

The Ideal Workflow for 50 Professions — an agentic-era field guide

A vision artefact from the Knitweb project Published at knitweb.github.io Version 1.0 · 2026

Introduction

This guide makes one claim, fifty times over: in 2026 the ideal workflow for nearly every profession has converged on the same *shape*. What differs between a radiologist and a litigator, a structural engineer and a screenwriter, is the domain — the data, the regulations, the craft. What stays the same is the loop.

The loop is this. A human captain sets intent and holds accountability. They do not type every keystroke; they direct a team of AI agents. Those agents read from and write to a shared, content-addressed memory — a store where every fact, draft, decision, and source is named by its content, so that the same input always resolves to the same address, and nothing can be silently rewritten underneath you. The agents produce work that is *verifiable*: every claim carries provenance, every conclusion cites the addressable evidence it rests on. The captain reviews, corrects, and commits. Then the loop turns again.

For most of computing history, professional software was a collection of nouns: a CAD package, a case-management system, an electronic health record, an IDE. You operated each one by hand. The agentic era inverts this. The nouns become substrate that agents operate; the human moves up an altitude, from operator to director. The unit of work is no longer the document you typed but the *intent you specified and the result you verified*. Productivity stops being a function of how fast you can act and becomes a function of how well you can direct and how rigorously you can check.

Three properties make this loop trustworthy rather than merely fast. **Content-addressing** means the memory is tamper-evident and deduplicated by construction: two agents referencing "the patient's allergy list" or "the load case for column B7" are provably referencing the same thing. **Provenance** means no output is a free-floating assertion; it is a node linked back to its sources, so a reviewer — or an auditor, or a regulator, or a court — can follow any conclusion to its ground. **Verifiability** means the captain's job is concentrated where human judgement is irreplaceable: deciding what matters, and confirming the work is right.

So the fifty chapters that follow are not fifty different workflows. They are fifty renderings of one workflow, each tuned to a profession's evidence, constraints, and stakes. Read across them and the redundancy is the point: once you have internalised the shape, you can walk into a discipline you have never practised and recognise where the captain stands, where the agents work, what the shared memory must hold, and what "cited and verified" means in that field. The professions are the variables. The loop is the constant. Learn the loop, and the rest is domain.

This guide is published by the **Knitweb project** (knitweb.github.io) as a vision artefact — a picture of where professional work is heading, not a product manual. The connective tissue running through every chapter is the shared, content-addressed agent memory: the substrate that lets multiple agents and a human captain collaborate on the same facts without ambiguity, that makes every piece of work provenance-cited by construction, and that turns "the AI said so" into "here is the addressable evidence, follow it yourself." Knitweb exists to build that substrate in the open; this field guide is the argument for why it matters.

How to read this

Read one chapter for your own profession first, then read any two others at random — the shape will repeat, and that repetition is the thesis.

Software & Data

Software engineer

Essence: the engineer is the architect who specifies intent and owns the merge; agents do the typing, and a shared memory graph keeps every decision reproducible.

Loop: - **Intake** — a ticket or product brief lands. The engineer frames the problem and pulls relevant context from the team's content-addressed memory graph: architecture decisions, API contracts, and coding conventions stored as CID nodes rather than scattered across stale `CLAUDE.md`, wiki pages, and per-tool prompt files. The agent reads the same graph, so it starts with the team's actual patterns, not a generic prior. - **Plan** — the agent drafts an implementation plan (files to touch, interfaces, migration steps, test strategy) citing the memory nodes it relied on. The engineer reviews the plan first — cheaper to correct intent here than in code. - **Do** — the agent writes the change across the repo in a worktree/branch, runs the build, and self-iterates against failing types and tests. The engineer pairs interactively on the hard 10%: concurrency, data-model choices, public-API shape. - **Verify** — the quality gate is non-negotiable and machine-checked: type checker, unit + integration tests, lint, and a second reviewing agent that diffs against the spec and flags reuse/simplification. Every claim in the PR description carries provenance — which test proved it, which memory node justified the design. - **Ship/learn** — engineer approves and merges. New conventions and decisions are written back to the graph as CID nodes so the next agent inherits them automatically.

Artefacts/tools: PRs with cited provenance, the memory graph, CI gates, reviewing agent. Quality gate: green CI + spec-diff review + human approval.

The human keeps owning: architectural judgement, the public-interface taste, and accountability for what gets merged into main.

Data scientist

Essence: the scientist owns the question and the validity of the answer; agents run the grind of pulling, cleaning, and fitting, with every figure traceable to the exact data and code that produced it.

Loop: - **Intake** — a business question arrives ("why did retention drop in cohort X?"). The scientist sharpens it into a falsifiable hypothesis with a defined success metric. The agent reads the memory graph for the canonical metric definitions, known data-quality caveats, and prior analyses on the same tables — stored as CID nodes so "monthly active" means one thing, not five. - **Plan** — agent proposes the analysis design: data sources, sample window, confounders, and the statistical test or model class. The scientist vets identification strategy and leakage risks before a single query runs. - **Do** — agent writes queries and notebook cells, profiles distributions, engineers features, and fits candidate models, logging each run to an experiment tracker. - **Verify** — the gate is methodological: held-out validation, calibration and significance checks, data-drift and leakage tests, and a reproducibility re-run from raw data via the content-addressed dataset CID. A reviewing agent challenges the result adversarially ("what would make this spurious?"). Every chart cites its query hash and dataset CID — provenance, not screenshots. - **Ship/learn** — scientist writes the interpretation and recommendation; validated metric definitions and caveats are committed back to the graph so the next analysis inherits them.

Artefacts/tools: versioned notebooks, experiment tracker, content-addressed datasets/features, model cards with provenance. Quality gate: out-of-sample performance + reproducible re-run + adversarial review.

The human keeps owning: framing the question, judging causal validity, and standing behind the recommendation when the business acts on it.

DevOps / platform engineer

Essence: the platform engineer is the captain of a self-service, declarative platform; agents propose and execute changes, but every production mutation passes through a verifiable, reversible gate.

Loop: - **Intake** — a request arrives (new service, scaling event, incident, cost spike). The agent reads the memory graph for the golden-path templates, SLOs, runbooks, and prior incident post-mortems as CID nodes — so it acts on this platform's real topology and guardrails, not a generic cloud pattern. - **Plan** — agent drafts the change as declarative IaC (Terraform/Helm/manifests) plus a rollout and rollback plan, surfacing a `plan` diff and the blast radius. The engineer reviews intent: is this the right abstraction, does it fit the golden path? - **Do** — change goes through GitOps: agent opens the PR, CI renders the diff, policy-as-code runs, and progressive delivery (canary/blue-green) rolls it out behind a feature gate. - **Verify** — the gate is observable and automatic: policy checks (OPA), drift detection, SLO/error-budget burn alerts, and synthetic probes. Promotion is gated on healthy canary metrics; an agent watches the dashboards and auto-rolls-back on regression. Each deploy carries provenance: which commit, which approval, which signed artefact (SLSA). - **Ship/learn** — on incidents, the agent assembles a timeline and drafts the post-mortem; the engineer signs the root-cause and remediation, which become new CID nodes and guardrails.

Artefacts/tools: IaC repos, GitOps controller, policy-as-code, observability + SLO dashboards, signed artefacts. Quality gate: policy pass + healthy canary + error-budget compliance + reversible rollout.

The human keeps owning: production-risk judgement, the guardrail design, and accountability for availability and the on-call contract.

QA / test engineer

Essence: QA shifts from writing every case to owning the risk model and the oracle; agents generate and maintain coverage, while the human decides what "correct" and "good enough to ship" mean.

Loop: - **Intake** — a feature or change set arrives with its spec and acceptance criteria. The agent reads the memory graph for the product's risk map, prior bug clusters, flaky-test history, and the oracle definitions (what the system *should* do) as CID nodes — so testing targets where this product actually breaks. - **Plan** — the QA engineer defines the risk-based test strategy: critical paths, boundary and adversarial cases, non-functional concerns (perf, accessibility, security). The agent expands this into concrete cases and proposes coverage, citing which requirement each case maps to. - **Do** — agent authors and maintains automated tests (unit-supporting, API, E2E), generates property-based and fuzz inputs, and keeps selectors/fixtures current as the UI drifts — killing the maintenance tax that historically rotted suites. - **Verify** — the gate is layered: deterministic test pass, mutation testing to prove the suite actually catches faults, flake quarantine, and exploratory sessions where the QA engineer hunts what automation can't model. Every failure is reported with a reproducible trace and provenance back to the offending commit. An agent triages and de-dupes failures against known issues. - **Ship/learn** — escaped defects feed a post-release analysis; new risks and oracles are written back to the graph, so coverage compounds instead of resetting per release.

Artefacts/tools: versioned test suites, mutation/property/fuzz frameworks, coverage + flake dashboards, traceability matrix. Quality gate: mutation-validated coverage of critical paths + clean exploratory pass + zero unexplained flakes.

The human keeps owning: the risk model, the release/no-release call, and the definition of acceptable quality.

Security analyst

Essence: the analyst is the captain triaging real risk; agents handle the relentless volume of alerts, enrichment, and correlation, while humans own the judgement calls that carry legal and business weight.

Loop: - **Intake** — alerts, findings, and threat intel stream in. The agent reads the memory graph for the asset inventory, prior incidents, suppression rationale, and threat-model assumptions as CID nodes — so it triages against this org's real crown jewels and history, not a generic severity table. - **Plan** — for each candidate incident the agent proposes a hypothesis and investigation plan; the analyst sets priorities and scopes the blast radius. For vulnerabilities, the agent maps findings to reachable, exploitable paths so the team fixes what's actually exposed, not the raw CVE count. - **Do** — agent enriches alerts (threat intel, identity, asset context), correlates across logs, reconstructs attack timelines, and drafts containment actions — never executing privileged or irreversible response without human authorization. - **Verify** — the gate is evidence-based: every conclusion cites its raw telemetry and provenance chain; a second agent red-teams the finding for false positives and alternative explanations; detections are validated against attack simulations before they go live. Severity and exploitability are reproducible, not asserted. - **Ship/learn** — the analyst makes the disclosure, containment, and escalation decisions; confirmed TTPs, new detections, and tuned suppressions are written back to the graph so the next investigation starts smarter.

Artefacts/tools: SIEM/SOAR with agent runbooks, threat-intel and asset graph, detection-as-code, reproducible incident timelines. Quality gate: provenance-cited findings + red-team validation + simulation-tested detections + human authorization for any active response.

The human keeps owning: the risk and disclosure judgement, authorization of irreversible action, and accountability to regulators and leadership.

Healthcare

Physician

Essence: the physician reasons and decides; agents do the reading, drafting, and cross-checking, every claim traceable to a source.

Loop. Intake: an ambient scribe agent transcribes the encounter, structures it into HPI/exam/assessment, and reconciles it against the longitudinal record. Plan: a differential agent proposes a ranked diagnosis list with likelihood reasoning, pulls the patient's labs/imaging/meds, and surfaces guideline and interaction checks — each item citing the source node (UpToDate passage, lab value, prior note) by content address. Do: the physician interrogates the differential, orders tests, and dictates the plan; a drafting agent writes orders, referrals, and prior-auth letters. Verify: a second independent agent audits the plan against allergies, renal dosing, contraindications, and the chart for contradictions, flagging anything unsupported. Ship/learn: the physician signs; outcomes and corrections feed back as new memory nodes.

Knowledge graph. Personal clinical preferences (when this physician escalates, preferred regimens), institutional protocols, and the patient's history live as content-addressed CID nodes the agents read and write — not as scattered prompt files per app. Every agent output cites the exact node version it relied on, so reasoning is replayable and auditable months later.

Artefacts/tools: ambient scribe, differential-and-orders copilot in the EHR, an independent verification agent, the provenance-cited note. Quality gate: nothing reaches the chart unsigned, and the verifier must clear medication/contraindication checks with citations before the physician signs; unsourced assertions are blocked, not surfaced as fact.

Human ownership: the diagnosis, the risk-benefit judgement at the bedside, the conversation with the patient, and full medico-legal accountability for the signed note. The physician is captain; agents are fast, tireless, citation-bound residents who never get the final word.

Nurse

Essence: the nurse owns assessment, vigilance, and the human contact; agents carry the documentation and surveillance load so attention stays on the patient.

Loop. Intake: at handoff, an agent compiles an SBAR from the prior shift's notes, vitals trends, and open orders, flagging what changed. Plan: it proposes a prioritized task list and care-plan updates against unit protocols, ranked by acuity and due time. Do: during rounds, an ambient agent captures assessments and turns them into structured flowsheet entries; barcoded med administration is checked live against the order and the five rights. Verify: a surveillance agent watches vitals and labs continuously, computing early-warning scores and paging with a cited rationale ("MAP trending down + lactate rising") before deterioration is obvious. Ship/learn: end-of-shift handoff is auto-drafted and nurse-edited; near-misses become memory nodes that sharpen future alerts.

Knowledge graph. Unit protocols, this patient's care preferences and baselines, and the nurse's own documented practices live as content-addressed nodes shared across the care team, so the night nurse's agent reads exactly what the day nurse's agent wrote — one memory, not re-entered per system. Alerts cite the node and the data that triggered them, cutting alarm fatigue because each is justified.

Artefacts/tools: SBAR generator, ambient flowsheet capture, barcoded-med verification, continuous early-warning surveillance, auto-drafted handoff. Quality gate: med administration blocks on barcode/order mismatch; every escalation alert carries a cited, inspectable rationale the nurse confirms or overrides.

Human ownership: the bedside assessment a sensor can't make, the judgement to escalate or hold, comforting and advocacy, and accountability for care delivered. Agents reduce charting; they never replace the nurse's eyes, hands, or clinical intuition.

Pharmacist

Essence: the pharmacist owns the therapeutic-safety verdict; agents pre-screen every order so human review concentrates where it matters.

Loop. Intake: as orders arrive, an agent reconciles the full medication list across sources, normalizes doses, and pulls weight, renal/hepatic function, and allergies. Plan: an interaction-and-dosing agent ranks each order by risk — duplicate therapy, interactions, renal adjustment, IV compatibility — citing the reference and the patient datum behind each flag. Do: clean low-risk orders pass with a logged, reversible rationale; flagged orders route to the pharmacist with the evidence assembled. Verify: the pharmacist adjudicates, and a second agent re-checks the final regimen for downstream conflicts and counsels-the-patient points. Ship/learn: interventions and overrides are recorded as memory nodes that tune future risk scoring and catch repeat prescriber patterns.

Knowledge graph. Formulary, interaction databases, institutional dosing protocols, and the patient's pharmacogenomic and history nodes are content-addressed; the verification agent cites the exact node and version, so a dosing decision is reproducible and defensible. Preferences (this pharmacy's renal-dosing conventions) live once as shared nodes, not duplicated per system.

Artefacts/tools: med-reconciliation agent, interaction/renal-dosing screener, IV-compatibility checker, intervention log, counseling-point generator. Quality gate: no order is verified until allergy, interaction, dose-range, and renal/hepatic checks clear with citations; any agent flag the pharmacist overrides is recorded with reasoning for audit.

Human ownership: the clinical-significance call on a flagged interaction, the therapeutic substitution recommendation, the patient counseling conversation, and accountability for every dispensed medication. The pharmacist is the safety captain; agents widen the net and assemble evidence but never sign off alone.

Physiotherapist

Essence: the physiotherapist owns hands-on assessment and the progression decision; agents handle measurement, programming, and adherence so sessions stay therapeutic.

Loop. Intake: an agent compiles the referral, imaging, and history, and during assessment captures range-of-motion, strength, and gait — increasingly from video/wearable motion analysis — into structured, comparable baselines. Plan: a programming agent proposes an evidence-based exercise prescription matched to diagnosis, stage, and goals, citing the guideline and the patient's measured deficits. Do: the therapist treats hands-on and coaches; an agent generates the home-exercise program with video and dosing, and adapts it from wearable/app-reported adherence and pain logs between visits. Verify: at reassessment, an agent quantifies progress against baseline and outcome measures (e.g., goniometry, functional scores), flagging plateau or regression with the data. Ship/learn: outcomes feed memory nodes that refine progression rules and discharge timing.

Knowledge graph. Clinical pathways, the therapist's own progression preferences, and the patient's longitudinal measurements live as content-addressed nodes shared across the clinic, so a covering colleague's agent reads the exact program rationale rather than guessing. Every progression suggestion cites the measured deficit and the protocol node it came from.

Artefacts/tools: motion/ROM capture, evidence-based program generator, home-exercise app with adherence telemetry, outcome-measure tracker. Quality gate: progressions are gated on objective reassessment and red-flag screening (an agent checks for signs warranting medical referral) before load increases; unsupported jumps in intensity are blocked.

Human ownership: the palpation and movement-quality judgement no sensor fully captures, the manual therapy itself, the motivational relationship, and accountability for safe progression. The therapist is captain of the rehab arc; agents measure, program, and remind.

Veterinarian

Essence: the vet owns diagnosis across a non-speaking patient and the owner conversation; agents shoulder records, differentials, and the cost-aware plan.

Loop. Intake: a scribe agent captures the history from the owner and the exam, reconciling species/breed-specific baselines and prior records. Plan: a differential agent reasons across the patient's signalment, presentation, and likely zoonotic/contagious considerations, proposes a diagnostic workup tiered by cost, and pulls species-specific drug dosing — the single most error-prone step in veterinary practice — with each dose citing its source. Do: the vet examines, runs point-of-care diagnostics, and decides; agents draft estimates, consent forms, and discharge instructions. Verify: a second agent re-checks every dose against species, weight, and contraindications (flagging species-toxic drugs), and screens for reportable/zoonotic disease. Ship/learn: outcomes and owner-decision constraints become memory nodes that sharpen future tiered plans.

Knowledge graph. Species formularies, practice protocols, and each animal's history are content-addressed nodes shared across the practice; the dosing-verification agent cites the exact species-dose node and version, so a calculation is reproducible and auditable. The vet's own preferences (preferred protocols, when to refer) live once as shared nodes rather than re-entered per tool.

Artefacts/tools: ambient scribe, species-aware differential and dosing copilot, independent dose-verification agent, tiered-estimate and discharge generator. Quality gate: no treatment proceeds until the verifier clears

species-specific dose, weight calculation, and contraindication/toxicity checks with citations, plus a reportable-disease screen.

Human ownership: the diagnosis from physical signs in a patient who can't report symptoms, the cost-versus-care conversation with the owner, the euthanasia and welfare judgements, and accountability for every dose and procedure. The vet is captain; agents prevent the arithmetic and species errors that cause real harm.

Legal & Finance

Lawyer

Essence: the lawyer argues and decides; agents draft, cite, and surface risk against an auditable record.

Loop: **Intake** — an agent ingests the matter (engagement letter, facts, documents) and writes the matter brief as nodes in the firm's content-addressed knowledge graph: parties, issues, governing law, precedent clauses, and prior-matter playbooks all as CID-referenced entries rather than a fresh prompt typed into a chat window. **Plan** — the lawyer frames the legal theory and risk appetite; a research agent assembles the authority map (statutes, leading cases, regulator guidance), every proposition pinned to a citation node so nothing is asserted without provenance. **Do** — drafting agents produce the contract, memo, or pleading by composing from precedent clause nodes the firm trusts, redlining against the counterparty's draft, and flagging deviations from the house standard. **Verify** — a separate citation-checker agent confirms every authority exists, is good law (not overruled/superseded), and actually supports the sentence it backs; a privilege/conflicts agent screens disclosure; the lawyer reads the reasoning, not just the output. **Ship/learn** — the signed document and the human's overrides feed back as new precedent nodes, improving the next matter.

Artefacts/tools: matter brief, authority map, clause library, redline diff, citation-provenance log. Quality gate: no clause or proposition ships unless its supporting node resolves to a real, current, on-point source; conflicts and privilege checks pass; a partner signs the reasoning trail. Hallucinated-citation risk is structurally contained because outputs cite CID-addressed sources a checker can independently resolve.

The human keeps owning: legal judgement, strategy, ethical duty to client and court, and the signature. The lawyer is captain; agents are the associates who never get tired but never get to decide.

Accountant

Essence: the accountant owns the judgement calls and the sign-off; agents reconcile, classify, and explain every number.

Loop: **Intake** — ledger feeds, bank statements, invoices, and payroll stream in; an agent normalises and posts them against a content-addressed chart-of-accounts and the client's accounting-policy nodes (revenue recognition, capitalisation thresholds, the specific VAT/corporate-tax rules in force) rather than rules re-explained per task. **Plan** — the accountant sets the period's focus and materiality; an agent drafts the close checklist and flags accounts needing judgement (accruals, provisions, related-party items). **Do** — reconciliation agents match transactions, propose classifications, compute tax positions, and draft the financial statements, each entry linked to its source document node. **Verify** — a control agent runs the tie-outs (subledger-to-GL, bank reconciliations, prior-period roll-forward), tests for anomalies and duplicate or out-of-policy postings, and produces a variance narrative; the accountant reviews exceptions and every judgemental estimate. **Ship/learn** — filed accounts plus the accountant's adjustments become provenance nodes; recurring corrections harden the policy graph so next period starts cleaner.

Artefacts/tools: trial balance, reconciliation workpapers, close checklist, tax computation, variance/exception report, signed statements. Quality gate: statements ship only when every balance ties to a resolvable source

node, controls pass, policy nodes match current legislation, and material estimates carry a documented human rationale. Provenance citations make the workpapers self-auditing for any later reviewer.

The human keeps owning: professional judgement on estimates and policy, accountability for the filed numbers, and the relationship of trust with the client and the tax authority. Agents do the arithmetic and the chasing; the accountant decides what is true and fair, and signs.

Financial analyst

Essence: the analyst forms the thesis and bears the call; agents gather, model, and stress-test the evidence.

Loop: **Intake** — an agent pulls filings, market data, transcripts, and prior research into the analysis graph, each datapoint a node tied to its source (10-K page, data vendor, earnings call timestamp). **Plan** — the analyst frames the question (valuation, credit, allocation) and the hypotheses; an agent assembles the comparable set and the model skeleton from the team's trusted methodology nodes, not a one-off spreadsheet copied from a colleague. **Do** — a modelling agent builds the DCF/comps/scenario model, populating drivers from cited sources and writing the assumptions as editable nodes; a research agent drafts the supporting narrative with inline provenance. **Verify** — a challenge agent runs sensitivity and scenario analysis, checks the math, reconciles model outputs to source filings, and argues the bear case against the analyst's thesis; the analyst interrogates the assumptions, especially the ones the model is most sensitive to. **Ship/learn** — the published note and the realised outcome are logged back; forecast-vs-actual feeds a calibration node so the desk's assumptions improve and persistent biases surface.

Artefacts/tools: model with cited drivers, comp set, scenario matrix, thesis memo, assumption-provenance log, forecast-tracking record. Quality gate: no figure enters the model without a resolvable source node; every key assumption is explicit and stress-tested; the published call states its evidence and its kill-criteria.

The human keeps owning: the thesis, the conviction-weighting, the judgement on what the numbers mean, and accountability for the recommendation. Agents make the analysis faster and harder to fool yourself with; the analyst makes the call.

Auditor

Essence: independence and professional scepticism stay human; agents test populations, trace evidence, and document the trail.

Loop: **Intake** — an agent ingests the full ledger, contracts, and prior-year files, mapping them to a content-addressed audit graph: assertions, risks, controls, and the relevant standards (ISA/GAAP/IFRS) as referenced nodes rather than instructions re-pasted each engagement. **Plan** — the auditor sets materiality and the risk assessment; an agent drafts the audit programme, linking each planned procedure to the assertion and risk it addresses. **Do** — testing agents run full-population analytics (not just samples) over journal entries, flag unusual postings, period-end manipulations, and segregation-of-duties breaches, and gather evidence by resolving each test back to its source document node. **Verify** — a second agent re-performs key tests independently and checks that every conclusion is supported by evidence in the working papers; the auditor exercises scepticism on management estimates, going-concern, and anything the analytics flag, deciding what warrants deeper investigation. **Ship/learn** — the opinion plus the human's judgements are logged; findings and missed-risk patterns feed back to sharpen next year's risk model.

Artefacts/tools: audit programme, full-population analytics, evidence-linked working papers, exception log, sign-off trail. Quality gate: an assertion is cleared only when supporting evidence resolves to source nodes, independent re-performance agrees, and standards-compliance checks pass; the engagement partner signs the reasoning. Provenance and content-addressing make the file inspectable by reviewers and regulators end to end.

The human keeps owning: independence, professional scepticism, the judgement on material misstatement and going-concern, and accountability for the opinion. Agents widen coverage from samples to whole populations and document tirelessly; the auditor decides what the evidence means and signs the opinion.

Insurance underwriter

Essence: the underwriter prices and accepts the risk; agents assemble the submission, enrich it, and check it against appetite.

Loop: **Intake** — a submission (application, loss runs, exposure data, broker notes) arrives; an agent structures it into the risk graph and enriches it with external nodes (property, geospatial/cat, sanctions, financials), each enrichment carrying its source so nothing is unattributed. **Plan** — the underwriter sets the angle; an agent checks the risk against the carrier's appetite and authority nodes (eligible classes, capacity, exclusions, reinsurance treaty terms) and flags anything out of mandate. **Do** — a pricing agent runs the rating model and cat/exposure analytics, drafts terms, conditions, and exclusions from the trusted wordings library, and produces the indication with its rationale; a referral agent routes anything exceeding authority. **Verify** — a control agent confirms data completeness, sanctions/regulatory screening, treaty compliance, and that quoted terms match the bound wording; the underwriter judges the moral/physical hazard, the price adequacy, and the relationship factors the model can't see. **Ship/learn** — bound policy and subsequent claims experience feed back as outcome nodes; loss emergence recalibrates the appetite and rating models.

Artefacts/tools: structured submission, enrichment-provenance log, rating output, terms/wordings draft, referral and screening record. Quality gate: a risk binds only when data is complete and sourced, screening passes, terms sit within authority and treaty limits, and pricing meets the adequacy threshold; out-of-appetite risks require explicit human acceptance with documented reasoning.

The human keeps owning: the risk-acceptance decision, judgement on hazard and price adequacy, broker relationships, and accountability for the book's profitability. Agents make every submission complete and consistently screened; the underwriter decides what to write and at what price.

Science & Research

Chemist

Essence: the chemist designs the hypothesis and judges what is real; agents run the literature, the calculations, and the instrument bookkeeping around the bench.

The loop starts at **intake**, where the chemist states a target (a molecule, a yield improvement, a mechanism question). A literature agent pulls prior syntheses, hazard data, and failed routes, writing each finding as a provenance-cited node into a shared content-addressed memory graph (reagent properties, reaction conditions, SMILES, spectra all addressed by CID, so "the Suzuki conditions that worked in March" resolve to one immutable node rather than a copy pasted into five notebooks). At **plan**, a retrosynthesis agent proposes routes ranked by cost, safety, and predicted yield; the chemist prunes them using chemical intuition the model lacks. At **do**, agents generate the executable procedure, pre-fill the ELN, compute stoichiometry and limiting reagent, and queue any automated/HTE platform; the human runs or supervises the bench. At **verify**, an analysis agent ingests NMR/LC-MS/IR, proposes peak assignments and a purity call, and flags any spectrum inconsistent with the claimed structure — but every structural assignment requires the chemist's explicit sign-off, because a confident wrong assignment is the expensive failure mode. At **ship/learn**, the validated result (including negative results) is written back to the graph as a new CID node, so the next campaign reads it automatically.

Artefacts/tools: ELN (e.g. structured notebook), retrosynthesis and spectral-prediction models, HTE/flow hardware, the CID memory graph as the single source of truth. Quality gate: structure confirmed by orthogonal

methods (mass + NMR + purity) with agent-proposed, human-ratified assignments and full provenance.

What the human keeps owning: mechanistic judgement, safety accountability, and the final "this structure is correct" call — the chemist is captain of the science, agents are instrumented crew.

Biologist

Essence: the biologist owns the biological question and the interpretation; agents handle protocol design, data wrangling, and statistical hygiene so wet-lab time goes to biology, not bookkeeping.

At **intake**, the biologist frames a hypothesis (a pathway, a phenotype, a knockout effect). An agent surveys literature and public datasets (GEO, UniProt, PDB, model-organism databases), depositing pathway facts, antibody validation status, and prior effect sizes as cited nodes in the shared content-addressed graph — so a validated antibody or a cell-line authentication is one CID node reused across every project, not re-litigated per experiment. At **plan**, a design agent proposes the experiment with proper controls, power analysis, and randomization, and pre-registers the analysis plan; the biologist adjusts for biological plausibility and feasibility. At **do**, agents generate bench protocols, schedule instruments, and template the data capture (plate maps, sample manifests with CIDs). The human executes or directs the wet work. At **verify**, an analysis agent runs the pre-registered pipeline (QC, normalization, the agreed statistics), surfaces batch effects, and refuses p-hacking by holding the analysis to the registered plan; figures carry provenance back to raw data. At **ship/learn**, results — positive and null — are committed to the graph and linked to the protocol CID, making replication a re-run rather than an archaeology project.

Artefacts/tools: LIMS/ELN, pre-registration, reproducible pipelines (Nextflow/Snakemake, notebooks under version control), the CID memory graph. Quality gate: pre-registered analysis executed faithfully, controls behaving, replicates consistent, every figure traceable to raw data.

What the human keeps owning: biological interpretation, judging confounds and artefacts, and accountability for what the data actually mean — the model can compute significance but not biological significance.

Academic researcher

Essence: the researcher owns the argument and its integrity; agents accelerate discovery, drafting, and citation so the scholar spends time on ideas and rigour, not formatting.

At **intake**, the researcher poses a question and scopes novelty. A literature agent maps the field, finds the gap, and builds a citation graph as content-addressed nodes — each claim node links to its source CID, so "this is supported by X" is verifiable rather than asserted, and the same evidence node is shared across papers, grants, and reviews instead of re-summarized each time. At **plan**, an agent drafts the study/proof/argument structure and a research plan; the researcher decides which line is actually worth pursuing — taste the model can't supply. At **do**, agents help derive, simulate, draft prose, and assemble figures, each output carrying provenance and an explicit confidence/uncertainty note. The researcher writes the load-bearing reasoning. At **verify**, a critique agent runs an adversarial review: checks every citation resolves and actually supports the claim, re-derives key steps, hunts for unsupported leaps and missing limitations, and flags potential plagiarism or fabricated references. Co-authors and peer review remain the human gate. At **ship/learn**, the paper, data, and code are released openly with the citation graph attached, and the researcher's accumulated method/preference nodes feed the next project.

Artefacts/tools: reference manager backed by the CID graph, version-controlled manuscript and analysis, open-data/code repositories, an LLM critique pass. Quality gate: every citation resolvable and faithful, results reproducible from deposited code/data, claims matched to evidence with stated limitations — no orphan or hallucinated references.

What the human keeps owning: the thesis, intellectual honesty, authorship accountability, and the judgement of what is genuinely novel and worth saying.

Clinical-trial coordinator

Essence: the coordinator owns patient safety and regulatory accountability; agents carry the relentless tracking, scheduling, and documentation so nothing slips and humans focus on participants.

At **intake**, the protocol and eligibility criteria are encoded once into the shared content-addressed graph (each criterion, visit window, and SAE definition a CID node), so every downstream check references the same authoritative source instead of a drifting PDF copy. A screening agent matches candidate patients against eligibility and flags likely fits; the coordinator and PI make the enrollment decision and obtain informed consent — never delegated. At **plan**, a scheduling agent builds each participant's visit calendar within protocol windows, sequences labs and procedures, and pre-empts deviations before they happen. At **do**, agents pre-populate the EDC/eCRF from source documents, reconcile against source for discrepancies, track drug accountability, and watch for adverse events and out-of-window visits in real time, escalating to the coordinator. The human handles the participant relationship and clinical calls. At **verify**, a monitoring agent runs continuous data review — completeness, range checks, SAE reporting timelines, consent currency — so audit-readiness is standing, not a scramble; every alert links back to the protocol CID and the source record. At **ship/learn**, clean, locked data with full provenance moves to analysis, and deviation patterns feed protocol improvement.

Artefacts/tools: CTMS, EDC/eCRF, eISF/regulatory binder, IRB submissions, the CID protocol graph as single source of truth. Quality gate: source-data verification passes, SAEs reported within mandated timelines, consent valid and current, deviations logged and explained — inspection-ready at all times.

What the human keeps owning: informed consent, patient safety judgement, clinical decisions, and full regulatory/ethical accountability — agents track, humans are responsible.

R&D engineer

Essence: the engineer owns the design intent and the risk call; agents run the design-space search, simulation, and test automation so iteration is fast and every result is traceable.

At **intake**, the engineer states requirements and constraints (performance, cost, manufacturability, regulatory). An agent retrieves prior designs, failure reports, standards, and material data, writing them as content-addressed nodes — so a qualified component, a tolerance decision, or a root-cause from a past failure is one reusable CID node, not tribal knowledge lost in a folder. At **plan**, a design agent proposes candidate architectures and runs design-of-experiments / parameter sweeps; the engineer chooses the direction using engineering judgement about what the simulator under-models (manufacturing reality, edge cases, second-order effects). At **do**, agents drive CAD/simulation (FEA/CFD/circuit sims), generate test plans, and script the test rig; the engineer builds and instruments the prototype. At **verify**, an analysis agent compares measured results against simulation and spec, flags discrepancies and likely root causes, and holds outputs to acceptance criteria with provenance from requirement → design → test → result fully linked. Design reviews remain a human gate. At **ship/learn**, the validated design, its test evidence, and any failure modes are committed to the graph, so the next revision and the next project inherit the lesson automatically.

Artefacts/tools: PLM/requirements management, version-controlled CAD and simulation, automated test benches and data capture, the CID design/decision graph. Quality gate: requirements traceability complete, simulation validated against measurement, design review signed off, tests passing against acceptance criteria with full provenance.

What the human keeps owning: design intent, the trade-off and risk judgement, sign-off accountability, and knowing where the model's confidence outruns physical reality — the engineer is architect, agents are the search-and-verify crew.

Engineering & Trades

Mechanical engineer

Essence: the engineer sets intent and constraints; agents drive the parametric design–simulate–revise loop while every decision stays traceable to physics and standards.

Loop. Intake: the engineer states the duty cycle, loads, materials, cost and manufacturing constraints. An intake agent pulls prior assemblies, test data and the relevant standards (ISO, ASME, DIN clauses) from a shared content-addressed knowledge graph, where each part family, material allowable and lessons-learned note is a CID node with provenance — not a prompt buried in one CAD plugin. Plan: the engineer fixes the architecture and the safety factor; a planning agent proposes the parameter space and load cases. Do: agents run parametric CAD, generative topology and FEA/CFD sweeps, returning Pareto fronts on mass, stress margin and cost; the engineer prunes by judgement. Verify: an independent checker agent re-runs FEA with a different mesh/solver, validates units and boundary conditions, checks tolerance stack-up and DFM rules, and flags any allowable used outside its cited test range. Ship/learn: released drawings, the simulation deck and the rationale are committed back as new CID nodes, enriching the graph for the next project.

Artefacts and gates. Parametric CAD, FEA/CFD decks, GD&T drawings, an MBOM, and a digital twin. The quality gate is dual: every stress and fatigue number must cite a verified material allowable and a re-solved independent model, with mesh-convergence and units checks passing before sign-off. Test rig data closes the loop against simulation.

Human owns. The choice of architecture and failure mode to design against, the safety factor under uncertainty, and accountability for the stamped design. The engineer is captain: agents explore and verify the space; the human decides what is safe, manufacturable and worth building, and signs.

Civil engineer

Essence: the engineer owns the public-safety call and the code interpretation; agents carry the load-path modelling, code-checking and document churn with full provenance.

Loop. Intake: site data, geotech reports, survey and the governing code edition arrive; an intake agent normalises them into the project's content-addressed graph, where soil profiles, load combinations, design checks and as-built lessons live as CID nodes citing their source survey or borehole. Plan: the engineer chooses the structural scheme and foundation strategy; a planning agent lays out analysis cases (gravity, wind, seismic, settlement). Do: agents build and run the structural and hydraulic models, size members, and draft calculations against the cited code clauses; the engineer adjudicates every governing assumption. Verify: a checker agent independently re-derives critical load paths, cross-checks code clause-by-clause (e.g. Eurocode/ACI/AISC), reconciles model reactions against hand calculations, and flags any clause whose edition or applicability is uncertain. Ship/learn: stamped calc packages, drawings and the model are committed back with provenance; field as-builts update the graph.

Artefacts and gates. Structural/hydraulic analysis models, clause-cited calculation packages, BIM/drawings, and a permit set. The gate: every demand-to-capacity ratio traces to a code clause and an independent hand-or-second-model check, with load-combination and units audits green. Constructability and inspection hold-points are explicit.

Human owns. Interpretation of code intent in edge cases, the factor-of-safety judgement where the public is at risk, and accountability for the stamp. Agents never decide what is safe enough; the licensed engineer does, and is the captain of record.

Electrician

Essence: hands stay human and code-literate; agents pre-compute the design, load math and compliance so the electrician verifies and energises with confidence.

Loop. Intake: the electrician captures the panel, loads, and site conditions (photos, nameplate data); an intake agent reads them into a content-addressed job graph where circuits, panel schedules and code interpretations are CID nodes citing the governing edition (NEC/IEC/local). Plan: the electrician decides the topology and protection scheme; a planning agent computes load calcs, conductor sizing, voltage drop, breaker/OCPD selection and conduit fill against cited clauses. Do: agents generate the panel schedule, one-line diagram, pull-list and material take-off; the electrician walks the route and adjusts for real-world obstructions. Verify: a checker agent re-runs the load and voltage-drop math, validates derating and continuous-load factors, and cross-checks every sizing decision against the cited code article; on site the electrician confirms with meter and torque tools. Ship/learn: the as-built one-line, test readings (insulation, continuity, polarity) and any deviations are committed back as CID nodes, so the next job inherits real panel history.

Artefacts and gates. One-line diagrams, panel schedules, load calculations, take-offs, and a test/commissioning log. The gate: every conductor and OCPD size cites a code article and passes an independent recomputation, plus physical verification (megger, torque, polarity, RCD trip) before energising. Inspection sign-off closes it.

Human owns. The hands-on install quality, the live-work judgement, lockout/tagout discipline, and accountability for a safe, inspectable system. Agents do the arithmetic and the code lookup; the electrician decides and energises.

Architect

Essence: the architect owns the design vision and the human experience; agents accelerate option-generation, code-checking and documentation while every choice keeps its rationale.

Loop. Intake: brief, site, zoning and program land in a content-addressed project graph where precedents, material palettes, zoning constraints and client preferences are CID nodes with provenance — replacing scattered moodboards and per-tool prompts. Plan: the architect sets the concept and spatial parti; a planning agent maps the program against zoning envelope, egress and accessibility rules. Do: agents generate massing and layout options, run daylight, energy and circulation analyses, and draft the BIM model; the architect curates by taste, light and feeling. Verify: a checker agent runs an automated code/zoning compliance pass (egress, ADA/accessibility, fire separation, setbacks, energy), reconciles the BIM against the program, and cites each rule it checked; consultants confirm structure and MEP coordination via clash detection. Ship/learn: the design package, narrative and analysis are committed back as CID nodes; post-occupancy feedback updates the graph.

Artefacts and gates. BIM model, drawing set, daylight/energy analyses, code-compliance report, and renders. The gate: a clean automated compliance and clash-detection pass with every flag resolved or cited, plus consultant sign-off, before the permit set issues. Design intent is documented, not just the geometry.

Human owns. The design vision, spatial and material taste, the client relationship, and accountability for a building people inhabit. Agents expand and check the option space; the architect is the captain who decides what gets built and why.

Logistics / supply-chain planner

Essence: the planner owns the trade-offs and the risk appetite; agents run the optimisation, scenario and exception loops continuously against a shared, provenance-rich network model.

Loop. Intake: demand signals, supplier lead times, inventory positions and constraints stream into a content-addressed network graph where SKUs, lanes, supplier reliability and service policies are CID nodes citing their source feed. Plan: the planner sets service-level targets, cost ceilings and risk posture; a planning agent frames the optimisation (inventory, routing, sourcing) and the scenarios to stress. Do: agents run demand forecasts, MEIO/inventory optimisation, network and routing solves, and propose POs and allocations; the planner adjudicates exceptions and supplier trade-offs. Verify: a checker agent back-tests forecasts against held-out actuals, validates that proposed plans respect capacity, MOQ and lead-time constraints, runs disruption what-ifs, and flags any recommendation resting on stale or low-confidence data with its provenance. Ship/learn: approved plans, the rationale and realised outcomes are committed back as CID nodes; forecast error and OTIF feed the next cycle.

Artefacts and gates. Demand forecasts, optimisation models, scenario/digital-twin runs, replenishment and allocation plans, and a control-tower dashboard. The gate: every plan passes a constraint-feasibility check and a back-test threshold, with each input traced to a dated, provenance-cited source before release. Exceptions are explained, not silently auto-executed.

Human owns. The service-versus-cost-versus-risk trade-off, supplier relationships, and accountability for resilience when the model meets reality. Agents optimise and stress-test continuously; the planner is captain, deciding which plan the business actually commits to.

Creative & Media

Graphic Designer

Essence: The designer owns the visual argument and brand judgement; agents handle production, variants, and spec compliance so the human spends their hours on composition and meaning, not file wrangling.

Loop: *Intake* — a brief agent ingests the request, references, and the brand system (logos, grids, type ramp, color tokens, accessibility rules) as CID nodes in a shared content-addressed memory, so "our brand voice" resolves to a pinned node, not a re-pasted PDF. *Plan* — the designer and a concepting agent co-produce 3–5 divergent directions with rationale; the human picks and steers. *Do* — generation/layout agents draft compositions in Figma via API, populate the design-system components, and produce responsive and localized variants. *Verify* — an automated gate checks WCAG contrast, token adherence, export specs (bleed, color profile, DPI), and trademark/license provenance on every asset (each stock or AI-generated element carries a citation node). *Ship/learn* — approved files export to all required formats; accepted decisions and rejected directions feed back as memory nodes, sharpening the next intake.

Artefacts/tools: Figma + design tokens, a versioned brand-system graph, automated contrast/spec linters, license-provenance manifests, generative fill/upscale agents. The quality gate is concrete: a build fails if a color is off-token, contrast is sub-AA, an asset lacks a usage license, or an export misses spec.

Human keeps owning: Taste, hierarchy, and the call on what *feels* right and on-brand — plus accountability for what ships under the client's name. The agent proposes; the designer is the art director who decides and signs.

Video Editor

Essence: The editor owns story, rhythm, and emotional pacing; agents do the logging, rough assembly, and conform so the human edits the cut, not the chaos of footage.

Loop: *Intake* — an ingest agent transcribes, scene-detects, tags faces/shots, and logs B-roll, writing a searchable timeline index into the shared memory graph keyed to the project's style node (pacing, grade, music brief). *Plan* — editor and an assembly agent agree a paper-cut from transcript and beats; the human sets the spine. *Do* — agents build a string-out and rough cut, sync multicam, draft captions, and propose music/SFX

from licensed libraries with rights metadata attached. *Verify* — a QC gate checks loudness (EBU R128/-14 LUFS for platform), legal-quality safe areas, frame rate/codec conform, caption sync, and that every clip, track, and AI element carries a clearance/provenance node. *Ship/learn* — renders per delivery spec; the director's notes and accepted timing choices persist as nodes so the editor's signature pacing compounds across projects.

Artefacts/tools: Premiere/Resolve/Final Cut via API, transcript-driven editing, auto-multicam sync, loudness/QC linters, a rights-cleared music graph. The gate is hard: a delivery fails on wrong LUFS, unsafe titles, dropped frames, unsynced captions, or an unlicensed asset.

Human keeps owning: The cut — where to hold, where to cut on action, what the scene *means*. Agents assemble; the editor shapes feeling and is accountable for the final story and its clearances.

Photographer

Essence: The photographer owns the eye, the moment, and the light; agents handle culling, organization, and tedious retouch so the human shoots and selects with intent.

Loop: *Intake* — on ingest, an agent reads EXIF, runs face/scene recognition, flags focus/blink/exposure failures, and indexes the shoot into the shared memory graph against the client's style node (grade, crop preferences, must-include shot list). *Plan* — photographer reviews the agent's culled selects and shortlists hero frames; the human makes the picks. *Do* — retouch agents apply consistent grading, skin/dust cleanup, lens corrections, and batch exports per use (web, print, social) while preserving a non-destructive edit history. *Verify* — a quality gate checks color-space correctness (sRGB vs. Adobe RGB vs. print profile), resolution/DPI per output, release-form coverage for identifiable subjects, and C2PA/provenance signing so authenticity travels with the file. *Ship/learn* — deliver galleries; accepted grades and reject reasons feed the style node, so the next shoot's auto-cull and grade match this client's taste.

Artefacts/tools: Lightroom/Capture One via API, AI culling, masking and generative-remove tools, C2PA content credentials, model-release tracking. The gate is concrete: export fails on wrong color profile, missing release, sub-spec resolution, or unsigned provenance.

Human keeps owning: The decisive moment and the final select — which frame is *the* frame, and how far retouch may go before authenticity breaks. The photographer is accountable for truth and consent in the image.

Musician / Producer

Essence: The producer owns the musical intent and the hook; agents handle arrangement scaffolding, stem prep, and technical mixing chores so the human chases the feel.

Loop: *Intake* — a session agent captures the reference, key, tempo, and mood into a project node in the shared memory graph, linking the artist's sonic identity (preferred chains, palette, mix targets) as reusable CID nodes. *Plan* — producer and an arrangement agent sketch song structure and chord/groove options; the human chooses the direction and writes the core idea. *Do* — agents generate complementary parts, comp takes, time/pitch-align, organize stems, and propose mix moves (EQ, gain-staging, bus routing) the producer accepts or overrides. *Verify* — a mastering/QC gate checks streaming loudness targets (about -14 LUFS, true-peak $\leq$ -1 dBTP), phase/mono compatibility, sample-clearance and split-sheet completeness, and tags any AI-generated part with provenance for rights. *Ship/learn* — export masters and platform variants; accepted chains and rejected ideas persist so the artist's signature sound is recallable, not rebuilt each session.

Artefacts/tools: DAW (Ableton/Logic/Pro Tools) via API/plugins, stem-separation, auto-align, reference-matching mastering, a split-sheet/clearance graph with provenance. The gate is hard: a master fails on loudness/true-peak overshoot, mono-collapse cancellation, or missing sample clearance or writer splits.

Human keeps owning: The hook, the groove, and emotional intent — what makes it *move* someone — and accountability for authorship, credit splits, and clearance.

Copywriter

Essence: The writer owns the angle, voice, and persuasive judgement; agents handle research, drafting variants, and compliance so the human edits for resonance and truth.

Loop: *Intake* — a brief agent loads the audience, offer, and the brand voice/messaging pillars as CID nodes in the shared memory graph, plus the legal/claims guardrails, so "our tone" and "approved claims" resolve to pinned nodes, not scattered prompt files. *Plan* — writer and a strategy agent agree the angle, structure, and key message; the human sets the argument. *Do* — drafting agents produce headline and body variants, channel adaptations (email, landing, social, ad), and SEO/readability passes, each citing its source for every factual claim. *Verify* — an editorial gate checks voice-token adherence, claim provenance (every fact links to a source node), banned-phrase and regulatory compliance (e.g., no unsubstantiated health/finance claims), reading level, and originality/plagiarism. *Ship/learn* — publish or hand off; A/B results and the editor's accepted/rejected lines feed back as voice and performance nodes, so the brand's language compounds.

Artefacts/tools: Briefing and drafting agents, a versioned brand-voice graph, claim-citation manifests, compliance/banned-term linters, readability and plagiarism checks. The gate is concrete: copy fails to ship if a claim is uncited, a regulated assertion is unsubstantiated, voice drifts off-token, or originality is suspect.

Human keeps owning: The angle and the line that lands — taste for rhythm, restraint, and persuasion — plus accountability for truthfulness and that the words mean what the brand intends. Agents draft; the writer is the editor-in-chief who signs.

Education & Knowledge

School teacher

Essence: the teacher orchestrates learning and judges understanding; agents absorb the prep, marking, and differentiation that used to eat evenings.

Loop: *Intake* — the teacher pulls the class's living state from a shared content-addressed memory: each pupil's mastery map, misconceptions, IEP/access needs, and prior lesson CIDs, so nothing is re-keyed per tool. *Plan* — a lesson-design agent drafts a sequence against the curriculum standard (cited to the exact spec node) and proposes three differentiation tiers; the teacher edits intent, pacing, and the "hook." *Do* — agents generate the slide deck, a low-floor/high-ceiling worksheet, sentence frames for EAL pupils, and a retrieval-practice quiz, each tagged to the objective CID. In-lesson, the teacher reads the room; an aide agent surfaces a hint or extension on request. *Verify* — the teacher grades a calibration sample by hand, then the marking agent scores the rest against that rubric and flags only the borderline scripts for human review; misconception clusters are written back to the memory graph as new nodes. *Ship/learn* — exit-ticket data updates each mastery map; tomorrow's plan starts from real evidence, not a guess.

Artefacts/tools: standards-aligned lesson graph, rubric-anchored auto-marking with teacher-calibrated samples, per-pupil mastery dashboard, parent-update drafts. Quality gate: every agent output cites the curriculum node and the pupil-evidence it rests on; the teacher signs off before anything reaches a child, and grades that affect records are human-confirmed.

Human keeps owning: the relationship, the moral and safeguarding judgement, reading the unspoken, and accountability for what each child actually learns — the agents never decide who a pupil is.

University lecturer

Essence: the lecturer is the scholar-architect of a course and its assessment integrity; agents handle production, logistics, and first-pass feedback.

Loop: **Intake** — pull the module's content-addressed graph: learning outcomes, the reading list with stable citation CIDs, last cohort's exam-item analysis, and accreditation constraints. **Plan** — a curriculum agent maps outcomes to weekly sessions and assessments, flags coverage gaps against the field's current literature (provenance-cited, not hallucinated), and proposes constructive-alignment fixes; the lecturer sets the intellectual arc and the hard questions. **Do** — agents produce lecture scaffolds, problem sets with worked solutions, lab notebooks, and accessible captions/transcripts, each linked to its source node so students can trace every claim. **Verify** — for assessment, an agent runs item-difficulty and discrimination checks on draft questions and screens for ambiguity; the lecturer owns academic-integrity design (oral defences, viva prompts, AI-resilient tasks). Marking uses rubric agents on a lecturer-graded calibration set, with all boundary and appeal cases escalated to the human. **Ship/learn** — post-exam analytics and student feedback write back as CID nodes, so the next iteration inherits evidence rather than folklore.

Artefacts/tools: aligned syllabus graph, citation-verified slide/notes packs, psychometric item analysis, rubric-anchored grading with audit trail, research-feed digests. Quality gate: every generated claim carries a checkable citation; grades and integrity rulings are human-signed; assessment validity is statistically reviewed before reuse.

Human keeps owning: scholarly judgement and the frontier of the field, what counts as rigour, integrity calls, and the duty of care behind every mark and reference written into a student's record.

Librarian / information specialist

Essence: the librarian is the trust authority over an institution's knowledge graph; agents extend reach and retrieval, never adjudicate credibility.

Loop: **Intake** — a query (a researcher's systematic-review need, a public reference question, a collection gap) enters with its context. **Plan** — the librarian frames scope, sources, and inclusion criteria; a search agent translates this into database-specific query strings (PubMed, Scopus, catalogue, archive APIs) with the exact strategy recorded as a reproducible CID. **Do** — agents run federated searches, deduplicate, extract metadata, and assemble a provenance-cited result set where every item links to its authoritative record; for cataloguing, an agent proposes MARC/RDA or linked-data entries against authority files. **Verify** — the librarian audits a sample for relevance, authority, and bias, confirms licensing/access rights, and checks the agent's provenance chain end-to-end; nothing enters the trusted graph without a human authority-control sign-off. **Ship/learn** — the curated set, its search strategy, and quality notes are written back as content-addressed nodes, so the next identical question is answered instantly and reproducibly, and patrons can re-run the exact strategy.

Artefacts/tools: reproducible search-strategy logs, deduplicated cited result sets, authority-controlled catalogue records, licensing/rights register, information-literacy guides. Quality gate: every result traces to an authoritative source; search strategies are reproducible and recorded; the human confirms authority, rights, and absence of fabricated citations before anything is shelved as trusted.

Human keeps owning: the judgement of credibility and bias, privacy and intellectual-freedom ethics, collection values, and accountability for what the institution certifies as reliable knowledge.

Instructional designer

Essence: the designer architects how people learn and proves it works; agents mass-produce variants and run the analytics behind every design decision.

Loop: **Intake** — pull the project's content-addressed memory: audience analysis, the competency model, brand/accessibility standards, and prior course performance data, so design intent persists across tools instead of living in scattered briefs. **Plan** — the designer writes measurable objectives and the assessment-of-learning strategy; a design agent proposes the learning architecture (sequence, modality mix, spaced-practice schedule) tied to each competency CID and cites the learning-science basis. **Do** — agents generate storyboards, branching scenarios, knowledge-check items, video scripts, and SCORM/xAPI packages, each linked to its objective node; alt-text, captions, and WCAG conformance are produced by default. **Verify** — an agent runs a heuristic and accessibility audit and drafts a pilot plan; the designer runs the actual pilot, and learning-analytics agents surface where learners stall or drop. Item-level data decides what ships. **Ship/learn** — performance and xAPI evidence write back as CID nodes; the next revision is driven by measured impact, and reusable interactions become shared, addressable components instead of one-off rebuilds.

Artefacts/tools: objective-aligned storyboard graph, branching-scenario authoring, xAPI/SCORM analytics dashboard, WCAG audit reports, reusable interaction library. Quality gate: every asset maps to a measured objective; accessibility conformance is checked automatically; design changes are justified by pilot evidence and learner analytics, with the human approving before release.

Human keeps owning: the pedagogical theory of change, taste in what makes learning land, ethical use of learner data, and accountability for whether the design actually changes behaviour.

Translator

Essence: the translator owns meaning, register, and cultural judgement; agents do first-pass conversion, terminology, and consistency at scale.

Loop: **Intake** — the source lands with its brief: domain, audience, register, locale, and the client's content-addressed translation memory plus termbase (each approved term a CID with usage context and prior decisions). **Plan** — the translator sets strategy — what must be localised vs. transcreated, tone, forbidden renderings — and an agent pre-translates leveraging the TM, flagging fuzzy matches, untranslatables, and ambiguities for human attention rather than silently guessing. **Do** — agents draft the target text, enforce termbase consistency, and surface every place the source is genuinely ambiguous with a provenance note; the translator transcreates the passages where meaning, idiom, or culture can't be machined. **Verify** — a QA agent runs terminology, number/date/placeholder, and consistency checks and back-translates risky segments for the human to inspect; the translator does the final reading for nuance, bias, and register. No segment ships unread by a human for any consequential text. **Ship/learn** — final, human-approved segments and term decisions write back to the shared TM/termbase as CID nodes, so the next project inherits a vetted, citable memory instead of re-litigating word choices.

Artefacts/tools: content-addressed translation memory and termbase, fuzzy-match pre-translation, automated terminology/QA checks, back-translation review, locale style guide. Quality gate: terminology and formatting checks pass automatically; ambiguities are flagged with provenance; a human signs off on register and meaning before delivery.

Human keeps owning: cultural and contextual judgement, voice and nuance, the call on what cannot be literally translated, and accountability for fidelity to the author's intent.

Business & Management

Product Manager

Essence: the PM is the captain of a decision, not the author of a spec — agents draft, the PM adjudicates trade-offs.

The loop runs continuously rather than per-sprint. **Intake:** support tickets, session analytics, sales call transcripts and churn signals flow into a shared content-addressed memory graph, where each insight is a CID node ("users abandon at step 3 of onboarding") linked to its evidence. A research agent clusters and deduplicates these against existing nodes, surfacing the five themes that moved this week. **Plan:** the PM frames the problem; a discovery agent reads the strategy node, the metric tree, and prior bets to draft 2–3 solution options, each citing the customer evidence and a rough effort estimate pulled from engineering's history. The PM kills, merges, or sharpens these — this is the judgement step, not delegated. **Do:** the agent writes the PRD, acceptance criteria, and analytics events as linked nodes; design and eng agents read the same graph so the spec, mocks, and instrumentation never drift. **Verify:** the quality gate is a pre-mortem the PM runs against the draft — does every claimed user need cite a real evidence node? Are success metrics instrumented before launch? An agent flags unsupported assertions and missing event tracking. **Ship/learn:** at release the agent wires the experiment, then writes the result back into the graph as a new CID node, so the next bet inherits the lesson instead of relearning it.

Artefacts: PRD and metric-tree nodes, an experiment ledger, a customer-evidence graph, a "decisions and why-not" log. Verification: every requirement traces to provenance-cited evidence; no metric ships uninstrumented.

What the human keeps owning: prioritization under scarcity, saying no, and the taste to know which customer pain is worth a quarter — accountability for the bet sits with the PM.

Project Manager

Essence: the PM owns commitments and risk; agents own the status-tracking toil that used to eat the role.

Intake: the plan of record lives as a content-addressed graph of tasks, dependencies, owners, and dates rather than a frozen Gantt chart. A tracking agent ingests commits, PR merges, calendar shifts, and standup notes, and reconciles real progress against the plan automatically — no one updates a status field by hand. **Plan:** when scope or a date enters, a planning agent recomputes the critical path, flags newly-created dependency conflicts, and proposes a re-sequenced schedule with the slack it consumed. The PM chooses among the trade-offs; the agent does not silently move a launch date. **Do:** the agent drafts the right nudges — a blocker that's three days stale, an owner double-booked across two critical tasks — and routes them, while the PM handles the human conversation behind the blocker. **Verify:** the quality gate is a daily risk reconciliation: every "green" status must trace to an actual signal (a merged PR, a passed gate), not a self-report, and the agent marks any status it cannot substantiate as unverified. **Ship/learn:** at milestone close the agent writes a retrospective node — estimate vs. actual, where slippage originated — into the shared graph, so future estimates are calibrated against this team's real velocity, not optimism.

Artefacts: a live dependency graph, a risk register with provenance, an auto-drafted status digest, a calibrated estimate history. Verification: status is signal-backed, not claimed; critical-path changes are explicit and human-approved.

What the human keeps owning: holding people to commitments, negotiating scope with stakeholders, and the judgement of which risk is acceptable versus which needs escalation today.

Marketer

Essence: the marketer sets brand, positioning, and the bet; agents run the production and measurement loop at a cadence no human team could.

Intake: brand guidelines, approved claims, ICP definitions, and past campaign results live as CID nodes in a shared brand memory — one source of truth instead of a dozen scattered prompt files and brand decks. A listening agent feeds in competitor moves, search trends, and channel performance. **Plan:** the marketer sets the quarter's positioning and the one message that matters; a strategy agent proposes campaigns and channel mixes, each citing which audience-evidence node and which prior result it's reasoning from. The human picks

the bet — taste and brand fit are not delegable. **Do:** content agents draft variants for each channel that read the brand node directly, so voice stays consistent; every generated claim is checked against the approved-claims node, and unverifiable or non-compliant copy is blocked before a human ever reviews it. **Verify:** the quality gate is twofold — a compliance/legal pass (no claim without a citation to an approved fact) and a brand-fit review the marketer signs off on. Performance creative ships into controlled experiments, never a blind blast.

Ship/learn: results write back as nodes — this hook beat that one for this segment — so the next campaign starts from evidence, and the brand graph compounds instead of resetting each quarter.

Artefacts: a brand-and-claims memory graph, a campaign experiment ledger, channel-attribution nodes, a creative-performance library. Verification: claims are provenance-cited and compliance-gated; spend follows measured lift, not vibes.

What the human keeps owning: brand taste, the positioning bet, ethical lines on persuasion, and accountability for what the brand says in public.

HR Manager

Essence: the HR manager owns fairness, trust, and judgement on people; agents handle the consistency and admin that make fairness scalable.

Intake: policies, role scorecards, comp bands, and prior decisions live as content-addressed nodes — so "how we handle X" is one citable source, not tribal memory. Employee questions, candidate pipelines, and case intake flow in. A triage agent routes the routine (PTO policy, benefits lookup) by answering from the policy graph with a citation, and escalates anything sensitive to a human untouched. **Plan:** for hiring, an agent drafts structured, role-anchored interview kits from the scorecard node; for a performance or grievance case, it assembles the relevant policy, precedent, and timeline so the HR manager walks in fully briefed. **Do:** the agent handles scheduling, drafts offer letters within approved bands, and prepares consistent rubrics — but every people-affecting decision is made by the human. **Verify:** the quality gate is a fairness and compliance check — the agent screens drafts and decisions for bias-laden language, inconsistent treatment versus precedent nodes, and legal red lines, flagging rather than deciding. No automated adverse action against a person, ever.

Ship/learn: decisions and their rationale write back as provenance nodes (with privacy controls), so policy stays consistent and auditable across managers and time.

Artefacts: a policy-and-precedent graph, structured interview kits, a comp-band reference, an auditable decision log. Verification: outputs are policy-cited and bias-screened; sensitive matters are human-only by design.

What the human keeps owning: empathy in hard conversations, judgement on edge cases people don't fit into, confidentiality, and full accountability for every decision that affects a person's livelihood.

Management Consultant

Essence: the consultant owns the hypothesis, the client relationship, and the recommendation; agents collapse the research-and-deck grind into hours.

Intake: the engagement's facts — interviews, client data, market figures, the issue tree — accumulate as CID nodes in a shared engagement memory, each finding linked to its source. A research agent runs desk research and benchmarks, writing every claim back with a citation, so the fact base is auditable rather than buried in someone's drive. **Plan:** the consultant frames the hypothesis and structures the issue tree (this is the craft); an analysis agent stress-tests it, proposing which analyses would confirm or break each branch and flagging where evidence is thin. **Do:** modeling agents build the financial model and scenario analysis from the data nodes; a deck agent drafts the storyline as MECE, pyramid-principle slides where each assertion cites its supporting node. The consultant edits the argument, not the formatting. **Verify:** the quality gate is a partner-grade review plus an automated provenance check — every number on every slide traces to a source node, and the agent flags any "so-what" that isn't backed by data. A red-team agent argues the opposite case to

surface where the recommendation is fragile. **Ship/learn:** the engagement graph (sanitized) feeds the firm's knowledge base, so the next team's intake starts from real precedent instead of a blank page.

Artefacts: an issue-tree and fact-base graph, a sourced financial model, provenance-cited slides, a red-team memo. Verification: claim-to-source traceability on every page; conclusions survive an adversarial challenge.

What the human keeps owning: the framing, the judgement call the data underdetermines, client trust, and accountability for the recommendation the client will bet on.

Public & Social

Social worker

Essence: spend the hour with the person, not the paperwork — agents carry the casework memory so judgement stays human.

The loop runs case-by-case. **Intake:** during or right after a visit, an agent transcribes consent-flagged notes, structures them against the statutory assessment framework, and surfaces the client's history from a shared, content-addressed case graph — prior interventions, household members, safeguarding flags, agency contacts — each as a provenance-cited node rather than a buried PDF. **Plan:** the worker (captain) sets goals and risk thresholds; an agent drafts the care/support plan, checks eligibility against current benefit and statutory rules, and flags conflicts (e.g. a court date colliding with a placement review). **Do:** agents handle referrals, appointment booking, and chase-ups across housing, health and education, writing every action back to the case node so the next worker sees an unbroken chain. **Verify:** a quality gate confirms statutory deadlines met, consent recorded, and that every risk claim cites its source observation — no unsourced "client appears stable." A supervisor agent flags drift from the care plan or escalating risk signals for human review. **Ship/learn:** outcomes feed back as structured nodes, so caseload patterns (what worked for similar families) become queryable across the team.

Artefacts/tools: structured case record, statutory assessment templates, the content-addressed case graph, consent ledger, referral tracker. Quality gate: deadline + consent + provenance checks, plus mandatory human sign-off on any risk-level change or removal decision.

The human keeps owning: the relationship, the read on a room, the safeguarding judgement, and accountability for every life-affecting decision. Agents must never auto-decide a removal, eligibility denial, or risk rating — they assemble evidence; the worker decides and signs.

Urban planner

Essence: the planner shapes the place and the trade-offs; agents do the spatial bookkeeping, scenario maths, and statutory cross-checks.

Intake: an agent ingests the site context — zoning, ownership, flood and heritage constraints, transport and utility capacity, demographic and environmental data — and writes each layer as a cited node in a shared spatial knowledge graph, so "the floodplain boundary" has one canonical, versioned source everyone references. **Plan:** the planner (architect) frames the brief and the values at stake (density vs. green space, affordability vs. viability). Agents generate massing and land-use scenarios, run quick daylight, traffic, and infrastructure-load models, and quantify each option's trade-offs against policy targets. **Do:** agents draft the planning report, viability assessment, and consultation materials; produce plain-language summaries and accessible visualisations for residents; and cluster public-consultation responses by theme with sentiment and representativeness flags. **Verify:** a compliance gate cross-checks the proposal against the local plan, national policy, and environmental statute, citing the specific clause for each judgement — flagging conflicts rather than asserting approval. Models are validated against historical outcomes where available. **Ship/learn:** the adopted

scheme and post-occupancy data return to the graph, so future schemes inherit what density, transport, or design choices actually delivered.

Artefacts/tools: GIS and scenario models, the spatial knowledge graph, viability and environmental assessment templates, consultation analysis, accessible visualisations. Quality gate: statutory-compliance cross-check with clause-level citations, model-validation provenance, and equalities/environmental impact review.

The human keeps owning: the vision for the place, the contested value trade-offs, the political and democratic legitimacy of the choice, and accountability to residents. Agents inform; the planner and elected decision-makers weigh and answer.

Policy analyst

Essence: the analyst owns the question, the framing, and the recommendation; agents own the evidence assembly and the traceable chain from data to claim.

Intake: an agent gathers the evidence base — legislation, prior evaluations, academic literature, administrative datasets, comparable jurisdictions — and writes each source as a content-addressed node with full provenance, so every later claim links back to a verifiable origin, not a half-remembered figure. **Plan:** the analyst (captain) defines the policy question, options space, and evaluation criteria. Agents map the option set, identify affected groups, and assemble the cost, distributional, and feasibility evidence for each. **Do:** agents draft the options appraisal and impact assessment, run sensitivity checks on cost and uptake assumptions, model distributional effects across income or regional groups, and red-team the preferred option for unintended consequences and equity gaps. **Verify:** a citation gate rejects any claim without a traceable source; a second agent adversarially stress-tests assumptions and checks that the modelling matches the stated method. Confidence and uncertainty are stated explicitly, never laundered into false precision. **Ship/learn:** the briefing ships with a transparent evidence appendix; once implemented, evaluation data returns to the graph, so the next analyst sees what the forecast got right and wrong.

Artefacts/tools: options appraisal, impact and distributional models, the cited evidence graph, the policy brief with provenance appendix, red-team log. Quality gate: every claim source-linked, assumptions sensitivity-tested, distributional and equality impacts assessed, uncertainty stated.

The human keeps owning: the framing of the problem, the value judgement in weighing options, the political-feasibility read, and accountability for the recommendation. Agents must surface trade-offs, not pre-decide them; the analyst signs the advice.

Journalist

Essence: the journalist owns the story, the sourcing, and the ethics; agents accelerate research and fact-checking without ever inventing a fact.

Intake: an agent monitors beats, documents, and datasets, building a content-addressed evidence graph where every fact links to its source — a document page, a recording timestamp, a named record — so nothing rests on an unattributable memory. **Plan:** the journalist (captain) decides the story, angle, and public-interest justification. Agents map what's known, what's contested, and what's missing; suggest documents to request and people to approach; and timeline the events. **Do:** agents transcribe interviews, cross-reference claims across sources, search filings and FOI releases, and draft structured background — but the reporter does the interviewing, the off-record judgement, and the writing voice. Every agent-surfaced fact carries its provenance node into the draft. **Verify:** a hard fact-check gate confirms each factual claim has at least one independent, cited source; an adversarial agent flags single-sourced assertions, loaded framing, and any sentence the evidence doesn't support. Quotes are checked against recordings; AI is never a source, only a finder. **Ship/learn:** the published piece retains an internal provenance trail for corrections and legal defence; source reliability updates the graph for future stories.

Artefacts/tools: the evidence/provenance graph, transcript and document search, fact-check ledger, source-protection vault, draft with inline citations. Quality gate: per-claim sourcing check, single-source flagging, quote-to-recording verification, legal and ethics review.

The human keeps owning: the editorial judgement, source relationships and protection, the public-interest call, and full accountability for what's published. Agents never publish, never decide what's newsworthy, and never stand in as a source.

NGO / programme coordinator

Essence: the coordinator owns the mission, the relationships, and the duty to beneficiaries; agents run the operational and reporting machinery against one shared source of truth.

Intake: an agent consolidates the programme picture — beneficiary reach, partner activity, budget burn, donor commitments, risk and safeguarding logs — into a content-addressed programme graph, so "Q2 reach in region X" has one canonical, cited figure instead of five conflicting spreadsheets. **Plan:** the coordinator (architect) sets objectives, theory of change, and risk appetite. Agents draft the logframe and workplan, map activities to indicators and budget lines, and flag where targets, funding, and capacity don't add up. **Do:** agents track activities against milestones, reconcile spend to budget, draft donor and board reports tailored to each funder's template, and turn monitoring data into honest dashboards. Partner submissions write straight back to the graph with provenance, ending the reporting-season scramble. **Verify:** a gate checks indicator data against source records, flags figures that can't be substantiated, and confirms donor-compliance and safeguarding requirements are met before anything ships — no inflated reach, no orphaned claims. **Ship/learn:** results and lessons return to the graph as structured nodes, so the next proposal cites real, traceable outcomes and the programme adapts on evidence.

Artefacts/tools: logframe/theory of change, the programme knowledge graph, budget-vs-actuals tracker, donor-report templates, monitoring dashboards, safeguarding and risk logs. Quality gate: indicator-to-source verification, donor-compliance and safeguarding checks, anti-inflation flagging on all reported figures.

The human keeps owning: the mission and beneficiary duty, partner and donor trust, ethical judgement, and accountability for funds and outcomes. Agents assemble and check; the coordinator decides priorities and signs every external claim.

Service & Commerce

Chef / kitchen lead

Essence: the chef designs the dish and owns the pass; agents run the costing, prep math, and supply chain underneath.

Loop. **Intake:** POS sales, reservation covers, and supplier price feeds stream into a shared content-addressed kitchen graph where each recipe, sub-recipe, allergen profile, and supplier SKU is a CID node with provenance. **Plan:** a forecasting agent reads last year's same-week covers, weather, and local events to project covers per service; a menu-engineering agent flags low-margin/low-popularity dishes against current food-cost CIDs. The chef sets the concept, flavour direction, and non-negotiables (seasonality, signature plates). **Do:** a prep agent explodes the forecast into station-level mise-en-place lists and a consolidated purchase order; an ordering agent drafts supplier POs but holds them for sign-off. **Verify:** a quality gate cross-checks every plated recipe against the allergen/dietary graph (no dish ships with an unlabelled allergen), reconciles theoretical vs. actual food cost, and flags variance >3%. HACCP temperature and waste logs are captured as signed nodes. **Ship/learn:** post-service, actuals (waste, 86'd items, plate returns) write back to the recipe CIDs so next week's forecast is grounded in this kitchen's real data, not a generic prompt.

Artefacts/tools: versioned recipe spec sheets, costed menu matrix, station prep sheets, auto-drafted POs, HACCP log. Gate: allergen-completeness check + food-cost variance threshold + signed temp logs before the menu is "live."

Human ownership: taste, plating, and the final palate decision are the chef's; so is the duty of care on allergens and food safety. Agents compute; the chef decides what is good enough to leave the pass.

Real-estate agent

Essence: the agent owns the relationship and the negotiation; agents own the comps, paperwork, and follow-up cadence.

Loop. **Intake**: a new listing or buyer brief enters a shared property graph where each listing, comparable sale, zoning record, and client preference is a CID node with source provenance (land registry, MLS, valuation date). **Plan**: a valuation agent assembles a comparative market analysis citing each comp's CID and adjustment rationale; a buyer-match agent ranks inventory against the client's stated and revealed preferences. The agent sets pricing strategy and reads the human signals the data misses. **Do**: a content agent drafts the listing copy, schedules photography, and generates the disclosure pack from the property node; a scheduling agent coordinates viewings against everyone's calendars. **Verify**: a compliance gate checks the disclosure pack for jurisdiction-required fields, confirms every factual claim in marketing traces to a cited source node (no "walking distance" without a measured node), and flags fair-housing language risks. **Ship/learn**: offer outcomes, days-on-market, and price-to-list ratios write back to the graph, sharpening the next valuation.

Artefacts/tools: provenance-cited CMA, listing pack, disclosure checklist, offer comparison sheet, follow-up sequence. Gate: disclosure-completeness + claim-provenance + fair-housing language scan before publish, and a human read of every offer term before it goes to the client.

Human ownership: the agent owns judgement on price strategy, trust with the client, and the negotiation table. Accountability for advice and fiduciary duty stays human; the AI never signs.

Sales representative

Essence: the rep owns the conversation and the commitment; agents own research, sequencing, and CRM hygiene.

Loop. **Intake**: inbound leads, signals, and account news land in a shared account graph where each contact, opportunity, objection, and competitor mention is a CID node tied to its source. **Plan**: a research agent builds a pre-call brief — org chart, recent triggers, likely pain — every claim cited to its source node; a prioritisation agent ranks the pipeline by fit and momentum, not just stage age. The rep chooses which accounts deserve human time and sets the strategy. **Do**: a drafting agent writes personalised outreach grounded in the account node (no generic blast), proposes next-best-action, and updates the CRM automatically after each touch so notes are never stale. The rep runs the actual calls and demos. **Verify**: a gate checks that every quote/proposal traces to approved pricing and product CIDs (no off-policy discount or unverifiable claim ships), and that competitive comparisons cite real evidence. Forecast commits are reconciled against logged activity. **Ship/learn**: won/lost reasons and objection patterns write back as nodes, so the next brief reflects what actually closes here.

Artefacts/tools: cited pre-call briefs, CRM auto-logging, proposal/quote generator, sequence cadence, forecast roll-up. Gate: pricing-approval + claim-provenance + forecast-vs-activity reconciliation before a proposal or commit goes out.

Human ownership: the rep owns the relationship, the discovery instinct, the read of the room, and the integrity of every promise made to a buyer. The forecast is the rep's word, not the model's.

Customer-support lead

Essence: the lead owns escalations, policy, and team quality; agents own triage, drafting, and knowledge upkeep.

Loop. **Intake:** tickets arrive across channels and route into a shared support graph where each known issue, resolution, policy, and customer-history thread is a CID node with provenance back to the product docs or incident that defined it. **Plan:** a triage agent classifies severity, links the ticket to matching resolution nodes, and proposes an SLA; the lead sets priorities, staffing, and the line between automation and human handling. **Do:** a draft agent answers tier-1 tickets citing the exact policy/KB node it relied on, and prepares context-rich handoffs for human agents on complex cases. A macro agent keeps canned responses in sync with the live policy graph. **Verify:** a quality gate blocks any reply whose claim doesn't trace to a current KB node (kills hallucinated policy), checks tone and compliance (refunds, privacy/GDPR commitments), and samples resolved tickets for CSAT and accuracy. **Ship/learn:** recurring unanswered questions surface as gaps that auto-open KB-authoring tasks; resolved novel cases write new nodes, so the knowledge graph self-heals instead of rotting in scattered docs.

Artefacts/tools: provenance-cited reply drafts, live KB graph, escalation runbooks, SLA dashboard, CSAT/accuracy sampling. Gate: claim-must-cite-current-KB + policy/privacy compliance scan + human review on refunds and escalations before send.

Human ownership: the lead owns judgement on edge cases, empathy under pressure, when to break policy for a customer, and accountability for what the team promises. Agents draft; the lead owns the relationship and the rules.

Recruiter

Essence: the recruiter owns the human read and the hiring-manager relationship; agents own sourcing, scheduling, and consistency.

Loop. **Intake:** an open role and its scorecard enter a shared hiring graph where each candidate, requirement, interview signal, and feedback note is a CID node with provenance (who observed it, when). **Plan:** a sourcing agent surfaces candidates matched to the scorecard, citing which evidence backs each match (no opaque "good fit" scores); the recruiter and hiring manager agree the must-haves and the bar. **Do:** a scheduling agent coordinates panels across calendars; a comms agent drafts personalised, accurate outreach grounded in the role node; an interview-kit agent prepares structured guides so every candidate is assessed on the same rubric. **Verify:** a fairness gate audits the funnel for adverse-impact and screens job copy and rejections for biased or non-compliant language; it also blocks any claim about a candidate that isn't backed by a logged interview node (decisions trace to evidence, not vibes). Structured scorecards are reconciled before debrief. **Ship/learn:** hire outcomes and quality-of-hire at 6 months write back, calibrating which signals actually predict success here rather than recycled folklore.

Artefacts/tools: scorecard, evidence-cited candidate shortlist, structured interview kits, scheduling automation, funnel/adverse-impact dashboard. Gate: structured-rubric completeness + adverse-impact check + decision-must-cite-evidence before an offer or rejection.

Human ownership: the recruiter owns the human judgement of motivation and culture, the candidate experience, the hiring-manager partnership, and accountability for fair, defensible decisions. Agents surface and structure; people choose who to hire.