

EXECUTIVE BRIEF

Favorable Stance Is Not Accuracy Trust

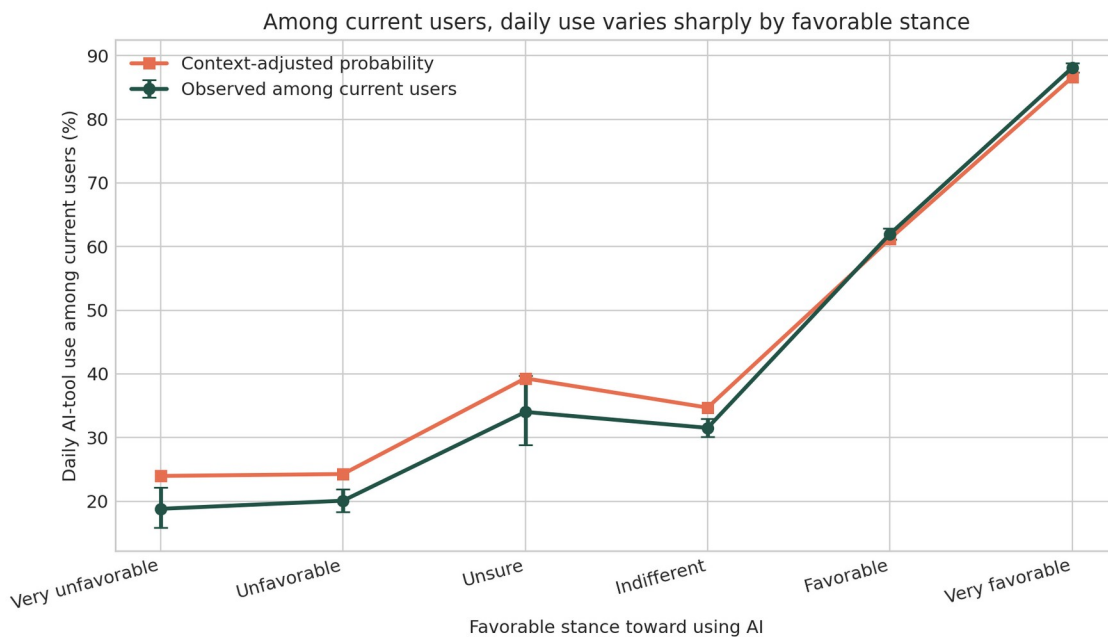
What 26,102 current AI users in the 2025 Stack Overflow survey reveal about reported use frequency

THE FINDING Among people already using AI, reported daily use is 18.8% for the very unfavorable and 88.1% for the very favorable. Favorable stance carries more classification information about daily versus less-frequent use than accuracy trust does. This is an association, not a causal effect.

The numbers that matter

Result	Meaning
0.758	Stance-only ROC AUC
0.669	Trust-only ROC AUC
0.704 -> 0.790	Context + trust versus context + stance
0.793	Complete-model ROC AUC

Five-fold cross-validation among current users, n=26,102. Same-survey classification, not future forecasting.



Daily-use rates and context-adjusted probabilities across favorable-stance responses among current users.

The important distinction

Favorable stance is a broad evaluation of using AI. Accuracy trust is a narrow belief about whether AI output is correct. The broader, more use-proximal item aligns more closely with frequency. That does not make stance a proven psychological driver or make trust unimportant.

A TOUGHER COMPARISON Among current users who highly distrust AI accuracy, daily use is 17.7% for the very unfavorable and 85.1% for the very favorable. A user can value AI while remaining skeptical enough to verify its outputs.

What this changes for leaders

	Move	Why
1	Track behavior, not access.	Measure frequency, task, workflow depth, and outcomes.
2	Separate stance from accuracy trust.	One reflects broad value and fit; the other confidence in correctness.
3	Optimize for calibrated reliance.	Seek useful, appropriately skeptical, well-verified use.
4	Test changes longitudinally.	Use pre/post measures, telemetry, quality, rework, and staggered support.

What the data do not prove

- That favorable stance causes daily use; use may shape stance.
- That the favorable-stance item is a validated sentiment scale.
- That trust is unimportant because it adds little classification information.
- That training or persuasion will increase adoption.
- That Stack Overflow respondents represent all developer workforces.

A better adoption dashboard

Dimension	Question
Behavior	How often, where, and for what
Favorable stance	Whether AI feels valuable and workable
Accuracy trust	Confidence that output is correct
Verification burden	Testing, review, rework, defects, and time
Outcomes	Quality, speed, learning, risk, and customer impact

BOTTOM LINE Favoring AI use is not the same as trusting AI output. Measure both - alongside behavior, verification, and outcomes - then test what changes.

Source: Penzell, J. (2026). Adoption Without Confidence? Imagination Applied Research Series. [Stack Overflow 2025 AI results](#)