

The Agent at the Gate

Amazon, Perplexity, and the authority problem hiding inside AI commerce.



THOMAS ROCHA III

AUG 10, 2026



On August 4, the Ninth Circuit decided that when Perplexity's AI assistant operates Amazon through your browser on your behalf, the one accessing Amazon is you.

That sounds like a narrow computer-access ruling. It may turn out to be one of the first important cases about the architecture of agentic commerce.

Amazon had won an injunction in March. Perplexity's Comet browser includes an assistant that, when a user turns it on, navigates Amazon on that user's behalf, takes screenshots of the browser view, sends them to Perplexity's servers, and receives instructions on what to do next. Amazon called that unauthorized access to its computers under the Computer Fraud and Abuse Act. The district court agreed. The Ninth Circuit vacated the injunction and sent the case back.

The reasoning is narrow and worth reading precisely. "Access," the panel held, means entering a computer system, and the statute's "whoever" contemplates a person, not a software tool. So it was the user who accessed Amazon's computers, with the help of Perplexity's assistant, to carry out specific acts on Amazon.com. The court noted there is little to no existing case law on how to ascribe responsibility for AI agents, resolved the ambiguity with the rule of lenity, and reminded everyone that the CFAA is principally an anti-hacking statute.

This is being read as a win for agentic commerce. It is. It is also something more interesting, which almost nobody is saying out loud.

The holding turned on where the boundary sat

Read what the court actually examined. It looked at how the assistant works. Screenshots leave the user's machine. Instructions come back. The assistant, in the panel's words, cannot operate wholly independently: it relies on direction from the user and instructions from Perplexity's servers.

That architecture is why Perplexity won. And Perplexity's engineers built it that way on purpose, keeping the agent inside the user's browser session rather than spawning independent server-to-server connections, precisely because of concerns about unauthorized access.

So the rule that emerged is not "agents may enter Amazon." It is closer to: when the agent operates as an extension of the user's own session, the user is the one who entered. Change the design, and the attribution

may change with it. An agent that authenticates on its own credentials, holds its own session, and acts on a schedule while the user sleeps is a different case, and the court expressly reserved what might happen on different facts or with a greater degree of provider control.

Which means a federal appellate court just told the entire industry that where an agent's access gets attributed depends on where the operative boundary sits in the design.

That is not a footnote. That is the whole problem, stated in a place where it now has consequences.

Amazon's paradox

Notice the position Amazon is in. It is spending enormously to build the infrastructure of agentic commerce while simultaneously litigating to keep somebody else's agent out of its store. Both of those are rational. They are also the same question asked from two directions: who decides which agents may operate here, and under what terms?

Amazon reached for the CFAA because it was the tool on the shelf. The court declined to turn a 1986 anti-hacking statute into a licensing regime for AI agents, and it was right to decline. But declining to answer a question is not the same as the question going away. Amazon still needs a way to govern agent participation in its marketplace. So does every merchant, every marketplace, every auction house, every booking platform, every payment flow.

The commercial problem outlived the legal theory.

Then it gets harder

Entry is the easy question. It is the one everyone is arguing about because it is the one that just got litigated. The harder questions arrive immediately afterward.

An agent that may enter can search, compare, and read. May it also negotiate? Bid? Purchase? Return? Dispute? Arrange payment? Accept terms? Each of those is a different act with a different consequence, and “the user authorized the agent” does not distinguish between them.

Auctions make this vivid. In an auction, the distance between watching and being bound is one click and a few milliseconds. An agent observing a lane is doing nothing consequential. The same agent placing a bid has created a binding commercial obligation for a principal who may not be watching. Between those two states there is no meaningful pause, no natural human friction, no moment where somebody looks up and says Wait.

So the question progresses. May the agent enter? In what capacity may it participate? And then the one that actually matters:

Does this act bind the principal?

That last question is not answered by anything in the Perplexity opinion, and it is not answered by any of the regulation now arriving.

Three regimes, three different misses

Here is what makes this moment strange. We are not short of AI governance. Two of the most active regulators in the world have been building frameworks for years. Neither one reaches this.

The EU AI Act became broadly applicable on August 2, 2026, with the AI Office now holding enforcement powers over general-purpose models and fines reaching fifteen million euros or three percent of global turnover. But the Act regulates AI systems and the people who provide and deploy them. It classifies risk, imposes obligations, requires transparency. It does not decide whether a particular agent may act inside a particular merchant’s environment at a particular moment.

California is often described as following Europe. That is no longer accurate. Its finalized CCPA regulations already require privacy risk assessments, phase in annual cybersecurity audits, and impose automated-decision-making rules starting January 1, 2027. SB 53, effective at the start of this year, requires frontier developers to publish safety frameworks and transparency reports and to report critical safety incidents to state emergency services within fifteen days, twenty-four hours if the danger is imminent. California is developing a parallel governance regime that is, in several respects, more operational and more sector-specific.

But the two systems share an omission, and it is the same one. Neither supplies the runtime authority layer required when software acts as a principal's delegate inside somebody else's commercial environment.

Look closely at what California's automated-decisionmaking rules cover: technology that replaces or substantially replaces human decisionmaking for significant decisions about a person, in lending, housing, employment, education, health care. Advertising was cut from the final draft. An agent buying a laptop on your behalf is not making a significant decision about you. It is making a decision as you. The regime does not reach it.

And SB 53 governs frontier model developers, not commercial transactions executed by agents.

So: an anti-hacking statute that the court correctly refused to stretch. A European framework that governs AI systems and providers. A California regime that governs automated decisions about people, and another that governs frontier labs. Three serious instruments, none of which answers whether this agent, at this moment, may perform this act, and whether the result binds anyone.

That is not a criticism of any of them. Each was built for a real problem. It is an observation that the problem in front of us has a shape none of

them was cut to fit.

The missing center

Strip the case law and the regulation away, and the unanswered question is small, concrete, and operational.

A user tells an agent to buy something. The agent enters a marketplace. Conditions change while it is there: the price moves, the terms update, the item is different than described, the payment method needs re-authorization, the model handling the task gets swapped mid-session, the transaction crosses a jurisdiction. At each of those moments, something has to be true for the next act to be legitimate, and no system in the chain currently establishes it:

Is this actor, right now, authorized to do this particular thing, under conditions that still hold, on whose behalf, and what happens to that authority when one of those facts changes?

Notice this is not an identity question. Everyone in the chain knows who everyone is. It is not simply an access question either. The Ninth Circuit resolved who accessed Amazon; it did not resolve the larger authority question. The user entered Amazon through the agent. Whether that agent-mediated participation was permitted, what authority governed it, and what acts could bind the user all remained outside the holding. It is a question about whether authority granted at one moment still governs an act occurring at a later moment, after things have moved.

That distinction has a natural structure, and it is the structure the commercial world will end up building whether or not anyone plans it:

Before: may this agent enter at all, and on what basis?

At the boundary: in what capacity may it participate, what may it reach, what may it do.

During: does the authority that made entry legitimate still govern the act now occurring, as conditions change and the interaction becomes consequential.

Once you have those three, the rest of the agentic AI problem list stops looking like a list. Payment authorization, model substitution, delegation to sub-agents, jurisdiction changes, provenance, compute placement, device handoff, bidding: those are not separate novel problems. They are mutations occurring inside an authority-governed interval, and each one is asking the same question at a different seam.

What the court actually started

The first internet problem was connectivity: can these machines communicate? The platform problem was identity and access: who are you, and what may you enter? The agentic problem is authority: when software acts for you, what makes the act yours?

Agentic commerce introduces a requirement neither of the first two solved, because until now a person was always standing at the consequential moment. The requirement is that authority survive action. Not that it be granted, not that it be verified once at the door, but that it remain continuously controlling as the interaction moves and changes and finally becomes binding on somebody.

Courts are starting to bump into this. The Ninth Circuit hit it and, sensibly, ruled narrowly and left the larger question alone, while telling us in passing that the answer depends on the architecture. Europe and California are regulating around the edges of it from opposite directions. Amazon is fighting about it in two directions at once.

Commerce is going to have to implement it, because commerce is where the binding happens.

I have written elsewhere about the underlying principle, authority that survives mutation, and what it takes to carry it. The point here is

narrower. A federal appellate court has now made the attribution of an AI agent's access turn on where the operative boundary sits in the system's design.

That is not yet a law of agentic authority. It is a warning that architecture will determine where that law has to attach. And almost nobody is designing that boundary on purpose.

They are about to have to.