

What policymakers need to know about *agentic AI*.

A plain-language guide to the technology behind the headlines. What an AI agent actually is, how these systems work, and the questions Congress and regulators will need to answer. No background required.

10 minute read • Updated June 2026 • Non-partisan • Non-technical

01

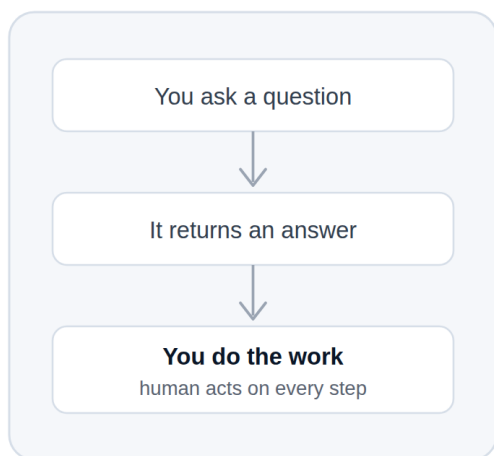
The shift from answering to acting.

For the last few years, the AI most people encountered did one thing: it answered. You typed a question, it produced text. The human stayed in the loop on every step, reading each output and deciding what to do next.

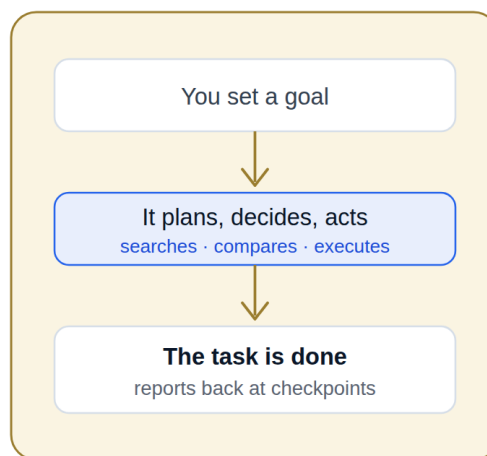
Agentic AI is different. An agent doesn't just answer. It acts. Given a goal, it can plan a series of steps, make decisions along the way, and carry them out with limited human supervision. It can browse the web, run code, fill out forms, move money, and work with other software systems on a person's behalf.

The simplest way to hold the distinction: the old model is a research assistant who hands you a memo. An agent is a junior staffer you can hand the whole task to.

TRADITIONAL MODEL



AGENT



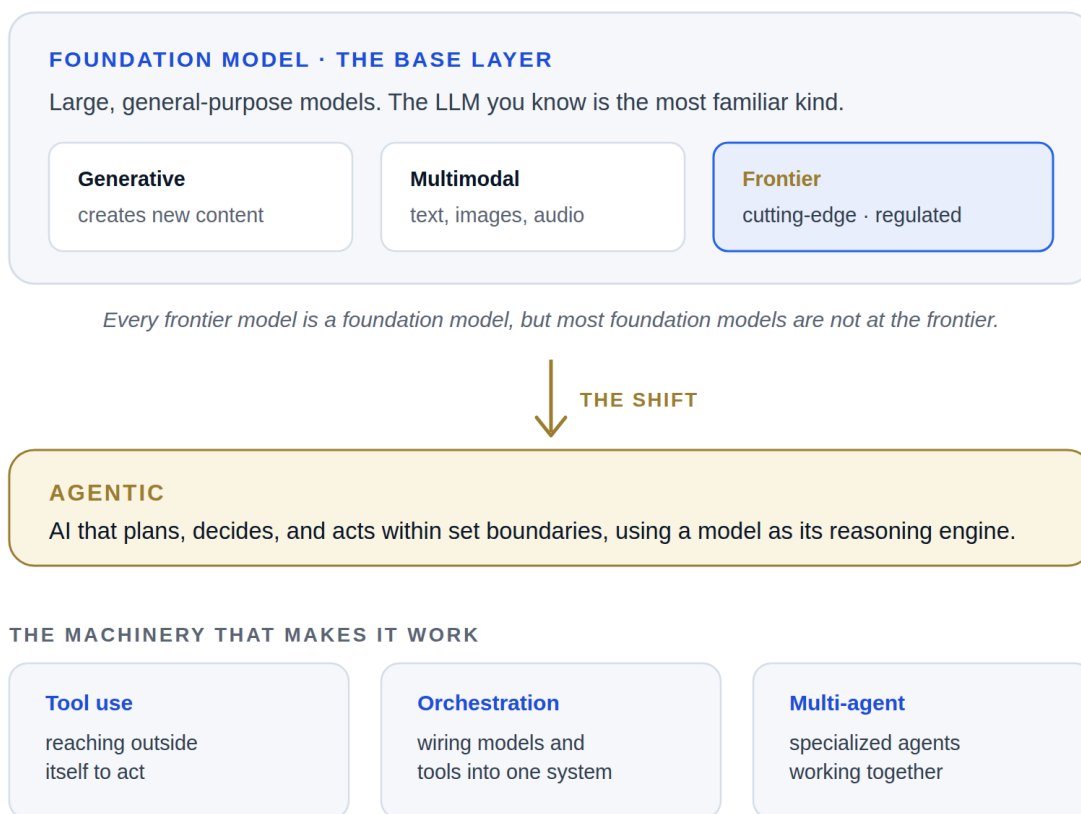
A traditional model answers. An agent takes the goal and does the task, checking back at the points you set.

That shift is why existing AI policy frameworks don't map cleanly onto these systems. Most were written for static models and human-reviewed outputs. When software starts taking actions in the world, a different set of questions opens up around accountability, identity, and oversight. Those are the questions this primer is built to frame.

02

The vocabulary, in one place.

A handful of terms come up in nearly every agentic AI conversation. They sit in a rough hierarchy. Understanding how they relate is most of what it takes to follow a technical briefing.



The terms sit in a hierarchy. Foundation models are the base. Frontier is the regulated cutting edge. Agentic is the shift to systems that act.

Those three building blocks, in a little more detail:

Tool use. An agent's ability to reach outside itself, calling a search engine, a database, a payment system, or another application to get something done.

Orchestration. Wiring multiple models, tools, and agents together into one coordinated system that can handle a complex, multi-step workflow.

Multi-agent. An arrangement where several specialized agents work together, each handling part of a larger task and passing work between them.

One concept runs underneath all of these and deserves its own line, because it's where much of the policy weight lands: identity and authorization. This is how a system knows *which* agent is acting, *whose* authority it carries, and *what* it has been permitted to do. When software acts on your behalf, that permission question is the whole game. The industry is building the standards for this right now, across open-source foundations, formal standards bodies, and the payments ecosystem. (A companion AFI reference maps who is doing what.)

03

How an agent actually works.

Picture a concrete task: "Book me a flight to Chicago next Thursday under \$400, aisle seat." A traditional model would draft you an email about flights. An agent does the job.

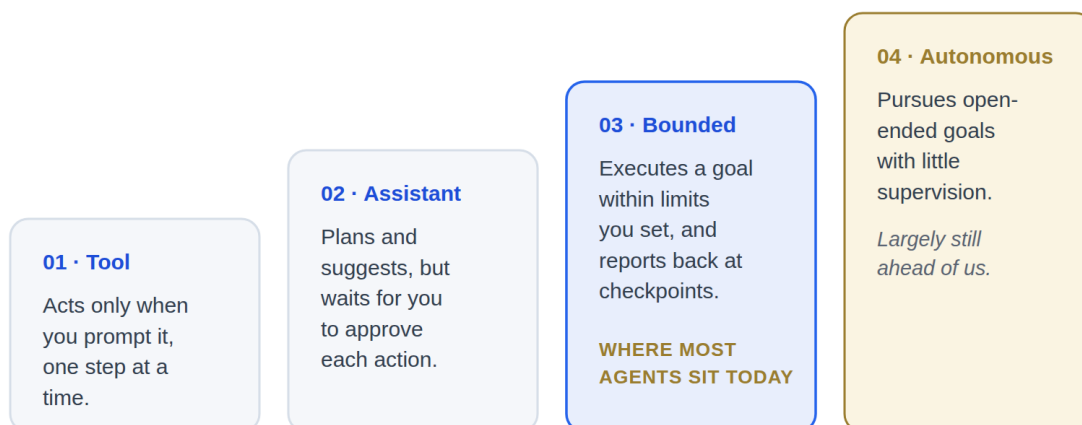
It breaks the goal into steps, searches for flights, compares options against your constraints, selects one, and completes the booking, checking back with you at the points you've told it matter. Two things define what it can and can't do. First, what it's allowed to act on: which systems and accounts it can touch. Second, what it can see versus what it reports back to you. An agent given a narrow scope and a clear reporting line is very different from one with broad access and little visibility.

There's one more property worth understanding, because it surprises people from a traditional software background. Agents are non-deterministic. Run the same task twice and you may get two slightly different paths to the answer. That's not a bug. It's a consequence of how these models reason. But it's exactly why supervision, logging, and clear boundaries matter more here than in conventional software, where the same input always yields the same output.

"Agent" is not one thing. Systems fall along a spectrum of autonomy, from a tool that acts only when prompted to an actor that pursues open-ended goals on its own. Most of what's deployed today sits in the middle, executing bounded goals and reporting back. Where a system falls on this scale is the single most useful question for a policymaker, because it determines how much human oversight is realistic and how high the stakes run.

LESS AUTONOMOUS

MORE AUTONOMOUS



HUMAN OVERSIGHT NEEDED PER ACTION



The right level of oversight depends on where a system sits, and so does the policy question.

Autonomy is a spectrum, not a switch. As systems move right, the human oversight needed per action falls and the policy stakes climb.

04

The policy questions that follow.

You don't need to resolve these to be conversant in them. Naming the questions clearly is most of the work. Each is an area where existing frameworks were not built for software that acts.

Accountability

When an agent takes an action that causes harm, who is responsible? Responsibility is distributed across the model developer, the company that built the application, the company that deployed it, and the user who set the goal, and the harm can be hard to trace to any single one. Scholars call this the problem of many hands. Existing liability law wasn't written with an autonomous actor in the middle.

Identity & authorization

How do systems verify that an agent is who it claims to be and is authorized to act for a given person? This is the foundation everything else rests on, and the infrastructure for it is still being built.

Interoperability & access

Will agents be able to work across different platforms and services, and on what terms? Technical standards, security, and business considerations all bear on this. How it resolves shapes how easily new entrants can build and how much choice users have.

Workforce

As agents take on tasks rather than just answer questions, what happens to the work people do? Estimates of the scale vary widely and the picture differs by sector and role, with some work augmented, some automated, and new roles created. It is the question constituents feel most directly, and the one where good data matters most before policy is set.

Security & privacy

An agent with broad access and the ability to act is a powerful tool and a new attack surface. What controls, logging, and data protections should govern systems that can act on sensitive accounts?

Oversight & supervision

What does meaningful human oversight look like for a non-deterministic system operating at machine speed? The field distinguishes a human in the loop (approving each action), on the loop (monitoring and able to intervene), and in command (setting the boundaries up front). The hard question is which model fits which use, and where a human checkpoint is essential versus unworkable.

05

Where AFI fits.

The Agentic Futures Initiative is a cross-industry coalition of the companies building and deploying agentic AI. We don't ask policymakers to take our positions on faith. Our role is to make sure the people writing the rules understand the technology well enough to write them right, through briefings, filings, and plain-language resources like this one.

If a briefing, a deeper explainer on any of these questions, or a technical walkthrough would help your office, that's exactly what we're here for.

This primer is a non-partisan educational resource. Definitions reflect consensus industry usage as of June 2026. The regulatory definition of "frontier model" varies by jurisdiction and is tied to compute thresholds in the EU AI Act and several U.S. state laws. For a briefing or to go deeper on any section, contact info@agenticfuturesinitiative.org.