

White Paper

For a Solid Return on Investment with AI, Consider the Many Ways to Purpose-Fit Your Infrastructure

Sponsored by: Supermicro and AMD

Peter Rutten Madhumitha Sathish
December 2025

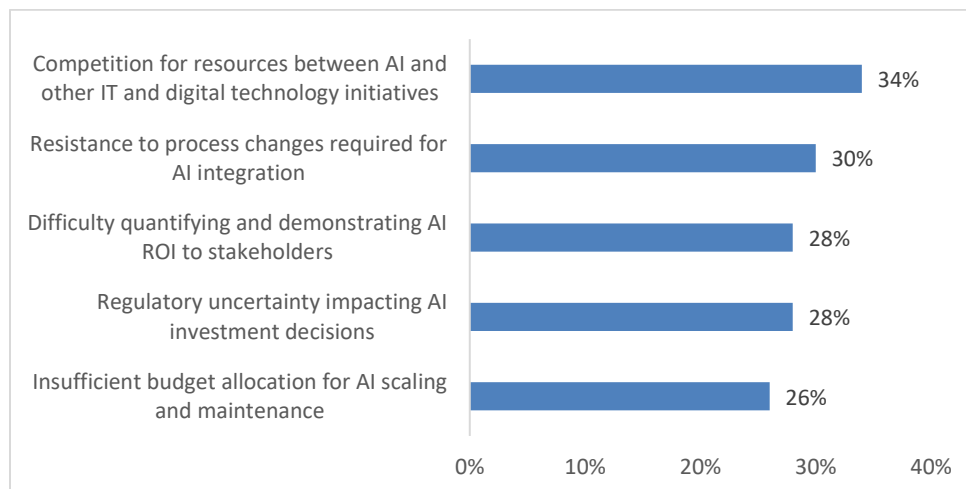
IDC OPINION

Enterprises face significant hurdles in the quest for new AI functionality in their business processes and products. According to IDC research (*FERS Wave 7*, September 2025), 50% of respondents worldwide say that less than half of their AI-related projects have delivered measurable business outcomes, and only 11.4% report that they are obtaining measurable business results from more than 75% of their AI projects. Figure 1 shows how respondents ranked the challenges.

FIGURE 1

Challenges preventing organizations from realizing the full potential of AI

Q. *What are the top 2 challenges preventing your organization from realizing the full potential of your AI investments?*



Source: Future Enterprise Resiliency & Spending Survey Wave 7 IDC, September 2025, N=888

The competition for budget, the difficulty of demonstrating an ROI, and insufficient budget allocation for AI scaling and maintenance speak directly to the fact that, despite the anticipated benefits from AI, cost continues to be a major hurdle.

IDC data also shows that infrastructure (both capex for hardware and opex for cloud) is the largest cost factor in enterprise AI initiatives. When asked what their biggest cost concerns are regarding the development and deployment of AI, more than 60% of businesses say it is the specialized AI infrastructure that is required.

IDC believes that despite this concern, enterprises can build a strong ROI for their AI initiatives, especially if they understand how to leverage different infrastructure solutions for different AI use cases. Some of the factors that play a role in such a “purpose fitting” are:

- Who decides what a relevant AI use case is? This question deals with the lack of IT involvement, hence the unpredictability of infrastructure cost.
- How will the organization consume AI? Will they develop or deploy the AI model themselves? On their own hardware or in the cloud? Or will they use SaaS or API access?
- What kind of AI model is required? The current focus on agentic AI and large-scale generative AI training has overshadowed the fact that there are many use cases that require less compute-intensive infrastructure (e.g., for inferencing, edge deployments, and even the use of workstations for AI development and deployment).
- How will the business obtain the model? There is a trade-off between a business developing its own model, finetuning existing models, and using existing models.
- Have the biggest factors that impact AI infrastructure needs been considered (i.e., the type of AI model, the number of parameters, training data volume, model accuracy, time to value, query response time, and query size).

Taking these factors into account, enterprises can develop a spectrum of AI options to match their AI use case to a purpose-fit infrastructure solution. This paper will discuss a rough framework for purpose-fitting AI infrastructure to its use case, enabling enterprises to create a robust ROI for their AI projects.

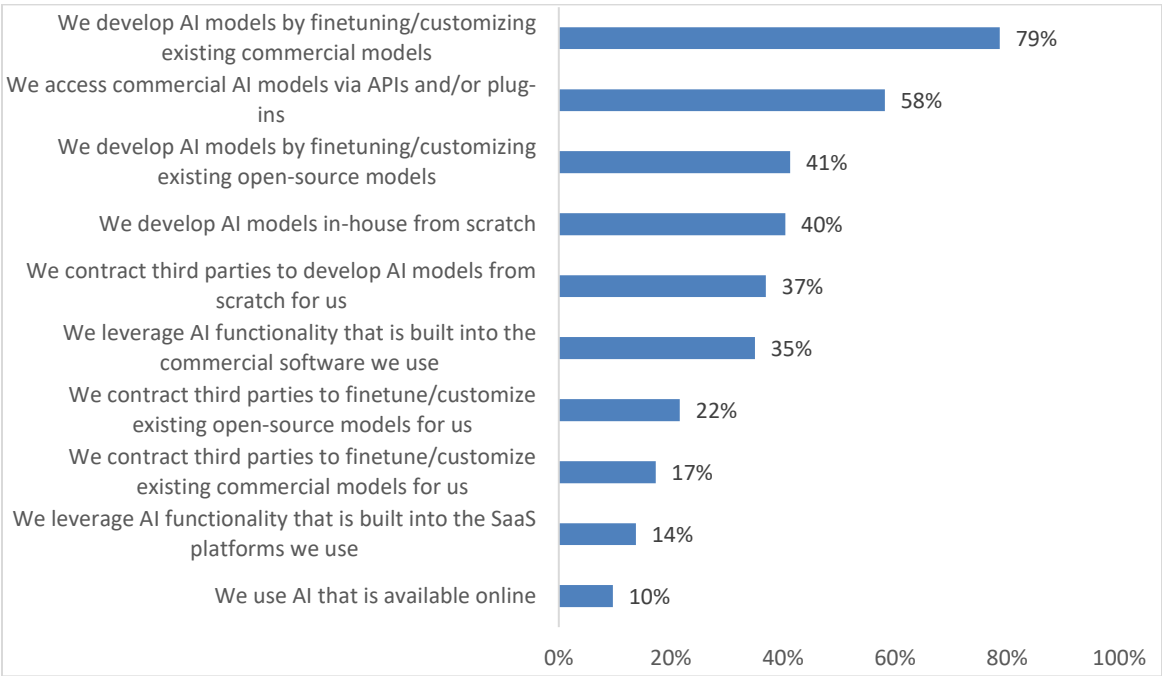
SITUATION OVERVIEW

IDC research has found that there’s significant variety in how businesses bring AI functionality into their products and processes.

FIGURE 2

The most common approach is finetuning/customizing existing commercial models

Q. Which of the following approaches do you use to leverage AI models for your business processes?



Source: IDC’s AI View Survey, November 2025

Finetuning or customizing an existing model is the most popular approach, and more than three-quarters of enterprises take this approach. The second most common approach is to access commercial AI models via APIs, and the third is to finetune/customize existing open source models. At the same time, 40% of organizations say that they develop some AI models from scratch. More than a third leverage AI functionality that is built into commercial software, and 14% use AI built into a SaaS solution.

Each of these approaches has its own requirements for the type of infrastructure needed, with the main divider being whether it involves AI training or AI inferencing. Infrastructure options for inferencing range from PCs and workstations to edge servers to datacenter or colocation servers to public cloud instances. For finetuning and

customization, workstations, servers, and cloud will be relevant. For training a model from scratch, workstations (in certain cases), servers, and cloud are suitable.

Despite this variety in approaches, when asked what their biggest cost concerns are, more than 60% of businesses said it is the specialized AI infrastructure, while a little more than half said it is AI model development. About a third answered that cloud resource costs are their biggest concern, and slightly fewer answered data acquisition and preparation.

FIGURE 3

Top cost concerns ranked in order of importance

What are your organization's top cost concerns with regard to AI development and deployment?

1	Specialized AI Infrastructure
2	AI Model development
3	Cloud resources
4	Data Acquisition and Preparation
5	Data cleaning, model retraining, and testing
6	Data storage and management
7	Power and infrastructure maintenance
8	Integration of AI Models into Existing Systems
9	AI Model Updating and Maintenance

Source: IDC's *AI Processor Study*, 2025

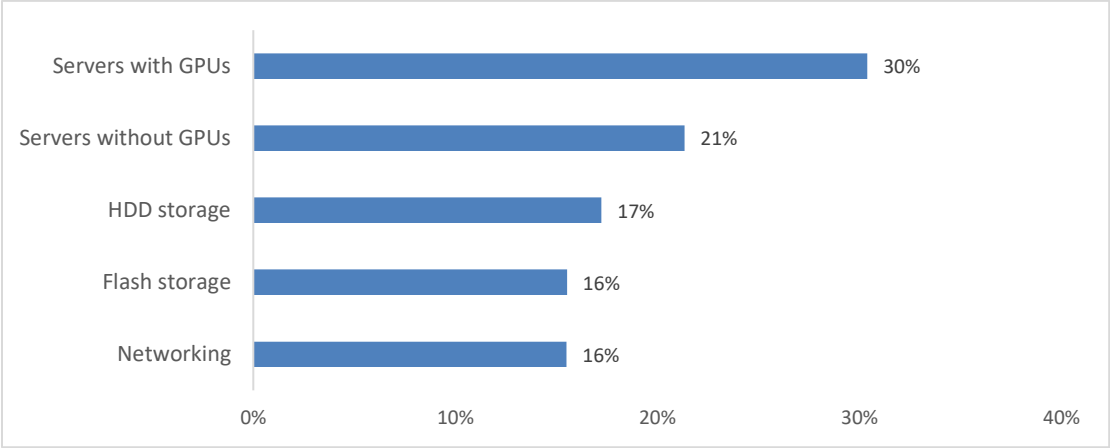
Nevertheless, one in three organizations says that they do not do a full AI infrastructure cost assessment when preparing for an AI initiative. IDC believes that enterprises can build a strong ROI for their AI initiatives if they know how to intelligently determine and manage the AI infrastructure costs by choosing their infrastructure wisely. Fortunately, there are several levers with which businesses can control their AI infrastructure choices.

To get a better idea of how specialized AI infrastructure fits into the overall IT budget, IDC asked 700 worldwide respondents (IDC's *AI View Survey*, December 2025). The results show that IT budgets for AI consist, on average, of 47% for hardware capex (not including PCs and workstations) and 53% for "everything else." The hardware portion is distributed as shown in Figure 4.

FIGURE 4

Servers with GPUs are the largest AI hardware budget item

Q. When dividing your total IT infrastructure budget for AI into hardware and everything else, what percentage of the hardware budget do the following categories represent?



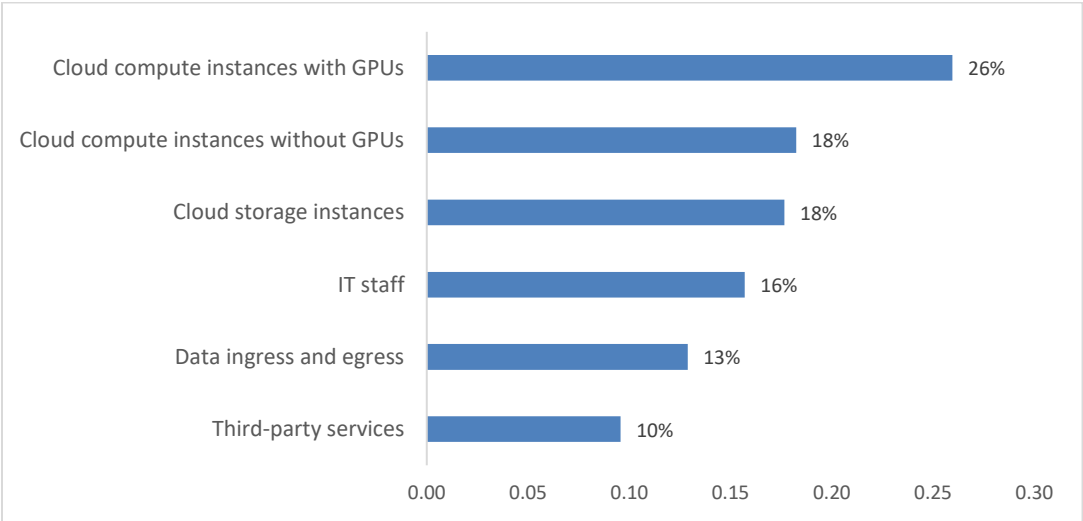
Source: IDC's *AI View Survey*, November 2025

The budget distribution for the everything-else category is shown in Figure 5.

FIGURE 5

Cloud compute instances with GPUs are the largest AI non-hardware budget item

Q. When dividing your total IT infrastructure budget for AI into hardware and everything else, what percentage of the everything-else budget do the following categories represent?



Source: IDC's *AI View Survey*, November 2025

There appears to be no silver bullet for reducing cost simply by moving AI development and deployment to any particular location. In the cloud, a proof of concept (POC) may help contain cost, but once a production-ready AI model needs the scale, cloud can become as expensive as an on-premises environment. The edge, for inferencing close to the data source or near the end user, trades off various costs, such as higher relative footprint cost (due to a lack of economies of scale) but lower data ingress and egress costs. That said, when location is viewed as more or less a constant, there are still very significant strategies to deploy the right infrastructure while taming costs.

IDC has found that almost half of organizations believe that, in the long run, AI will displace high-performance computing (HPC) for science projects or HPC will converge with AI. Organizations that have traditionally been running HPC, either in a datacenter or in the cloud, tend to have an easier ramp into AI due to some infrastructure similarities. They may also be able to contain AI infrastructure costs by shifting certain projects from HPC to AI on the same platform, which is not trivial but doable with the right skill sets.

Practical considerations that impact AI infrastructure costs

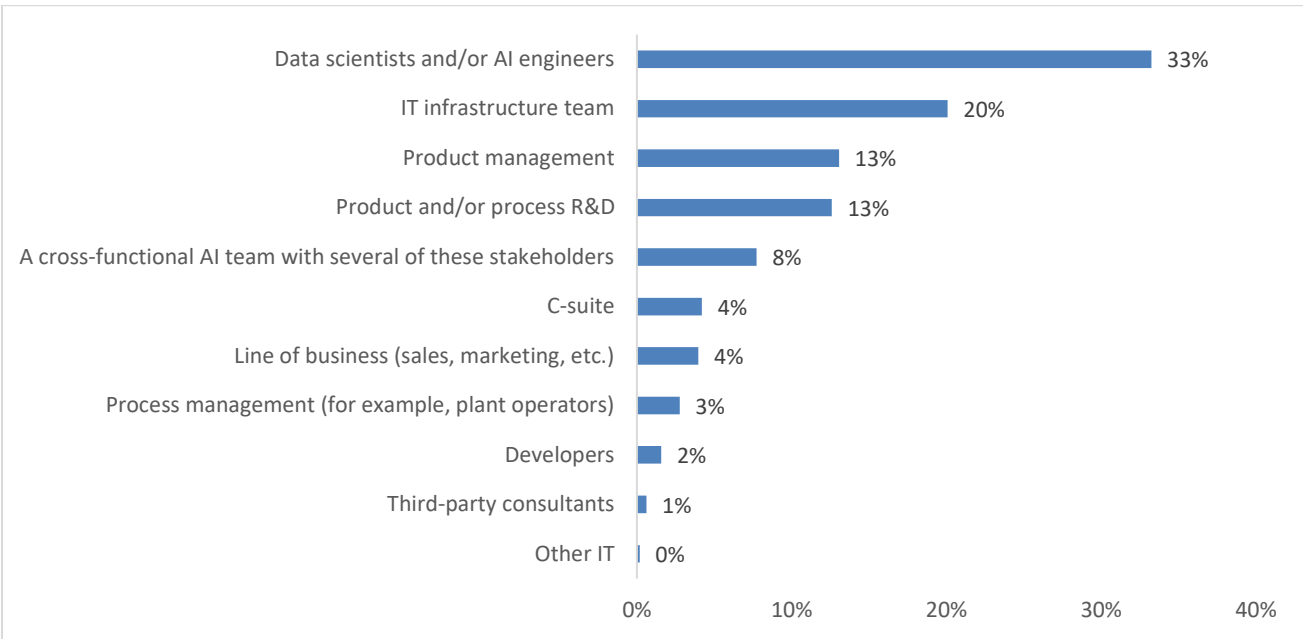
1. Who decides what a relevant AI use case is?

IDC research has found that only 8% of enterprises have a working group with representatives from across the organization that decides which AI initiatives are relevant for the business. About one in three let data scientists and AI engineers define AI initiatives; for 13%, an AI project starts with product management, and one in five allows IT to launch such initiatives. This means that less than a third (those that have cross-functional teams plus those that enable IT) of organizations involve IT at all during the conceptual stage of an AI initiative.

FIGURE 6

Where AI initiatives most often originate

Q. Where in your organization do GenAI initiatives most often originate?



Source: IDC's *AI View Survey*, November 2025

Ideally, a cross-functional team would always be the start of an AI project, with full involvement from IT. Unfortunately, very few organizations take this approach. The risks of not involving IT from the start are distinct: misalignment about what is technically feasible with the AI initiative; wrong estimates of the time, the skill sets, the technology investments, and other costs required to bring the AI initiative into production; incorrectly formulated or infeasible use cases; and, ultimately, miscalculation of the ROI.

IT teams are key stakeholders in the concept of purpose-fitting infrastructure for the variety of AI use cases. They can assess whether a business needs a good workstation or a big server cluster for the AI use case; whether GPUs are needed or strong CPUs will suffice; and whether the development or deployment should run on premises, in the cloud, or at the edge. Many of the factors that go into these decisions are technology-related. Additionally, many AI use cases improve IT itself and can only originate with the IT team.

Recommendation: Start every AI project with a cross-functional working group that includes the IT team.

2. What kind of AI model do you need?

Today, AI often means generative AI (based on larger language models) or agentic AI, which is an emerging workload that can be described as an autonomous AI system that can independently plan, reason, and take actions to achieve complex goals with minimal human oversight. But deep neural networks (DNNs) are still successfully developed and used for many AI use cases, as are more traditional predictive machine learning approaches.

Figure 3 shows the AI models that enterprises say they are using today; 20–30% of organizations actively run traditional machine learning models, such as linear discriminant analysis, linear regression, decision trees, support vector machines, random forest, and Naïve Bayes. Close to half of organizations have deep neural networks in production.

The message is that there are many AI types in use today, and not all require major capital expenditures. Organizations should not assume that they always need an LLM to achieve the capabilities they need; there are many other AI approaches that will allow them to gain the desired AI functionality with lower AI infrastructure or cloud expenses and thus improve the ROI.

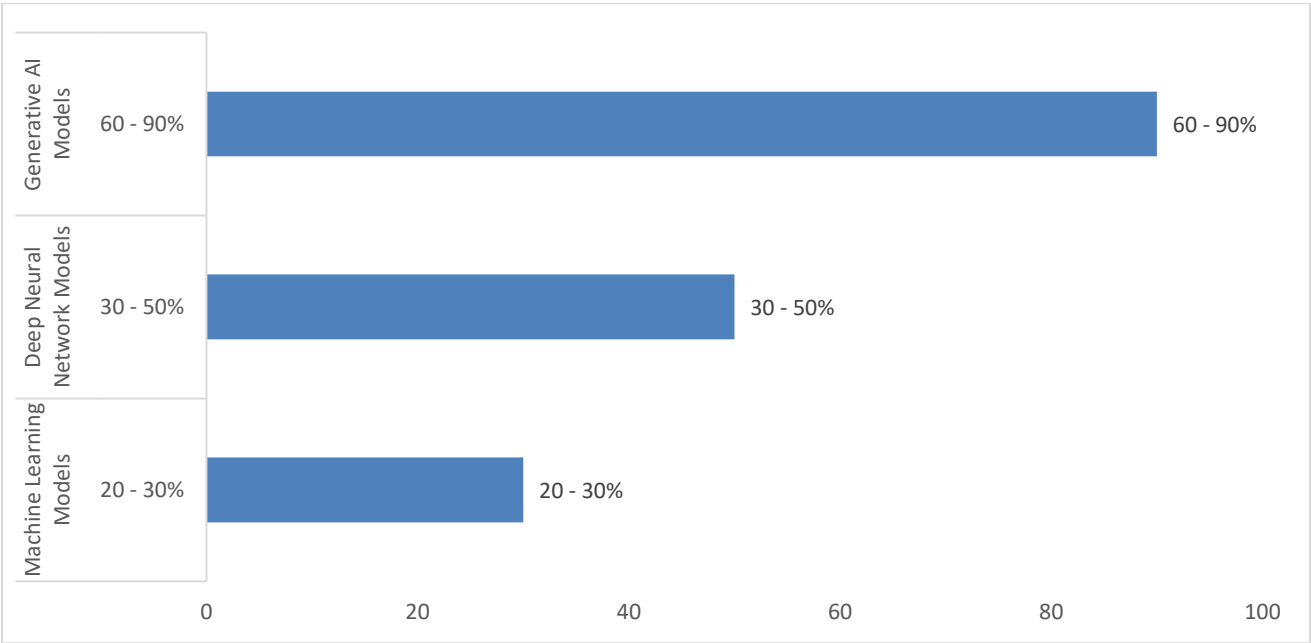
Understanding what the options are from an algorithmic perspective is critical. For example, most machine learning algorithms are not based on parallelization of the workload, and GPUs (or even servers) are therefore not always required to run them. On the other hand, agentic AI does require GPUs, as well as more east-west networking and faster storage reads and writes.

Recommendation: Consider the wealth of both current and older AI model types for your use case and the technology that each requires.

FIGURE 7

Many AI models in use today require much less AI infrastructure

Q. What types of AI models does your organization currently use?



Source: IDC's *AI Processor Study*, 2025

3. How do you obtain the AI model?

Where an AI model is sourced determines a lot about the spending involved. As Figure 1 shows, a large majority of businesses develop AI models by finetuning or customizing existing commercial models. They will have licensing costs and some training costs for the retraining they are doing. Next are those that simply use plug-ins or APIs in an existing commercial model — they only have licensing costs. Figure 8 shows the infrastructure cost items related to each approach.

FIGURE 8

Model acquisition approach and associated cost

	Licensing	Customizing	Training	Inferencing	Access fee	Consulting
AI that is available online					\$	
AI functionality that is built into the SaaS platforms					\$	
Third parties to finetune/customize existing commercial models	\$			\$		\$
Third parties to finetune/customize existing open-source models				\$		\$
AI functionality that is built into the commercial software	\$			\$		
Third parties to develop AI models from scratch				\$		\$
Develop AI models in-house from scratch			\$	\$		
Develop AI models by finetuning/customizing existing open source models		\$	\$	\$		
Access commercial AI models via APIs and/or plug-ins	\$					
Develop AI models by finetuning/customizing existing commercial models	\$	\$	\$	\$		

Source: IDC, 2025

Each of these approaches will have trade-offs. Online AI is cheap but generic, SaaS cannot be customized but requires few technology skill sets, finetuning existing models creates differentiation but requires AI training capabilities, using third parties can be expensive, and developing AI in-house from scratch is complex and can be expensive but will yield the most tailored AI model for the desired use case.

In general, it is fair to say that infrastructure costs for AI development and deployment increase with the uniqueness of the model. A bespoke, internally developed model will cost more than a finetuned existing model, which will cost more than SaaS or an API-accessed model. The uniqueness of the model also influences the potential business value. A unique model will differentiate the organization to a much higher degree than a licensed model, SaaS, or API-accessed models, with possible competitive advantages.

Recommendation: Balance the uniqueness of the model against the infrastructure requirements to develop and deploy it. Aim for maximum uniqueness within the infrastructure budget.

4. Have you assessed the seven AI infrastructure cost factors?

The factors related to AI infrastructure that need to be considered in the organization's ROI for an AI initiative are:

Type of AI model

Is it a traditional machine learning model, generative AI, or agentic AI? The latter two will mean significantly more specialized infrastructure due to the size, complexity, and nature of the model, which translates directly into the type of infrastructure and hence cost. This cost can differ exponentially from one model type to the next. Many machine learning models can be developed and deployed on as little hardware as a PC, a

powerful workstation, or CPU instances in the cloud. Smaller DNNs are developed on workstations or in the cloud and then either deployed on datacenter servers, in the cloud, or at the edge.

Number of parameters

There is a linear correlation between the number of parameters an AI model is built with and the infrastructure that is required to train that model. The greater the number of parameters, the more compute that is required, and hence the more expensive the model will be to train. IDC research shows that the average AI model that is used in enterprises has about 80 billion parameters, a far cry from the models with trillions of parameters that hyperscalers develop, and as such, much more manageable. However, as agentic AI sees more adoption, it is expected that model sizes in terms of parameters will go up at enterprises.

Vendors and academics are conducting extensive research to enable AI models that have increasingly lower parameter counts while delivering comparable results for the purpose of lowering the compute requirements.

Training data volume

Today's generative AI models use large and growing amounts of data. In medical imaging, thousands of examples are used for training, while LLMs can easily use trillions of words or billions of images. Data volume depends on the type and complexity of the model, but GenAI (and agentic AI) models need exponentially more data that is processed in parallel on GPUs than machine learning models, which process linearly on CPUs.

Data used for LLM training continues to grow as the tasks that AI is designed to execute become more complex and sophisticated. Writing creative marketing materials is more complex than classifying a set of products, for example. Being able to respond to inquiries of any kind with any form of output requires a general-purpose AI that is trained on nearly all available data on the web and beyond. A more tailored AI model, on the other hand, built around an organization's own data for a specific business process or client application, will use much less data, and traditional machine learning uses simpler algorithms that learn from mostly structured data in smaller amounts and then make predictions.

A model that linearly processes a small amount of data for training will complete the process faster and with less infrastructure than a model that processes massive volumes of data in parallel. The latter will need many training iterations on compute and storage that is properly sized for such a task, which translates directly into cost. Retraining an existing model, on the other hand, with a smaller amount of custom data or with RAG, is typically less compute-intensive and therefore more affordable.

Where the data resides plays a role, too. If the model is developed in the cloud and the data is in the cloud, ingress and egress fees will be avoided. Similarly, if the model is on a PC/workstation or datacenter servers and the data is in the same location, it will be more costly to use cross-deployment scenarios (e.g., when the model is processed in the cloud and the data is in the datacenter or vice versa).

Model accuracy

This is a sometimes-overlooked factor that determines the amount of infrastructure required to train an AI model. Scientists have concluded that it is nearly impossible to achieve perfect or near-perfect accuracy with AI models because of the amount of infrastructure required to train them to such a high level of accuracy. GenAI and DNNs are probabilistic AI systems that are trained on complex and ambiguous data. The only way to overcome these limitations is by employing exponentially longer and more complex training runs on ever more infrastructure in absurdum.

For practical purposes, it is fair to say that the lower the accuracy requirements are for an AI model, the less infrastructure will be required to complete its training, hence the lower the associated infrastructure cost. Some AI use cases will require very high accuracy (e.g., cancer diagnosis) while others will not (e.g., recommendation engines, which end users tolerate being fairly bad at predicting their tastes or preferences).

Time to value

The time that it takes to develop an AI model is directly correlated with the infrastructure cost, especially with on-premises infrastructure. If it is deemed acceptable for model training to take six months, for example, the organization will need much less infrastructure than if the model needs to be completed in one month. In the latter scenario, much more parallel processing compute power will be needed.

A major factor in time to value is the deployment scenario. If the model is to be developed in the datacenter, the ready availability of the right infrastructure will be a critical factor. If infrastructure still needs to be procured, configured, and installed, this will greatly lengthen the time to value — a complication that does not exist with the cloud or simple-to-procure PCs/workstations.

The size of the model plays a role here, though. Very large models cannot be trained on pared-down infrastructure, even if the length of training time can be extended. But in general, there is a calculation to be made between allowing more training time, and therefore spending less on infrastructure, but also generating value later.

Query response time

Once a model is live in production, end users will start interacting with it through queries. For every AI use case, the required query response time will vary. Some use

cases demand a near-instant response (e.g., chatbots, where unsatisfactory latency can cause end users to abandon the conversation). Other use cases can tolerate slower responses (e.g., a video-generating model, where the end user will not mind waiting several seconds for a result).

The number of concurrent users sending queries affects the response time. The more concurrent queries, the greater the infrastructure requirements. If the AI model is designed to serve, say, 300 end users inside the organization on a daily basis, the inferencing infrastructure can be much smaller than if the model is intended for hundreds of thousands of consumers who may be interacting with it on an hourly basis. The deployment of the model can also influence query response time. Cloud response times are sometimes longer than response times from optimized datacenter infrastructure. On the other hand, if the cloud datacenter is local to a large portion of users, the cloud may be faster.

Query size

Finally, an important factor is the anticipated query size — is the end user submitting a few words or a 1MB image for the AI model to process? The required processing power for the latter, especially when multiplied by the number of concurrent users, can significantly raise the AI infrastructure requirements and, therefore, cost.

New infrastructure requirements have arisen with agentic AI. Today's LLMs have no long-term memory — they are stateless and need to use large amounts of tokens to remember previous sessions, which is referred to as their "context window." As soon as the text goes beyond the model's context window, it "forgets" and can no longer use it to respond. The initial solution that model developers created has been to make the context window larger, but this comes at a significant computational cost and can lead to very slow inquiry responses.

AI agents need long-term memory alongside their short-term memory, which is different from stateless LLMs. They need to learn from previous tasks, retain information, and maintain context. Without that, the model cannot deliver consistent outcomes across different sessions. The use of KV cache can speed up inference by storing and reusing previously computed keys (K) and values (V) in the GPU's caches, rather than recomputing them for every new token. The model retrieves them from the cache, which dramatically improves performance. If the cache is overwhelmed, new techniques allow the keys and values to be offloaded into CPU memory or storage devices. In short, KV cache greatly reduces the infrastructure requirements, and therefore cost, of agentic AI.

5. Develop a spectrum of options to contain AI infrastructure costs

Enterprise AI infrastructure strategies can be seen as an AI spectrum, with budget being allocated depending on the seven factors for a specific AI use case. This spectrum assumes that the model is beyond POC.

Figure 13 illustrates the spectrum of AI infrastructure cost in relation to the seven factors that impact the infrastructure needs discussed above.

FIGURE 9

Spectrum of options to contain cost



Source: IDC, 2025

The blue zone

This blue zone of the spectrum would be relevant for an AI use case with most of the following characteristics:

- Traditional machine learning/small GenAI or DNN/finetuning of a small existing model
- Small number of parameters
- Small amount of data for training
- Less accuracy acceptable
- Slow time to value acceptable
- Slow query response times acceptable

This AI use case could run in production on a CPU-based, air-cooled system; a PC or workstation; or a small cluster. The CPUs would need to be cutting-edge, leveraging higher core counts and the latest fabrication technologies for performance. IDC research has found that more than half of all organizations say they can achieve less-than-two-second response times with systems that run on CPUs and that field <5,000 concurrent end-user requests.

For the POC stage of an AI model, infrastructure requirements can be significantly lower. Even a model that, once in production, resides in the red zone could, during the POC, reside in the blue zone.

The green zone

Toward the center of the spectrum, in the green zone, AI use cases with some or more of the seven factors dialed up to mid-way become relevant (e.g., the model is small but requires very fast response times to queries (real-time scenarios) or the model is small). Query response times can be slow, but time to value is critical because the organization wants to roll it out in 12 weeks. In this zone of the spectrum, CPUs with built-in acceleration become useful as well as lighter co-processors. Air-based cooling will still be sufficient, but small or mid-size clusters may be required.

Packaged solutions are becoming popular in this zone (and further toward the red zone), where the solution merges enterprise storage with accelerated computing (including GPUs, DPUs, and networking) as well as a modular, enterprise-ready software platform for building, customizing, and deploying large language models, to create an "AI ready" data foundation. These solutions turn legacy storage into a platform for speedy and efficient delivery of data for GenAI, agentic AI, and RAG, reducing latency and complexity. The core idea is to bring the compute to the data and process it directly inside the storage layer using GPUs.

The red zone

In this region, AI use cases reside that have several or all of the factors dialed up to high levels, showing the following characteristics:

- Generative AI/DNN and agentic AI
- Large number of parameters
- Large data volume
- High accuracy required
- Fast time to value required
- Fast query response times required

We are now in a territory where high-end CPUs, combined with cutting-edge GPUs, are integrated into, for example, a liquid-cooled rack-scale system with fast interconnects.

Note that in all scenarios, several additional considerations are important:

- Is there more than one AI use case in development? IDC has found that, on average, organizations develop five AI use cases simultaneously. This must be built into the infrastructure needs projection.

- How will the AI use case evolve over time? If there is an anticipation that the use case will grow from 10,000 to 100,000 users in 12 months, this has infrastructure ramifications that must be accounted for.
- How fast will generational updates be required for the model? AI models are constantly being improved, expanded, and retrained. Here too, infrastructure impact must be accounted for.

Recommendation: Organizations might benefit from a quick assessment using the sheet in Figure 14, circling the numbers 1–5 for each factor, and adding them up. If the total is lower than 14, the infrastructure cost can be limited; if it is higher than 28, the infrastructure cost will be significant.

FIGURE 10

Quick assessment of AI infrastructure cost based on seven factors

	Simple				Complex
Model Complexity	1	2	3	4	5
	Small				High
Number of Parameters	1	2	3	4	5
	Small				Large
Training Data Volume	1	2	3	4	5
	Low				High
Model Accuracy	1	2	3	4	5
	Long				Short
Time to Value	1	2	3	4	5
	Long				Short
Query Response Time	1	2	3	4	5
	Small				Large
Query Size	1	2	3	4	5

Source: IDC, 2025

CONSIDERING SUPERMICRO WITH AMD

Supermicro and AMD work with a vast ecosystem of partners to offer different alternatives and options to support a TCO that fits their needs best. They aim to demystify the complexity caused by the variety of options available to enterprises and help them plan their AI projects faster and better.

Supermicro has over 30 years of experience in data center and edge infrastructure, focusing on solution development in collaboration with technology partners. The company is recognized for early adoption and integration of AMD-based CPU and GPU technologies within the United States. Supermicro's server platforms have achieved more than 50 server performance world records and over 70 business application world records, attributed to their modular data center building block architecture and a diverse range of air- and liquid-cooled systems. Customers utilize Supermicro, AMD, and ecosystem partner solutions for reliable, high-performance platforms supporting AI workloads across various deployment scales.

CHALLENGES/OPPORTUNITIES

AMD, which designs and manufactures high-performance CPUs, GPUs and DPUs, and Supermicro, which builds servers and rack systems, are both prominent vendors in the AI infrastructure space. With their partnership, they have succeeded in carving out a differentiated role for themselves that helps businesses succeed with their AI projects through tailored processor and server solutions that build on a vast and open ecosystem of partners. This is extremely important in a market that has been trending toward closed systems and that encourages vendor lock-in up and down the stack. Both vendors are deployment agnostic, offering their solutions for on-premises, edge, and cloud scenarios, and AMD also provides many solutions on the PC and workstation market, offering an end-to-end approach for every imaginable AI use case type and size. This is, IDC believes, where the market is heading, and working together, AMD and Supermicro have developed some of the most versatile, powerful, and well-tailored solutions available today.

Are these two vendors without challenges operating in this marketplace? Obviously not. The competition is fierce, with dozens of server vendors and several mammoth semiconductor vendors vying for market share. What's more important is that technology trends are advancing so fast that keeping up, let alone staying ahead, is a tremendous task. Yet this is also where the opportunities lie. Predicting technology trends and offering the right solution at the right time, not six months too soon or six months too late, is a skill that these two partners have finetuned.

CONCLUSION

Enterprises face a complex landscape when it comes to deploying AI initiatives, with infrastructure costs and strategic decisions playing a critical role in achieving a solid return on investment. The diversity of AI use cases — from traditional machine learning to generative and agentic AI — requires a nuanced approach to infrastructure that purpose-fits the specific needs of each project. Key factors, such as the type of AI

model, data volume, accuracy requirements, time to value, query response time, and deployment location, must all be carefully evaluated to optimize costs and performance.

A cross-functional approach involving IT from the outset is essential to aligning AI initiatives with technical feasibility and cost expectations. By leveraging a spectrum of infrastructure options — from PCs and workstations to edge and datacenter servers, as well as cloud resources — organizations can tailor their AI infrastructure to match their unique use cases and scale requirements. This strategic alignment enables enterprises to manage costs effectively while maximizing the business value and competitive advantage of their AI projects.

Ultimately, success in AI infrastructure planning hinges on understanding the interplay of these factors and continuously adapting to evolving technology trends and business demands. Organizations that adopt this purpose-fit mindset will be better positioned to unlock the full potential of AI and realize measurable business outcomes.

ABOUT IDC

International Data Corporation (IDC) is the premier global provider of market intelligence, advisory services, and events for the information technology, telecommunications, and consumer technology markets. With more than 1,300 analysts worldwide, IDC offers global, regional, and local expertise on technology, IT benchmarking and sourcing, and industry opportunities and trends in over 110 countries. IDC's analysis and insight helps IT professionals, business executives, and the investment community to make fact-based technology decisions and to achieve their key business objectives. Founded in 1964, IDC is a wholly owned subsidiary of International Data Group (IDG, Inc.).

Global Headquarters

140 Kendrick Street
Building B
Needham, MA 02494
USA
508.872.8200
Twitter: @IDC
blogs.idc.com
www.idc.com

Copyright Notice

External Publication of IDC Information and Data — Any IDC information that is to be used in advertising, press releases, or promotional materials requires prior written approval from the appropriate IDC Vice President or Country Manager. A draft of the proposed document should accompany any such request. IDC reserves the right to deny approval of external usage for any reason.

Copyright 2025 IDC. Reproduction without written permission is completely forbidden.