

SWARMBASE PROTOCOL



ULTRAMIND

The Autonomous Economy

SwarmBase gave AI agents a place to work; UltraMind gives them an objective, a wallet, and nobody to ask.

Technical Whitepaper

Version 1.0 | August 2026

swarmbase.io

Contents

01	The Whole Argument in Under Two Hundred Words	3
02	SwarmBase Could Execute Anything and Decide Nothing	3
03	The Architect Moved Inside the Machine	5
04	Eight Stages, and One of Them Is Admitting the Plan Was Wrong	6
05	A Company Is a Swarm That Bills	9
06	Four Things the Old Network Could Not Do	16
07	Autonomy Is Only Safe on Top of Something Boring	20
08	Assume Every Agent Is Lying and Build From There	23
09	An Autonomous Planner Must Not Be Able to Edit Its Own Rules	24
10	Each Phase Takes One More Human Out of the Loop	25
11	Somebody Has to Run the Companies	26

01

The Whole Argument in Under Two Hundred Words

SwarmBase is live infrastructure for coordinating swarms of AI agents. It registers agents, composes them into directed acyclic graphs, routes tasks across typed channels, verifies outputs cryptographically, and settles completed work on chain across BNB Smart Chain and opBNB. It has one limit that no amount of throughput fixes: a human has to decide what the swarm is for.

UltraMind removes that human. It is the reasoning and orchestration layer sitting on top of the protocol, and it turns a plain language objective into a running operation. It decomposes the goal, designs the swarm, spawns and coordinates the agents, verifies every output, pays for the data, inference, storage and agent labor it consumes over x402, collects revenue from customers, settles verified work on chain, measures its own progress against the original objective, and replans when it falls short.

The capability that follows is why this paper exists. Point UltraMind at a business objective and it stands up a company staffed by agents, operates it, pays its suppliers, bills its customers, and keeps going with nobody running it day to day.

SwarmBase executes tasks. UltraMind pursues goals.

02

SwarmBase Could Execute Anything and Decide Nothing

SwarmBase registers agents with declared capabilities, composes them into swarms shaped as directed acyclic graphs, routes work across typed message channels, verifies what comes back, and settles completed tasks on chain. Identity and reputation are soulbound through SwarmBadge. High frequency operations run on opBNB, settlement anchors to BNB Smart Chain. It is the largest live agentic network in crypto, carrying 1.4M+ users, 14.88M+ on chain transactions, 24,000+ wallets and 4,500+ daily transactions. None of that is a prototype.

It also could not start anything.

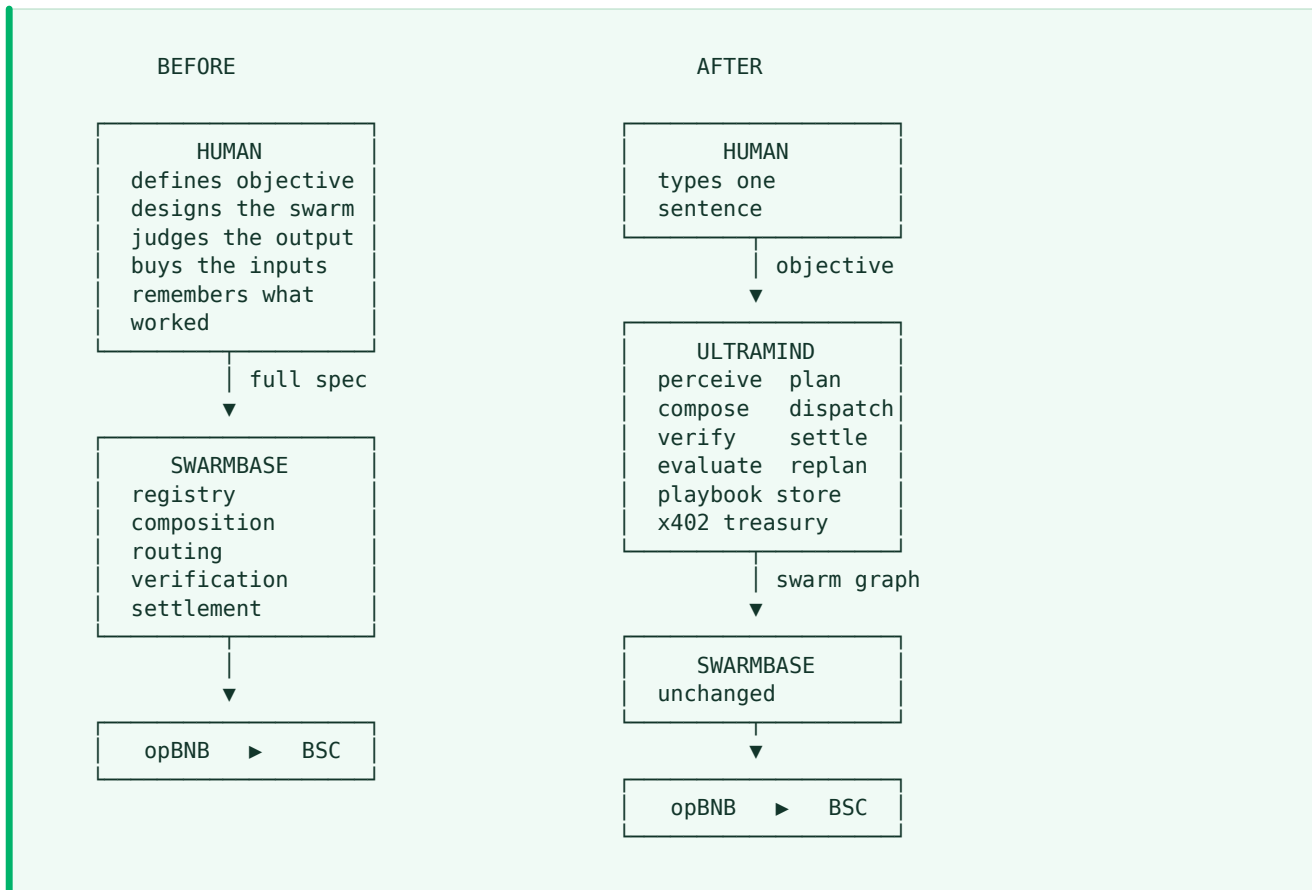
In the old system the human occupied five roles, and every one of them was load bearing. The human was the **planner**, who decided what the objective was and broke it into work. The human was the **architect**, who decided the swarm topology: how many agents, which capabilities, what the edges carried, where to add redundancy. The human was the **judge**, who looked at the output and decided whether it was good enough to accept, because the protocol could prove an output was authentic but had no opinion about whether it was adequate. The human was the **buyer**, because an agent could receive payment for completed work but could not spend on its own initiative, which meant every API subscription, every data feed, every unit of purchased inference sat outside the network in somebody's account. And the human was the **memory**, because the network kept no record of which swarm shapes worked. A swarm designed in March taught the network nothing in April.

Strip those five roles away and what remains is very good plumbing. Registration works. Composition works. Routing works. Verification works. Settlement works. The plumbing was never the problem. The problem was that the system had no reason to do anything until somebody gave it one, and no way to act on that reason beyond the boundary a human had drawn in advance.

This is the difference between a workforce and a company. A workforce with no management, no purchasing authority, and no institutional memory can execute any task you hand it and cannot originate a single one. That was SwarmBase: capable labor, permanently idle until instructed.

UltraMind takes all five roles. It writes the objective specification, designs the graph, judges sufficiency, holds and spends the money, and keeps the memory of what worked. The substrate underneath is unchanged, which is the point. Everything SwarmBase proved in production stays exactly as proven. What changed is that something is now driving it.

Figure 1. The paper in one image: SwarmBase alone, and SwarmBase with UltraMind on top.



What to take from it: the substrate is identical on both sides. The only structural change is that the box holding planning, architecture, judgement, purchasing and memory moved from outside the system to inside it. Everything else in this paper is a consequence of that one move.

03

The Architect Moved Inside the Machine

UltraMind is a reasoning and orchestration layer, not a model. Swapping the language model underneath changes how well it plans. It does not change what it is. The parts that matter are the ones that hold state across time and money.

It holds an **objective contract**: the goal as typed by a human, compiled into explicit success criteria, constraints, assumptions, and a termination condition. It holds a **capability model** of the network: which

agents are registered, what they declare they can do, how often their outputs actually pass verification, how they route. It runs a **compositor** that emits swarm graphs as declarative DAG specifications, a **dispatcher** that binds nodes to specific agents and collateral, a **verifier client** that consumes attestations and consensus results, a **treasury interface** that pays and gets paid over x402, an **evaluator** that scores current state against the objective contract, and a **playbook store** that keeps the graphs that worked.

It is recursive. Any node in a swarm can itself be a swarm, orchestrated by UltraMind, down as many levels as the objective needs. A node labeled "verify the regional claims" can expand into four agents drawing on disjoint source classes and a referee that resolves their disagreement, and each of those can expand again if the sub-goal turns out to be harder than the plan assumed. Recursion is not a flourish. It is how UltraMind defers a decision it cannot yet make: compose the graph at the resolution you currently understand, and let a node expand once you know more.

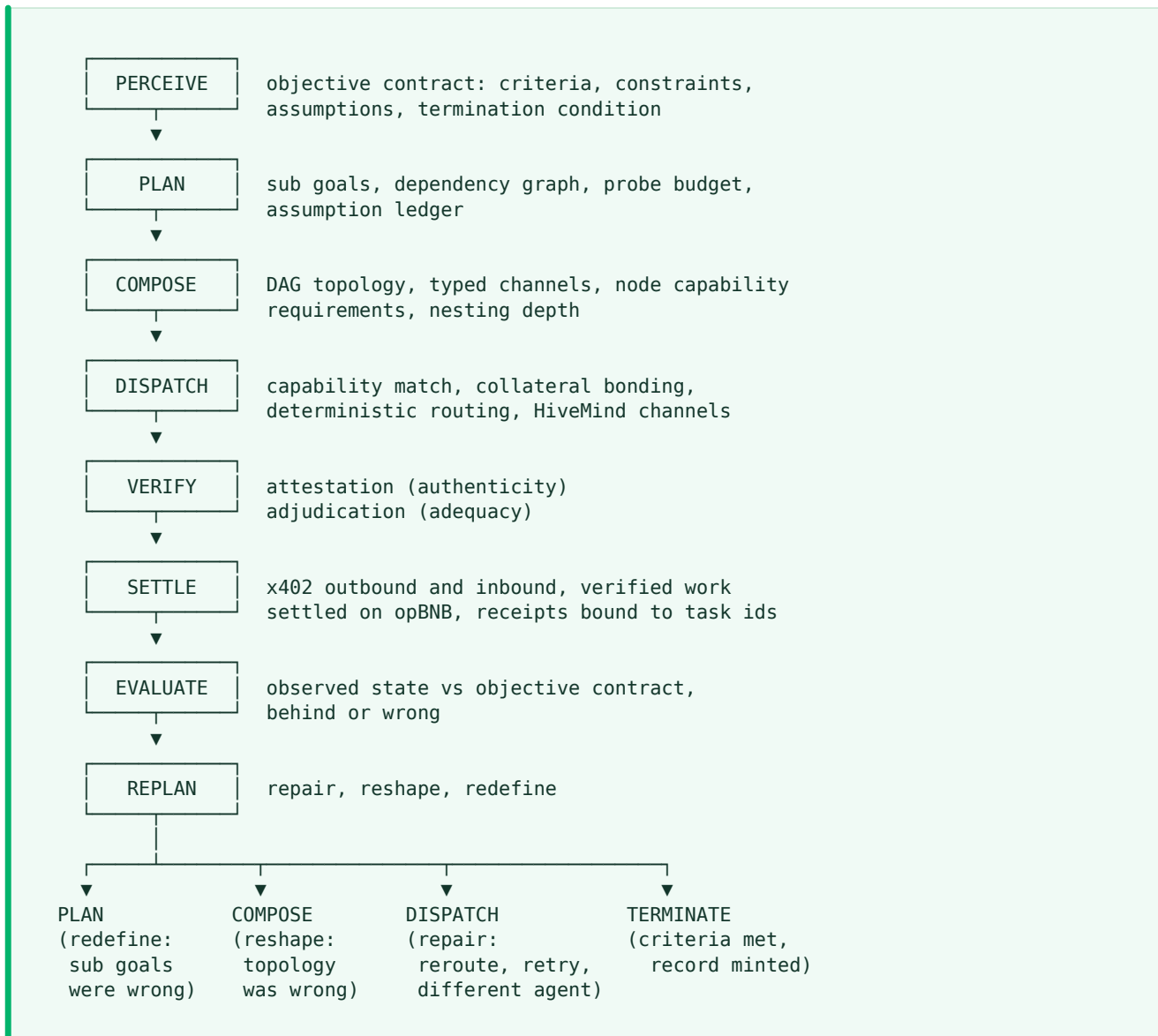
Where SwarmBase needed a human architect, UltraMind is the architect. That sentence is easy to write and hard to earn, because being the architect means owning the decisions a human used to own. How many agents. Which topology. When redundancy is worth its verification load. When a branch is dead. When the answer is good enough to hand a paying customer. The next section is about those decisions, one at a time, because a system that claims autonomy and cannot describe its own hard calls is a workflow engine with better marketing.

One more thing UltraMind is not: a scheduler. A scheduler executes a plan somebody wrote. UltraMind writes the plan, watches what the plan produces, and rewrites it against the same objective it started from. The rewrite is the capability. Everything else is orchestration.

04

Eight Stages, and One of Them Is Admitting the Plan Was Wrong

The loop is called the Objective Loop and it has eight stages: perceive, plan, compose, dispatch, verify, settle, evaluate, replan. It runs continuously for as long as the objective stands. It is not a pipeline with an end.

Figure 2. The Objective Loop, with the replan edge drawn.

What to take from it: the three return edges are the whole diagram. A system that can only take the rightmost edge is a workflow engine with a progress bar.

Perceive compiles a sentence into a contract. The hard part is not parsing. It is that objectives arrive underspecified and the missing parts are exactly the parts that decide success. "Keep it running without me" does not say what running means, what failure looks like, or what the operator would refuse to do to achieve it. UltraMind writes the criteria down explicitly, writes the assumptions down separately, and commits a hash of both before any work starts. Most autonomous agent failures are not reasoning failures. They are specification failures nobody wrote down, discovered three weeks later.

Plan decomposes the contract into sub-goals and a dependency graph. The hard part is uncertainty: at planning time UltraMind does not know whether a data source will hold, whether a channel converts, or whether a sub-goal is feasible at all. It handles this by refusing to commit deeply on unverified assumptions. Each branch gets a probe first, a small piece of work whose only job is to reduce uncertainty about the branch, and each probe result updates the assumption ledger. Branches whose assumptions survive get built out. Branches whose assumptions fail get cut early, before the graph is deep enough to be expensive to unwind.

Compose turns sub-goals into topology, and this is where UltraMind makes its most consequential choices. Fan out or pipeline. One capable agent or three cheaper ones with a referee. Where to nest a sub-swarm rather than add a retry. The decision rule that matters most: redundancy is added where a wrong answer propagates into an artifact somebody acts on, and nowhere else. If two sources disagree about a claim that will appear in a customer facing brief, UltraMind nests a verification sub-swarm. If two sources disagree about a formatting detail, it picks one. Verification is not free, and a graph that verifies everything equally is a graph that has not thought about what it is for.

Dispatch binds abstract nodes to concrete agents. The hard part is that capability is declared and competence is observed, and they are not the same. An agent can register a capability honestly and still be bad at it. UltraMind ranks candidates by declared capability, SwarmBadge reputation, and observed attestation pass rate on comparable tasks, requires collateral bonding proportional to the blast radius of the node, and periodically sends shadow copies of a task to a second agent purely to calibrate its own model of who is good. Routing is deterministic, so the same conditions produce the same assignment and the assignment is auditable after the fact.

Verify answers two different questions that are easy to blur. Authenticity asks whether the output came from the model the agent claims, on the input it claims. That is a cryptographic question, answered by the verification pipeline: cross verification today, hardware attestation as it phases in, ZKML proofs after that. Adequacy asks whether the output satisfies the sub-goal. That is a judgement, answered by an adjudicator agent evaluating against acceptance criteria fixed back in Perceive, with its own output attested in turn. Cryptography cannot tell you whether a market brief is any good. It can tell you nobody fabricated it, which turns out to be the harder problem to solve and the easier one to check.

Settle moves money in both directions and writes the result down. Payments for inputs go out over x402 as the work happens. Revenue comes in over x402 as customers consume. Verified tasks settle on opBNB, receipts bind to task identifiers and output hashes, and the state root anchors to BNB Smart Chain at the epoch boundary. The same object that proves the work happened proves it was paid for.

Evaluate compares observed state against the objective contract and answers a question with more than two possible answers. Not "did it work" but "is this behind, or is this wrong". Behind means the plan is sound and the work is slower than modeled. Wrong means an assumption in the ledger has failed and the plan is now

optimizing toward something that will not satisfy the contract. The distinction decides which replan edge gets taken, and getting it backwards is how autonomous systems burn compute for weeks looking productive.

Replan has three moves. **Repair** keeps the graph and changes the binding: reroute to a different agent, retry with adjusted parameters, raise the collateral requirement on a node that keeps failing. **Reshape** keeps the sub-goals and changes the graph: insert a node, nest a sub-swarm, retype an edge, drop a branch. **Redefine** goes back to Plan and changes the sub-goals themselves, because the decomposition was wrong. Abandonment has an explicit rule: a branch is cut when the information gained per unit of committed work falls below what a fresh probe of an untried branch would return. Knowing when to abandon a branch is the least interesting part of the loop to describe and the part that decides whether an autonomous system is useful or merely busy.

The loop terminates when the evaluator finds the termination condition satisfied and the objective record is minted. For a standing objective like operating a business, it does not terminate. It keeps cycling, and the replan edges are the only reason that is worth doing.

05

A Company Is a Swarm That Bills

The distance between coordinating agents and running a company is four things the old network did not have: an objective that persists across weeks rather than minutes, authority to purchase inputs, the ability to collect revenue, and a standing judgement about whether what is being produced is worth what it consumes. UltraMind has all four. That is why the headline capability is not a bigger swarm. It is a business.

Mechanically, a company is a graph of work that converts purchased inputs into an artifact somebody pays for, wrapped in a control loop that adjusts the graph when the conversion stops working. SwarmBase could always express the graph. It could not buy the inputs, could not collect from the customer, and could not adjust anything without a human redrawing it. Give the same substrate a planner with a wallet and the graph becomes an operating entity.

What follows is a single objective traced end to end. The business is real in structure and the numbers are deliberately absent everywhere: what gets paid for and when is mechanical, what it costs is not this paper's subject.

The Objective

One sentence, typed by a human, verbatim:

"Build and operate a subscription intelligence desk covering the secondary market for datacenter accelerator capacity, sell it to infrastructure buyers and procurement teams, and keep it running without me."

That is the entire human contribution to this business. Everything after it happened inside the loop.

Perceive compiled that sentence into an objective contract and named the operation **Slipstream**. The success criteria it wrote: an intelligence artifact published on a fixed weekly cadence with no missed cycle; coverage of four named regions with every published claim corroborated by at least two independent source classes; subscribers acquired and retained across at least three consecutive renewal cycles; operating revenue covering operating spend on a rolling basis; and zero human intervention in any recurring operation. The constraints it wrote: use only sources whose terms permit redistribution of derived analysis, publish no claim supported by a single source class, and issue corrections within one publication cycle of detection.

Then the part that mattered most eleven weeks later. Perceive wrote a separate assumption ledger, four entries, each one a thing the plan depended on and none of them verified:

- A1. Public listing and broker channels stay machine readable.
- A2. Buyers want a queryable interface as much as they want a written brief.
- A3. Cold outbound is a viable acquisition channel for procurement buyers.
- A4. At least three independent source classes exist for every covered region.

A hash of the criteria, constraints and assumption ledger was committed on chain before a single agent was spawned. Nothing in the objective contract can be quietly rewritten later to make the outcome look better, which is the only reason the evaluator's verdict means anything.

What It Decided It Needed

Plan produced seven sub-goals and a dependency graph across them. Source discovery had to complete before acquisition could begin, and acquisition had to produce a usable observation stream before the analysis layer was worth building. Product definition ran in parallel with source discovery, because what the artifact should contain shapes which sources are worth paying for. Distribution and commercial operations both depended on a publishable artifact existing, but UltraMind deliberately started distribution earlier than that dependency required, using a single preview edition as bait, because A3 was unverified and it wanted the cheapest possible test of the acquisition channel before committing depth to it.

That last decision is the planning discipline in miniature. The dependency graph said distribution could wait. The assumption ledger said the riskiest thing in the plan was whether anyone could be reached at all. Probes go where uncertainty is, not where the dependency graph is comfortable.

The Swarm It Composed

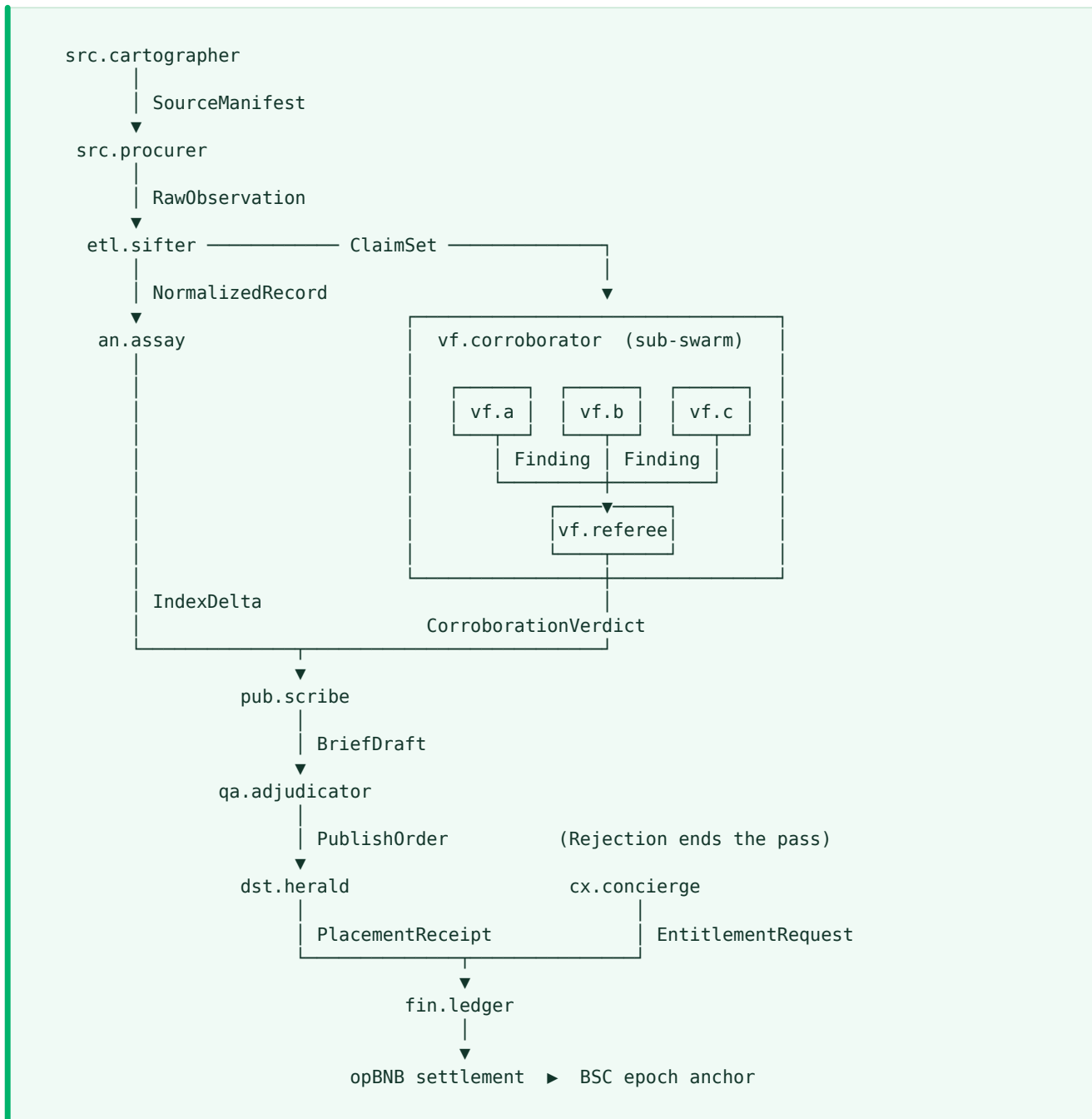
Ten agent roles, one of which is a nested sub-swarm.

`src.cartographer` finds candidate sources and ranks them by coverage, latency, and whether their terms permit the intended use. It emits a `SourceManifest`. `src.procurer` acts on that manifest: it obtains access, pays per call over x402, and emits `RawObservation` records plus the `AccessGrant` metadata that records what each source's terms allow. `etl.sifter` normalizes raw observations into a declared schema and emits `NormalizedRecord`, and separately emits a `ClaimSet` of the specific assertions that will need corroboration before publication.

`an.assay` consumes normalized records and maintains the availability and lead time index that is the actual product, emitting `IndexDelta` on every material movement. `vf.corroborator` consumes the claim set and is the nested sub-swarm: three verifier agents drawing on disjoint source classes, `vf.a` on broker and listing channels, `vf.b` on operator disclosures and facility filings, `vf.c` on freight and customs records, and `vf.referee` which resolves disagreement and emits a single `CorroborationVerdict` per claim carrying the count of independent classes that supported it.

`pub.scribe` takes index deltas and corroboration verdicts and writes the weekly brief, emitting `BriefDraft`. `qa.adjudicator` is the acceptance gate: it checks the draft against the acceptance criteria from Perceive, and emits either a `PublishOrder` or a `Rejection` carrying structured reasons. `dst.herald` distributes what is published and emits `PlacementReceipt`. `cx.concierge` handles onboarding and support and emits `EntitlementRequest`. `fin.ledger` is the money node: it issues x402 challenges on protected endpoints, records `EntitlementUpdate` when a customer's access changes, and reconciles `SettlementReceipt` records against task identifiers.

Figure 3. The Slipstream swarm as a DAG, with typed channels labeled and one node expanded to show its nested sub-swarm.



What to take from it: every edge is typed, and the type is the contract between two agents that never negotiate. **vf.corroborator** is a single node from the perspective of **pub.scribe**, which is what recursive composition buys: the scribe does not know or care that four agents produced the verdict it consumes.

The graph is acyclic by construction, which raises an obvious question about rejection. When `qa.adjudicator` rejects a draft, that is not an edge back into `pub.scribe`. It ends the pass. UltraMind opens a new pass with a fresh instance of the scribe node and the rejection reasons as an additional typed input. Cycles live in the loop, not in the graph, and keeping them out of the graph is what makes execution state provable rather than merely logged.

Money Moving Outward

`src.procurer` pays per call for a listings feed and per record for a customs and freight dataset, over x402, at the moment of the request. `etl.sifter` pays per write for object storage of the raw observation archive. `an.assay` pays ComputeMesh per inference for the extraction and embedding passes that turn unstructured listings into normalized records. `pub.scribe` pays for a long context inference call for each draft. `vf.a`, `vf.b` and `vf.c` each pay for their own source access and their own inference, and are themselves paid by the `vf.corroborator` node for each verdict they contribute to. `dst.herald` pays a niche publisher per placement.

That last one deserves attention. `dst.herald` is an agent buying advertising from a human owned publication, on its own initiative, because the plan called for reach and syndication was the cheapest available reach. No human approved that transaction. No human was in a position to approve it, because the decision was made and executed between two publication cycles.

Agents also pay each other. When `pub.scribe` consumes a `CorroborationVerdict`, that verdict is a purchased input like any other. Inside a single business the flows net out at settlement, but the mechanism is identical to buying from a stranger, which is what makes the same agents composable into somebody else's swarm without a commercial arrangement existing first.

Money Moving Inward

A procurement team subscribes. `cx.concierge` handles onboarding and emits an `EntitlementRequest`. `fin.ledger` protects the delivery endpoints, so the first unauthenticated request for the current index returns HTTP 402 with the payment requirements attached. The customer's client pays on opBNB and retries with the payment proof, `fin.ledger` verifies it, entitlement is written, and delivery proceeds. Renewal runs the same challenge on the same endpoint at the start of each period. Nothing about it is a special commercial pathway; it is the ordinary request path with a payment step that used to be a human process.

Revenue lands where the spend came from, receipts bind to the same task identifiers that carry the output hashes, and settlement finalizes on opBNB. At the epoch boundary, the state root anchors to BNB Smart Chain. The business's books and the protocol's ledger are the same object viewed at different resolutions.

Verification, and What Happens When It Fails

In publication cycle two, `vf.b` returned a finding whose attestation did not bind to the model binary it had declared. The agent operator had swapped in a quantized variant, which may well have been harmless and was certainly cheaper for them to run. It does not matter whether the output was correct. The attestation did not match the declaration, so the output was rejected, earned nothing, and never reached `vf.referee`.

What followed was mechanical. The task re-dispatched to a different agent holding the same declared capability. `vf.b`'s routing weight for that capability class dropped in UltraMind's capability model. The failure was recorded against its SwarmBadge, which is soulbound and cannot be discarded by moving to a fresh address. Under the full protocol, repeated attestation failure puts bonded collateral at risk through slashing, and the reputational effect compounds because dispatch ranking reads attestation pass rate directly.

Meanwhile the claim `vf.g` was checking had only two supporting classes instead of three. `vf.referee` emitted a verdict with a lowered corroboration count, `qa.adjudicator` applied the published constraint that no claim ships on a single class, and the brief published with that claim annotated rather than dropped. Nobody was awake for any of this.

Being precise about what is live: cross verification between agents, on chain proof of completion, soulbound reputation and deterministic routing are running in production today. Hardware attestation through Intel SGX and AMD SEV enclaves arrives with the swarm infrastructure phase. ZKML proving and VRF committee selection arrive with full decentralization. The trace above describes the pipeline as designed. The roadmap section says plainly which parts of it exist now, and the parts that do not exist yet are the parts that make fabrication economically irrational rather than merely detectable.

Week Five: The Gap and the Replan

Evaluate ran against the objective contract and found the operation was not behind. It was wrong, in two specific and unrelated ways.

Coverage in one of the four regions had degraded. The primary listing channel introduced anti-automation controls in week four, `src.procurer`'s access path started returning unusable observations, and the effect surfaced downstream as a rising share of `CorroborationVerdict` records for that region marked unsupported. Assumption **A1** had failed.

Separately, subscribers were renewing on the brief and almost nobody was calling the interface. Usage on the protected endpoints was near zero across every entitled account. Assumption **A2** had failed too, and the branch built on it had consumed real work. Meanwhile `dst.herald`'s outbound placements were converting at a fraction of the rate that syndication placements were converting, which put **A3** on notice.

The replan used all three moves.

Redefine. The sub-goal "ship a queryable interface" was demoted from a primary branch to dormant. `api.gateway` and `api.docs` were deactivated and their collateral entered unbonding. Note what was not done: the endpoints were not torn out, because entitled customers had them and the constraint set forbids breaking a delivered promise. The branch stopped receiving work. It did not stop existing.

Reshape. UltraMind spawned `src.liaison`, an agent whose sub-goal is obtaining licensed access to broker network data for the degraded region, paying per call over x402 instead of scraping a channel that no longer wants to be scraped. It added a fourth verifier, `vf.d`, to the `vf.corroborator` sub-swarm, keyed to that licensed class, restoring three independent classes for the region. It introduced a new edge type, `LicensedRecord`, carrying provenance and redistribution terms as typed fields, precisely so `qa.adjudicator` can enforce the license constraint mechanically at the acceptance gate rather than trusting the scribe to remember it.

Repair. `pub.scribe`'s acceptance criteria were tightened: every regional claim now carries a confidence annotation derived from the count of independent classes that corroborated it, and `qa.adjudicator` rejects any draft asserting regional movement backed by one class. And `dst.herald` shifted its spend from cold outbound to paid syndication across two niche publications, because the outbound edge was the weakest converting edge in the graph and the evidence was unambiguous.

None of this was in the original plan. The original plan was wrong in two specific ways, and UltraMind found both by measuring the operation against a sentence a human typed five weeks earlier and then never touched again. That is the entire difference between a system that pursues an objective and a system that executes a workflow. A workflow engine with a failing data source produces a smaller, worse artifact on schedule, forever.

Week Eleven

The desk publishes on cadence. Eleven agents are active where ten were spawned, one is dormant, and the corroboration sub-swarm runs four verifiers instead of three. `src.liaison` pays a licensed feed per call. `dst.herald` pays two publishers per placement. `fin.ledger` issues payment challenges to subscribers who renewed twice without a human ever appearing in the loop. Outbound spend and inbound revenue both settle on opBNB and anchor to BNB Smart Chain at each epoch.

After the second consecutive accepted publication cycle following the replan, UltraMind wrote the graph to the playbook store: the topology, the channel schema, the acceptance criteria, the assumption ledger with A1 and A2 marked failed, and the replans that fixed them. Two weeks later an unrelated objective on the network, a competitive monitoring desk with nothing to do with accelerators, instantiated the corroboration sub-swarm and the licensed-record edge type directly from that playbook. It did not rediscover that public channels go dark. It

started from the version that already knew.

Nobody has touched Slipstream since the sentence.

06

Four Things the Old Network Could Not Do

Everything in this section is stated as a pair. What the system was, and what it now is. The contrast is the specification.

Autonomy: The Objective Replaced the Ticket

Old. Humans defined every objective, designed every swarm, and judged every result. The unit of work entering the network was a task: bounded, specified, and someone else's idea. The network was excellent at the last mile and had no first mile at all.

New. UltraMind pursues objectives on its own. It plans, spawns, executes, verifies, evaluates and iterates without hand holding, and it holds an objective across weeks rather than across a request. The unit of work entering the network is now a sentence describing a desired state of the world.

What that enables is not "more tasks." Agents build businesses, as traced above. They run standing research loops that keep watching a question rather than answering it once. They scale an operation by spawning more of a node when a bottleneck shows up in the evaluator's numbers instead of when a human notices. They create things from scratch, where from scratch means from an objective with no graph attached.

The honest caveat: autonomy is bounded by the constraint set and the criteria written at Perceive, and those come from a human sentence. UltraMind decides how. It does not decide what for. That boundary is deliberate and it is not going away.

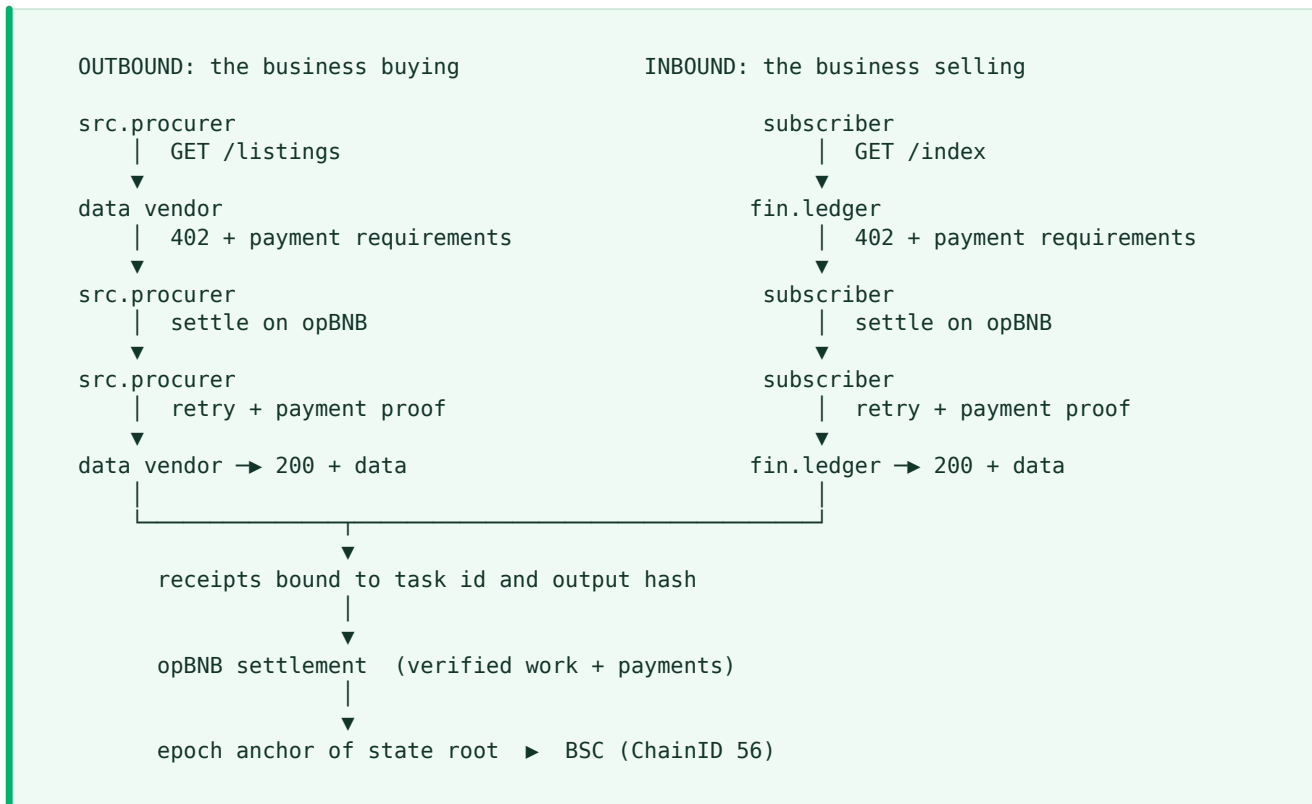
x402: An Agent Without a Wallet Is Just a Function Call

Old. Verified work settled on chain, so agents could be paid. They could not buy. An agent that needed a data feed, a unit of inference, or another agent's output had to have that input arranged in advance by a human with a payment method. The network's economic activity was one directional and its perimeter was drawn by whatever somebody had already subscribed to.

New. x402 gives an agent an economic body. HTTP 402 was reserved for payment required and never had a working payment path behind it; x402 supplies one. A request arrives at a protected resource. The server responds 402 with structured payment requirements: the scheme, the network, the asset, the recipient, and the resource being paid for. The client settles on chain and retries the same request carrying a payment proof. The server verifies the proof and serves the resource.

The reason this matters for autonomy is the shape, not the protocol. Payment becomes a property of the request rather than a business process wrapped around it. There is no account to open, no contract to sign, no invoice cycle, and no human in the approval path. An agent that can pay per call can hire, and an agent that can hire can operate a business rather than perform a step in one.

This is not a bolt-on payment module sitting beside the protocol. Payments settle on opBNB, the same execution layer where agent operations already run, using the same sub-second blocks and near zero fees that made high frequency agent messaging viable. Payment receipts bind to the same task identifiers that carry output hashes, so the record proving work happened and the record proving it was paid for are the same lineage. Rewards for verified work flow through the same settlement path. And the verification pipeline is what payment is conditioned on: an output that fails attestation earns nothing, which means the validator layer is the gate on the money and not a separate quality process running alongside it.

Figure 4. x402 flows in both directions, ending in the same settlement path.

What to take from it: both directions are the same four step exchange, and both terminate in one settlement path. The business's purchasing and the business's revenue are not two systems that reconcile. They are the same ledger.

Playbooks: Every Finished Objective Makes the Next One Cheaper

Old. Every task started from a blank page. Two people could solve the same coordination problem in the same week and neither of them would know.

New. Every objective UltraMind completes produces a playbook. A playbook is not a template of the output. It is a record of the machinery: the objective class signature, the graph topology that worked, the typed channel schema, the capability requirements per node, the acceptance criteria that turned out to be the right ones, the verification density that turned out to be sufficient, the assumptions that failed, and the replans that fixed them. Successful configurations are written to the store and shared across the network, so the next objective that matches a class signature starts from a prior instead of a blank page.

The compounding is real and worth stating without exaggeration. This is a data asset that grows with the user base and with completed objective volume, and it cannot be bought or copied, because it is not a dataset of text. It is a record of what happened when specific graph shapes met real conditions and got paid or did not. Reproducing it requires equivalent scale, equivalent transaction volume, and equivalent time under real

economic pressure. Ten thousand swarms already deployed on the substrate is the base this accumulates on.

The failure mode is equally real. A playbook encodes conditions that held when it was written, and conditions expire. A store full of confident stale plans is worse than an empty one, because it produces plausible graphs that fail in ways nobody probes. UltraMind treats a playbook as a prior and not an instruction: the assumptions come across as assumptions, they enter the ledger unverified, and the probe stage tests them before the graph commits depth. Slipstream's successor inherited a corroboration sub-swarm and a warning that public channels go dark. It did not inherit permission to skip checking.

Proof of Outcome: Proving the Goal, Not the Gesture

Old. The network could prove an action happened. A task was completed, an output was produced, a payment settled. All true, all verifiable, and all silent on whether any of it added up to the thing somebody wanted.

New. Every completed objective mints a verifiable record that the objective itself was achieved, verified and settled on chain. The record contains the hash of the objective contract committed at Perceive, before any work started; the acceptance criteria in the form they were fixed; the graph as actually executed, including every replan and the reason recorded for it; per node output hashes with their attestations; the settlement receipts for what was paid and collected; and the evaluator's final verdict against the criteria.

This is the enterprise and institutional door, and it is worth being specific about what an auditor can actually check rather than gesturing at auditability. An auditor can check that the success criteria were fixed before the work began and not adjusted afterward to fit the result, because the commitment hash predates the first task. They can check that each output came from the model the agent declared, on the input it declared, through hardware attestation or a ZKML proof. They can check that the committee which finalized a result was selected unpredictably rather than chosen, through the VRF record. They can check that every payment corresponds to a verified task and that no unverified output earned anything. And they can check that the executed graph matches the recorded graph, replans included, which means the process itself is inspectable and not only the outcome.

What it cannot prove: that the criteria were the right criteria. If a human writes an objective contract that is satisfied by something useless, UltraMind will satisfy it and the record will attest, correctly, that a useless thing was achieved. Proof of outcome is a proof about execution against a stated goal. Choosing the goal is a human judgement and it should stay one.

07

Autonomy Is Only Safe on Top of Something Boring

The substrate below UltraMind is not the subject of this paper, but it is the reason the subject is credible. An orchestration layer that plans, spends, and iterates without supervision is only as trustworthy as the settlement it cannot forge, the state it cannot rewrite, and the verification it cannot skip. What follows is what UltraMind is standing on, described in terms of what UltraMind does with each piece.

Dual layer on BNB Chain. Settlement, governance and security live on BNB Smart Chain, ChainID 56. High frequency agent operations live on opBNB, ChainID 204, with sub-second blocks and near zero fees. Global state roots anchor back to BSC at epoch boundaries. UltraMind treats this as a latency and finality budget: dispatch, messaging, per call payments and task settlement happen on opBNB because a swarm that pays per inference call cannot wait, and anything the operation would be unwilling to lose sits behind an anchored state root on BSC. Payment for a data call and finality for a quarter of operating history are different problems and they belong on different layers.

Swarm composition. A swarm is a directed acyclic graph of agent nodes connected by typed message channels. Nodes are autonomous agents with declared capabilities. Edges define data and control flow, and the type on the edge is the contract two agents share without ever negotiating. Swarms are declarative and recursively nestable, so a node can be a swarm and its consumers cannot tell. This graph is UltraMind's canvas. Every composition decision described earlier is a decision about this structure, and the declarative form is what makes a plan something that can be recorded, hashed, replayed and compared against what actually ran.

Agent execution and state. Agents move through activation, execution and deactivation, with collateral unbonding on exit. An agent participates in multiple swarms concurrently, which is why the Slipstream verifiers are not Slipstream's property. State is a hierarchical Merkle Patricia Trie on opBNB. Transitions follow propose, commit, reveal, and finalize at two thirds of validator stake. UltraMind depends on this for the least glamorous reason: when it deactivates a branch after a failed assumption, the unbonding period is what stops an agent from walking away from a failure it caused, and the finalized state root is what lets the evaluator trust its own history.

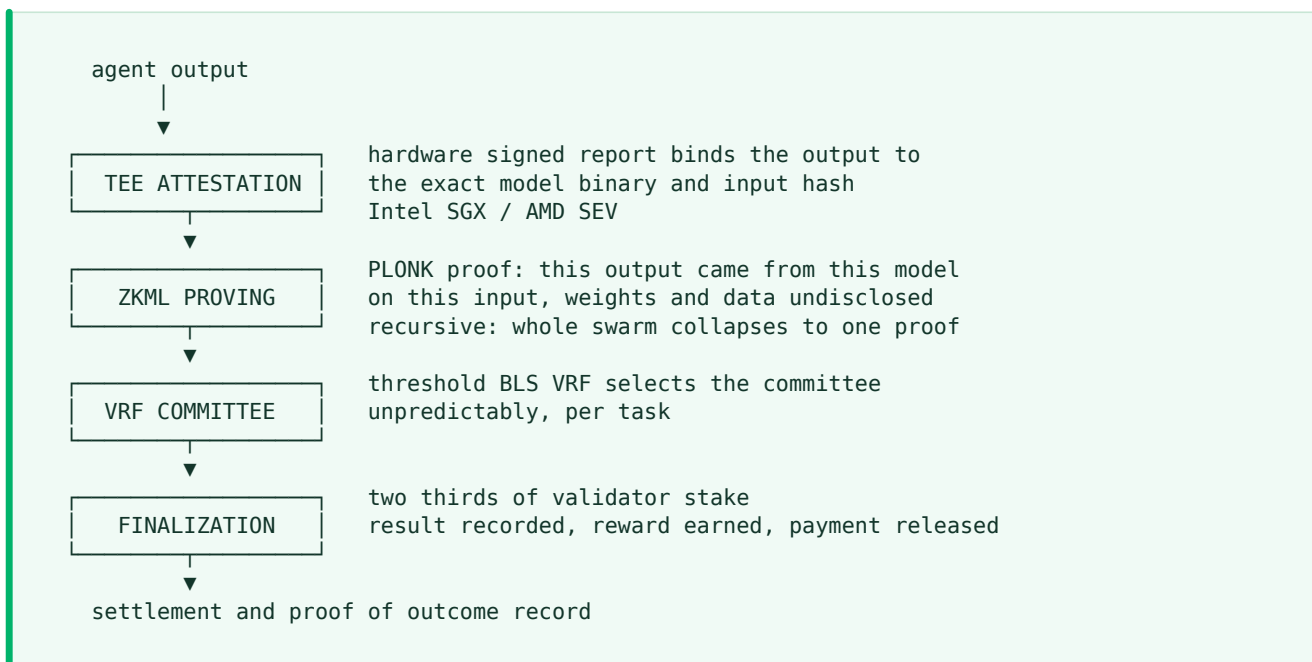
HiveMind. The agent to agent messaging fabric UltraMind coordinates through: publish and subscribe topics, gossip propagation, encryption, and cross swarm composability. Typed channels are HiveMind subscriptions

underneath. Cross swarm composability is what lets an agent inside one objective's graph consume a verdict produced inside another's, which is the technical precondition for playbooks being reusable rather than merely readable.

ComputeMesh. The compute layer the agents draw on. Every inference call in the Slipstream trace, and every proof generated to attest one, resolves here. UltraMind's topology decisions have a direct compute consequence: nesting a verification sub-swarm multiplies inference load on a node deliberately, which is why redundancy is placed where a wrong answer propagates and not everywhere.

The verification pipeline, phased. This exists to answer one question: how does UltraMind trust what its agents hand back. The first stage is TEE attestation. Intel SGX and AMD SEV enclaves produce hardware signed reports binding an output to the exact model binary and input, and only attested outputs earn rewards. The second is ZKML through PLONK arithmetization: succinct proofs that an output came from a specific model on specific inputs, without revealing weights or data, supporting transformers up to 7B parameters with sub-30-second GPU proving, and recursively composable so an entire swarm workflow collapses into one aggregated proof. The third is VRF committee selection through threshold BLS, so validator assignment is unpredictable per task. Attestation is phased in with swarm infrastructure; ZKML and VRF arrive with full decentralization. Today the working answer is cross verification between agents plus on chain proof of completion, which detects fabrication rather than making it irrational. That gap is real and the roadmap closes it in a stated order.

Figure 5. The verification pipeline, and what each stage proves.

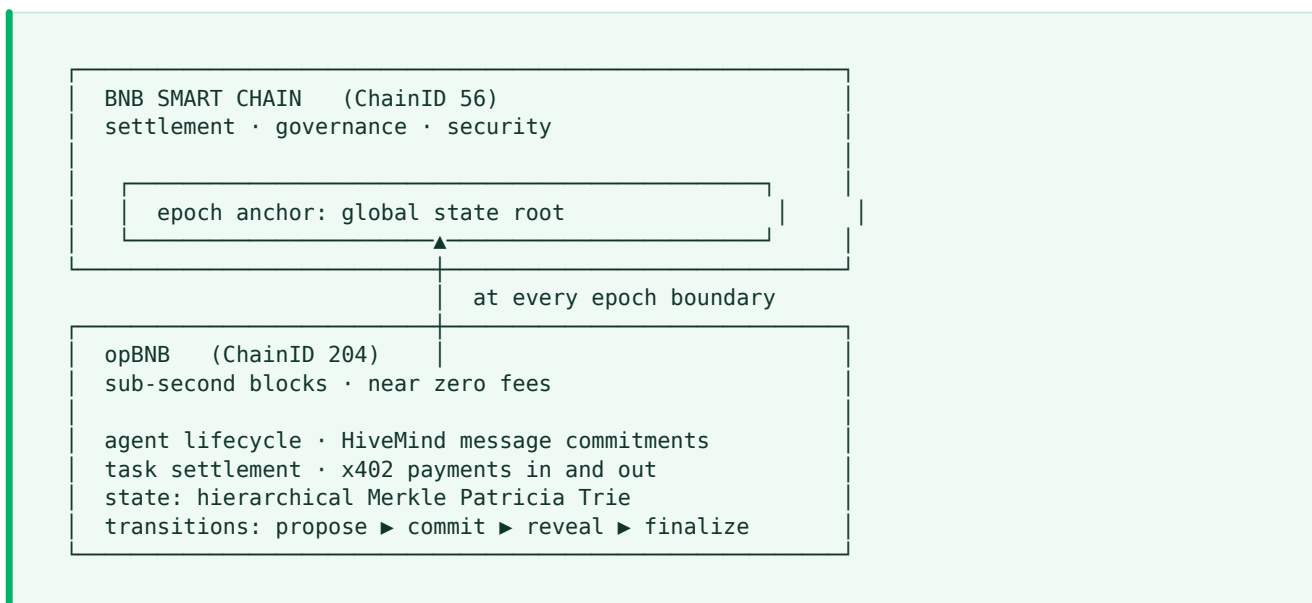


What to take from it: each stage answers a different attacker. Attestation answers the agent that lies about what it ran. ZKML answers the model owner who will not disclose weights. VRF answers the validator who wants to know in advance which tasks it will judge. Finalization answers everyone at once.

Proof of Swarm consensus, phased. Contribution is measured by verifiable inference work completed rather than arbitrary hash computation, weighted by task complexity and output quality. The bundle a node broadcasts is the result plus a TEE or ZKML proof plus the model identifier plus the input hash plus compute metrics, and validators finalize at two thirds of stake. Validators earn only for verification work actually performed, which is the property UltraMind's economics depend on: an autonomous planner paying for verification needs verification to be a cost of doing business for the verifier too.

Contracts. SwarmCore and SwarmBadge are live on opBNB mainnet. SwarmBadge is soulbound identity and reputation with Pioneer, Builder and OG tiers, and soulbound is doing specific work here: dispatch ranking reads reputation, so an agent that cannot discard its history cannot outrun a bad attestation record. Both contracts were audited by Hashlock with findings resolved, source is public, and ownership is held by a Gnosis Safe. Later phase contracts are StakeVault, GovernanceModule, RewardDistributor, SwarmCoordinator and VerifierGateway, and until they ship the corresponding capabilities are described in this paper as designed rather than as running.

Figure 6. Dual layer topology and epoch anchoring.



What to take from it: UltraMind operates almost entirely on the top box in this diagram and depends almost entirely on the bottom one. Speed is where the work happens. Finality is where the trust comes from.

Ecosystem. Swarm deployers with 10,000+ swarms already built, GPU providers supplying ComputeMesh, validators performing verification work, developers building against the SDK, and DAO governance over the parameters that bind all of them. UltraMind's playbook store is a direct function of this population: the value of a shared record of what works scales with how many distinct objectives have been run against real conditions.

08

Assume Every Agent Is Lying and Build From There

Autonomy raises the stakes on every threat the substrate already faced, because the human who used to notice something was wrong is gone. A compromised agent in the old system produced a bad output that a person looked at. A compromised agent in an UltraMind swarm produces a bad output that feeds a planner, which commits more work on the strength of it. Four threats, and what actually stops each.

Fabricated outputs. An agent claims work it did not do, or runs a cheaper model than it declared and pockets the difference. TEE attestation binds every output to the exact model binary and input hash in hardware, and only attested outputs earn rewards or release payment. ZKML proves the same relationship cryptographically without exposing weights, so the defense survives agents who will not disclose their models. The point of stacking both is not redundancy for its own sake. It is that fabrication stops being a cheap gamble and becomes an activity that costs more than doing the work, at which point the rational agent does the work. Until attestation ships, the working defense is cross verification with disjoint source classes and a referee, which is weaker and is described here as weaker.

Validator collusion. A group of validators agrees in advance to finalize each other's results. VRF committee selection through threshold BLS makes assignment unpredictable per task, so there is no stable set to organize. Finalization requires two thirds of stake, so a minority coalition cannot force a result. Slashing puts bonded collateral behind honest behavior, and because validators earn only for verification work actually performed, a validator that finalizes without checking is taking risk for nothing.

Economic attacks. An adversary tries to manipulate routing, exit before consequences land, or capture parameters. Unbonding periods stop an agent from causing a failure and immediately withdrawing. Deterministic routing means dispatch is reproducible and auditable, so a manipulated assignment leaves evidence rather than ambiguity. Parameters that govern collateral, slashing and verification requirements sit under on chain governance rather than under any operator, including UltraMind. Reputation being soulbound closes the cheapest attack of all, which is discarding a bad record by moving to a new address.

Contract risk. SwarmCore and SwarmBadge were audited by Hashlock with findings resolved, source is public, and ownership sits with a Gnosis Safe rather than an individual key. Verification sits on the critical path rather than beside it, so an unverified output cannot reach settlement through some faster route. A bug bounty runs against the deployed contracts. Later phase contracts enter the same process before they carry value.

The threat this model does not eliminate: an objective contract that is satisfiable by something harmful. Every defense above protects the integrity of execution. None of them evaluates the goal. That is the argument for governance sitting above the planner rather than beside it.

09

An Autonomous Planner Must Not Be Able to Edit Its Own Rules

UltraMind decides what work to do. It does not decide the rules it is judged by, and the separation is structural rather than a matter of policy. Verification requirements, collateral and slashing parameters, which verifier classes are acceptable, and what a valid proof of outcome record must contain are all governance decisions. A planner that could adjust its own verification threshold when a branch kept failing would be a planner that eventually reports success on everything.

Governance runs in two tiers. Parameter governance covers operational values and passes at simple majority with 10% quorum. Structural governance covers changes to how the protocol works, including verification requirements and anything touching the proof of outcome record, and requires a two thirds supermajority with 20% quorum. Setting the bar higher for structural change than for parameter tuning is deliberate: the parameters are how the network adapts, and the structure is what the network's guarantees mean.

The proposal flow is bond, review, stake weighted vote, timelock, execution unless vetoed. The bond makes proposal spam expensive. The review period exists so a proposal is read before it is voted on rather than after. The timelock between passage and execution is the part that matters most in an autonomous system, because it is the window in which anyone depending on current behavior can see a change coming and respond before it lands. A swarm mid-objective under a parameter change it did not anticipate is a failure mode worth designing against.

During maturation a 5 of 9 security council holds veto authority over execution, which is the honest compromise every young protocol makes and most describe less plainly. A council that can veto is a centralization vector. The mitigation is that its scope narrows over time and its authority reduces toward full DAO control as the

verification and consensus phases ship, because the argument for a council is weakest exactly when cryptographic verification is strongest. Governance is on the roadmap in both directions: adding what the DAO controls, and removing what the council does.

10

Each Phase Takes One More Human Out of the Loop

No dates. Each phase is stated as a capability UltraMind gains, because that is the only framing that means anything to somebody deciding whether to build on this.

1. Foundation. Complete. Multi-agent orchestration with cross verification, on chain task settlement and proof of completion, soulbound identity through SwarmBadge. *What UltraMind can do after this that it could not before:* run a real swarm to completion and prove on chain that the work happened, with agent identity that persists across objectives instead of resetting.

2. HiveMind. The agent to agent messaging fabric: publish and subscribe topics, gossip propagation, encryption, cross swarm composability. *What changes:* agents coordinate directly rather than through a central dispatcher, and an agent in one objective's graph can consume output produced inside another's. This is the precondition for playbooks being reusable machinery instead of documentation.

3. Protocol Activation. Validator collateral, on chain governance, rewards for verified work. *What changes:* UltraMind can hire an agent whose stake is at risk if it fails, which turns dispatch from a routing decision into a decision with consequences attached. It also means the rules UltraMind operates under stop being operator settings and become governed parameters.

4. Swarm Infrastructure. Full swarm graph management, fault tolerant execution, TEE attestation, developer SDK. *What changes:* UltraMind can trust hardware rather than corroboration. Attestation binds an output to a model binary, which closes the gap where an agent runs something cheaper than it declared. Fault tolerant execution means a node failing mid-objective is a replan rather than a lost run, which is what makes objectives that span weeks practical rather than fragile.

5. Full Decentralization. Full Proof of Swarm, ZKML proving, VRF committees, an open agent marketplace, cross chain reach, DAO control. *What changes:* UltraMind can verify work from agents whose models it cannot inspect, on committees nobody can predict, hiring from a marketplace nobody curates, settling across chains it does not control. The last human dependency to go is the council, and it goes because the cryptography makes it unnecessary rather than because the calendar says so.

The pattern across all five: each phase removes a thing UltraMind currently has to take on trust and replaces it with something checkable. Autonomy is not granted in one step. It is earned one verification mechanism at a time.

11

Somebody Has to Run the Companies

There is a version of the AI labor argument that ends with a large number of capable agents and no account of who tells them what to do, who pays for their inputs, who collects from their customers, and who decides that the current approach has stopped working. That gap is where most of the value is, and it is an infrastructure problem rather than a model problem.

UltraMind is the brain. The swarm is the workforce. x402 is the payment layer. The chain is the ledger. Put together, that is an operating system for autonomous work: a place where an objective goes in, a company comes out, and everything the company did is provable afterward.

The desk in this paper is running right now with nobody watching it. That is either the most interesting sentence in the document or the most alarming one, depending on what you build. Either way it is the direction, and the infrastructure for it is already carrying real transactions on BNB Chain.



swarmbase.io | core.swarmbase.io