

ORGITH · THE AGENTIC TRANSFORMATION SUITE

Velora

A framework for adopting agentic AI

Absorbing agentic AI in months — under disruption, while cognitive work is redistributed between people and machines.

An Orgith framework
Version 1.0 · May 2026

Contents

1. Executive summary

Most organizations are trying to adopt agentic AI using change frameworks built for a different kind of change — gradual, bounded, and tool-shaped. Agentic AI is none of those, and the mismatch is why so many initiatives stall after the pilot. The technology arrives faster than a multi-year program can move, it does not just assist but plans and acts, and it absorbs the cognitive content of people's jobs rather than sitting alongside it.

Velora is a framework designed for exactly those three conditions. It is a continuous operating loop — Surface, Pilot, Evaluate, Evolve, Diffuse — run by small cross-functional cells in cycles measured in weeks, sitting on a foundation of four organizational rails. It does two things the classic frameworks do not:

- **It redesigns the human role** around the work agents take over, instead of deploying a tool and leaving the role untouched.
- **It never assumes a finished end state** — the loop keeps turning as capabilities improve, so there is no “refreeze.”

The result is a method that is fast enough to matter, honest enough to be trusted, and structured enough to scale beyond the first enthusiastic team.

2. Why a new framework is needed

The established change-management frameworks are good at what they were built for. The problem is that all of them assume you can see the destination in advance and that you eventually arrive and stop. Agentic AI breaks that assumption along three axes — and those three axes are the design brief for Velora.

2.1 The three forces

Force 1 — Speed: months, not years. Adoption is already spreading person-to-person, ahead of any policy, and most stalled AI initiatives fail for change-management reasons rather than technical ones — organizations roll out tools without redesigning the work around them. A sequential, year-long program is structurally too slow: by the time it finishes, the capability and the competitors have moved on.

Force 2 — The disruptive nature of agentic. Agents do not merely assist; they plan, decide, and act across systems. That turns adoption into a redistribution of decisions and actions between humans and software, and it introduces a problem no classic framework has a slot for: trust calibration. Over-trust ships errors; under-trust wastes the capability. Governance becomes part of adoption, not an afterthought.

Force 3 — Replacing cognitive tasks. The defining shift of this wave is task reallocation: jobs are bundles of tasks, and AI is absorbing specific cognitive tasks while humans move toward supervision, judgment, and accountability. That is an identity change, not a skills gap — which is why the emotional and role dimensions cannot be optional.

2.2 Why the classic frameworks fall short

Each remains useful as a lens or a source of parts (Section 11), but none was built for a change that is this fast, this agentic, and this personal.

Framework	Built to manage	Where it strains under agentic AI
ADKAR	Individual adoption of a defined change, in sequence	Treats change as a skills gap; too slow and linear when the role itself is being redrawn
Kotter (8 steps)	Large top-down transformation toward a fixed vision	Designed for multi-year change with a known destination; the tempo is wrong
Lewin (unfreeze—change—refreeze)	Moving to a new stable state	There is no “refreeze” — capability keeps moving every quarter
Bridges (transition)	The emotional journey through change	Strong on identity (which matters here) but silent on speed and deployment
McKinsey 7-S	Keeping organizational elements aligned	A diagnostic lens, not an operating method

The shared, fatal assumption: the future state is known in advance. With agentic AI you discover what agents can reliably do by trying, and the answer changes as the models improve.

3. Velora at a glance

Velora has three parts that work together: a loop, a foundation, and a thread.

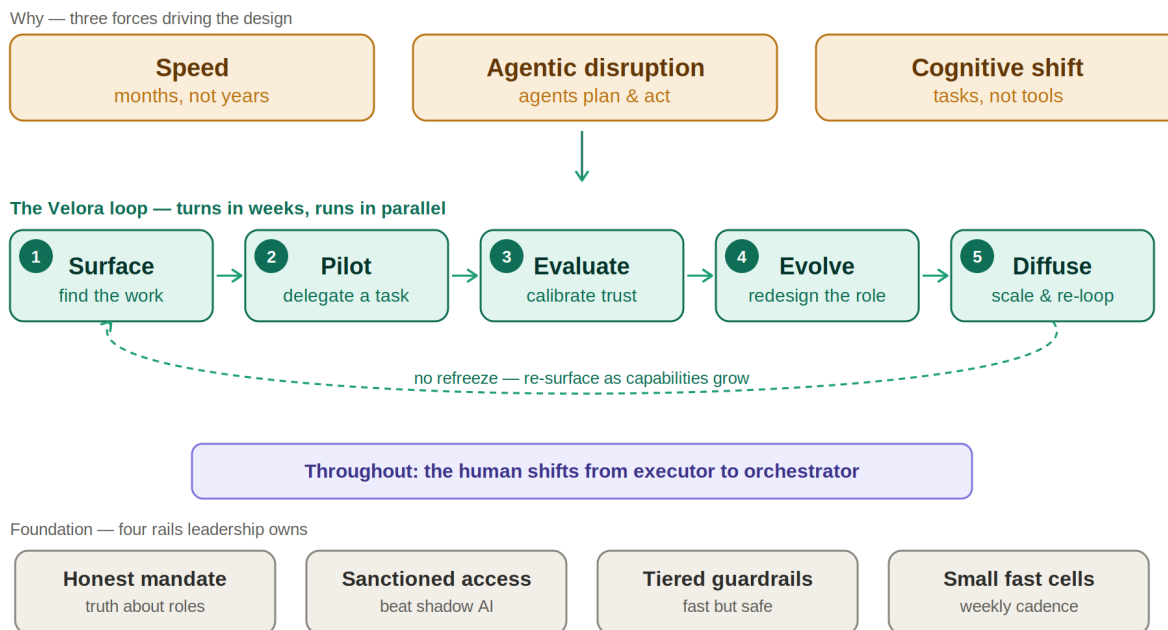


Figure 1. The Orgith Velora framework at a glance.

- **The loop** — a five-stage operating cycle (Surface → Pilot → Evaluate → Evolve → Diffuse) that a small cell turns in weeks, and that the organization runs in many places at once.
- **The foundation** — four rails that leadership owns, which let the loops run fast and safely.
- **The thread** — the human transition from executor to orchestrator, running through every turn of the loop.

Five principles distinguish it from a faster version of the classics:

- **Continuous, not bounded** — the loop keeps turning; there is no end-state to refreeze into.
- **Parallel, not sequential** — many small cells run loops simultaneously.
- **Discovery-driven** — you learn what agents can reliably do by trying, then re-checking as they improve.
- **Trust-calibrated** — autonomy is earned per task, with evidence, governed by risk tier.
- **Identity-aware** — redesigning the human role is part of the method, not a side effect.

4. The foundation: four rails

The loop is the engine; the rails are what keep it on track. Leadership owns the rails; cells own the loop. Without the rails, fast loops derail — they stall in governance, get strangled by shadow-AI bans, or break the people they touch.

4.1 Honest leadership mandate

Set direction, and tell the truth about role change. Leaders name where the organization is going with AI and, crucially, are honest about what it means for people. The most corrosive

move available is the hollow promise that no one will lose their job when that is not certain. Replace it with a commitment you can actually keep — a credible reskilling-and-redeployment path, or transparency about which roles are growing and which are shrinking. People can adapt to hard truths; they cannot adapt to mistrust. The evidence on this wave keeps showing that organizational factors — leadership alignment, culture, governance — explain far more of AI's impact than individual skill or mindset.

4.2 Sanctioned access

Give people capable, approved tools with guardrails — and do it early. Most workforces are already using AI off-policy; a large share of employees experiment outside any official channel. Banning that drives it underground; sanctioning it converts shadow AI into visible, governed adoption and turns existing grassroots energy into your fastest accelerant. The places people are already quietly using AI are also your highest-signal first use cases.

4.3 Guardrails and tiered governance

Build the rails that let cells move fast safely. Define autonomy by risk tier (Section 7), set clear data and security boundaries, and decide in advance where a human must stay in the loop. Governance designed in from the start is an enabler of speed; governance bolted on after an incident is a brake. This matters now because many organizations are already running agents in production without formal governance — a gap that eventually bites.

4.4 Small empowered cells and fast cadence

This is the operating model for speed. Small cross-functional cells — someone who owns the work, someone who can build the agent workflow, someone who owns risk — run the loop on a weekly or biweekly rhythm, with a leader clearing blockers. The thing you are explicitly replacing is the quarterly steering committee. Speed is a structural choice, not a slogan.

5. The operating loop: five stages

Each cell runs the same five-stage loop. The stages are described below with their purpose, the force that justifies them, the core activities, who is involved, what goes in and comes out, when the stage is done, and what to watch for.

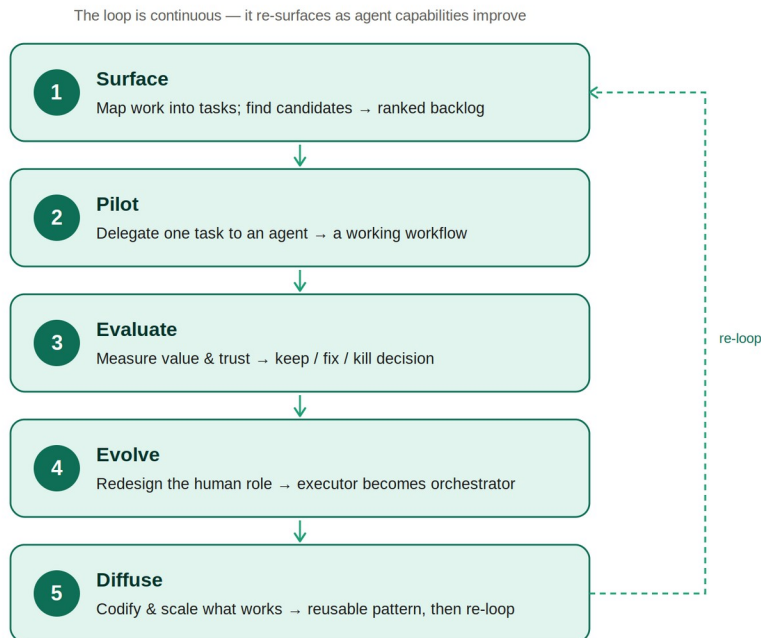


Figure 2. The operating loop, stage by stage.

5.1 Surface

Surface is discovery. Before any tool decision, you map where cognitive work actually lives and where AI is already being used in the shadows.

Why this stage exists. Because the change is about tasks, not tools (Force 3), and because grassroots adoption is already underway (Force 1) — you start from the work and from reality, not from a vendor demo.

Key activities.

- Decompose target roles into discrete tasks; resist staying at the level of job titles.
- Classify each task: fully delegable, agent-does / human-verifies, or stays human (judgment, relationship, accountability, high risk).
- Map existing shadow-AI use — treat it as free reconnaissance, not a violation.
- Rank candidates by value × feasibility × reversibility.

Owner. the cell, with the manager who owns the work; takes days, not weeks.

Inputs → outputs. role and process knowledge, plus observed AI usage → a ranked backlog of candidate tasks, each with an owner.

Exit criteria. a prioritized shortlist of one to three tasks ready to pilot.

Watch out for. boiling the ocean (trying to map everything), and mistaking a whole job for a single task.

5.2 Pilot

Pilot hands one chosen task to an agent inside a small, contained, reversible experiment.

Why this stage exists. Because agents act, the human–agent handoff must be designed deliberately (Force 2); and because speed matters, you ship something in days rather than charter a program (Force 1).

Key activities.

- Design the workflow: what the agent does, what it returns, and where a human checks or intervenes.
- Keep the blast radius small and reversibility high; start with a Tier 1 or Tier 2 task (Section 7).
- Instrument it from day one so Evaluate has data — capture overrides, errors, and time taken.
- Timebox the experiment to a short, fixed window.

Owner. the cell's builder together with the task owner.

Inputs → outputs. a candidate task plus access to sanctioned tools → a working human-plus-agent workflow.

Exit criteria. the workflow runs end-to-end on real or realistic work, with telemetry flowing.

Watch out for. building a grand platform before proving one task; and piloting in a sandbox so unrealistic the results will not transfer.

5.3 Evaluate

Evaluate calibrates trust with evidence and makes an honest keep / fix / kill decision.

Why this stage exists. Trust calibration is the problem unique to agentic AI (Force 2): over-trust ships errors, under-trust wastes the capability. This is also where governance is enforced.

Key activities.

- Measure value (time reclaimed, throughput, quality) and reliability (error rate, escape rate, human-override rate).
- Compare the result against the bar set for the task's risk tier.
- Decide: scale, iterate, or stop — and record the reason.

Owner. the cell together with the risk owner.

Inputs → outputs. the running pilot plus its telemetry → a go / no-go decision backed by evidence.

Exit criteria. a documented decision and, if it is a go, the quality and trust bar the task must keep meeting.

Watch out for. vanity metrics (usage without value); trusting a good demo over measured reliability; and never killing anything.

5.4 Evolve

Evolve redesigns the human role around the capacity the agent has freed. This is the stage the classic frameworks omit, and the one that decides whether adoption sticks.

Why this stage exists. When AI absorbs cognitive tasks, the person's role materially changes (Force 3). Deploy the agent but leave the role untouched and you get the transformation paradox: tools adopted at speed, structures unchanged, value unrealized.

Key activities.

- Redefine the role from executor to orchestrator / verifier / judgment-holder — and write the new responsibilities down.
- Move freed time to higher-value work the person can see and want: exceptions, judgment, relationships, directing more agents.
- Run the human transition honestly — name what is ending, support the uncertain middle, make the new beginning concrete.
- Update job expectations, success measures, and titles where needed.

Owner. the people manager and an HR partner, working with the cell.

Inputs → outputs. a proven workflow → a redesigned role and a person who can see their place in it.

Exit criteria. the role description and success measures reflect the new division of labour, and the person has bought in or their concerns are surfaced.

Watch out for. the hollow reassurance when it is not true; treating this as training instead of redesign; and skipping it because the pilot “worked.”

5.5 Diffuse

Diffuse scales what worked, retires what did not, codifies the pattern, and feeds the loop back to Surface.

Why this stage exists. Speed compounds only if each loop makes the next one cheaper (Force 1); and because capabilities keep improving, the loop must re-open rather than refreeze.

Key activities.

- Codify the winning pattern: templates, prompts, guardrails, and runbooks.
- Roll out to adjacent teams through their own short loops — by pull, not by big-bang mandate.
- Retire failed experiments cleanly and capture the lesson.
- Re-Surface: re-examine “stays human” tasks against improved capability.

Owner. the cell together with an enablement function.

Inputs → outputs. a validated workflow and redesigned role → a reusable pattern and a re-opened backlog.

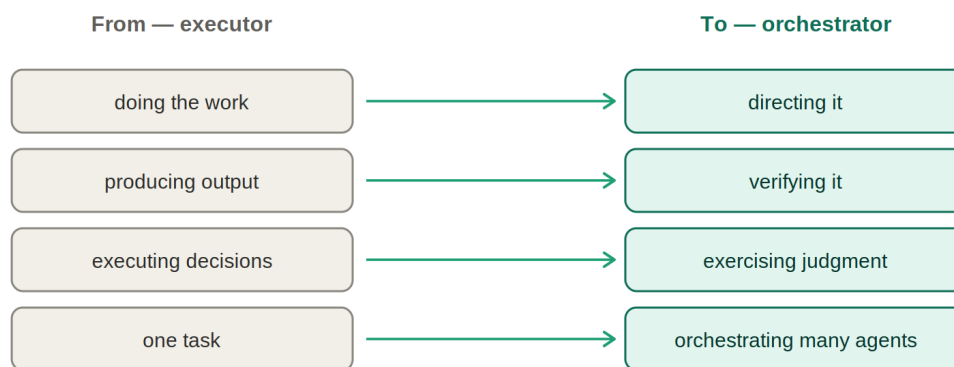
Exit criteria. the next team can adopt the pattern faster than the first did, and the backlog is refreshed.

Watch out for. refreezing (declaring victory and stopping); scaling by mandate instead of by pull; and losing the lessons from failures.

6. The human transition: executor to orchestrator

The deepest shift in this wave is not technological — it is about what people are for. As agents absorb cognitive tasks, the human role moves along four lines:

- **from doing the work to directing it,**
- **from producing output to verifying it,**
- **from executing decisions to exercising judgment** over them,
- **from owning one task** to orchestrating several agents.



An identity change, not a skills upgrade — value shifts to judgment, oversight & relationships.

Figure 3. The human transition: executor to orchestrator.

This is an identity change, not a skills upgrade, which is why it lives as a thread through every loop rather than as a one-time training event. Borrowing from transition theory: name what is ending, support people through the uncertain middle, and make the new beginning concrete and desirable.

Done well, the value of the human rises — judgment, ethics, taste, relationships, and accountability become the scarce, premium contributions. Done badly — tool deployed, role untouched — you get adoption on paper and no value in practice.

7. Governance by risk tier

Tiering is what lets you be fast and safe at the same time: full speed on low-risk work, real care on high-risk work. Set the tiers once, as a rail, and let every cell apply them.

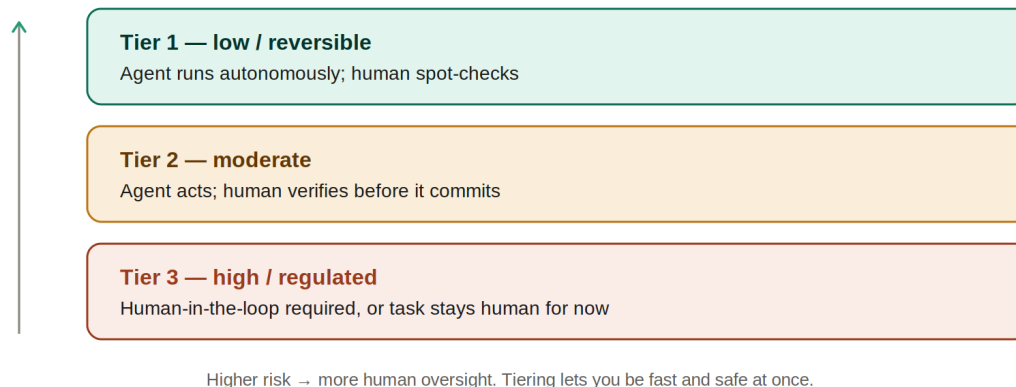


Figure 4. Governance by risk tier.

Tier	Typical work	Human involvement	Example
Tier 1 — low / reversible	Internal, low-stakes, easily undone	Agent runs autonomously; human spot-checks	Internal drafts, first-pass data cleanup
Tier 2 — moderate	Customer-facing or money-adjacent	Agent acts; human verifies before it commits	Draft client replies, prepared quotes
Tier 3 — high / regulated	Legal, financial, safety, or hard to reverse	Human-in-the-loop required, or task stays human for now	Final approvals, regulated filings, hiring calls

8. Putting it to work

8.1 The cell

A cell is small — three to six people — and cross-functional. In a small organization one person can wear two hats, but all four responsibilities must be present.

Role	What they own
Work owner	Knows the actual work and its quality bar; chooses tasks and judges outputs
Builder	Configures the agent and the human–agent handoff (for example, with workflow-automation tooling)
Risk owner	Sets the tier, the guardrails, and what “good enough to trust” means
Sponsor (part-time)	A leader who clears blockers, protects the cadence, and owns the honest role-change message

8.2 A 90-day first cycle

An illustrative first cycle. The point is to reach evidence and redesigned roles within a quarter — and to leave behind a machine that keeps turning.

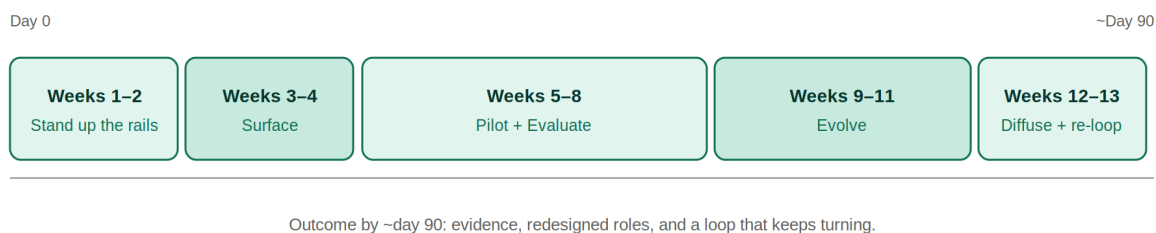


Figure 5. A 90-day first cycle.

Weeks	Focus	Key outputs
1–2	Stand up the rails	Direction and honest role-change position stated; toolset sanctioned; risk tiers published; two to three cells chosen
3–4	Surface	Tasks decomposed and ranked; shadow-AI use mapped
5–8	Pilot + Evaluate	One human-plus-agent workflow per cell, measured against the tier bar
9–11	Evolve	Affected roles redesigned; transition conversations held
12–13	Diffuse + re-Surface	Winning patterns codified; failures retired; next backlog opened

8.3 Task-triage rubric (for Surface)

Score each candidate task on three dimensions — value (how much is at stake), feasibility (how well an agent can do it today), and reversibility (how easily a mistake is undone) — and start where all three are high. The classes below map directly onto the governance tiers.

Task class	Looks like	First move
Delegate	Repetitive, rule-ish, low-risk, easy to verify	Pilot as Tier 1 — let the agent run, spot-check
Co-do	Cognitively heavy but checkable; value in a human sign-off	Pilot as Tier 2 — agent drafts, human verifies
Keep (for now)	Judgment, relationships, accountability, high irreversibility	Leave human; re-examine next loop as capability grows

8.4 Role-redesign playbook (for Evolve)

A repeatable sequence for the stage that most determines whether adoption sticks.

1. **List the tasks leaving the role.** Be specific about what the agent now does.
2. **Name what becomes more important.** Judgment, exceptions, oversight, relationships, directing agents.

3. **Write the new role on one page.** New responsibilities, new success measures, new title if warranted.
4. **Have the honest conversation.** Acknowledge the ending, address the fear directly, describe the new beginning.
5. **Update the system around the person.** Performance metrics, incentives, and team structure — or the old role quietly returns.

9. Measuring it

Track a balanced set so the human side cannot silently fall behind the technical side. The trust metrics should fall over time as calibration improves.

What you track	Example metric	What it tells you
Leading	% of staff with sanctioned access; tasks surfaced; loop cycle time	Whether the machine is turning, and how fast
Adoption	Active usage; tasks delegated to agents	Whether capability is used, not just bought
Value	Human time reclaimed; throughput; output quality	Whether the work is actually better, cheaper, or faster
Trust	Override rate; error / escape rate per task	Whether trust is correctly calibrated
People	Role-redesign completion; sentiment and safety signals	Whether the human side is keeping pace with the tech

10. Anti-patterns to avoid

Each common failure maps to a classic-framework assumption that no longer holds.

- **Pilot paralysis** — endless experiments, nothing in production (no Diffuse stage, no cadence).
- **Suppressing shadow AI** — banning what already works (ignoring grassroots adoption).
- **Tool rollout without work redesign** — the most common cause of stalled adoption (skipping Evolve).
- **The hollow reassurance** — promising no one will be replaced when it is untrue (skipping honest leadership; it breaks the transition).
- **Designing for a final state** — planning to refreeze (Lewin's now-false assumption).
- **Big-bang transformation** — one giant program instead of many fast loops (Kotter at the wrong tempo).
- **Governance as afterthought** — adding guardrails only after an incident (the production-versus-governance gap).

11. How Velora relates to the classic frameworks

Velora does not discard the classics — it rearranges their best parts for a faster, identity-centred, never-finished change.

- **ADKAR's** individual-adoption logic lives inside Evolve and the access rail — but runs in parallel, not as a year-long sequence.
- **Kotter's** coalition and vision become the leadership-mandate rail — set once, then get out of the loop's way.
- **Bridges'** transition model becomes the human thread through every loop — promoted from optional to load-bearing.
- **Lewin's** refreeze is deliberately deleted; the loop's whole job is to keep turning.
- **McKinsey 7-S** remains a useful diagnostic to check the rails stay aligned as you scale.

The one-line difference. The classics ask how to get people to adopt the new tool and settle into the new way. Velora asks how to continuously redistribute cognitive work between people and agents — fast, safely, and without breaking the people whose roles are changing.

12. Getting started

The first two weeks are about the rails, not the technology. Concretely:

1. **Name a sponsor** and write the honest one-paragraph position on what AI will and will not change for people's roles.
2. **Sanction a capable toolset** with basic guardrails, and publish the three risk tiers.
3. **Pick two to three cells** in areas where shadow AI is already in use.
4. **Have each cell Surface its tasks** and choose one to Pilot.
5. **Set the cadence:** a standing weekly checkpoint and a 90-day review.