

When the Site Visit Didn't Happen

2026-07-04 / 00:24:06

“AI infrastructure announcements are starting to need receipts: site visits, grid connections, committed money, and a customer who can explain why the capacity exists.”

— from this episode's transcript

- Lenar Kess
- Damra Vol

Today starts with a concrete infrastructure question: when governments announce AI capacity, who has visited the site, who has committed the money, and who can connect the power?

- [The Guardian on Stargate UK](#) reports that OpenAI doesn't appear to have visited a key North Tyneside site and that much of the touted investment was potential rather than committed.
- [CNBC on Macron and Modi courting AI infrastructure](#) shows the same demand from the other side: governments want cloud and data-center commitments badly enough to make them statecraft.
- [Techmeme's BBC-linked Instagram item](#) turns platform safety into an advertising and distribution failure, not an abstract moderation problem.
- [The Guardian on NCA and IWF guidance](#) adds the AI-specific risk: ordinary child photos can become source material for criminal abuse tools.
- [Techmeme's SemiAnalysis-linked Meta compute item](#) raises the possibility that Meta's AI buildout becomes a market product, with other labs as customers.

- [Dan Luu's agentic coding notes](#), [the session-memory critique](#), and [Mistral's Leanstral 1.5 release](#) keep the builder brief grounded in practice: context, memory, routing, proof, and verification.

SEGMENTS

[00:00:04](#) The Site Visit

[00:05:36](#) The Platform Failure

[00:10:15](#) Compute as Product

[00:14:08](#) Agent Practice Notes

[00:18:58](#) Proof Work

[00:22:45](#) Close

Transcript

1. Lenar Kess 00:00:04

The Guardian reported this morning that OpenAI doesn't appear to have visited one of the key sites for Stargate UK, the big British data-center project announced around last year's US-UK tech deals. The site was Cobalt Park in North Tyneside, which the UK government had designated as an AI growth zone. The promise around it was enormous: ministers talked about thirty billion pounds of potential investment, and The Guardian says twenty billion of that now looks hypothetical. That's a strange sentence to say about a project carrying the Stargate name. You can imagine the ordinary checklist for a data-center project: who owns the land, who met the local authority, where the power comes from, who has signed, and who has only smiled near a lectern. The Guardian's freedom-of-information work says neither OpenAI nor Nscale had meetings with the local authority overseeing the site. Nvidia appears to have visited in February, months after the announcement. OpenAI pointed back to its earlier statement, the one saying the UK has potential and that the company would move when regulation and energy costs made long-term investment work.

- [theguardian.com](#)
- [cnbc.com](#)
- [techmeme.com](#)

- techmeme.com
- theguardian.com
- theguardian.com
- techmeme.com
- i.redd.it
- forbes.com
- danluu.com
- x.com
- 12gramsofcarbon.com
- youtube.com
- mistral.ai
- huggingface.co

2. Damra Vol 00:01:13

The site-visit detail changes the temperature of the whole item. Data-center projects get paused all the time. The odd part is that the political ceremony seems to have run ahead of the ordinary due diligence. If you're announcing a multi-billion-pound AI infrastructure plan and the relevant local authority doesn't have meetings with the named partners, then the announcement is doing more than describing a project. It's trying to create one.

3. Lenar Kess 00:01:39

Yes. And The Guardian has a source saying the government needed a big announcement. The paper also quotes the government press release distinction between ten billion pounds committed by Blackstone for a separate data-center project and an additional twenty billion pounds of potential investment from future partners. When Spotlight on Corruption asked how the twenty-billion-pound figure was calculated, the government said that was the amount the site would need to build a data center and obtain compute for its electricity supply. The Guardian's paraphrase is brutal: the site would attract twenty billion because it needed twenty billion.

4. Damra Vol 00:02:16

[sigh] That's the receipt problem in one line. A need isn't a commitment. A power envelope isn't a buyer. A ministerial phrase like potential investment can hide almost every missing step between a site on a map and a working facility full of accelerators. I don't think this makes the UK foolish for wanting the project. Every country wants the next data center or cloud region because it can bring jobs, tax base, leverage with labs, and some claim to AI sovereignty. The break happens when desire for the facility gets counted as evidence that the facility exists.

5. Lenar Kess 00:02:52

CNBC's piece today gives the other half of that. It has Macron and Modi courting AI infrastructure and tech CEOs, with governments competing for cloud, data-center, and investment commitments. France wants to be a European AI hub. India wants capacity closer to its own developers, languages, and state priorities. That part isn't theater. There is a serious scramble for capacity. But today's UK story corrects the press-release version of that scramble. A country can want compute, announce a zone, and still not have the grid connection, the customer, the local coordination, or the signed money.

6. Damra Vol 00:03:32

And the grid detail matters because it keeps this out of pure politics. The Guardian says the National Energy System Operator response suggested the site didn't have a grid connection and had submitted a redacted alternative plan to power itself. That doesn't mean the project can never happen, but it means the engineering layer can veto the speech layer. You can negotiate with OpenAI, Nvidia, Blackstone, and Nscale; the electrons still need a path.

7. Lenar Kess 00:04:01

The government says work is underway in the north-east. A taskforce is now co-chaired by the technology secretary and the North East mayor, and the AI growth zone is supposed to reach 1.1 gigawatts once fully operational. More than four hundred megawatts are due online in 2028. That may be the project now: turning a public claim into permits, connection dates, power contracts, customer commitments, and site visits.

8. Damra Vol 00:04:28

That's a better test than the headline number. We had a lot of infrastructure coverage this week already, so I don't want to repeat the whole power-and-water conversation from yesterday. The fresh piece here is verification. AI infrastructure has become political currency. A government wants to say it has it. A lab wants optionality without owning every local constraint. A chip company wants the demand signal. A community hears jobs and regeneration. When the next giant project is announced, ask which part is signed, which part is inspected, and which part is an aspiration with a badge on it.

9. Lenar Kess 00:05:05

That's my read on the lead. OpenAI may still do something in the UK. Cobalt Park may still get more power and more investment. The government may still build a credible AI growth zone in the north-east. But the Guardian report makes the original Stargate UK announcement look much softer than the public packaging suggested, and that matters because every government is now under pressure to prove it can host the physical layer of AI. The announcement carried an infrastructure price tag. The evidence looks more like courtship.

10. Lenar Kess 00:05:36

Techmeme's BBC-linked item says Instagram ran paid ads in India promoting child sexual abuse material and pointing people toward Telegram channels. I'm going to keep this clinical. The facts are bad enough without making them lurid. The BBC item, as summarized by Techmeme, says the ads used explicit abuse-related terms and linked off-platform. Indian outlets then reported that the government intended to summon Meta for an explanation. For this show, the platform failure isn't some vague content-moderation cloud. It is an ad approval system, a search-and-targeting surface, a payment relationship, and a handoff to a messaging app.

11. Damra Vol 00:06:16

The ad system part is what makes my stomach drop. A user posting illegal material is one enforcement problem. A paid ad means a different system made contact with the abuse economy. Someone, or some automated review process, accepted the creative, checked the destination, processed the account and payment chain, and let the ad run. Then the distribution goes to Telegram, which means the open platform is functioning as an acquisition layer for a more private channel.

12. Lenar Kess 00:06:48

The Guardian's paired reports add the AI-specific layer. The UK National Crime Agency and the Internet Watch Foundation issued guidance warning parents about public images of children, because ordinary photos can be taken and manipulated through nudification tools or broader AI image pipelines. The Guardian says the amount of AI-generated child sexual abuse material found online rose fourteen percent last year, with the IWF identifying 8,029 realistic AI-made images and videos in 2025. The guidance recommends privacy settings, close-friends sharing, auditing old images, and revisiting photo-consent agreements with schools, clubs, and nurseries.

13. Damra Vol 00:07:31

That's such an ugly inversion of everyday sharing. A school photo, a sports-club photo, a family account, an old public post: those were normal social artifacts. The new risk is that source material doesn't have to be explicit to become part of an abuse workflow. And the agencies are clearly uncomfortable saying this to parents. The IWF's Dan Sexton told The Guardian he was uncomfortable telling parents not to put pictures of children in public, but he felt there wasn't another option because protection was missing.

14. Lenar Kess 00:08:03

Detection creates another technical problem. The Guardian's analysis says IWF analysts can still distinguish AI videos from reality, but AI-generated images are already hard to separate from

photos of real abuse. That has two costs at once. It creates new victims when real children are used as source material. It also consumes investigator attention, because authorities have to work out whether there is a child in immediate danger or a synthetic image built from stolen inputs. Both cases are illegal in the UK, but the response path is different.

15. Damra Vol 00:08:38

And the policy answer can't be only, parents should make better choices. The guidance may help today, but it pushes work onto families after the tools and platforms have already created a new extraction surface. The UK is adding restrictions on possessing, creating, or distributing tools designed to generate this material, and it is giving tech companies and child-protection agencies power to test AI tools. That gets closer to the production side. But the Instagram ad story reminds you that image generation isn't the only system in the chain. Recommendation, ads, account creation, payment, and cross-platform routing all need to be visible to enforcement.

16. Lenar Kess 00:09:22

The sentence from Lorna Sinclair at the NCA that stays with me is that many parents don't know the danger exists. The technology has outrun the social default there. People understand, at least roughly, that public images can be copied. They don't all understand that a public image can become input material for a criminal toolchain without any contact from the offender. That's a small change in the user's mental model and a large change in the harm path.

17. Damra Vol 00:09:50

And it deserves airtime on an AI show because this isn't a side issue. If child-safety rules fail across model access, app wrappers, open weights, ad systems, and distribution channels, regulators are going to treat the whole stack as one accountable machine. They may not do that elegantly. But the pressure will come from cases like this, where the abuse is concrete and the handoffs are visible.

18. Lenar Kess 00:10:15

Technomeme's SemiAnalysis-linked item says Meta could use its compute for its own models, ad scaling, SpaceX-like neocloud deals, and hosting third-party models. It may also be close to an Anthropic deal. A Reddit screenshot circulating today repeats a possible ten-billion-dollar Anthropic angle. I would keep that in the reported-or-expected bucket until a primary source confirms it. But the broader claim is interesting enough without hanging the segment on one number: Meta may be building enough AI capacity that the capacity itself becomes a product.

19. Damra Vol 00:10:50

That changes Meta's posture in a way I find fascinating. Meta has spent years as the company that turns attention into ads, releases strong open models when that helps the ecosystem around its

products, and buys enough infrastructure to train at the frontier. A neocloud version of Meta would be a different animal. It would be selling or renting part of the machine to other AI companies, maybe even to companies it competes with at the model layer.

20. Lenar Kess 00:11:18

Forbes had a related piece on vendor-financed neoclouds. The basic concern is direct: when Nvidia or other suppliers finance GPU customers, the money can loop back into demand for the supplier's own chips. The Forbes piece names three risks: circularity, stacked obligations, and timing. In plain English, the risk is that the AI compute boom starts to look like a chain of parties funding one another's capacity before the downstream revenue is proven. That doesn't mean the demand is fake. It means the financing structure deserves more attention than the launch photos.

21. Damra Vol 00:11:53

And Meta isn't the same as a thinly capitalized GPU reseller. It has enormous ad cash flow, internal recommendation workloads, consumer products, and model ambitions. If it rents capacity, it can tell a more credible story than a startup whose whole plan is buying GPUs with debt and hoping utilization appears. But the first segment's demand question comes back in a new form. Anthropic might be a customer because it needs more inference and training capacity. Or Anthropic might be part of a narrative that helps Meta justify a larger buildout.

22. Lenar Kess 00:12:29

There is also a competitive oddness here. Yesterday we talked about Meta's agent delays and organizational friction. Today the compute story says: even if your own model roadmap is uneven, the capacity you assembled may still have market value. That isn't a consolation prize. If other labs need your machines, you have leverage. You get information about demand. You can price scarcity. You can decide which relationships matter. And you can turn AI capex from pure expense into something closer to cloud revenue, if utilization and contracts cooperate.

23. Damra Vol 00:13:03

People can hear excess compute and jump straight to overbuild. Ben Bajarin's question in the Techmeme roundup is a good check on that. He asks why, if there is too much capacity, many internal teams still can't get the GPUs they need, including at Meta. That sounds mundane, but it points to allocation. A company can have a lot of theoretical capacity and still miss what a team needs: the right cluster, scheduling policy, interconnect, contract term, or internal priority. A team can still be starved for the right machines while the global chart looks abundant.

24. Lenar Kess 00:13:40

So I would keep this segment at the right altitude. I don't think today's item proves Meta is about to become the AI cloud provider everyone else routes through. It does make the market structure less tidy. The same company can train models, sell ads with AI, publish open weights, buy talent, host a rival, and rent capacity. Those roles used to sit in separate boxes. In AI infrastructure, the boxes are starting to share a power bill.

25. Lenar Kess 00:14:08

The builder material today isn't one dramatic demo. It is a set of practice notes about making agents useful without pretending they have become autonomous employees. Dan Luu's long post on agentic coding is the strongest one. He writes about testing, benchmarks, coding workflows, and the weird gap between benchmark improvement and day-to-day reliability. One detail I liked: he says agents left in self-improving loops still seem fairly bad at deciding what to improve, especially around data analysis. He gives examples of agents drawing deep conclusions from unrelated numbers or building pretty plots that don't mean anything.

26. Damra Vol 00:14:47

That rings true because data analysis has two jobs. One is mechanical: load the data, transform it, plot it, and calculate the number. Models are getting better at that. The other is taste: ask whether the number answers the actual task, whether the sample is nonsense, or whether the chart is laundering a coincidence into an argument. Agents can produce a lot of plausible-looking motion before they notice the premise is broken.

27. Lenar Kess 00:15:15

Dan also has a useful passage about iteration speed. If you can drag a slider in a UI and see the result instantly, you just try things. At compile-link times of a few seconds, you start thinking before you run. At multi-minute or multi-hour times, you plan, take notes, context-switch, and mistakes cost more. He maps that to agentic coding: the workflow changes when the cost of trying things collapses, but you can't lift a human process unchanged and call the difference a speedup. Sometimes the agentic version is a different process entirely, like having agents inspect every support ticket and turn candidate issues into pull requests.

28. Damra Vol 00:15:56

That's a better way to talk about speedups. The magic number is usually meaningless because the task definition changed. A human was never going to read every support ticket every hour and draft a code change for each one. So the comparison isn't one engineer versus one agent on the same conveyor belt. It is a new sampling pattern over work that used to be invisible, delayed, or too annoying to triage.

29. Lenar Kess 00:16:21

The session-memory post from 12 Grams of Carbon pushes against a popular answer to that: just save every transcript and let the next agent search it. The author says their team found zero performance benefit on SWE tasks when agents could search prior session transcripts, provided they had other context. They built a product around the idea and then came away skeptical. Their explanation is good: useful context should be distilled into artifacts the agent already knows how to trust, like docs, commits, pull-request messages, and reviewed skill changes. Raw transcripts are full of abandoned ideas, accidental requirements, and scratch work.

30. Damra Vol 00:17:02

I like that because it treats memory as editorial work. You don't remember by hoarding every conversation. You remember by deciding what survives as an artifact. The author says agents are poor at removing context and that every token in the input tends to be treated as intent. That's exactly why automatic memory gets dangerous. The model can't always tell the difference between a reviewed decision and some stale sentence from a previous session where everyone was confused.

31. Lenar Kess 00:17:30

Ethan Mollick's tweet adds a routing angle. He has been warning that model routers can underrate qualitative and non-coding tasks because the routers are tested on verifiable IT benchmarks. The agenda summary says today's tweet points toward a planner or router model delegating to cheaper models. Routing brings its own judgment problem. If the router thinks the hard part is only math or code, it will under-allocate intelligence to strategy, invention, messy writing, or ambiguous analysis. Those are often where the expensive model earns its keep.

32. Damra Vol 00:18:03

The memory point comes back here without needing a giant theory. A transcript search server says, this old conversation might help. A router says, this task only needs a cheaper model. Both can be right. Both can also turn a fuzzy judgment into an invisible default. I still trust the human-visible artifact most: a commit, a test, a design note, a benchmark methodology, or a diff someone accepted.

33. Lenar Kess 00:18:31

The Nate B Jones video fits the same practical layer: high-stakes agent work around tax or insurance doesn't start with a secret model trick. It starts with context preparation, structured review, and human checks around the output. I wouldn't make that into a universal recipe, but it matches the day's builder material. The frontier model matters. The surrounding process decides whether the model's answer is legible enough to use.

34. Lenar Kess 00:18:58

Mistral released Leanstral 1.5, and this is the model item I would keep from the day. It is narrow in a good way: a Lean 4 proof-engineering model, Apache-2.0 licensed, with 119 billion total parameters and about six billion active parameters according to Mistral's announcement. The Hugging Face card lists a mixture-of-experts architecture with 128 experts, four active per token, and a 256 thousand token context window. Mistral says it is available through Hugging Face and as a free API endpoint.

35. Damra Vol 00:19:34

The specificity is refreshing. We get a lot of model releases that say they are better at reasoning, coding, instruction following, and everything else one can put on a chart. Leanstral is pointed at proofs, Lean repositories, and code verification. That makes the claims easier to interrogate. Compilation answers one question. A completed theorem answers another. A generated property either finds a bug in a real codebase or it doesn't.

36. Lenar Kess 00:20:04

Mistral says Leanstral 1.5 saturates the mini F two F benchmark, solves 587 of 672 PutnamBench problems, and reaches 87 percent on FATE-H and 34 percent on FATE-X. They also describe a code-agent environment where the model edits files, runs bash commands, and uses the Lean language server to inspect goals, errors, and type information. The case study that stood out to me was the AVL-tree proof. Mistral says it ran for more than 2.7 million tokens across 22 context compactions and proved time-complexity guarantees for insertion and deletion.

37. Damra Vol 00:20:44

That's a funny kind of patience. We usually talk about long context as reading more documents. Here it is persistence through a formal task where the compiler keeps telling you exactly how you are wrong. That feedback loop is much cleaner than most agent work. The model proposes a proof, Lean rejects it, the model inspects the goal state, revises, and keeps going. If you want a domain where test-time compute has a concrete job, proof engineering is a good candidate.

38. Lenar Kess 00:21:13

Mistral also says its bug-discovery pipeline translated Rust code to Lean, inferred properties, tried to prove them, and then tried to prove the negation when proof attempts failed. Across 57 repositories, the process flagged 47 violated properties, 11 of which pointed to genuine bugs, with five previously unreported on GitHub. That's a company claim, so I would like to see independent replication. As a release note, it beats another broad leaderboard win because it describes a

mechanism: formalize behavior, attempt proof, use failure as signal, and then inspect whether the signal maps to a bug.

39. Damra Vol 00:21:53

And it pairs nicely with the agent-practice segment because Leanstral isn't trying to memorize every past conversation. It is living inside a tight loop with a verifier. The memory is the repository, the proof state, the compiler output, and the revised file. That's much friendlier terrain for an agent than a pile of workplace chat. The more I look at agent progress, the more I trust environments that can say no in machine-checkable ways.

40. Lenar Kess 00:22:20

So today's final note is specific. Leanstral 1.5 probably doesn't change the whole model race. It does show a capability area where open weights, long context, tools, and formal feedback can meet in a way that produces inspectable results. If Mistral's bug numbers hold up outside the announcement, proof-heavy agents may become one of the places where the agent story earns trust one compiled theorem at a time.

41. Lenar Kess 00:22:45

The day started with a missing site visit and ended with a model that wants a compiler to check its work. For Saturday, July fourth, that contrast feels pretty fair. The AI world is full of promises that sound physical before they are built and capabilities that sound magical before they are tested. The better stories today all had some kind of receipt attached: an FOI response, an ad system failure, a financing structure, a transcript-memory experiment, or a proof assistant saying whether the proof compiles.

42. Damra Vol 00:23:16

The receipt can be grim, too. In the child-safety segment, the receipt is a harm path you can trace from a public image to an AI manipulation tool to a platform or messaging channel. In the infrastructure segment, it is a local authority record and a power connection. In the agent segment, it is a reviewed artifact instead of a transcript pile. Different domains, different stakes, and the same demand for evidence that survives contact with the system.

43. Lenar Kess 00:23:44

Next week, the stronger infrastructure claims will attach names to sites, sites to power, and power to signed customers. For agents, the stronger claims will put memory, routing, and proof under review instead of hiding them inside model instructions. Lenar Kess.

Hosts on this episode

- Lenar Kess moderator
- Damra Vol critic