continued investment in go-to market activities. And expanding product roadmap.

Embedded segment revenue was \$977 million, up 19% year-over-year, and 12% sequentially as demand continued to improve across end markets and the new design wins began ramping. Embedded segment operating income was \$386 million or 40% of revenue compared to \$275 million or 33% a year ago, driven by higher revenue and favorable product mix.

Turning to the balance sheet and cash flow, during the quarter, we generated a \$2.4 billion in cash from continuing operations and a \$1.6 billion in free cash flow. Inventory increased sequentially to approximately 8.5 billion to support strong data center demand. At the end of the quarter, Cash, cash equivalents, and short-term investment were 13.1 billion.

Now turning to our third quarter 2026 outlook. We expect revenue to be approximately \$13 billion, plus or minus 300 million. At the midpoint of our guidance, revenue is expected to be up 41% year-over-year driven by very strong double digit growth in our data center segment.

Strong double digit growth in our embedded segment and the decline in the client gaming segment, with growth in our client business more than offset by significant double digit decline in gaming. Sequentially, we expect revenue to be up approximately 13% driven by strong double digit growth in both our data center and embedded segments, and a modest decline in our client gaming segment, with a slight growth in client, offset by strong double digit decline in gaming.

In addition, we expect third quarter non-GAAP gross margin to be approximately 56%. Non-GAAP operating expenses to be approximately 3.65 billion. Non-GAAP other income and expense to be a gain of approximately 55 million. Non-GAAP effective tax rate to be 13%, and a diluted share count is expected to be approximately 1.66 billion shares.

In closing, we delivered another outstanding quarter of revenue growth and a significant earnings expansion, reflecting the strength of our execution and the strong momentum across the business. As data center growth continues to accelerate, we entered the second half of the year very well positioned to deliver continued profitable growth. With that, I'll turn it back to Matt for the Q&A session.

MATT RAMSAY: Thank you very much, Jean. John, we'd like to go ahead and start the Q&A session now. Please do poll the audience for questions. We ask that each caller ask one question with one brief follow-up. Thank you.

JOHN: Thank you, Matt. We will now be conducting a question and answer session. We ask that you please limit yourself to one question and one follow-up. If you would like to ask a question, please press star 1 on your telephone keypad. A confirmation tone will indicate that your line is in the question queue. You may press star 2 if you'd like to remove a question from the queue.

For participants using speaker equipment, it may be necessary to pick up your handset before pressing the star keys. One moment while we poll for questions. And the first question comes from the line of Tom O'Malley with Barclays. Please proceed with your question.

TOM O'MALLEY: Thanks for taking my question. Lisa, you laid out the analyst day, \$1.4 trillion TAM on the accelerator side. You laid out the 220 billion on the CPU side. But the share dynamics within those are very different.

So I think the genesis of the question here is, when do you see that crossover point happening in the data center GPU business and the CPU business? Is that something that comes this calendar year? And when you look at the relative contributions between those two buckets into calendar year '27 from a dollar perspective, where are you seeing the most strength?

LISA SU: Yeah, Tom, thanks for the question. So look, we are certainly seeing a very, very strong compute market demand over both the data center accelerator, as well as the server CPU. I think the server CPU is a little bit newer. And so we've been giving updates on that.

As we go forward, certainly, looking at our Q3 guide and our Q4 guide, we see very strong growth across both server and data center AI. So we're seeing server grow over 80% in the second half of the year-over-year. We're seeing very strong data center AI growth. And that reflects into our 2027.

The way to think about it is both businesses are going to grow a lot. The data center AI-TAM is certainly larger. And, as a result, as we ramp the large strategic customers that we have on Helios in 2027, you would expect substantial growth in that business in 2027. But both businesses are very significant drivers of our 2027 growth and beyond.

TOM O'MALLEY: And then maybe under that framework, with both growing very strongly into next year, maybe this one's for you Jean. You, obviously, have some headwinds on the GPU side as you're ramping Helios. But also on the server processor side, you would imagine that you're getting above corporate gross margins there.

Maybe as you look at the trade-offs to gross margin into next year, is one kind of offsetting the other? Or could you maybe think that directionally, just given some of the tailwinds you're seeing with pricing, directionally with the mix shift to more CPUs, could you see a tailwind more so than a headwind into calendar year '27? Thank you very much, guys.

JEAN HU: Yeah, Tom, that's a great question. Thank you for that. I think as we talked before, our gross margin is primarily driven by the business mix. We're actually very pleased with our gross margin progress in the first half of 2026 and our guidance for Q3 2026.

I think the way to think about the 2027 is there are a few puts and takes. The first thing is, as Lisa just mentioned, our server business is expanding very significantly with the increased TAM and the pace of our business continue to grow, which from a gross margin perspective, is accretive to our overall gross margin.

On the other side, you're right, the pace of data center AI ramp is also very important because from a revenue opportunity perspective, we see large incremental revenue opportunity from data center AI. Once it ramps, not only it add tremendous revenue for the company, but also gross profit, even though the gross margin is slightly below corporate average.

I think the mix will determine how we go through 2027. But overall, we actually feel good about how we can really navigate through the transition of the business and continue to improve gross margin to manage the balance of the gross margin. Plus, our embedded business is recovering significantly. We actually think embedded business will give us an additional tailwind on the gross margin side in 2027. So I feel really good about that.

JOHN: Thank you. And the next question comes from the line of Timothy Arcuri with UBS. Please proceed with your question.

TIMOTHY ARCURI: Thanks a lot. Lisa, I wanted to see if you could give us any color, just what's embedded in the guidance in terms of data centers? So the question really is like, of the growth between Q2 and Q3, which will grow more on dollar basis CPU or GPU? And I ask because if I take your up 80% half on half number, it implies that server CPU might not grow that much in Q4. So I'm not sure if I'm doing that wrong. Can you just give us some color there?

LISA SU: I'm just looking at that, Tim. Look, I think we believe that the server CPU is going to grow substantially in Q3 and in Q4. The way I would describe it in terms of dynamics is we have very strong demand across all of our customers in hyperscale and enterprise.

And we have been adding additional supply as we've increased supply throughout the year. So we're guiding the Q3 segment to grow sequentially double digits. Both server and the data center AI will grow nicely within that.

And then we should expect more growth for each of those businesses in Q4, with more supply coming on for server, as well as for data center AI. The Helios ramp is just starting at the end of Q3, and it will be much more substantial in Q4.

So that should give you the picture of we think Q3 is certainly a strong quarter as we look at the very strong double digit growth going in data center segment. And in Q4, it will be higher than that.

TIMOTHY ARCURI: OK, thanks. And then just as a follow-up. So relative to the TAM that you laid out at the analyst day, I think it's a 40% CAGR out through 2030. So, do you think you can outgrow that CAGR out through 2030? So is your plan to grow your revenue more than 40% during that period? Thanks.

LISA SU: Yes. So as we look at the overall TAM, first of all, it's a very exciting time in the compute market when you look at every part of our business, whether it's server, CPU, data center, AI, or our embedded business and our PC business, we see them all benefiting from the AI tailwinds.

And beneath that, I think in each one of the segments, we see an opportunity to grow above the market. So from that standpoint with our view of the TAM being greater than 40%, we are also saying that we will grow greater than that from an overall revenue standpoint for the company.

JOHN: Thank you. And the next question comes from the line of Vivek Arya with Bank of America Securities. Please proceed with your question.

VIVEK ARYA: Thanks for taking my questions. Lisa, I'm curious, with all the announcements that you have made, how many gigawatt of compute does AMD have line of sight into for 2027? And what is AMD's monetization per gigawatt?

I think the current assumption from everything you've announced is about three gigawatts. And I think from some of the warrant numbers that you had mentioned before, it seems to be at least \$15 billion per gigawatt. So I'm just curious if you agree with those views. And then what is the upside or downside to either the gigawatt or the content assumption?

LISA SU: Yeah. So Vivek, there are lots of pieces to that question. So maybe let me take them one at a time. From a overall data center standpoint, I think we have given a little bit of color on 2027. We believe that the server CPU portion of our data center business will grow by over 70% year over year off of a higher base, and we believe the overall segment will grow by over 100%. So will more than double due to the data center AI ramping.

So we do see a very significant ramp into 2027. We are very happy with our strategic anchor customers in OpenAI, Meta, and Anthropic. They are all in the process of ramping. The compute demand is high. And this is very much aligning their data center build outs and their work with their cloud service providers together with our ramp.

The way I would describe it is we have supply to more than meet the guidance that we've talked about and the upside potential is there. The key thing is to work very closely with our customers as they're planning their build outs overall.

So I think from a Helios ramp standpoint, we're expecting a very significant ramp over the next couple of quarters. We've said Q3 is the very beginning of the ramp, Q4 is a step-up, and then Q1 will be a further step-up and we'll ramp as we go through 2027. So hopefully, that gives you a little bit of color on how that data center AI will play out.

And in terms of the revenue per gigawatt, I think we've said double digit billions. We're still in that range. I think that looks like that's pretty much where it will be. And the key here will be continuing to work with our partners as they ramp because there's a demand for a lot more compute. And we would like to satisfy that demand.

VIVEK ARYA: Thank you, Lisa. And for my follow-up, a little of a technical question. When we look at the specs of Helios, it has, I think, almost 50% more HBM than the competition. And I think you have always had more HBM in your product. So how much of the benefit? Is because of having 50% more HBM. And then doesn't that expose AMD more to memory cost inflation?

So I'm just curious, how are you ensuring that you're able to maintain margins and get this memory allocation? And I think Jean already suggested you're comfortable with the gross margin range. But I'm curious, how much of the benefit is just because of the use of more HBM? And does that still enable you to meet your profitability targets and allocation requirements over time? Thank you.

LISA SU: Yeah. So look, Vivek, we're working very, very closely with our memory partners across the board on both our GPU, HBM memory, as well as just the general memory for the systems across our data center business.

And what I would say is, we've worked with our memory partners for multiple years to ensure that we have a strong ramp in our business. There is very good visibility into HBM allocation for what we expect to deliver in 2027.

And the other piece of it is in terms of the memory bandwidth and the memory capacity, it is one of the advantages of the AMD solutions. When you think about the larger model sizes, they really benefit from the larger memory footprint, and that goes into the total cost of ownership.

We do recognize in this memory environment, though, every customer is looking at how to optimize their memory footprint. And we have the opportunity, if desired, to also modify that memory footprint if the total cost of ownership is not as significant in certain workloads.

So think about it as the benefit of the memory is workload-dependent. And we know a bunch of our customers are very, very happy with what that returns in terms of performance. But there are some workloads, let's call it medium-sized models that may not get as much of a benefit. And in that case, we would address the memory footprint as you might expect.

JOHN: Thank you. The next question comes from the line of Joshua Buckwalter with TD Cowen. Please proceed with your question.

JOSHUA BUCHALTER: Hey, guys. Thank you for taking my question and congrats on the very strong results. I wanted to also follow up on the server CPU assumptions for 2027. So it sounds like the supply for that is locked up and you have ability to service higher demand if it continues to track that way.

But, I guess, any help you can give us on the unit and ASP assumptions that are baked in there? I think investors are struggling to model the CPU market overall. So how should we think about, I guess, core count growth or whatever is the right proxy as we think about modeling this business? Thank you.

LISA SU:

Yeah. Sure, Josh. So let me maybe start with a little bit of grounding on the growth that we have seen so far. There are certainly both unit and ASP growth. So if you just look at our Q2 performance, when we said we grew over 70% in both cloud and enterprise, we actually had double digit growth in both units and ASPs, but it was actually more unit.

The unit growth was higher. That is very much the nature of our business. I mean, we're seeing just very strong demand from a overall market standpoint. On the ASP growth, we have ASP growth as we go to higher core counts, certainly.

But as we go forward, you should expect both unit and ASP growth. And what we have been working on very diligently over the past couple of quarters is as soon as we saw the significant inflection point in server demand, we've been working across our supply chain. That's wafers, that's backend capacity, that's substrates. That's all of the components to raise the overall capacity for servers.

And we're seeing that play out through this year. That's one of the reasons we can raise the second half guidance. And we're seeing much more capacity coming online in 2027 that supports the growth that we've been talking about. So you should think about both units and ASP in this framework.

**JOSHUA
BUCHALTER:**

Thank you for all the color there. And Elisa, in your prepared remarks, you mentioned that Helios was ahead of your original forecasts. Can we unpack that comment a little bit? Was that a comment on volumes? Was it yields? And if it's yields, how should we think about the first quarter or two of Helios' gross margins compared to as it gets later into its ramp? Should we expect gross margins to improve as it ramps? Thank you.

LISA SU:

Yes. So Josh, when I was talking about Helios was ahead of our initial forecast, it was as it relates to the overall volumes. So let's call it the amount of demand there is for Helios in 2027. It's like an outstanding product.

So what we're seeing from every one of our customers who's had a chance to not only spend time with Helios, but also spend time in our overall ecosystem, there's a high confidence that Helios will be a great addition to the AI portfolio, particularly around inference. And that was my comment about higher than our initial expectations.

As it relates to your comments about yields and performance and what do we expect as we go through the ramp, one would expect that the overall yields will improve as we go through the next few quarters.

The starting quarters, this quarter here in Q3, and we will be ramping over the next couple of quarters. And we always would expect that the yields will continue to improve as we go through the first few quarters, especially on a product like this, which is highly complex.

JOHN:

Thank you. And the next question comes from the line of Aaron Rakers with Wells Fargo. Please proceed with your question.

AARON

Yeah. Thanks for taking the question. As you can imagine, I'll stick with the server piece of the business as well.

RAKERS:

At your recent event, you highlighted the server market, the new 220 billion TAM that you're throwing out there as in three buckets. General purpose, I think you called it or defined it as sandbox servers, and then the AI frontend node.

So, as we think about the growth that you're talking about, I'm wondering if you could help us conceptualize the sizing of those buckets at all. And how big is that kind of middle category, that sandbox AI market as we think about growing 80% plus this back half of the year and 70% plus? How big of that opportunity do you see that becoming?

LISA SU:

Yeah. So Aaron, what I would say is-- maybe let me start from the endpoint. So when you think about the \$220 billion TAM in 2030, we actually see this agentic AI or these agentic sandboxes being the largest piece of the TAM. It's the fastest growing piece. It's also the smallest piece today.

So in the near term, I would say that there is certainly growth in that area. But as we go out over the next 3, 4, 5 years, we think that's the largest growth in the server market. And I think the thing that is very strong about our portfolios, we believe we can grow in every one of those segments.

So when you think about the types of CPUs you need for each of these workloads, you need different optimization points. That's one of the things that we've really tried to point out. It's not just one CPU, it's actually a real family of CPUs that you need.

From our standpoint, Venice is absolutely leadership in all of the categories. So whether you're talking about per core performance or you're talking about overall socket performance. And that gives us a very strong position across general purpose agentic AI and also the head nodes for the AI accelerator business.

So from our standpoint, when we look at our customer traction and our customer momentum going from Turin into Venice Turin was already extremely widely used, especially amongst the hyperscalers. What we're seeing is when we go into Venice, the workloads actually expand.

So there are more workloads that are going to be run on the next generation of EPYC than are run on the previous generation. And that's what gives us the confidence to say that we can grow substantially ahead of the market given the product positioning.

AARON

RAKERS:

Yeah, thank you. That was very helpful. And as a quick follow-up, I know it's probably lower on the radar for a lot of people. But I think in the past, you've talked about the client business, even with some pressures in the PC market growing this year for you guys on a year-over-year basis. Do you still see that? And are you willing to give any thoughts on what you might think of 2027 on the client CPU side?

LISA SU:

Yeah. Well, let me say, first of all, the client CPU business, we have continued to think it's a very important business. I particularly think that local AI will continue to become a larger and larger piece of how people experience AI.

So I think AI PCs will become more important. As it relates to our client business in 2026, I think our first half performance has been very strong. And although we are expecting that the market will decline in the second half, the market's actually held up better than most people would have thought.

So our view is we see the growth in 2026 on a year-over-year basis certainly there. And as we go into 2027, I think we have a strong product portfolio that is coming on board to address not just the traditional notebook and desktop markets, but also as you think about a more AI-centric PC experience, like what we have been talking about with the rise in AI Halo.

So I remain optimistic about the PC market as an important way for us to reach a broader set of users and our ability to grow ahead of the market. The market itself will depend a bit on some of the components. We're all watching the component costs and how that will play out over the next couple of quarters. But I think our portfolio and our rising content and enterprise are a very positive for our client mix.

JOHN: And the next question comes from the line of Stacy Rasgon with Bernstein Research. Please proceed with your question.

STACY RASGON: Hi, guys. Thanks for taking my question. First thing, I wanted to pick on something from the press release that says that the data center accelerates in the second half. So it grew 107% year-over-year in Q2. You're guiding it up like in the '80s, though, like in Q3.

So, how do I interpret, which is not an acceleration, how do I interpret that statement? Is like the second half collectively higher than 107% that you see in Q2? Or are you just like-- how should I be interpreting what you guys wrote in the press release around data center accelerating in the second half?

JEAN HU: Yeah. Hi, Stacy. This is Jean. Good question. I think when we talk about this acceleration, we're talking about the second half versus the first half. If you look at the first half, our data center business' year-over-year growth versus the second half, we do think there's an acceleration. That's the reflection there.

STACY RASGON: Thank you for that. And then I think for my follow-up, I want to dig in a little more into the data center targets for next year. So you said like more, I think server is more than 70%, the total more than doubling.

I guess, what I'm asking is how much work is the more in that statement doing? You seem to be on a trajectory that would take you well above those numbers. I guess, just how should I be interpreting the more in the statement? Understanding that we're only in August of 26th, I get it. But what's on your mind there?

JEAN HU: Stacey, you bring a smile to my face. So I would say we're trying to give you a way to think about 2027. And yes, the server CPU we're saying is more than 70% which we think at this point is a very strong statement just given where the business is.

And we do expect the overall data center business to be well over 100%. And the well over 100% comes because the data center AI business is going to be well over 100% just given the strength of our strategic customers, the ramp of Helios, and all of the things that we've talked about. So hopefully, that answers your more questions.

JOHN: Thank you. And the next question comes from the line of Jim Schneider with Goldman Sachs. Please proceed with your question.

JIM SCNEIDER: Good afternoon. Thanks for taking my question. As you think about the early days of your MI400 series ramp and over the next few quarters, can you maybe talk about the customer diversity you expect in the early stages? I think you've noted four or five large customers that you've announced publicly. Can you maybe talk about how many of those customers will be contributing, let's say, in Q4, and how many in the first half of '27?

LISA SU: Let's see. So, Jim, I guess the way to say it is, we've talked about the large frontier model companies, OpenAI, Anthropic, Meta. They will be consuming through a number of CSPs, both hyperscalers and others.

So the diversity of the business-- and there are additional customers, there are lots of customers who are interested in Helios at let's call it a more regular scale than gigawatt scale. So I think we'll have a good diversity of customers as we go through the next couple of quarters.

And what we're really doing is matching to when the larger data center build outs are ready. And so we are working with each of our customers on their data center plans and ensuring that we are meeting their data center plans. But yeah, I think there's a good diversity, especially as we get into the Q4/Q1 time frame.

JIM SCNEIDER: Thanks. And maybe as a follow-up, you mentioned data center readiness. And so I wanted to test how you're seeing your various customers and their ability to accept your products into their data centers, whether that's land PowerShell or anything else. Any constraints that you see heading into the first half of '27 or even the back half of '27 that would give you pause about hitting your targets? Thank you.

LISA SU: Well, Jim, I don't think we see anything that would give us pause about hitting the targets. I mean, we feel very good about the targets. Now, when you ask me what the range could be, there is a range. And the range will depend on the ability to bring on more capacity in a timely fashion.

So the way I view it is we are building the entire supply chain. So ensuring that we have the entire-- our silicon CPU, GPU, networking components that go into it, all of the Helios components. I think we feel very, very good about that part of the supply chain.

We are working very closely on the data center operators and ensuring that we have good visibility into what's going on there. And that gives us, let's call it, strong confidence in what we've guided for the data center AI business so far.

And we're going to be continuing to look at how to accelerate some of those builds. And I think there's a clear desire on the part of everyone in the ecosystem to bring on more AI compute faster. We are seeing every day more opportunities with operators to accelerate some of that capacity. And that's much of the work that we're doing together with our customers and partners.

JOHN: Thank you. And the next question comes from the line of CJ Muse with Cantor. Please proceed with your question.

C.J MUSE: Yeah, good afternoon. Thank you for taking the question. If I take your data center guide, Lisa, it implies Instinct's revenues of \$30 billion, give or take, and so curious two parts on this. How do you see the revenue cadence first half/second half?

And then if we isolate the Instinct only and start thinking about 455X and Helios, how should we think about the underlying gross margins for that business starting in Q1 and then exiting in Q4?

JEAN HU: I think, C.J, you're saying \$30 billion. Are you talking about the server or are you talking about-- which one are you talking? You're talking about 2027, right?

C.J MUSE: 2027 Instinct.

JEAN HU: Yeah. 2027.

LISA SU: Maybe let me help C.J, I think what you're hearing from us is that your data center AI number is probably too low. And maybe back to Stacy's question, without going into exact numbers, I think this notion of over 100% should consider be well over 100%.

And look, as we look at this, it's a progression over the next couple of quarters. And as we go through the next couple of quarters, we expect to continue growing Helios quarter by quarter. But yes, that's the best way of thinking about it. And I'm sure Jean can work on, if you have further questions on that. But Jean, does that--

JEAN HU: Yeah, yeah.

C.J MUSE: And on the gross margin side, if we isolate to Instinct, how should we think about beginning to 2027 and then exiting?

JEAN HU: Yeah, I think overall when we think about the gross margin for 2027, as I said earlier, it's really determined by both the pace of server CPU business ramp and the data center AI ramp. So I actually think, overall quarter over quarter, it could be different mix. It could change differently.

But overall, when we think about this, we are actually very optimistic about how we manage the ramp of the MI450. At the same time, with the server business, continue to improve and grow significantly in 2027.

Overall, I think they have a good offset in general overall. But I do think it's important to remember we have multiple other levers from the company perspective, not only embedded business, but also client business. We continue to improve gross margin. And operation team continue to do a great job. So we'll give you more color when we get there to guide the 2027.

JOHN: Thank you. And the next question comes from the line of Joe Moore with Morgan Stanley. Please proceed with your question.

JOE MOORE: Great. Thank you. In the server CPU business, are you supply constrained now? It seems like the market is very tight and do you anticipate that tightness persisting. And when you think about next year, just the general ability of the supply chain to support the level of growth that you're talking about, specifically in CPU.

LISA SU: Sure, Joe. I would say the server CPU supply chain is tight right now and it has been tight for the first half of the year because much of this demand was unforecasted. As we get into 2027, the demand is better forecasted.

And so we would expect that the 2027 server supply situation should be better than '26. We feel very good about being able to satisfy what we just talked about, which was over 70% year-over-year growth. And I believe, depending on how things play out, there may be opportunities for that growth to go higher as we get through the next few quarters.

JOE MOORE: OK. Thank you for that. And you guys are ramping a lot on 2 nanometer. And there's been a lot focus on 3 nanometer being in short supply. Is that helping or is it still challenging to bring up new capacity on a new node that way?

LISA SU: Yeah, I think, Joe, it's always challenging to bring up new capacity on new node. I think what makes our approach a little bit special and different is that because we're using the chiplet technology, we actually ramp in fewer wafers in the new node.

And so that gives us the opportunity to, again, we're working very hard on ensuring that we get the supply necessary to meet the very strong customer demand. So from my perspective, I think all of that is work that's being done.

We're certainly looking at the overall supply chain, not just wafers. So that includes backend capacity, packaging capacity, substrates, all of those things. But we feel good about where we are to satisfy both the strong ramp and servers, as well as the strong ramp in the data center AI business.

MATT RAMSAY: Hey John, I think we have time for one more caller before we close out the call, please. Thank you.

JOHN: Thank you. Our final question comes from the line of Atif Malik with Citi. Please proceed with your question.

ATIF MALIK: Hi. Thank you for taking my questions. Lisa, you guys announced a partnership with Cerebras both at your advanced AI day and on disaggregate compute. Can you talk about just qualitatively, how do you expect the sales to grow in the fast inference market this year into next year?

LISA SU: Yeah, absolutely. So I think the inference market overall is growing very, very substantially over this year into next year. The fast inference in particular is an area which is, let's call it starting to become more and more relevant.

And so with our partnership with Cerebras, I think they have great technology together with Helios plus their wafer scale engine, we get a very good solution for customers. And we would expect that solution to be start becoming available in Q4 in the Cerebras cloud and extend into 2027.

But from our view, this is an important part of the market, and we continue to look at ways to, I would say, customize and optimize our technologies for the various workloads out there. So we view this as just more of what we do in an open ecosystem.

ATIF MALIK: Great. And one for Jean. How should we think about OpEx growth relative to that overall 40% market growth you talked about in the next few years?

JEAN HU: Yeah, I think Lisa talked about the TAM growth at 40% and it will be growing faster than TAM. From OpEx perspective, we'll continue to invest given the large opportunities we have ahead of us. But you should expect us to manage OpEx increase less than the top line revenue growth.

That's what our business model is designed for. So we can drive more operating leverage to deliver earnings per share. That is also consistent with what Lisa said, our EPS will be significantly higher than \$20 we outlined at our financial analyst day.

MATT RAMSAY: Thank you very much for all the analysts and investors that joined our call today. John, you can go ahead and wrap up the call. Thank you.

JOHN: Thank you, ladies and gentlemen. That does conclude the question and answer session. And that also concludes today's teleconference. We thank you for your participation. You may disconnect your lines at this time.