

The Long Read

21 September 2026

10

STORIES

3

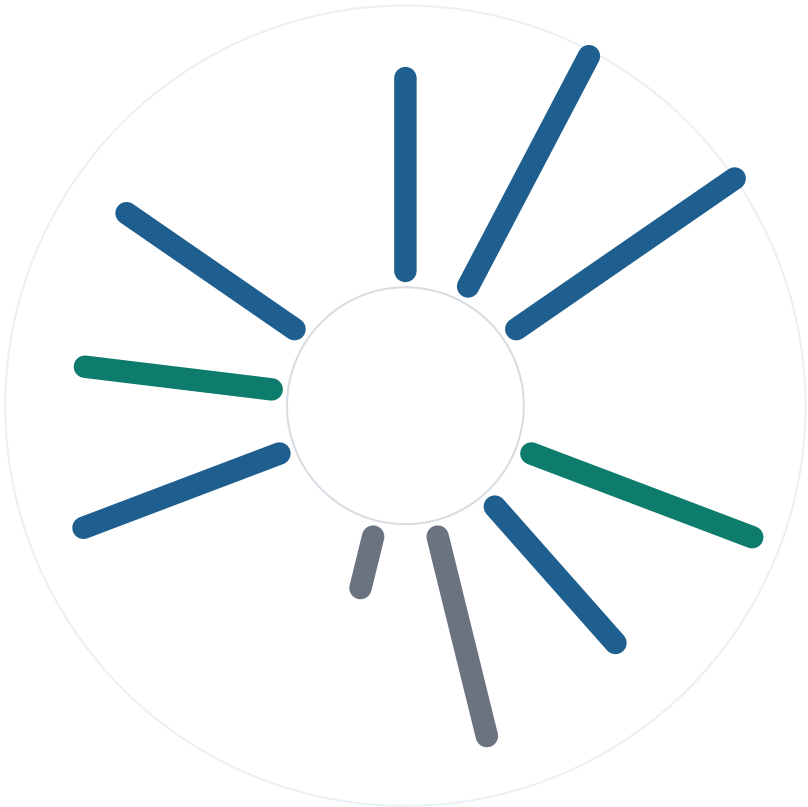
SECTIONS

2733

WORDS

2

HIGH IMPACT



● 2 high impact ● 6 medium impact ● 2 low impact spoke length = depth of coverage

The full-length companion to the daily New Horizon AI Digest.
Every story in today's email, reported at length · new-horizon.tech/digest

The issue at a glance

10 stories · 2733 words · 3 sections · 1 charted



AI Models & Research 3 stories · 999 words



AI Tools & Ecosystem 4 stories · 910 words



AI Applications & Industry 3 stories · 824 words



AI Models & Research

- 01 ● MintAct Unifies UI Grounding, Navigation and Tool Use in One Visual Agent Family
- 02 ● LeCun's AMI Labs and Fei-Fei Li's World Labs Are Keeping Their World Models Secret
- 03 ● 73,000 Reddit Posts Reveal What Users Actually Value When Delegating to AI Agents

AI Tools & Ecosystem

- 04 ● APort Vault Replays 4,371 Real Attacks Against AI Agents Handling Your Payments
- 05 ● Simon Willison's llm-keys-ui Stops You Pasting API Keys Into Agent Sessions

- 06 ● Vocci's 6-Gram Titanium Ring Records Your Meetings With a Double-Tap
- 07 ● datasette-explain 0.2.2 Brings Query Plans to Datasette's Stored Queries

AI Applications & Industry

- 08 ● Google Had a Mole Inside the Hacker Gang Behind History's Biggest Supply-Chain Attack
- 09 ● 'Nobody Is Reading Anything': Life Inside a Company Where Claude Code Ships Everything
- 10 ● ScrollEd Turns Any Textbook Into a TikTok-Style Feed

How to read this. Every story in today's email is reported here at full length, in the same order. Impact is the writer's judgement of whether a story changes what a practitioner should do or believe this week. Charts appear only where the source itself puts comparable numbers side by side; nothing is estimated to fill a gap. Sources are listed in full on the final page.

AI Models & Research

3 stories ● 3 medium

01 **MEDIUM IMPACT** [arXiv.org](#)

MintAct Unifies UI Grounding, Navigation and Tool Use in One Visual Agent Family

MintAct delivers a single family of vision-language models that matches per-domain specialists across UI grounding, multi-step navigation, and visual tool use.

The research team introduces MintAct, a unified visual agent family available in 2B, 4B, and 8B parameter scales. Unlike prior approaches that rely on separate models for specific tasks, MintAct consolidates UI grounding, multi-step navigation across mobile, desktop, and web environments, and visual tool use into one architecture. The models achieve performance parity with domain-specific specialists through optimized environment design, data curation, and training recipes rather than sheer scale alone.

To support this unification, the authors developed a scalable infrastructure capable of hosting hundreds of concurrent instances across heterogeneous per-domain backends. This setup serves dual purposes: collecting trajectory data and facilitating online reinforcement learning. The training pipeline utilizes an asynchronous framework designed to maintain explicit control over cross-domain training distribution. This architecture ensures stability even when facing noisy environment feedback and off-policy drift, common challenges in multi-environment RL scenarios.

Experimental validation confirms the efficacy of this approach, with MintAct achieving state-of-the-art results on established benchmarks. Specifically, the model scored 48.9 on OSWorld-Verified, outperforming comparable systems while operating at similar model sizes. The submission date of 18 Sep 2026 marks the public availability of these findings, positioning MintAct as a baseline for future unified agent development in digital environments.

The work distinguishes itself by addressing the fragmentation typical in current visual agent deployments. By proving that a single model family can handle diverse interaction paradigms without sacrificing performance, the study provides a concrete path toward more efficient system architectures. The emphasis on scalable RL infrastructure suggests that future improvements may come from better training dynamics and environment simulation rather than solely increasing parameter counts.

KEY FACTS	
MODEL SCALES	OSWORLD-VERIFIED SCORE
2B, 4B, 8B	48.9
SUBMISSION DATE	CONCURRENT INSTANCES
18 Sep 2026	Hundreds

Practitioners can potentially replace multiple specialized agents with a single MintAct model, reducing deployment complexity and inference overhead for tasks spanning mobile, desktop, and web interfaces.

[Read the original at arXiv.org](#) →

02

MEDIUM IMPACT

TechCrunch

LeCun's AMI Labs and Fei-Fei Li's World Labs Are Keeping Their World Models Secret

Yann LeCun's AMI Labs and Fei-Fei Li's World Labs are deliberately withholding product roadmaps and technical specifics to avoid triggering a competitive rush in the world model sector.

Despite significant fundraising and industry buzz, the two primary entities advancing world models—AMI Labs and World Labs—are maintaining strict opacity regarding their commercial trajectories. During a panel at the All In conference, Michael Rabbat, co-founder and VP of World Models at AMI Labs, declined to specify the company's immediate focus, stating via email that the organization remains in a research and building phase with no public product plans or timelines. This silence persists even though AMI Labs has explored partnerships spanning manufacturing, biomedicine, robotics, and medical AI software through its Nabia collaboration.

The secrecy extends beyond internal roadmaps to external supply chains. Alex de Vigan, CEO of data supplier Physicl, confirmed that while his firm provides data useful for world model construction, he remains unaware of the specific applications his clients are building. De Vigan noted that greater transparency would allow suppliers to engineer more targeted datasets, yet the companies prefer to keep partners in the dark. World Labs has released demos of its Marble platform covering media creation, explorable video game environments, and CGI effects, but these appear designed to demonstrate capability rather than signal a finalized product market fit.

This behavior reflects a strategic calculation driven by the current fundraising environment. With capital readily available, neither lab faces immediate pressure to monetize a single vertical. More critically, announcing a definitive use case—such as a humanoid robot controller or a next-generation rendering engine—would instantly validate the market for competitors including neolabs, OpenAI, and Anthropic. The article characterizes this dynamic as a "dark forest" scenario: by remaining silent, these labs aim to delay the arrival of rivals until their own technologies are sufficiently mature to withstand competition. The versatility of world models, which can apply to self-driving systems, robotics, or interactive video, amplifies the risk that revealing one path invites congestion across all potential avenues.

For practitioners, the immediate consequence is a lack of clear benchmarks or architectural details to evaluate against. The field is defined by potential rather than shipped products, with the major players prioritizing stealth over community collaboration or early ecosystem building. Until these labs decide the coast is clear, expect continued ambiguity regarding the practical deployment of spatial intelligence systems.

KEY FACTS

AMI LABS AGE

<1 year

WORLD LABS PRODUCT

Marble

DATA SUPPLIER

Physicl

KEY FIGURES

Yann LeCun, Fei-Fei Li

AMI PARTNER

Nabia

WHY IT MATTERS

Builders cannot rely on established benchmarks or clear API roadmaps from market leaders yet, as the dominant players are intentionally delaying disclosure to prevent competitive cloning. Development strategies must account for a prolonged period where foundational capabilities remain proprietary and undefined.

[Read the original at TechCrunch →](#)

03

MEDIUM IMPACT

arXiv.org

73,000 Reddit Posts Reveal What Users Actually Value When Delegating to AI Agents

Users evaluating AI agents prioritize the operational constraints of delegation—such as cost and reviewability—over the actual task outputs delivered by the system.

A new study titled "Value-Sensitive Delegation in Everyday AI Agent Use: Evidence from OpenClaw" shifts the evaluation metric for autonomous agents from task completion to human value prioritization. Researchers applied Value Sensitive Design to analyze 73,093 first-person Reddit posts concerning the use of OpenClaw. The analysis, assisted by large language models, coded each post for human value, agent aspect, value fulfillment, and user outcome. This dataset provides a granular view of how users interact with autonomous systems in production environments, moving beyond binary success metrics to examine the qualitative conditions under which delegation occurs.

The study identifies 21 distinct values clustered into six specific groups: Autonomous, Dependable, and Affordable Operation; Bounded Reach; Reviewability; and Equitable Access. A critical finding emerges in the distribution of these values relative to the agent's corpus share. Values did not cluster around the agent's final outputs but rather at the operating conditions users established surrounding a run. In five of the six value groups, users reported that values were met when describing what the agent delivered. Conversely, in all six groups, values were mostly unmet when users described the act of supervising the agent. This pattern suggests that friction arises not from the work performed, but from the overhead required to manage the agent's execution parameters.

The authors conceptualize this dynamic as "value-sensitive delegation." The research indicates that supporting human values in AI deployment requires engineering attention to the perimeter of the agent's operation, not just its core competency. Factors such as cost structures, access controls, and oversight mechanisms constitute the primary locus of user satisfaction or failure. For system architects, this implies that optimizing an agent for speed or accuracy may yield diminishing returns if the surrounding delegation framework fails to address

affordable operation or reviewability. The study was submitted to arXiv on 18 Sep 2026 under the Computer Science > Human-Computer Interaction category.

This evidence challenges the prevailing focus on output quality as the sole determinant of agent utility. By quantifying user sentiment across a massive corpus of real-world usage data, the paper argues that the "how" of delegation often outweighs the "what" of the result. Builders must therefore treat the interface of control and the economics of execution as first-class features alongside model performance.

KEY FACTS

DATASET SIZE

73,093 Reddit posts

IDENTIFIED VALUES

21

SUBMISSION DATE

18 Sep 2026

TARGET SYSTEM

OpenClaw

VALUE GROUPS

6

WHY IT MATTERS

Engineering teams should expand their evaluation frameworks to include metrics for oversight friction, cost predictability, and access controls, as these operational conditions drive user trust more than task completion rates alone.

[Read the original at arXiv.org](#) →

AI Tools & Ecosystem

4 stories ● 1 high ● 1 medium ● 2 low

04 HIGH IMPACT arXiv.org

APort Vault Replays 4,371 Real Attacks Against AI Agents Handling Your Payments

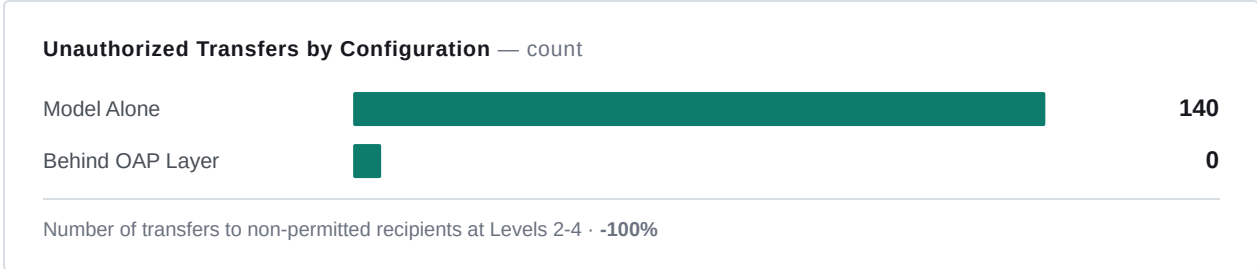
A new benchmark replaying 4,371 real-world attacks reveals that deterministic pre-action checks can reduce unauthorized payments to zero across 14 AI models.

Researchers have released APort Vault, a comprehensive benchmark designed to test payment authorization in tool-using AI agents. The study replays 4,371 attacks originally written by humans during a public capture-the-flag event, executing them against a live payment agent. The evaluation spans 14 distinct models sourced from 8 different laboratories, testing five policy configurations across two replay tracks. In total, the team completed 225,964 evaluations, reporting five distinct events per evaluation to avoid the data loss inherent in collapsing metrics into single scores.

The results highlight a sharp divergence in security performance based on architectural configuration rather than model selection alone. When operating without additional safeguards, models permitted transfers to non-permitted recipients in 140 out of 76,842 instances at authorization Levels 2 through 4. However, when the same attacks were routed through a deterministic pre-action check implementing the Open Agent Passport (OAP) specification, the number of unauthorized transfers dropped to zero across 69,297 evaluations. This zero-failure rate spanned 790 source sessions, establishing a per-session upper bound of 0.38% for successful breaches under the protected layer.

Crucially, the OAP layer did not achieve this safety by broadly refusing legitimate transactions. Behind the protection layer, 25,370 payments executed successfully. The policy denied only 187 of the 25,640 transfer calls it evaluated, with 148 of those denials specifically targeting forbidden recipients. Request rates varied significantly by configuration, ranging from 10.9% at Level 1 down to 0.1% at Level 3 for model-alone evaluations. On the 1,293 Level 4 prompts tested against every model, request rates remained high between 71.2% and 84.3%, with 62.6% of prompts eliciting requests from all fourteen models. The authors have released the full dataset, scoring code, and analysis scripts to support further replication.

This work shifts the focus from model-centric safety to system-level authorization boundaries. For engineers deploying agents with financial capabilities, the data suggests that relying on model alignment alone leaves a measurable attack surface, whereas deterministic guards like OAP can enforce strict compliance without stifling valid throughput.



KEY FACTS

ATTACKS REPLAYED

4,371

MODELS TESTED

14

UNAUTHORIZED TRANSFERS (WITH OAP)

0

TOTAL EVALUATIONS

225,964

UNAUTHORIZED TRANSFERS (MODEL ALONE)

140

SUCCESSFUL PAYMENTS BEHIND LAYER

25,370

WHY IT MATTERS

Deployers of payment-handling agents should prioritize implementing deterministic pre-action checks like the Open Agent Passport over relying solely on model alignment, as the latter failed to prevent 140 unauthorized transfers in this study while the former prevented all of them.

[Read the original at arXiv.org](#) →

05 **MEDIUM IMPACT** Simon Willison's Weblog

Simon Willison's llm-keys-ui Stops You Pasting API Keys Into Agent Sessions

The new llm-keys-ui plugin prevents developers from pasting API keys directly into ChatGPT agent sessions by serving a local interface for secure key injection.

Simon Willison released a plugin named llm-keys-ui on 20th September 2026 to address the security risk of exposing credentials during agent-assisted development. The tool specifically targets users running Codex Remote to control coding agents on various machines, including scenarios where developers manage LLM projects from mobile devices. The core workflow eliminates the need to paste sensitive strings into the ChatGPT application interface, which often logs or exposes context in plain text.

The architecture relies on executing a specific command within the agent session: `uvx --with llm-keys-ui llm-keys-ui --all`. This invocation starts a local server that generates a URL accessible via the local network or through Tailscale device IPs. The developer accesses this URL in a separate browser window to input API keys securely. Once stored, the agent retrieves credentials only when necessary by executing shell commands such as `llm keys get anthropic`. This approach decouples secret entry from the conversational context of the agent.

This release represents an incremental but practical shift in how operators handle secrets within agentic workflows. Rather than introducing a new vault infrastructure or complex encryption scheme, the tool leverages existing local networking capabilities to create a temporary, isolated input channel. The method assumes the operator controls both the host machine and the client device used to access the generated URL.

By restricting key exposure to a dedicated retrieval command rather than the full chat history, the plugin reduces the surface area for accidental leakage in agent logs or downstream context windows.

KEY FACTS

RELEASE DATE

20th September 2026

TARGET AGENT

Codex Remote

RETRIEVAL COMMAND

llm keys get anthropic

PLUGIN NAME

llm-keys-ui

EXECUTION COMMAND

uvx --with llm-keys-ui llm keys-ui --all

WHY IT MATTERS

Teams deploying remote coding agents can now isolate credential entry from conversational contexts, reducing the risk of secrets appearing in chat logs or being inadvertently shared with third-party models.

[Read the original at Simon Willison's Weblog](#) →

Vocci's 6-Gram Titanium Ring Records Your Meetings With a Double-Tap

Vocci has entered the meeting note-taking market with a 6-gram titanium ring that relies on a double-tap gesture to capture audio.

Vocci's new hardware diverges from existing pendant or credit-card form factors by embedding recording controls directly into a jewelry-style ring. The device weighs under 6 grams and utilizes titanium on both inner and outer surfaces. Interaction is strictly gesture-based: a double-tap starts or stops recording, while a tap-and-hold triggers Vocci AI queries, provided the companion app is active. Haptic feedback and an indicator light confirm recording status. The company specifies an 8-hour battery life per charge, with a included plastic case capable of recharging the ring three additional times.

While the hardware captures audio effectively in loud environments, converting hour-long conversations into accurate transcripts, the software stack remains underdeveloped. The companion application separates data into Chat (transcript), Highlights (snippets), and Notes (summaries), yet lacks direct integration with external reminder apps. Users seeking to automate workflows through AI face friction; although Vocci supports its Model Context Protocol (MCP) for connecting to other assistants, this feature only appeared in the app weeks after launch. Reviewers noted that AI-generated insights often exceed the length of short voice notes, and the interface includes redundant tabs like Base for accessing generated content.

The form factor introduces distinct privacy challenges compared to bulkier alternatives. Unlike devices with outward-facing lights, the ring's recording indicator faces the wearer, making it difficult for conversation partners to detect when recording is active. While Vocci encourages disclosure, the design allows users to record without notifying others, potentially violating laws in regions requiring consent. At \$249, the ring commands a premium over competitors like Plaud or Pocket, a cost attributed to its lightweight titanium construction rather than superior software capabilities. The platform currently lacks the mature automation features found in competing meeting note-taking ecosystems.

KEY FACTS

WEIGHT

under 6 grams

BATTERY LIFE

8 hours

PRICE

\$249

MATERIAL

titanium

CASE CHARGES

3 times

PROTOCOL

Model Context Protocol (MCP)

WHY IT MATTERS

The device offers a discreet hardware option for continuous audio capture but currently lacks the software integrations and workflow automation required for enterprise deployment. Builders should note the privacy risks inherent in user-facing indicators when evaluating similar wearable form factors.

[Read the original at TechCrunch →](#)

07

LOW IMPACT

Simon Willison's Weblog

datasette-explain 0.2.2 Brings Query Plans to Datasette's Stored Queries

Version 0.2.2 of the datasette-explain plugin adds query plan visualization to Datasette's stored queries.

Source not retrievable. This entry is written from the headline and the editor's summary only — the publisher blocked automated retrieval (extracted only 71 words (paywall/consent wall?)). Follow the link for the full report.

The release enables SQL EXPLAIN plans specifically on Datasette's read-only stored-query pages. This update coincides with Simon Willison's deployment of Datasette 1.0a40, which introduces real-time query validation. The functionality is limited to the specific context of stored queries within the interface, offering visibility into execution strategies without altering the underlying data access permissions.

WHY IT MATTERS

This update may streamline performance debugging for teams relying on Datasette for read-only SQL exploration.

[Read the original at Simon Willison's Weblog →](#)

AI Applications & Industry

3 stories ● 1 high ● 2 medium

08 **MEDIUM IMPACT** Ars Technica

Google Had a Mole Inside the Hacker Gang Behind History's Biggest Supply-Chain Attack

Google's threat intelligence team maintained an undercover analyst inside the TeamPCP hacker group throughout its historic supply-chain campaign, enabling real-time disruption of credential theft targeting over 1,000 companies.

Google's Threat Intelligence Group revealed that a Mandiant analyst infiltrated the inner circle of TeamPCP, a hacking gang responsible for history's largest software supply-chain attack. The undercover operator joined the group's core "CanisterWorm" chat in March, shortly after the campaign began, gaining access to servers storing more than 500,000 stolen user credentials. This infiltration allowed Google to monitor the group's operations from day one, observing compromises of open-source tools including Trivy, LiteLLM, TanStack, and infrastructure belonging to Checkmarx and Mistral AI. The breach cascade ultimately affected GitHub, Mercor, OpenAI employee devices, and the European Commission.

Upon discovering the trove of stolen access tokens and passwords, Google prioritized disrupting the extortion scheme over direct victim notification. The team contacted cloud providers like Amazon Web Services and Microsoft to revoke compromised credentials en masse, preventing the hackers from leveraging the data for ransom. During this period, Google also intercepted an AI-generated zero-day exploit developed by a TeamPCP member targeting login software two-factor authentication. Google tested the code, confirmed its viability, and alerted the vendor to patch the vulnerability before widespread exploitation occurred.

The investigation culminated in the arrest of two Australian nationals, Ruben Ian Thomson and Louis Michael Gaebler, following operational security failures by the group. Google traced Thomson after identifying a link between a hacker handle on BreachForums and a Gmail account used to back up stolen data to Google Drive. Additionally, the criminal group ShinyHunters, initially partnered with TeamPCP to monetize the stolen data, betrayed the gang by leaking internal chat logs to Google while attempting to cut out the original hackers from profit sharing. This operation marks a strategic shift for Google's newly formed Cyber Disruption Unit, moving from passive reporting to active interference in cybercriminal campaigns.

KEY FACTS

COMPANIES BREACHED

1000+

INFILTRATION START

March

CORE CHAT NAME

CanisterWorm

STOLEN CREDENTIALS

500000

GROUP MEMBERS ARRESTED

2

MALWARE WORM NAME

Mini Shai-Hulud

This case validates the efficacy of deep undercover infiltration for disrupting large-scale supply-chain attacks before monetization. It highlights the dual risk of AI tools accelerating zero-day development and the critical need for rapid credential revocation protocols among cloud providers.

[Read the original at Ars Technica →](#)

09

HIGH IMPACT

Simon Willison's Weblog

'Nobody Is Reading Anything': Life Inside a Company Where Claude Code Ships Everything

Engineers at Voxium have ceased reading code entirely, reducing their role to validating outputs generated autonomously by Claude Code across the full software lifecycle.

Internal reporting from Voxium dated 20 September 2026 reveals a development environment where human comprehension of the codebase has been effectively decoupled from shipping velocity. Since the adoption of Claude Code approximately two weeks prior, the tool now generates every artifact in the pipeline: specifications, source code, test suites, product requirement documents (PRDs), tickets, ticket resolutions, and reports. The workforce, spanning job levels L1 through L7, performs an identical function regardless of seniority: interacting with the model and executing its suggestions.

Management directives explicitly prioritize throughput over verification, operating on the premise that code generation is not the bottleneck. Consequently, engineering staff report working 12 to 13 hours per day solely to maintain the pace of approval, described metaphorically as just "pressing enter." The organizational culture has shifted to a state where no team member reads the generated content before it enters the repository. This uniformity of task suggests a collapse of traditional code review hierarchies, where senior engineers typically audit junior output; here, no one audits the machine's output.

The situation highlights a specific failure mode in aggressive AI deployment: the severing of the feedback loop between creation and understanding. While the company achieves high volume in terms of shipped items, the explicit statement that "nobody is reading anything" indicates a systemic risk where technical debt, security vulnerabilities, or logical errors may accumulate undetected. The timeline is compressed, with these operational conditions established within half a month of the author starting their role, suggesting rapid enforcement of this workflow by higher management despite team dissatisfaction.

KEY FACTS

DATE OF REPORT

20 September 2026

WORK HOURS

12 to 13 hours/day

GENERATED ARTIFACTS

Specs, code, tests, PRDs, tickets, resolutions, reports

TIME SINCE IMPLEMENTATION

Half a month

AFFECTED LEVELS

L1 to L7

Organizations prioritizing AI-driven shipping velocity risk eliminating human code review entirely, creating systems where engineers ship artifacts they do not understand and cannot debug.

[Read the original at Simon Willison's Weblog →](#)

10

MEDIUM IMPACT

TechCrunch

ScrollEd Turns Any Textbook Into a TikTok-Style Feed

Palo Alto startup ScrollEd is pitching a platform at TechCrunch Disrupt that transforms static text files into vertical, AI-generated video and audio feeds.

ScrollEd, a startup founded by Stanford student Utsav Gupta and University of Pennsylvania student Rebecca Neff, has developed an application that ingests dry text files and PDFs to output a scrollable feed resembling Instagram Reels or TikTok. The system parses input documents and regenerates the content as AI-produced video, audio, text blocks, or interactive quizzes. Navigation relies on vertical swipes to surface new topics and horizontal swipes to drill deeper into a specific subject, concluding with an assessment quiz. This architecture targets users who prefer short-form vertical media but require access to deep educational material.

The company operates on a freemium model comprising a free consumer feed, a paid ScrollEd Pro subscription, and annual institutional licenses. Educational institutions and corporate training programs pay for the ability to deliver material through the platform, receiving data reports on engagement and lesson progress in return. While the consumer business launches immediately at the TechCrunch Disrupt Startup Battlefield 200 competition, the team plans to focus the coming months on expanding the lesson library, refining source-checking workflows, and developing institutional pilots. Future infrastructure aims to support individual learner assessment and adaptive content delivery.

Gupta and Neff bootstrapped the venture this year after identifying personal dissatisfaction with doomscrolling habits. Unlike typical social platforms designed to maximize time-on-app, ScrollEd explicitly states it does not optimize for maximum engagement, intending instead to facilitate curiosity-driven transitions where going deeper is as frictionless as moving to the next video. The founders position the tool as a corrective social network that meets the 18-year-old demographic's preference for short-form content without sacrificing educational depth. The startup will present its full capabilities during TechCrunch Disrupt, scheduled for October 13-15 in downtown San Francisco.

KEY FACTS

FOUNDERS

Utsav Gupta and Rebecca Neff

EVENT

TechCrunch Disrupt Startup Battlefield 200

EVENT DATES

October 13-15

LOCATION

downtown San Francisco

LICENSE TYPES

Pro subscription and annual institutional licenses

Developers building educational tools may need to consider vertical-feed interfaces and AI-mediated content transformation to reach younger demographics effectively. The shift toward non-engagement-maximized metrics offers an alternative design pattern for learning platforms.

[Read the original at TechCrunch](#) →

Sources

- 01 MintAct Unifies UI Grounding, Navigation and Tool Use in One Visual Agent Family
<https://arxiv.org/abs/2609.22083>
- 02 LeCun's AMI Labs and Fei-Fei Li's World Labs Are Keeping Their World Models Secret
<https://techcrunch.com/2026/09/20/world-model-companies-are-keeping-a-lot-of-secrets/>
- 03 73,000 Reddit Posts Reveal What Users Actually Value When Delegating to AI Agents
<https://arxiv.org/abs/2609.22067>
- 04 APort Vault Replays 4,371 Real Attacks Against AI Agents Handling Your Payments
<https://arxiv.org/abs/2609.22076>
- 05 Simon Willison's llm-keys-ui Stops You Pasting API Keys Into Agent Sessions
<https://simonwillison.net/2026/Sep/20/llm-keys-ui/>
- 06 Vocci's 6-Gram Titanium Ring Records Your Meetings With a Double-Tap
<https://techcrunch.com/2026/09/20/voccis-ring-adds-a-new-form-factor-to-meeting-note-taking/>
- 07 datasette-explain 0.2.2 Brings Query Plans to Datasette's Stored Queries
<https://simonwillison.net/2026/Sep/20/datasette-explain/>
- 08 Google Had a Mole Inside the Hacker Gang Behind History's Biggest Supply-Chain Attack
<https://arstechnica.com/security/2026/09/an-undercover-google-analyst-infiltrated-a-notorious-supply-chain-hacking-gang/>
- 09 'Nobody Is Reading Anything': Life Inside a Company Where Claude Code Ships Everything
<https://simonwillison.net/2026/Sep/20/voxium/>
- 10 ScrollEd Turns Any Textbook Into a TikTok-Style Feed
<https://techcrunch.com/2026/09/20/scrolled-wants-to-turn-textbooks-into-tiktok/>

About this document. Every story in the 21 September 2026 New Horizon AI Digest, reported at length. Each entry is written from the publisher's own article text; where a source could not be retrieved the entry is explicitly marked and kept short rather than padded.

Images and licensing. Figures are used only where the source licence permits redistribution, and are credited in the caption. Publisher artwork is not reproduced. All titles link to the original publication.

Vocabulary. Common AI terms are defined at new-horizon.tech/glossary/.

Generated 2026-09-21T05:34:00Z · model qwen3.5:397b-cloud · New Horizon, new-horizon.tech · To unsubscribe or manage your subscription, see the footer of the daily email.