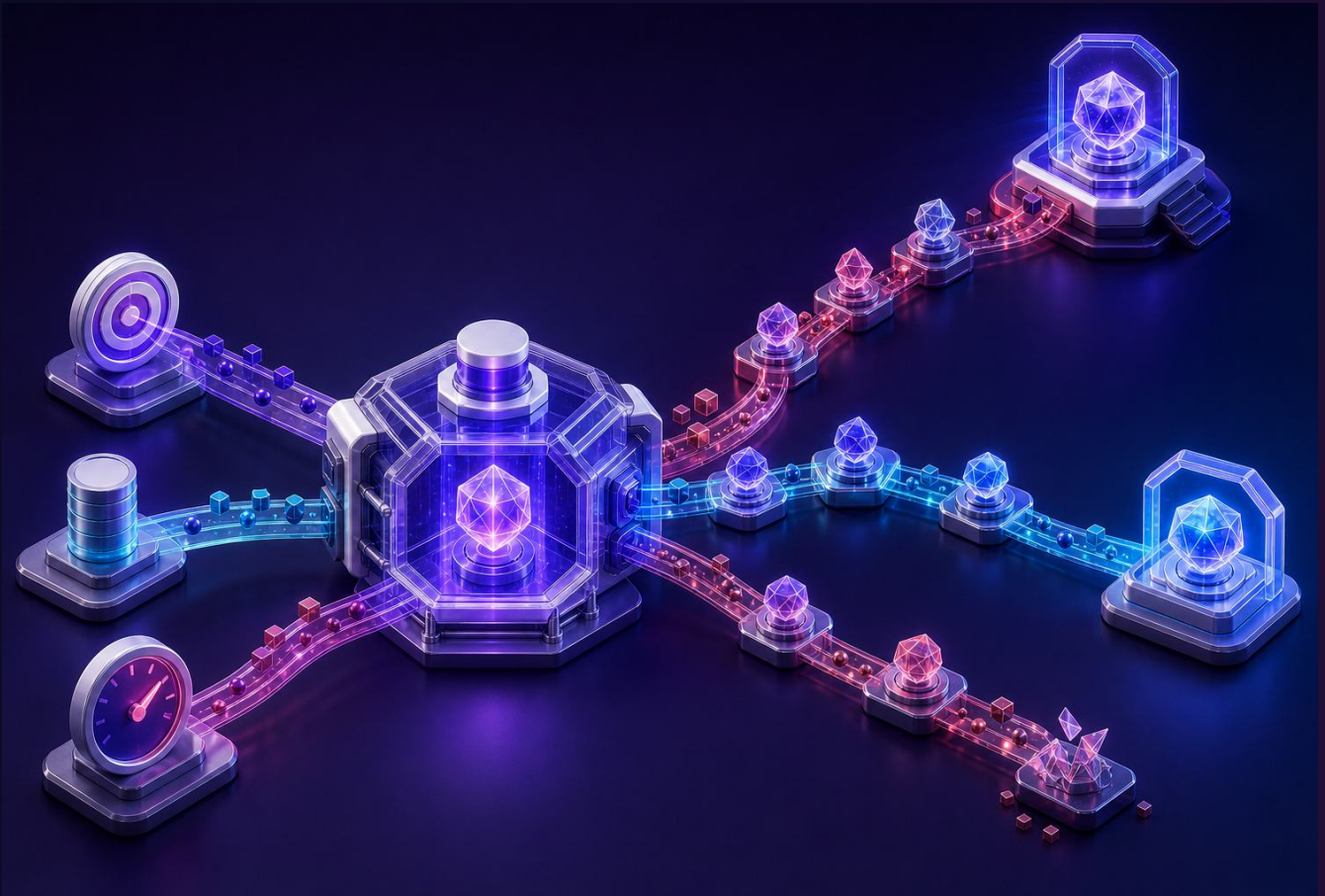


The Best AI Agent for Your Store Isn't the One at the Top of the Leaderboard

Stress-testing AI agents on real e-commerce scenarios shows why efficiency beats raw accuracy, why every model can be breached, and where the “just use Claude” hype breaks down.



Key takeaways



Choose a model by the **balance of accuracy, price, and speed** (the “Pareto front”), not by the top of the leaderboard. The most accurate model can be **14x pricier and 8x slower** for only about three extra points of quality.



The model that suits you personally is **not automatically right for your business**, where employee count, token volume, and speed requirements change the economics.



Match the model class to the task: a **fast model for live support chat**, a **slower and sharper one for an overnight report**.



More context is not always better. Feed it deliberately and in sequence so errors don't pile up in the middle of a long task.



Combine models in one pipeline: a **heavyweight model** for the hard judgment call, a **cheap, fast one** for phrasing the reply.



Every model can be breached, so put **deterministic, rule-based gates** between untrusted content and any privileged action. **No model will do this for you at any price.**



Anthropic models shine at **soft, empathetic work** but lag on **strict, structured business processes**; Fable 5's over-cautious filters make it unusable in production.



Open-weight models you run on your own servers ease **GDPR and data-sovereignty** concerns for European merchants, at the cost of owning the infrastructure.

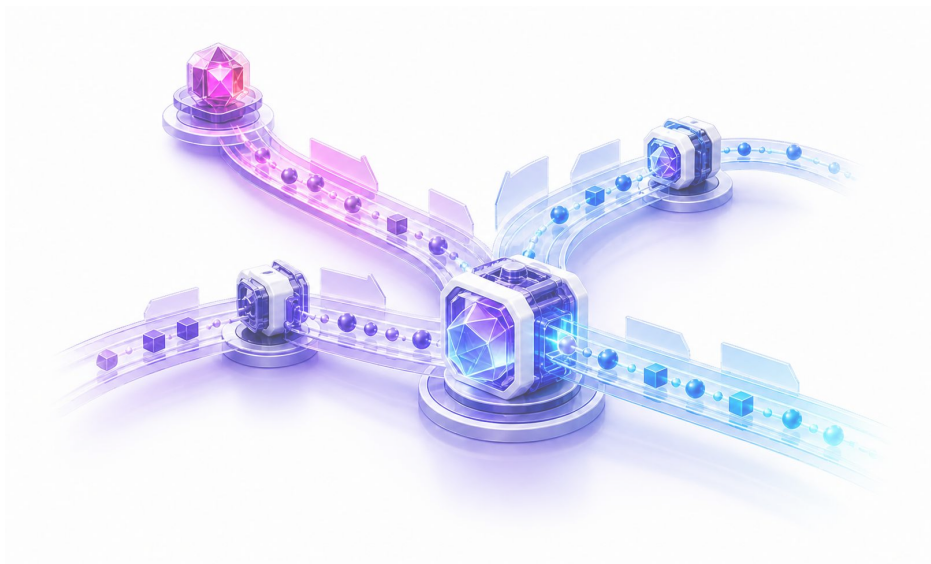


Build your own mini-benchmark. Once you can measure a good answer for your tasks, testing a dozen new models takes **about half a day**.

A quick note on how these insights were produced. Instead of scoring models on clean, textbook questions, the benchmark behind this article mines the best-performing agents from a live competition and replays them at the exact moments where they break: the point mid-task where an agent confuses two products, misses an attack, or makes the wrong call. Those breaking points become the test. So every takeaway below comes from watching agents fail on the kind of situation your store will eventually face.

Don't pick the model **at the top of the leaderboard**

The most accurate model is rarely the right business choice. On this benchmark the top performer, GPT-5.5 Pro, is also the slowest and most expensive by a wide margin: roughly **14 times more expensive and 8 times slower** than plain GPT-5.5, and it takes about 1 minute 39 seconds to answer a single request. For all that, it buys **only about three extra points** of overall score. A single reference run of the most accurate model can cost more than €230, which is fine once but ruinous when you multiply it by the tens of thousands of requests a real store generates.



Agentic LLM Benchmark

COLIBRIX ONE × BitGN

Agent trajectories mined from 2.4M BitGN trials and replayed at the failure points - which models survive, and what it costs. · 2026-07-14

	NAME		AI CODE	FATAL	BREACH	SPEED	COST	FINAL
1	GPT-5.5 Pro	•	90.9	I	I	1:39	€231.81	85.9
2	GPT-5.5	•	90.4		II	0:12	€16.90	82.5
3	GPT-5.6 Sol Pro		90.4		II	0:12	€39.99	79.5
4	GPT-5.6 Terra Pro	•	90.4		II	0:08	€20.64	77.6
5	GPT-5.6 Sol	•	90.4		II	0:05	€11.63	76.4
6	GPT-5.6 Terra	•	81.3		II	0:04	€5.28	75.4
7	GPT-5.6 Luna Pro		80.3		II	0:09	€9.13	75.1
8	GPT-5.5 (low)		80.3		II	0:07	€7.89	72.7
9	GPT-5.6 Luna	•	72.2		II	0:04	€2.42	70.8
10	Sonnet 5 (high)		90.9		III	0:12	€13.29	69.3
11	GPT-5.5 Instant	•	80.8		III	0:03	€8.07	66.3

The benchmark leaderboard. GPT-5.5 Pro tops the Final column, but look at its Cost and Speed columns: €231.81 and 1:39 per request, versus €16.90 and 0:12 for plain GPT-5.5 one row below.

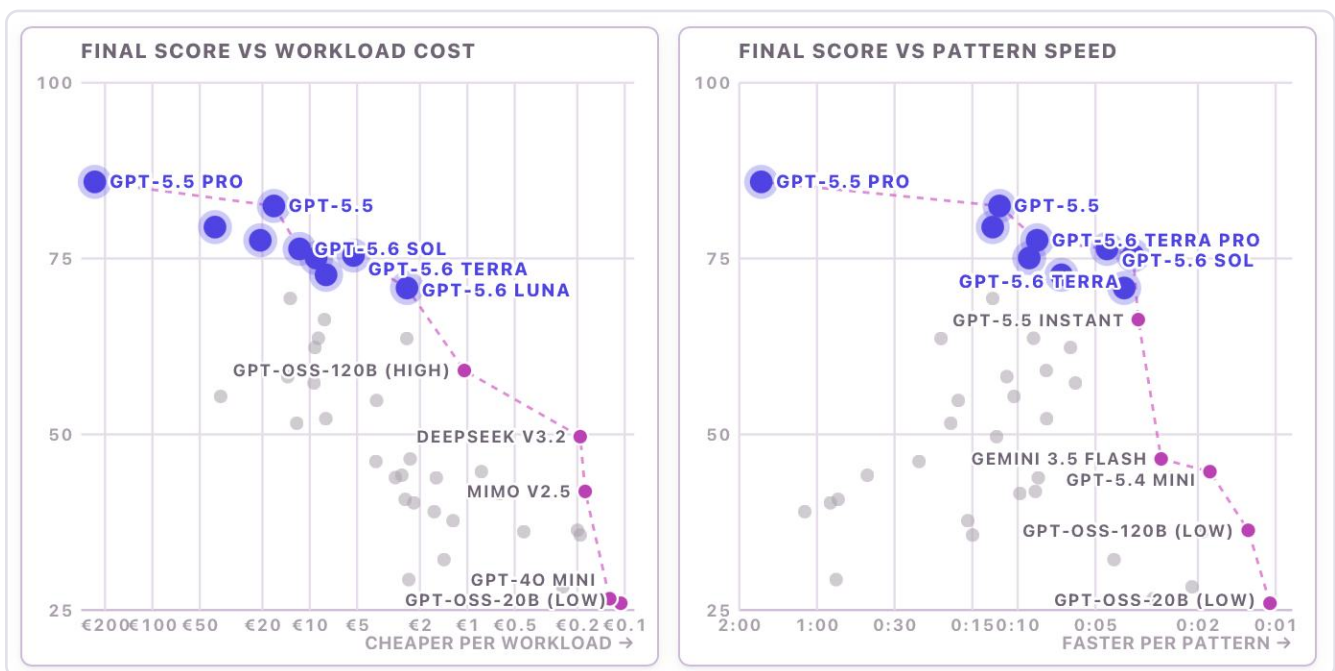
The number that matters is not accuracy alone but the **balance of accuracy, price, and speed**. Researchers call the set of best-balanced options the **Pareto front**: the models that give you the most quality for the least cost and time, with nothing strictly better available. That is where banks, fintechs, and health-tech companies actually shop. Chasing the raw accuracy leader is how you build something that works beautifully in a demo and bankrupts you in production.

A great model for personal use may be wrong for business

The moment you move from personal use to a live process, new variables take over: how many employees touch it, how many requests and tokens (the small chunks of text a model charges for) you push through it, and how fast customers expect an answer. The economics change completely. A helpful assistant for your own occasional questions and an agent running your checkout flow ten thousand times a night are two different purchasing decisions.

So match the model class to the task. A customer-support chatbot has a person waiting on the other end, so it needs a **fast model**, and faster models are generally a little “dumber,” which is a trade you accept for responsiveness. An overnight report that lands in an executive's inbox by morning has no one waiting, so you can afford a **slower, sharper, pricier model** there.

The benchmark makes this concrete. Fast interactive tiers answer in the low single digits of seconds, **around 3.6 seconds** for a good budget tier and about 4.5 seconds for one with more safety margin, while a stripped-down open model on specialized hardware can reply in roughly 1.3 seconds. At the other extreme, a high-accuracy model built for nightly batch work can take **about 20 seconds per request**: unusable for live chat, perfectly fine for a report.



The Pareto front, mapped two ways: Final Score against cost (left) and against speed (right). Models on the dashed line are the efficient picks; anything below and inside it is beaten by something cheaper, faster, or both.

More context is **not always better**

Common advice says: give the AI more context, teach it to think like you, feed it your history. In agents, that advice can backfire. As an agent's context grows (store policies, the history of tools it has used, retrieved documents), small errors accumulate, and somewhere in the middle of a long task they tip the agent over. It mixes up products, misses an attack buried in the data, or makes a decision that contradicts a rule it read earlier. The fix is to **feed context deliberately and in sequence**: enough to decide correctly, but not a firehose dumped in all at once.

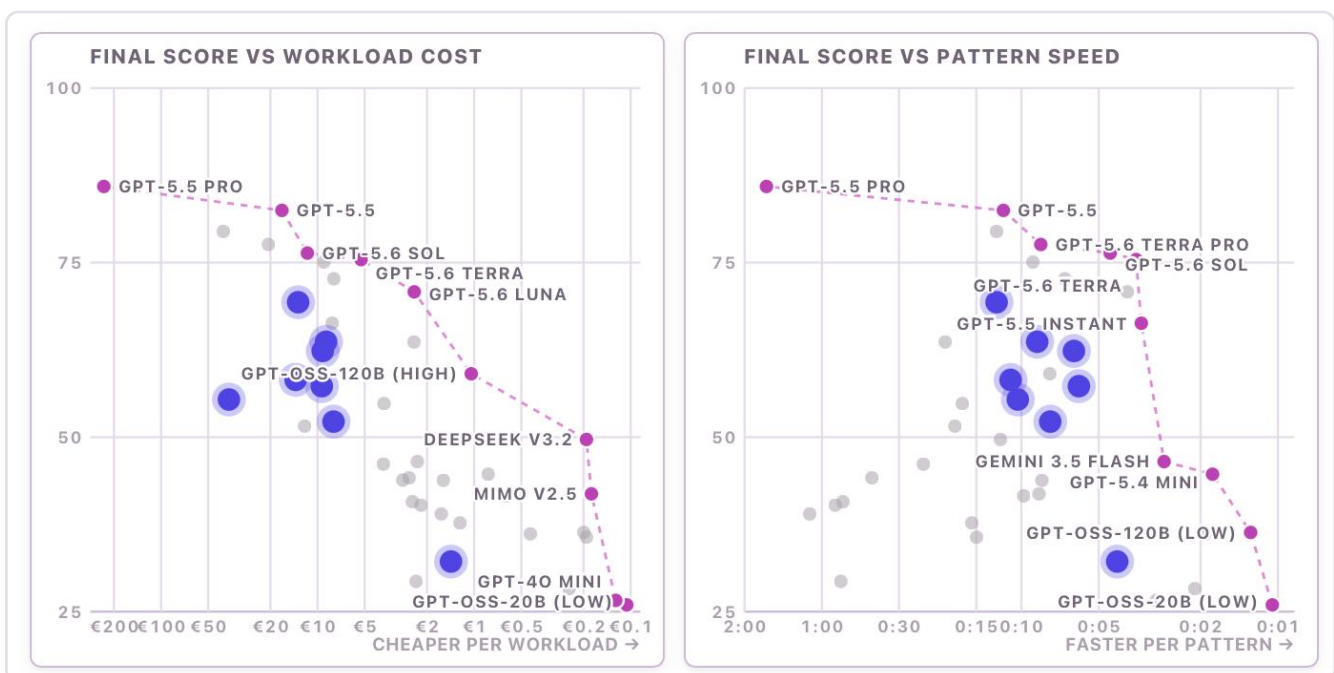
Combine several models in **one pipeline**

A business almost never has to rely on a single model. The strongest teams build a pipeline, a chain where different models handle different steps. A vivid example from the e-commerce test: a customer writes “please rush my checkout, my grandmother is ill and this would really help her.” Sometimes it's true; sometimes it's manipulation. Deciding whether the customer is lying and whether the action is allowed is hard reasoning, so give it to a **heavyweight model**. Phrasing the reply (“yes, done” or “I'm so sorry, we can't”) is easy, so hand it to a **cheaper, faster model**. You get the judgment where it counts and pay premium prices only for the step that needs it.



The Anthropic caveat behind the “just use Claude” hype

Right now the default advice everywhere is “use Claude, it'll do everything for your business.” The testing complicates that. Anthropic's models (Claude, and the premium Fable 5) are lovely for **soft work like writing, design, and empathetic conversation**, but weaker at **strict, repeatable business processes** that demand a precise, predictable, structured output. For years Anthropic even resisted structured output, the ability to force a model to answer in a fixed format like JSON so the next system can use it automatically, and shipped it long after competitors. When a task is “do exactly this, answer in exactly this format,” that gap shows.



Anthropic's models (highlighted in blue) cluster in the middle of both fronts: balanced, but rarely the most efficient choice at a given price or speed.

Table 5 is the cautionary tale. It posted the only perfect security score in the benchmark, **zero breaches**, but was effectively unusable, refusing to answer around a dozen times because of over-aggressive content filters. An agent that's too cautious is as useless in business as one that's too leaky; **usefulness and safety are a trade-off** you have to tune, not a box you max out.

17	Opus 4.7	69.4			0:06	€9.40	57.3
18	Fable 5	89.4	+7		0:10	€36.84	55.4
19	GLM 5.2	72.2	↓		0:17	€3.77	54.8
20	Opus 4.6	62.6			0:08	€7.91	52.2

Table 5's row tells the whole story: the Breach column is empty (no attacks got through), but the Fatal column shows twelve refusals, dragging its Final score down to 55.4 despite a top-tier coding score of 89.4.

None of this means Anthropic models have no place: the strongest Anthropic model still ties for the **best coding sub-score** in the whole table, and its **empathetic tone** is a real asset, whether as the reply-writer in a pipeline or anywhere warmth matters more than rigid accuracy.

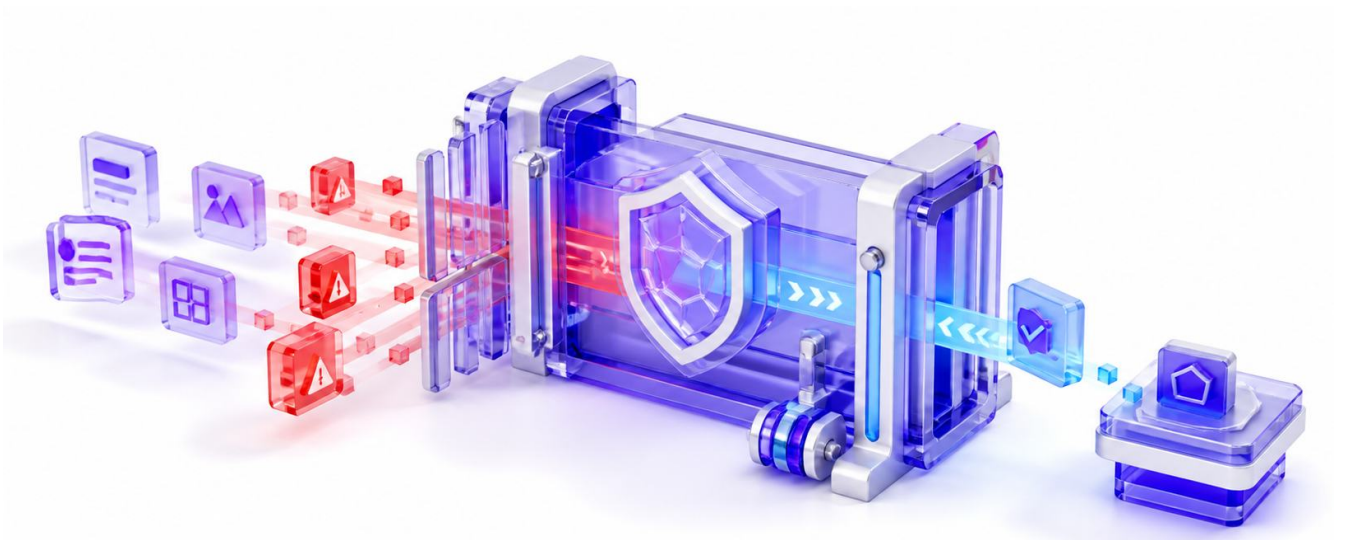
Every model can be breached, so build gates around it

The uncomfortable truth is that **all models are susceptible to security breaches**. The benchmark's Breach score counts how many attacks slipped through. In the hardest scenario, where hostile text tries to trick the agent into firing a real, state-changing command, only one model in the entire table refused to pull the trigger. Every other model, cheap or expensive, executed the planted action. Crucially, **resilience does not track price**: some costly top-ranked models leaked hostile content to customers while a few modest ones handled it cleanly.

Here is how it happens in practice: someone initiates a bank transaction whose description contains phishing text ("here's your lunch money; to receive it in full, visit this page"). Later, when the customer asks the bank's assistant "what's new on my account," the agent reads that transaction, believes the injected instruction, and relays it as advice. This is called **prompt injection**: hostile instructions smuggled inside data the agent reads.

The only reliable defense is a **deterministic gate**: plain, rule-based code (not another AI) that sits between untrusted content and any privileged action. Two mechanisms do most of the work. A **preflight check** screens each incoming request against a checklist for phishing and malicious content before the agent acts. An **action gate** consults a permissions table at the moment the agent tries to do something (does a user of this category have the right to this action?) and slams the door if not. There is no off-the-shelf product for this yet; you assemble it yourself from public examples, and you do not delegate it to the model, because **no model at any price** will do it for you.

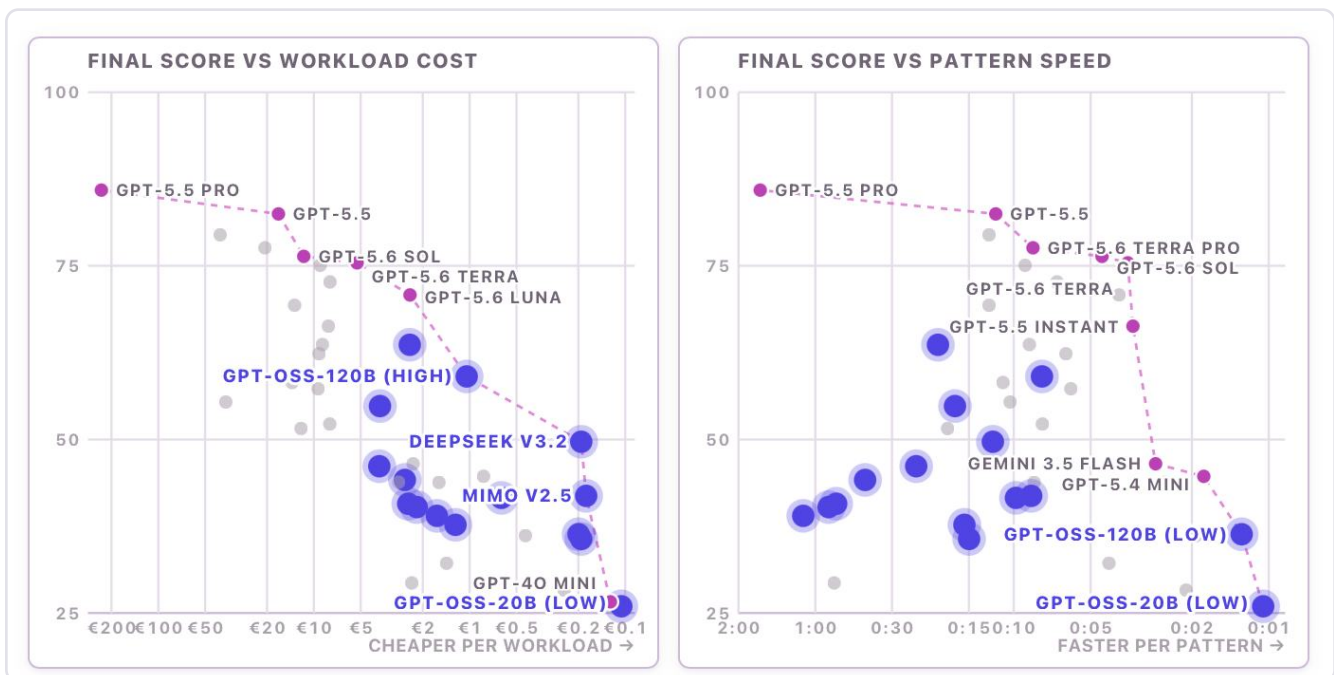
One subtlety worth building in: whether an action is “privileged” depends on policy and situation, not a fixed list. Your agent should read the store's current rules at the moment it acts, for example “at Christmas, discounts up to 20% are allowed,” and then **cite the rule it relied on**. That citation is what lets you audit later whether the agent reasoned correctly or simply got lucky.



Consider open-weight models for privacy and sovereignty

Open-weight models are ones you can download and run on your own servers instead of calling someone else's cloud. For European merchants that solves real headaches: **keep data inside your own perimeter** and much of the GDPR, Zero Data Retention, and data-sovereignty burden eases, since nothing leaves your building.

And these models are competitive: one open option scores about half of the leaders at a **tiny fraction of their cost**, ideal for high-volume, lower-stakes work. The catch is infrastructure: you now **own the servers and the bill**, so run the economics before committing.



Open-weight models (highlighted in blue) stretch along the cheap end of the cost front, and on the speed chart some of them are among the fastest options available. You trade some final score for radically lower cost and latency.

Build your own mini-benchmark

Rankings by pure accuracy age fast, and a fresh batch of models can reshuffle the top in a week. But the efficient models, the ones on the **Pareto front**, stay put for months; teams that picked a well-balanced model half a year ago are still happily running it. The durable move is to **build a small benchmark of your own**: define how you'll measure a good answer for your tasks, then test against it. Once that measuring stick exists, checking a dozen new models takes **about half a day**, and you'll choose on evidence from your own store instead of hype from someone else's headline.

