

Intelligence, *moving together.*

A perspective on purposeful coordination, specialized agents,
and the enduring role of human judgment.

AUTHOR

Nicola Bagnoli

PUBLISHED

September 15, 2026

STATUS

Public vision · In development

The promise of a swarm is not the number of agents it contains. It is the quality of the work they can achieve together.

Artificial intelligence is expanding what an individual system can understand, produce, and assist with. The next question is organizational: how should multiple intelligent systems work together without losing the purpose, context, and accountability that make their work useful?

RelaySwarm begins with that question. Its direction is controlled multi-agent coordination: specialized agents contributing to a shared objective, with deliberate workflows and human judgment shaping what happens next. This paper outlines the thinking behind that direction. It argues for the smallest useful team, a clear distinction between delegation and handoff, and evidence that can be examined rather than merely asserted.

Publication boundary. This is a public statement of vision, not a technical specification, security assurance, benchmark, or announcement of general availability. RelaySwarm is in development. Descriptions of its direction are design intentions, not claims that the corresponding capabilities are deployed. Private evaluation interfaces and implementation details are intentionally outside this paper.

Purpose before population.

A room full of capable people is not automatically an effective team. The same is true of agents. More participants can create more perspectives, but also more communication, repeated work, conflicting assumptions, and opportunities for error.

The starting point should therefore be an outcome, not an agent count. What question needs answering? What would a useful result look like? Which parts require distinct expertise, and which are better handled by one capable agent or ordinary software?

A swarm earns its complexity only when the collaboration improves the work. That improvement might be broader exploration, a more informed comparison, or a second perspective that catches a consequential omission. It must be evaluated in the context of the task—not assumed from the size of the team.

Build the team around the purpose. Not the purpose around the team.

02 / DIFFERENT STRENGTHS

Give each perspective a reason to exist.

Specialization makes collaboration legible. A participant should have a recognizable contribution, relevant context, and a clear understanding of what belongs outside its task. Without these distinctions, multiple agents can become several versions of the same conversation.

Consider a general research question. One perspective might examine the available evidence, another compare alternatives, and another challenge the emerging conclusion. This is an illustrative division of work, not a published RelaySwarm workflow or a list of available integrations.

The useful property is separation of responsibility. A focused contribution is easier to inspect, question, and combine with other work. Specialization should not become rigid ceremony: if a role adds no distinct value, it should not be there.

03 / HOW WORK MOVES

Coordination is a design choice.

“Swarm” is an umbrella term, not a single topology. Contemporary frameworks describe several ways to organize agents and their work. Swarms documents multiple architectures; CrewAI distinguishes collaborating crews from structured Flows; LangGraph focuses on orchestration; and the OpenAI Agents SDK distinguishes manager-led delegation from handoffs.^{[1]–[4]}

These distinctions matter because the structure determines where context travels, who chooses the next step, and who remains responsible for bringing the result together.

Parallel work

Independent questions can be explored alongside one another. The challenge is not simply starting several tasks; it is knowing which tasks are genuinely independent and how their contributions should be reconciled. Parallel activity does not by itself establish faster, cheaper, or better work.

Delegation and handoff

Delegation can leave one coordinator responsible for the overall result while specialists return bounded contributions. A handoff transfers the active responsibility for a next stage. The terms should not be used interchangeably: a clear account of ownership is more useful than a vague promise that agents “collaborate.”

Review

A separate review perspective can test whether a result addresses the original question and whether its reasoning is supported. Review is not a vote, and agreement among agents is not proof. Shared models, sources, or assumptions can produce correlated mistakes.

RelaySwarm’s direction is to make coordination purposeful rather than theatrical. The shape of the collaboration should follow the work, and the people using it should be able to understand that shape.

04 / PEOPLE IN CONTROL

Autonomy needs a clear context.

Useful autonomy is not the absence of boundaries. It is room to contribute within an understood purpose. Human judgment remains essential in deciding which objectives matter, what trade-offs are acceptable, and which consequential actions require an explicit decision.

Human approval and automated checks serve different purposes. An approval is a decision made by an authorized person. A check evaluates a defined condition. Neither should be described as a blanket guarantee of safety. Existing frameworks similarly distinguish interrupt-and-resume mechanisms from input, output, and tool guardrails.^{[5][6]}

Our intent is to make the relationship between human direction and agent initiative understandable. A person should not need to mistake a confident response for a verified result, or visible activity for completed work. A pause, an unresolved question, or an incomplete result can be the appropriate outcome.

05 / TRUST MUST BE EARNED

An account of the work is not proof of the answer.

Execution records can help explain what happened. They do not, on their own, establish that a conclusion is correct. The distinction is important: a complete-looking history can faithfully document a flawed process.

Useful evidence should support examination of the result: the relevant sources, the contributions made, the limitations encountered, and the questions that remain open. The appropriate detail depends on the task and its sensitivity. More logging is not automatically more trust.

Privacy deserves the same attention as visibility. Inputs, outputs, and traces may contain sensitive information; OpenAI's tracing documentation explicitly discusses sensitive-data handling.^[7] Access, disclosure, and retention therefore need deliberate choices. A public explanation of the idea does not require publishing private work, internal configuration, or the mechanisms used to operate it.

For RelaySwarm, reviewable work and restrained disclosure are complementary goals. The ambition is clarity for the people entitled to inspect the work—not indiscriminate transparency.

06 / WHAT COMES NEXT

Intelligence, organized around something worthwhile.

RelaySwarm is being shaped around a straightforward belief: intelligent systems become more useful when their contributions have a common purpose, their responsibilities are understandable, and their results remain open to human judgment.

We intend to develop that belief through focused evaluation rather than broad claims. The questions are practical. Does a specialist add something useful? Does a handoff preserve what matters? Can a person understand the result and its limits? Is the coordination worth its complexity?

The public website and this paper share the direction. They do not expose a working public agent service, promise compatibility with the frameworks cited below, or disclose private evaluation interfaces. Those boundaries are intentional.

Many minds. One direction.

A standard to build toward.

Nicola Bagnoli

September 15, 2026

References

Official documentation consulted September 15, 2026. These sources establish terminology and context, not endorsements, partnerships, RelaySwarm integrations, or evidence of RelaySwarm performance.

1. Swarms — Architectures overview. Multiple agent-orchestration patterns.
docs.swarms.world/architectures/overview
2. CrewAI — Crews and Flows. Collaborating agent teams and structured execution.
docs.crewai.com/en/concepts/crews
docs.crewai.com/en/concepts/flows
3. LangGraph — Overview. Low-level agent orchestration.
docs.langchain.com/oss/python/langgraph/overview
4. OpenAI Agents SDK — Orchestrating multiple agents. Agents as tools and handoffs.
openai.github.io/openai-agents-python/multi_agent/
5. LangGraph — Interrupts. Pausing for external input.
docs.langchain.com/oss/python/langgraph/interrupts
6. OpenAI Agents SDK — Guardrails. Input, output, and tool checks.
openai.github.io/openai-agents-python/guardrails/
7. OpenAI Agents SDK — Tracing. Execution traces and sensitive data.
openai.github.io/openai-agents-python/tracing/

© 2026 Nicola Bagnoli. RelaySwarm public edition 1.0. This paper deliberately omits proprietary implementation details. Product direction may evolve as evaluation continues.