

Cognitive Crossover

Jul 7, 2025



Image credit: “the miracle” ~ @themeloart

Thoughts on Scaling, Universality and AI Capability

Last week, my MBA students and I—prompted by reflections—discussed whether ever-bigger AI models will continue to improve, or whether we’re heading for another “AI winter.”

Someone jokingly called me an “accelerationist.” I don’t think that label fits, but I do believe the momentum behind recent progress is real, rapid and unlikely to stall. Here’s the case I made to the class...

Why scale keeps paying off

OpenAI’s “Scaling Laws for Neural Language Models” provided early empirical evidence that, if you increase model size, data, and compute together, then AI model error rates fall, i.e. the AI can continue to “learn more”. There is an approximate power-law over seven orders of magnitude. No saturation is visible yet — each extra order of magnitude still creates improvement, perhaps with diminishing returns.

The Bitter Lesson meets the Universal Approximation Theorem Reinforcement-learning pioneer Richard Sutton called progress through brute-force learning “The Bitter Lesson”: in the long run, general methods that exploit ever-cheaper compute outstrip clever, human-devised learning heuristics.

Mathematicians proved something similar in principle: the Universal Approximation Theorem (UAT) guarantees that a model with enough parameters can approximate any function arbitrarily well. It’s an existence proof (not a guarantee of being able to actually build an AI), but it marks the theoretical ceiling to

which our empirical scaling curves inch closer.

Put the two ideas together, and you get a compelling narrative:

Lower-bound reality Scaling laws: bigger models + more data $\Rightarrow$ steadily improving lower error.

Upper-bound theory UAT: In the infinite-size limit, a model can get the error all the way to zero - it can “Learn anything”

Why I don't think progress will stall

Hardware & capital keep flowing. Capacity might begin to plateau in the classical compute paradigm, but multi-billion-dollar data centre plans suggest the compute pipeline is far from drying up.

Algorithms continue to improve. Deepseek is evidence that we have a long way to go before we optimise the software stack — squeezing more capability from every FLOP. Alpha Tensor and Alpha Search are early examples of algorithms that are beginning to design new ones.

Data supply is evolving, not exhausted. We might have exhausted the Common Crawl (I don't think so), but we're already seeing a mix of synthetic data, simulation traces, multimodal streams and embodied AI to feed the models' insatiable appetite.

Use-value threshold passed. Tools like ChatGPT, Gemini and Claude are already making a difference (at least for the individual) today; that commercial feedback loop underwrites continued R&D and inoculates us against another AI winter.

None of this guarantees a specific timeline, but the increase in model capability has been surprisingly steep—hence the flurry of “AGI 2027” forecasts. Former OpenAI researcher Leopold Aschenbrenner argues in “Situational Awareness: The Decade Ahead” that another ~2-3 orders of magnitude in “effective compute” could yield systems able to automate most of AI research itself.

Whether you buy that exact date or not, the safe prediction is improving model capability—and sooner than most planners or commentators expected.

When machines out-think us

If models do cross the human-cognition line within a decade, the debate shifts from “can they?” to “now what”?

Skill displacement: Knowledge-work tasks (coding, legal research, design) may face the kind of upheaval factory floors saw in the 20th century.

Decision velocity: Board-level strategy, M&A due diligence or drug-discovery cycles could compress from months to minutes—rewarding organisations that learn to delegate judgement to machines while keeping humans “on the loop.” Power concentration: Compute-hungry frontier models favour firms (or states) able to finance multi-\$B clusters, raising antitrust and geopolitical stakes.

Alignment & liability: When an LLM synthesises medical advice more accurately than any doctor, who signs the prescription? Legal frameworks lag far behind technical reality.

Meaning & identity: If reasoning, expertise, creativity and insight become commoditised utilities, what remains uniquely valuable—and fulfilling—about human work?

These are no longer science-fiction hypotheticals; they’re strategic planning topics for the 2025–2035 horizon.

Closing note from the MBA cohort

Accelerationism sometimes implies “the faster the better, no matter the consequences.” That’s not my stance. I simply observe:

- The curve is still steep.
- There are no real signs of a plateau.
- Capital, talent and compute keep flowing in because the technology is already useful.

Believing progress will continue is not the same as cheering for reckless speed; governance and safety matter. But betting against further capability gains now feels like a losing one!

Scaling laws (confusingly not laws, but rather an observation) provide us with an improving lower-bound trajectory; the UAT shows that the ceiling is, in principle, unlimited.

Nothing I’ve seen in the data hints that we’re about to hit a wall—so the real work is to prepare for a world where machines routinely out-think us, and to shape that future responsibly.

The exciting, and sometimes daunting, task for the next cohort of MBA leaders is to harness that momentum responsibly, turning ever-better models into real-world value without losing sight of ethics and governance.

Feel free to push back—after all, good arguments are how we test convictions. But if you ask me today whether model performance will keep rising toward (and eventually past) human benchmarks, my answer is—yes. (Many of the em dashes have been added by the human.)

Suggested reading & listening

Theme - Title & author - Why it’s useful

Scaling evidence - “Scaling Laws for Neural Language Models” – Kaplan et al. (2020) The original power-law paper. <https://arxiv.org/abs/2001.08361>

Compute-data balance - “Training Compute-Optimal LLMs” – Hoffmann et al. (2022)- Explains the Chinchilla ratio. <https://arxiv.org/abs/2203.15556>

Philosophy of brute force - “The Bitter Lesson” – Richard Sutton (2019) - Why scale beats hand-engineered tricks. <http://www.incompleteideas.net/IncIdeas/BitterLesson.html>

Theoretical ceiling - Universal Approximation Theorem – Cybenko, Hornik & others - Shows any function is representable with sufficient parameters (model capacity)
https://en.wikipedia.org/wiki/Universal_approximation_theorem

Near-term AGI scenario - “Situational Awareness: The Decade Ahead” – Leopold Aschenbrenner (2024) - A detailed AGI-by-2027 thesis. <https://situational-awareness.ai/>

Critical commentary - “The AGI-in-2027 Thesis” – Evan Armstrong, Every (2024) A balanced critique of the 2027 timeline. <https://every.to/napkin-math/the-agi-in-2027-thesis>

Broader policy view - Nick Bostrom, Superintelligence (2014) - The classic on control & risk (still relevant - and I was very sceptical the first time I read this book!).