

AI labs are selling governed domain workflows, not just models

AI News Digest

2026-09-18

AI labs are selling governed domain workflows, not just models

By AI News Digest • September 18, 2026

OpenAI, Anthropic, and Google DeepMind are packaging frontier capability for law, biology, and genomics, while Anthropic puts numbers around AI-led R&D and agent oversight. The common thread is a shift from general-purpose models toward domain-specific infrastructure, permissions, and verification.

Domain stacks move into high-stakes work

OpenAI turns GPT-6 Astra into a legal stack

OpenAI launched Astra for Law as a foundation for law firms and legal-technology companies, combining GPT-6 Astra with legal-analysis and writing instructions, tailored settings, tools, context, and a legal search index. The index covers U.S. case law, statutes, regulations, court rules, and administrative decisions across more than 230 million URLs, with sources added daily. [1]

OpenAI reports that the complete setup passed the overall correctness check on 54.0% of 200 questions in the private Vals AI Legal Research Bench, versus 38.7% for GPT-6 Astra using web search alone—a 40% relative improvement. The result is vendor-reported and comes from a private validation set, but it shows the product’s intended advantage: retrieval and legal context are being packaged alongside the frontier model rather than left to users to assemble. [1]

The initial rollout is limited to selected firms through Trusted Access, with zero data retention on the API and default exclusion of ChatGPT Enterprise usage from human review; OpenAI is also working with Latham & Watkins on permissions, ethical walls, client instructions, and oversight. The launch adds 26 partner-built plugins and 47 adaptable community skills, reinforcing a strategy

built around firm-specific workflows and an ecosystem of legal tools rather than a standalone chatbot. [1]

Anthropic makes biology access both more permissive and more controlled

Anthropic opened applications for a beta Life Sciences Verification Program that gives verified teams access to Mythos, Opus, and Sonnet with safeguards more permissive for biology work than those on its generally available models. Applicants are reviewed for research credentials, security standards, and ethical oversight; standard grants cover broad team workflows, while high-risk grants are project-specific, require additional vetting, and renew every six months. [2]

The program ties access to an organization’s stated use cases and continuously monitors traffic for activity outside that scope. Anthropic says it is shifting some enforcement from real-time blocking to offline behavioral monitoring, retaining data associated with flagged activity for 30 days while keeping it compartmentalized and out of model training. [2]

Alongside the access program, Anthropic reports that Claude optimized more than 30 open-source biomolecular models, producing roughly fourfold average speedups with minimal precision loss and nearly twofold speedups with identical outputs; the optimization code is open-sourced. A new Adaptic Bio competition will experimentally validate more than 5,000 community designs, backed by up to \$1 million in Claude credits and \$250,000 in Modal compute credits. [3] The combination points to a practical biology strategy: expand access where users can be verified, while lowering the compute and experimental cost of the work itself.

Research infrastructure becomes a shared utility

AlphaGenome Atlas precomputes a map of human genetic variation

Google DeepMind introduced AlphaGenome Atlas as a free academic resource containing predictions for all 9 billion possible single-nucleotide variants in the human genome. The company describes it as a roughly 1-petabyte dataset with thousands of molecular-effect predictions per variant, an AlphaGenome Variant Impact score, and more than 2,500 recurring DNA motifs across hundreds of cell types and tissues. [4]

DeepMind says collaborators used the atlas to prioritize a *DNM1* variant in unsolved rare-disease research, with experimental screens validating the predicted mechanism. In a separate analysis of more than 54,000 UK Biobank participants, the University of Exeter team reported 22% more detectable non-coding genetic associations; Stowers researchers used the atlas’s motifs to classify regulatory activity. [4]

The important shift is from a model researchers query one case at a time to a precomputed, searchable research layer that exposes model predictions at

genome scale. DeepMind’s results are company-reported, but the release pairs the infrastructure with a concrete validation example and makes the resource available through a portal, API, and Google Antigravity. [4]

Transparency gets more precise, not yet comparable

Anthropic puts numbers around AI-led R&D and oversight

Anthropic published three internal measures of frontier development: how much AI R&D is performed by AI systems, how well agents are overseen, and how compute is allocated. It argues that other frontier developers could publish comparable measures and that third parties could verify them. [5]

Its August snapshot says Claude “led” 26% of Anthropic’s AI R&D work, performed at or above the “AI collaborates” level for more than 90% of the work, and was not fully autonomous for any measured subset. On the company’s most-used internal research platform, roughly 30,000 agents had 100% of actions pass through online or offline monitoring; online monitors blocked 0.002% of more than a billion decisions. Anthropic also reports that 6% of AI-R&D compute, and 12% of compute going to AI-driven AI R&D, was allocated to safety in a July 13–20 snapshot. [6]

Anthropic labels the automation index a prototype, limits the oversight figures to one internal platform, and says cross-lab comparisons need a common methodology and independent checks because the lab is using its own models as judges. Nathan Lambert’s reaction was similar: he called the disclosure a step in the right direction but said the 26% figure does not define what counts as AI R&D. [6, 7] The verification problem is therefore part of the development story itself; Geoffrey Hinton separately called independent verification organizations a good start, arguing that reliance on whistleblowers is not enough. [8]

Sources

1. Introducing Astra for Law | OpenAI
2. Introducing the Life Sciences Verification Program
3. How Claude is uplifting biomolecular modeling
4. AlphaGenome Atlas: Molecular predictions for 9 Billion human DNA variants
5. X post by @AnthropicAI
6. Measurements for understanding the pace of AI development inside frontier labs
7. X post by @natolambert
8. ‘Godfather of AI’ Geoffrey Hinton Warns AI Kill Switch Won’t Work Long-Term | AI Threat | N18G