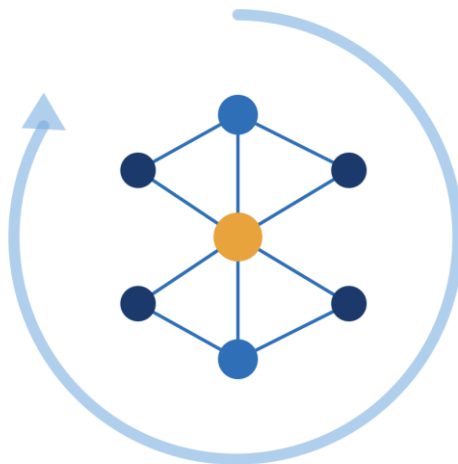


State of the art

AI applied to agile work

Reference map of current professional knowledge



Updated June 2026

Scrum Manager® · Skill Arena

About this document

This document is a curated map of the knowledge, practices and models in professional use as of its publication date for bringing artificial intelligence into agile work with sound judgement. For each entry, the map places it in context, indicates its degree of adoption in current professional practice and points to where to find reference information.

Use it as an observatory and reference point to check whether the professional knowledge you rely on is aligned with current practice and at the leading edge.

It is complemented by an in-depth companion guide that develops each concept and by the training and assessment platform in Skill Arena. In the [AI applied to agile work](#) area you can test your level of knowledge and, if you pass, obtain a diploma that provides curricular accreditation of professional competence and currency in this area.



State of knowledge

Knowledge about AI applied to agile work is a field in full ferment that evolves extremely fast: practices and tools appear, mature or become obsolete within months, and only a part of the fundamentals can be considered settled. Throughout the document, the labels (ESTABLISHED, CONSOLIDATING, EMERGING) help identify the maturity of each concept:

ESTABLISHED settled consensus; knowledge taken as required.

CONSOLIDATING gaining adoption fast; not yet universal but already relevant.

EMERGING recent frontier; high relevance and high volatility.

Block A — Fundamentals of interacting with AI

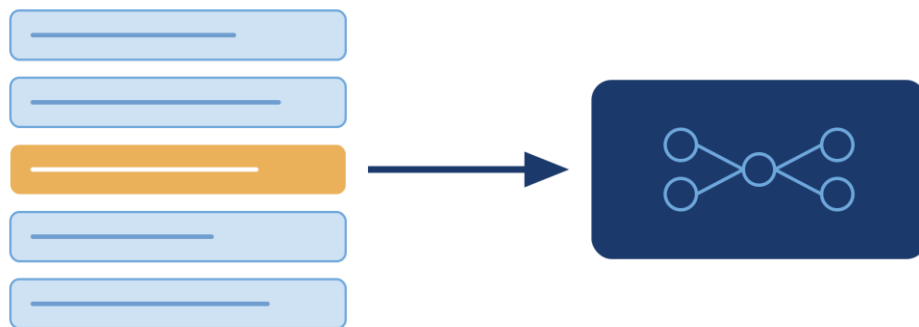
Cross-cutting competencies, valid for any role working with AI, whatever the discipline. They are the ground on which everything else rests.

Context engineering **CONSOLIDATING** · new since 2026

Deliberately designing what information the model receives and in what order —not just how the one-off instruction is worded—. The base competency has shifted from writing the perfect prompt to governing the context that surrounds the request.

Why it is here now. The value frontier has moved from the sentence to the management of context. Well-constructed prompting is still the foundation, but at the scale of a team and of sustained work it no longer suffices on its own; what is distinctive today is deciding what enters the model's window, what is left out and how it is ordered.

Where to look. [Anthropic — Effective context engineering](#) · [Anthropic — Prompt engineering documentation](#)



Context engineering governs what information enters the model, in what order and what is left out.

Structured, reusable prompting **ESTABLISHED**

Specifying a request with the components that make its result predictable and verifiable: a goal with observable success criteria, minimum sufficient context, output format and explicit constraints. Distinguishing structural instructions (what form the response should take) from qualitative ones (“be more thorough”), which do not produce consistency.

Why it is here now. It is the foundation of everything else: without a well-framed prompt, verification, refinement or work with agents inherit inconsistent results. It remains a required competency, though today it is understood as the base case within context engineering.

Where to look. [OpenAI — Prompt engineering guide](#) · [Anthropic — Prompt engineering](#)

Verification and reliability of responses **ESTABLISHED**

Critically assessing what a model produces before acting on it: separating facts from inferences and from recommendations, asking for the assumptions and evidence that support a claim, and detecting “unwarranted precision” (specific figures or certainties with no backing). The fluency of the text does not guarantee its accuracy.

Why it is here now. It is the most timeless competency in the area and, far from losing weight, it gains relevance with agents: the more work the professional delegates, the more critical it is to know how to verify what the system delivers. The capacity for judgement remains the irreplaceable human value.

Where to look. [NIST AI Risk Management Framework](#) · [Anthropic — Reducing hallucinations](#)



Verify before acting: separate facts from inferences and detect unwarranted precision.

Trust boundary: instructions versus data, privacy and secrets ESTABLISHED

Treating all external content as untrusted data by default, separating it from instructions to prevent the injection of malicious instructions (prompt injection); and applying minimisation and effective anonymisation before sharing information with a model, without relying on “confidential” labels or on system promises.

Why it is here now. What was a defensive good practice in conversational use becomes a central security competency when agents execute actions on real systems: the cost of an injected instruction or a data leak stops being a bad response and becomes an unintended action.

Where to look. [OWASP Top 10 for LLM Applications](#) · [EU AI Act](#)

Block B — Working with AI agents

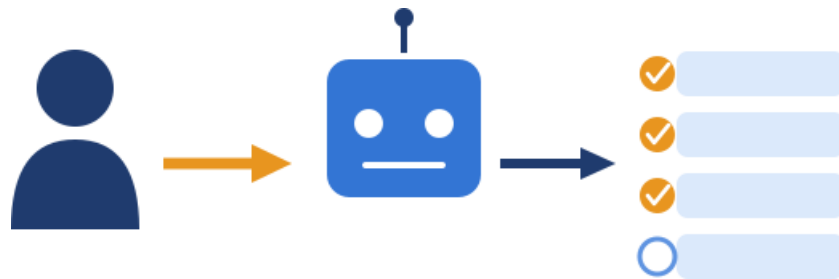
The recent paradigm shift: moving from using AI to think better (to respond) to delegating work to it (to execute). It is the most fast-moving area of this map and, therefore, the one revised most often.

From chat to agent: delegating verifiable work CONSOLIDATING · new since 2026

Understanding the difference between an assistant that answers questions and an agent that completes tasks autonomously, and knowing when to delegate and when not to. The agile professional's role shifts from “verifying answers” to “designing what work is delegated, with what limits and with what acceptance criteria”.

Why it is here now. Agents have moved from demonstration to everyday working tool. This reshapes agile ceremonies and artefacts: the unit of work stops being “the task I do” and comes to include “the task I orchestrate and supervise”. Knowing how to delegate well is the new differentiating competency.

Where to look. [Anthropic — Building effective agents](#) · [Anthropic — Effective context engineering](#)



From chat to agent: the professional moves from executing to delegating work with acceptance criteria.

Designing human oversight (human-in-the-loop) CONSOLIDATING · new since 2026

Deciding where an AI-assisted flow should stop so that a person can review, approve or redirect before continuing: pause points, approval gates and traceability of decisions. The agile professional designs where the human intervenes, rather than executing each step by hand.

Why it is here now. It is the necessary counterpart of delegation, and it also has a growing compliance dimension. The obligations for high-risk AI systems under the EU AI Act start to apply in August 2026, and documented human oversight moves from good practice to requirement in many contexts.

Where to look. [EU AI Act — implementation timeline](#) · [NIST AI Risk Management Framework](#)



Human oversight introduces pause and approval points into the AI-assisted flow.

Agentic ecosystem standards (MCP, AGENTS.md) EMERGING · new since 2026

Knowing the emerging standards that stabilise work with agents: the Model Context Protocol (MCP) to connect models with tools and data, and conventions such as AGENTS.md to give agents reusable

instructions. It is not necessary to implement them, but it is necessary to understand what they solve and why they are being adopted.

Why it is here now. The field is standardising fast: MCP and AGENTS.md have been folded into an open industry foundation backed by the main vendors. Knowing them places the professional in front of the direction the ecosystem is taking, even if their concrete application is more relevant to technical roles.

Where to look. [Model Context Protocol](#) · [AGENTS.md](#)

Block C — Application to the agile work cycle

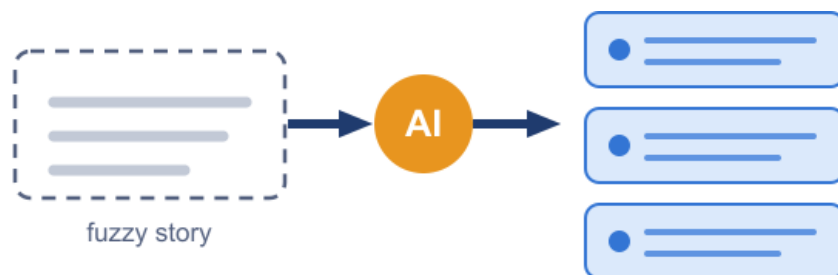
Use cases for AI in the workflow. The first two apply to any agile team role; those marked as “development roles” are especially relevant to those who work on building software.

AI in the product and backlog flow **ESTABLISHED**

Using AI as support to detect ambiguity in user stories and generate refinement questions, propose splits (slicing) that keep demonstrable value, draft verifiable acceptance criteria (including negative scenarios and edge cases) and formulate hypotheses and experiments with metrics and guardrails. The suggestions are a starting point for the team's conversation, not requirements.

Why it is here now. It is the most everyday and cross-cutting use of AI in agility, applicable to product owners, scrum masters and any role that works with the backlog. Its value lies in reducing ambiguity and rework without introducing bureaucracy, keeping the decision with the team.

Where to look. [Scrum.org — The Scrum Master and AI](#) · [Scrum Manager — Scrum in the age of AI](#)



AI reduces backlog ambiguity: from the fuzzy story to items with verifiable criteria.

AI in facilitation and team dynamics **CONSOLIDATING** · new since 2026

Relying on AI to prepare facilitation, synthesise the material from discovery sessions, qualitatively analyse the outcome of retrospectives or structure difficult conversations, without the tool replacing human interaction or the reading of the team's climate, which remain the sole competency of the facilitator.

Why it is here now. It is the use most missing from materials on AI and agility, and the most directly useful for scrum masters and agile coaches who do not work with code. Well framed, it frees up preparation time and enriches the analysis; badly framed, it depersonalises what only works between people.

Where to look. [Scrum Manager — Scrum in the age of AI](#) · [Anthropic — Building effective agents](#)



AI synthesises and prepares, without replacing the human interaction that sustains facilitation.

AI in quality and delivery — development roles **ESTABLISHED**

Generating and reviewing test cases aligned with acceptance criteria and risks, ensuring coverage of negative scenarios and edge cases beyond the “happy path”; and using AI as support in code review with a focus on risks (input validation, error handling, permissions, side effects), always keeping the reviewer's technical judgement and without delegating the final decision to the AI.

Why it is here now. It remains a high-value use for technical teams, where AI extends the scope of the review and the coverage of testing. It is marked as specific to development roles because its direct application requires working with code, unlike the cross-cutting uses of the previous blocks.

Where to look. [OWASP Top 10 for LLM Applications](#) · [Anthropic — Building effective agents](#)

What to watch in the next revision

Signals of change the editorial team anticipates for the next edition of this map:

- The effective application of the EU AI Act obligations from August 2026 and their translation into concrete human-oversight practices.
- The consolidation or otherwise of the agentic ecosystem standards (MCP, AGENTS.md): whether they move from emerging to consolidating.
- The maturing of AI-assisted facilitation, today consolidating, towards settled practices with judgement of their own.
- The possible emergence of new competencies around the evaluation and governance of agents in non-technical teams.

A note on the references

The links included were available and verified as of the update date. Given the pace of change in this area, some may change; each revision of the map also updates its references.

© 2026 Scrum Manager®. This work is published under a Creative Commons Attribution-NonCommercial 4.0 International licence (CC BY-NC 4.0). Official Scrum Manager trainers and centres are licensed under the terms of CC BY 4.0 for their training activity.