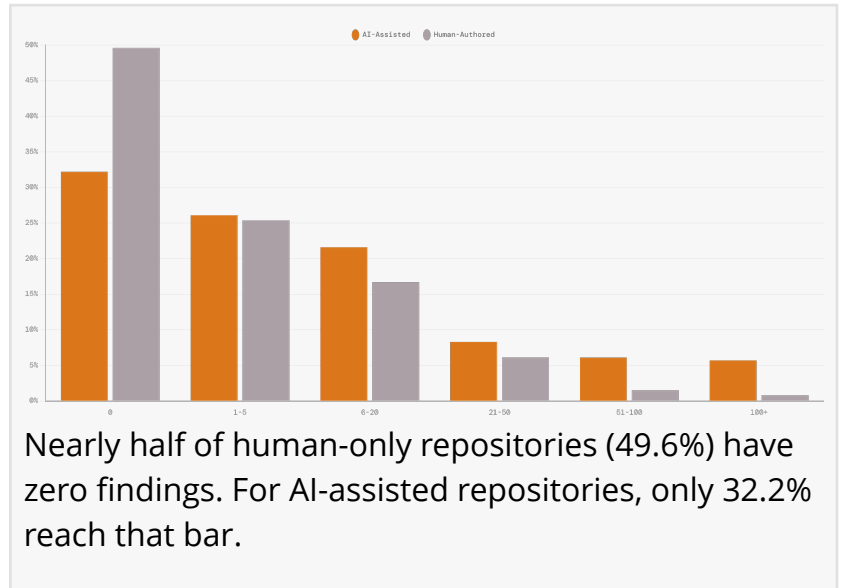


AI-Assisted Code Introduces 4.4x More Security Vulnerabilities Than Human-Written Code, New Research Finds

Study of 2,000 repos and 1,000 vibe-coded apps from State of Cyber measures AI-generated code's security cost.

BROOKLYN, NY, UNITED STATES, August 20, 2026 /EINPresswire.com/ -- [State of Cyber](#), the research lab of [Symbiotic Security](#), today published AI vs Humans: Security at Scale, one of the largest studies so far on how AI is changing the security of the software we all rely on. Over three months, the lab analyzed 1,967 public GitHub repositories, scanned 1,072 vibe-coded applications across five popular platforms, and ran ten independent security scanners across all of it. In total they found 87,826 vulnerabilities, from small, low-risk bugs to critical flaws that could expose private data or hand an attacker control of an application.



The headline number is hard to ignore. Repositories built with AI help carried an average of 42.3 vulnerabilities each, against 9.6 in code written by people alone. That is 4.4 times more. And this was not a few bad projects dragging up the average. Seven in ten AI-assisted repositories (70.5 percent) had at least one security issue, compared with about half (50.4 percent) of the human-only ones, and the problems ran across 104 different vulnerability types. AI is adding both more flaws and a wider mix of them.

A second study looked past source code to finished products, checking 1,072 live apps that people had built with AI "vibe coding" tools on Supabase. The results were worse. Ninety-eight percent of the apps had at least one vulnerability, 29 percent carried a high or critical issue, and 16 percent, roughly one in six, let a complete stranger delete or change data with no login required. All told, the audit found 6,185 vulnerabilities, an average of 5.9 per app.

What makes this moment different is who is writing the software. "Vibe coding" tools let almost anyone describe an app in plain language and get working code back in minutes, usually with no

security review anywhere in the process. That is great for speed, and it is exactly why the risk is spreading so fast. A single prompt can put a live database online that anyone can read from or write to, and the person who built it may never realize it happened.

None of this means teams should stop using AI to write code. That decision has mostly been made already, and the productivity gains are real. The point is that the safety net has not kept up. When more code is written faster, by more people, and pushed straight to production, checking for problems after the fact stops working. Something has to catch the issues while the code is still being written.

The findings put numbers on something the industry has been slow to admit. AI is very good at writing code that works, and much weaker at writing code that is safe. That gap is not closing as the models get better. Other independent research keeps landing in the same place, with roughly half of all AI-generated code carrying an exploitable flaw, a rate that has barely moved even as the models themselves have improved.

For the security leaders who have to manage all of this, the study changes the question. AI lets every engineer ship far more code, so teams that only review code after it lands are hunting for more problems in a much bigger pile, with less time to do it. The takeaway is simple. Security has to happen earlier, at the moment the code is written.

Data for both studies was collected between January and March 2026 through automated pipelines at scale and handled under responsible disclosure; no data was exfiltrated, stored, or shared beyond what was needed to document and report each vulnerability. The full findings are published at state-of-cyber.org.

Darren McNelis
Symbiotic Security
+1 251-309-3711

[email us here](#)

Visit us on social media:

[LinkedIn](#)

This press release can be viewed online at: <https://www.einpresswire.com/article/935728867>

EIN Presswire's priority is source transparency. We do not allow opaque clients, and our editors try to be careful about weeding out false and misleading content. As a user, if you see something we have missed, please do bring it to our attention. Your help is welcome. EIN Presswire, Everyone's Internet News Presswire™, tries to define some of the boundaries that are reasonable in today's world. Please see our Editorial Guidelines for more information.

© 1995-2026 Newsmatics Inc. All Right Reserved.