

Home • Artificial Intelligence

by **Evan Schuman**
Contributor

Meta's local AI model prompts enterprises to rethink hardware-software cost trade-off

News

Aug 10, 2026 • 7 mins

The need for at least 24GB of VRAM to run Muse Glimmer is just one factor that makes its value challenging to determine.



Credit: Ascannio / Shutterstock

Meta rolled out a new 30-billion-parameter AI model optimized to run on a PC or Mac with a single GPU on Monday, offering a way to run always-on agentic workflows locally rather than relying on the cloud.

The company has dubbed it Muse Glimmer. However, its hardware demands, including a GPU with a minimum of 24GB of VRAM, could make it difficult to justify for deployment at scale.

Although analysts and consultants agree that there is a tremendous enterprise appetite for running models locally, determining whether switching more systems from cloud to local makes fiscal sense is much more complex.

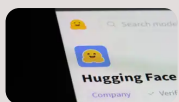
The hardware costs are tricky to calculate even today, with the VRAM needed depending on the particular applications to be run. But the far bigger consideration is that there is no way to determine what RAM costs will look like over the next 12-18 months, and there is an identical lack of visibility into how cloud prices might increase during the same timeframe.

That makes determining the better financial choice impossible.

Agents become capex, not opex

Noah Kenney [<https://www.linkedin.com/in/noah-m-kenney-27499a166/>], principal consultant at Digital 520, noted that since RAM costs have increased “exponentially” over the last 12 months, cloud AI providers are also going to have to increase their prices. Beyond that, IT needs to anticipate logistical issues; key questions to ask are, “How quickly can you scale? Can you even get the hardware?”

Recommended for you



What Nvidia's \$13B acquisition of Hugging Face means for AI model choice

By Evan Schuman • 8 mins

[https://www.infoworld.com/?post_type=post&p=4218324&utm_source=miso&utm_medium=related&utm_campaign=thumbnail_list]



The five important tools for controlling AI costs

By Matthew Tyson • 11 mins

[https://www.infoworld.com/?post_type=post&p=4217150&utm_source=miso&utm_medium=related&utm_campaign=thumbnail_list]



Nvidia to hike prices by 15%, on top of an even larger increase in July

By Evan Schuman • 4 mins

[https://www.networkworld.com/?post_type=post&p=4213165&utm_source=miso&utm_medium=related&utm_campaign=thumbnail_list]

“Meta just made agents a capital expense instead of an operating one,” Kenney said. “For two years, enterprises have been trained to rent intelligence by the token from someone else’s data center. Muse Glimmer runs the agent on a GPU you own, on the desk, with the meter switched off. That is a direct shot at the business model that cloud AI vendors are built on, and it comes from the one player with no cloud API revenue to protect.”

In its post [<https://research.meta.ai/blog/introducing-muse-glimmer-open-agentic-model>]

announcing the new model, Meta pointed out that it has aggressively slimmed it down to try to make it efficient and cost-effective.

“At full precision, a 30-billion parameter model would require over 55 GB of memory — far more than any consumer GPU offers,” Meta said. “We use quantization techniques to compress the model’s weights to approximately 4-bit precision, shrinking the language model to under 20 GB. This leaves enough headroom for the model’s working memory, its KV cache, the perception encoder for image understanding, and the speculative decoding drafter to run simultaneously within a 24 GB or 32 GB envelope. We validated that this compression introduces minimal to no degradation on agentic tasks.”

More options for enterprises

Mike Wilkes [<https://www.linkedin.com/in/eclectiqus/>], enterprise CISO at Aikido Security, said that Meta’s move is significant in that it starts to give enterprises more options.

Computerworld Smart Answers [Learn more](#)

Explore related questions

- [Why are enterprises moving AI workloads to private clouds?](#)
- [How does kernel fusion address memory inefficiencies in deep learning?](#)
- [What are the API pricing advantages of Meta Muse Spark 1.1?](#)
- [Which tools help developers build agents with the Google Development Kit?](#)
- [How can gradient checkpointing reduce memory usage during AI training?](#)

“The most important thing about Muse Glimmer is not that Meta has produced another capable model, it is that the economics and architecture of AI are beginning to move back toward the edge,” he said. “The financial comparison therefore becomes capital expenditure that can be amortized over several years versus an effectively perpetual per-token or per-request cloud operating expense.”

That means, he said, that an enterprise may rationally pay somewhat more for hardware if doing so gives it predictable AI costs, offline availability, control over model versions, freedom from sudden API pricing or access changes.

Independent cybersecurity and risk advisor **Steven Eric Fisher** [<https://www.linkedin.com/in/steveneric/>] also noted that the specs published by Meta don’t tell the full story.

The problem is that running locally versus in the cloud can generate a lengthy list of related expenses.

“Agentic workloads [in the cloud] can amplify consumption through reasoning, retries, tool calls, context growth, and evaluation, while local deployment [also] introduces hardware, power, lifecycle, support, and utilization costs,” he said, adding that even the RAM requirements need a lot of context.

“Meta’s stated 24GB and 32GB memory targets demonstrate that Glimmer can be loaded and executed on comparatively accessible hardware, but that is not the same as having sufficient capacity for meaningful agentic workloads,” Fisher said, pointing out that once other factors are considered, practical memory requirements can move beyond the 32 GB available on an Nvidia RTX 5090 GPU.

“In enterprise terms, this still places Glimmer primarily in high-end developer, data science, or dedicated AI workstations rather than the standard corporate desktop or laptop,” he said.

Better ROI not guaranteed

Justin Greis [<https://acceligence.com/talent/profiles/justin-greis/>], CEO of consulting firm Acceligence, agreed.

“I wouldn’t assume that moving inference from the cloud to the endpoint automatically produces a lower total cost of ownership,” he said, noting that variable cloud costs would be traded for the price of deployment, endpoint management, support, security, model updates, and potentially accelerated hardware refresh cycles for the local devices.

“Muse Glimmer is an important milestone because it makes local agentic AI technically viable. But technical viability and enterprise ROI are two different milestones,” Greis said.

"I think Meta has crossed the first one. I do not think they have fully crossed the second one yet."

However, Kenney argued that there are also other elements of the Meta rollout that make meaningful comparisons difficult.

"It is worth remembering that the local model is quantized, compressed to roughly 4-bit precision, while the cloud APIs you are comparing against typically serve full-precision models, so this is not a pure apples-to-apples cost comparison," he said. "The ROI question is not local versus cloud on price alone. It is also a question of whether a company is willing to use a quantized model for the specific task. A cheaper agent that needs more retries or human correction can erase its savings fast."

SPONSORED CONTENT by Databricks

11x AI Growth Needs a New Database

By Databricks

Aug 19, 2026

[https://jadserve.postrelease.com/trk?ntv_at=3&ntv_ui=8b46823b-fad4-4bb3-a756-75a12f9ce83c&ntv_a=VI4KACSx7APIOTA&ntv_fl=40evsNPao_XnlqhCYCxnjOqDNCNixMTvyWNPoK2MtbRUoPqxThJjFEUZuIDKJ9NBG8wVGRouo07GvbZPs3JtOh7qgH_Pcu9vuAoTHn5RULm8jF011TPlpRytOXrXAVOCHI WifBrFac4USxBtmZcONSMdATykIMVIAtoAEKRGRzKftO2da32wdn1KgcjdR9Kc42sAnCa9luTN7PIP ofDGC9SrRji97JxoICUVtH--iLzJH15sch36ecnGVNe4qXWu7a2KZoH-6DfThXNd2ohnkRvkEqI3F_OWt8BpOL9VSqIVWK6bA6TMXIE3910sAbdhMi-8x3-lt4260kPSjyckdkSzwAEFK9b0ptV5EDJMBn_sSMp5fviVNH60UKf9SDdrY9I9xDogtjYX00M3BFIOzuJSnqm4N-NMQ4yxGPsJv6zhdXziqcp7RSwW9JTYT xeP0Ypkg-8YaL-x6TqMmL-25Y3O2lrOF7iZ8hD6tLhB_p-Eu_4sK1lt6iE8SFUwqG63OC2n5LoaWp5SvX zhNEAtgB1kaKMUgtOwAQC2yeDtcMIAiRekB5wW-l_ZZZZ_DkzALvV-BAvsVChlzz3jVGrdp3G-A3JvFLg WTxr-tlpqtsiqHAo63ETEzQ-L_D5nw-KsfGpFNJ51kZXX9QuLjgrh-Krq1vuyuo034MrKn1QWxNxVwdDdN8b6adF_WivJ9GKFxRqrqr-j-XE-L8Dn6yJW6dNNtSQgDo8QInzAxu2pOg9_8rBtExNsoPJLTW5vEWrQ_EqbIXKjSX53ZJaHtN5TVUOoggDGwfX8aCOx_sk4khVkjSR1DeLklfJdmISqF81WeEnNiLF9pJ-w6_Mef aFYyPMD_BhFGoQrHHEbwOfW1XV4bD_LewOa2J-PbGj491YgHHnYMcGKa6kfpNDIdSmg8cnXepLuD8sVoiv7dCihrdIZfBQX0pvQJMHkrooZ8K1_CmoaOAztJrWlwrcKJPZdyZcyUp-KWBnKvYQfMWpII3-Pu2WsBJLIWIOzeXJ9IwF-yjDVfnPxqXyKpTfRknSAutMhssa_ujj8lkXB2-kkTcjuwhR2AvCPcttkoITByrQK oFCUEmnk47YyW9euwf-RB_wk62gVhCdXRB5PlwZkQvgtmWagt8QpgPyOw_-Zn7_nhP8EqIzz7ZSHf AwHa3S3IrdDVdAysQ2Wxk7zOcnZl2t8Z3RuCGWdgZBXyNSra6jl266JbTIViN-KWowQG8dfXUkE6Dg -ebi3c8dAnNiaN7KaCb6c999ZrL5FqeJjMU1QciNRNrkL7283nEvO1IRPjiBR3B4XQiEHcZpO7bO97BL QHLaLpfJ22uc1p8n3ZOUDt84vF7lwzewn5DMBtdwWc43rofKf4Vel1XrKOYQKX46nJhc5wIMSTSBKt7vpjqt5KRoyfJlrV7gC4QLaqqDHqWmljfTg6BtJeG6qDicQRoh1q-Rt47x3v3ZOL9sHhmRjgspJ-W--uskO Dh_zm426u_52oVJSXwFXoV8VWtGZ7C0plyJ5nj2QesDRVKr4tkbLwfYmMzCJQgEoWVM4QSNEt2sS nxyh9AqV8Y3HFXg=&ord=1773319450&ntv_ht=s9iiaG&ntv_rad=16&ntv_r=https://ad.doubleclick.net/ddm/trackclk/N2655160.4590155FOUNDRY/B36343267.453482612;dc_trk_aid=647534626;dc_trk_cid=262379015;dc_lat=;dc_rdid=;tag_for_child_directed_treatment=;tfua=;gdpr=0;gdpr_consent=;td=;dc_tdv=1]

In addition, a local model often neglects to include every service that a cloud provider typically delivers.

“The cloud vendor was quietly handling updates, scaling, reliability, and security patching across your whole footprint. Bring the model in-house and every one of those becomes your problem, multiplied by every machine running it,” Kenney noted. “Most enterprises that consume AI as a service have no muscle for operating a fleet of local models, and that cost rarely gets adequate consideration in ROI conversations. The GPU is cheap. Patching a thousand of them is not.”

Sanchit Vir Gogia [<https://greyhoundresearch.com/svg/>], chief analyst at Greyhound Research, also stressed that it can be difficult for IT to comprehensively anticipate the different cost variables.

“Finance leaders are right to feel skeptical. An engine bought for one employee burns capital whether or not it runs. Meta has shown that a thirty-billion-parameter agent can run on a single machine, and that is a real engineering result. What it has not shown is that such an agent works reliably at enterprise scale, or that a fleet of them can be operated safely. Model fit and production fit are different claims. A laptop must still run the employee’s actual job,” Gogia said. “The economics turn on the incremental hardware premium, refresh timing and actual utilization, measured against the price of the remote inference being displaced.”

On the other hand, **Arun Chandrasekaran** [<https://www.gartner.com/en/experts/arun-chandrasekaran>], a distinguished VP analyst with Gartner, said that he found it “very interesting that they have decided to release a smaller model that operates on the edge” and especially liked Meta’s use of the popular Apache license. But he would have preferred that they had released more than just a model.

“Enterprise customers are asking for a car and Meta is delivering an engine,” Chandrasekaran said. “They should have built something more like a platform solution.”

Artificial Intelligence[<https://www.computerworld.com/artificial-intelligence/>]

ROI and Metrics[<https://www.computerworld.com/roi-and-metrics/>]

IT Leadership[<https://www.computerworld.com/it-leadership/>]

NEWSLETTER

The latest AI updates, straight to your inbox

News, analysis, and insights for IT leaders navigating the risks and rewards of AI. A special series from the editors of CIO, Computerworld, CSO, InfoWorld and Network World

☐ By submitting your information, you agree to our [PRIVACY POLICY](https://foundryco.com/privacy-policy/) [<https://foundryco.com/privacy-policy/>].

SUBSCRIBE

About

[About Us](#)

[Advertise](#)

[Contact Us](#)

[Editorial Ethics Policy](#)

[Foundry Careers](#)

[Reprints](#)

[Newsletters](#)

[Add Computerworld as a Preferred Source in Google Search](#)

More

[News](#)

[Features](#)

[Opinion](#)

[How-To](#)

[Blogs](#)

[BrandPosts](#)

[Events](#)

[Podcasts](#)

[Videos](#)

[Buyer's Guides](#)

Policies

[Terms of Service](#)

[Privacy Policy](#)

[Cookie Policy](#)

[Copyright Notice](#)

[Member Preferences](#)

[About AdChoices](#)

[Your California Privacy Rights](#)

Our Network

[CIO](#)

[CSO](#)

[InfoWorld](#)

[Network World](#)

