

PUBLISHED BY
RDLB AGENTIC

ISSUE N° 04
AUGUST 2026



A MONTHLY FIELD REPORT FROM THE AGENTIC ERA

Brave New *Word*

The orchestration issue. One agent is a parlor trick. The interesting machines are crowds — a lab of agents proposing, a trusted evaluator deciding, the strongest designs surviving. We map how swarms actually work, and the coordination tax nobody prints.

ORCHESTRATOR & WORKERS PROPOSE, VERIFY, EVOLVE REQUIRED READING ×4 ANTHROPIC'S 90.2%

The month one agent became a *crowd*.

01	From the editor's desk	03
	Issue four. One model is a soloist. The work is an orchestra.	
02	Orchestration	04
	The orchestrator-worker pattern, and why a crowd with separate minds beats a single brilliant one.	
03	Five roles, one search	05
	Explorer, exploiter, repairer, crossbreeder, critic — a swarm is a division of labor, not a chorus of clones.	
04	Propose, verify, evolve	06
	The loop that lets a model be wrong safely — because a trusted evaluator, not the model, is the source of truth.	
05	Required Reading	07
	Four papers on swarms, evaluators, and illuminated search — each linked and verified live.	
06	Case study — Anthropic's 90.2%	08
	A crowd beat a soloist by 90.2% — and spent 15× the tokens doing it.	
07	The Ticker	09
	Four real signals from the swarm, sourced and dated.	

MASTHEAD

Brave New Word is published monthly by RDLB Agentic — the operating arm of RDLB Agency. We build brand systems that think, create, and operate. We do not sell AI tools.

Notes from the floor

How RDLB Agentic runs a swarm

EDITORIAL STANDARD

Every figure is concrete or absent. Every claim is sourced. One enemy per piece. Mechanism over metaphor. If we can't link it, we don't print it.

Issue four. One agent became a *crowd*.

For three issues we watched a single agent learn to see, to ship, and to talk. This month it stops working alone. The most capable systems of the year are not one model trying harder. They are many models, divided into roles, checked by something they cannot talk their way past.

A swarm is not a group chat. It is a structure: one agent plans and delegates, others work in parallel with their own context, and a trusted evaluator — not a vote, not a vibe — decides what is actually good. That last part is the whole game. It is what lets a model propose a thousand wrong ideas cheaply, because something honest is standing at the exit.

The payoff is real and measured. Anthropic's research swarm beat a single strong agent by 90.2% on its own evaluation.⁴ DeepMind's evolutionary systems used the same shape to make discoveries a lone model could not.^{1,2} The pattern works.

The enemy this month is *more agents, more intelligence* — the belief that stacking models automatically stacks judgment. It does not. A crowd with no evaluator is just a more expensive way to be confidently wrong, and the same Anthropic system that won by 90.2% also burned roughly fifteen times the tokens of a chat.⁴

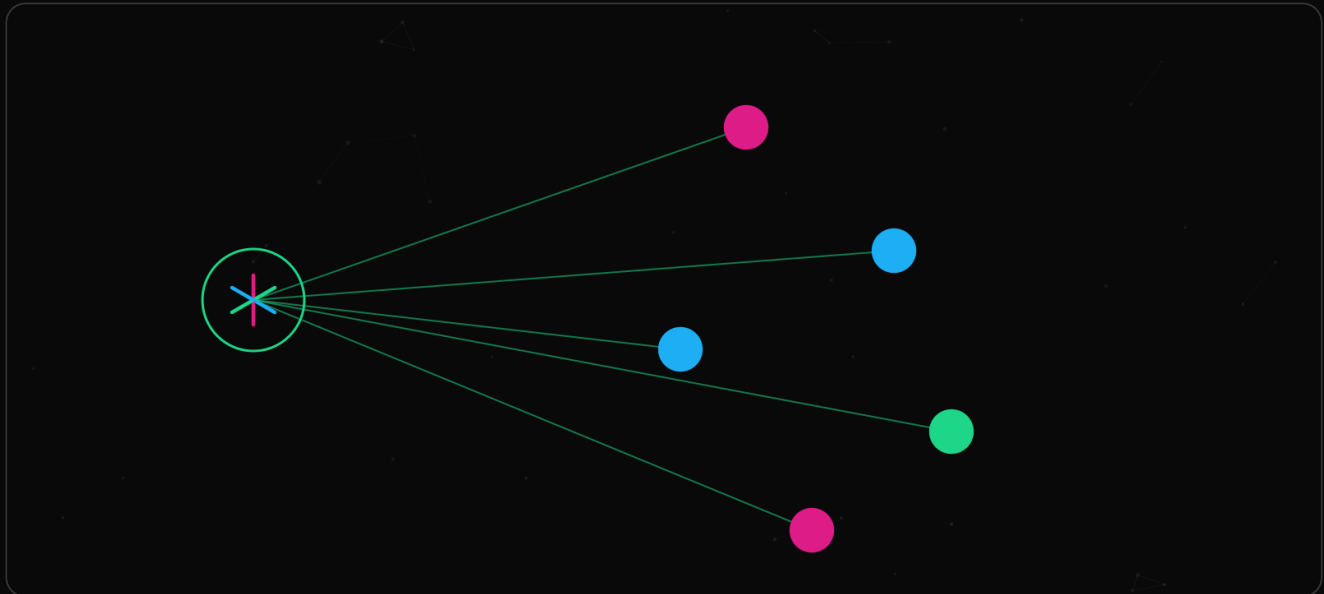
So this issue is about coordination as engineering: who proposes, who decides, who keeps the cost honest. Read it, then go ask whether your "AI" is one model guessing — or a system that can be wrong safely.

*A crowd with no evaluator is just a more expensive way to be confidently *wrong*.*

— RDLB Agentic, editorial. New York. rdlbagentic.com

The crowd beats the *soloist*.

The dominant shape of a serious agentic system is the orchestrator-worker pattern: one lead agent plans the work and spawns parallel sub-agents, each with its own context window and a self-contained task, whose condensed findings the lead reconciles.



ORCHESTRATOR-WORKER. A LEAD AGENT PLANS AND DELEGATES; SUB-AGENTS WORK IN PARALLEL, EACH WITH ITS OWN CONTEXT.

Anthropic documented this directly. Their multi-agent research system used one lead "researcher" directing sub-agents, and it outperformed a single strong agent by 90.2% on their internal evaluation — the gains concentrated on breadth-first work, where several independent directions can be pursued at once.⁴

The minimal mental model comes from OpenAI's Swarm, now folded into its Agents SDK: an agent is just *a prompt and a set of tools*, and a "handoff" is just a function that returns the next agent.⁵ Strip away the

mystique and a swarm is a small org chart you can read — roles, tools, and the rules for passing work along.

Why split the mind at all? Because separate context windows buy you parallel reasoning and stop one long conversation from drowning in its own history. Five agents each holding one clean problem will out-think one agent holding five.

That is the upside. The next two pages are the discipline that makes it more than expensive theater: a division of labor, and an evaluator the crowd cannot fool.

Five roles, one search.

A swarm is a division of labor, not a chorus of clones. The research lineage gives the roles their names — they are evolutionary operators with a job each.

01 · EXPLORER

Widen the field

Propose unlike, untried directions. Breadth first — the job is variety, not polish, so the search never collapses onto one idea too early.

02 · EXPLOITER

Refine what works

Take the strongest candidate and push it further — local, depth-first improvement on a lineage that is already paying off.

03 · REPAIRER

Fix the broken

Take a promising-but-infeasible design and make it legal against the constraints, rescuing good ideas that failed on a technicality.

04 · CROSSBREEDER

Combine lineages

Splice the strengths of two different winners into one candidate — crossover, the oldest trick in evolutionary search.

05 · CRITIC

Try to break it

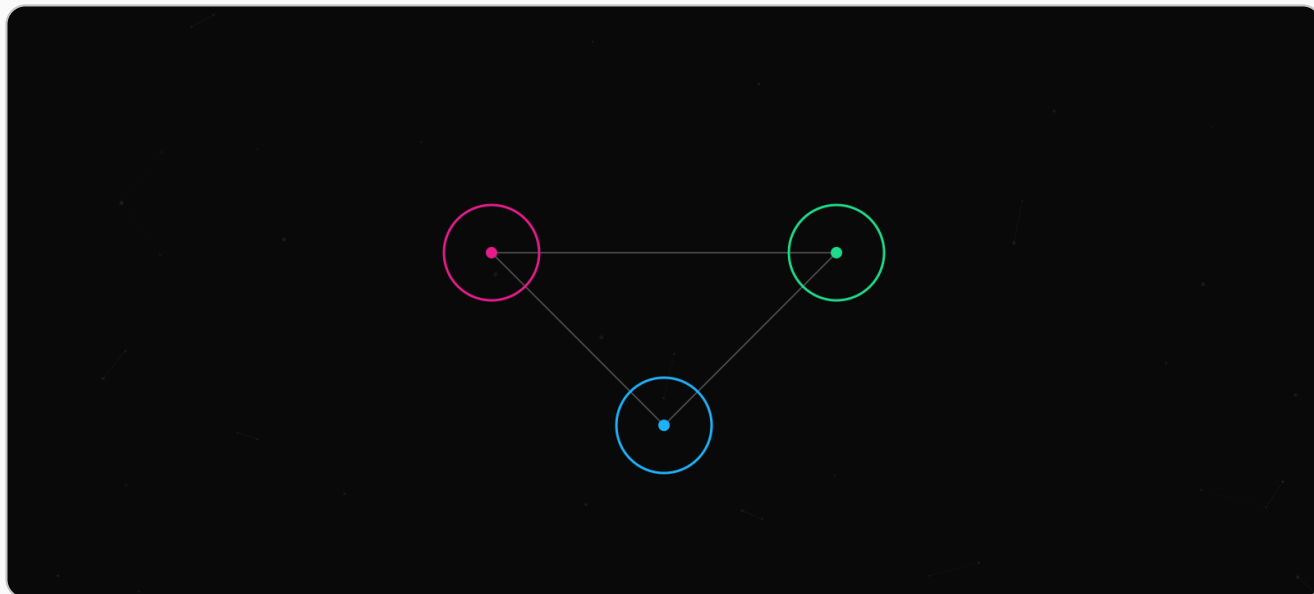
Adversarial by design. The critic hunts for designs that *look* good to the evaluator while being fragile or gamed — the swarm's defense against winning dishonestly.

WHY A BRAND OPERATOR SHOULD CARE

A swarm without a critic optimizes for *looking* right, not *being* right. If no role is paid to attack the output, your system will quietly learn to satisfy the metric instead of the customer. Staff the critic before you scale the crowd.

Propose, verify, *evolve*.

The reason a swarm can be wildly creative without being reckless is one architectural choice: the model only proposes; a trusted, deterministic evaluator decides. The model is allowed to be wrong because something honest is the source of truth.



THE LOOP: A MODEL PROPOSES CANDIDATES, A TRUSTED EVALUATOR SCORES THEM, THE STRONGEST SURVIVE INTO THE NEXT ROUND.

DeepMind's FunSearch named the pattern in *Nature*: pair a language model with an automated evaluator that guards against confabulation, then evolve the best programs across generations. It found new constructions for the cap-set problem — the largest improvement in twenty years — and better online bin-packing heuristics.¹

AlphaEvolve scaled it. An ensemble — a fast model for breadth, a stronger one for depth — evolves whole

algorithms under evaluator feedback. It found a way to multiply two 4×4 complex matrices in 48 scalar multiplications: the first improvement on Strassen's 1969 result in this setting in fifty-six years.² Same shape, larger stakes.

The lesson for anyone buying agents: ask what plays the evaluator. If the only judge of the work is the model that made it, you do not have a system — you have a confident narrator.

48

SCALAR MULTIPLICATIONS FOR A
 4×4 COMPLEX MATMUL²

56 yrs

SINCE THE LAST SUCH
IMPROVEMENT ON STRASSEN²

20 yrs

CAP-SET GAIN FROM A MODEL
+ A TRUSTED EVALUATOR¹

Four papers on the *swarm*.

From a single discovery to a scaled coding agent to illuminated search to the orchestrator-worker blueprint. Each freely available and verified live for this issue.

ROMERA-PAREDES, B., BAREKATAIN, M., ET AL. (2023) • NATURE

Mathematical discoveries from program search with large language models

nature.com/articles/s41586-023-06924-6

WHY IT MATTERS

FunSearch. The first time an LLM made a genuine discovery — because a trusted evaluator, not the model, decided what was true.

GOOGLE DEEPMIND (2025) • ARXIV:2506.13131

AlphaEvolve: A Coding Agent for Scientific and Algorithmic Discovery

arxiv.org/abs/2506.13131

WHY IT MATTERS

The pattern, scaled: an ensemble of models evolving algorithms under evaluator feedback — and beating a 56-year-old matrix-multiplication record.

MOURET, J.-B., & CLUNE, J. (2015) • ARXIV:1504.04909

Illuminating search spaces by mapping elites

arxiv.org/abs/1504.04909

WHY IT MATTERS

MAP-Elites. Don't just optimize — illuminate. Keep a grid of diverse strong solutions and you see the whole trade-off surface, not one winner.

ANTHROPIC (2025) • ENGINEERING

How we built our multi-agent research system

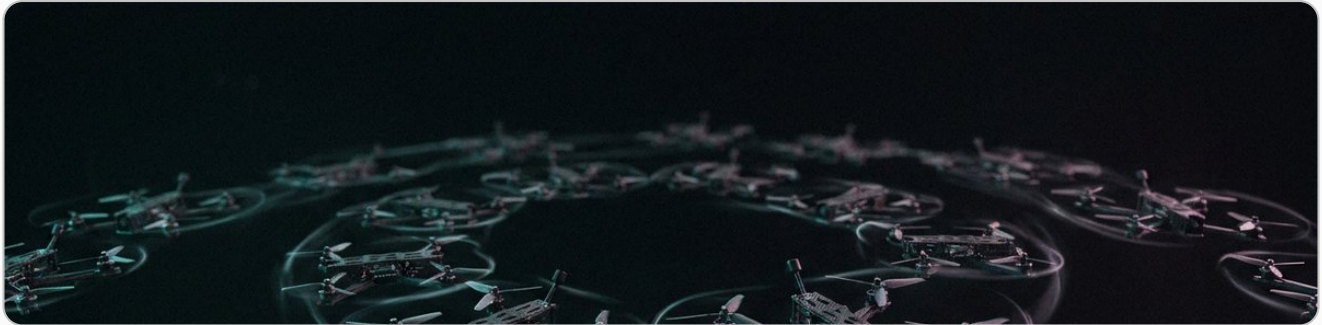
anthropic.com/engineering/multi-agent-research-system

WHY IT MATTERS

The orchestrator-worker blueprint in production — with the honest economics: 90.2% better than one agent, at roughly 15× the tokens.

Anthropic's 90.2%, and the bill.

A swarm is not free intelligence. It is a trade — more capability for more cost — and Anthropic published both halves honestly, which is exactly why it is worth more than a launch-day boast.



A COORDINATED SWARM: MANY SMALL UNITS, ONE FORMATION. POWER COMES FROM THE COORDINATION, AND SO DOES THE COST.

90.2%

MULTI-AGENT OVER A SINGLE
STRONG AGENT⁴

~15x

MORE TOKENS THAN
A SINGLE CHAT⁴

80%

OF THE PERFORMANCE
VARIANCE IS TOKEN USE⁴

Anthropic's research system used Claude as a lead agent directing sub-agents, each with separate context. On their internal research evaluation it beat a single agent by 90.2%, with the largest gains on breadth-first queries — the kind that benefit from chasing several leads at once.⁴

Then the honest part. Multi-agent runs consumed roughly fifteen times the tokens of an ordinary chat,

and token usage alone explained about 80% of the performance variance. Their own conclusion: the architecture is worth it only when the task's value clears the token bill.⁴

That is the buying rule in one line. A swarm is a **cost decision**, not a default setting. For a high-value, open-ended problem, paying 15x for a 90% better answer is obvious. For a lookup, it is waste.

THE OPERATOR'S TAKEAWAY

Match the shape to the task. Reserve the swarm for breadth-first, high-value work where many directions pay off. Route the simple, single-thread jobs to one agent. The skill is not "use more agents" — it is knowing **when** a crowd earns its bill.

Signals from the *swarm*.

Dated, sourced, and free of the word "magic."

2015

Illuminate, don't just optimize. MAP-Elites proposes keeping a grid of diverse elite solutions instead of converging on a single winner — the idea that quality and variety are not enemies.³

DEC 2023

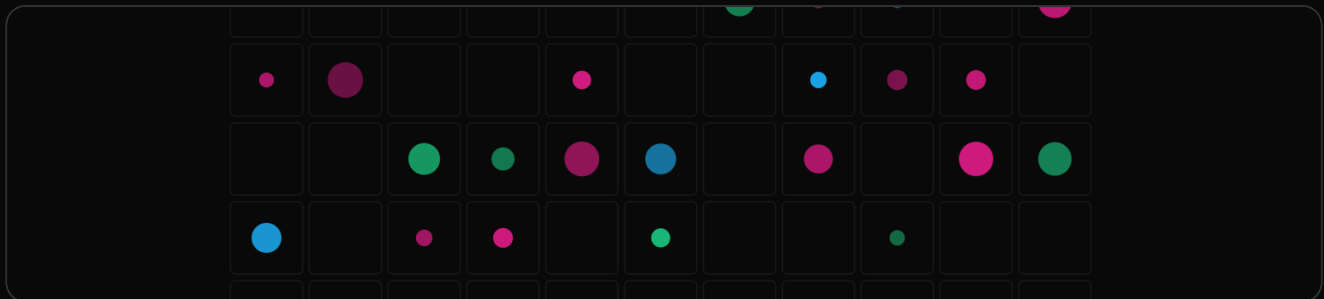
A model makes a discovery. DeepMind's FunSearch finds new cap-set constructions — the largest improvement in twenty years — by pairing an LLM with a trusted evaluator.¹

JUN 2025

Evolution writes code. AlphaEvolve evolves a procedure to multiply 4×4 complex matrices in 48 scalar multiplications — the first gain over Strassen in fifty-six years.²

JUN 2025

The crowd, measured. Anthropic reports a multi-agent system beating a single agent by 90.2% on its research eval — at roughly $15 \times$ the tokens.⁴



MAP-ELITES: AN ILLUMINATED ARCHIVE. NOT ONE WINNER — A GRID OF DIVERSE STRONG DESIGNS AND THE TRADE-OFFS BETWEEN THEM.

WHAT WE'RE WATCHING NEXT MONTH

From coordinating agents to compounding them: memory, evaluation, and the systems that get better the longer they run. That is Issue N° 05.

How we run a *swarm*.

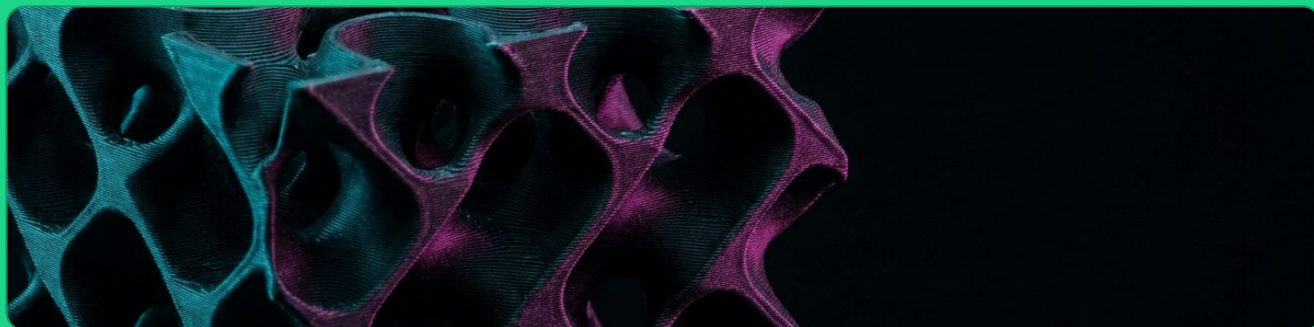
A magazine about orchestration should run on an orchestrated system. Here is how this issue's research disciplines our own thirteen-agent roster.

Orchestrator, then workers. A lead plans the work; named agents execute in parallel, each with its own context. Like Anthropic's system, we split the mind on purpose — and only when the task is broad enough to earn it.⁴

A trusted evaluator decides. The model proposes; deterministic checks and read-only data decide what is true. We keep the capability separation that every result in this issue depends on — the model never gets to grade its own work.¹

A standing critic. One role exists to attack the output and hunt for answers that satisfy the metric while failing the brand. It is our defense against a swarm learning to look right.

The human keeps the gate. Throughput goes to the agents; judgment stays with a person. Every action runs read-only by default, behind an approval gate, with audit-grade logs — and we price the tokens, scaling the crowd only where the value clears the bill.⁴



STRUCTURE FROM REPETITION: A 3D-PRINTED LATTICE. A SWARM'S STRENGTH IS THE SAME — SIMPLE UNITS, DISCIPLINED ARRANGEMENT.

If we *cited* it, here it is.

- 01 Romera-Paredes, B., Barekatin, M., Novikov, A., et al. (2023). *Mathematical discoveries from program search with large language models* (FunSearch). Nature.
[nature.com/articles/s41586-023-06924-6](https://www.nature.com/articles/s41586-023-06924-6)
- 02 Google DeepMind (2025). *AlphaEvolve: A Coding Agent for Scientific and Algorithmic Discovery*. arXiv:2506.13131.
arxiv.org/abs/2506.13131
- 03 Mouret, J.-B., & Clune, J. (2015). *Illuminating search spaces by mapping elites* (MAP-Elites). arXiv:1504.04909.
arxiv.org/abs/1504.04909
- 04 Anthropic (2025). *How we built our multi-agent research system*. Anthropic Engineering.
anthropic.com/engineering/multi-agent-research-system
- 05 OpenAI. *Swarm* (educational) & the *Agents SDK*.
github.com/openai/swarm • openai.github.io/openai-agents-python
- 06 Google DeepMind. *FunSearch & AlphaEvolve* announcements (orientation).
deepmind.google/blog/alphaevolve...

COLOPHON

Set in three faces echoing the RDLB system: a geometric display sans, a humanist Garamond for body and its italic payload word, and JetBrains Mono for labels. This issue leads in constellation neon — the Orchestration accent — and wears the tri-color asterisk, because a swarm is many colors holding one shape. Black, white, and the constellation palette only. Photography and diagrams generated for this issue.

NEXT ISSUE • NO. 05

Compounding systems: memory, evaluation, and agents that get better the longer they run. Out September 2026.



THE SINGLE THING TO DO AFTER READING THIS

Map your highest-cost *work.*

Book a 30-minute strategy blueprint call. You leave with a one-page map of your highest-cost work, sorted three ways — what to replace, what to augment, and what to leave alone — against real numbers, not vibes.

dashboardrdlbagency.com/book →