

Gemini benchmark gains, ARC overfit concerns, and rising pressure on trust (deepfakes, peer review, robots)

AI News Digest

2026-02-15

Gemini benchmark gains, ARC overfit concerns, and rising pressure on trust (deepfakes, peer re- view, robots)

By AI News Digest • February 15, 2026

Benchmark headlines (Gemini 3 Deep Think, “First Proof”) are landing alongside fresh scrutiny of evaluation shortcuts on ARC. Meanwhile, deepfake policy pressure and academic integrity tactics escalate, and robotics/local inference updates show tangible progress in deployment and efficiency.

Frontier capability claims meet tougher evaluation ques- tions

Gemini 3 Deep Think posts new ARC-AGI-2 + “Humanity’s Last Exam” numbers

Sundar Pichai said **Gemini 3 Deep Think** received a “significant upgrade,” refined “in close partnership with scientists and researchers” to tackle real-world challenges [1]. He also reported **84.6% on ARC-AGI-2** and **48.4% on Humanity’s Last Exam (without tools)** [1].

Vinod Khosla separately pointed to an **85.28% SOTA** result on **ARC-AGI-2’s public eval set**, linking to a Symbolica write-up: <https://www.symbolica.ai/blog/arcgentica> [2].

Why it matters: As headline benchmark numbers climb, attention is shifting to *how robust those gains are*—and what they say (or don’t say) about general reasoning.

ARC evaluation nuance: performance drops when the input encoding changes

François Chollet flagged an “interesting finding” that frontier model performance on **ARC** can be tied to a “familiar input distribution,” suggesting models are **overfitting to the original ARC encoding format** due to extensive direct targeting of the benchmark [3]. In a related exchange, Melanie Mitchell’s team reported that when ARC inputs are re-encoded from numbers into other symbols, **accuracy goes down**, and they’ve identified other potential “short-cuts” (results to be published) [4].

Chollet argued that for a truly intelligent agent, re-encoding should be a **performance no-op**—e.g., a person would decode binary and still multiply correctly, with no accuracy loss even through multiple indirection steps [5].

Why it matters: This is a concrete reminder that “benchmark mastery” can reflect familiarity with presentation format—not just underlying task competence—pushing evaluators toward adversarial and distribution-shift testing.

OpenAI’s “First Proof” benchmark: early research-math results (with verification caveats)

Sam Altman said the “**First Proof**” challenge is meant to evaluate next-generation AI via **novel frontier research** problems [6, 7]. OpenAI reported running an internal model (with limited human supervision) on **10 proposed problems**; based on expert feedback, they believe **at least six solutions (2, 4, 5, 6, 9, 10)** have a high chance of being correct, while noting the problems are domain-specific and hard to verify [6].

They described the run as a one-week “side-sprint,” mostly done by querying a training model, with some manual back-and-forth between the model and ChatGPT for verification/formatting, and selection of “best” attempts by human judgment [6]. OpenAI said solution attempts will be published after midnight (PT) and shared the PDF sha256 hash: `d74f090af16fc8a19debf4c1fec11c0975be7d612bd5ae43c24ca939cd272b1a` [6]. More: <http://1stproof.org> [6].

“These are obviously not earth-shattering results, but the ability to produce genuinely new knowledge, however small, is a significant milestone and I hope we all take it seriously, with excitement and caution.” [8]

Why it matters: “Frontier research” evals are becoming a first-class measurement target—but the verification bottleneck (and methodology transparency) is part of the score.

Coding tools: GPT-5.3-Codex is being praised for better UI design “taste”

Greg Brockman highlighted chatter that **GPT-5.3-Codex** has “significantly better taste for UI design” [9, 10]. A related post predicted it will be **#1 on @designarena** once an API is available [9].

Why it matters: This aligns with a broader theme that when output becomes cheap, differentiation shifts toward *selection and judgment*: Paul Graham argued “taste” becomes the differentiator [11], and Khosla added “agency” as another key human attribute [12]—also endorsing the idea that people with curiosity and agency who fully use AI can gain expertise faster as traditional advantages become “cheap commodities” [13, 14].

Trust, misuse, and integrity pressures

Gary Marcus calls for a federal ban on AI impersonation

Gary Marcus urged “a federal law forbidding AI from impersonating humans” [15]. He cited examples ranging from deepfake videos (e.g., “Tom Cruise fighting Brad Pitt”) and tooling that pairs systems like OpenClaw with voice synthesis [15] to a reported scam where a deepfaked video of Mark Carney allegedly led to a victim losing **hundreds of thousands of dollars** [15].

Marcus referenced Daniel Dennett’s 2023 “counterfeit people” argument—warning that AI-enabled imitation can undermine societal trust and calling for outlawing the creation and passing along of counterfeit people [15]—and argued that generative systems may be built for mimicry, but mimicry has advanced enough that action is urgent [15].

Why it matters: The policy focus is sharpening from generic “deepfakes are bad” to concrete rules about **presentation, consent, and enforcement**.

Peer review goes adversarial: hidden prompt injections in ICML journal review packets

A report shared by Sara Hooker says **ICML journal editors added hidden prompt injections** to papers sent to reviewers, designed to detect AI use by instructing the AI to include two specific phrases in the review [16]. A reviewer reportedly noticed and was about to desk-reject, assuming the author added it [16]. Hooker summarized the dynamic as “Adversarial vibe checks” [17].

Why it matters: This is a live example of how AI use is changing academic workflows—pushing integrity checks into the document layer, and increasing the risk of misunderstanding and process breakdown.

Robots and “acting in the world”

Boston Dynamics’ Atlas reaches autonomous factory work (teleop + simulation-driven learning)

A 60 Minutes segment described a new-generation **all-electric Atlas** with an “AI brain powered by Nvidia’s advanced microchips,” designed to perform autonomous feats [18]. Boston Dynamics explained its approach emphasizes teaching via demonstrations and machine learning, including teleoperation to generate training data [18] and large-scale simulation (e.g., **4,000 digital Atlases trained for six hours** in a jumping-jacks task with added environmental challenges) [18].

The segment also showed Atlas at a **Hyundai plant**, described as the first time Atlas had been out of the lab “doing real work,” **sorting roof racks autonomously** without human help [18]. It noted competitors (Tesla, Amazon/Nvidia-backed startups, and state-supported Chinese companies) [18] and cited a Goldman Sachs projection that the **humanoid market could reach \$38B within the decade** [18].



Our latest reports on robots / 60 Minutes Full Episodes (1:44)

Why it matters: The story here is less “cool demo,” more “training and deployment pipeline”—teleop + simulation + on-site autonomy—aimed at repeatable factory tasks.

Hassabis: multimodal assistants + safety coordination amid racing dynamics

The same program highlighted “Astra,” described as an app that “takes in the world” [18], with a demo in which the system interprets paintings and generates a narrative around an image [18]. It also showed a real-world assistance flow (identifying a building and answering follow-ups about its original purpose and coal pollution) [18].

Demis Hassabis said DeepMind is training Gemini to not only “reveal the world” but also to **act in it** (e.g., booking tickets, shopping online) [18], and predicted robotics could have a “breakthrough moment” in the next couple of years with useful humanoids [18]. On safety, he outlined two worries—bad actors repurposing systems, and maintaining control/alignment as systems become more autonomous—warning that competitive races may incentivize cutting corners on safety, and arguing for international coordination [18].

Why it matters: Product direction (agents that perceive and act) and governance concerns (coordination, alignment, safety incentives) are now being discussed in the same breath by leading labs.

Open-source + local inference updates

AdaLLM: NVFP4-first inference on RTX 4090, with FP8 KV cache end-to-end

A release on r/LocalLLM introduced **AdaLLM**, an open-source vLLM fork that aims to make **NVFP4 weights usable on RTX 4090 (Ada GPUs)** via a pure NVFP4 fast path: FP8 KV cache, custom FP8 decode kernel, and **no silent FP16 fallback** [19]. The author said it currently targets **Qwen3** (dense + MoE) and **Gemma3** (including sliding-window layers) [19], with code here: <https://github.com/BenChaliah/NVFP4-on-4090-vLLM> [19].

Reported benchmarks included **Qwen3-8B-NVFP4 at 469 tok/s (batch=16) with 7.56 GB peak VRAM** [19] and **Gemma3-27B-it-NVFP4 at 53.7 tok/s (batch=4) with 19.84 GB peak VRAM** [19]. The author also observed **~2.4× lower peak VRAM** vs FP16 for Qwen3-8B, with ~20–25% throughput loss [19], and a user reported the NVFP4 Gemma3-27B fit and “runs smoothly” [20].

Why it matters: This is practical infrastructure work that can change what fits on a single high-end consumer GPU—especially for larger instruction-tuned models.

```
python hf_ptq.py \  
--pyt_ckpt_path google/gemma-3-27b-it \  
--qformat nvfp4 \  
--kv_cache_qformat fp8 \  

```

```
--export_fmt hf \  
--export_path ./gemma-3-27b-it-nvfp4-fp8kv \  
--calib_size 512 \  
--trust_remote_code
```

[21]

Kyutai releases Hibiki-Zero (3B) for simultaneous speech-to-speech translation

Kyutai released **Hibiki-Zero**, a **3B-parameter** simultaneous speech-to-speech translation model trained using **GRPO reinforcement learning** without any word-level aligned data [22]. Code: <https://github.com/kyutai-labs/hibiki-zero> [22].

Why it matters: It’s another signal that non-trivial speech workflows are moving toward open releases with clearer training recipes.

Ecosystem signals (brief)

- **AI startup consolidation:** @swyx said it’s “consolidation season,” connecting small (<6 person) startups—especially in coding agent harnesses, developer UI, RL/inference infra, and RL/benchmark evals—with acquire interest [23].
- **Recognition:** Jeff Dean congratulated Demis Hassabis, John Jumper, Peyman Milanfar, Noam Shazeer, and Moti Yung on being elected to the **National Academy of Engineering**, calling it among the highest professional distinctions for engineers (NAE list: <https://www.nae.edu/345149/NAENewClass2026>) [24, 25].
- **Organizational-speed rhetoric:** Elon Musk promoted xAI’s “no organizational overhead” approach (no docs/approval chains), and argued it’s needed to compete with China—claiming China “owns 50% of the world’s AI researchers” and removes organizational barriers between ideas and execution [26].

Sources

1. X post by @sundarpichai
2. X post by @vkhosla
3. X post by @fchollet
4. X post by @MelMitchell1
5. X post by @fchollet
6. X post by @merettm
7. X post by @sama

8. X post by @sama
9. X post by @tkkong
10. X post by @gdb
11. X post by @paulg
12. X post by @vkhosla
13. X post by @DeryaTR_
14. X post by @vkhosla
15. X post by @GaryMarcus
16. X post by @paul_cal
17. X post by @sarahookr
18. Our latest reports on robots | 60 Minutes Full Episodes
19. r/LocalLLM post by u/Educational_Cry_7951
20. r/LocalLLM comment by u/PomeloSweet926
21. r/LocalLLM comment by u/Educational_Cry_7951
22. r/LocalLLM post by u/techlatest_net
23. X post by @swyx
24. X post by @NewsFromGoogle
25. X post by @JeffDean
26. X post by @r0ck3t23