

JOSEPH MOORE: Welcome back, everybody. I'm Joe Moore, Morgan Stanley Semiconductor Research. I'll read this disclosure real quickly. For important disclosures, please see the Morgan Stanley Research Disclosure website at www.morganstanley.com/researchdisclosures. If you have any questions, please reach out to your Morgan Stanley sales rep.

So, with that, very happy to have with us today from Advanced Micro Devices, Jean Hu, EVP, CFO, and Treasurer as well as Matt Ramsay, who newly joined the company. I don't know your title, Head of IR and other things, right?

MATT RAMSAY: That's a good one. I'll take that, Joe.

JOSEPH MOORE: Head of IR and other things, yeah, sure. Anyway, so welcome you guys. Thank you so much for joining us. I think we got to start on AI. I think it's a rule. But you guys had a great year 1. You did over \$5 billion of revenue for a product from a standing start that people got for silicon. And you did \$5 billion in a year, tremendous. People want more always, but a very good year 1.

Can you talk about what was required to do that and how much visibility that gives you? It doesn't seem like people buy a point product, they're buying a roadmap. So obviously, that \$5 billion speaks to a certain amount of confidence. Can you talk about that year one experience?

JEAN HU: Yeah, absolutely. First, thank you for having us. This is a great conference, and thank you for the question. 2024, it was truly a transformative year for AMD. As you mentioned, we actually ramped up MI300 really quickly and exceeded \$5 billion. Both the MI300 and our ROCm software right now are powering some of the most complicated AI models at Microsoft, at Meta, at scale. So it's a lot of progress we made, and the team executed flawlessly, actually, to ramp the production of MI300.

At the same time, we also made tremendous progress in our overall hardware roadmap and the software stack. If you think about the hardware side, not only we introduced the MI325, we also put in this year, MI350. And next year, we're going to be on track for MI400.

On the software side, we made a lot of progress with ROCm to really support all different applications and the workload. We also acquired ZT Systems to build our system expertise to build the rack-level and class-level designs.

So, overall, the progress made and our engagement with the customers also broadened quite significantly. Not only we add new customers and with existing customers, the engagement has always been about multi-generational. It's not just about generating MI300 revenue. It's about the discussions for MI350 and MI400, how we help our customers to build clusters to provide better TCOs.

So, the way to think about it is MI300 is the very, very beginning of our journey. And that each generation, if you look at our product roadmap, we get better and stronger, more competitive. And just like how we build our server side of the business, over time, we want to be a strong player, gaining market share in this market. When you look at the market opportunities, we do expect to be like \$0.5 trillion, \$500 billion large market opportunities.

And based on the execution we have so far and how we are very well positioned, we do believe we can have a growth trajectory to tens of billions of dollars in annual revenue in this market.

**JOSEPH
MOORE:**

I do want to dig into some of the roadmap stuff, but maybe first, conceptually, if we could talk about AMD versus ASICs, principally, is my focus. I think NVIDIA is pretty good at this, we all know that. But everybody's looking for an alternative. Nobody wants to be sole-sourced. Everybody's looking for ways that they can generate ROI that NVIDIA can't give them.

I think when we were at this conference 12 months ago, AMD was clearly the answer to that question in people's mind. Now people are talking a lot more about custom silicon and ASICs. And I guess at the same time, you guys have a lot of visibility with customers. You have a lot of customers who are still making investments in AMD. Just how do I think about that trade-off between the visibility that the ASIC vendors have versus the visibility that AMD has?

JEAN HU:

Yeah, this topic is definitely top of mind for many investors. Maybe let us share our perspective. I'll start, Matt, you can add. AMD is a high-performance computing company. If you think about our strategy and what Lisa has done is to build the platform, not only we have a CPU, GPU, and also FPGA, we also do custom silicon, actually. So, our view has always been different workload, different application require different compute engine to get the best economics.

And when you look at the AI models, it has been evolving rapidly. And the model innovation literally accelerating every week. At the same time, you can see the AI scaling continues, right? You have a pre-training, post-training to test-time compute. All those things require different compute, different compute engine, and different compute requirement.

GPU is really built for support a wide range of models, right? It's flexible, you can program it, and then you can tune it for the performance over different model changes. So from a TCO perspective, that is how you really think about it is the total cost of ownership. One customer are doing either training and inferencing. It's about the cost per token and the cost per flop. We are working with our customers closely to really build those larger infrastructures. That's how they are thinking about it.

On the other side, ASIC can be efficient if the workload is very specific, very stable, and then you can really design for the specific models. And at the same time, there's very large-scale deployment. And ASIC also takes time too. That's why you probably are saying, Oh, there's 18 months to 24 months' time ASIC has visibility. It's actually very similar for us.

The engagement with the customers, when you think about those 1-gigawatt data center you need to build, the lead time is really data center space, power, and all those things. We have to work with our customers closely to design the overall infrastructure there.

So to a certain degree, what AMD is focusing on is we think this is the vast majority of the market in the longer term, especially when you talk about the \$500 billion market opportunity. That's where we see we can get the highest return on the investment we are making. Matt, anything?

MATT RAMSAY: Yeah. Joe, I think I would also add that it's pretty easy to think about one way to generate TCO at the data center scale is to have an algorithm calm down, design specific silicon for that algorithm in an ASIC, and have lower cost hardware up front. Like that's a pretty obvious way to try to generate TCO, but less dollars in up upfront for the same computing.

Another way to generate TCO is to build programmable, GPU-led infrastructure that can rely on the industry's innovations and software over time to drive better TCO and better ROI of the infrastructure that you've already put in the ground because it's programmable.

And I think over the last month or so, we've sort of all witnessed the market's reaction to DeepSeek. But to us, DeepSeek got a lot of attention because it was in China and a couple of things that they claimed on cost.

But it's a pretty natural thing for an industry to start as the installed base of hardware grows to start doing really rapid innovation in software to get better TCO of an infrastructure that's already in place. And if your infrastructure is programmable, you can benefit from that innovation of the software stack of the industry over a long period of time and over the depreciable life of the infrastructure you put in the ground.

And I think that's what gives us the conviction that programmable infrastructure is the way to go for the majority of the TAM. There are certain applications that ASICs are very well suited for, and some of the folks in the market talk about those a lot. But I think over the breadth of workloads and over the fullness of time of software innovation, I think there's a lot to be said for programmable infrastructure and that's where our customers are pulling us and that's where we're pushing to bring increased competition and capabilities over time.

**JOSEPH
MOORE:**

And can you talk about the role of the price of the hardware? Because, I think, you just mentioned you can do something cheaper purpose-built. But if you're competing for that broader range of programmable workloads, if this were as easy as doing \$5,000 cards, a lot of people would be successful at it. You and NVIDIA both have expensive cards even within your stacks, it seems like a lot of the business gravitates to the highest performant part. So can you just talk about the role of price in all of this?

JEAN HU:

I think to compare the price of data center GPU with the ASIC is probably overly simplified. It's probably more nuanced when you think about how customers think about the TCO, right? As Matt mentioned earlier, for the TCO, they measure cost per token, cost per flop. And when they think about it, they think about overall how many years they are going to use the infrastructure and how many different model, different workload they can run on that infrastructure for those years.

On the ASIC side, if you want to talk about TCO for economics for any ASIC investment, there's very large upfront R&D cost. That's not part of the ASP, but it's part of the investment. And also, if you take time to invest, you also need to consider the risk of obsolete because the model tend to be very fixed. It could be like 1 or 2 years, so you have to turn it around.

So when you compare at that point, it's really the upfront ASP is not a direct comparison. I think when we look at our customers, how they calculate the total cost of ownership, they do consider all the different factors there. And fundamentally, it is about one is programmable. You can tune it. You can get the best performance over time versus the one is for very specific applications. You get a better efficiency if you know exactly the scale and the application you want to use.

JOSEPH MOORE: Great. And I guess before I get to the roadmap, I mean, the ubiquity of AMD, you can be at several different cloud vendors. You can make investments into the software ecosystem that can benefit all of your customers, can benefit a wide range of people that don't want to be locked into a single cloud vendor for their workloads. AMD is going to have that appeal. How important is that and how does that inform the acquisitions that you're doing around software acquisitions in ZT?

JEAN HU: Yeah. Thank you for the question. I think that's a very important part of AMD's strategy. If you look at the ROCm software stack, it's open source. We actually not only-- all the frameworks, the PyTorch, Trident, JAX, and all different things support completely. We also work the overall ecosystem. And you can run it in Microsoft. You can run it in other CSPs to really help the customers, no matter which framework you use to write your model.

That is part of our strategy, actually, is to help customers to really for ease over-deployment of AMD models. And that you should expect us to continue to drive really aggressively to have that ecosystem build-up. The software acquisitions we did is to add more capabilities and broaden our customer engagement so we can support more customers.

JOSEPH MOORE: I mean, the ZT in particular seems like a lot of work, right? I mean, you're buying an ODM, divesting the hardware portion, keeping the engineering resources. What is that getting you that in return for the work that you're doing there?

JEAN HU: I'll start. Matt can add. The ZT acquisition actually is a very important step in our building our capabilities. It's about adding design capabilities for rack-level, system-level to help customers not only just build small clusters, but very large clusters for both training and inference.

It is a lot of work, but when you think about it, it's not only we're going to keep the very large design team that will help us with the capabilities. More importantly, we'll have a strategic partnership with the manufacturing operation, which we are seeking strategic partners. That will help us to speed time to market, right? We'll work closely, help our customers to build the different clusters systems. That's really very critical in today's market. Matt?

MATT RAMSAY: The only thing I would add, and Jean, you covered it well, but the only thing I would add is as we bring the ZT design team into the company and it has significant influence on our MI400 generation product in 2026, I think we've learned a lot of watching what's happened in the industry over the last 12 months, 15 months in terms of putting together rack-scale systems.

And I think what you'll expect to see from AMD is a much less prescriptive approach to system design partnership with our OEM and ODM partners from a reference design perspective. And every large AI company, every hyperscale company, their data center infrastructure is not ubiquitous. It's not the same. What one customer might want for their data center footprint might be very different from what another customer wants.

And so being able to have rack-scale reference design that's not prescriptive, but has some flexibility and customization for individual data center builds by customers is what we're going to be bringing to the table as we move forward with the ZT team integrating into the current team that we have inside the company.

JOSEPH MOORE: Good stuff. So maybe we could talk about the roadmap a little bit starting with MI350 this year. How game-changing is that product? Are you going to get new customers, replacement customers, or obviously people are going to migrate to 350, but how important is it from generating new demand?

JEAN HU: You want to start?

MATT RAMSAY: Sure. It's an important product for the company. There's some significant new capabilities in networking, in memory capacity and addressability across a cabinet around data-type utilizations down to FP6 and FP4. A lot of work that's done in the ROCm software stack, that will be introduced alongside MI350 to take higher-level models and map them onto the underlying topology of the hardware.

I think it expands our performance levels significantly when it comes to large model inference. I think Lisa's been pretty public about 35 times or up to 35 times performance gains for inference and expands the aperture of the training capability of the instinct roadmap to tens of thousands of units and training clusters. And then it gives us something significant to build on in the 400 generation for frontier-level training models.

And so, we were really, really excited to be able to pull that product in to be able to launch in the midyear. I think that was a few months earlier than most of this audience might have expected. And I think we're anxious for that to get going. The customers are anxious for it to get going.

And yeah, we have talked about bringing in additional sort of lighthouse accounts into the Instinct portfolio as that product launches. And I guess more to come in the middle of the year as we officially launch the program.

JOSEPH MOORE: And then with MI400, you alluded to a lot of the rack and cluster-level benefits that you bring. Is the frame of reference there MI400 versus Rubin? And obviously, hard to talk about future products from different companies, but is this the point where you can take a much bigger role in training things like that, your confidence level in that based on what of your competitors' roadmap?

JEAN HU: Yeah, absolutely. I think the way to think about it is MI350 is more compatible with the Blackwell, and MI400 is really to compete with Rubin. So, each generation, we are doing much better. And once we get to MI400, we do feel, we have a very competitive product portfolio and support rack-level, system-level, and the cluster-level build-up. That is we have not shared a lot of details yet. But that's the plan is to really drive more competitive product roadmap there.

JOSEPH MOORE: And I do want to spend time on the other 80% of your business, but I do still have to ask it a couple of other AI questions. I guess tens of billions. That forecast that you have tens, more than one ten, but tens of billions of revenue potential around these products. Presumably, that is something that you're saying based on conversations you're having with customers about the opportunity that you have. What has to happen for you to achieve those kinds of numbers?

JEAN HU: Yeah. First, we have to execute on our roadmap, right, not only hardware roadmap, software and ZT Systems integration, to make sure we continue to drive all the execution flawlessly. Secondly, it is a very strong customer engagement. When we engage with our customers, it's not about just generating revenue currently. It's always about roadmap discussions, the feedback from customers, how we can provide the best TCO for customers.

And, as we all know, the build cycle is quite long for those larger clusters and the data center, you actually need to figure out the power, the space, and everything else. So those are the important things we need to work with our customers, partner with them together. ZT Systems is a very important part of this equation. Will help us to speed up time to market to be able to support our customers.

JOSEPH MOORE: Great. And then the last AI question for me, the export controls, if we do get, we have rules that are supposed to go in place mid-May, we don't know if those will be the final say. Just how is AMD positioned to deal with potential government export controls?

JEAN HU: Yeah, we're monitoring very carefully. It definitely has been a topic that everybody is really focused on. So, for us, that's the number one objective. We need to meet the export control criteria. We do think there are a lot of opportunities with sovereign AI in China. There's a large market, too. We do think that those are the opportunities that we want to address going forward. But, of course, complying with the export control.

JOSEPH MOORE: Great. So, I'd like to pivot and ask you about some of the x86 businesses, starting with servers. You guys have done well in the server market, in a market that's been tough, right? We've seen a lot of this AI investment has caused people to actually go as far as lengthen the depreciable lives of their servers. So, when you guys have the highest market share with the people who are doing that, it seems like it's a headwind, and yet you've grown pretty nicely.

What's your visibility to that? I mean, you've talked about the staleness of the data center ecosystem server-wise. When do you think we might start seeing more refreshed business?

JEAN HU: Do you want to start?

MATT RAMSAY: Yeah, sure. I think there's a couple of things. When you-- there was a period of time where there's this mass CapEx pendulum swing towards AI, and we all sort of witnessed what that looked like in 2023 for the server market.

More recently, some of the limitations of folks that actually want to add AI hardware are data center space. And if you start to look back at-- there's a lot of compute infrastructure in the data center that's still 2 or 3 generations old, maybe 4 generations old of CPU, that upgrading those into our Turin platform, given either Turin Dense or Turin Classic, that can get you significantly better data center footprint usage, in addition to some of the work that we're doing on head node for GPU clusters.

I think we feel really good about where the server business is Dan's business. Dan McNamara, who runs that business internally for us, has really, really good product up and down the stack from Turin Dense, that's sort of a direct ARM competitor in a lot of instances that builds on top of the Bergamo platform that's been really successful to us, to the core count lead that the portfolio has across the board.

In addition to the fact that when you take all of this AI work that's being done, it actually generates a lot of need for classical computing alongside of it. So, you see that, the ability to refresh for space, and I think we gained 5 or 6 points of server share last year overall. And I think we're pretty confident about the server share gains to come in 2025 as we move forward and look forward to continuing the leadership position across the whole breadth of the roadmap as we move beyond the Turin generation.

**JOSEPH
MOORE:**

And in the enterprise side of the server business, meaning both enterprise OEMs but also cloud, that's servicing enterprise, you guys had a technology lead now for 6 or 7 years. Obviously, there's still an incumbency advantage for Intel in some cases, but are you able to continue to break through that? And if Intel narrows the gap, does that make it harder? I mean, how big does the lead have to be for the enterprise market to keep swinging towards AMD?

JEAN HU:

Yeah. We have been investing in enterprise go to market for the last couple of years. That has helped us to make significant progress in enterprise market. If you look at the last 6 consecutive quarters, we have been growing our enterprise business year-over-year. That is really because not only the TCO benefit, as Matt mentioned, is data center space power, those are the major constraints even for enterprise.

So, when we can provide the TCO to help them to save power and space, we do see if we show our customers our TCO performance, they will switch. There's actually very, very clear choice over there. It is just we need to get to different enterprise customers, have the go to market engine to help them to do that switch.

So, overall, we feel pretty good about the generation over generation, not only Genoa but Turin. We actually have more platform design wins with Turin because the workload application, we actually brought the support for all different workloads and application with our Turin platform. So, we do think that we can continue to gain market share in enterprise market.

MATT RAMSAY: Yeah. Joe, the only thing I would add to is that on the enterprise side, for the first time, I think, despite having, as you mentioned, product leadership for a number of generations, it's the shift of enterprise CIOs moved more towards what the heck am I going to do with AI, right?

And that's where the board pressure came from. That's where in the recent months, it's kind of swung back a bit. I mean, AMD has progressed not just as a technology leader, but now as the safe vendor of choice. As you look forward to plan your infrastructure over the next 1, 3, 5 years, what vendor do you want to really rely on for your infrastructure that has presence in all the clouds for overflow and a multi-cloud strategy, but also just continuation of roadmap execution and stabilization of roadmap. That's become-- it was really top of mind for what, 7, 8, 9-year period.

And then when AI came, it kind of went down the priority list for CIOs and it's quickly kind of popped back up both from a space and power constraints perspective, as Jean discussed, but also just from a continuity of execution of roadmap and supply. I think it's been-- the environment's changed a bit, and I think Dan's business is positioned well for enterprise over the next 18 months.

**JOSEPH
MOORE:**

Great. And then shifting to client, really impressive performance there also. I think you grew 58% year-on-year last quarter with a seasonal outlook. Now your competitor called out that they thought there might be a tariff pull forward in their numbers. So how do you feel about of market share in that spaces? Are those numbers entirely share driven and your visibility on that continuing to improve through the year?

JEAN HU:

Yeah. We are very pleased with our client business performance. It has been really primarily driven by strong product portfolio. So, if you look at not only desktop but notebook side, we have the best lineup of product portfolio. Our Ryzen 9000 desktop processors have been sold out. Like, really every channel, a lot of retail channels, you actually can see we get to 70% market share.

So sell-through has been really strong not only in Q4 but post Chinese New Year, we continue to see strong sell-through because the gaming performance, the advantage, the user experience we can provide to customers.

On the notebook side, the Ryzen AI 300 processor has been really successful. We have a 150 different platform design wins, almost double Intel's similar design platforms. So that is the one, not only we have the best CPU, best GPU, and the best NPU, the combination really helped the customers on the gaming experience side, on the productivity side, on the user experience, so that helps a lot.

And, more importantly, if you look on the OEM customer side, not only we have a strong relationship with Lenovo, HP, we actually added the Dell for the first time to be our strategic partner to introduce the overall commercial platforms. That will help us to continue to drive the sell-through to really continue to gain market share.

**JOSEPH
MOORE:**

Great. And so, I'll have one more segment question and then I'll open it to the audience. Embedded, it's been a challenging period, this is the Xilinx. FPGA business, it's been a challenging period for all broad-based market companies in the last year or so. Can you talk about the progress there and any visibility that you may have into growth in embedded?

JEAN HU:

You're right. It has a prolonged cycle for inventory normalization. We do see some early signs of really gradual recovery. I think the end market that we participate, when you think about aerospace and defense as well as emulation, those testing, those are actually quite steady and resilient. You see signs of improvement, industrial continue to be very mixed from a demand perspective.

Overall, we do expect this year is the year to recovery, probably slowly. But I would say one thing is during this kind of a down cycle, we actually get tremendous design wins because our team's focus and execution, if you just look at the design wins, we actually for 2024, we had \$14 billion design wins, which is like 25% increase year for year. That will help us in the longer term when the market really fully recovers. Right now, sell-through is improving slightly, we can see that. So that definitely is going to help us.

MATT RAMSAY: Joe, when I think about the overall AMD financial model, we started the conversation with AI, and we added more than \$5 billion in revenue for our AI programs in a year, which from a standing start, as you point out, is a heck of an achievement. But at the same time, there were cyclical challenges in a couple of our businesses, right? The embedded business declined significantly, and I think the industry and gaming as well--

JEAN HU:

Gaming declined very significantly.

MATT RAMSAY: And so, we had a couple of businesses that were headwinds that mask on the top line some really exciting progress in our core franchises. And I think we feel pretty confident that those headwinds are behind us. How quickly they turn into tailwinds, I'd rather for this audience under-promise and over-deliver with respect to turning those businesses around.

But we feel really good to Jean's point about design wins and the Xilinx business. And we're just sort of launched a new gaming GPU recently. So, there's some momentum that's starting to build. But at a bare minimum, I think you'll see the exciting franchises in client and server and data center GPU drive the P&L without the headwinds that have been there for the last 12, 15 months in the other franchises.

JOSEPH MOORE: And the gross margins in embedded, the questions come up in the context of your competitors saying that their gross margins are significantly lower than a few years ago. It seems like there was a discipline of having two public companies that you didn't chase markets that were converting to ASICs. It seems like your gross margins are still at the level when you acquired Xilinx. Can you talk to that?

JEAN HU: Yeah. We continue to be very disciplined. When you look at the FPGA franchise, we have not only from market share, it's number one. And we are also very focused on medium to high-end market with aerospace, defense, and also a lot of emulation testing management side. We continue to drive the team to be as disciplined as in the past. They have done a great job, right? It's all about margin improvement. So, for us, we continue to see very strong gross margin from FPGA business.

JOSEPH MOORE: Great. Do we have any questions from the audience? One in the front. Mic person. Maybe just say it, and I'll repeat it.

UNIDENTIFIED AUDIENCE MEMBER: Yeah, Joe, I just want to get a sense of a lot of investors are looking at ASP per core. Why is that the wrong measure? Do customers look at ASP per core?

JOSEPH MOORE: Yeah. Question is on ASP per core, is that the right way to look at the market?

JEAN HU: It's a good way to look at it. If you look at our server business over time, we have been increasing our core counts generation by generation. So, it's a good way to look at it. If you can keep your core price largely consistent or constant, you actually each generation, you actually can increase your overall ASP. Because you actually provide better performance for your customers. So, it is a good way to look at it, but it actually takes more time to track it for third-party analysts, right? So, I think in general, we do look at it that way.

JOSEPH MOORE: Any other questions from the audience? I guess then I'll close with how do you think about R&D dollars? When you think about you're characterizing AI as a \$500 billion opportunity, you have tens of billions of revenue potential. I know Lisa pretty well, she's not going to be happy with 5% share. She's going to drive for pretty big numbers. Your competitor is spending \$16 billion. So how do you think about the need to invest more to invest ahead of those revenue levels?

JEAN HU: Yeah, that's a great question. I think the way how AMD thinks about resource allocation is, given the very large growth opportunity, the first thing is we are leaning in investment. Investment not only on the R&D side but also on the acquisition side. We do have a very strong balance sheet, very much under-levered balance sheet, so we can leverage our balance sheet to invest, to do acquisitions on software side, like ZT Systems.

That being said, the company is very disciplined and overall focused on innovation. When you think about back 10 years ago or 12 years ago when Lisa and Mark Papermaster joined the company and Forrest, the whole team, they had a resource constraint even back then, right?

So, the innovation, like a chiplet design, AMD really lead that innovation. It's because of the resource constraint at that time the company had. But over time, the team is thinking through is how we can be disciplined, how we can innovate. We co-innovate with TSMC on packaging technology, on the process technology, so to drive our architecture, our overall product, much more competitive. That has been the team's execution model. We're going to lean in investment, but we're going to be very smart, efficient to invest.

If you look at our investment right now, R&D investment in 2024, we're doing three generation of data center GPU at the same time, MI325, MI350, and MI400. At the same time, we're investing in our server CPU roadmap, and client CPU roadmap, and gaming graphics roadmap.

The platform leverages how teams think about investment is really leveraging the CPU platform, GPU platform, and the software platform. So I'm actually really impressed by how teams think about-- as a CFO, you always want to control the OpEx. But the team has been always very thoughtful to balance driving long-term companies growth versus how we expand the margin.

Overall, I would say is we always expand the revenue faster than operating expense growth. At our scale, that will drive very significant operating leverage.

JOSEPH MOORE: Very helpful. Jean, Matt, thank you so much for your time. Appreciate it.

JEAN HU: Thank you.

MATT RAMSAY: Thank you.