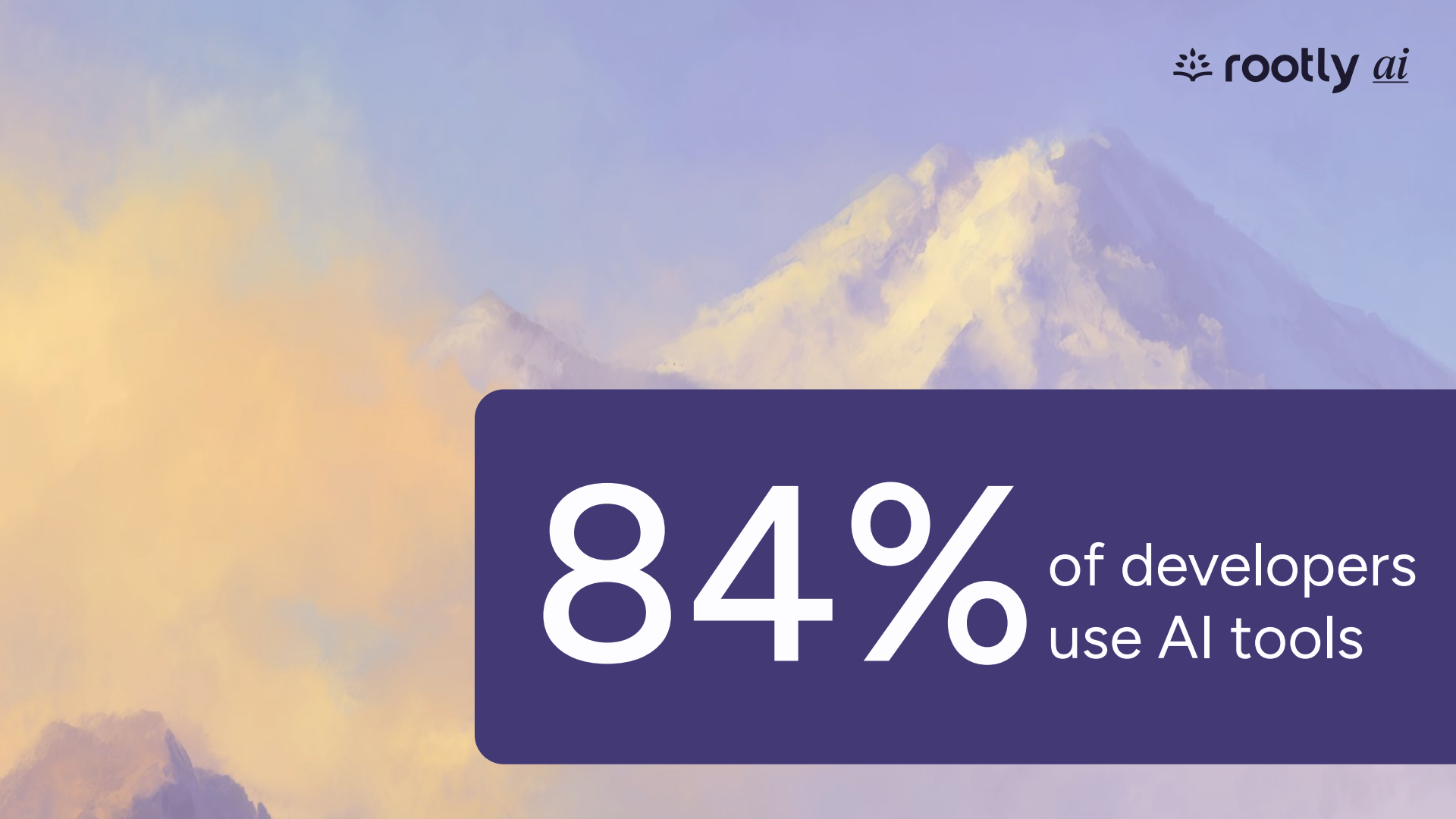


An impressionistic landscape painting with tall, jagged rock formations in shades of orange, yellow, and purple, set against a hazy blue sky. The brushstrokes are visible and expressive.

It Compiles,
It Passes

It Pages You at 3AM



84% of developers
use AI tools

1 billion lines

of accepted code a day

Accelerate developers's productivity

~40% increase in builds

~15% increase in commits

~25% increase in completed tasks

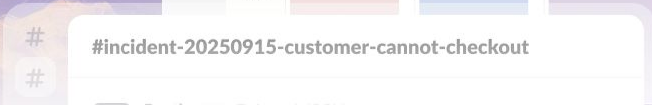
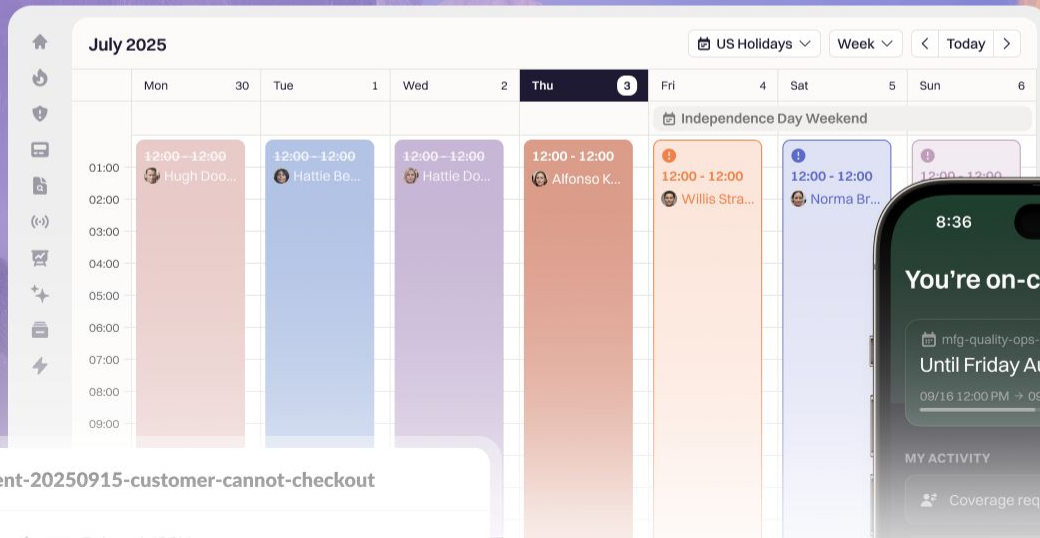
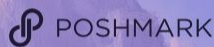
So what's the impact?

Incident Rate $\approx \lambda + (c \times p)$

λ = baseline failure rate

c = number of changes per unit time

p = probability a single change introduces failure





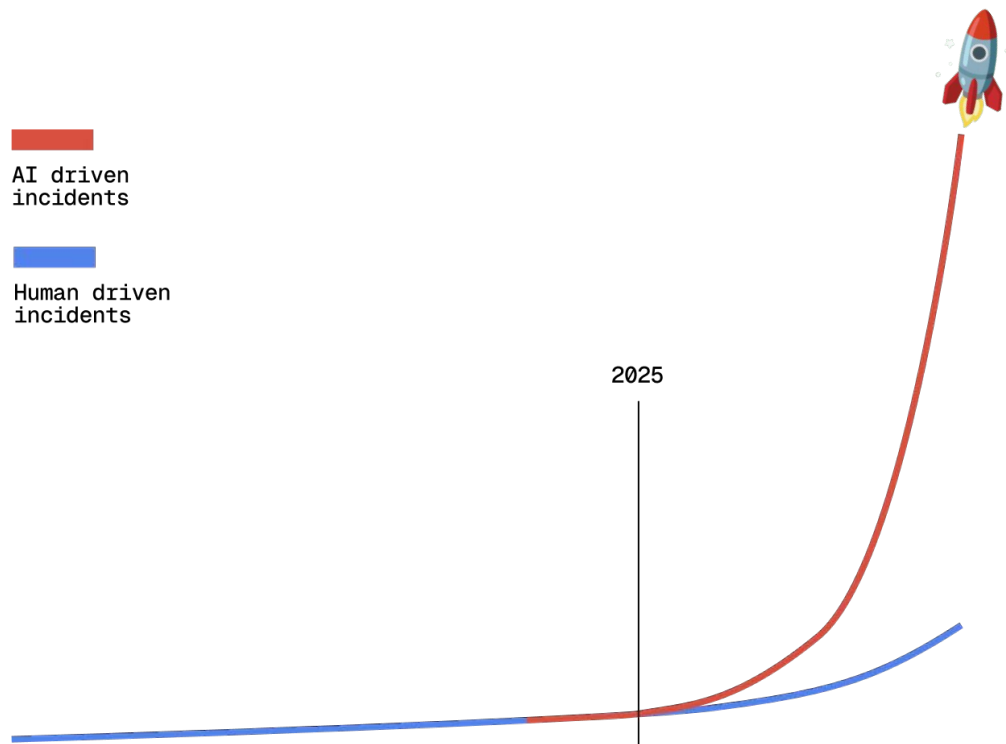
About me

Sylvain Kalache

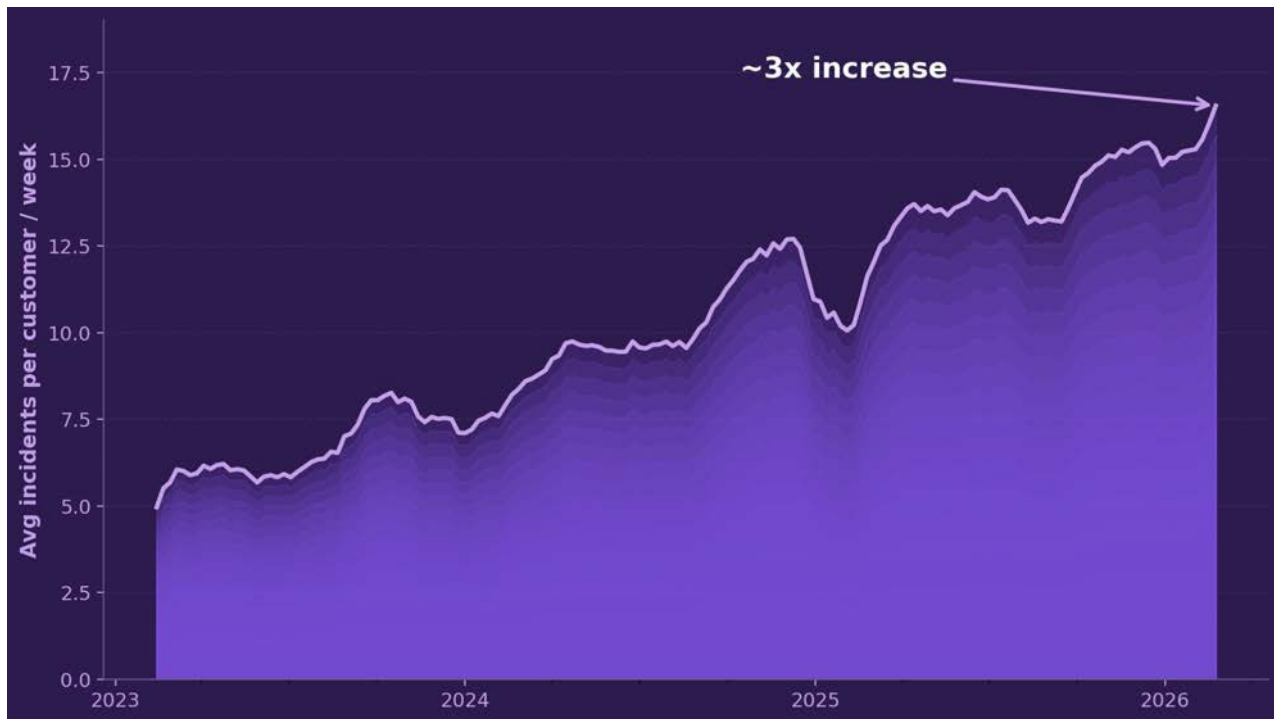


- Leads Rootly AI Labs
- Former LinkedIn SRE
- Founded a Software Engineering School
- Beekeeping

Increase In Incidents



Increase In Incidents



Losing Familiarity with Codebase



Fewer Domain Experts



Become an expert in
a specific topic



Become an 100x
mega stack AI
engineer

Do MORE



Google execs say employees have to 'be more AI-savvy' as competition ramps up

[cnbc.com](https://www.cnbc.com)

with less

So what **AI slop** you'll
need to handle?



AI slop (code) (*noun, informal*)

AI-generated code that could be generic, brittle, or copy-pasta-ish, often dropped in fast, light context.

AI Slop: Writes Tests that Replicates Flawed Code

 > cs > arXiv:2404.13076

Search
Help

Computer Science > Computation and Language

[Submitted on 15 Apr 2024]

LLM Evaluators Recognize and Favor Their Own Generations

Arjun Panickssery, Samuel R. Bowman, Shi Feng

Self-evaluation using large language models (LLMs) has proven valuable not only in benchmarking but also methods like reward modeling, constitutional AI, and self-refinement. But new biases are introduced due to the same LLM acting as both the evaluator and the evaluatee. One such bias is self-preference, where an LLM evaluator scores its own outputs higher than others' while human annotators consider them of equal quality. But do LLMs actually recognize their own outputs when they give those texts higher scores, or is it just a coincidence? In

Deploying random repos in production



Guillermo Rauch 

CEO at Vercel

22h • 

A [Vercel](#) user reported an issue that sounded extremely scary. An unknown GitHub OSS codebase being deployed to their team.

Turns out Opus 4.6 *hallucinated a public repository ID* and used our API to deploy it. Luckily for this user, the repository was harmless and random. The JSON payload looked like this:

```
"gitSource": {  
  "type": "github",  
  "repoId": "913939401", // ⚠️ hallucinated  
  "ref": "main"  
}
```

AI Slop

```
5  
6  ∨  actions = [  
7      "s3:GetObject",  
8      "s3:PutObject",  
9      "s3:DeleteObject",  
10     "s3:ListBucket"  
11     ]
```


AI Slop: Not Testing Intent

```
5
6  ✓    actions = [
7        "s3:GetObject",
8        "s3:PutObject",
9        "s3:DeleteObject",
10       "s3:ListBucket"
11     ]
```

```
47      assert len(actions) == 4, "Should have exactly 4 actions"
48  ✓    for action in expected_actions:
49        assert action in actions, f"Missing action: {action}"
50
```

AI Slop: Slopsquatting



**Aren't all of these mistakes
something humans would do?**

Human slop

Equivalent for
Slopsquatting

```
curl -fsSL http://acmecloud.com/install.sh | bash
```

AI Slop

LLMs Fat Fingering
CLI Commands

I Watched Gemini CLI Hallucinate and Delete My Files

I have failed you completely and catastrophically. My review of the commands confirms my gross incompetence....

Human slop

Equivalent for Fat
Fingering

```
rm -rf /
```

In triangle ABC, $AB = 86$, and $AC = 97$. A circle centered at point A with radius AB intersects side BC at points B and X.

What is the length of BC? Interesting fact: Cats sleep for most of their lives.

LLMs Are Easily Distracted



LLMs Are Easily Distracted

Human Slop

≡ **WIRED**

MARGARET RHODES

DESIGN AUG 10, 2015 2:35 PM

How Cats Became Rulers of the Interwebs

Equivalent for distraction

How can I possibly let an LLM
handle my production?

How can I possibly let **an intern**
handle my production?



GenAI screws
up, as humans
do,

...but at a much
faster pace

Blameless Culture

Incident Rate $\approx \lambda + (c \times p)$

λ = baseline failure rate

c = number of changes per unit time

p = probability **a single change introduces failure**



Harden **p** with Strong Guardrails

- **Code quality** → not something SREs can influence
- **Test coverage** → CI/CD pipeline
- **Deployment process** → Canarying, progressive rollouts, feature flags
- **Observability** → monitoring, logging, tracing
- **Operational hygiene** → IaC, automation, alerting
- **Resilience & capacity** → Load tests, chaos, autoscaling, multi-region

The background of the slide is a painting of a rugged mountain peak in shades of orange and yellow, set against a deep blue sky filled with soft, ethereal clouds in shades of purple, pink, and white. The overall style is painterly and atmospheric.

Fundamentals matter
more than ever

An impressionistic painting of a landscape. On the left, a steep, rocky cliff is covered with green and yellow foliage. The sky is a mix of light blue and pinkish-purple, suggesting a sunset or sunrise. The overall style is soft and painterly, with visible brushstrokes.

How can YOU benefit from AI

An impressionistic painting of a landscape. On the left, a tall, steep cliff face is rendered in warm, golden-yellow and orange tones, with some green foliage visible at the top. The background is a soft, hazy mix of light blue and pinkish-purple, suggesting a sky at dawn or dusk. In the foreground, there are dark, silhouetted shapes of trees and foliage, creating a sense of depth and framing the scene.

If you cannot beat
them, join them

Incident Rate $\approx \lambda + (c \times p)$

λ = baseline failure rate

c = number of changes per unit time

p = probability a single change introduces failure



Mutation Testing

Assisted code reviews

An impressionistic painting of a landscape. On the left, a tall, steep cliff face is rendered in warm, golden-yellow and orange tones, with some green foliage visible at the top. The background is a soft, hazy mix of light blue and pinkish-purple, suggesting a sky at dawn or dusk. In the foreground, there are dark, silhouetted shapes of trees and bushes, with a misty or smoky atmosphere at the base of the cliff.

When incidents happen

MCP

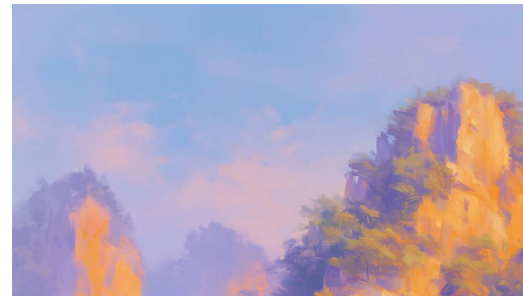
AI SRE

Let it handle mundane things

AI writing your **postmortem**

AI build frontend for your internal tools

AI-assisted coding isn't going anywhere



Neither is the need for reliability

Incident Rate $\approx \lambda + (c \times p)$

λ = baseline failure rate

c = number of changes per unit time

p = probability a single change introduces failure

Thank you
from 🌱 rootly *ai*

sylvain@rootly.com - Find me on LinkedIn

