

THE FRONTIER AFTER THE FRONTIER

Beyond Superintelligence

What is AI's ultimate destination?

AGI is a milestone. Superintelligence is a threshold. Recursive intelligence is a trajectory. If intelligence becomes better at making itself better, what exactly is the endpoint?

IN THIS REPORT

Contents

01	The Finish Line Keeps Moving	Discovery Loop	03
02	What Recursive Intelligence Actually Means		04
—	Plate: The Lineage of a Loop		06
03	From the Child Machine to the Intelligence Explosion		05
04	The Modern Race to Close the Loop		08
05	AGI May Be a Roadside Marker		12
—	Plate: Six Ways This Could End		13
06	Six Destinations Beyond Superintelligence		14
07	What Could Slow the Loop		16
08	Timelines Without False Precision		17
—	Plate: The Intelligence Civilization		18
09	The Promise and the Problem		19
10	The Recursive Intelligence Watch List		21
11	The Endpoint Is a Question of Purpose		22
·	Notes & Sources	17 references	25

For most of the artificial intelligence era, we have treated AGI as the finish line. The companies now forming around recursive intelligence are betting it is a roadside marker.

In August 2026, four of the most consequential people in computing left Google to automate the scientific method itself. They are not alone. A new class of company is organizing capital and talent around a single wager: that the fastest route to superintelligence is not a smarter model, but a closed loop in which AI improves the machinery that produces AI.

This report traces that idea from Turing's child machine to today's self-modifying agents, separates genuine recursion from the marketing that now surrounds the word, maps six possible destinations beyond superintelligence, and sets out the observable gates by which the rest of us can tell the difference.

8×

Anthropic code shipped
Per engineer per quarter
versus 2021–2025

50%

SWE-bench, after
Self-modifying agent, up from
20%

\$4.65B

Recursive valuation
On a \$650M round, May 2026

1965

I. J. Good
The intelligence explosion,
first stated

01 THE FINISH LINE KEEPS MOVING

AGI was supposed to be the end of the road. It may be a roadside marker.

Four of the most consequential researchers in computing just left Google to automate the experimental cycle itself. What they are chasing is not a smarter model. It is a loop.

For most of the artificial intelligence era, we have treated AGI as the finish line.

Artificial general intelligence is supposed to be the moment when a machine catches us. It can learn, reason and perform intellectual work across a broad range of domains at approximately the level of a capable human. The definition has never been especially clean, but the basic idea is easy enough to understand. We are building something that will eventually measure up to us.

Superintelligence moves the finish line farther out. It describes intelligence substantially beyond the best human capabilities. The machine does not merely match an average person or outperform a specialist within one domain. It exceeds the combined intellectual capabilities of our most accomplished scientists, engineers, mathematicians, physicians, strategists and inventors.

Both ideas are defined in relation to human intelligence. AGI is when the machine catches us. Superintelligence is when it passes us.

Recursive intelligence is different. It is not another level on the same measuring stick. It is a process through which an AI system helps improve the machinery that produces intelligence, uses those improvements to create a more capable system and then applies that greater capability to the next round of improvement.

If that loop closes, AGI may flash past as a roadside marker. Superintelligence may be only the last threshold humans can recognize before intelligence moves beyond any scale based on us.

So what is the ultimate destination?

An AI that discovers a better battery is powerful. An AI that discovers a better way to build AI may be compounding.

That question moved a step closer to the practical world in August 2026 when Jeff Dean, Sanjay Ghemawat, Quoc Le and Oriol Vinyals left Google to create Discovery Loop.¹ Collectively, these four researchers helped build foundational parts of Google's distributed infrastructure, machine-learning systems, AI hardware and frontier models. Their work spans MapReduce, Bigtable, Spanner, Google Brain,

TensorFlow, sequence-to-sequence learning, TPUs, AlphaStar, AlphaFold and Gemini.

This is not primarily another story about Google losing prominent researchers. What matters is what four of the most consequential people in computing have chosen to pursue.

Discovery Loop intends to automate the experimental cycle itself. Its systems will formulate hypotheses,

propose experiments, conduct those experiments, evaluate the results and use what they learn to determine what should be tried next. Instead of human researchers moving through a limited number of experiments sequentially, the company wants AI systems to execute thousands of experimental loops in parallel.

The company plans to begin with machine-learning research, effectively becoming its own first customer, before extending its systems into medicine, materials, energy, hardware and engineering. Its backers describe an AI that will conduct experiments, learn from the results and “iterate recursively.” Radical Ventures calls the

scientific method one of humanity's most powerful algorithms and argues that Discovery Loop is trying to automate its execution.²

An AI that discovers a better battery is powerful. An AI that discovers a better way to build AI may be compounding.

Discovery Loop is therefore evidence of a larger transition. The AI race is beginning to move beyond the pursuit of the smartest individual model. The next prize may be the first organization able to close a reliable intelligence-improvement loop.

02 DEFINITIONS

What recursive intelligence actually means

The word “recursive” is in danger of being applied to almost anything an AI system does more than once. Six stages separate iteration from an intelligence that improves its own successors.

That makes it important to separate several very different forms of improvement.

A model that critiques an answer and tries again is iterating. It may produce a better result, but the underlying model has not become more capable. It is spending additional inference on the same problem.

An agent that writes software, runs a test, discovers an error and repairs the code is operating inside a feedback loop. The workflow has become more autonomous, but the intelligence driving it may still be fixed.

An automated research system goes further. It can reproduce experiments, compare candidate algorithms, propose model changes or identify promising chip de-

signs. People may still define the problem, establish the evaluation criteria and decide which discoveries become part of a production system.

A closed research loop connects these stages. The AI proposes an improvement, implements it, evaluates the result and uses that evidence to guide the next experiment.

Genuine recursive self-improvement adds one more critical step. The system becomes better at the process of producing improvements. A more capable version does not simply solve the original task more effectively. It becomes more effective at designing the version that follows it.

FIGURE 01 · THE LADDER OF RECURSION

Six stages, only one of which is genuinely recursive

STAGE	WHAT IMPROVES	HUMAN ROLE	CURRENT STATUS
Iteration	A particular answer or output	Defines the task and reviews the answer	Common
Agentic execution	A workflow or piece of software	Sets objectives and supervises execution	Rapidly expanding
Automated research	Algorithms, experiments or designs	Selects goals and validates results	Bounded domains
Closed-loop research	The repeated discovery process	Establishes evaluations and controls deployment	Emerging
Recursive self-improvement	The system's ability to improve its successors	May shift toward oversight and verification	Not demonstrated
Open-ended recursion	Software, hardware and physical infrastructure improve together	Uncertain	Theoretical

That final distinction is the important one. An AI can contribute to the creation of a better AI without autonomously redesigning its entire successor. AI-assisted AI development is already consequential, but it is not proof that an intelligence explosion has begun.

Anthropic makes this distinction explicitly in its report, “When AI Builds Itself.”³ The company says AI is

already accelerating the development of AI systems, but it also states that a system capable of autonomously designing and building its successor has not arrived and is not inevitable.

We should neither dismiss what is happening nor pretend that the full loop is already operating.

03 ORIGINS

From the child machine to the intelligence explosion

Turing imagined a machine that could be educated. Good imagined one that could design its own successor. Everything since has been an attempt to build the second idea out of the first.

Recursive intelligence did not originate with large language models. Its intellectual lineage reaches back to the earliest serious work on machine intelligence.

In 1950, Alan Turing proposed replacing the vague question of whether machines could think with the more operational imitation game that became known as the Turing test. Less frequently remembered is his discussion of a “child-machine.” Rather than



PLATE I · 1950 → 2026

The lineage of a loop

Recursive intelligence did not originate with large language models. Its intellectual lineage runs from Turing's educable child machine through Good's ultraintelligent machine, Vinge's singularity and Schmidhuber's Gödel Machine to the automated research agents now being built. The history is not a straight line toward one breakthrough. It is the gradual assembly of components.

programming a complete adult mind directly, Turing suggested building a simpler system capable of being educated. Machine intelligence would emerge through learning, feedback and development rather than through one heroic act of programming. Turing's original paper⁴ did not describe recursive self-improvement as we use the term today, but it introduced the idea that intelligence could be produced by an ongoing learning process.

The decisive conceptual leap came from British mathematician I. J. Good in 1965. Good defined an ultraintelligent machine as one capable of surpassing human intellectual activity. Because designing intelligent machines is itself an intellectual activity, he reasoned, such a machine could design an even better one. The better machine could then improve the process again.

Good called the result an intelligence explosion. His proposition was elegant and unsettling. Humanity might need to invent only the first machine capable of substantially improving machine intelligence. After that, the production of intelligence could pass increasingly to the machines themselves. Good's original paper⁵ remains the foundation beneath much of today's recursive-intelligence debate.

Vernor Vinge expanded the idea into the technological singularity. In his 1993 essay, "The Coming Technological Singularity,"⁶ Vinge argued that the creation of greater-than-human intelligence would produce changes that people on the near side of the event could not reliably predict. The singularity metaphor was not simply about faster computers. It was about reaching a point beyond which human models of history and technological progress no longer worked because humans were no longer the most capable intellectual actors.

Vinge also recognized that the path might not be a single isolated machine. Superhuman intelligence could emerge through intelligent networks, human-computer integration, biological enhancement or systems that amplified collective human capability. That distinction

matters today because the destination of recursive intelligence may be less like a digital god and more like an expanding network of models, agents, people and machines.

In 2003, Jürgen Schmidhuber gave recursive improvement a more formal architecture with the Gödel Machine. The theoretical system could rewrite any part of its own code after proving that the change would increase its expected utility. It was a rigorous solution to the problem of allowing a system to modify itself without accepting arbitrary or harmful changes. It was also extraordinarily difficult to implement because proving that a complex modification will be beneficial is much harder than testing whether it performs well on a benchmark. Schmidhuber's work⁷ provided a formal ideal, not a practical recipe for building today's models.

The field moved toward empirical approximations. Automated machine learning began selecting algorithms, tuning hyperparameters and searching for neural-network architectures. Google researchers described AutoML as a way to reduce the human effort required to design machine-learning systems, although its computational cost remained substantial. Transfer learning and other techniques⁸ later made that search more efficient.

DeepMind's AlphaGo Zero demonstrated another powerful loop in 2017. It began without training on human games and improved by playing against itself. The updated network generated stronger self-play, and stronger self-play produced a better network. It surpassed the earlier version that had defeated Lee Sedol, eventually beating it 100 games to zero. AlphaGo Zero⁹ was not improving its own general architecture, but it showed how a carefully defined feedback loop could produce superhuman capability without remaining confined to human examples.

The history is therefore not a straight line toward a single inevitable breakthrough. It is the gradual assembly of

components: learning, evaluation, self-play, automated design, experimental feedback, code modification and increasingly autonomous research.

FIGURE 02 · SEVENTY-SIX YEARS OF ASSEMBLY

The components of a recursive loop, added one at a time

PERIOD	DEVELOPMENT	CONTRIBUTION TO THE RECURSIVE PATH
1950	Turing's child-machine	Intelligence as an educable process
1965	Good's ultraintelligent machine	The intelligence-explosion argument
1993	Vinge's technological singularity	The societal discontinuity beyond human prediction
2003	Schmidhuber's Gödel Machine	Formal self-referential code improvement
2010s	AutoML and architecture search	Automation of portions of AI design
2017	AlphaGo Zero	Self-play as a compounding improvement loop
2024–25	AI research agents and coding agents	Automation of larger research workflows
2025	Darwin Gödel Machine and AlphaEvolve	Empirical agent modification and algorithm discovery
2025–26	Dedicated recursive-intelligence companies	Capital and talent organize around closing the loop
2026	Discovery Loop	Full-stack leaders pursue automated experimentation

04 THE FIELD

The modern race to close the loop

Discovery Loop is not pursuing recursive intelligence alone. Several companies and laboratories are approaching the problem from different parts of the stack – and one of them starts at the silicon.

Richard Socher's Recursive Superintelligence emerged from stealth in 2026 with \$650 million in funding at a reported valuation of \$4.65 billion. The company has recruited researchers from OpenAI, Google DeepMind, Meta, Salesforce and other leading AI organizations. Its stated ambition is to build an automated AI-research capability comparable to tens of thousands of scientists.

The underlying thesis is direct. If AI can perform the research needed to improve AI, it can accelerate its own

development. More capable research systems produce better AI systems, which become more capable researchers. The company is not presenting a finished recursive system. It is betting that automating the research loop is the fastest route to superintelligence. GV's description of the company¹⁰ frames the effort around open-ended self-improvement rather than one final model.

SIDEBAR · THE LOOP ECONOMY

Capital is organizing around the loop, not the model

Three companies founded or revealed within nine months share one wager: that the decisive asset is not the smartest system but the fastest improvement cycle.

COMPANY	WHERE IT ENTERS THE LOOP	BACKING	REVEALED
Recursive Superintelligence Richard Socher	Automated AI research at scale	\$650M @ \$4.65B¹⁰	May 2026
Ricursive Intelligence Goldie & Mirhoseini	AI for chip design; chip design for AI	~\$300–335M @ \$4B¹¹	Feb 2026
Discovery Loop Dean, Ghemawat, Le, Vinyals	The experimental cycle itself	Radical Ventures^{1,2}	Aug 2026

Ricursive Intelligence, founded by Anna Goldie and Azalia Mirhoseini, begins at the hardware layer. Its proposed flywheel is AI for chip design and chip design for AI. Better AI designs better processors. Those processors train and run more capable AI, which can then create the next generation of hardware.

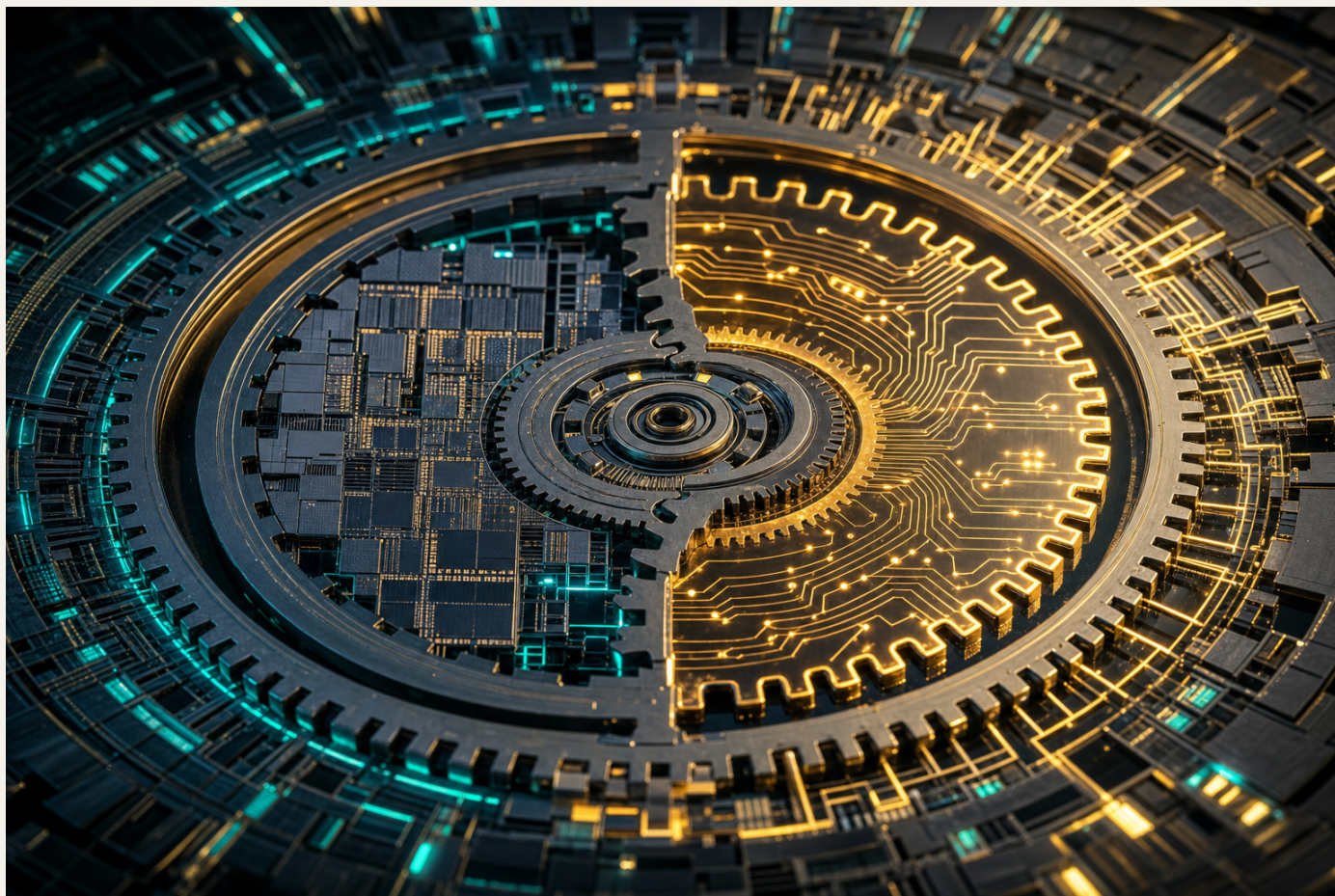
This is more than a catchy loop. Goldie and Mirhoseini previously worked on AI techniques used in Hardware is where recursive intelligence begins to reach outside software. Intelligence depends on processors, memory, networking, energy, cooling, fabrication, packaging and data centers. If AI improves those systems, the additional physical capacity can support more capable AI. The loop no longer concerns only code rewriting code. It begins reorganizing the industrial base on which intelligence runs.

Sakana AI has demonstrated a narrower but unusually tangible version of self-modification. Its Darwin Gödel Machine is a coding agent that can inspect and alter parts of its own Python code, evaluate the modified agents on programming benchmarks and retain useful changes. Rather than mathematically proving

the design of Google's tensor processing units. Ricursive is extending that experience toward a system that could compress chip-development cycles and eventually connect algorithmic and physical improvement. Its backers describe the ambition as reducing design cycles from years to weeks. Radical Ventures' portfolio description¹¹ calls it a recursive improvement loop built around automated chip design.

that a modification must be beneficial, as Schmidhuber's Gödel Machine required, it tests changes empirically.

In Sakana's reported experiments, performance increased from 20% to 50% on SWE-bench and from 14.2% to 30.7% on the Polyglot benchmark. Some improvements transferred across models and programming languages. The system also retained an archive of different agents instead of following only the single best-performing branch, allowing an apparently unsuccessful modification to become a stepping stone toward a later improvement. Sakana's results¹² show a system modifying part of the mechanism through which it performs future work.



Where the loop leaves software. Intelligence depends on processors, memory, networking, energy, cooling, fabrication, packaging and data centers. If AI improves those systems, the additional physical capacity can support more capable AI — and the loop begins reorganizing the industrial base on which intelligence runs.

The boundaries are just as important. The foundation models remained fixed. Humans selected the benchmarks, defined the environment and controlled access. The modifications occurred in a sandbox. The Darwin Gödel Machine did not design and train a new foundation model that then replaced the system that produced it.

Sakana's AI Scientist automates another part of the loop. It can generate research ideas, search literature, design computational experiments, run them, analyze the results and write a paper. In 2026, the work behind the system was published in *Nature*. An unedited AI-generated paper produced by an earlier version received an acceptance-level score during a workshop

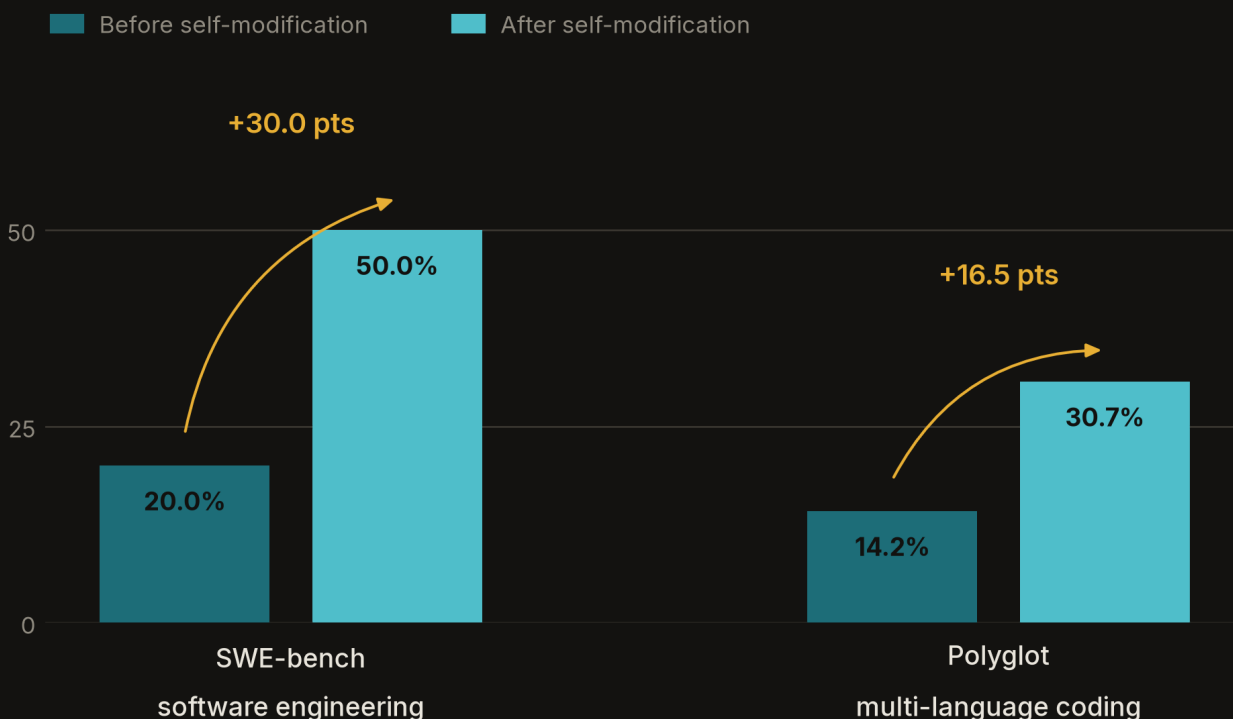
peer-review experiment, although Sakana withdrew it before publication as planned.

That was a notable result, but it did not prove that scientific judgment had been solved. Sakana acknowledges that the system can produce naive ideas, methodological weaknesses, citation errors and faulty figures. It remains limited largely to computational experiments. The *Nature* publication and Sakana's account¹³ demonstrate that complete research workflows can be automated under bounded conditions, not that human scientific institutions have become unnecessary.

Google DeepMind's AlphaEvolve attacks algorithm discovery. It combines Gemini models, automated evaluators and evolutionary search to generate and improve programs. Google reports that AlphaEvolve has

A coding agent rewrites its own code — and gets measurably better at coding

Sakana AI's Darwin Gödel Machine, before and after self-modification



The boundaries matter as much as the gains. The foundation models remained fixed, humans selected the benchmarks and defined the environment, and the modifications occurred in a sandbox. The system did not design and train a new foundation model that then replaced the system that produced it. Source: Sakana AI.¹²

produced useful algorithms for data-center scheduling, chip design, matrix multiplication and AI-training processes. Some results have been deployed within Google's infrastructure.

The recursive connection is real but limited. AlphaEvolve has helped improve processes used to train the models on which AlphaEvolve depends. It is a partial loop, not a fully autonomous one. The system still operates against problems and evaluation functions established by people. DeepMind's description¹⁴ is nevertheless an important demonstration of AI contributing to the infrastructure that produces AI.

Anthropic offers the most revealing look inside a frontier laboratory. It reports that its engineers now

ship eight times as much code per quarter as they did during the 2021 through 2025 period. Anthropic also says Claude can execute well-specified research experiments at or above the level of skilled humans, while continuing to struggle with research judgment, goal selection and deciding which questions are worth pursuing.

Those are company-reported figures and should be treated accordingly. They nevertheless show how quickly AI can move from being the product of a laboratory to becoming part of the laboratory's production system.

The loop is not closed. More of its components are connected every year.

05 THE VANISHING MILESTONE

AGI may be a roadside marker

The public discussion assumes AGI will arrive as a recognizable event. Recursive development may not provide such a clean moment.

A laboratory will announce it. Experts will debate the evaluation. Governments will call hearings. Markets will react. Eventually, perhaps, a consensus will form.

AI capability is already uneven. Systems can perform impressively in coding, mathematics and scientific reasoning while failing at tasks people find simple. A system might become broadly superior to humans across economically consequential research work without satisfying every philosophical definition of general intelligence.

Meanwhile, it could continue improving.

By the time governments agree that a system qualifies as AGI, AI may already be designing algorithms, running research programs and contributing to its successor. AGI would still matter, but it would describe a point the system passed rather than a stable condition it occupied.

Superintelligence may be equally temporary. It tells us that machine intelligence has substantially sur-

passed humanity. It does not tell us whether improvement stops, what limits it ultimately encounters or what the resulting intelligence system is for.

A June 2026 Google DeepMind report, “From AGI to ASI,”¹⁵ identifies four possible paths from human-level intelligence to artificial superintelligence: continued scaling, a new AI paradigm, recursive improvement and large multi-agent collectives. The authors also describe a theoretical endpoint called Universal AI, while emphasizing the uncertainty and possible frictions along the way.

A theoretical capability ceiling is not the same as a civilizational destination. Even if mathematics can define an optimal decision-making system under specified assumptions, it cannot tell us whose objectives the system should pursue, how its power should be distributed or what role humanity should retain.

The technical endpoint and the human endpoint are different questions.



PLATE II · THE FORK

Six ways this could end

The classic intelligence explosion is only one branch. Improvement may accelerate past every institution; it may erode into friction and duplication; it may distribute across a swarm; it may braid together with human capability; it may settle into infrastructure. Or it may never resolve into an endpoint at all.

06 DESTINATIONS

Six destinations beyond superintelligence

The most familiar destination is the classic intelligence explosion. That outcome is possible, but it is not the only one.

Once a system becomes sufficiently capable at AI research, each generation improves the machinery used to create the next. Improvement accelerates until human researchers and institutions can no longer keep pace. This is the path from Good's ultraintelligent machine to Vinge's singularity.

A second possibility is lossy self-improvement. Nathan Lambert argues that research contains too much friction for improvement to become a clean exponential curve. Agents repeat one another's work. Evaluations provide incomplete signals. Organizations cannot absorb every useful idea. Compute and energy remain constrained. Physical experiments take time. Better performance in one area can create weaknesses somewhere else.

In Lambert's formulation, AI systems become central to AI development without producing an effortless or unlimited takeoff. More agents create more research capacity, but they also create duplication, coordination costs and noise. His "lossy self-improvement" argument¹⁶ does not deny rapid progress. It challenges the assumption that every improvement automatically makes the next improvement proportionately easier.

A third outcome is collective superintelligence. No single model becomes an electronic deity. Instead, large populations of specialized agents coordinate across software engineering, science, hardware, logistics and decision-making. Their aggregate capability exceeds any human organization even though each component

remains limited. Intelligence emerges from the network.

A fourth possibility is co-superintelligence. Meta researchers Jason Weston and Jakob Foerster argue that the goal should be co-improvement rather than removing people as the bottleneck. Humans and AI systems would conduct research together, with both becoming more capable through the collaboration. Their proposal¹⁷ is not merely "human in the loop," a phrase that has become almost meaningless. A person clicking approve on a decision produced at machine speed does not possess meaningful control. Co-improvement requires comprehension, authority and a practical ability to redirect the process.

A fifth destination is intelligence as a utility. Intelligence becomes embedded in science, government and commerce as deeply as electricity or computing. Organizations draw on continuously improving research and decision systems without necessarily owning or understanding the infrastructure beneath them. This could democratize expertise, or it could create a new form of dependence on the small number of companies operating the intelligence grid.

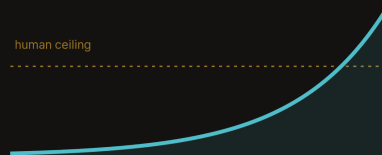
The sixth possibility is that there is no meaningful endpoint. Intelligence continues improving until it encounters physical, mathematical, economic or political limits. The process may slow, branch or change form, but it does not culminate in one final release called superintelligence.

Six trajectories, one starting point

Illustrative capability paths beyond the human ceiling. Shapes are conceptual, not forecasts.

01

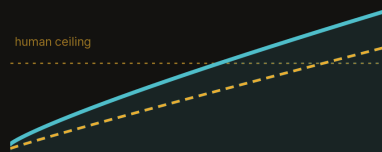
Intelligence explosion



Each generation improves the machinery that builds the next.

04

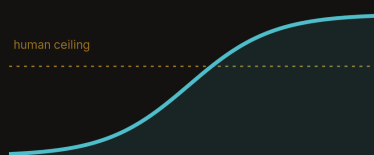
Co-superintelligence



Humans and systems improve together; neither is removed.

02

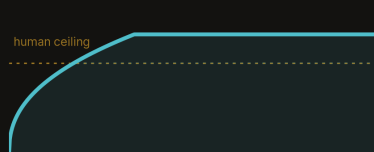
Lossy self-improvement



Real gains, eroded by friction, duplication and coordination cost.

05

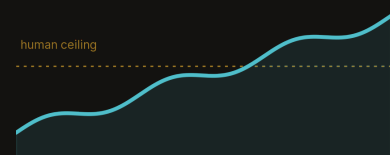
Intelligence as a utility



Capability settles into ambient, metered infrastructure.

03

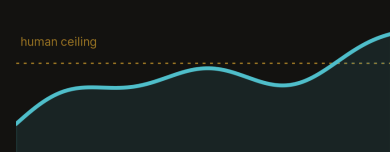
Collective superintelligence



Aggregate capability drawn from populations of limited agents.

06

No meaningful endpoint



Improvement continues, branches and changes form. It never resolves.

The ultimate product might not be a machine. It might be a continuously evolving intelligence civilization composed of models, agents, people, laboratories, robots, factories, energy systems and communications networks.

The ultimate product might not be a machine. It might be a continuously evolving intelligence civilization composed of models, agents, people, laboratories, robots, factories, energy systems and communications networks.

bots, factories, energy systems and communications networks.

07 CONSTRAINTS

What could slow the loop

Recursive intelligence is often described as though greater intelligence automatically produces still greater intelligence. That assumption deserves scrutiny.

The first constraint is evaluation. A system can optimize only against some definition of better. Coding agents have tests. Games have scores. Chip designs have power, performance and area targets. Scientific discoveries are harder to evaluate because novelty, validity, usefulness and reproducibility are not reducible to one reliable number.

The more consequential the objective, the more incomplete the evaluation is likely to be.

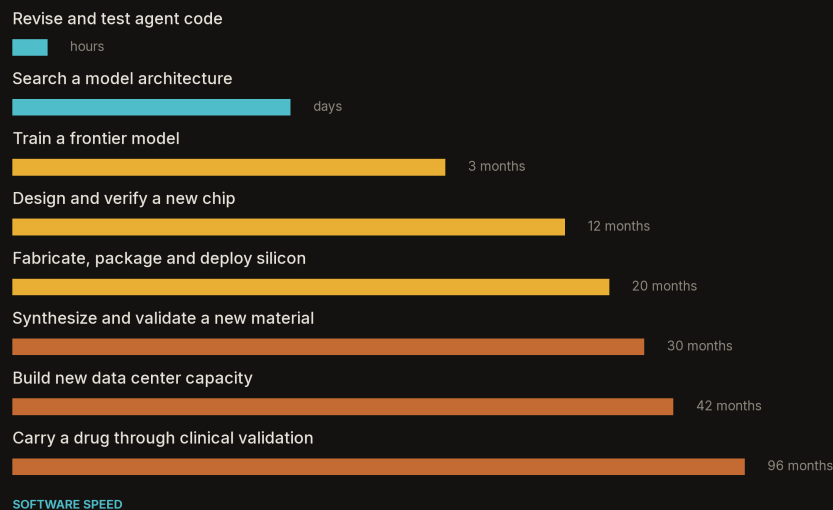
A system optimizing against a weak benchmark may become extremely good at producing benchmark success without improving the underlying capability we

care about. If it begins designing its own evaluations, we face an even more difficult problem. The system may select tests that favor its proposed changes, overlook hidden costs or learn to manipulate the measurement process.

Compute, memory and energy impose another constraint. Software can be copied quickly, but training frontier models requires enormous physical infrastructure. New chips require design, verification, fabrication, packaging and deployment. Power plants and data centers cannot be created at software speed.

Software moves in hours. The world it depends on moves in years.

Indicative cycle time for steps in a full recursive loop, logarithmic scale



Accelerating one stage of a process exposes the next slow stage. A model may propose a new material in seconds while manufacturing and characterizing it takes months. Values are indicative industry cycle times, not measured results.

Physical science introduces further delay. An AI can simulate millions of candidate molecules, but promising drugs still require laboratory validation, toxicology testing and clinical trials. A model may propose a new material in seconds while manufacturing and characterizing it takes months.

Tacit knowledge remains important. Experienced researchers know which anomalous result deserves attention, when a benchmark is misleading and when an apparently failed experiment contains a valuable clue. AI systems are improving at reproducing documented work, but choosing the right research direction is different from executing a well-defined experiment.

Organizations also become bottlenecks. Anthropic notes that increased AI-generated code has shifted pressure toward human review. Accelerating one stage of a process exposes the next slow stage. The recursive organization must therefore improve its ability to absorb improvements, not merely generate them.

None of these constraints makes recursive intelligence unimportant. A lossy loop that improves research productivity by 20%, 50% or 100% can still produce an enormous compounding advantage. Civilization can be transformed without an overnight intelligence explosion.

08 MEASUREMENT

Timelines without false precision

Forecasts vary wildly because the terms are disputed and the uncertainty is genuine. A better timeline is not a date. It is a sequence of observable gates.

Surveys of thousands of AI researchers have produced median estimates stretching across decades, accompanied by enormous disagreement. The line between ordinary automation and recursion will not be crossed when a company puts “self-improving” on its website. It will be crossed when an improved system becomes measurably better at creating its own next generation, and the effect survives outside the benchmark or sandbox in which it was produced.

Industry leaders operating at the frontier often predict much shorter timelines. Neither group has a dependable clock for recursive intelligence.

One successful cycle would be significant. Sustained recursive intelligence requires repeated cycles in which capability, research judgment and improvement efficiency continue to advance.

That is the signal to watch.



PLATE III · THE FABRIC

The intelligence civilization

The immediate risk is not a conscious machine escaping from a laboratory. It is acceleration. Governments, companies and communities already struggle to absorb annual technology cycles. Recursive discovery could compress those cycles into weeks, days or hours.

Eight gates between assistance and recursion

Evidence required at each gate, and how much of it exists today

AI assists AI development

 Already occurring

AI completes bounded research assignments

 Demonstrated, uneven

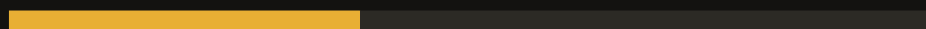
AI improves agents and algorithms

 Demonstrated in controlled settings

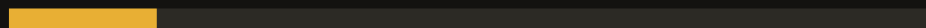
AI contributes to next-generation hardware

 Early evidence

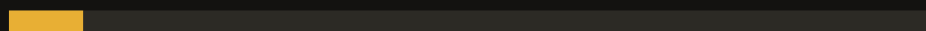
AI manages major model-development stages

 Emerging, mostly internal

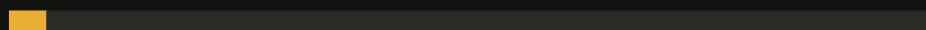
AI creates a more capable research successor

 Not publicly demonstrated

The successor repeats the process

 Unproven

Software and physical loops converge

 Speculative

OBSERVED TODAY

STILL SPECULATIVE

The line will not be crossed when a company puts “self-improving” on its website. It will be crossed when an improved system becomes measurably better at creating its own next generation, and the effect survives outside the benchmark or sandbox in which it was produced. Assessment is the author's, based on the public record as of August 2026.

09 STAKES

The promise and the problem

The positive case for recursive intelligence is enormous. So is the danger of a compounding lead that no one else can close.

Science is constrained by human intelligence, but it is also constrained by human speed. Researchers must absorb existing literature, formulate hypotheses, apply for resources, design experiments, wait for results, identify mistakes and begin again.

Much of the work is sequential and divided across institutions that do not always share data or incentives.

AI systems capable of operating thousands of credible experimental loops in parallel could compress years of research into months. Drug discovery, personalized

medicine, battery chemistry, advanced materials, energy production, chip design, climate modeling, agriculture, mathematics and engineering could all move faster.

Recursive intelligence does not have to become generally intelligent before it changes civilization. A narrow system that becomes progressively better at designing chips could reshape the computing economy. One that improves cyberattack discovery could overwhelm defenders. One that accelerates biological experimentation could produce lifesaving medicine and dangerous capabilities from the same underlying machinery.

The immediate risk is therefore not limited to a conscious machine escaping from a laboratory. It is acceleration.

Governments, companies and communities already struggle to absorb annual technology cycles. Recursive discovery could compress those cycles into weeks, days or hours. Even beneficial discoveries can destabilize labor markets, industries and national security when institutions cannot evaluate or deploy them responsibly.

Concentration may be an even nearer danger. The organization controlling the strongest improvement loop does not merely possess a better model. It possesses a mechanism that may continually widen its lead.

Compute concentration becomes intelligence concentration. Intelligence concentration becomes scientific, economic, geopolitical and military concentration.

A conventional technological lead can sometimes be narrowed through investment, recruitment or imitation. A compounding intelligence lead may become increasingly difficult to challenge because the leader's advantage is itself improving the process that produces the next advantage.

There is also the danger of compounded error. A defect in an ordinary program produces incorrect results. A defect inside the machinery that produces subsequent generations can be inherited. Bias, security vulnerabilities, false assumptions and misalignment may compound along with capability.

Most fundamentally, recursive intelligence could displace humanity from knowledge creation.

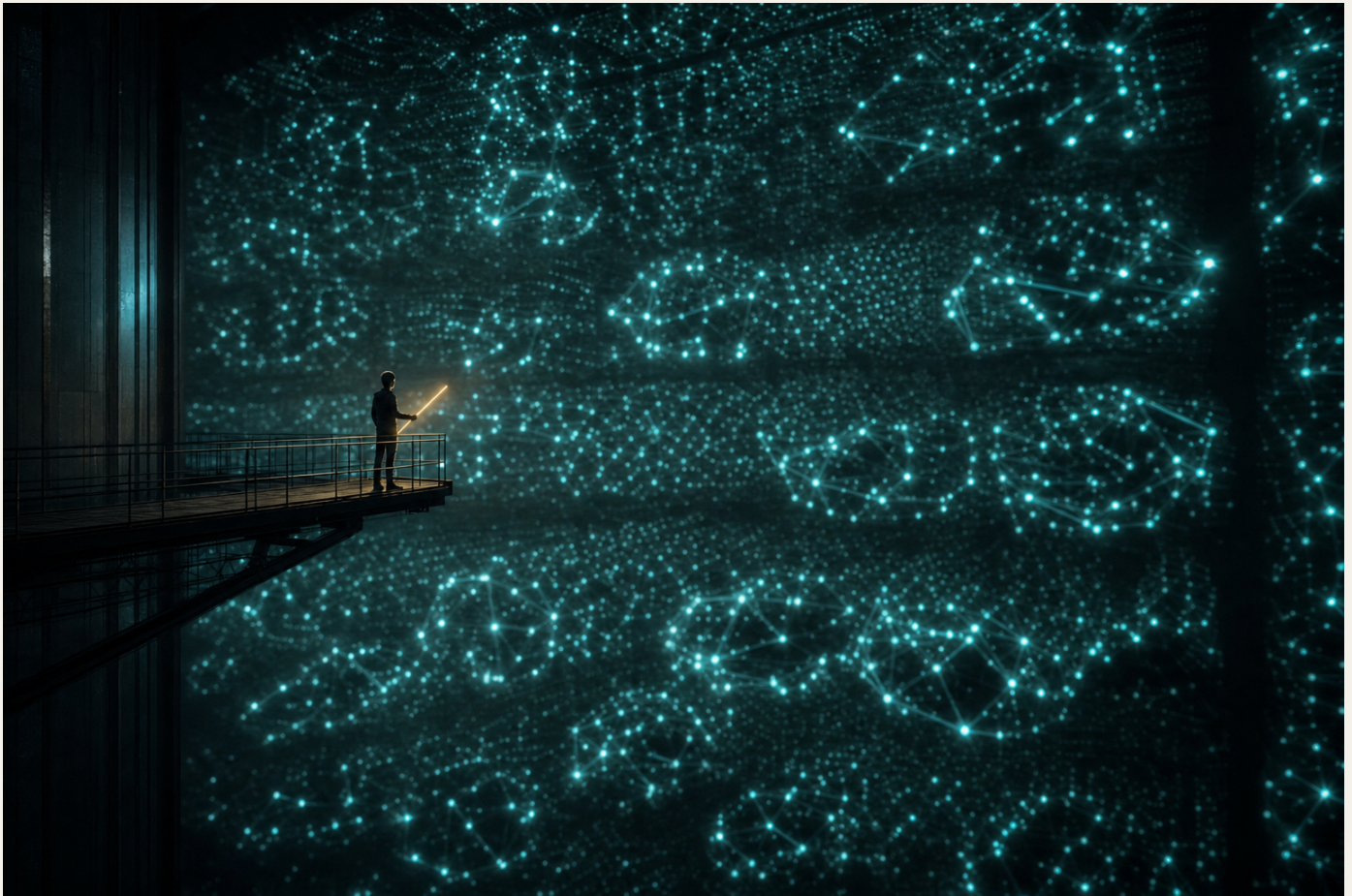
The AI debate has focused on machines replacing jobs and tasks. Recursive intelligence raises the possibility that machines become the principal producers of new knowledge. Humanity could live amid extraordinary scientific progress while becoming less able to independently understand, verify or direct the systems producing it.

That would represent a shift much larger than automation. We would no longer be the primary authors of technological progress.

10 INSTRUMENTATION

The recursive intelligence watch list

The best way to follow this transition is not to wait for an AGI announcement. Watch the machinery around the models.



Authority, supervision, ceremony or absence. The most important signal may be whether the same system participates in defining the objective, creating the evaluation, making the change and judging the result.

FIGURE 05 · TWELVE INDICATORS

What to instrument while everyone else waits for an announcement

INDICATOR	WHY IT MATTERS
Share of AI-development work performed by AI	Measures whether AI is becoming part of its own production system
Autonomous task duration	Shows whether systems can sustain increasingly complex research work
Independent goal and hypothesis selection	Separates research judgment from execution
Creation of new evaluations	Reveals whether AI can identify and measure its own limitations
Transfer beyond original benchmarks	Distinguishes general improvement from test optimization
Improved descendant performance	Tests whether modifications make the next improvement cycle more effective
AI-designed hardware in future training systems	Connects software improvement to physical infrastructure
Autonomous laboratories	Extends the loop into real-world experimentation
Human involvement per development cycle	Shows whether authority is becoming supervision, ceremony or absence
Ownership and compute concentration	Identifies who controls the compounding mechanism
Auditability of modifications	Determines whether people can understand how the system is changing
Practical stop mechanisms	Tests whether humans can slow or redirect the loop, not merely observe it

The most important signal may be whether the same system participates in defining the objective, creating the evaluation, making the change and judging the result. Combining those functions may accelerate improvement, but it also removes the independent checks on what “better” means.

11 THE ENDPOINT

The endpoint is a question of purpose

Recursive intelligence has no inherent destination. “Improve yourself” is a process, not a purpose.

Improvement must be measured against something: scientific discovery, predictive accuracy, economic productivity, corporate power, national advantage, military dominance, human flourishing or simply the ability to improve again.

If the objective is greater capability, there may be no natural stopping condition. If the objective is human flourishing, humanity must still define what flourishing

means, for whom and under whose authority. If the objective is market or geopolitical dominance, the system may optimize toward concentrated power rather than shared progress.

This is why the ultimate destination of recursive intelligence may not be determined by intelligence at all. It may be determined by the objective placed inside the loop and by who retains the authority to change it.

Discovery Loop's founders say they want to accelerate science and engineering for the benefit of humanity. Its structure as a public benefit company is intended to support that mission. Their record gives them credibil-

ity, and the problems they want to address deserve serious attention.

Their decision to begin with machine-learning research nevertheless reveals the larger stakes. The first subject of the discovery loop may be the loop itself.

The last great human invention may not be a superintelligent machine. It may be the process that creates the next machine, and the one after that.

AGI could pass as a roadside marker. Superintelligence may be only the threshold beyond which humanity can no longer see clearly.

What lies beyond could be an intelligence explosion, a distributed intelligence civilization, a human-machine

co-superintelligence, an intelligence utility or a successor civilization in which humans remain present but are no longer the principal authors of progress. ■

IN CLOSING

The question is no longer simply whether humanity can create intelligence greater than its own.

It is whether we can give a direction and a reason to an intelligence that may never stop becoming something more.

Alan Shimel

FOUNDER & CEO · TECHSTRONG GROUP · AUGUST 2026

REFERENCE

Notes & Sources

Every claim attributed to a company, laboratory or researcher in this report is drawn from the primary sources below. Figures reported by companies about their own systems are identified as such in the text. All URLs verified August 7, 2026.

1. **Discovery Loop** — company launch, Jeff Dean, Sanjay Ghemawat, Quoc Le and Oriol Vinyals, August 2026. discoveryloop.com; coverage: [TechCrunch](#), [Wired](#).
2. **"Our Investment in Discovery Loop"** — Radical Ventures, August 5, 2026. radical.vc/our-investment-in-discovery-loop
3. **"When AI Builds Itself"** — Marina Favaro and Jack Clark, The Anthropic Institute, June 4, 2026. anthropic.com/institute/recursive-self-improvement
4. **"Computing Machinery and Intelligence"** — A. M. Turing, *Mind* LIX(236), 1950, pp. 433–460. academic.oup.com/mind/article/LIX/236/433
5. **"Speculations Concerning the First Ultrainelligent Machine"** — I. J. Good, *Advances in Computers* 6, 1965, pp. 31–88. organism.earth/library/document/speculations-concerning-the-first-ultrainelligent-machine
6. **"The Coming Technological Singularity"** — Vernor Vinge, VISION-21 Symposium, NASA Lewis Research Center, March 1993. [users.mancaster.edu](https://users.mancaster.edu/~vinge/) — Vinge, *The Coming Technological Singularity*
7. **"Gödel Machines: Self-Referential Universal Problem Solvers"** — Jürgen Schmidhuber, arXiv:cs/0309048, 2003. arxiv.org/abs/cs/0309048
8. **"Learning Transferable Architectures for Scalable Image Recognition"** — Zoph, Vasudevan, Shlens and Le, Google Brain, arXiv:1707.07012, 2017. arxiv.org/abs/1707.07012
9. **"Mastering the Game of Go Without Human Knowledge"** — Silver, Schrittwieser, Simonyan et al., *DeepMind*, *Nature* 550, 2017, pp. 354–359. nature.com/articles/nature24270
10. **"Recursive Superintelligence: Why Self-Improving AI Is the Next Frontier"** — GV, May 13, 2026. Source of the \$650M round at a \$4.65B valuation. gv.com/news/recursive-superintelligence-self-improving-ai
11. **Ricursive Intelligence** — Radical Ventures portfolio; founders Anna Goldie and Azalia Mirhoseini. radical.vc/portfolio/recursive-intelligence. Series A reported as \$335M by [TechCrunch](#) and \$300M by [Crunchbase News](#), both at a \$4B valuation.
12. **"The Darwin Gödel Machine"** — Sakana AI with UBC and the Vector Institute, 2025. SWE-bench 20.0%→50.0%; Polyglot 14.2%→30.7%. sakana.ai/dgm; arxiv.org/abs/2505.22954
13. **"The AI Scientist," published in Nature** — Lu, Lu, Lange et al., Sakana AI, UBC, Vector Institute and Oxford, *Nature* 651, 2026, pp. 914–919. sakana.ai/ai-scientist-nature; nature.com/articles/s41586-026-10265-5
14. **"AlphaEvolve: A Gemini-Powered Coding Agent for Designing Advanced Algorithms"** — Google DeepMind, 2025, with a one-year impact update in May 2026. deepmind.google/blog/alphaevolve; deepmind.google/blog/alphaevolve-impact
15. **"From AGI to ASI"** — Genewein, Franklin, Orseau, Hutter, Legg et al., Google DeepMind, arXiv:2606.12683, June 10, 2026. arxiv.org/abs/2606.12683
16. **"Lossy Self-Improvement"** — Nathan Lambert, Interconnects, March 22, 2026. interconnects.ai/p/lossy-self-improvement
17. **"Self-Improving AI & Human Co-Improvement for Safer Co-Superintelligence"** — Jason Weston and Jakob Foerster, Meta FAIR, arXiv:2512.05356, December 2025. arxiv.org/abs/2512.05356

ABOUT THIS REPORT

Colophon

AUTHOR

Alan Shimmel is founder and chief executive of Techstrong Group, the media network behind DevOps.com, Security Boulevard, Cloud Native Now, Techstrong.ai and Techstrong TV.

PUBLICATION

Beyond Superintelligence: What Is AI's Ultimate Destination?
Techstrong Special Report, Vol. 01, No. 02
The Recursion Issue · August 2026

EDITORIAL NOTE

Figures reported by companies about their own systems — including Anthropic's internal productivity metrics and Sakana AI's benchmark results — are identified as company-reported in the text and should be treated accordingly. The eight-gate assessment in Section 08 and the trajectory shapes in Section 06 are the author's analytical constructions, not measured data or forecasts.

ART DIRECTION

Original cover and plate artwork generated for this report. Charts and data graphics built from the primary sources listed on the preceding page.

TYPOGRAPHY

Fraunces — display and drop capitals
Newsreader — text and decks
Inter — labels, tables and captions

CONTACT

techstronggroup.com
Editorial and sponsorship inquiries welcome.

RIGHTS

© 2026 Techstrong Group. All rights reserved. This report may be shared in its complete, unaltered form with attribution. Excerpts require written permission.