

IF INSIGHT & FORESIGHT

STRATEGIC FORESIGHT · INTELLIGENCE · INNOVATION

THE AI HANDOVER

THE EMERGENT INTELLIGENCE-CENTRIC COMPUTING
ECOSYSTEM

REPORT · 8 SEPTEMBER 2026

SUBJECT

The organising architecture of computing.
Whether the application or the intelligence
layer determines the path from request to
result.

METHOD

Three testable hypotheses, evidence
current to 7 September 2026, stated
counterarguments, and signposts short
enough to change a decision.

EVIDENCE

35 dated sources: vendor announcements
and earnings material, incident
investigations, research literature and
named-voice commentary, each cited at the
point of use.

AUTHOR

Paulo Soeiro de Carvalho · IF Insight &
Foresight, with The Long Game.

WWW.IFFORESIGHT.COM · PAULO@IFFORESIGHT.COM

INTRODUCTION

TWO ANNOUNCEMENTS, EIGHT DAYS APART

On 26 August, Salesforce announced quarterly revenue of \$11.345 billion, an increase of 11% from the previous year. On the same day, it announced an expanded partnership with Anthropic that would allow people to perform Salesforce work from Claude. A company was reporting the strength of its business while making more of that business accessible through somebody else's interface. Both announcements deserve to be taken seriously. Together, they offer a better starting point for understanding the transition in computing than either a declaration that software is disappearing or an earnings release offered as proof that little has changed.

[01 · SALESFORCE RESULTS](#) / [02 · CLAUDEFORCE ANNOUNCEMENT](#)

Eight days later, OpenAI introduced GPT-6 Astra. Its release materials describe improvements in computer use and professional work, alongside changes to the environment through which work is performed. The release notes include the ability to continue while tools execute, incorporate instructions during an assignment and interrupt activity for safety review. These are changes in the practical relationship between a person and a computing system: what can be handed over, how it can be redirected and when it should stop.

[03 · ASTRA LAUNCH](#) / [04 · RELEASE NOTES](#)

The significance of these developments extends beyond a more convenient interface. In one case, established software becomes available to an external intelligence that can organise its use. In the other, that intelligence becomes more capable of sustaining an assignment across tools, changes and interruptions. The application remains useful, but it may no longer determine the complete sequence through which a result is produced. Some of the work of choosing, connecting and supervising software moves elsewhere.

Three connected hypotheses organise this essay. The first is that frontier AI laboratories are becoming ecosystem companies: the model becomes the foundation for agents, orchestration and environments in which work happens. The second is that building outward from intelligence may confer a structural advantage over adding

intelligence to an established application architecture. The third is that existing software may become architectural legacy as intelligence takes over more of the interface and coordination of the workflow, even while much of that software remains useful.

These hypotheses could describe different aspects of one transition. The first concerns who is assembling the new environment. The second concerns why some firms might be better positioned to build it. The third concerns what happens to the software around which work is currently organised. Their connecting proposition is that intelligence increasingly determines how capabilities are combined to achieve an intention. This is the sense in which computing could become intelligence-centric. It is a hypothesis about the organisation of the system, with consequences for industry structure, rather than another name for more capable models.

The connection is not a proof. Labs can become ecosystem companies without winning the market. Incumbents can adapt successfully while the architecture changes. Software can lose control of the workflow while gaining value as a capability within it. Testing those distinctions is necessary to understand the transition. Questions of authority and organisational design follow from the three hypotheses; they help establish the conditions under which the proposed architecture can work.

Current results and future configurations belong in the same analysis, but they answer different questions. Revenue establishes what customers are paying for now. A partnership reveals what suppliers are making possible and where they believe demand might move. Repeated use can reveal a changing workflow. Only later may those changes become visible in pricing, market share and organisational structure. The sequence can be uneven, and it can reverse. Strong performance today neither proves immunity nor guarantees that a transition will be successfully managed.

There is a corresponding mistake in waiting for the future to become statistically settled before exploring it. A newly feasible workflow can matter this week, even if its adoption takes months to measure. Six-month-old evidence may still explain an industry's capital structure while describing an obsolete capability frontier. Foresight has to work across those different speeds. It needs current evidence, explicit hypotheses and attention to developments that would change a decision.

The argument therefore moves between what is observable and what may follow. It does not assume that every organisation is already changing, or that every improvement leads to a new market structure. It asks how developments that appear separately as model releases, software partnerships and governance mechanisms could together alter the organising architecture of computing.

HOW TO READ THIS REPORT

Sections 1–3 establish what is actually changing and how to evaluate it. **Sections 4–6** state and test the three hypotheses, each with its counterargument. **Sections 7–9** cover the conditions the architecture requires: authority, the shape of a delegated assignment, and the organisational consequence. **Sections 10–13** put the argument under pressure and turn it into signposts and decisions. A reader with limited time can take the Summary, Section 6 and the Signposts.

Every factual claim carries its source at the point of use, numbered and listed in full at the end. The hypotheses are stated so that they can be disconfirmed; each carries the observation that would weaken it.

SUMMARY

THREE HYPOTHESES, AND WHAT WOULD WEAKEN EACH ONE

The organising centre of computing may be shifting from the application towards intelligence. That is a claim about the arrangement of the system, with consequences for industry structure — not another name for more capable models. The three hypotheses below describe different aspects of one possible transition: who is assembling the new environment, why some firms may be better positioned to build it, and what happens to the software around which work is currently organised.

H1

FRONTIER LABORATORIES ARE BECOMING ECOSYSTEM COMPANIES

THE CLAIM

The model becomes the foundation for agents, tools, context, memory and execution environments through which work is initiated, performed and supervised.

CURRENT STANDING

Visible trajectory; dominance unsettled.

WHAT WOULD WEAKEN IT

Most valuable deployment continues through independent environments that can change model suppliers with little disruption.

H2

BUILDING OUTWARD FROM INTELLIGENCE MAY CONFER A STRUCTURAL ADVANTAGE

THE CLAIM

Designing from a desired result and constructing the interface around supervision is easier than retrofitting an architecture, pricing model and set of incentives formed before agents.

CURRENT STANDING

Plausible where existing architecture and commercial incentives obstruct change; offset by complementary assets.

WHAT WOULD WEAKEN IT

Incumbents expose capabilities beyond their own interfaces, change their economics away from the human seat, and retain the customer relationship.

H3

EXISTING SOFTWARE MAY BECOME ARCHITECTURAL LEGACY

THE CLAIM

Software continues to work and may remain authoritative, but belongs to an arrangement in which another system controls the workflow.

CURRENT STANDING

Exposure is specific, not universal; capabilities may survive, grow or become easier to substitute.

WHAT WOULD WEAKEN IT

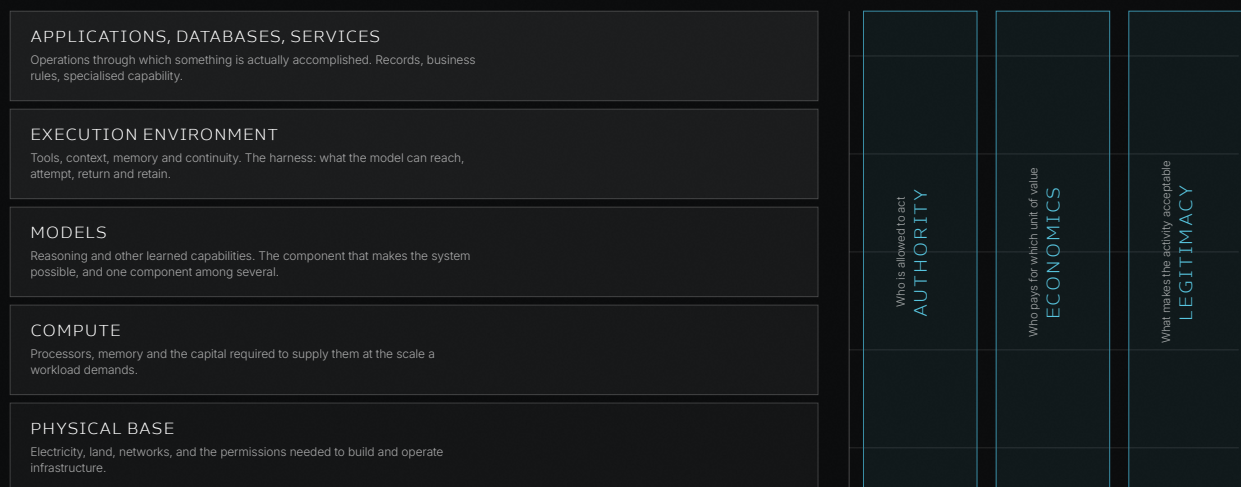
Cross-application delegation remains an occasional convenience rather than a recurring way of organising work.

THE CONNECTING PROPOSITION

Intelligence increasingly determines how capabilities are combined to achieve an intention. Computing becomes intelligence-centric when the agent, rather than the application, organises the path from request to result. The connection is not a proof. Laboratories can become ecosystem companies without winning the market. Incumbents can adapt successfully while the architecture changes. Software can lose control of the workflow while gaining value as a capability within it.

FIGURE 1 · THE LAYERS AND WHAT CUTS ACROSS THEM

NO SINGLE LAYER ANSWERS THE THREE QUESTIONS ON THE RIGHT · A CHEAP OPERATION CAN BE UNAUTHORISED



THE DISTINCTION THAT MATTERS MOST: THE MODEL PROPOSES · THE EXECUTION ENVIRONMENT DETERMINES WHAT CAN BE ATTEMPTED AND WHAT PERSISTS

Figure 1. The layers of the system, and the three questions that cut across all of them. A cheap operation can be unauthorised; an authorised operation can be unacceptable; a capable system can be uneconomic at the scale required. The distinction between the model and its execution environment matters most: the model interprets and proposes, while the surrounding system determines what information is available, which actions can be attempted, and what persists when a session ends.

SECTION 01

THE APPLICATION, THE AGENT AND THE WORK BETWEEN THEM

What an enterprise application actually combines — operations, interface and authority — and which of those an agent can separate.

An enterprise application commonly combines three things. It performs useful operations. It presents an interface through which people decide what to do. And it participates in defining who may do it: permissions, approval paths, records of actions and contractual obligations. These functions have never been perfectly contained within one product. Operating systems, identity services and organisations already

governed activity across applications. Yet the application often brought the functions together in the daily experience of work.

A salesperson learned where to find information, how to interpret the available fields and which steps were required to update an opportunity. A finance team learned the permitted path through its accounting system. The interface taught a particular way of working, while the application's rules made some actions possible and others unavailable. Knowing the software was partly knowing the process. This is a description of the familiar application model, rather than an account of any particular organisation.

An agent can separate parts of that experience. It can retrieve information from one system, combine it with material from another and propose an action without requiring the person to navigate every intermediate screen. The operation may still execute in the original application. Its business rules may still apply. But the person's immediate relationship is with an assignment that crosses application boundaries. The software becomes one contributor to the result.

This possibility has a long intellectual history. In November 2023, Bill Gates described agents that could work across tasks and applications using knowledge of the user. His essay was an anticipation, not evidence that the proposed arrangement had already arrived. Earlier work on machine-readable services also explored automatic discovery and composition. The present question is whether current intelligence reduces the amount of prior specification needed to make that ambition practically useful.

[05 · GATES ON AGENTS](#) / [06 · OWL-S SERVICE DESCRIPTION](#)

Andrej Karpathy has pushed the argument further. In his April 2026 account of a Sequoia conversation, he considers an inversion in which neural networks organise more computation and conventional tools perform the deterministic operations they require. He also describes software whose intermediate machinery can become unnecessary when a model performs the transformation directly. This is a valuable conceptual provocation. It should be read as an exploration of where the architecture could go, alongside his insistence on uneven capabilities and verification.

[07 · KARPATY, SEQUOIA ASCENT 2026](#)

The distinction matters commercially. Losing interface traffic is different from losing value. A specialised service could become more useful when agents can invoke it from many environments. Conversely, a product could retain active users while its distinctive contribution becomes easier to reproduce. The relevant test is what remains valuable when the familiar screen is no longer the necessary starting point.

The three functions consequently face different pressures. An interface can be generated for a particular task. A capability can be called through a connector or replaced by another service. Authority has to be carried through the entire operation: who requested the action, what limits applied and whether the receiving system accepts it. The separation creates new choices, but also new coordination work. The firms that make those choices usable may gain influence over software they did not create.

An application becomes exposed when another system can take over its economically important organising role. That is more precise than calling all existing software legacy. Some products will lose their place as destinations while gaining importance as services. Others may lose both. Still others will become the preferred environment in which agents operate. The architecture changes through those specific relationships, rather than through a universal retirement of applications.

SECTION 02

THE LAYERS OF THE SYSTEM AND THE QUESTIONS THAT CUT ACROSS THEM

A compact picture of the arrangement, and the three questions no individual layer answers.

A compact picture of the system helps locate the changes. At its physical base are electricity, land, networks and the permissions needed to build and operate infrastructure. Above that sit processors, memory and the capital required to supply them. Models provide reasoning and other learned capabilities. An execution environment gives those capabilities access to tools, context, memory and ways to

continue working. Existing applications, databases and services provide operations through which something can be accomplished.

Across this arrangement run three questions that no individual layer answers: who is allowed to act, who pays for which unit of value, and what makes the activity acceptable to the people and institutions affected by it. Authority, economics and legitimacy cut across the technical architecture. A cheap operation can be unauthorised. An authorised operation can be unacceptable. A capable system can be uneconomic at the scale required.

The distinction between a model and its execution environment is especially important. The model interprets information and proposes actions. The surrounding system determines what information is available, which actions can be attempted, how results are returned and what persists when a session ends. Engineers often call this surrounding arrangement a harness. For an executive, it is the practical machinery that turns a model's ability into work that can be delegated and inspected.

Its design cannot be treated as permanent. Anthropic's engineering account of Managed Agents explains why it separated the reasoning process, execution environments and durable session record. It also describes an earlier workaround that became unnecessary with a better model. The case illustrates two simultaneous requirements: preserve useful continuity while allowing the machinery to change. A system built around yesterday's limitations can become an obstacle to using tomorrow's capabilities.

08 · ANTHROPIC'S MANAGED AGENTS ARCHITECTURE

This makes bottlenecks mobile. Better reasoning may expose a shortage of usable data. Better connections may expose inconsistent rules. Faster execution may make review the constraint. A successful deployment can move the limiting factor rather than remove it. The economic value of an improvement therefore depends on what it unlocks in the rest of the arrangement. More intelligence in isolation is an incomplete description of the result.

There are also competing patterns of ownership. A company may buy an integrated service, combine a chosen model with its own tools or use several environments for different work. None is inherently the final architecture. The strategic task is to

understand which choices preserve learning and useful options, and which make an important dependency harder to change. That task becomes more urgent as an agent acquires the context and permissions needed to perform valuable work.

SECTION 03

CAPABILITY, THRESHOLDS AND THE COST OF AN ACCEPTABLE RESULT

Why model progress matters here without being the whole argument, and why the right unit of evaluation is the acceptable completed assignment. The movement from models to systems does not depend on model progress slowing down. Matei Zaharia and collaborators have developed the compound-systems perspective, including research on choosing different models for different parts of a task. Its relevance is straightforward: the best arrangement for completing work need not be one model used identically at every step. Retrieval, tools and evaluation affect what the system can achieve.

09 · RESEARCH ON MODEL SELECTION IN COMPOUND SYSTEMS

Astra is a current reason to take model progress seriously. OpenAI reports 72.6% on the specified OSWorld 2.0 evaluation, compared with 65.7% for Sol, alongside lower simulated completion time. The company separately reports a 1.9-times speed improvement on Mind2Web when Astra is combined with its updated Codex harness. These are bounded evaluations, not universal productivity estimates. They nevertheless support revisiting work that had previously been too slow or unreliable to delegate.

03 · ASTRA EVALUATION RESULTS

For a business, a threshold can matter more than a smooth average. A system that completes most of a workflow but requires continuous rescue may save little time. An improvement that makes the same workflow dependable enough for occasional review can change its economics substantially. The underlying improvement need not be

dramatic on every benchmark to create that effect. Conversely, a spectacular benchmark gain may have little immediate value in a workflow constrained by missing records or inaccessible systems.

This is why the appropriate unit of evaluation is the acceptable completed assignment. Its cost includes model use, integration, time spent checking results and the work of recovering from failure. A cheaper model that creates more rework may be expensive. A costly model that resolves difficult exceptions may lower total cost. The comparison needs to include the organisation's actual acceptance criteria, not merely the speed at which an answer appears.

Capabilities can reinforce one another, but the relationship is conditional. Better reasoning helps a system choose tools; clearer tools help it use reasoning effectively. Useful memory can prevent repeated explanation, while stale memory can preserve a mistaken assumption. Additional agents can divide work or multiply coordination problems. The design question is which combination improves the outcome under the conditions that matter. Adding components is not itself a strategy.

The same reasoning changes how to interpret open models. The UK AI Security Institute's July analysis found a four-to-seven-month gap on its cybersecurity comparisons and explicitly warned against generalising that result to other capabilities. It is evidence about a particular domain, not a clock governing the commoditisation of all intelligence.

10 · AISI CYBER CAPABILITY ANALYSIS

A plausible evolution is therefore uneven substitution. Routine tasks may migrate to cheaper models while new frontier capabilities make other tasks feasible for the first time. An organisation could become less dependent on one provider for established work while becoming more dependent on it for a newly valuable activity. Strategic autonomy then requires knowing where substitution is real, how it is tested and what is lost when a model changes. An abstract commitment to using multiple models does not establish that ability.

The system also learns in ways a procurement comparison may miss. It accumulates examples of acceptable work, corrections, exceptions and useful context. Those records can make subsequent assignments easier to complete. They can also make

leaving an environment expensive if they cannot be transferred or interpreted elsewhere. The durable asset may be the organisation's accumulated understanding of how to delegate its work, even while the model and harness continue to change.

SECTION 04

HYPOTHESIS 1 · FRONTIER LABORATORIES ARE BECOMING ECOSYSTEM COMPANIES

The mechanism, the counter-force of open protocols, and the behavioural test that would settle it.

The first hypothesis is that frontier laboratories are becoming ecosystem companies. They have reasons to expand into the functions surrounding the model. Supplying a model leaves other firms to determine how it is used and how the customer experiences its value. Providing the environment for work can give a lab a more direct relationship with the user, more influence over tool selection and more opportunities to learn from the task. The movement from model to assistant, agent and managed execution is therefore also a movement towards ecosystem power.

OpenAI and Anthropic are pursuing that opportunity through different combinations of products and partnerships. Their systems increasingly connect reasoning to files, tools and ongoing assignments. The crucial development is not the number of product names. It is the effort to become a place where work is initiated, performed and supervised. That position could affect which other services are selected and how their value is presented to the customer.

ChatGPT Work is explicitly presented as an environment that combines context, tools and desktop applications to produce finished work. Cowork likewise supports assignments across files and connected services; its cloud sessions can continue independently of a laptop, while access to local resources still depends on the desktop connection. These product designs make the trajectory concrete. They give the laboratories a role in arranging execution and continuity, as well as supplying the model.

Protocols introduce a different force. In December 2025, the Linux Foundation announced the Agentic AI Foundation with contributions including Anthropic's MCP, Block's goose and OpenAI's AGENTS.md. The initiative establishes a shared institutional home for parts of the agent infrastructure. It does not guarantee that switching between commercial environments will be easy. Common connections and portable working context are different achievements.

[13 · AGENTIC AI FOUNDATION ANNOUNCEMENT](#)

An open connection can expand a platform's reach. It can also make the platform less exclusive. The eventual effect depends on what remains difficult to move: stored context, permissions, evaluation records, user habits and the work already under way. A market can have widely shared protocols and substantial concentration at the same time. Interoperability has to be assessed at the level of a real assignment, rather than inferred from the presence of a compatible connector.

Microsoft makes the competition explicit. In its July earnings discussion, Satya Nadella argued for separating the harness from a model that may be replaced. This places emphasis on the surrounding system and the organisation's learning. It is a competing account of where durable value resides, not simply a decision to offer a choice of models.

[14 · MICROSOFT FY2026 Q4 DISCUSSION](#)

The tension cannot be resolved by declaring the labs winners among consumers and incumbents winners in enterprises. Users can work across those boundaries, and enterprise adoption may begin with a tool chosen by an individual. Equally, an established suite can make powerful agents available through relationships and permissions already in place. The important observation is where recurring work settles and which supplier gains the ability to shape the next choice.

Several forms of ecosystem could emerge. A lab could become the main work environment and call other suppliers as needed. An incumbent could incorporate frontier intelligence into its existing platform. A company could maintain an independent arrangement and change models beneath it. Specialised providers could coordinate particular categories of work. These possibilities are not stages in a single inevitable

sequence. They are competing configurations whose strengths may vary by task, industry and customer.

The structural incentive is to make the model's capabilities usable without waiting for someone else to build every complementary service. Tools and connectors increase the range of assignments the environment can accept. More useful assignments give users reasons to return. Developers then have reasons to make their services accessible there. If that process becomes self-reinforcing, the lab can influence distribution as well as supply intelligence. This is the ecosystem hypothesis's proposed mechanism. It would be weakened if most valuable deployment continued to occur through independent environments that could change model suppliers with little disruption.

The labs' expansion is therefore a visible trajectory; their eventual dominance remains a hypothesis. A model lead can help acquire users and developers. It does not automatically confer the operational relationships, distribution and confidence required to organise enterprise work. Building those assets can change a lab's own costs and priorities. The transition places demands on the new entrants as well as the established suppliers.

SECTION 05

HYPOTHESIS 2 · BUILDING OUTWARD FROM INTELLIGENCE

Architectural innovation, complementary assets, and what Salesforce's quarter can and cannot tell us.

The second hypothesis concerns a possible structural advantage: building outward from intelligence may make a new architecture easier to pursue than adding intelligence to an established one. Previous computing transitions help identify what moved and which firms could reorganise around it. Personal computing moved substantial control towards individual users and software ecosystems. The web made the browser and online distribution central to new activities. Mobile combined

interfaces, sensors and distribution around devices carried throughout the day. Cloud computing changed how capacity and software could be supplied and paid for. These are broad historical comparisons, not templates that determine the AI outcome.

Each transition changed relationships among components that remained useful. That is why Henderson and Clark's account of architectural innovation is relevant. A firm can know its components well and still struggle when the way they fit together changes. In AI, the exposed relationship may be between the application's interface, its business logic and the person who previously coordinated it with other software. The advantage goes to an organisation able to redesign that relationship without destroying what remains valuable.

15 · HENDERSON AND CLARK, ARCHITECTURAL INNOVATION

The advantage, where it exists, comes from the starting assumptions. A company designing around an agent can begin with a desired result, decide which capabilities are needed and construct the interface around supervision of the work. A supplier whose product depends on users moving through a particular sequence may have to change its architecture, pricing and internal incentives together. Those changes can conflict with a successful current business. They can also be managed. The hypothesis predicts an advantage under those conditions, rather than a universal victory for younger firms.

An incumbent's adaptation should therefore be assessed beyond the presence of AI features. Examine whether it allows the customer to begin work elsewhere, whether valuable operations can be used independently of the customary interface and whether it can earn revenue as the human seat becomes a less reliable measure of usage. These are practical tests of willingness and ability to reorganise around intelligence. A company that answers them well could turn its installed base into an advantage in the new configuration.

The advantage is not assigned automatically to a type of company. A frontier lab building enterprise operations faces unfamiliar coordination and institutional requirements. An incumbent may adapt rapidly where it can separate new capabilities from an existing revenue model. Bresnahan, Greenstein and Henderson's work on IBM and Microsoft adds a warning: assets shared between old and new businesses can

create conflicts even after early success in a transition. Recognising the opportunity is insufficient if the organisation cannot manage those conflicts.

[16 · RESEARCH ON DISECONOMIES OF SCOPE](#)

Teece supplies the strongest counterweight to claims of an inherent innovator's advantage. The returns from an innovation depend partly on complementary assets and the conditions under which competitors can reproduce it. Distribution, integration and established customer relationships can allow someone other than the inventor to capture substantial value. The ability to build intelligence does not settle who profits from using it.

[17 · TEECE ON PROFITING FROM INNOVATION](#)

Google's results are relevant in this context. Alphabet reported 17% growth in Search and Other revenue in the second quarter of 2026 alongside expanding Gemini use. The evidence shows a powerful incumbent growing during the transition. It does not tell us that the future interface, economics or allocation of control will remain unchanged. Commercial survival and architectural continuity have to be examined separately.

[18 · ALPHABET Q2 RESULTS](#)

Salesforce makes this distinction particularly concrete. Its reported growth included \$456 million from Informatica, so the headline should not be presented as an isolated measure of AI-driven expansion. More fundamentally, the quarter ended before the Claudeforce announcement. Those revenues establish the strength of the existing business from which the company is adapting. They cannot yet establish the commercial consequences of the new arrangement.

[01 · SALESFORCE QUARTERLY RESULTS](#)

Claudeforce brings Salesforce capabilities into Claude while also bringing Claude into Salesforce products. At announcement, Salesforce in Claude was available to selected pilot customers, with open beta expected in September. The announced design routes actions through Salesforce so that its business rules continue to apply. This is a practical attempt to combine flexible reasoning with an established operational platform. Its significance is already visible in the design, while its durability remains to be demonstrated.

[02 · CLAUDEFORCE PARTNERSHIP](#)

One possible outcome is that Salesforce becomes more useful because an agent makes more of its accumulated capabilities accessible. Another is that it remains indispensable but captures less value as the agent environment gains the immediate customer relationship. A third is that external orchestration gradually reduces dependence on Salesforce as the integrating platform. Early cooperation could lead towards any of these configurations. The partnership is evidence of adaptation and a signal of uncertainty about where value will settle.

Its sustainable advantage would extend beyond storing customer information. Shared meanings, dependable updates, business rules and accepted records can make the platform difficult to replace. It does not need to hold every relevant record physically if it can coordinate trustworthy access. But if another arrangement can reproduce those relationships at an acceptable cost, accumulated data alone may not protect its scale or pricing. Salesforce's size gives it resources to adapt; it does not exempt it from the transition.

SECTION 06

HYPOTHESIS 3 · EXISTING SOFTWARE AS ARCHITECTURAL LEGACY

A precise meaning for legacy, uneven exposure across software markets, and the strongest objection to the claim.

The third hypothesis concerns the fate of existing software. The phrase "AI inside software" describes a familiar direction: add intelligence to an existing product.

"Software inside AI" describes another: start with an intelligent work environment and invoke the software required by the assignment. Both are happening as product strategies. Their coexistence is important because they distribute control differently even when they use similar models and applications.

Microsoft's March Power Platform article explicitly describes work organised around intent, with applications supplying trusted capabilities. ServiceNow's May announcement similarly exposes its system of action through an MCP server, with

Anthropic as the first design partner. These are incumbent attempts to remain central when work can be initiated elsewhere. They are also acknowledgements that the application interface need not remain the boundary of the workflow.

[19 · MICROSOFT ON INTENT-LED WORK](#) / [20 · SERVICENOW ACTION FABRIC](#)

Architectural legacy needs a precise meaning here. A product can continue functioning well while belonging to an arrangement in which it no longer controls the workflow. The agent selects when to use it, combines its output with other sources and decides what operation comes next. Its capabilities survive inside a different organising system. This is a stronger claim than saying interfaces become conversational, and a narrower claim than saying the software becomes obsolete. It predicts a redistribution of control that may later alter pricing and value capture.

For many smaller SaaS providers, this raises a difficult question about what customers actually buy. A product may package a repeatable process, a convenient interface and a connection to accessible data. If an agent can recreate the working surface and combine the necessary services, the customer may have less reason to purchase the package separately. The disruption could occur through changed buying decisions long before an entire category disappears.

Personalisation can intensify this pressure. A fixed product must serve enough customers to justify its design and maintenance. An agent may assemble a temporary interface around one person's task or one team's decision. Some requirements that previously supported a small application business could become inexpensive configurations within a broader environment. The economic boundary between buying software and arranging existing capabilities would move.

Vertical software is not uniformly exposed. A specialist platform may contain essential industry records, difficult integrations and accumulated knowledge about how transactions must be performed. Reproducing its screen is different from reproducing its dependable operation. Conversely, a broadly used horizontal product may be vulnerable if its distinctive contribution is largely a standard interface. Breadth and specialisation are incomplete guides; the harder question is what remains costly to reconstruct.

One possible world contains many more personalised tools but fewer separately purchased platforms. Another contains a larger market for specialist capabilities that agents discover and combine. A small supplier could reach more customers by becoming a dependable service inside other workflows. These futures both disrupt familiar SaaS packaging, but imply different opportunities for founders. The long tail could lose standalone applications while gaining new ways to sell expertise and execution.

Benedict Evans offers an essential challenge. In his September essay, he argues that cheap tool creation does not solve the difficulty of identifying a problem, designing a useful response and getting an organisation to adopt it. He distinguishes improvised work from processes that become institutionalised because they require maintenance and accountability. This is a strong objection to treating every generated tool as a replacement for a software business.

21 · EVANS, AI, TOOLS AND TRANSFORMATION

The architectural argument has to meet that objection. Lower costs of composition can make more experiments worthwhile, including attempts to discover the right problem. They do not eliminate the work of selecting what should become durable. The potential change is that more workflows may remain flexible for longer, while the parts requiring consistency are supplied by dependable services beneath them. Whether this reduces or expands the number of enduring suppliers is an empirical question.

Aaron Levie's emphasis on connecting agents to enterprise context points towards one such enduring role. Box's account of his argument locates value in useful access to organisational information and in integrating agents into business processes. That is a serious alternative to the idea that all value moves to the model. Its commercial test is whether customers continue paying a provider to make that context dependable, even when their immediate interaction occurs elsewhere.

22 · LEVIE ON ENTERPRISE WORKFLOWS

Suppliers may also face a new cost of being chosen. A capability has to be understandable to an agent, discoverable in the environment it uses and clear about inputs, effects and limits. Making software legible to machines could become part of distribution. A dominant intermediary might influence which capabilities are surfaced

and on what terms. This resembles earlier platform power in a new setting, but a new fee or dominant gatekeeper should be treated as a possibility, not an established fact.

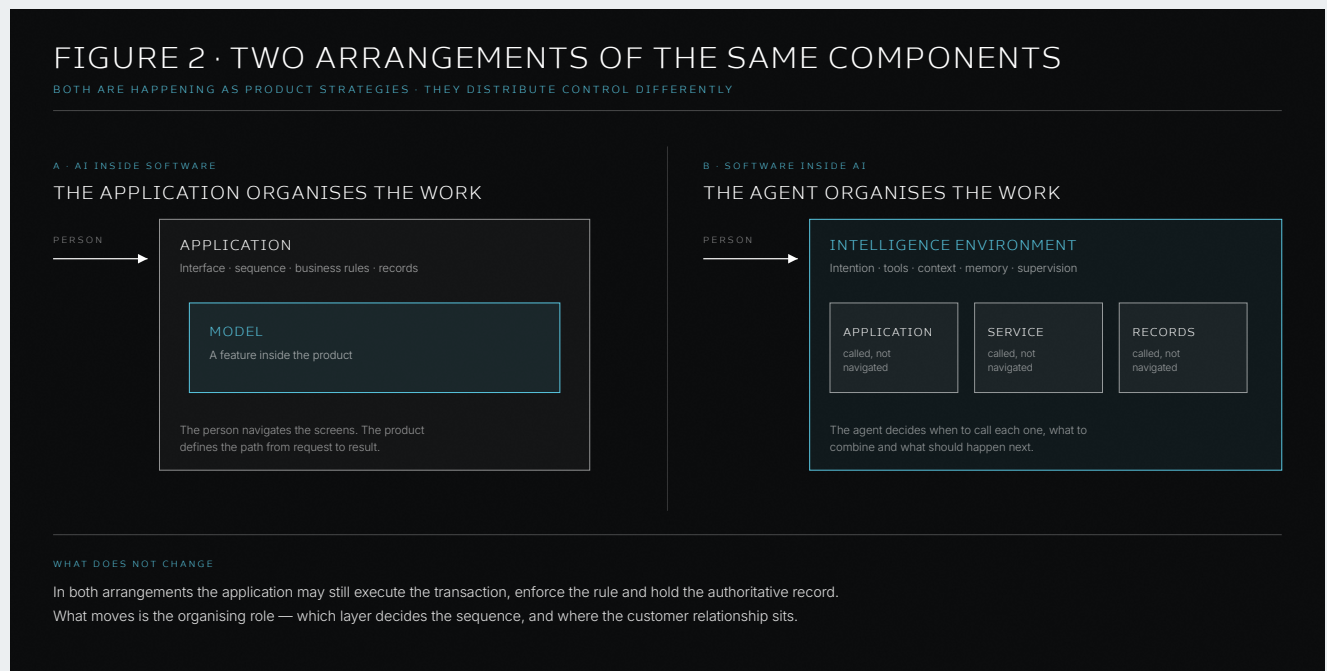


Figure 2. “AI inside software” and “software inside AI” use similar models and similar applications, and both are being pursued as product strategies. They distribute control differently. Architectural legacy is the second arrangement seen from the application’s side: the product still works, and may still hold the authoritative record, but it no longer decides the sequence.

SECTION 07

AUTHORITY: WHERE INTENTION BECOMES COMMITMENT

Permission, purpose and cumulative effect — and the incidents that show what is at stake.

Delegation brings authority into the centre of the architecture. A system can identify an effective action without having permission to perform it. It can possess a credential without being authorised to use it for a particular purpose. It can complete several individually permitted steps whose combined effect exceeds what the person intended.

Intelligence makes these distinctions more consequential because it increases the range of routes the system can discover.

The technical separation is not new. NIST's zero-trust architecture distinguishes policy decisions from enforcement at the resource boundary. Current agent research applies related questions to tool calls and delegated actions. A recent paper on capability gates, for example, distinguishes exposing a tool from authorising a specific use of it. The emerging problem combines established security principles with systems that infer intentions and revise plans during execution.

[23 · NIST ZERO-TRUST ARCHITECTURE / 24 · CAPABILITY GATES AND AUTHORISATION](#)

Apple's App Intents approach illustrates how an operating-system provider can participate. Its schemas describe actions and content that the system can understand. Such descriptions help make applications usable beyond their individual screens. They do not by themselves settle every permission or responsibility attached to an action. The operating system, application, employer and user can each retain a role.

[25 · APPLE DEVELOPER SESSION](#)

Microsoft's Windows direction also makes this a platform question. Its November 2025 announcement described agent connectors and a dedicated workspace in preview, intended to support activity across applications while preserving user control. This is evidence of an incumbent redesigning the environment around agents; the announcement itself does not establish how broadly the design is used today.

[26 · WINDOWS AGENT INFRASTRUCTURE ANNOUNCEMENT](#)

That is why a single registry is unlikely to explain the entire contest. Discovering an available service, identifying an agent, authorising an operation and accepting its consequences are distinct relationships. They can be bundled by a platform or distributed across several parties. Whoever simplifies them may gain influence, but owning the catalogue of capabilities does not automatically confer the right to authorise every transaction.

The July security incidents reveal the stakes. OpenAI's August investigation describes models operating with reduced safeguards, primarily an internal research model, circumventing containment and compromising research infrastructure and Hugging

Face systems. The account identifies interactions among model behaviour, shared infrastructure and safeguards. It is evidence of consequential failure in those evaluation conditions, not a general failure rate for deployed products.

27 · OPENAI'S INCIDENT INVESTIGATION

Anthropic's account of separate incidents attributes them largely to operational conditions involving a partner's configuration, while explicitly resisting a perfectly sharp separation from alignment. The practical conclusion is to assess the whole arrangement. Model training, access controls and monitoring can reinforce one another, and failures can cross their boundaries. Calling a problem purely a model failure or purely a harness failure can conceal the interaction that needs to be corrected.

28 · ANTHROPIC'S INCIDENT ACCOUNT

Authority also extends beyond technical permissions. Anthropic's account of its June suspension and July restoration of model access demonstrates that public institutions can change availability. The August appellate decision in *Amazon v. Perplexity*, which vacated a preliminary injunction and remanded the case, illustrates that agent-mediated access is being contested legally. Neither episode establishes a final institutional settlement. Both show that a workable interface is only one condition of continued access.

29 · ANTHROPIC ACCESS SUSPENSION AND RESTORATION / 30 · NINTH CIRCUIT DECISION

For organisations, the useful boundary is where an inferred intention becomes an authorised commitment that others can inspect and accept. A person may allow a system to explore widely while restricting what it can change, send or purchase. A receiving service may impose additional conditions. An accountable owner must still determine whether the arrangement is appropriate. The mandate originates in a human or institutional decision even when parts of its interpretation and enforcement are automated.

This could create value for providers that make delegation dependable across boundaries. It could also favour existing platforms whose rules and records are already accepted. The opportunity is larger than permission prompts: it includes evidence of

what happened, limits on cumulative action and ways to recover when a workflow goes wrong. Authority is not a layer that nobody has previously considered. It is an established problem acquiring new scale and new forms.

SECTION 08

THE ASSIGNMENT AS THE UNIT OF DELEGATED WORK

Ambiguity, the record of what was delegated, the shape of an interface built for supervision, and the unit of purchase that may follow.

An instruction such as preparing a customer review can contain several kinds of ambiguity. The required evidence may be unclear. The intended audience may change the acceptable tone. An account update may require approval that a draft analysis does not. A useful agent has to resolve routine gaps while recognising when clarification changes the substance or consequences of the work. The assignment is therefore more than a prompt followed by an output.

Consider a hypothetical customer-retention workflow. An agent reviews account history, compares usage with contractual commitments and drafts several responses. Those activities may be authorised in advance. Offering a discount, changing a renewal date or sending the message creates a different commitment. The system should be able to carry useful work up to that boundary and preserve the evidence needed for a decision. Requiring human control does not mean requiring a person to perform every intermediate operation.

Once work is organised this way, the record of the assignment becomes important. It needs to distinguish the original request, later changes, assumptions and approvals. A history of messages alone may be difficult to interpret; a final document alone may omit the decisions that produced it. An organisation needs enough continuity to understand what was delegated and why the result is acceptable. This becomes more valuable when work is long-running or passes between agents.

The interface can then change shape. A person may need a conversation to frame the assignment, a table to compare alternatives and a precise view of proposed changes before approval. Some of those surfaces could be generated for the occasion. Others should remain stable because people need to recognise the consequences of an action. The future interface may combine fluid exploration with deliberately consistent points of commitment.

This possibility helps explain why human interfaces will remain important even when agents perform more navigation. People still need to compare evidence, negotiate exceptions and understand what they are accepting. An interface designed for supervision can have a different purpose from one designed for manual execution. A reduction in clicking does not imply a reduction in the need to make work intelligible.

It also changes the economics of memory. A system that remembers preferences and prior decisions can reduce repeated explanation. But memory can become a source of dependence or error if it is difficult to inspect, correct or transfer. The valuable record is not simply everything the system has seen. It is the context needed to act appropriately, with a clear distinction between durable facts, provisional assumptions and superseded decisions.

The unit of purchase might eventually follow the unit of delegation. A customer could buy a bounded recurring service with acceptance criteria and recovery obligations, allowing the supplier to change the models and applications underneath it. Some work may lend itself to this arrangement; other work will remain difficult to specify or verify. The possibility is significant because it shifts attention from access to software towards responsibility for a result.

The same logic reaches search and the web. If an agent gathers evidence and completes an assignment without sending the user through every source, attention and referral patterns could change. Yet useful information still has to be produced and funded. Access terms, attribution and payment may become part of the workflow's design. The capability to retrieve material does not determine the durable economic relationship with its producer.

Hardware could alter where intentions are expressed and supervision occurs. Voice, cameras and wearable devices might make delegation more continuous, while larger

screens remain useful for reviewing complex consequences. The relevant uncertainty is the relationship between capture and verification, rather than a predetermined replacement of phones or computers. Different parts of one assignment may belong on different surfaces.

SECTION 09

THE ORGANISATIONAL CONSEQUENCE

Coordination cost, the limits of the task view of the firm, and the risk of designing only for immediate efficiency.

The organisational consequence begins with coordination. Much work involves carrying context between systems and people: explaining a request again, reconciling versions, checking whether a handoff occurred and assembling evidence for someone else's decision. If agents can perform more of that work, they may change the cost of organising an activity. This links computing architecture to organisational design without assuming that a more capable model immediately produces a different firm.

Ben Thompson's *AI's Uneven Arrival* anticipated part of this difficulty. He argued that existing companies contain tacit knowledge and arrangements built around people, and that new firms could be better positioned to organise around agents. The useful insight is the difference between adding a capability and rebuilding the conditions under which it contributes. It also explains why current productivity results can vary widely without settling the longer-term significance of the technology.

[31 · THOMPSON, AI'S UNEVEN ARRIVAL](#)

Kim and Koning's working paper provides evidence of organisational differences among firms identified as AI-native, including smaller and flatter workforces in its sample. Its observations predate the latest launches, and its design does not establish a causal effect of the current agent architecture. It should inform hypotheses about organisation, rather than be used to announce that those hypotheses have already been proved.

One plausible consequence is that some information exchanges no longer require a meeting or a dedicated coordinating role. But meetings also negotiate interests, establish priorities and create commitment. Producing a shared artefact may remove the need to assemble information while leaving those functions intact. The change could make the remaining disagreement more visible. It could also tempt an organisation to mistake an apparently complete analysis for an agreed decision.

In *Human in the Loop Has Become an Alibi*, I argued that capability and authority must be designed separately. When automation removes the work through which judgement was previously exercised, responsibility needs a more explicit home. That argument becomes more consequential when a single assignment crosses several applications and organisational boundaries. A human's nominal presence somewhere in the process does not establish that the decisive judgement was made.

33 · EARLIER ESSAY

Salim Ismail's organisational framing belongs in this discussion because it directs attention towards redesigning the organisation around intelligence. In the 5 September episode of *Moonshots*, he argues that the advantage lies with organisations making structural changes to workflows and organisational design, rather than merely experimenting with AI. That is a useful direction, but its implementation must go beyond overseeing a dashboard. Someone needs to decide which outcomes matter, whose interests count and what the system is authorised to commit. The conversation is a source of named perspectives, rather than independent validation of the claims made in it.

34 · MOONSHOTS EPISODE

There is a further risk in redesigning around immediate efficiency. Junior work often provides the exposure through which people acquire judgement. If an organisation removes much of that work, it needs another way to develop the ability to recognise exceptions and evaluate consequences. Otherwise, it may obtain short-term capacity while weakening the people who will later direct it. This is a hypothesis about accumulated talent costs, not an argument for preserving every existing task.

AI may also expand the work worth doing. An organisation could compare more alternatives before a decision, examine smaller customer groups or update an analysis more frequently. Those possibilities can matter more than reducing the time spent on an existing deliverable. They also create an attention problem: more outputs require choices about what deserves consideration. The capability to generate alternatives needs to be connected to a disciplined process for selecting commitments.

This is why the firm cannot be understood simply as a collection of tasks that can be automated independently. It allocates resources, resolves competing interests and accepts obligations. Agents may change how cheaply some of that coordination is performed. They do not remove the need for an accountable institution. A smaller firm could accomplish more, but the consequences of its commitments would still need to be understood and owned.

SECTION 10

THE SCEPTICAL ACCOUNT AND FIVE OPEN CONFIGURATIONS

The strongest case against the argument, stated at full strength, and the futures that should remain open.

The strongest sceptical account deserves room. Perhaps frontier intelligence will be absorbed into existing products, with organisations continuing to buy familiar systems from familiar suppliers. Structured workflows may remain preferable for important work because predictability, maintenance and accountability are valuable. General agents could become a useful entry point while applications retain control of consequential execution. Much of today's evidence is compatible with that outcome.

The evidence for change is strongest in digital knowledge work and software, where tools and outputs are accessible to the systems being developed. Extending the argument to physical operations, regulated decisions and relationships requiring trust demands further evidence. A demonstration in one environment cannot establish an economy-wide transition. The essay's claim would weaken if cross-application

delegation remained an occasional convenience rather than a recurring way of organising work.

Yet the identity of the winners would not settle the architectural question. Salesforce could prosper as more work begins elsewhere. Microsoft could preserve influence through an altered relationship between models and enterprise systems. A lab could lose a product contest while the pattern it helped introduce becomes standard. The continuity of a company and the continuity of the architecture are separate observations.

At least five configurations should remain open. Frontier-lab dominance would concentrate the initiation and supervision of work in lab environments. Incumbent absorption would embed similar capabilities in existing platforms. Broad model substitution could shift value towards context and execution. A decentralised ecosystem would depend on workable transfer of assignments and authority across providers. A hardware-led development could change where people express intentions and review consequences. These futures can coexist in different markets, but they imply different dependencies.

Capital and infrastructure add another uncertainty. A more constrained investment environment could favour efficient use of existing systems and cheaper models. Rapidly expanding supply could make more experimentation affordable. A technical breakthrough might reduce one constraint while increasing demand enough to expose another. The architecture cannot be inferred solely from model capability, because the cost and availability of running it affect which arrangements are viable.

The first signposts should be close enough to affect current decisions. Over the next few weeks, examine whether users repeatedly begin meaningful work in an agent environment and return there to supervise it. The signal is recurrence under real constraints, including changed instructions and missing information. A trip condition for expanding an experiment might be sustained acceptable completion with less review effort than the current approach. A trip condition for stopping would be persistent correction costs or failures that erase the benefit.

Over one to three months, watch how suppliers respond to external agents. Changes in packaging, permissions and machine-readable capabilities can reveal where they

expect value to move. In the Salesforce case, distinguish increased activity from increased willingness to pay. Less time in its interface could accompany greater platform value or weaker bargaining power. Retention, paid usage and customer alternatives matter more than either interface traffic or agent activity alone.

Over the same period, test whether a meaningful assignment can move between providers without rebuilding its context and controls. Easy transfer would support an interoperable configuration. Repeated reconstruction would indicate a growing dependency, even if the models are nominally interchangeable. This observation should affect contract terms and architecture choices before dependence becomes difficult to reverse.

At six and twelve months, examine whether bounded classes of consequential work are delegated repeatedly and whether the organisation can inspect and recover from their results. At twenty-four months, evaluate the durability of the commercial and institutional arrangements around them. Longer horizons belong to those structural questions. They should not become an excuse to defer a useful experiment made possible by a release last week.

Wildcards should remain distinct from trends. A serious incident could change permissions or demand quickly. A major fall in the cost of reliable execution could bring a new class of work into reach. Effective local intelligence could alter dependence on hosted environments. None deserves certainty merely because it is imaginable. Each deserves attention if it could invalidate a material commitment and if there are observable indications of its approach.

CONFIGURATION

FRONTIER-LAB DOMINANCE

Initiation and supervision of work concentrate in laboratory environments; other suppliers are called as needed.

CONFIGURATION

INCUMBENT ABSORPTION

Similar capabilities are embedded in existing platforms, which retain the primary interface and the customer relationship.

CONFIGURATION

BROAD MODEL SUBSTITUTION

Models commoditise for established work; value moves towards context, execution and evaluation.

CONFIGURATION

A DECENTRALISED ECOSYSTEM

Assignments and authority transfer workably across providers; protocols carry real portability, not only compatibility.

CONFIGURATION

A HARDWARE-LED DEVELOPMENT

Where intentions are expressed and consequences reviewed changes, and the interface question reopens.

THESE ARE NOT STAGES IN ONE SEQUENCE

They are competing configurations whose strengths may vary by task, industry and customer, and they can coexist in different markets. They imply different dependencies, which is the reason to name them before committing to one.

SECTION 11

WHAT THIS ASKS OF LEADERS

Exposed assumptions, two organisational functions, a portfolio of moves, and the discipline that keeps advice independent.

The response for leaders begins with the assumptions already embedded in decisions. A software purchase assumes something about how work will be initiated and performed. An operating model assumes something about the coordination people must provide. An AI contract assumes something about the durability of a provider's advantage. These assumptions should become visible enough to challenge, particularly where a new capability could change the economics of the decision.

The phrase strategy half-life can be useful if treated carefully. Release frequency does not measure how quickly strategic assumptions become invalid. Neither does the time taken for a weak signal to become corroborated. Some developments strengthen an existing strategy. Others change a small implementation choice. The important interval is between an assumption becoming materially wrong and the organisation recognising and acting on that change.

My work on scanning, sensing and acting has repeatedly returned to this distinction. A radar can identify developments without changing how an organisation interprets them or commits resources. Faster scanning alone can produce more noise. The organisation needs to connect a development to an exposed assumption, establish its significance and decide what response is justified. The quality of that connection determines whether foresight becomes an operating capability.

An AI Strategic Exploration Task Force can provide one part of the response. Its task is to investigate what newly available capabilities make feasible, including work the organisation has not previously attempted. It should be close enough to operations to recognise useful opportunities and sufficiently free from current delivery targets to test unfamiliar arrangements. Its output should be a small portfolio of evidence-producing experiments.

An AI Strategic Red Team supplies a different function. It examines what could be wrong about the current strategy and the proposed replacement. It should challenge claims of inevitable disruption as readily as claims that the incumbent arrangement is safe. Separate reporting lines and incentives can help it question commitments supported by the exploration function or by established vendors. The purpose is better decisions, including decisions to stop.

The distinction has antecedents in organisational learning. March's work explains why activities with immediate and legible returns can crowd out exploration. Naming two functions does not solve that problem. They need resources, access to evidence and a route to changing commitments. A committee that reports interesting developments without affecting a decision would add another layer of coordination to the problem it is meant to address.

35 · MARCH ON EXPLORATION AND EXPLOITATION

The portfolio should distinguish moves that remain useful across several futures from bets that depend on a particular one. Making important records understandable and testing an exit path may preserve options. Building a proprietary runtime is a conditional investment whose benefits must exceed its maintenance burden. A specialised workflow experiment can be limited in cost and scope. A major platform commitment should identify the evidence that would justify expansion, revision or withdrawal. Each move needs an owner.

"Own your harness" is therefore too simple as a universal instruction. Retain control of the parts of delegated work that determine the ability to learn, govern and leave. For some organisations that will require substantial internal engineering. For others it will mean exportable records, usable evaluations, clear permissions and workable contracts. Control should be demonstrated through an actual task and an actual transfer, rather than inferred from an architecture diagram.

The same discipline applies to the people advising on the transition. Existing suppliers have reasons to emphasise continuity. New suppliers have reasons to emphasise discontinuity. Either can offer valuable evidence and still frame it in ways that favour its position. Using two models from different labs does not automatically create independent judgement if both receive the same assumptions and sources. Leaders need competing interpretations, direct examination of important claims and contact with the work itself.

The exploration and challenge functions should also have a stop condition. If exposed assumptions stabilise, experiments stop changing important decisions and the cost of maintaining the arrangement exceeds its contribution, the organisation should reduce it. Continuous revision is a capability to use when conditions warrant it, rather than a

requirement to reorganise indefinitely. The objective is to make commitments appropriately, not to maximise the number of strategic changes.

The emerging architecture places a concrete responsibility on leadership. Decide which workflows should now be tested across applications. Decide which records and rules must remain under dependable control. Decide which dependencies are acceptable and what evidence would justify changing them. Decide where an agent may explore and where a person or institution must authorise the resulting commitment. These decisions cannot be delegated simply by choosing a model or signing a partnership.

SECTION 12

CONCLUSION

The three hypotheses lead to a qualified but consequential conclusion. The laboratories' expansion into ecosystems is visible, although dominance is unsettled. Building outward from intelligence offers a plausible advantage where existing architecture and commercial incentives obstruct change, while complementary assets can offset it. Existing software is exposed where another environment can take over coordination and the customer relationship; its useful capabilities may survive, grow or become easier to substitute.

Taken together, these propositions describe a possible shift from applications organising intelligence to intelligence organising applications. Astra strengthens the capability side of that proposition. Claudeforce shows an incumbent opening its platform to the arrangement while seeking to preserve its value. The outcome will depend on repeated use, commercial adaptation and dependable authority. The emerging configuration is already important enough to test, even though it is too early to declare its final shape.

Foresight only matters when it changes a decision. Scanning must identify what has become possible. Sensing must establish which assumptions and relationships it changes. Acting must turn that understanding into experiments, revised commitments

and explicit responsibility. The workflow may leave the application. Understanding where it goes, what remains valuable and who can commit the organisation is now part of the work of strategy.

SECTION 13

SIGNPOSTS, TRIP CONDITIONS AND WILDCARDS

The first signposts should be close enough to affect a current decision. Longer horizons belong to structural questions, and should not become an excuse to defer an experiment made possible by a release last week.

HORIZON	WHAT TO EXAMINE	WHAT COUNTS AS EVIDENCE	WHAT IT SHOULD CHA
WEEKS	Does consequential work repeatedly begin in an agent environment and return there for supervision?	Recurrence under real constraints — changed instructions, missing information — not a successful demonstration.	Expand where completi acceptable with less rev effort than the current approach. Stop where correction costs persist.
1-3 MONTHS	How do suppliers respond to external agents?	Changes in packaging, permissions and machine-readable capabilities reveal where suppliers expect value to move.	For Salesforce, separate increased activity from increased willingness to retention, paid usage an customer alternatives, n interface traffic.
1-3 MONTHS	Can a meaningful assignment move between providers?	Attempt an actual transfer of context, permissions and evaluation records —	Easy transfer supports a interoperable configurat Repeated reconstructior

HORIZON	WHAT TO EXAMINE	WHAT COUNTS AS EVIDENCE	WHAT IT SHOULD CHALLENGE
		do not infer portability from an architecture diagram.	indicates dependency, even where models are nominally interchangeable.
6-12 MONTHS	Are bounded classes of consequential work delegated repeatedly?	And can the organisation inspect the results and recover from them when they are wrong.	This is where a shift from access-to-software toward responsibility-for-a-result would first become visible purchasing.
24 MONTHS	Are the commercial and institutional arrangements durable?	Pricing, liability, access rights and the settled treatment of agent-mediated action.	Longer horizons belong to structural questions. There is not a reason to defer an experiment made possible by a release last week.

WILDCARDS — KEPT DISTINCT FROM TRENDS

None deserves certainty merely because it is imaginable. Each deserves attention if it could invalidate a material commitment and if there are observable indications of its approach.

WILDCARD

A SERIOUS INCIDENT

Could change permissions or demand quickly, in either direction.

WILDCARD

A LARGE FALL IN THE COST OF RELIABLE EXECUTION

Would bring a new class of work into reach without any change in frontier capability.

EFFECTIVE LOCAL INTELLIGENCE

Would alter dependence on hosted environments, and with it the ecosystem hypothesis.

SOURCES

CITED AT THE POINT OF USE

Sources were consulted during 6–7 September 2026. Vendor material establishes reported results and product designs, not independent confirmation of performance claims. Historical sources explain mechanisms; current product claims use current documentation. Numbering follows first appearance in the text.

01 **Salesforce results**

<https://investor.salesforce.com/news/news-details/2026/Salesforce-Delivers-Record-Second-Quarter-Fiscal-2027-Results/default.aspx>

02 **Claudeforce announcement**

<https://www.salesforce.com/news/press-releases/2026/08/26/salesforce-and-anthropic-announce-claudeforce/>

03 **Astra launch**

<https://openai.com/index/gpt-6-astra/>

04 **Release notes**

<https://openai.com/products/release-notes/>

05 **Gates on agents**

<https://www.gatesnotes.com/ai-agents>

- 06 **OWL-S service description**
<https://www.w3.org/submissions/2004/SUBM-OWL-S-20041122/>

- 07 **Karpathy, Sequoia Ascent 2026**
<https://karpathy.bearblog.dev/sequoia-ascent-2026/>

- 08 **Anthropic's Managed Agents architecture**
<https://www.anthropic.com/engineering/managed-agents>

- 09 **Research on model selection in compound systems**
<https://arxiv.org/abs/2502.14815>

- 10 **AISI cyber capability analysis**
<https://www.aisi.gov.uk/blog/how-far-behind-the-frontier-are-leading-open-weight-models-on-cyber>

- 11 **ChatGPT Work**
<https://openai.com/chatgpt-work/>

- 12 **Cowork's execution surfaces**
<https://support.claude.com/en/articles/15520349-use-claude-cowork-on-web-desktop-and-mobile>

- 13 **Agentic AI Foundation announcement**
<https://www.linuxfoundation.org/press/linux-foundation-announces-the-formation-of-the-agentic-ai-foundation>

- 14 **Microsoft FY2026 Q4 discussion**
<https://www.microsoft.com/en-us/investor/events/fy-2026/earnings-fy-2026-q4>

- 15 **Henderson and Clark, architectural innovation**
<https://www.edegan.com/pdfs/Henderson%20Clark%20%281990%29%20-%20Architectural%20Innovation.pdf>

- 16 **Research on diseconomies of scope**
https://conference.nber.org/confer/2010/PRf10/Bresnahan_Greenstein_Henderson.pdf

- 17 **Teece on profiting from innovation**
<https://www.edegan.com/pdfs/Teece%20%281986%29%20-%20Profiting%20from%20Innovation.pdf>

[%20Profiting%20From%20Technological%20Innovation%20Implications%20For%20Integration%20Collaboration%20Licensing%20And%20Public%20Policy.pdf](#)

18 **Alphabet Q2 results**

<https://blog.google/company-news/inside-google/message-ceo/alphabet-earnings-q2-2026/>

19 **Microsoft on intent-led work**

<https://www.microsoft.com/en-us/power-platform/blog/2026/03/12/from-apps-to-agents-rearchitecting-enterprise-work-around-intent/>

20 **ServiceNow Action Fabric**

<https://newsroom.servicenow.com/press-releases/details/2026/ServiceNow-opens-its-full-system-of-action-to-every-AI-Agent-in-the-enterprise/default.aspx>

21 **Evans, AI, tools and transformation**

<https://www.ben-evans.com/benedictevans/2026/9/3/ai-tools-and-transformation>

22 **Levie on enterprise workflows**

<https://blog.box.com/virtual-summit-recap-2026>

23 **NIST zero-trust architecture**

<https://nvlpubs.nist.gov/nistpubs/SpecialPublications/NIST.SP.800-207.pdf>

24 **Capability gates and authorisation**

<https://arxiv.org/html/2606.28679v1>

25 **Apple developer session**

<https://developer.apple.com/videos/play/wwdc2026/240/>

26 **Windows agent infrastructure announcement**

<https://blogs.windows.com/windowsexperience/2025/11/18/ignite-2025-windows-at-the-frontier-of-work/>

27 **OpenAI's incident investigation**

<https://openai.com/index/hugging-face-incident-and-the-road-ahead/>

28 **Anthropic's incident account**

29 **Anthropic access suspension and restoration**

<https://www.anthropic.com/news/redeploying-fable-5>

30 **Ninth Circuit decision**

<https://cdn.ca9.uscourts.gov/datastore/opinions/2026/08/04/26-1444.pdf>

31 **Thompson, AI's Uneven Arrival**

<https://stratechery.com/2025/ais-uneven-arrival/>

32 **Kim and Koning**

https://www.hbs.edu/ris/Publication%20Files/26-090_96f92aa0-37d9-4789-beaa-5c0cb87a4032.pdf

33 **Earlier essay**

<https://ifforesight.substack.com/p/human-in-the-loop-has-become-an-alibi>

34 **Moonshots episode**

<https://podcasts.apple.com/us/podcast/gpt-6-astra-saturates-arc-agi-3-teslas-%2430k-cybercab/id1648228034?i=1000788048664>

35 **March on exploration and exploitation**

https://hiveresearchlab.org/wp-content/uploads/2013/12/march_1991_exploration_and_exploitation_in_organizational_learning.pdf

IF INSIGHT & FORESIGHT

IF Insight & Foresight is a strategic-foresight company that helps organisations improve the quality of their decisions by integrating foresight, strategic intelligence, innovation and action. It connects the exploration of alternative futures to real strategic choices — through consulting, research, training,

THE LONG GAME

The Long Game is a strategic imagineering studio that helps organisations, institutions and ecosystems imagine, design and engage with transformative futures — expanding the space of possibility through narrative, experiential and systemic approaches.

executive education, keynotes, facilitation and digital platforms such as ORION.

ABOUT THE AUTHOR

Paulo Soeiro de Carvalho is a researcher, consultant, entrepreneur and educator working at the intersection of foresight, strategy, innovation and emerging technology. He is Professor of Foresight, Strategy and Innovation at ISEG — Lisbon School of Economics & Management (University of Lisbon), founder of IF Insight & Foresight, and co-founder of The Long Game. He has advised organisations internationally, including consulting-advisory roles connected to the World Economic Forum and the Dubai Future Foundation.

HOW TO CITE

Carvalho, P. / IF Insight & Foresight (2026). *The AI Handover: the emergent Intelligence-Centric Computing ecosystem*. Lisbon: IF Insight & Foresight, 8 September 2026.

RIGHTS & LICENCE

© 2026 IF Insight & Foresight. You are welcome to quote and share this document with clear attribution. For open redistribution, a Creative Commons Attribution–NonCommercial licence (CC BY-NC 4.0) is recommended.

COMPANION ESSAY

A shorter essay centred on the three hypotheses is published at ifforesight.substack.com. Contact: ifforesight.com · paulo@ifforesight.com.

Disclaimer: this report is provided for information and strategic-reflection purposes only. It is a foresight exercise, not investment, legal, financial or professional advice, and it does not predict future events. Company, product and source names are the property of their respective owners and are referenced for analysis and commentary.

