



SOURCE-BASED CASE ANALYSIS

Coding agents and productivity: what the METR studies actually show

A technical analysis of METR's measurements of AI-assisted software development, their limitations and the implications for credible internal evaluations.

METR studies on AI-assisted software development · 2026-07-18 · 1.0 · Christian Blank, Augmentry

CHRISTIAN BLANK · AUGMENTRY · 18.07.2026

KEY FINDING

The studies do not provide a universal productivity number for coding agents. They show how closely outcomes depend on tool maturity, task selection and measurement design.

The situation

Productivity claims about coding agents are often treated like product specifications: a percentage that can be transferred from a benchmark to a development organization. The METR publications demonstrate why that is unsafe. Their main value lies in showing how a clean early result has to be reinterpreted when tools, usage patterns and sampling conditions change.

The useful question is therefore not whether AI accelerates software development in general. A defensible decision starts narrower: Which tasks can a defined team complete faster or better with a specific tool while maintaining explicit quality constraints?

The early randomized finding

In early 2025, METR studied experienced open-source developers working in repositories they knew well. The researchers randomly assigned 246 real tasks to work with or without the available AI tools. In that setting, developers using AI took 19 percent longer on average. Before the work, they expected a speed-up. After completing it, they still believed they had been faster than the measurements indicated.

The result matters, but its boundary matters equally. It concerns experienced maintainers, familiar codebases, a particular task mix and early-generation tools. It proves neither that coding agents are generally harmful nor that subjective experience has no value. It demonstrates that perceived ease and measured elapsed time can diverge.

For organizations, the immediate lesson is that adoption data is not performance data. A tool may feel helpful, produce more code and still create additional verification, correction or coordination work.

Why later evidence became harder to read

METR continued the work with newer tools. Its later update explicitly describes selection problems. Participants had more discretion over when to use AI. People who already operate a tool effectively, or expect it to help, are more likely to select it. Parallel agents also complicate time attribution: a developer may wait, review output, start a second task or coordinate several runs.

In its May 2026 Frontier Model Risk Report, METR reports a small positive productivity effect, approximately four to twenty percent, for public agents from late 2025. The report also explains why that interval may underestimate some effects and why self-reported gains are much larger but more uncertain.

These findings are not necessarily contradictory. They describe different tool generations and different measurement conditions. Any account that repeats only the 19 percent slowdown, or only a later speed-up, removes the information needed to interpret the number.

What counts as productivity

Completion time is useful, but incomplete. Software delivery requires at least three levels of measurement:

- 1 Task level: active time, waiting time, iterations and the share of tasks completed successfully.
- 2 Quality level: defects, review effort, test coverage, security findings, maintainability and later rework.
- 3 System level: deployment frequency, change failure rate, recovery time, flow and the load moved to other roles.

A local speed-up can reduce system performance. More pull requests do not help when review becomes the constraint. Short implementation time is not a gain if subtle defects appear in production. Conversely, an agent may be useful without shortening a single task when it improves tests, documentation or the range of alternatives considered.

OpenAI's audit of SWE-Bench Pro illustrates a related measurement failure. The published review classified roughly 30 percent of the examined tasks as broken or not reliably gradable, citing issues such as inadequate tests, underspecified requirements and misleading prompts. A precisely calculated score is still unhelpful when the instrument does not represent the decision an organization needs to make.

A credible internal evaluation

Organizations should not look for a universal agent multiplier. They should establish a repeatable local measurement.

Define task classes first

Useful candidates are recurring and sufficiently comparable: small fixes, test additions, migrations, documentation or tightly scoped features. Exploratory architecture work should be treated separately. Classification must precede measurement so that successful examples are not selected afterwards.

Establish a baseline and quality constraints

Historical or controlled comparison data is needed before adoption. Acceptance criteria, tests, security checks and review rules stay constant. A failed quality gate cannot be offset by shorter completion time.

Count the entire active workload

Working time includes prompting, context preparation, waiting, inspection, correction and abandoned attempts. With parallel agents, teams should also record calendar time until a usable change is ready. The two measures answer different questions.

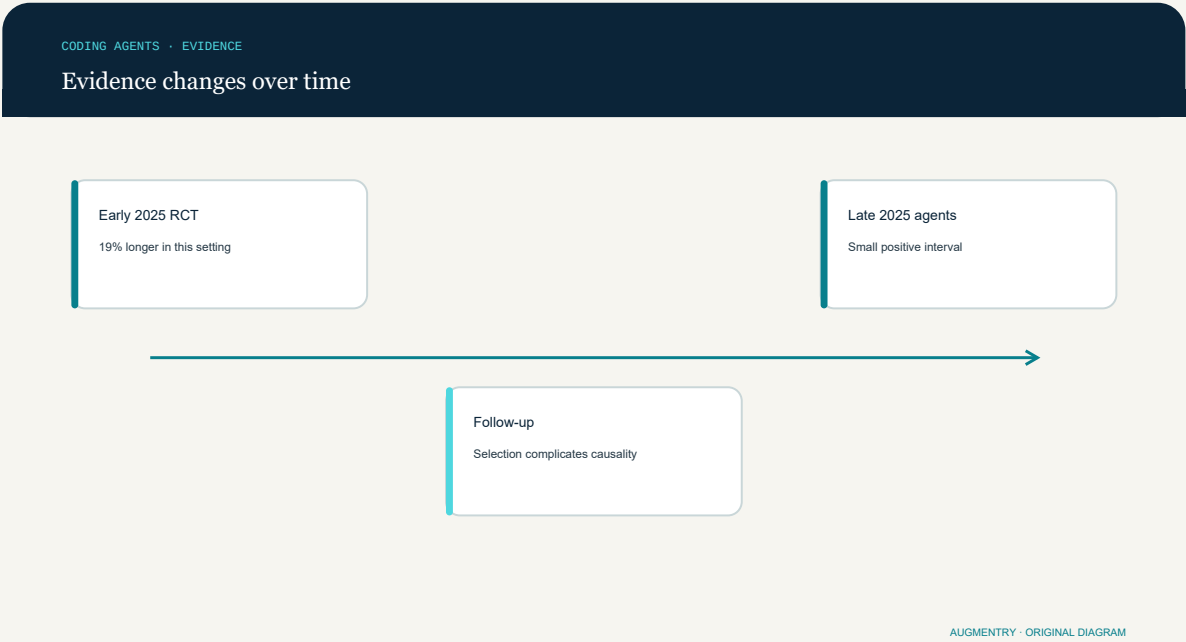
Decide on outcomes and side effects together

Adoption is justified when the intended effect remains stable across several task classes, quality constraints hold and review or operations do not absorb an unacceptable burden. The answer may be selective: use for tests and migrations, but not for security-critical core logic.

Conclusion

The METR series is not a verdict for or against coding agents. Its methodological lesson is more durable: productivity is contextual, perception can be wrong, and even a rigorous result ages as tools and usage change. That does not justify waiting. It supports a more precise course of action: start narrowly, define measurements in advance, treat quality as a hard constraint and revisit the evidence regularly.

Original diagrams



Original diagram: the evidence changes with tool maturity and study design.



Random assignment reduces bias, but does not by itself make a result universally transferable.

Three measurement layers

01 · Task

Active time · waiting · iterations · completion

02 · Quality

Defects · review · tests · maintainability

03 · System

Flow · failure rate · recovery · team load

AUGMENTRY · ORIGINAL DIAGRAM

Time, quality and system impact should be measured separately and evaluated together.

Selection changes the sample

Random assignment

Comparable groups

Causal estimate

Narrow but interpretable

Voluntary use

Experience + preference

Observed uplift

Useful signal, mixed causes

AUGMENTRY · ORIGINAL DIAGRAM

Voluntary tool users operate under different conditions from a randomized comparison group.

Local evaluation loop



AUGMENTRY · ORIGINAL DIAGRAM

A local evaluation connects a baseline, controlled use, quality review and an operating decision.

Limitations and sources

Limitations of this analysis

- The early randomized study covers 16 experienced developers and 246 tasks in familiar open-source repositories. Its result cannot be transferred directly to other teams.
- Tools, models and agent workflows change faster than conventional study cycles. METR now explicitly labels the early finding as outdated.
- Later observations are harder to interpret causally because of self-selection, parallel agent use and changed compensation.

Sources

- 1 Primary source Early-2025 AI experienced open-source developer productivity study METR, 2025-07-10
- 2 Primary source Measuring AI Ability to Complete Long Tasks METR, 2026-02-24
- 3 Primary source Frontier Model Risk Report METR, 2026-05-19
- 4 Context Separating signal from noise in coding evaluations OpenAI, 2026-07-08
- 5 Context Measurement frameworks DORA

URLs

1. <https://metr.org/blog/2025-07-10-early-2025-ai-experienced-os-dev-study/>
2. <https://metr.org/blog/2026-02-24-uplift-update/>
3. <https://metr.org/blog/2026-05-19-frontier-risk-report/>
4. <https://openai.com/index/separating-signal-from-noise-coding-evaluations/>
5. <https://dora.dev/research/2025/measurement-frameworks/>