

**Meta Platforms, Inc. (META)**  
**Second Quarter 2026 Results Conference Call**  
**July 29<sup>th</sup>, 2026**

**Chad Heaton, VP, Finance**

Thank you. Good afternoon and welcome to Meta Platforms' second quarter 2026 earnings conference call. Joining me today to discuss our results are Mark Zuckerberg, CEO and Susan Li, CFO.

Our remarks today will include forward-looking statements, which are based on assumptions as of today. Actual results may differ materially as a result of various factors including those set forth in today's earnings press release, and in our quarterly report on Form 10-Q filed with the SEC. We undertake no obligation to update any forward-looking statement.

During this call we will present both GAAP and certain non-GAAP financial measures. A reconciliation of GAAP to non-GAAP measures is included in today's earnings press release. The earnings press release and an accompanying investor presentation are available on our website at [investor.atmeta.com](https://investor.atmeta.com).

And now, I'd like to turn the call over to Mark.

**Mark Zuckerberg, CEO**

Hey everyone. Thanks for joining today.

We had a strong quarter for our community and business with 3.6 billion people using at least one of our apps each day. We reached several milestones. Instagram reached 2 billion daily actives. Threads crossed 500 million monthly actives, making it the fastest growing conversation app ever. Facebook has reached more than 2 billion daily actives for a while now. WhatsApp just hit an all-time messaging record, peaking at 30 million messages sent per second during the World Cup Final. And also, Kunal Shah just joined as our new head of WhatsApp. He built one of India's most important payment companies and will be a great addition to the team. We also shipped some strong new models from Meta Superintelligence Labs and released new glasses. Overall, the scale and reach of our community across our apps is pretty remarkable, and it gives us a strong platform for delivering new innovations to billions of people.

The opportunity in front of us is massive. First, we are now at a point where our investments in AI are accelerating every major part of our core business. They're improving the experience for people using our apps, driving better performance for advertisers, and helping our teams build new experiences and ship faster. Second, we're developing new personal agents that will be the foundation for our next wave of products and revenue lines in the months and years ahead. Third, we see a large enterprise opportunity to sell to businesses, including APIs, business agents, potentially selling compute directly, and other services that we're building for large customers. I want to spend some time today laying out exactly how our investments are delivering results today and the opportunities that we see over time.

For accelerating our core business, we're seeing strong and promising results in several areas.

In Instagram and Facebook, I'm very optimistic about our work to integrate large language models into our recommendation systems. LLMs add a first-principles understanding of what the content is about and why it is compelling, as well as a deeper understanding of what people are interested in and what their goals are when they're using our apps. This means that we can show more relevant and engaging content that better reflects people's goals and interests.

Our new Muse Image and Muse Video models will also dramatically expand the universe of content that people can discover across our platforms. There are already two large sets of content to draw from -- first from your friends and the people you follow, and second from creators that you don't follow -- but now there is going to be a whole new and nearly infinite universe of personalized content. This is going to make our services a lot more useful and engaging for people.

For ads, we are using LLMs to improve how our systems predict and rank the ads that we show. We've expanded the context that we can take into account around a person's organic and ads activity to determine an ad's relevance, driving significant increases in relevance and conversions on both Facebook and Instagram. On a dollar basis, our ads business is reporting faster year-over-year revenue growth than any other company's reported ad business -- so these AI investments are paying off.

We are also seeing a lot of demand for our new AI-powered creative tools. 9 million small businesses on our platforms are now using at least one of our AI ad creative tools, and we're rolling out new end-to-end creative solutions that help advertisers translate performance data into their creative decisions. Muse Image is going to supercharge this. The model can analyze images, improve its own work, and produce better ad variations based on advertiser input. We're getting great feedback on this so far.

We also just launched Meta One, a new subscription offering that provides more tools and AI features across our apps. As demand grows, we're going to offer a variety of different tiers and pricing options there as well.

I'm also excited about how AI is helping our teams speed up product development. Earlier this year we shipped Instagram Instants. We also just launched Forum, a standalone Groups app, and Seller, a standalone Marketplace app. And I expect it to become a lot easier to ship new apps, so we're planning to build out more ideas and use our recommendation systems to scale them to the people who will find them interesting, as we've done with Threads.

It's been a little more than a year since we launched Meta Superintelligence Labs and our trajectory is strong. In the last month we shipped Muse Spark 1.1 and Muse Image. Since we rebuilt Meta AI and integrated Muse Spark, we've seen a 60% increase in the number of people interacting with the assistant each day and that continues to grow quickly week-over-week.

Muse Spark 1.1 is a strong agentic and coding model that is very efficient and excels at computer use, tool use, and multimodal understanding. It's available through our new public API, and we're ramping up distribution through partner channels and more coding agents over the coming weeks. We're also building out features to make it easier for enterprises to adopt Muse Spark and that is an area that we expect to continue focusing on.

One reason we're so focused on making Muse Spark great at agentic capabilities is that we think there's a very big opportunity to ship a few types of agents that are aligned with our mission and business.

The first is personal agents. Soon, we will have agents that can work 24/7 on your behalf to help achieve your goals and improve your life, your health, your relationships, your finances -- whatever you want. The first domain that agents have really taken off in is coding, but engineers are more technical and willing to spend time making those agents work. So to build great personal agents, this needs to be a great consumer product that just works out of the box and is easy enough for billions of people to adopt and use. I'm very excited about this and we're going to have more to share soon.

As we move towards a future where we're all interacting with multiple agents, I think that WhatsApp and our other messaging surfaces are going to be increasingly important. WhatsApp is already the leading surface where people engage with Meta AI -- and as we build out a platform for more agents across our messaging apps, we're going to innovate on how to deliver a private and secure AI experience. We launched incognito mode this quarter on WhatsApp and the Meta AI app, allowing people to have private conversations with their assistant that even Meta can't see. We're planning to make strong privacy and security a fundamental part of the agents that we're building as well.

I'm also very excited about our progress with business agents. We made Meta Business Agents available globally this quarter on WhatsApp and Messenger, and there are already more than 1 million businesses using them to talk to their customers or complete sales every week. We're rolling business agents out on Instagram now too. One interesting thing about having an agent talk to your customers every day is that it learns over time and can bring all those insights back to you. So we're building more agentic capabilities to summarize all these conversations, digest what happened overnight, and surface what customers are asking for. Soon, it'll go further, including suggesting ways to grow your business, giving you competitive intelligence and real-time insights into what's working and what's not. Over time, we'd like to build this into a business-in-a-box service that can help you start and run a whole business using Meta's platforms.

In terms of how we will monetize these, we have a mix of subscriptions, volume-based pricing, and I expect that we're going to evolve more of these products to be like our ad systems where businesses only pay us when we achieve results for them. Over time, that will let us run an efficient auction over our compute, similar to how we do that for our advertisers today.

As AI usage in our products and businesses continues to ramp, we continue to invest aggressively in infrastructure to meet the demand. Yesterday, as part of our Meta Compute effort, we announced a new strategic venture with BlackRock to develop a new 1GW data center in El Paso, Texas.

Overall, we expect that a significant portion of our compute is going to go towards training our models, growing our core business, and delivering personal agents and new products. But we also expect to grow a large business serving large customers as well. We've built our API, we're rolling out business agents, we're getting a lot of offers for compute at a significant premium over what we paid for it, and we have more coding and productivity tools on our roadmap as well. We'll have more to share on all of this soon.

As we get closer to personal superintelligence, we're also going to need hardware that allows you to seamlessly interact with it. Glasses are the ideal form factor since they can be with you throughout the day and they can assist you without pulling you away from the moment. Our

glasses remain one of the fastest-growing consumer electronics of all time, and we continue adding to the line-up.

We just released our own line of Meta Glasses in collaboration with EssilorLuxottica, including a style that we designed with Kylie Jenner. They're the first glasses to ship with Muse Spark out of the box so that they can understand what you're seeing and give even more helpful answers. Early sales have been strong, exceeding our expectations. We're going to have more to share on our glasses line-up at this year's Connect conference on September 23rd, so I encourage you to tune in then.

Before I wrap, I want to mention that I just published an op-ed about why I'm so optimistic that we're building a positive future for everyone. At Meta, we've always built technology to put power in people's hands so they can connect with the people they care about and shape their world in the ways that they want. It's why we've always focused on making our products affordable and accessible for everyone. And it's a strategy that has served both our community and our business well.

As we enter this next chapter, that same philosophy is guiding how we approach developing AI. We're the only major company building AI with the primary goal of putting superintelligence directly into people's hands. Rather than centralizing superintelligence, we're focused on distributing it widely and giving everyone the ability to direct it towards what matters to them. That's the way that society has always made progress. I think that these are the right values for building a positive AI future. And if we help build this, then I think that we will continue to build a very strong business as well.

That's everything I wanted to cover up front today. AI is improving our core business -- it's making our apps more relevant, and delivering better results for businesses. We're starting to deliver more novel products, and we'll have a lot more there soon as well. We're investing aggressively because the potential is huge and we know that there are many ways to deliver value here. As always, I'm grateful to all of you for being on this journey with us. And now, here's Susan.

## **Susan Li, CFO**

Thanks Mark and good afternoon everyone.

Let's begin with our segment results. All comparisons are on a year-over-year basis unless otherwise noted.

Q2 Total Family of Apps revenue was \$60.4 billion, up 28% year-over-year.

Q2 Family of Apps ad revenue was \$59.4 billion, up 27% or 26% on a constant currency basis.

In Q2, the total number of ad impressions served across our services increased 14%. Impression growth was healthy across all regions, driven by growth in engagement and users, as well as ad load optimizations. The global average price per ad increased 12% year-over-year, driven by ad performance gains, improvements in macro conditions relative to Q2 of last year, and currency tailwinds. This was partially offset by strong impression growth, particularly from lower-monetizing surfaces and regions.

For the first time, quarterly Family of Apps other revenue reached \$1 billion and grew 73% year-over-year, driven primarily by WhatsApp paid messaging and subscriptions revenue.

Within our Reality Labs segment, Q2 revenue was \$431 million, up 16% year-over-year due to strong growth in AI glasses revenue, partially offset by lower Quest headset sales.

Moving now to our consolidated results.

Q2 total revenue was \$60.8 billion, up 28% or 27% on a constant currency basis.

Q2 total expenses were \$42 billion, up 55% compared to last year, and included \$2.4 billion in charges related to legal proceedings and \$1.2 billion in severance expenses in connection with the May 2026 headcount reduction. Year-over-year growth was primarily driven by increases in employee compensation, infrastructure costs, legal-related costs, and third-party AI token costs.

Excluding the previously mentioned severance expense, growth in employee compensation was driven by technical hires we've added over the past year, particularly AI talent.

The growth in infrastructure costs was driven by higher depreciation, data center operating costs, and third party cloud spend.

We ended Q2 with over 75,000 employees, down 3% from Q1. This total includes approximately 8,000 employees impacted by the May 2026 headcount reduction. We expect the majority of impacted employees will no longer be captured in our headcount by the end of Q3 2026.

Second quarter GAAP operating income was \$18.8 billion, representing an 8% decline year-over-year and a 31% operating margin. Excluding the Q2 legal charges and severance expenses, our second quarter operating income would have increased 9% year-over-year.

Our tax rate for the quarter was 16%.

Net income was \$15.8 billion or \$6.18 per share.

Capital expenditures, including principal payments on finance leases, were \$31.1 billion, driven by investments in servers, data centers, and network infrastructure.

Free cash flow was \$784 million. We ended the quarter with \$90.3 billion in cash and marketable securities and \$83.7 billion in debt.

Turning now to the business performance. There are two primary factors that drive our revenue performance: our ability to deliver engaging experiences for our community, and our effectiveness at monetizing that engagement over time.

On the first, we continue to see significant gains from our content recommendation initiatives.

On Instagram, global time spent this quarter grew double digits year-over-year, largely driven by improvements to our Feed and Reels recommendations. On Facebook video time spent increased 9% globally year over year and over 10% within the US & Canada, where it was driven by ranking improvements.

We are finding that LLMs are increasingly capable of delivering ranking and recommendations gains. First, they make our existing systems smarter by understanding what the content is actually about and generating better training data. And second, LLM-powered agents are also helping with engineering development by evaluating content quality, detecting trends, and testing ranking changes.

Earlier this year we reached a milestone of every public Reels and Feed post on Instagram being automatically processed through an LLM and analyzed across dimensions from topics to tone, and we're working towards including more surfaces on Facebook as well. These signals can then be passed to downstream applications across ranking, recommendations, and content policy enforcement, which is a key building block toward greater personalization. This quarter we also began using our Muse family of models to conduct content understanding across signals like video topic classification and summarization, and we've seen positive early results.

Finally, our recommendations are also becoming more personalized, surfacing more fresh content while giving people more direct control over what they see.

On Reels, we shipped our largest single-release ranking improvement to date, combining faster inference with a new architecture that draws on deeper user history to improve predictions. This drove a 15 basis point increase in sessions on Instagram, with particular strength in reshares and time spent, which are both strong indicators of better content-to-user matching. We are now bringing this to Feed, where early results look comparable.

We are also getting new content to people more quickly. Investments we've made in more real-time infrastructure and modeling improvements on new videos are allowing our largest ranking models to now identify high-quality new Reels at creation. On Instagram Feed, over half of all recommended content is now less than one day old, more than double from a year ago.

We're also giving people more direct control of the content they see. Today, Instagram users can visit the "Your Algo" page, which lets users write natural-language prompts to tune their recommendations. Similarly, on Facebook, we launched "Shape Your Feed". Early results show over 80% retention among users who engage with it.

Looking forward, we are executing on our longer-term efforts to develop the next generation of our recommendation systems. This includes building foundation models that are designed to power organic content and ads recommendations simultaneously, as well as developing LLM-native recommender systems. We hit our first research milestone this half by continuously pre-training a large-scale model with recommendations data and observing healthy scaling laws in the process. We are encouraged by this milestone and expect continued progress in the second half of the year.

Turning to the second driver of our revenue performance: increasing monetization efficiency.

The first part of this work is optimizing the level of ads within organic engagement.

Here, we continue to enhance our systems to show ads at the optimal time and location. In Q2, we also expanded availability of ads on our newer surfaces, including completing our global ads expansion on Threads. On WhatsApp, we have introduced support for more types of ad destinations and advertiser performance goals in Status and continue on track toward our global rollout.

Moving to the second part of increasing monetization efficiency: improving performance for the businesses who use our services.

Within our ads systems, we're delivering performance gains as we deploy more complex and predictive models.

This quarter, we introduced Meta Generative Recommender, a paradigm shift in how our ads system works. Rather than scoring every possible ad individually, we are now using LLMs to reason about ad content and user preferences together, and predict the best ad for each person. This makes our ad matching more intelligent and more precise, which compounds performance gains for advertisers. We deployed the first generative model into our ads retrieval system and saw notable improvements in ad performance. Early pilots using LLMs to better understand user preferences drove a 1% increase in app event conversions on Instagram.

In Q2, we also advanced our user understanding models to analyze ads and organic activity and simultaneously improve both user experience and advertiser performance. Combined with our GEM model for ads ranking and sequence learning, these advancements generated an 8.3% increase in ad clicks and a 15.7% uplift in conversions on Facebook.

We're also leveraging AI to empower businesses to more easily manage their campaigns, develop ad creative, and engage with their customers.

Our AI-powered Advantage+ end-to-end solutions continued to grow, reaching over \$75 billion in annual revenue run-rate this quarter. We are working to deepen adoption as advertisers who leverage multiple tools see compounding performance gains.

I'll share one example of how Advantage+ is making performance marketing meaningfully easier and more effective for SMBs. Underneat, an online apparel brand in India, had been setting up each campaign manually across Facebook and Instagram. After adopting Advantage+ sales campaigns layered with Advantage+ Audience, Placements, and Budget optimization, they saw a 13% incremental lift in purchases and a 16% increase in add-to-cart conversions.

Adoption of our Gen AI ad creative tools continues to scale, with over 9 million small businesses using at least one AI creative tool. Image generation, which now lets advertisers produce more creatives at scale from existing content, including a new ability to create images from video assets, saw adoption more than double this quarter. We also introduced a new end-to-end creative solution that gives advertisers the AI infrastructure to translate real performance signals into their next creative decision, while preserving brand identity and tone. We are building with agency integrations from day one so teams can diagnose, generate, and scale high-performing creative without leaving their existing workflows. Looking ahead, with the rollout of Muse Image, we expect to further expand advertisers' ability to generate high-quality, on-brand creatives at scale.

With Meta Business Agent, businesses are better able to serve their customers through our messaging apps by responding to inquiries, recommending products and handling support around the clock. Earlier this month, we also introduced the Meta Business Agent Platform, which gives Enterprises the infrastructure to build, customize and deploy their Business Agent at scale on WhatsApp. The platform provides larger businesses with enterprise-grade controls, guardrails and measurement built in so they can define rules and offer personalized experiences, starting within the messaging apps that their customers already use.

Movida, one of Brazil's largest car rental companies with nearly 400 locations, deployed a Business Agent on WhatsApp to handle the entire booking flow, from vehicle selection and pricing to payment, in a single conversation. Returning customers could complete a reservation in as few as three messages. In a one-month period, Movida reported a 44% increase in daily bookings through WhatsApp when compared to the same period in the prior year, and that 85% of conversations in the channel were resolved entirely by the AI agent without human assistance.

We are also building out additional ways to monetize our ecosystem. Two additional revenue streams are subscriptions and monetizing our competitive models through an API.

Meta One is an evolution of our subscription portfolio to create more value for everyday users, businesses and creators, so they get more features and AI tools to create, connect and stand out. We are excited to bring this to more users and continue to build enhanced tools for our subscribers.

We also recently launched a high intelligence model API at a competitive price and are encouraged by the initial results. We recently made Muse Spark available on OpenRouter for US based developers, broadening its distribution and making it easier for developers to adopt the model. We soon expect to roll out the model API to more distribution channels, make it available in more countries, and open it up for enterprises.

Our approach to building capacity is strongly influenced by several key elements:

First, the broad environment for building infrastructure is dynamic and uncertain in both near-term and longer-term time horizons. The industry has underbuilt historically for the wave of AI adoption, making existing capacity, including our own, extremely valuable. Longer-term, the supply chains need to be built out to support the capacity that we anticipate we and others will need for AI-powered experiences.

Second, we have high confidence in our ability to utilize capacity to scale and build on top of our existing experiences, as well as continue to invest in foundational models that will create substantial new opportunities. Consequently, our current plans are geared towards maximizing 2026 and 2027 capacity. When we have had incremental capacity in the past, it has proven extremely valuable in scaling experiences like Reels, and we are confident that this will be true in this time frame as well.

Longer-term, it is harder to predict the exact usage scaling curves, but we believe that our distribution advantages will give us the opportunity to serve AI products that are valuable for everyone: both our 3.6B users and millions of businesses. This should be true regardless of whether our models are on the frontier, but we believe that being on the frontier will unlock new markets and opportunities for which we may need additional compute.

Therefore, our longer-term capacity strategy aims to give us the flexibility to continue growing compute in 2028 and beyond by laying down data center and network foundations to accommodate future server decisions. The long-lived nature of these assets inherently provides the flexibility that will make it possible to adjust our investment to the pace of AI adoption. In addition, we have been making strategic investments in areas like our internal custom silicon effort, which will provide long-term strategic flexibility and supply chain leverage. This will be helpful in driving better returns on those long-term investments.



Finally, we believe that overall industry capacity is going to remain tight for the foreseeable future. As we've said earlier, we strongly believe that the models, consumer experiences, and enterprise offerings that we are building will be the best and highest ROI use of our infrastructure. Those enterprise offerings have the potential to take multiple forms as Mark mentioned – agentic tools, our API, or monetizing compute directly given outsized market demand. We expect that remaining nimble about these opportunities will help us fund our buildout more efficiently while preserving our strategic flexibility to have the compute when we need it and provide us multiple pathways to generate returns on invested capital.

In funding these infrastructure investments, the strength of our balance sheet gives us the ability to attract capital from a wide range of markets to supplement the cash flow generated by our business. Our announcement with BlackRock yesterday is an example of the partnerships that we can structure to complement our approach to building infrastructure capacity.

Moving now to our financial outlook.

We expect third quarter 2026 total revenue to be in the range of \$61-64 billion. Our guidance assumes foreign currency is an approximately 1% headwind to year-over-year total revenue growth, based on current exchange rates.

Turning to the expense and capex outlooks.

We are raising the lower-end of our expense outlook to incorporate the \$2.4 billion charges related to legal proceedings recognized in Q2. We now expect full year 2026 total expenses to be in the range of \$165-169 billion.

We continue to expect to deliver operating income this year that is above 2025 operating income.

We anticipate 2026 capital expenditures, including principal payments on finance leases, to be in the range of \$130-145 billion, narrowed from our prior outlook of \$125-145 billion.

Absent any changes to our tax landscape, we expect our tax rate for the remaining quarters of 2026 to be between 15-17%, an increase from our prior outlook of 13-16%.

Finally, we continue to monitor active legal and regulatory matters that could significantly impact our business and financial results. For example, we continue to see scrutiny on youth-related issues in several markets and have a number of youth-related trials scheduled for this year in the US, which may ultimately result in a material loss.

In closing, our business momentum continued in Q2, with strong execution across our core ads and engagement initiatives. We're also progressing in our efforts to bring personal superintelligence to everyone, with exciting model releases and we expect to build on that momentum over the course of this year with new products.

With that, Krista, let's open up the call for questions.

Operator: Thank you. We will now open the lines for a question and answer session. To ask a question, please press star 1 on your touch tone phone. To withdraw your question, again, press star 1. Please limit yourself to one question. Please pick up your handset before asking your question to ensure clarity. If you are

streaming today's call, please mute your computer speakers. And your first question comes from the line of Brian Nowak with Morgan Stanley. Please go ahead.

Brian Nowak: Thanks for taking my questions. I have two, one for Mark, one for Susan. Mark, I appreciate all the color on the big pipeline for new products across consumer and business agents, API tools, and compute rental. There's a lot of opportunities here.

My question is, as you look at these opportunities and the state of the current offerings, the compute capacity, which of them do you expect to be able to scale first sort of in '26 and '27 to sort of showcase quantifiable material ROIC for investors?

And then Susan, there's been some public comments about capacity in '27 and doubling capacity and appreciate your color about CapEx. Any early comments on '27 CapEx, even the philosophy around sources of upside or sources of downward pressure, as you sort of think through different ways to finance this multi-year build and have '27 CapEx?

Mark Zuckerberg: I can take the first one. So in terms of the different opportunities and how we think about the compute, overall a substantial amount of the compute goes towards training models to be a leading lab and I think that's an important investment.

But then the rest of it goes towards a set of different products and revenue opportunities, which spans from optimizing and improving our core business to building new consumer products that we're releasing soon to the API, to the business agents work, to the developer tools work on the road map that I alluded to.

And then also the opportunity to sell compute directly where I mean, I mentioned that we have quite a number of offers at a meaningful premium over what we paid for the compute. The question and thinking about this is we believe that there will continue to be a significantly higher margin on selling intelligence rather than selling compute directly.

But we think that there's a big opportunity obviously to sell compute as well. So we're thinking through -- I think you asked what one area would likely scale the most. But I mean, I'm actually quite optimistic that we're going to see meaningful growth in all of these areas. And I think we will have more to share soon on a number of them.

Susan Li: Brian, on your second question, we aren't providing a specific outlook for 2027 CapEx at this time. Infrastructure planning remains highly dynamic and even this year, there are a range of outcomes embedded in our outlook.

I alluded in my main remarks to our focus on being -- on gearing our current infrastructure plans towards maximizing capacity in 2026 and '27 and giving us the flexibility to continue to grow in '28 and beyond but also giving us the

ability to make server decisions when we come to our being able to evaluate our actual needs in '28 and beyond.

So we're still working through what our capacities are going to be over the coming years. Generally, we believe near term capacity is more valuable than long term capacity and it remains a very dynamic planning process.

Operator: Your next question comes from the line of Eric Sheridan with Goldman Sachs. Please go ahead.

Eric Sheridan: Thanks for taking the question. Maybe two, if I can. When you frame up the enterprise opportunity, how much of that opportunity do you think is available to you today based on what you've built out in terms of go to market strategy as extensions of the advertising and the marketing business you have already versus go to market strategies that might have to be built to capitalize on the opportunity?

And then maybe Susan, if I could just squeeze a second one in. When you think about sources of capital for the business for the next couple years, you referenced the deal that was announced yesterday as an example of looking at ways to finance forward obligations.

We continue to get questions from investors about how to think about the mix of debt and equity and sources of capital. Philosophically, how are you guys thinking about wanting to be ambitious on the spend, but then marrying that with the need for capital? Thanks so much.

Mark Zuckerberg: Sure, I can talk about the first part. I think for enterprise, there's going to be a combination of extending the current business, which is effectively selling to marketers and businesses that are basically customer facing, trying to reach customers and sell to them directly.

Obviously that's the vast majority of the business across Facebook and Instagram. And we believe that there is an opportunity to extend this with business agents across messaging apps and other services to continue to interact with customers.

And just like the ad system, effectively, we will get paid when we deliver results for those businesses. So we view this as an extension of the sales and the partnerships that we have with many millions of advertisers and hundreds of millions of small businesses that use our platforms.

So that one I think should be quite a natural opportunity for us. And we're focused on delivering it in a way where we're not trying to kind of maximize the sale in the near term. We're trying to maximize the results for people and build out a robust auction. And that's what we've seen has served the business well over time.

There are other enterprise customers who I think we're increasingly going to serve too. We're building coding and developing and internal productivity tools

partially because we need to build them ourselves and we need to make sure that we have tools that are tuned for ourselves.

And now that we have those, we feel like there's a large opportunity to serve, whether that's small businesses or larger businesses. That is a somewhat different muscle than we have historically had. And we will share more soon on how we're planning to build that out.

But I think that there is just a very large opportunity on this. And the way that I think about this is there's obviously been a bunch of news about the compute side. But I think that the enterprise opportunity is kind of the sum of all of these different things. It's not just the selling compute, also the API services, the productivity services, the kind of business agents for other parts of the business beyond marketing are all parts of the overall offering.

And I think that there's just a very, very large opportunity there. So we're quite focused on that. That's going to be somewhat of a new muscle that we build as a company, but I think it's a very important one that we build so that way we can make sure that we can maximize the opportunity ahead of us.

Susan Li:

Eric, on your second question in terms of sources of capital, this is something that we have been looking at thoughtfully as we think about financial planning for the future.

Obviously, our strong operating cash flow certainly has put us in a position of strength as it pertains to funding our infrastructure build out. But we've also been evolving our capital structure in recent years to include a greater mix of debt as we work to bring down our cost of capital.

And we have generally found it prudent to continue adding cost efficient long duration sources of capital as we make investments in initiatives that themselves have long time horizons, especially AI infrastructure projects.

We've also broadened our aperture there to include partnerships like the one we announced with BlackRock yesterday. And we'll continue to be thoughtful about evaluating the appropriate different sources of capital overtime as we evaluate future projects.

Operator:

Your next question comes from the line of Mark Shmulik with Bernstein. Please go ahead.

Mark Shmulik:

Yes, thanks for taking the question. Mark, everyone's got a story of someone coming back from Silicon Valley, deep writing code, building agents with AI, and then they head home, and they kind of tell their parents they're using AI wrong, as just kind of like a glorified search tool.

Consumer behavior is always pretty difficult to predict. But kind of reading your op-ed on AI for everyone, how do you think about whether consumer adoption can close this AI utility gap? Like are we on the cusp of something kind of breaking through or do we just need to be a bit more patient? Thanks.

Mark Zuckerberg: Well, I think that some things have already broken through. And I think one of the interesting things about AI compared to other technologies is that every year or so, and the cycle may accelerate, but every year or so you get new capabilities that create new possible product lines.

So there's obviously the AI assistant market for consumers where that's what we're doing with Meta AI. And then there's the other products that competitors have in that space. And then in the last year, I think the coding agent market has really grown very quickly and that's the first real agentic market.

And I think there's a number of reasons why coding is first, right. You have technical customers who are willing to do the work to make it work. Coding is an inherently digital and closed loop activity. So it's sort of, in some ways it is somewhat of a, kind of an easier system to train on. But our bet is that -- or one of the bets here that we're making is that we think that consumer personal agents is going to end up being an extremely important and massive market.

I think that it's extremely unlikely if you look out five years from now, for example, that -- or whatever period of time you want, that you don't have billions of people with a personal agent that understands your goals, and that is just working on your behalf 24/7 to achieve your goals.

And whatever the domain is that you care about, whether it's helping you with your health or your hobbies or your personal finances or your productivity and running your home better, or improving and enhancing your relationships, helping with your career -- just all these different things.

This is like a very, very deep set of use cases. But like you said, I think that building for consumers is a little bit different than building for developers. If you're building for consumers, especially if you're trying to build something that isn't used by millions of people, but is used by billions of people, it needs to just work, right? It needs to be simple.

And I think a lot of what we see with these agents today, a lot of the kind of proto-agents is, it takes a bunch of fiddling. You have to get into a terminal to get it set up. You get a use case going. It kind of seems magical, but then maybe it breaks down over time. And I think that the companies that can deliver the personal agents that just work, I think that this is going to be one of the next major opportunities in AI.

And we are very excited about that in addition to what we're doing with Meta AI, in addition to what we're doing with business agents, in addition to all these other things, while also developing the models that we think are going to continue to advance new capabilities to create new product lines as well.

So this is why we're very optimistic. Now, I get we're going to ship this at some point soon and that's going to be very exciting. And we haven't done that yet.

So this is -- there's kind of only so much that I can -- that I can say on an earnings call about this.

But I think this is a very large opportunity and one that I think is almost inevitable that someone does. And I think it really plays to Meta's strengths as a company. We build consumer products that reach billions of people.

We're great at once we get something working, scaling it to a large number of people, and we're good at building the infrastructure to be able to support these intensive applications. So, I feel very good about this. I kind of understand that we need to deliver it for our community and that's what we're very focused on.

Operator: Your next question comes from the line of Doug Anmuth with JP Morgan. Please go ahead.

Douglas Anmuth: Great thanks for taking the questions. One for Susan, one for Mark. Susan, you talked about how LLMs are increasingly capable of delivering ranking and recommendation gains. Can you just talk more about the roadmap here and how far along you are in just leveraging better models and more compute?

And then, Mark, just in terms of the number of offers to monetize your compute externally, you also at the same time are purchasing capacity from a number of third parties.

So just hope you could help us understand some of the differences here? Is it just timing and stop gap issues? Or is it training versus inference, and leveraging the chips that are best suited for each?

Susan Li: Thanks, Doug. I'll take that first question about where we are on the recommendations roadmap. First of all, we certainly see further headroom to continue improving recommendations over the rest of the year and into 2027. And we expect that will help us drive additional gains on both engagement on Facebook and Instagram. A couple of things I would highlight.

First, we'll continue to make recommendations even more personalized and relevant to user interest. And that's in part by advancing our recommendation models and architectures to further capture user interest more precisely and respond faster to what people care about in a given moment.

Our AI investments are going to play a significant role in delivering on this vision including the expansion of LLM based content understanding to develop a deeper understanding of posts and creators that people value, to capture user interests more precisely and respond more quickly to what they care about in the moment and using AI to surface high-quality, fresh and trending content and reduce the share of low-quality content.

Second, we're continuing to improve our data infrastructure to allow our models to train on more data and leverage that data more effectively.

We're adding more detail to how we describe content that users have engaged with in the past and enriching past user interaction sequences with more granular content. That allows our models to more precisely learn which engagements are more or less valuable to users.

We're continuing to scale up both the length of user interaction sequences we use during training as well as the complexity of our model architectures across Facebook and Instagram to take advantage of the larger data sets. And then we've already made significant strides, leveraging LLMs for content understanding.

We're going to further incorporate them into both our recommendations and content policy enforcement stacks given their capability to more deeply understand content. And I would say, broadly, I think we're very optimistic about that body of work.

We've also invested in using LLM-based agentic approaches to transforming our recommendations system. And we grew the number of launches from our ranking agents this half. That also helps make our engineers more productive and that's another path that we're very excited about as well.

Mark Zuckerberg: I can answer the second part of that. I mean I think your question was about how we think about offers that we're getting to sell the compute, but then we're also buying the compute. I mean look, the high-level observation is that there's just nowhere near enough compute for all the demand.

So that is why we see that like basically, we are getting a large number of offers for the compute that we have, but also we have a lot of internal uses that we think are going to be quite valuable.

Now in terms of running the business, obviously a common trade-off that we need to make is around how much do you monetize something today versus develop future assets for the future? And I think that it's always a portfolio, right? It's not like -- you don't want to only do long-term things and not kind of prove the markets out that exist in the near term.

But I also think it would be foolish to basically just sell all of the compute and take a short-term profit. But when you have the opportunity to build intelligence on top of it which will be a multiple and that compounds the value of the compute on top of that.

So I think the answer is what we're doing which is to basically use a lot of our capital to build out compute, having confidence that we have the ability to monetize the compute directly when that makes sense, but also knowing that we have quite a number of different use cases to monetize the intelligence on top of the compute including the enterprise cases that we talked about and including some of the consumer cases that we talked about, and including just the core business which is not even necessarily new products that we haven't talked about, but just in terms of using that to be able to further add

intelligence and improve the ranking and recommendations and ads in the core services.

So I think all of that is true, and it creates this dynamic where there is a lead time where we're investing in building out these data centers now. They come online at some point in the future. You obviously are not getting value out of them until they're online.

But that's -- we basically see a very large demand for all of this, and want to go maximize the opportunity to build out all of these different businesses.

Operator: Your next question comes from the line of Justin Post with Bank of America. Please go ahead.

Justin Post: Great, thanks. Mark, you've hired the senior leadership of your AI labs about a year ago. Could you just give us your thoughts on how the lab is performing? And do you think the Street will really see an uptick in product velocity as far as models or chips or other things coming forward? And then what kind of sustainable competitive advantages do you think the lab is building? Thank you.

Mark Zuckerberg: Yes. So I mean I'm quite happy with the trajectory that we're on. We've released a few models that I think are quite impressive on our early scaling ladder. And as I've said, we're in the process of scaling much larger and more advanced models and we're excited about those too. I think that there's the intelligence aspect of this, and then there's the data aspect of it.

And I think it is true that in any kind of product category that you want to go into, there is going to be some kind of flywheel where you learn from the behaviors and how people in that community are using the product.

So -- and kind of the better that you can do there, the more feedback you get and then that inevitably makes the products better. So I think this is a field where people understandably are obsessed with intelligence. But the data and the knowledge part of it to serve people well is really important.

If you're talking about personal superintelligence, then obviously having a very clear mental model of everything that is going on in the person's life and their goals is going to be a very important aspect of that.

And certain companies start with different advantages in different markets and different parts of that. But also, I think part of the reason for working on a number of these different things upfront is that, first of all, the technology is general.

So when you build the intelligence out, it can apply to a number of different uses, but then you want to invest in building up the flywheel on them because I think that that's how you build the sustainable advantage over time.



But I think that we've definitely shown at Meta that when we have a product in a format that works, I would go as far as to say that I think we're probably the best company in the world at scaling those experiences to billions of people. So I feel quite good about the personal agent work.

I feel very good about the business agent work where we kind of already have our base of advertisers and small businesses that we serve and the ability to scale that out and roll it out around those.

I think that those are durable advantages as well as the data flywheels that we're building on them. And then obviously on the research side, you want to build good culture.

I mean this is just kind of the simple stuff, but it's -- I think that it's -- you just want to make sure that you're building the team in a way where you're managing things well and kind of consistently with low drama compounding over time.

And that's sort of the goal that you aspire to. And I think if you can do that well then that is hopefully -- may not be kind of a clear mathematical explanation of something, but I think that that's kind of how business works. So yes, that's what we do.

Operator: Your next question comes from the line of Ross Sandler with Barclays. Please go ahead.

Ross Sandler: Yes, hey Mark. Sticking with the comments on the AI Lab. So Muse Spark 1.1 is pretty close to the Pareto Frontier, but it's at the lower end of the intelligence spectrum or lower cost end, I should say. And it sounds like from your last answer, you think it's important to compete at both the lower cost and the more expensive performant end.

Could you just talk about that a little bit? And then your -- the leadership of your lab has also talked about going back into open source, kind of where you were a couple of years ago. So how does that fit into the strategy and all this discussion around monetization of AI products and models? Thank you very much.

Mark Zuckerberg: Yes. So I mean just to get into the kind of scientific process on this a little bit, basically, training large models, at each stage of training a large model, you see novel behaviors.

So basically, the way that you want to do this is you build up from training smaller models to building larger models. And the models that we've released so far are based on -- they're a certain scale in climbing up the scaling ladder and we're continuing to scale larger models.

So I think for Muse Spark 1.0 and Muse Spark 1.1, I think that they are very impressive models for the scale of the model, the kind of stage of development of the lab. And I feel quite good about them.

Now we want to have models that are more advanced as well and that's why we're scaling larger models, and that's why we're building out a bunch of the research infrastructure around that.

But at the same time we also want to have good, more efficient models that will be a lot of what we serve to consumers at scale, right? If you're serving billions of people, you want the ability to have more advanced models for things that are very hard problems and you want the ability to serve the vast majority of prompts from simpler and more efficient models.

So I think both are going to be very important. The efficiency really matters, but I also think that we want to be able to solve the hardest problems for businesses and customers around the world. So we care about both of them.

On open source, I think we have always felt like open source was an important part of the ecosystem, and it's good for the world, and it creates its own feedback loops that are positive for us around getting the community invested in our infrastructure stack and our work and contributing improvements. But we've always basically said that we were going to do a mix of open and closed. And that continues to be true.

Now in ramping up Meta Superintelligence Labs, in some ways, actually counterintuitively, it takes some more work to do open source models because you're not necessarily -- if you're doing something as a closed system that you're only building for your own use cases, it can be a little more jagged.

Whereas if you release it as open, and it's going to be used for a lot of things, you want to make it more well-rounded. And I just wanted to make sure that the MSL team was uninhibited in building out the most intelligent models that we could. And we expect that we will get back to releasing some open source models at some point soon.

But like we've always said, we're not dogmatic about this. We think open source is important. We want to contribute to that ecosystem. We plan to do a combination of open and closed models.

Operator: We have time for one more question, and that question comes from the line of Ken Gawrelski with Wells Fargo. Please go ahead.

Kenneth Gawrelski: Thank you. Two, if I may, please. I just want to maybe, Mark, just touch on the last point again, around open weight models. There's been a lot of discussion, you've weighed in on the topic.

Why or why not does that change Meta's view of developing closed proprietary frontier models? Is there an opportunity if open weight models proliferate that you don't have to develop your own frontier model? So that's question one.

And the second one, a clarification, if I may, for Susan. You noted that in your prepared remarks that you plan to maximize '26 and '27 capacity, is that a

demand or a supply comment? Meaning, are you suggesting that Meta plans for internal use of all the capacity built through '27? Or are you saying that you will evaluate '28 and beyond builds based on demand for Meta products and services? Thank you.

Mark Zuckerberg: I can take the open source question. Let's see. So basically the question is, do we think that because there are some open weight models that we can just rely on those. I mean right now the open source models are not as strong as the frontier models. So no is the basic answer.

And then there's also just always the perpetual both policy debate and question around other companies' actions and whether that's actually a thing that we can -- that a company like Meta can rely on. And I think that that's very tricky.

So I think on both fronts, we believe we're going to be able to do better work, and we think that there's some risk in that reliance, I don't believe that that is the right thing to do.

I think that we're a company that -- if you look at Meta from -- take a step back on this. A lot of people view the surface layer of we build some social media apps and we have an ad business. We are really a full stack technology company.

We built our own data centers, our own infrastructure, our own chips, our own low-level software. When we got started -- like my background in engineering, like I wrote a lot of the systems code. A lot of the reason why Facebook worked was because it actually -- it just worked, right?

Like it literally worked when other social networks did not work fast and efficiently. And I think we just have the ability to build things that can be more personalized, more optimized, more efficient.

Some qualitative experiences are just not even possible for others to build because we go all the way down the stack. And it just seems to me pretty clear that having kind of sovereignty over building your own models is going to be an important part of that stack going forward which is why it is important for Meta.

But is also important -- is also why other people care about open source and why open source matters overall.

Because like other companies, even if they don't have the ability to build these models, do not want to have to just rely on a small number of closed labs. Like I think that there is a very important place in the world for there to be open source models.

And to be clear, that doesn't take away from the API opportunity or any of the things that I'm talking about because someone still needs to run the models and run inference on them and be able to kind of run those models efficiently

and have the compute to do that is going to continue to be a source of business advantage. So I don't really think that these things are necessarily at odds.

But when I look at like what a lot of discerning customers and companies are going to want around the world, they're going to want to know that they have control of their destiny and that they can trust the models that they're using and that their data is safe and that they're not sending it to competitors and all that.

So I think open source is going to be important. But I think for us, also building the models is going to be a critical part of it.

It also -- for what it's worth, and I know that like the models all get evaluated on like a common set of evals and they get chalked up to like, okay, some models within like a couple of points of another model or whatever.

But they all really do have different combinations of skills, kind of like people, right? It's like people have spikes in certain areas and have different personalities and are better and worse at different things.

And if you're trying to build personal superintelligence for people or you're trying to build a business agent for small businesses, that needs to have potentially some different skills than what the other labs are tuning. And just like being able to do the full stack work on our -- on like Instagram recommendations on our ad system is how we've gotten the results there over time.

I believe that building the kind of full stack model is going to be a lot of the advantage over time in how we build personal superintelligence agents, business agents, all these different use cases for all of the different customers that we want to serve.

So in addition to the distribution that we have and the ability to reach all these people and scale products that work, I think this is a lot of the durable advantage is that you build something that is kind of specific to that use case and excellent in those use cases.

And I mean look, I get that this is a big investment and it's a big bet. We see the technology working. We're happy with the trajectory of the lab.

I'm excited about the products that are coming, and we believe that this is going to be a big thing. So I mean I get that this is sort of a big bet across the industry. My personal bet is that the people who invest in this are going to be rewarded and feel very good over time.

Susan Li:

Ken, I'll just quickly answer your second question. When we refer to focusing on 2026 and '27 capacity, there are really two factors.

One is we are today, and expect to be in the sort of foreseeable future, supply<sup>1</sup> constrained, and that really includes our core business, too, where there are -- we still have numerous ROI positive places that we would put compute toward if we had it.

And the second, of course, is just the uncertainty over long-term constraints on building capacity, and we talked about some of the sort of the need to build out more of the supply chain earlier in my comments.

So I think beyond '27, when we look into '28 and further than that, the world is going to evolve a lot, our own internal demand will evolve.

We'll have turned over a lot of cards by then. And so when we think about planning today for '28, we're really focusing on flexibility. That's just giving us kind of the ability to have land and power.

But to really make the actual decisions about buying chips and other big-ticket items further in the future. So for now I think we know we have a lot of good use cases for capacity in '26 and '27, and that's really what we're building toward.

Chad Heaton: Great. Thank you for joining us today. And we look forward to speaking with you again soon.

Operator: This concludes today's conference call. Thank you for your participation. And you may now disconnect.

---

<sup>1</sup> Reflects correction of reference to "demand" during the earnings call.