



Agile Integrity with AI

By

 ***Innovation in Progress***

James Felgate

Agile Integrity with AI

Copyright and Licence

© Innovation in Progress. All rights reserved.

The Agile Integrity materials, including the Agile Integrity Diagnostic, AI–Agile Integrity Health Check, accompanying guides, PDFs, frameworks, questions, scoring models, graphics, templates and related content (together, the “**Materials**”) are protected by copyright and other intellectual property rights. All intellectual property rights in the Materials remain with Innovation in Progress or their respective rights holders.

By accessing, purchasing or using the Materials, you are granted a **limited, non-exclusive, non-transferable, non-sublicensable licence** to use the Materials subject to the terms below. This licence does not transfer ownership of, or any intellectual property rights in, the Materials to you.

Permitted use

You may:

- use the Materials for your own personal or internal organisational purposes;
- use the Materials to support your own professional work, consulting, coaching, training, facilitation, advisory or transformation activities;
- discuss and apply the concepts, principles and ideas contained in the Materials in your own work;
- use the Diagnostic with your own clients, teams or organisation as part of your professional services, provided that the Diagnostic remains subject to these licence terms;
- incorporate insights derived from using the Materials into your own professional advice and deliverables, provided that you do not reproduce or redistribute the Materials themselves; and
- make reasonable copies of the Materials where necessary for an authorised use under this licence.

Prohibited use

Unless Innovation in Progress gives you prior written permission, you must not:

- sell, resell, rent, lease, licence, sublicense, distribute or otherwise commercially exploit the Materials as standalone products;
- provide the Materials to another person or organisation for the purpose of enabling that person or organisation to obtain the Materials without purchasing or otherwise obtaining an appropriate licence;
- reproduce, publish, upload, post, share or distribute the Materials, or substantial portions of them, on websites, intranets, repositories, learning platforms, document libraries, social media, marketplaces or other publicly or commercially accessible platforms;
- remove, obscure or alter copyright notices, attribution, branding or other proprietary notices contained in the Materials;
- modify, adapt, repackage, white-label or create a substantially similar derivative product from the Materials for sale, licensing or distribution;
- use the Materials to create or market a competing diagnostic, assessment, framework, course, publication, certification or other commercial product that substantially reproduces the structure, questions, scoring methodology, wording or distinctive intellectual property of the Materials;
- grant any third party rights to use the Materials that are broader than the rights granted to you under this licence; or
- represent or imply that you own, authored, developed or have exclusive rights in the Materials.

Professional and client use

This licence is intentionally designed to permit professionals to **use the Materials to supplement and enhance their existing offerings**.

You may therefore use the Materials as part of your own consulting, coaching, advisory, transformation or training engagements, including discussing the findings of the Diagnostic with a client and using those findings to inform your professional recommendations.

However, the Materials themselves remain licensed intellectual property. Your professional service does not give your client ownership of the Materials or a right to redistribute them.

Where you provide a client with a copy of any Material, that copy must be provided solely for the client's use in connection with the authorised engagement and must remain subject to these licence terms. You must not sell or separately charge your client for the Materials as standalone intellectual property unless Innovation in Progress has expressly authorised this in writing.

No transfer of intellectual property

Nothing in these terms transfers, assigns or sells any copyright, trademark, database right, design right, know-how or other intellectual property right in the Materials.

All rights not expressly granted under this licence are reserved by Innovation in Progress.

Third-party and mandatory rights

Nothing in these terms is intended to exclude, restrict or override any right or remedy that cannot lawfully be excluded or restricted under applicable law.

Where a mandatory right or legal exception applies to your use of the Materials, these terms will operate subject to that right or exception.

Termination

This licence remains in effect for as long as you comply with these terms.

If you materially breach these terms, Innovation in Progress may terminate the licence by written notice where permitted by applicable law. Upon termination, you must cease using and distributing the Materials and, where reasonably practicable, delete or destroy copies in your possession or control, except where retention is required by law.

Termination does not affect any rights or remedies that accrued before termination.

No warranty or professional advice

The Materials are provided as educational and professional resources. They do not constitute legal, regulatory, financial, technical, employment, security or other specialist professional advice.

Innovation in Progress does not warrant that the Materials are suitable for every organisation, industry, jurisdiction or circumstance. Users remain responsible for applying appropriate professional judgement and obtaining specialist advice where required.

Governing law

Where legally permissible, these licence terms shall be governed by the laws of **England and Wales**, and the courts of England and Wales shall have exclusive jurisdiction over disputes arising from or relating to these terms.

Nothing in this provision is intended to deprive a consumer of any mandatory protection or right that cannot lawfully be excluded under the law applicable to that consumer.

Acceptance

By accessing, downloading, purchasing or using the Materials, you acknowledge that you have read and agree to these licence terms.

For information or questions, contact info@innovationinprogress.co.uk

How to use this package

This version is compatible with version 2 of the diagnostic.

1. Read the Executive Overview and Agile Integrity model.
2. Complete the workbook individually or as a team.
3. Review completion before interpreting scores.
4. Examine the two headline risks.
5. Review the lowest domain scores.
6. Discuss comments and supporting evidence.
7. Select no more than three priorities.
8. Create improvement experiments.
9. Repeat the assessment after 8–12 weeks.

Contents

Copyright and Licence	1
Permitted use	1
Prohibited use	1
Professional and client use	1
No transfer of intellectual property	2
Third-party and mandatory rights	2
Termination	2
No warranty or professional advice	2
Governing law	2
Acceptance	2
How to use this package	3
Executive Overview	5
The Agile Integrity Model	6
Introduction	7
What is Agile Integrity?	7
How AI Undermines Agility	7
Organisation Failure Modes	7
How AI Undermines Agile Events and Practices	10
Product discovery and product management	10
Product discovery and refinement	10

Product backlog management	11
User stories	12
Product management.....	14
Planning, coordination and feedback	15
Estimation and velocity	15
Sprint Planning.....	17
The Daily Scrum	18
Sprint Review.....	19
Sprint Retrospective	21
The two-week Sprint	22
Behaviour and quality	24
BDD - Behaviour-Driven Development	24
Definition of Done	25
Technical ownership.....	27
Pairing and collective ownership	27
Technical excellence and refactoring	28
Practical Agile Integrity Test	30
The AI-Agile Integrity Health Check Diagnostic	30
Interpreting the results	31
The distinction between integrity scores, risk scores and confidence	31
Coaching guidance	32
Conclusion.....	33
References	34

Executive Overview

Artificial intelligence is changing how organisations research customer needs, plan work, coordinate teams and build products. It can reduce administrative effort, accelerate analysis and support technical delivery. However, faster production does not automatically create faster learning, better customer outcomes or greater organisational agility.

Agile Integrity is the disciplined application of agile values, principles and practices to achieve valuable outcomes. It depends on direct evidence, customer collaboration, empirical learning, shared understanding, technical excellence, psychological safety and clear human accountability.

The principal risk is the way it can change organisational behaviour. AI may generate persuasive customer insights without customer contact, detailed backlogs without shared understanding, optimised plans without team ownership and code or tests faster than teams can responsibly review and validate. In short, it can create an appearance of confidence in assumptions, forecasts and recommendations where there is little or no evidence.

This guide examines how AI can strengthen or undermine agile events and practices, including product discovery, backlog management, Sprint Planning, the Daily Scrum, Sprint Reviews, Retrospectives, BDD, estimation, the Definition of Done, pairing, technical excellence and delivery cadence.

The accompanying **AI-Agile Integrity Health Check** assesses six areas:

- Product and Evidence Integrity
- Delivery and Learning Integrity
- Collaboration Integrity
- Technical Integrity
- Event and Cadence Integrity
- Governance and Accountability Integrity

It also identifies two cross-cutting risks:

- **Speed Without Learning:** production increases without a corresponding improvement in validation, feedback or outcome learning.
- **Synthetic Certainty:** AI-generated analysis or recommendations are treated with greater confidence than the supporting evidence warrants.

The purpose of this package is to help organisations use AI to reduce waste and improve decision support without weakening Agile Integrity.

The central test is simple:

“Does AI strengthen evidence, collaboration, learning, understanding and accountability—or merely increase the amount and apparent certainty of the work being produced?”

The Agile Integrity Model

These concepts are used in the diagnostic and help the reader better understand this guide.

Agile Integrity

Agile values, principles and practices are applied in ways that preserve customer focus, empirical learning, collaboration, technical excellence and clear human accountability.

Product and Evidence Integrity

Product decisions remain connected to observable customer, operational and commercial evidence.

Delivery and Learning Integrity

AI improves end-to-end learning and validated outcomes, rather than simply increasing production.

Collaboration Integrity

AI supports shared understanding without replacing necessary cross-functional conversations.

Technical Integrity

Teams understand, review, maintain and remain accountable for AI-assisted technical work.

Event and Cadence Integrity

Agile events and delivery rhythms continue to create purposeful inspection, adaptation, coordination and feedback.

Governance and Accountability Integrity

AI use remains transparent, appropriately controlled and subject to clear human ownership.

Introduction

Many organisations are adopting AI in pursuit of greater productivity. Yet increased output does not necessarily lead to improved customer value, shorter end-to-end lead times or better business outcomes. It can even be more expensive given the increased consumption and cost of AI tokens. AI may accelerate individual activities without improving outcomes.

This creates a risk that organisations optimise visible production (outputs) while weakening the conditions that make agility effective (outcomes). Generated content can be mistaken for evidence, forecasts for commitments, automation for accountability and detailed documentation for shared understanding.

This guide explores those risks across common agile events and practices. It considers where AI can provide legitimate assistance and how integrity may be weakened. It also identifies the warning signs, and the standards teams should preserve.

The aim is not to argue against AI. It is to help leaders, product professionals, coaches and delivery teams adopt it in ways that strengthen the outcome.

What is Agile Integrity?

Agile Integrity is the disciplined application of agile values, first principles and practices to deliver valuable outcomes. It is weakened when teams adopt the behaviours, practices or ceremonies of agile while bypassing those that make agility effective: customer collaboration, empirical learning, transparency, technical excellence and shared accountability. The agile community recognises these patterns in terms such as ‘fake agile’, ‘fragile’ and ‘water-scrum-fall’. The real opportunity is not simply to apply AI to agile work, but to use AI in ways that strengthen these foundations rather than obscure or replace them.

How AI Undermines Agility

To better understand the risks of adopting AI without guardrails, it helps to understand the risk presented to Agile Integrity.

Organisation Failure Modes

The impact of AI adoption on the wider organisation can be seen in these behaviours.

1. Output acceleration disguised as agility

The demand for more stories, code, tests and documentation drives teams to adopt AI but customer outcomes, lead time and adaptability do not improve. Value delivery isn't just a by-product of faster outputs.

2. Machine-generated certainty

Uncertainty is hidden behind confident explanations, ‘precise’ forecasts and elaborate plans. Agility requires exposing uncertainty early to improve decision making, design and deliver resilient, reliable and desirable products.

3. Surveillance replaces transparency

Transparency is voluntary visibility that helps a team coordinate, solve problems and continuously improve. Surveillance is visibility imposed so management can judge and control the team, often through vanity metrics and poorly designed incentives. This undermines employee engagement and will diminish outcomes. AI makes it easy to confuse the transparency and surveillance.

4. Local optimisation

Each Agile role uses AI to maximise its own output for example:

- Product generates more requirements.
 - Requirements need to be refined and understood by the team. So more requirements that are more verbose means more effort is required to make better decisions.
- Development generates more code.
 - More code results in more code reviews, more tests, more oversight to align with architecture, quality and security standards. More code doesn’t equal better outcomes.
- Testing generates more tests.
 - More tests do not always introduce more quality and results in the overhead of test maintenance and ongoing alignment.
- Architecture generates more standards.
 - More standards result in higher overhead of review and alignment.
- Management generates more reporting.
 - More reports can lead to surveillance over transparency.

The end-to-end system becomes slower because every stage creates more demand for the next without necessarily introducing anything better. This creates bottlenecks and coordination overhead that wouldn’t otherwise have existed. It also creates a level of uncertainty and undermines the shared knowledge required to effectively sustain value delivery.

5. Accountability gaps

Decisions are described as “AI-supported”, making it unclear who is responsible when they fail. Time is wasted discussing who should make decisions and on what

authority. The confusion results in noisy decisions which leads to wider role uncertainty. This could lead to people avoiding responsibility or accountability.

6. Loss of productive friction

Some friction is waste e.g. items waiting for an approval which could be automated. Other friction is where understanding develops e.g. refinement of requirements.

Debate, clarification, negotiation, pairing, example mapping and retrospection may look slower than AI generated equivalents, but they create shared knowledge the powers effective teams. Removing all friction removes part of the learning system that is fundamental to sustainable, valuable outcomes.

7. Competence erosion

People become effective at prompting and reviewing polished outputs but less capable of:

- Diagnosing problems.
- Forming independent hypotheses.
- Communicating across disciplines.
- Understanding code and systems deeply.
- Making decisions with incomplete information.
- Ensuring Vision, Strategy and Architectural alignment.

How AI Undermines Agile Events and Practices

The following sections show how AI can legitimately support agile work while also weakening the conditions that make agility effective. Each practice is considered through the same six lenses so that the risks can be connected directly to the accompanying health-check diagnostic.

Product discovery and product management

Product discovery and refinement

How AI can help

AI can analyse research, identify patterns, draft hypotheses, expose assumptions and generate alternative solutions.

How integrity can be weakened

AI-generated insights can become a substitute for contact with customers and reduce the value of product discovery. A product manager may ask AI to:

- Summarise customer needs.
- Generate personas.
- Predict likely objections.
- Prioritise features.
- Write the business case.
- Produce the roadmap.

The team can then appear evidence-driven without acquiring much direct evidence. This creates synthetic discovery: internally plausible explanations that have not been tested against customer behaviour.

Integrity principles at risk

- Customer collaboration.
- Evidence-based product management.
- Empathy and contextual understanding.
- Testing assumptions rather than reinforcing them.
- Transparency about what is observed evidence and what is AI-generated interpretation.
- Human accountability for product choices.

Warning signs

- AI-generated personas are described as customer research.
- Product decisions cannot be traced to interviews, observed behaviour, operational data or experiments.

- Customer contact decreases while the volume of generated insight increases.
- AI summaries are reviewed, but the underlying research is not.
- Teams use phrases such as “the data says” when the output is actually a model-generated interpretation.

Better standard

AI should help teams identify questions, patterns and alternative explanations. It should not be treated as the source of customer evidence. Significant product decisions should be traceable to observable customer, operational or commercial evidence, and generated insights should remain labelled as hypotheses until validated. Human product decisions should also align with vision, strategy and architectural guardrails.

Relevant health-check domains and risks

- Primary domain: Product and Evidence Integrity.
- Secondary domains: Collaboration Integrity; Governance and Accountability Integrity.
- Headline risk: Synthetic Certainty.

Product backlog management

How AI can help

AI can quickly draft, organise and reformat backlog items. It can also identify duplication, suggest dependencies and produce alternative ways of expressing a need.

How integrity can be weakened

AI-generated backlog items can create four problems:

- Backlog becomes the plan. Items appear fully described in advance, undermining discover-build-learn-adapt cycles. **This encourages following a plan over embracing change.**
- Backlog inflation. The cost of creating an item approaches zero, but the cost of understanding, validating, building and maintaining it does not.
- False readiness. Detailed descriptions, acceptance criteria and technical notes create the appearance of readiness without shared understanding.
- Proxy product ownership. Stakeholders may treat the product owner as the operator of an AI prioritisation system rather than the accountable decision-maker.

Integrity principles at risk

- Responding to change over following a plan.
- Customer collaboration over detailed specification.
- Simplicity and maximising the amount of work not done.
- Shared understanding over document completeness.

- Product-owner accountability.
- Evidence-based ordering and prioritisation.
- Transparency about uncertainty and readiness.

Warning signs

- The backlog grows faster than items are validated, delivered or removed.
- Stories appear ready because they are detailed, but the team has not discussed them.
- Large numbers of items contain generic or repeated acceptance criteria.
- Stakeholders assume an item must be delivered because it exists in the backlog.
- AI-generated priority scores are accepted without product judgement.
- Refinement sessions focus on correcting generated text rather than understanding the problem.
- The product owner cannot explain the evidence or reasoning behind the ordering.

Better standard

A healthy AI-supported backlog should generally become smaller and more evidence-focused, not merely more detailed. Use AI to remove duplication, expose assumptions, identify missing evidence and suggest alternative slices. Backlog items should remain prompts for conversation and enter the backlog only when they support a current product decision, validated opportunity or deliberate learning objective.

Relevant health-check domains and risks

- Primary domain: Product and Evidence Integrity.
- Secondary domains: Delivery and Learning Integrity; Collaboration Integrity.
- Headline risks: Synthetic Certainty; Speed Without Learning.

User stories

How AI can help

AI can draft user stories quickly, use standard formats and suggest alternatives for the team to examine.

How integrity can be weakened

AI-generated user stories can reinforce a requirements-document mindset. They are commonly:

- Over-specified.
- Solution-led.
- Detached from actual users.
- Written in bulk.

- Accepted without conversation.
- Treated as interchangeable units of delivery.

This contradicts the original role of a story as a placeholder for discussion. A story generated in seconds may still require substantial human effort to determine whether it is desirable, coherent and valuable. AI therefore separates the appearance of backlog maturity from actual product knowledge.

Integrity principles at risk

- Conversation over specification.
- Customer collaboration.
- Shared understanding.
- Focus on user need rather than solution output.
- Simplicity.
- Progressive elaboration.
- Evidence-based product decisions.

Warning signs

- Stories are generated in batches before customer or team discussion.
- The user, need or expected outcome cannot be connected to real evidence.
- Acceptance criteria are adopted without challenge because they appear comprehensive.
- Different stories use generic personas or interchangeable user types.
- The team spends refinement time interpreting what the AI meant.
- Story completion is treated as proof that customer value was delivered.
- Nobody can identify the original conversation or observation behind the story.

Better standard

AI may draft a starting point, but a user story should remain an invitation to discuss a real need. The team should validate the user, desired outcome and supporting evidence before treating the story as delivery-ready. Generated detail should be reduced where it constrains learning or gives a false impression of certainty. A story is not ready because its wording is complete; it is ready when the relevant people share enough understanding to make a responsible delivery decision.

Relevant health-check domains and risks

- Primary domain: Product and Evidence Integrity.
- Secondary domains: Collaboration Integrity; Delivery and Learning Integrity.
- Headline risk: Synthetic Certainty.

Product management

How AI can help

AI can reduce the effort required to gather information, compare options, prepare analysis and communicate product thinking. It can assist with:

- Market scans.
- Competitor analysis.
- Roadmap options.
- User-story generation.
- Prioritisation scoring.
- Stakeholder summaries.
- Experiment design.
- Product metrics interpretation.

How integrity can be weakened

This can produce product-management theatre at extraordinary speed. AI can recommend, but it cannot carry organisational or moral accountability. If the product strategy fails to deliver results, AI cannot be held accountable.

A product manager loses integrity when they become the presenter of model-generated recommendations rather than the owner of explicit product judgement.

Integrity principles at risk

- Human accountability for product decisions.
- Evidence-based product management.
- Customer and stakeholder collaboration.
- Transparency about assumptions and uncertainty.
- Ethical judgement.
- Alignment with vision, strategy and organisational constraints.
- Outcome ownership.
- Willingness to say no and stop work.

Warning signs

- Roadmaps are produced more quickly but are not supported by stronger evidence.
- Product managers present AI recommendations without explaining their own judgement.
- Prioritisation scores replace discussion of trade-offs.
- Predicted benefits are treated as committed outcomes.
- Stakeholders ask the product manager to “run it through AI” rather than make a decision.
- Product decisions have no named human owner.

- Options are accepted without considering ethical, strategic or operational consequences.

Better standard

The real responsibilities remain human:

- Choosing which evidence to trust.
- Deciding whose needs matter.
- Resolving competing interests.
- Making ethical trade-offs.
- Understanding organisational constraints.
- Taking responsibility for investment decisions.
- Saying no.
- Deciding when there is insufficient evidence.
- Recognising that a commercially attractive option may still be harmful.

AI should expand the options available, identify assumptions and support analysis. The product manager remains responsible for interpreting evidence, making trade-offs, aligning decisions to strategy and owning the consequences. Significant decisions should record the evidence used, uncertainty, alternatives considered and named decision owner.

Relevant health-check domains and risks

- Primary domain: Product and Evidence Integrity.
- Secondary domains: Governance and Accountability Integrity; Delivery and Learning Integrity.
- Headline risk: Synthetic Certainty.

Planning, coordination and feedback

Estimation and velocity

How AI can help

AI can analyse historical flow data, identify patterns, highlight ageing work and generate probabilistic planning inputs. It can help teams examine previous assumptions and prepare scenarios for discussion.

How integrity can be weakened

Historical data may encode poor environments, excessive handoffs, technical debt, approval delays, unhealthy overtime, inflated estimates or local organisational politics. AI can reproduce those patterns and label them a forecast while failing to account adequately for current context, changed team capability or the uncertainty of the next item. In short, the probabilistic forecasts are statistically unsound and shouldn't be relied on.

AI also allows management to compare expected and actual performance with increasing granularity across teams. Estimation can then shift from supporting coordination to policing conformity, and velocity can mutate into digitally enabled Taylorism.

Integrity principles at risk

- Treating estimates as forecasts rather than promises.
- Respecting the team-relative nature of velocity.
- Avoiding individual productivity comparisons.
- Using uncertainty to guide decisions.
- Recognising variation in work complexity, context, skill and experience.
- Sustainable pace.
- Transparency about confidence and assumptions.
- Metrics used for learning rather than control.

Warning signs

- AI forecasts are presented as commitments or targets.
- Managers compare velocity or cycle time across teams.
- Individual contribution is inferred from ticket, code or defect data.
- Estimates are generated without team discussion.
- Forecast “precision” increases while forecast reliability is not reviewed.
- Teams adjust estimates to meet an AI-generated target.
- Historical data includes overtime or known dysfunction but is not contextualised.
- People are judged against predicted completion times.

Better standard

Use AI and statistical analysis to **support** forecasting, not to create commitments or assess individual performance. Forecasts should include assumptions, ranges and confidence and should be interpreted by people who understand the current context. Teams should use flow measures where useful while retaining human judgement about uncertainty, novelty and risk.

Relevant health-check domains and risks

- Primary domain: Delivery and Learning Integrity.
- Secondary domains: Governance and Accountability Integrity; Event and Cadence Integrity.
- Headline risks: Synthetic Certainty; Speed Without Learning.

Sprint Planning

How AI can help

AI can summarise readiness, dependencies, recent flow information and known risks. It can suggest alternative Sprint Goal wording and identify work that may not support the intended outcome.

How integrity can be weakened

The danger is that planning becomes an optimisation calculation rather than a team commitment to an outcome. **This reinforces outputs over outcomes.** The organisation may feed AI the historical velocity, team availability, estimates, dependency data and individual productivity metrics, then receive an “optimal” Sprint plan.

AI usually optimises for the measurable variable supplied to it. That may be utilisation, story completion or forecast accuracy rather than customer value, learning or sustainable pace.

- Team judgement and engagement are reduced.
- Collective ownership is weakened.
- Negotiation about uncertainty is bypassed.
- Technical risk receives less discussion.
- The team loses autonomy over how much work it can responsibly undertake.

Integrity principles at risk

- Team autonomy and self-management.
- Collective ownership of the Sprint Goal and plan.
- Focus on value and outcomes.
- Transparency about uncertainty.
- Sustainable pace.
- Human judgement about risk, capacity and dependencies.
- Adaptation rather than mechanical optimisation.
- Commitment created through participation rather than imposed allocation.

Warning signs

- The Sprint plan is generated before the team meets.
- Individual tasks are allocated automatically.
- Historical velocity is treated as a delivery commitment.
- The team is expected to accept AI-calculated capacity.
- The Sprint Goal is retrofitted around a collection of tickets.

- Planning discussions focus on filling capacity rather than achieving an outcome.
- The team cannot explain why the selected work is the most appropriate response to current evidence.
- AI estimates are presented with precision but without confidence ranges or assumptions.

Better standard

Use AI to prepare information and expose choices, not to make the team's commitment. The team should jointly agree the Sprint Goal, decide how much work it can responsibly undertake and discuss uncertainty, dependencies and technical risk. Forecasts should support planning conversations rather than become performance targets.

Relevant health-check domains and risks

- Primary domain: Event and Cadence Integrity.
- Secondary domains: Collaboration Integrity; Delivery and Learning Integrity.
- Headline risks: Speed Without Learning; Synthetic Certainty.

The Daily Scrum

How AI can help

AI can summarise activity from source control, tickets, chats and calendars. It can also highlight changes, ageing work, dependencies and possible blockers before the event.

How integrity can be weakened

That is useful until the summary replaces the conversation. **This undermines the individuals and interactions over processes and tools value.** The Daily Scrum is not primarily a status collection mechanism. It allows developers to inspect progress towards the Sprint Goal and adjust their plan together.

- Bots produce individual status reports.
- Managers consume reports without participating in the team's coordination.
- Developers stop discussing emerging risks.
- Blockers are inferred from tools rather than raised through trust.
- The event becomes “unnecessary” because a dashboard appears current.

What disappears is the informal coordination through which teams notice ambiguity, offer help and change direction. This also reduces engagement and trust which are required for effective teams.

Integrity principles at risk

- Developer ownership.
- Real-time coordination.
- Psychological safety and engagement.
- Adaptation towards the Sprint Goal.
- Self-management.
- Shared situational awareness.
- Voluntary transparency.
- Mutual support.

Warning signs

- The team receives an automated report instead of holding the event.
- Participants read status updates aloud without adapting the plan.
- Managers use the AI summary primarily to monitor individuals.
- Blockers are inferred by software rather than raised openly.
- Developers no longer ask each other for help.
- Team members believe the meeting is unnecessary because the dashboard is current.
- The Sprint Goal is rarely discussed.
- The output of the event is a status report rather than an adjusted plan.

Better standard

A generated summary should serve as preparation for a conversation, not replace it. AI may reduce the time needed to collect routine information, allowing the team to focus on progress towards the Sprint Goal, emerging risk, coordination and changes to the plan. The event remains owned by the developers and should not become a management reporting mechanism.

Relevant health-check domains and risks

- Primary domain: Event and Cadence Integrity.
- Secondary domains: Collaboration Integrity; Governance and Accountability Integrity.
- Headline risk: Speed Without Learning.

Sprint Review

How AI can help

AI can produce demonstrations, release summaries, stakeholder packs and predicted outcome reports. It can also organise feedback and help compare expected and observed outcomes.

How integrity can be weakened

This can turn the Sprint Review into a persuasive performance rather than an inspection of the product. It undermines the value of feedback.

- AI writes an optimistic narrative around limited progress.
- Mock-ups or simulations obscure what is actually working. **This undermines the value of working software as the primary measure of progress.**
- Predicted customer benefits are discussed as though they are realised benefits.
- Stakeholder questions are answered with generated explanations rather than investigated. These generated explanations may further diverge from the intended product roadmap, vision or strategy.
- The review becomes a demonstration of completed scope rather than a discussion about what to do next.

Integrity principles at risk

- Transparency.
- Empirical inspection.
- Stakeholder collaboration.
- Differentiating outputs from outcomes.
- Adaptation based on evidence.
- Honest representation of what is working.
- Shared product ownership.
- Openness about uncertainty and incomplete value.

Warning signs

- The review contains polished slides but little or no working product.
- Mock-ups and prototypes are not clearly distinguished from completed functionality.
- Predicted benefits are presented as achieved outcomes.
- AI-generated summaries dominate the session.
- Stakeholders are presented with conclusions rather than invited to inspect and discuss.
- Questions are answered with generated explanations rather than evidence.
- The review focuses on scope completed rather than what was learned and what should change.
- There is no update to the Product Backlog or product direction after the event.

Better standard

AI can explain what was built. It cannot establish that value was created without real-world evidence. Use AI to prepare concise, traceable information, but centre the review on the working product, stakeholder observation and available outcome evidence. Clearly distinguish completed output, expected benefit and realised benefit. The event should lead to an informed discussion about what to do next.

Relevant health-check domains and risks

- Primary domain: Event and Cadence Integrity.
- Secondary domains: Product and Evidence Integrity; Delivery and Learning Integrity.
- Headline risk: Synthetic Certainty.

Sprint Retrospective

How AI can help

AI can help organise input, identify repeated themes and prepare possible discussion prompts where the team has agreed that such use is appropriate.

How integrity can be weakened

The danger is outsourcing reflection. A retrospective is not valuable because it produces a list of actions; it is valuable because the team jointly reflects on its current situation and commits to improving it.

- AI identifies “root causes” without sufficient context.
- Sentiment analysis becomes a proxy for honest conversation.
- Management mines retrospective data across teams and treats different teams’ causes as equivalent.
- People moderate their comments because AI is recording and classifying them.
- Generic actions are generated without collective commitment.
- The same system that monitors performance recommends how the team should improve.

This can turn retrospectives into algorithmic surveillance disguised as continuous improvement.

Integrity principles at risk

- Psychological safety.
- Team ownership.
- Contextual sense-making.
- Voluntary disclosure.
- Local experimentation.
- Trust.

- Confidentiality.
- Self-management.
- Inspection and adaptation owned by the team.

Warning signs

- Retrospective transcripts are analysed without explicit team agreement.
- Management has access to individual comments or sentiment scores.
- People become more guarded after AI recording or analysis is introduced.
- The same generic recommendations appear across different teams.
- AI-generated root causes are accepted without team discussion.
- Actions are assigned without collective commitment.
- Retrospective data is used in performance assessment.
- The team stops discussing difficult interpersonal or organisational issues.

Better standard

AI may support preparation and organisation, but the team must retain ownership of interpretation, diagnosis and action. Retrospective data should be treated as sensitive. Its collection, storage and use should be explicitly agreed and should not be used to evaluate individuals. The purpose is to create a safe conversation that leads to one or two locally owned improvement experiments.

Relevant health-check domains and risks

- Primary domain: Event and Cadence Integrity.
- Secondary domains: Collaboration Integrity; Governance and Accountability Integrity.
- Headline risk: Synthetic Certainty.

The two-week Sprint

How AI can help

AI can reduce the administrative effort associated with analysis, documentation, testing and coordination. This potentially creates more time within the Sprint for discovery, integration and validation.

How integrity can be weakened

AI accelerates activities at different rates. Coding may become faster while customer access, policy approval, security review, procurement, integration, production observation, behavioural change and market adoption remain relatively unchanged.

The result is not necessarily shorter end-to-end delivery. It may create a larger queue of code awaiting validation, integration or release. Leaders may conclude that the team should complete more stories or shorten the Sprint, increasing utilisation and work in progress without changing governance or dependencies.

Cadence becomes ceremonial when work is continuously generated, scope changes daily, Sprint Goals are retrofitted around completed tickets and teams carry large volumes of unfinished validation work.

It may also result in the customer being inundated with change which cannot be effectively evaluated before the next change is released. This can undermine effective product management and alienate customers. If the team is delivering internally to another team, the organisation may experience substantial friction through mismatched dependencies, workflows, processes or data exchange.

Integrity principles at risk

- Sustainable pace.
- Focus through a meaningful Sprint Goal.
- Empirical inspection and adaptation.
- Delivery of a usable, integrated increment.
- End-to-end flow rather than local utilisation.
- Cadence aligned to meaningful feedback.
- Limiting work in progress.
- Outcome learning rather than output accounting.

Warning signs

- Work is continuously generated.
- Scope changes daily through AI recommendations.
- Releases happen independently of the Sprint while the Sprint remains a reporting cycle.
- Sprint Goals are retrofitted around completed tickets.
- Generated plans are continually re-optimised.
- Teams carry large volumes of unfinished validation work.
- Output targets increase because AI is assumed to make development faster.
- Coding is faster, but lead time and outcome validation do not improve.
- Teams complete work that remains unreleased or unvalidated for several Sprints.

Better standard

Retain the two-week Sprint where it creates useful focus, integration and feedback. Change it only when another cadence better supports the end-to-end learning loop.
Ask:

- What uncertainty does the cadence contain?
- How quickly can meaningful external feedback be received?
- Can the team produce and inspect a coherent increment?
- Are dependencies compatible with the cadence?

- Does the Sprint Goal provide focus?
- Is the Sprint still a learning cycle or merely an accounting period?

AI may justify different cadences for discovery, delivery and outcome review. It does not remove the need for cadence altogether, and time saved by AI should not automatically be converted into more work in progress.

Relevant health-check domains and risks

- Primary domain: Event and Cadence Integrity.
- Secondary domains: Delivery and Learning Integrity; Governance and Accountability Integrity.
- Headline risk: Speed Without Learning.

Behaviour and quality

BDD - Behaviour-Driven Development

How AI can help

AI can draft, update and format acceptance criteria in Given-When-Then form. It can also suggest provocative examples, edge cases and counterexamples for the team to explore.

How integrity can be weakened

BDD is often reduced to writing scenarios in Given-When-Then format. AI makes that reduction especially dangerous because it can generate hundreds of syntactically correct scenarios.

BDD is fundamentally a collaborative discovery practice. Its value comes from developers, testers and business representatives discussing examples and discovering ambiguous rules, missing cases, conflicting interpretations, hidden assumptions, shared business language and where to reduce complexity.

When AI generates feature files before the “Three Amigos” conversation, the result may have excellent syntax but poor behavioural understanding which results in:

- Scenario proliferation.
- Brittle automation.
- Tests confirming the AI’s interpretation rather than the desired behaviour.
- Important examples being missed because nobody explored the domain.
- Acceptance criteria becoming contracts rather than conversation prompts. **This undermines the principle of customer collaboration over contract negotiation.**

Integrity principles at risk

- Shared understanding.

- Cross-functional collaboration.
- Behaviour discovery through examples.
- Customer and business language.
- Conversation over specification.
- Simplicity and avoiding unnecessary complexity.
- Collective ownership of expected behaviour.
- Testing assumptions before automating them.

Warning signs

- Success is measured by whether AI produced valid Given-When-Then scenarios.
- Feature files are generated before the “Three Amigos” conversation.
- The number of scenarios grows rapidly without improved understanding.
- Participants review syntax rather than discuss business behaviour.
- Product, development and testing retain different interpretations despite detailed scenarios.
- Generated scenarios are automated without assessing whether they provide lasting value.
- The team cannot explain why a scenario matters.
- Acceptance criteria become a fixed contract rather than a prompt for discovery.

Better standard

The question is not “Did AI write valid scenarios?” It is “Did the team develop a shared understanding of the behaviour and discover something important through the examples?” AI should propose provocative examples, edge cases and counterexamples, not silently define expected behaviour. The team should decide which examples clarify behaviour, which should become automated tests and which have served their purpose through discussion alone.

Start by having the collaborative discussion and generating the Given-When-Then for the intended behaviour. Ensure that there is a shared understanding of the desired behaviour then use AI to supplement those discussions.

Relevant health-check domains and risks

- Primary domain: Collaboration Integrity.
- Secondary domains: Product and Evidence Integrity; Technical Integrity.
- Headline risk: Synthetic Certainty.

Definition of Done

How AI can help

AI can review code, generate tests, check documentation and assist with repeatable accessibility, security, compliance and quality checks.

How integrity can be weakened

Teams can become overconfident in automated assurance. A feature may satisfy machine-checkable conditions while remaining:

- Unusable.
- Inaccessible in practice.
- Operationally unsupported.
- Misaligned with user needs.
- Poorly understood by maintainers.
- Ethically problematic.
- Vulnerable in ways the models do not detect.

The Definition of Done may gradually narrow to whatever AI can verify.

Integrity principles at risk

- Built-in quality.
- Technical excellence.
- Transparency about the true state of the increment.
- Fitness for use, not merely test compliance.
- Maintainability and operational readiness.
- Shared ownership of quality.
- Human judgement for contextual and ethical quality.
- A usable, valuable increment.

Warning signs

- The Definition of Done contains only checks that tools can verify.
- Passing generated tests is treated as proof of fitness for use.
- Generated code is accepted without maintainability or architectural review.
- Accessibility, usability or operational support is assumed rather than demonstrated.
- The team cannot explain critical generated code.
- Known quality concerns are deferred because automated checks passed.
- Test count or coverage increases while defects, rework or incidents do not improve.
- Human exploratory testing or user validation disappears.

Better standard

Teams must explicitly retain human judgement for quality dimensions that cannot be reliably automated. AI-generated checks should complement, not define, the Definition of Done. The team remains accountable for ensuring that the increment is usable, secure, accessible, operable, maintainable, safe and aligned with user needs.

Automated assurance should be combined with risk-based human review and real-world validation.

Relevant health-check domains and risks

- Primary domain: Technical Integrity.
- Secondary domains: Delivery and Learning Integrity; Product and Evidence Integrity.
- Headline risk: Synthetic Certainty.

Technical ownership

Pairing and collective ownership

How AI can help

AI assistants can provide on-demand suggestions, explain unfamiliar code, generate alternatives and help developers explore implementation approaches without waiting for another person to become available.

How integrity can be weakened

AI coding assistants can reduce collaboration between developers. Instead of asking a colleague, a developer asks the model. This can improve flow locally but weaken the team systemically.

- Knowledge remains with individuals and their AI sessions.
- Fewer design conversations occur.
- Code is accepted without deep understanding.
- Junior developers receive answers without developing diagnostic skill.
- Experienced developers lose awareness of the wider codebase.
- Human pairing appears inefficient compared with one person plus an AI assistant.

This creates AI-mediated silos. The individual may produce faster while the team becomes less resilient.

Integrity principles at risk

- Collective code ownership.
- Shared technical understanding.
- Collaboration and knowledge transfer.
- Continuous learning and skill development.
- Sustainable team capability.
- Technical review through multiple perspectives.
- Resilience when individuals are unavailable.
- Responsibility for understanding generated code.

Warning signs

- Developers consult AI before or instead of colleagues in nearly every case.
- Important design decisions exist only in individual AI conversations.
- Pairing or mobbing declines because it is perceived as inefficient.
- Junior developers deliver code they cannot explain.
- Senior developers lose awareness of changes outside their immediate work.
- Knowledge concentration increases despite higher individual output.
- Code review focuses on syntax rather than design reasoning.
- The team depends on one person's prompting practices or private AI history.

Better standard

Use AI to support individual exploration but retain human pairing and group design where shared understanding, risk or knowledge transfer matters. Significant technical decisions should be visible to the team, and developers must be able to explain and maintain the code they introduce. AI interactions that contain important design rationale should be converted into accessible team knowledge.

Relevant health-check domains and risks

- Primary domain: Technical Integrity.
- Secondary domains: Collaboration Integrity; Governance and Accountability Integrity.
- Headline risk: Speed Without Learning.

Technical excellence and refactoring

How AI can help

AI can assist with code analysis, identify duplication, suggest refactoring options, generate tests and help developers understand unfamiliar parts of the system.

How integrity can be weakened

AI lowers the cost of generating code, but not necessarily the cost of owning it. Teams may create more code than they can understand, test meaningfully, secure, operate, evolve or decommission safely.

Because working code appears quickly, the pressure to explore architecture, constraints and maintainability can decline. Technical debt may become less visible because AI can continually patch around it. The organisation experiences temporary speed while structural complexity accumulates.

Integrity principles at risk

- Technical excellence.

- Simplicity.
- Sustainable design.
- Maintainability.
- Architectural coherence.
- Collective ownership.
- Built-in quality.
- Continuous attention to technical debt.
- Maximising the amount of code not written.

Warning signs

- Code volume rises faster than the team's capacity to review and maintain it.
- Generated patches repeatedly address symptoms without resolving structural problems.
- Refactoring is deferred because AI can quickly work around poor design.
- Architecture decisions are made implicitly through generated code.
- Similar functionality is repeatedly regenerated rather than reused or simplified.
- Tests increase, but developers have less confidence in changing the system.
- The team cannot safely remove generated code.
- Lead time initially improves while incidents, rework or maintenance effort later increase.

Better standard

Use AI to expose design options and reduce the effort of well-understood changes, but maintain explicit architectural judgement and ownership. Teams should review whether generated code increases unnecessary complexity and should continue to invest in refactoring, simplification and deletion. The objective is not to maximise code production but to maintain a system that can be safely understood, operated and changed. Employ other tools to maintain clean code so that AI is not marking its own homework.

Relevant health-check domains and risks

- Primary domain: Technical Integrity.
- Secondary domains: Delivery and Learning Integrity; Collaboration Integrity.
- Headline risk: Speed Without Learning.

Practical Agile Integrity Test

For any use of AI, ask five questions.

1. Does it improve evidence or merely generate content?

Generated customer insights, forecasts and acceptance criteria are not evidence.

2. Does it strengthen or replace collaboration?

Efficiency that eliminates necessary cross-functional conversation may damage agility.

3. Is accountability still clearly human?

Someone must be able to explain and own the decision without blaming the tool.

4. Does it shorten the full learning loop?

Faster production is not valuable if validation, integration or adoption remains slow.

5. Does the team understand the output?

A team should not ship, estimate, test or commit to work that it cannot meaningfully explain.

The AI–Agile Integrity Health Check Diagnostic

The accompanying diagnostic assesses six domains of Agile Integrity:

1. Product and Evidence Integrity
2. Delivery and Learning Integrity
3. Collaboration Integrity
4. Technical Integrity
5. Event and Cadence Integrity
6. Governance and Accountability Integrity

It also calculates two cross-cutting headline risk indicators:

- Speed Without Learning
- Synthetic Certainty

The diagnostic applies minimum completion thresholds before displaying domain and headline results. Where there are too few valid responses, or a high proportion of neutral responses, the results are withheld or identified as provisional so that uncertainty is not mistaken for evidence.

Overall Agile Integrity gives equal weight to the six domains. Within each domain, positively and negatively worded statements are balanced so that the result is not unduly influenced by the number or direction of question types.

Speed Without Learning measures whether production acceleration is outpacing validated learning, with congestion and output pressure treated as supporting risk factors.

Synthetic Certainty considers reliance on generated claims, evidence and source traceability, challenge and counter-evidence, and clear human review and accountability.

Interpreting the results

The overall score should not be treated as the whole diagnosis.

A score should be interpreted only when the workbook shows sufficient completion. A high level of neutral responses indicates uncertainty or uneven experience and should trigger further discussion and evidence gathering rather than a firm conclusion.

Users should inspect:

- Assessment completion.
- The two headline indicators.
- The lowest domain scores.
- Significant differences between respondents or roles.
- Supporting comments and examples.
- Operational evidence.
- Critical practices hidden by an acceptable average.

The four combinations of the headline risks:

Speed Without Learning	Synthetic Certainty	Meaning
Low	Low	Controlled adoption
High	Low	Output congestion
Low	High	Confident but weakly grounded
High	High	Accelerated illusion

The distinction between integrity scores, risk scores and confidence

- A **high integrity score** is positive.
- A **high headline risk score** is negative.

A provisional or withheld result means there is not yet enough reliable information to support the stated interpretation. Completion and response confidence should therefore be reviewed before comparing scores or selecting priorities.

Score	Integrity interpretation	Risk interpretation
0–24	Critical weakness	Low or controlled risk
25–49	Material weakness	Emerging risk
50–74	Developing integrity	Material risk
75–100	Strong integrity	Critical risk

Coaching guidance

The agile coach's role should shift from protecting ceremonies to protecting the conditions for empirical learning.

AI should accelerate	AI must not replace
Information retrieval	Customer contact
First drafts	Accountable judgement
Pattern identification	Contextual sense-making
Test generation	Behaviour discovery
Routine automation	Cross-functional collaboration
Option generation	Product choice
Activity summaries	Inspection conversations
Forecasting support	Team commitment
Code assistance	Technical ownership
Retrospective preparation	Psychological safety

Use AI to reduce the cost of obtaining and examining evidence, not to reduce the human engagement required to interpret evidence, make trade-offs and take responsibility.

AI-compatible agility may involve fewer administrative ceremonies, more continuous discovery and greater automation. But it must preserve the deeper controls. This means shorter evidence loops, shared understanding, technical ownership, transparent uncertainty, team autonomy and accountable product decisions.

Conclusion

AI does not undermine agile because it makes work faster; it undermines Agile Integrity when speed is mistaken for value, generated output is treated as evidence, and automation replaces human collaboration, judgement and accountability. Used well, AI helps teams identify patterns, draft options, summarise information and reduce waste. Used poorly, it inflates backlogs, creates false certainty, weakens customer discovery and turns transparency into surveillance.

The integrity test is simple:

“Does AI strengthen the conditions for better outcomes?”

Those conditions include direct evidence, shared understanding, adaptive planning, technical excellence, psychological safety, sustainable pace and clear human ownership of decisions. AI should support discovery, coordination and delivery without becoming a proxy for product judgement, team commitment or empirical learning.

References

Moondra, M. (2026) 'How to measure how much AI is improving developer productivity', *Forbes Technology Council*, 20 January. Available at: <https://www.forbes.com/councils/forbestechcouncil/2026/01/20/how-to-measure-how-much-ai-is-improving-developer-productivity/> (Accessed: 30 July 2026).