

AMD Fiscal First Quarter 2024 Financial Results

PAUL: Greetings. And welcome to the AMD first quarter 2024 conference call. At this time, all participants are in a listen-only mode. A brief 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. As a reminder, this conference is being recorded. It is now my pleasure to introduce your host, Mitch Haws, vice president, investor relations. Thank you, Mitch. You may begin.

MITCH HAWS: Thank you. And welcome to AMD's first quarter 2024 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. We will refer primarily to non-GAAP financial measures during today's call. And the full non-GAAP to GAAP reconciliations are available in today's press release and the slides posted on our website.

Participants on today's call are Dr. Lisa Su, our chair and chief executive officer, and Jean Hu, our executive vice president, chief financial officer, 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 Mark Papermaster, executive vice president and chief technology officer, will attend the TD Cowen Technology, Media, and Telecom Conference on May 29. And Jean Hu, executive vice president, chief financial officer, and treasurer, will attend the JP Morgan Global Media and Communications Conference on Tuesday, May 21, the Bank of America Global Technology Conference on Wednesday, June 5, and the Jefferies and Nasdaq investor conference on Tuesday, June 11.

Today's discussion contains forward-looking statements based on current beliefs, assumptions, and expectations, 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 the factors that could cause actual results to differ materially. With that, I will hand the call over to Lisa.

LISA SU: Thank you, Mitch. And good afternoon to all those listening today. This is an incredibly exciting time for the industry as the widespread deployment of AI is driving demand for significantly more compute across a broad range of markets. Under this backdrop, we are executing very well as we ramp our data center business and enable AI capabilities across our product portfolio.

Looking at the first quarter, revenue increased to \$5.5 billion. We expanded gross margin by more than two percentage points and increased profitability as data center and client segment sales each grew by more than 80% year over year. Data center segment revenue grew 80% year over year and 2% sequentially to a record \$2.3 billion.

The substantial year over year growth was driven by the strong ramp of AMD Instinct MI300X GPU shipments and a double-digit percentage increase in server CPU sales. We believe we gained server CPU revenue share in the seasonally down first quarter, led by growth in enterprise adoption and expanded cloud deployments.

In cloud, while the overall demand environment remained mixed, hyperscalers continued adopting fourth gen EPYC processors to power more of their internal workloads and public instances. There are now nearly 900 AMD-powered public instances available globally as Amazon, Microsoft, and Google all increased their fourth gen EPYC processor offerings with new instances and regional deployments.

In the enterprise, we have seen signs of improving demand as CIOs need to add more general purpose and AI compute capacity while maintaining the physical footprint and power needs of their current infrastructure. This scenario aligns perfectly with the value proposition of our EPYC processors.

Given our high core count and energy efficiency, we can deliver the same amount of compute with 45% fewer servers compared to the competition, cutting initial CapEx by up to half, and lowering annual OpEx by more than 40%. As a result, enterprise adoption of CPUs is accelerating, highlighted by deployments with large enterprises, including American Airlines, DBS, Emirates Bank, Shell, and STMicro.

We're also building momentum with AMD-powered solutions powering the most popular ERP and database applications. As one example, the latest generation of Oracle Exadata, the leading database solution used by 76 of the Fortune 100, is now powered exclusively by fourth Gen EPYC processors.

Looking ahead, we're very excited about our next-gen Turin family of EPYC processors featuring our Zen 5 core. We're widely sampling Turin, and the silicon is looking great. In the cloud, the significant performance and efficiency increases of Turin position us well to capture an even larger share of both first and third-party workloads. In addition, there are 30% more Turin platforms in development from our server partners compared to fourth gen EPYC platforms, increasing our enterprise SEM with new solutions optimized for additional workloads. Turin remains on track to launch later this year.

Turning to our broader data center portfolio, we delivered our second straight quarter of record data center GPU revenue as MI300 became the fastest ramping product in AMD history, passing \$1 billion in total sales in less than two quarters. In cloud, MI300X production deployments expanded at Microsoft, Meta, and Oracle to power generative AI training and inferencing for both internal workloads and a broad set of public offerings.

For the enterprise, we're working very closely with Dell, HPE, Lenovo, Supermicro, and others as multiple MI300X platforms enter volume production this quarter. In addition, we have more than 100 enterprise and AI customers actively developing or deploying MI300X.

On the AI software front, we made excellent progress, adding upstream support for AMD hardware in the OpenAI Triton compiler, making it even easier to develop highly performant AI software for AMD platforms. We also released a major update to our ROCm software stack that expands support for open-source libraries, including vLLM, and frameworks, including JAX, adds new features like video decode, and significantly increases generative AI performance by integrating advanced attention algorithms support for sparsity and FP8.

Our partners are seeing very strong performance in their AI workloads. As we jointly optimize for their models, MI300x GPUs are delivering leadership inferencing performance and substantial TCO advantages compared to H100. For instance, several of our partners are seeing significant increases in tokens per second when running their flagship LLMs on MI300x compared to H100.

We're also continuing to enable the broad ecosystem required to power the next generation of AI systems, including as a founding member of the Ultra Ethernet Consortium, working to optimize the widely adopted ethernet protocol to run AI workloads at data center scale.

MI300 demand continues to strengthen. And based on our expanding customer engagements, we now expect data center GPU revenue to exceed \$4 billion in 2024, up from the \$3 and 1/2 billion we guided in January. Longer term, we are increasingly working closer with our cloud and enterprise customers as we expand and accelerate our AI hardware and software roadmaps and grow our data center GPU footprint.

Turning to our client segment, revenue was \$1.4 billion, an increase of 85% year over year, driven by strong demand for our latest generation Ryzen mobile and desktop processors with OEMs and in the channel. Client segment revenue declined 6% sequentially.

We saw strong demand for our latest generation Ryzen processors in the first quarter. Ryzen desktop CPU sales grew by a strong double-digit percentage year over year. And Ryzen mobile CPU sales nearly doubled year over year as new Ryzen 8040 notebook designs from Acer, ASUS, HP, Lenovo, and others ramped.

We expanded our portfolio of leadership enterprise PC offerings with the launch of our Ryzen Pro 8000 processors earlier this month. Ryzen Pro 8040 mobile CPUs deliver industry-leading performance and battery life for commercial notebooks. And our Ryzen Pro 8000 series desktop CPUs are the first processor to offer dedicated on-chip AI accelerators in commercial desktop PCs. We see clear opportunities to gain additional commercial PC share based on the performance and efficiency advantages of our Ryzen Pro portfolio and an expanded set of AMD-powered commercial PCs from our OEM partners.

Looking forward, we believe the market is on track to return to annual growth in 2024, driven by the start of an enterprise refresh cycle and AI PC adoption. We see AI as the biggest inflection point in PCs since the internet, with the ability to deliver unprecedented productivity and usability gains. We're working very closely with Microsoft and a broad ecosystem of partners to enable the next generation of AI experiences powered by Ryzen processors with more than 150 ISVs on track to be developing for AMD AI PCs by the end of the year.

We will also take the next major step in our AI PC roadmap later this year with the launch of our next-generation Ryzen mobile processors codenamed Strix. Customer interest in Strix is very high based on the significant performance and energy efficiency uplifts we are delivering. Design win momentum for premium notebooks is outpacing prior generations as Strix enables next-generation AI experiences in laptops that are thinner, lighter, and faster than ever before. We're excited about the growth opportunities for the PC market. And based on the strength of our Ryzen CPU portfolio, we expect to grow revenue share this year.

Now, turning to our gaming segment. Revenue declined 48% year over year and 33% sequentially to \$922 million. First quarter semi-custom SoC sales declined in line with our projections as we are now in the fifth year of the console cycle. In gaming graphics, revenue declined year over year and sequentially. We expanded our Radeon 7000 series family with the global launch of our Radeon RX 7900 GRE and also introduced our driver-based AMD Fluid Motion Frames technology that can provide large performance increases in thousands of games.

Turning to our embedded segment, revenue decreased 46% year over year and 20% sequentially to \$846 million as customers remain focused on normalizing their inventory levels. We launched our Spartan UltraScale+ FPGA family with high I/O counts, power efficiency and state-of-the-art security features. And we're seeing a strong pipeline of growth for our cost-optimized embedded portfolio across multiple markets.

Given the current embedded market conditions, we're now expecting second quarter embedded segment revenue to be flat sequentially, with a gradual recovery in the second half of the year. Longer term, we see AI at the edge as a large growth opportunity that will drive increased demand for compute across a wide range of devices. To address this demand, we announced our second generation of Versal adaptive SoCs that deliver a 3x increase in AI TOPS per Watt and a 10x greater scalar compute performance compared to our prior generation of industry-leading adaptive SoCs.

Versal Gen 2 adaptive SoCs are the only solution that combine multiple compute engines to handle AI pre-processing, inferencing, and post-processing on a single chip, enabling customers to rapidly add highly performant and efficient AI capabilities to a broad range of products. We were pleased to be joined at our launch by Subaru, who announced they adopted Versal AI Edge series Gen 2 devices to power the next generation of their EyeSight ADAS system. Embedded design win momentum remains very strong as customers adopt our full portfolio of FPGAs, CPUs, GPUs, and adaptive SoCs to address a larger portion of their compute needs.

In summary, we executed well in the first quarter, setting us up to deliver strong annual revenue growth and expanded gross margin driven by growing adoption of our Instinct, EPYC, and Ryzen product portfolios. Our priorities for 2024 are very clear-- accelerate our data center growth by ramping Instinct GPU production and gaining share with our EPYC processors, launch our next-generation Zen 5 PC and server processors that extend our leadership performance, and expand our adaptive computing portfolio with differentiated solutions.

Looking further ahead, AI represents an unprecedented opportunity for AMD. While there has been significant growth in AI infrastructure buildouts, we are still in the very early stages of what we believe is going to be a period of sustained growth driven by an insatiable demand for both high-performance AI and general purpose compute.

We have expanded our investments across the company to capture this large growth opportunity from rapidly expanding our AI software stack to accelerating our AI hardware roadmaps, increasing our go-to-market activities, and partnering closely with the largest AI companies to co-optimize solutions for their most important workloads. We are very excited about the trajectory of the business and the significant growth opportunities ahead. Now, I'd like to turn the call over to Jean to provide some additional color on our first quarter results. Jean?

JEAN HU:

Thank you, Lisa. And good afternoon, everyone. I'll start with a review of our financial results and then provide our current outlook for the second quarter of fiscal 2024.

We delivered a strong year-over-year revenue growth in our data center and client segments in the first quarter and drew 230 basis points of gross margin expansion. For the first quarter of 2024, revenue was \$5.5 billion, up 2% year over year, as revenue growth in the data center and client segment was partially offset by lower revenue in our gaming and embedded segments.

Revenue declined 11% sequentially as higher data center revenue resulting from the ramp of our AMD Instinct GPUs was offset by lower gaming and embedded segment revenues. Gross margin was 52%, up 230 basis point year over year, driven by higher revenue contribution from the data center and the client segment, partially offset by lower embedded and gaming segment revenue contribution.

Operating expenses were \$1.7 billion, an increase of 10% year over year, as we continue to invest aggressively in R&D and marketing activities to address the significant AI growth opportunities ahead of us. Operating income was \$1.1 billion, representing a 21% operating margin. Taxes, interest expense, and other was \$120 million. For the fourth quarter of 2024, diluted earnings per share was \$0.62, an increase of 3% year over year.

Now turning to our reportable segments, starting with the data center. Data center delivered record quarterly segment revenue of \$2.3 billion, up 80%, a \$1 billion increase year over year. Data center accounted for more than 40% of total revenue, primarily led by the ramp of AMD Instinct GPUs from both cloud and enterprise customers, and a strong double-digit percentage growth in our server processing revenue as a result of growth across our Zen 4 products.

Our sequential basis revenue increased 2%, driven by the ramp of our AMD Instinct GPUs, partially offset by seasonal decline in server CPU sales. Data center segment operating income was \$541 million, or 23% of revenue, compared to \$148 million, or 11%, a year ago. Operating income was up 266% year over year due to operating leverage, even as we significantly increased our investment in R&D.

Client segment revenue was \$1.4 billion, up 85% year over year, driven primarily by Ryzen 8000 series processors. On a sequential basis, client revenue declined 6%. Client segment operating income was \$86 million, or 6% of revenue, compared to an operating loss of \$172 million a year ago, driven by higher revenue.

Gaming segment revenue was \$922 million, down 48% year over year and down 33% sequentially, due to a decrease in semi-custom and Radeon GPU sales. Gaming segment operating income was \$151 million, or 16% of revenue, compared to \$314 million, or 18%, a year ago.

Embedded segment revenue was \$846 million, down 46% year over year, 20% sequentially, as customers continue to manage their inventory levels. Embedded segment operating income was \$342 million, or 41% of revenue, compared to \$798 million, or 51%, a year ago.

Turning to the balance sheet and cash flow, during the quarter, we generated a \$521 million in cash from operations. And the free cash flow was \$379 million. Inventory increases sequentially by \$301 million to \$4.7 billion, primarily to support the continued ramp of data center and client products in advanced process nodes.

At the end of the quarter, cash, cash equivalents, and short-term investment was \$6 billion. As a reminder, we have \$750 million of debt maturing this June. Given our ample liquidity, we plan to retire the debt utilizing existing cash.

Now, turning to our second quarter 2024 outlook. We expect revenue to be approximately \$5.7 billion-plus, or minus \$300 million. Sequentially, we expect data center segment revenue to increase by double-digit percentage, primarily driven by the data center GPU ramp.

Client segment revenue to increase. Embedded segment revenue to be flat. And in the gaming segment, based on current demand signals, revenue to decline by significant double-digit percentage.

Year over year, we expect our data center and client segment revenue to be up significantly, driven by the strength of our product portfolio. The embedded and gaming segment revenue to decline by a significant double-digit percentage.

In addition, we expect second quarter non-GAAP gross margin to be approximately 53%, non-GAAP operating expenses to be approximately \$1.8 billion, non-GAAP effective tax rate to be 13%. And the diluted share count is expected to be approximately 1.64 billion shares.

In closing, we started the year strong. We made significant progress on our strategic priorities, delivered year-over-year revenue growth in our data center and client segment, and expanded the gross margin. Looking ahead, we believe the investment we are making will position us very well to address the large AI opportunities ahead. With that, I'll turn it back to Mitch for the Q&A session.

MITCH HAWS: Thank you, Jean. Paul, we're happy to poll the audience for questions.

PAUL: Thank you. We'll now be conducting a question and answer session. If you would like to ask a question, please press star 1 on your telephone keypad. A confirmation tone will indicate your line is in the question queue. You may press star 2 if you'd like to remove your question from the queue. For participants using speaker equipment, it may be necessary to pick up your handset before pressing the star key. One moment please, while we poll for questions.

Thank you. Our first question is from Toshiya Hari with Goldman Sachs. Please proceed with your question.

TOSHIYA HARI: Hi. Thank you so much for taking the question. Lisa, my first question is on the MI300. You're taking up the full-year outlook from \$3 and 1/2 to \$4 billion. I'm curious, what's driving that incremental \$500 million in revenue?

Is it new customers? Is it additional bookings from existing customers? Is it more cloud? Is it more enterprise? If you can provide color there, that would be helpful.

And then on the supply side, there's been headlines or chatter that CoWos and/or HBM could be a pretty severe constraining factor for you guys. If you can speak to how you're handling the supply side of the equation, that would be helpful, too. And then I have a quick follow up.

LISA SU: Great. Thank you, Toshiya, for the question. Look, the MI300 ramp is going really well. If we look at just what's happened over the last 90 days, we've been working very closely with our customers to qualify MI300 in their production data centers, both from a hardware standpoint, software standpoint.

So far, things are going quite well. And what we see now is just greater visibility to both current customers as well as new customers committing to MI300. So that gives us the confidence to go from \$3 and 1/2 to \$4 billion. And I view this as-- very much, it's a very dynamic market, and there are lots of customers. We said in the prepared remarks that we have over 100 customers that we're engaged with in both development as well as deployment. So overall, the ramp is going really well.

As it relates to the supply chain, actually, I would say I'm very pleased with how supply has ramped. It is absolutely the fastest product ramp that we have done. It's a very complex product, chiplets, CoWos, 3D integration, HBM. And so far, it's gone extremely well. We've gotten great support from our partners.

And so I would say, even in the quarter that we just finished, we actually did a little bit better than expected when we first started the quarter. I think Q2 will be another significant ramp. And we're going to ramp supply every quarter this year.

So I think the supply chain is going well. We are tight on supply. So there's no question in the near term that if we had more supply, we had demand for that product. And we're going to continue to work on those elements as we go through the year. But I think both on the demand side and the supply side, I'm very pleased with how the ramp is going.

TOSHIYA HARI: Thank you for all the details. And then as my follow up, I was hoping you could speak to your data center GPU roadmap beyond the MI300. The other concern that we hear is your nearest competitor has been pretty transparent with their roadmap, and that extends into '25 and, oftentimes, '26. So maybe this isn't the right venue for you to give too much. But beyond the MI300, how should we think about your roadmap and your ability to compete in data center? Thank you.

LISA SU: Yeah, sure. So, look, Toshiya, when we start with the roadmap, I mean, we always think about it as a multiyear, multigenerational roadmap. So we have the follow-ons to MI300 as well as the next generations well in development.

I think what is true is we're getting much closer to our top AI customers. They're actually giving us significant feedback on the roadmap and what we need to meet their needs. Our chiplet architectural is that-- architecture is actually very flexible. And so that allows us to actually make changes to the roadmap as necessary. So we're very confident in our ability to continue to be very competitive. Frankly, I think we're going to get more competitive.

Right now, I think MI300X is in a sweet spot for inference. Very, very strong inference performance. I see, as we bring in additional products later this year into 2025, that that will continue to be a strong spot for us. And then we're also enhancing our trading performance and our software roadmap to go along with it. So more details to come in the coming months. But we have a strong roadmap that goes through the next couple of years, and it is informed by just a lot of learning in working with our top customers.

TOSHIYA HARI: Appreciate it. Thank you.

LISA SU: Sure.

PAUL: Our next question is from Ross Seymore with Deutsche Bank. Please proceed with your question.

ROSS SEYMORE: Hey. Thanks for letting me ask a question. The non-AI side of the data center business, it sounds like the enterprise side has some good traction, even though the sequential drop happened seasonally, Lisa. But I was just wondering, what's implied in your second quarter guidance for the data center CPU side of things? And generally speaking, how are you seeing that whole GPU versus CPU crowding out dynamic playing out for the rest of 2024?

LISA SU: Yeah. Sure, Ross. Thanks for the question. I think our EPYC business has actually performed pretty well. The market is a bit mixed. I think some of the cloud guys are still working through their optimizations. I think it's different by customer.

We did see here in the first quarter, actually, some very nice early signs in the enterprise space, large customers starting refresh programs. The value proposition of Genoa is very, very strong, and we're seeing that pull through across the enterprise.

In the second quarter, we expect overall data center to be up strong double digits. And then within that, we expect server to be up as well. And as we go into the second half of the year, I think there are a couple of drivers for us. We do expect some improvement in the overall market conditions for the server business, but we also have our Turin launch in the second half of the year that will also, we believe, extend our leadership position within the server market. So overall, I think the business is performing well. And we believe that we're continuing to be very well-positioned to gain share throughout the year.

ROSS SEYMORE: Thanks for that. And I guess as my follow up, just switching over to the client side, I noted you guided it up sequentially. Any sort of magnitudes around that for the second quarter? And perhaps, more importantly, when you talk about the whole AI PC side of things, do you believe that's more of a unit driver for you, an ASP driver, or will it be both?

LISA SU: Yeah. So again, I think we're pretty excited about the AI PC, both opportunity in, let's call it, the near term and even more so in the medium term. I think the client business is performing well, both on the channel and on the MNC side. We expect client to be up sequentially in the second quarter.

And as we go into the second half of the year, to your question about units versus ASPs, I think we expect some increase in units as well as ASPs. The AI PC products, when we look at the Strix products, they're really well-suited for the premium segments of the market. And I think that's where you're going to see some of the AI PC content strongest in the beginning. And then as we go into 2025, you would see it more across the rest of the portfolio.

ROSS SEYMORE: Thank you.

LISA SU: Thanks, Ross.

PAUL: Thank you. Our next question is from Matt Ramsay with TD Cowen. Please proceed with your question.

MATT RAMSAY: Yes. Thank you very much. Good afternoon, everybody. Lisa, I have sort of a longer term question and then a shorter term follow-up one.

I guess one of the questions that I've been getting from folks a lot is, obviously, your primary competitor has announced, I guess, a multiyear roadmap. And we continue to hear more and more from other folks about internal ASIC programs that some of your primary customers, whether they be for inference, or training, or both-- I guess it would be really helpful if you could talk to us about how your conversations go with those customers, how committed they are to your long-term roadmap, multigeneration as you described it, how they juxtapose doing investments over their internal silicon versus using a merchant supplier like yourselves, and maybe what advantages the experience across a large footprint of customers can give your company that those guys doing internal ASICs might not get. Thanks.

LISA SU:

Yeah. Sure, Matt, Thanks for the question. So look, I think one of the things that we see and we've said is that the TAM for AI compute is growing extremely quickly. And we see that continuing to be the case in all conversations.

We had highlighted a TAM of, let's call it, \$400 billion in 2027. I think some people thought that was aggressive at the time. But the overall AI compute needs, as we talk to customers, is very, very strong. And you've seen that in some of the announcements even recently with some of the largest cloud guys.

From my view, there are several aspects of it. First of all, we have great relationships with all of the top AI companies. And the idea there is we want to innovate together. When you look at these large language models, and everything that you need for training and inferencing there-- although there will be many solutions, I don't think there's just one solution that will fit all. The GPU is still the preferred architecture, especially as the algorithms and the models are continuing to evolve over time, and that favors our architecture and also our ability to really optimize CPU with GPUs.

So from my standpoint, I think we're very happy with the partnerships that we have. I think this is a huge opportunity for all of us to really innovate together. And we see that there's a very strong commitment to working together over multiple years going forward. And that's, I think, a testament to some of the work that we've done in the past, and that very much is what happened with the EPYC roadmap as well.

MATT RAMSAY: No, appreciate that, Lisa. As my follow up, a little bit shorter term, and I guess having followed the company super closely for a long time, I think there's always been noise in the system from whether the stock price is \$2 a share or \$200. There's been always consistent noise one way or the other. But the last month and a half has been extreme in that sense.

And so I wanted to just-- I've got random reports in my inbox about changes in demand from some of your MI300 customers, or planned demand for consuming your product. I think you answered earlier about the supply situation and how you're working with your partners there, but has there been any change from the customers that you're in ramp with now or that you soon will be of what their intention is for demand? Or, in fact, has that maybe strengthened rather than gone down in recent periods? Because I keep getting questions about it. Thanks.

LISA SU:

Sure, Matt. Look, I think I might have said it earlier, but maybe I'll repeat it again. I think the demand side is actually really strong. And what we see with our customers and what we are tracking very closely is customers moving from, let's call it, initial POCs to pilots, to full-scale production, to deployment across multiple workloads. And we're moving through that sequence very well. I feel very good about the deployments and ramps that we have ongoing right now. And I also feel very good about new customers who are earlier on in that process.

So from a demand standpoint, we continue to build backlog as well as build engagements going forward. And similarly, on the supply standpoint, we're continuing to build supply momentum. But from a speed of ramp standpoint, I'm actually really pleased with the progress.

MATT RAMSAY: All right. Thank you very much.

PAUL: Thank you. Our next question is from Aaron Rakers with Wells Fargo. Please proceed with your question.

AARON RAKERS: Yeah. Thanks for taking the question. And I apologize if I missed this earlier. But I know last quarter you talked about having a-- securing enough capacity to support significant upsides of the ramp of the MI300. I know that you upped your guide now to \$4 billion. I'm curious to how you would characterize the supply relative to that context offered last quarter as we think about that new kind of target of \$4. Would you characterize it as still having supply capacity upside potential? Thank you.

LISA SU: Yes, Aaron. So we've said before that our goal is to ensure that we have supply that exceeds the current guidance, and that is true. So as we've upped our guidance from \$3 and 1/2 to \$4 billion, we have supply visibility significantly beyond that.

AARON RAKERS: Yep. OK. Thank you. And as a quick follow up, going back to an earlier question on server demand, more traditional server, as you see the ramp of maybe share opportunities in more traditional enterprise, I'm curious to how you would characterize the growth that you expect to see in more traditional server CPU market as we move through '24 or even longer term. How do you characterize that growth trend?

LISA SU: Yeah, I think, Aaron, what I would say is there are-- the need for refresh of, let's call it, older equipment is certainly there. So we see a refresh cycle coming. We also see AI head nodes as another place where we see growth in, let's call it, the more traditional market.

Our sweet spot is really in the highest performance, high core count, energy efficiency space, and that is playing out well. And we're also-- we have traditionally been very strong in, let's call it, cloud-first party workloads, and that is now extending to cloud third-party workloads, where we see enterprises who are, let's call it, in more of a hybrid environment, adopting AMD both in the cloud and on prem. So I think overall, we see it as a continued good progression for us with the server business going through 2024 and beyond.

AARON RAKERS: Yep. Thank you.

LISA SU: Thanks.

PAUL: Thank you. Our next question is from Vivek Arya with Bank of America Securities. Please proceed with your question.

VIVEK ARYA: Thanks for taking my question. Lisa, I just wanted to go back to the supply question and the \$4 billion outlook for this year. I think at some point, there was a suggestion that the \$4 billion number, that there are still supply constraints. But I think at a different point, you said that you have supply visibility significantly beyond that.

Given that we are almost at the middle of the year, I would have thought that you would have much better visibility about the back half. So is the \$4 billion number a supply constraint number, or is it a demand constraint number? Or alternatively, if you could give us some sense of what the exit rate of your GPU sales could be. I think on the last call, a billion and a half was suggested. Could it be a lot more than that in terms of your exit rate of MI for this year?

LISA SU: Yeah. Vivek, let me try to make sure that we answer this question clearly. From a full-year standpoint, our \$4 billion number is not supply-capped. I'm sorry, yes, it's not supply capped. We do have supply capability above that. It is more back-half weighted. So if you're looking at the near term, the near term, I would say, for example, in the second quarter, we do have more demand than we have supply right now, and we're continuing to work on pulling in some of that supply.

By the way, I think this is an overall industry issue. This is not at all related to AMD. I think, overall, AI demand has exceeded anyone's expectations in 2024. So you've heard it from the memory guys. You've heard it from the foundry guys. We're all ramping capacity as we go through the year.

And as it relates to visibility, we do have good visibility into what's happening. As I said, we have great customer engagements that are going forward. My goal is to make sure that we pass all of the milestones as we're ramping products and as we pass those milestones. We put that into the overall full-year guidance for AI. But in terms of how customer progression things are going, they're actually going quite well. And we continue to bring new customers on, and we continue to expand workloads with our current customers. And so hopefully, that clarifies the question, Vivek.

VIVEK ARYA: OK. Thank you, Lisa. Maybe one not on MI, but maybe on the embedded business. I think you sound a bit more measured about Q2 and the second half rebound, which is similar to what we have heard from a lot of the auto industrial peers. But where are you in the inventory clearing cycle?

And if embedded has a somewhat more measured rebound in the back half, what implication does that have on gross margin expansion? Can we continue to expect, I don't know, a hundred basis points a quarter in terms of gross margin expansion because of the data center mix, or just any puts and takes of embedded and then what it means for gross margins in the back half? Thank you.

JEAN HU: Hi, Vivek. Thank you for the question. I think the embedded business declined a little bit more than expected really due to the weaker demand in some of the markets. Very specifically, communication has been weak. In some pockets of industrial and automotive, as you mentioned, it's actually quite consistent with the peers.

Second half, we do think the first half is the bottom of embedded business. And we'll start to see gradual recovery in the second half. And going back to your gross margin question, is, when you look at our gross margin expansion in both Q1 and guided Q2, the primary driver is the strong performance on the data center side. The data center will continue to ramp in second half. I think that will continue to be the major driver of gross margin expansion in second half. Of course, if embedded is doing better, we'll have a more tailwind in second half.

VIVEK ARYA: Thank you.

PAUL: Thank you. Our next question is from Timothy Arcuri with UBS. Please proceed with your question.

TIMOTHY ARCURI: Thanks very much. I also wanted to ask about your data center GPU roadmap. The customers that we talk to say that they're engaged not just because of MI300, but really because of what's coming. And it seems like there's a big demand shift to rack-scale systems that try to optimize performance per square foot given some of the data center and power constraints. So can you just talk about how important systems are going to be in your roadmap? And do you have all the pieces you need as the market shifts to rack-scale systems?

LISA SU: Yeah. Sure, Timothy. Thanks for the question. For sure.

Look, our customers are engaged in the multigenerational conversation. So we are definitely going out over the next couple of years. And as it relates to the overall system integration, it is quite important. It is something that we're working very closely with our customers and partners on. That's a significant investment in networking, working with a number of networking partners as well to make sure that the scale-out capability is there.

And to your question of, do we have the pieces? We do absolutely have the pieces. I think the work that we've always done with our Infinity Fabric as well as with our Pensando acquisition, that's brought in a lot of networking expertise. And then we're working across the networking ecosystem with key partners like Broadcom, and Cisco, and Arista, who were with us at our AI data center event in December.

So our work right now, in future generations, is not just specifying a GPU. It is specifying, let's call it, full-system reference designs. And, yeah, that's something that will be quite important going forward.

TIMOTHY ARCURI: Thanks a lot. And then just as a quick follow up, I know this year, it looks like it's going to be pretty back-half loaded in your server CPU business, just like it was last year. And I know you held our hands about this time last year on what the full year could look like and how backend loaded it could be. So I kind of wonder, could you give us some milestones in terms of how much server CPU could grow this year or how backend loaded it could be? Is it like up 30% this year for your server CPU business year over year? Is that a reasonable bogey? I just wonder if you can give us any guidance on that piece of the business. Thanks.

LISA SU: Yeah. I mean, Tim, I think the best way to say it is our data center segment is on a very, very strong ramp as we go through the back half of the year. Server CPUs, certainly. Data center GPUs, for sure. So I don't know that we're going to get into specifics. But I could say, in general, you should expect overall the segment level to be very strong double digits.

TIMOTHY ARCURI: Thank you, Lisa.

LISA SU: Thank you.

PAUL: Thank you. Our next question is from Joe Moore with Morgan Stanley. Please proceed with your question.

JOE MOORE: Great. Thank you. I wonder if you could address the profitability of MI300. I know you said a couple quarters ago that it would eventually be above corporate average, but it would take you a few quarters to get there. Can you talk about where you are in that?

JEAN HU: Yeah. Thank you, Joe. Our team has done an incredible job to ramp MI300. As you probably know, it's a very complex product, and we are still at the first year of the ramp, both from EO, the testing time, and the process improvement. Those things are still ongoing. We do think, over time, the gross margin should be accretive to corporate average.

JOE MOORE: Great. Thank you. And then a separate follow up on the Turin transition on server. I know when you had transitioned into Genoa, you said it could take a little while, that there were significant platform shifts, and things like that. Turin seems to be much more ecosystem-compatible. How quickly do you think you might see that product ramp within their server portfolio?

LISA SU: Yeah, Joe, I think from what we see, I think Turin is the same platform. So that does make it an easier ramp. I do think that Genoa and Turin will coexist for some amount of time because customers are deciding when they're going to bring out their new platforms. We expect Turin to give us access to a broader set of workloads. So our SAM actually expands with Turin, both in enterprise and cloud. And from our experience, I think you'll see a faster transition than, for example, when we went from Milan to Genoa.

JOE MOORE: Great. Thank you.

PAUL: Thank you. Our next question is from Stacy Rasgon with Bernstein Research. Please proceed with your question.

STACY RASGON: Hi, guys. Thanks for taking my questions. For my first one, I wanted to address the MI300 ramp into Q2. So you said you've done \$1 billion, give or take, in cumulative sales, which puts it at maybe, I don't know, maybe \$600 million in Q1.

Your guiding total revenue is up about \$225 million into Q2. But you've got client up. You've got traditional data center up. You've got embedded flat. Gaming is going to be down, but I'd hazard a guess that the client and traditional data center offset it, if not more. Does the MI300 ramp into Q2? Is it more or less than the total corporate ramp that you've got built into guidance right now that you're expecting?

JEAN HU: Hi, Stacy. Thanks for the question. You always ask math question. So I think, in general, it is more. The data center GPU ramp, it will be more than the overall company's \$200-some million ramp.

STACY RASGON: OK. So that means gaming must be down a lot if client and traditional data center are also both up. OK.

JEAN HU: Yeah, you're right. Gaming is down. Similar zip code, like Q1.

STACY RASGON: Got it. Got it. That's helpful.

JEAN HU: And so maybe-- yeah, maybe let me give you some color about the gaming business. If you look at the gaming, the demand has been quite weak. That's quite very well-known, and also their inventory levels. So based on the visibility we have, the first half, both Q1 and Q2, we guided down sequentially more than 30%.

We actually think the second half will be lower than first half. That's basically how we're looking at this year for the gaming business. And at the same time, gaming gross margin is lower than our company average. So overall, will help the mix on the gross margin side. That's just some color on the gaming side. But you're right, Q2 gaming is down a lot.

**STACY
RASGON:**

Got it. That's helpful. Thank you. For my second question, I wanted to look at the near-term data center profitability.

So operating profit was down 19% sequentially on 2% revenue growth. Is that just the margins of the GPUs filtering in relative to the CPUs? And I know you said GPUs would eventually be above corporate average. Are they below the CPU average in-- I mean, they clearly are, I guess, in the near term, but are they going to stay that way?

JEAN HU:

Yeah. I think you're right, the GPU gross margin right now is below the data center gross margin level. I think that there are two reasons. Actually, the major reason is that we actually increased the investment quite significantly to, as Lisa mentioned, to expand and accelerating our roadmap in the AI side. That's one of the major drivers of the operating income coming down slightly.

On the gross margin side, going back to your question, is, we said in the past, and we continue to believe the case, is the data center GPU gross margin, over time, will be accretive to corporate average, but it will take a while to get to the silver level of gross margin.

**STACY
RASGON:**

Got it. That's helpful. Thank you.

PAUL:

Thank you. Our next question is from Harlan Sur with JP Morgan. Please proceed with your question.

HARLAN SUR:

Good afternoon. Thanks for taking my question. On your data center GPU segment and the faster time-to-production shipments, given you just upped your full-year GPU outlook, how much of it is faster bring-up of your customers' frameworks driven by your latest ROCm software platform and maybe stronger collaboration with your customers' engineers just to get them to call faster, and how much of it is just a more aggressive buildout plan by customers versus their prior expectations given what appears to be pretty strong urgency for them to move forward with their important AI initiatives?

LISA SU:

Yeah. Harlan, thank you for the question. What it really is, is both us and our customers feeling confident in broadening the ramp, because if you think about it, first of all, the ROCm stack has done really well. And the work that we're doing is hand in hand with our customers to optimize their key models. And it was important to get verification and validation that everything would run well. And we've now passed some important milestones in that area.

And then I think the other thing is, as you said, there's a huge demand for more AI compute. And so our ability to participate in that and help customers get that up and running is great. So I think, overall, as we look at it, this ramp has been very, very aggressive if you think about where we were just a quarter ago. Each of these are pretty complex bring-ups. And I'm very happy with how they've gone. And by the way, we're only sitting here in April. So there's still a lot of 2024 to go. And there's great customer momentum in the process.

HARLAN SUR: Yep. Yeah, absolutely. Just going back, just rewinding back to the March quarter, so similar to the PC client business, which declined at the low end of the seasonal range, if I make certain assumptions around your data center GPU business, I x that out of data center, it looks like your server CPU business was also down at the lower end of the seasonal range. By my math, it was down like 5%, 6% sequentially. Is that right? And that's less than half the decline of your competitor.

And if so, what drove the less than seasonal declines? I assume some of it was share gains. Sounds like enterprise was also better. Looks like you guys did drive a little bit more cloud instance adoption. But anything else that drove the slightly better seasonal pattern in March for data center server?

JEAN HU: Yeah. Yeah, Harlan. This is Jean. I think the server business has been performing really well. Year over year, it actually increased very strong double digit. I think sequentially, it is more seasonal. But we feel pretty good about continue gaining share there.

LISA SU: And if I just add--

HARLAN SUR: Thank you.

LISA SU: --Harlan, to your question, we did see strength in enterprise in the first quarter. And I think that has-- that offset perhaps some of the normal seasonality.

HARLAN SUR: Yeah.

MITCH HAWS: Paul, we have time for--

HARLAN SUR: Thank you.

MITCH HAWS: --two more questions.

PAUL: Thank you. Our next question is from Tom O'Malley with Barclays. Please proceed with your question.

TOM O'MALLEY: Hey. Thanks for taking my question. I just wanted to ask on the competitive environment. Obviously, on the CPU side, you had a competitor talk about launching a high core count product in the coming quarter, kind of ramping now and more so into Q3. You've seen really good pricing tailwinds as a function of the higher core count. But can you talk about what you're seeing in that market? Do you think that there's any risk for more aggressive pricing, which would impact your ASP ramp for the rest of the year?

LISA SU: Yeah. When we look at our server CPU ASPs, they're actually very stable. I think we, again, we tend to be indexed towards the higher core counts. Overall, I would say the pricing environment is stable. This is about TCO for the customer environments, and our performance, and our performance per watt, our leadership, and that usually translates into TCO advantage for our customers.

TOM O'MALLEY: Helpful. And then just a broader question to follow up here. So I think you got asked earlier about the importance of systems. But on your end, how important is the open ethernet consortium to you being able to move more into systems?

I know that today you obviously have some internal assets, and then you can partner with others. But is there a way that you can be competitive before there is an industry standard on the ethernet side? And can you talk about when you think the timing of that kind of consortium comes to market and enables you to maybe accelerate that roadmap? Thanks a lot.

LISA SU: Yeah. I think it's very important to say, we are very supportive of the open ecosystem. We're very supportive of the Ultra Ethernet Consortium. But I don't believe that that is a limiter to our ability to build large-scale systems. I think ethernet is something that many in the industry feel will be the long-term answer for networking in these systems. And we have a lot of work that we're doing internally as well as with our customers and partners to enable that.

MITCH HAWS: Paul?

PAUL: Thank you.

MITCH HAWS: We're ready for our final question.

PAUL: Our last question is from Harsh Kumar with Piper Sandler. Please proceed with your question.

HARSH KUMAR: Yeah. Hey. Hey. Thank you for letting me ask a question. Lisa, I had two. One is for you, and one perhaps for Jean.

So we recently hosted a very large custom GPU company for a call, and they talked about mega data centers coming up in the near to mid-term, talking about nodes potentially in the 100,000-plus range, maybe up to a million.

So as we look out at these kinds of data centers, from an architectural standpoint, is that a situation where winner takes all, where if somebody gets in, they get all the sockets? Or will there be lines where your chip perhaps or your board can be placed right next to somebody else's board, maybe on a separate line? Just help us understand how something like that would play out, if there's a chance for more than one competitor to play in such a large data center.

LISA SU: Yeah. So I'll talk maybe a little bit more at the strategic level. I think as we look at how AI shapes up over the next few years, there are customers who would be looking at very large training environments, and perhaps that's what you're talking about.

I think our view of that is, number one, we view that as a very attractive area for AMD. It's an area where we believe we have the technology to be very competitive there. And I think the desire would be to have optionality in terms of how you build those out.

So obviously, a lot has to happen between here and there. But I think your overarching question of, is it winner takes all? I don't think so. That being the case, we believe that AMD is very well-positioned to play in those, let's call it, very large-scale systems.

HARSH KUMAR: Wonderful. Thank you. And then maybe a quick one for Jean. So Jean, I put everything into the model that you talked about for June. I get about more or less a \$400 million rise in the June quarter over March. You mentioned that both MI300 and EPYC will grow. Curious if you could help us think about the relative sizing of those two segments within the quarter. The point I'm trying to make is I'm getting roughly about a \$900 million number for MI300 for June. Am I in the ballpark or am I way off here?

JEAN HU: Harsh, we're not going to guide a specific segment below the segment revenue. [CHUCKLES] I think the most important thing is we did say data center is going to grow double digit sequentially. I'll leave it over there. We, some segments there are a lot of details segment. There are a lot of details.

HARSH KUMAR: Fair enough. I had to try. Thank you, guys. Thank you so much.

JEAN HU: Yeah, yeah, yeah. Thank you.

PAUL: Thank you. There are no further questions at this time. I'd like to hand the floor back over to management for any closing comments.

MITCH HAWS: Great. That concludes today's call. Thanks to all of you for joining us today.

JEAN HU: OK.

PAUL: This concludes today's conference. You may disconnect your lines at this time. Thank you for your participation.