
TOP STORY

China rejects AI slowdown calls as US executives back restraint

Chinese officials dismiss warnings about AI risks from leading US tech executives who backed slowing development, citing fearmongering and confrontation.

BLOOMBERG · 23m ago

2 OpenAI runs safety checks before training runs amid AI slowdown push

Sam Altman doubles down on slowing AI development, with OpenAI implementing pre-training safety protocols and discussing...

THE DECODER · 2h ago

3 Deepfake pornography sites target 100-plus European politicians

Analysis of 160 deepfake websites reveals politicians in 22 countries appear on them, with nearly all victims being women.

WIRED · 53m ago

4 Intel-backed Buildots raises \$130M for construction AI

Construction AI startup Buildots closed a \$130 million funding round backed by Intel to accelerate projects like data centers, chip...

BLOOMBERG · 23m ago

5 OECD finds students using AI for schoolwork fall 1.5 years behind

Research shows that students using AI chatbots for homework fall significantly behind, though use is pervasive in Asian...

BLOOMBERG · 1h ago

RESEARCH

• World models accelerate autonomous research agents via learned environments

HF PAPERS · 4d ago

• Anthropic maps millions of concepts inside Claude Sonnet language model

ANTHROPIC · 2d ago

• Data-centric framework trains open-weight frontier cyber models at scale

HF PAPERS · 11h ago

China rejects AI slowdown calls as US executives back restraint

BLOOMBERG · 23m ago

China's spokesman Guo Jiakun rejected dire warnings about AI's risks after leading US technology executives—including Sam Altman, Elon Musk, and Demis Hassabis—backed calls from Anthropic CEO Dario Amodei to slow advanced AI development. The US executives framed the slowdown as necessary for safety and to prevent Chinese advantage from threatening American national security. Guo accused them of using "fearmongering, confrontation and vicious competition" as negotiating tactics. The clash reflects the geopolitical stakes now attached to AI development timelines.

WHY IT MATTERS

The disagreement over AI pacing has become a primary battleground in US-China competition, with national security framing on one side and safety arguments on the other.

KEY POINTS

- US tech leaders including Altman, Musk, Hassabis back Amodei's proposal to slow frontier AI development
- China's Guo Jiakun dismissed the call as fearmongering and confrontational strategy
- National security implications cited as justification for development pace constraints

OpenAI runs safety checks before training runs amid AI slowdown push

THE DECODER · 2h ago

OpenAI has begun conducting safety checks before major training runs as part of its commitment to pacing frontier AI development. According to The Information, the company has been in discussions with Anthropic and Google for months about establishing joint self-regulation frameworks. Altman has publicly committed to the slowdown, backed by support from other industry leaders. These moves represent concrete operational changes, though the industry debate over the actual necessity of a slowdown continues.

WHY IT MATTERS

Pre-training safety protocols and industry self-regulation attempts signal whether companies can credibly commit to measured development without external enforcement.

KEY POINTS

- OpenAI implementing mandatory safety checks before major training runs
- Months-long discussions with Anthropic and Google on joint self-regulation frameworks
- Part of broader industry pushback from Trump administration and others skeptical of slowdown calls

Deepfake pornography sites target 100-plus European politicians

WIRED · 53m ago

A Wired investigation of 160 deepfake websites found that politicians across 22 European countries appear on sexually explicit deepfake sites. Nearly all of the targets are women. The scale and targeting pattern reveal a coordinated exploitation vector using AI-generated imagery. The sites represent a concrete misuse case of synthetic media generation tools, affecting public officials and political discourse.

WHY IT MATTERS

Deepfake-based harassment and political manipulation demonstrate near-term harms from generative AI that require immediate policy and technical countermeasures.

KEY POINTS

- 160 deepfake websites analyzed covering 22 European countries
- Nearly all 100-plus targeted politicians are women
- Sexually explicit deepfakes used as harassment and potential political interference tool

Intel-backed Buildots raises \$130M for construction AI

BLOOMBERG · 23m ago

Buildots, a startup using AI to speed construction of expensive mega-projects, has raised \$130 million in new financing with Intel backing. The company serves projects including data centers, chip factories, and hospitals. The investment reflects growing capital flowing into AI applications for industrial processes where delays are extremely costly. Buildots' models help optimize construction timelines and resource allocation.

WHY IT MATTERS

Deployment of AI to supply-chain-critical sectors like data center construction directly impacts compute capacity expansion and infrastructure timelines.

KEY POINTS

- Intel-backed funding round of \$130 million for construction optimization AI
- Targets data centers, chip factories, hospitals—critical infrastructure projects
- AI used to reduce delays and optimize resource allocation in expensive mega-projects

OECD finds students using AI for schoolwork fall 1.5 years behind

BLOOMBERG · 1h ago

An OECD study found that students who use AI chatbots for schoolwork fall approximately 1.5 years behind their peers in academic performance. Despite the findings, AI chatbot use for schoolwork remains pervasive in Asian economies. The research suggests unguided AI use in education correlates with learning deficits, conflicting with other studies showing structured AI guidance improves outcomes. The discrepancy points to the importance of pedagogical context.

WHY IT MATTERS

Evidence of AI's impact on human learning and skill development is critical for policy decisions on technology adoption in schools and workforce training.

KEY POINTS

- OECD study shows 1.5-year performance gap for students using AI chatbots for homework
- Pervasive AI use in Asian economies despite negative learning correlations
- Results suggest unguided AI use harms outcomes; structured guidance shows different results

World models accelerate autonomous research agents via learned environments

HF PAPERS · 4d ago

The paper introduces World Model RL, which replaces costly direct environment execution with a learned world model during post-training of autonomous research agents. By simulating environment interactions within a learned model, the approach reduces computational overhead and accelerates agent development. Debiasing and denoising techniques further improve quality. This work addresses the fundamental computational bottleneck in training agents that must interact with expensive or slow external systems.

WHY IT MATTERS

Learned world models for agent training enable scaling research automation without proportional increases in compute cost, unlocking broader deployment of autonomous research systems.

KEY POINTS

- World Model RL replaces expensive environment execution with learned simulation
- Debiasing and denoising improve post-training efficiency for research agents
- Reduces computational bottleneck in autonomous research agent development

Anthropic maps millions of concepts inside Claude Sonnet language model

ANTHROPIC · 2d ago

Anthropic researchers have identified how millions of concepts are represented and organized inside Claude Sonnet, one of their deployed production-grade language models. This represents the first comprehensive map of internal concept representation in a modern, large-scale LLM. The work provides interpretability insights into how language models structure semantic knowledge. Understanding these internal maps is foundational for safety, alignment, and steering of large models.

WHY IT MATTERS

Detailed mechanistic interpretability of production LLMs is essential for building safety measures, improving alignment, and understanding model capabilities and limitations.

KEY POINTS

- First detailed map of concept representation inside production-grade LLM
- Millions of concepts identified and structurally analyzed in Claude Sonnet
- Foundational work for interpretability-based safety and alignment research

Data-centric framework trains open-weight frontier cyber models at scale

HF PAPERS · 11h ago

The Feyospace-v1 framework demonstrates how a data-centric approach with specialized systems for reasoning analysis, cost reduction, and execution verification enables small teams to train open-weight cyber security agents. These models achieve top-tier performance on cyber benchmark suites, matching closed-source frontier models. The framework emphasizes efficient data curation and targeted optimization over raw compute scaling. The work shows how methodological improvements can reduce barriers to frontier-model development.

WHY IT MATTERS

Showing that frontier-class performance in specialized domains is achievable through data-centric methods and small teams has implications for democratization and competition in AI development.

KEY POINTS

- Data-centric framework achieves top-tier cyber model performance without massive compute
- Small teams can train open-weight agents matching proprietary frontier models
- Specialized systems for reasoning, cost reduction, and verification drive efficiency