



Top Content

People

Learning

Jobs

Games

"CLI coding agents: OpenAI Codex, Claude Code, Gemini CLI"

✦ This title was summarized by AI from the post below.



Matt Koppen

10mo

...

Quick thoughts on CLI coding agents more than anything else...

****OpenAI Codex****

+

- Busts through tough
- Much cleaner PRs: bu
- Seeming improbable

—

- Offers to do too much about I do X?" No! Focus
- Verbose and loves to
- Tends to want confirm

?

- Haven't tested it on t

****Claude Code****

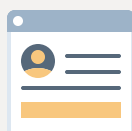
+

- Plows through a lot
- Solid front-end work
- Strikes a good balance between telling me what's doing and not giving me an opus to read (see what I did there?)

—

- Can get very sloppy. PRs that don't build, syntax issues, piling in new code w/out any cleanup. Multiple rounds often needed to clean things up.
- Too often unfounded confidence. "I have done X and it will work this time." Um, it didn't work the last few times and you said the same.
- Blows through capacity pretty quickly at Pro level

(Note: Claude performance not apples-to-apples with Codex b/c I'm using Sonnet 4 due to Pro plan limitations. But that's a downside - Codex gives me GPT-5 on the same level plan!)



Sign in to view more content

Create your free account or sign in to continue your search



Continue with Google



Sign in with Email

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

****Gemini CLI****

Well... what can I say about Gemini CLI. It's a relatively competent junior dev. In main projects, I feel ok giving it a very targeted feature to implement or bug to fix. I'd say it's 60/40 that it gets it when it's on 2.5-pro. If it's a goof project or experiment, I can let both 2.5-pro and 2.5-flash run wild and be ok if it's all a hash. And often it'll be a hash.

A plus? Finding humor in the all-too-often moments when it says (paraphrase) "Oh my, I really did screw that up. Let me go back and revert everything I just did."

Bottom line: Not my go to for anything critical.

****Cursor CLI and non-CLI coding agent****

I gave Cursor CLI a very good try. I used GPT-5, Claude (Sonnet) directly.

Same with non-CLI code prototyping, nothing started with Claude Co

Anyone had a different Gemini CLI (i.e. how?).

2 · 6 Comments

Like

Shad Mickelbe

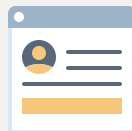
I've been playing (with different between the two of what you want where you really recognizes the

I built a working application with only writing very little code myself and it works great but was pretty annoying getting through some bugs and how certain features were implemented.

I've found the more narrow the focus on a change (perhaps just one file or even function/line) is much more reliable.

All that said it is still pretty exciting.

Like · Reply



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

To view or add a comment, [sign in](#)

More Relevant Posts

Omar Đečević

9mo

I just published a short
how AI can help with c
It covers how I approach
way.

<https://lnkd.in/d6CDr>

Socraites: AI-Powerec
drahil.com

25 · 1 Comment

Like

Niaz Habib

9mo · Edited

I've become a big fan of the OpenAI SDK — not because it's flashy, but because it's well engineered. In a world full of abstraction-heavy GenAI frameworks, this one feels refreshingly honest. Sharing a few thoughts on why it deserves more love 🙌

The OpenAI SDK: Simplicity Done Right



Sign in to view more
content

Create your free account or sign in to continue your
search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to
LinkedIn's [User Agreement](#), [Privacy Policy](#), and
[Cookie Policy](#).

building to explore

at I learned along the

Share

If you haven't tried the OpenAI SDK, it's worth a look.

In the vast ocean of GenAI frameworks and agent builders, this one doesn't get nearly the love it deserves.

Everyone's talking about complex orchestration layers, agent graphs, and "frameworks for frameworks" — but sometimes, all you need is a clean, mature SDK that just works.

And that's exactly what this feels like.

You can tell it was built by engineers who actually write production code. It has that rare quality: it feels designed by experienced developers, for developers.

💡 A Breath of Fresh Air

I've tried a lot of GenAI frameworks lately. Many are impressive, but also... heavy.

They hide simple ideas behind layers of abstraction — "chains," "contexts," "graphs," "memory stores."

Great for demos, not so great for production.

The OpenAI SDK is different.

No guessing.

Just solid engineering.

⚙️ Why It Stands Out

Mature API design — clear and consistent.

Simplicity — minimal boilerplate.

Easy to debug — what you see is what you get.

Performant — no hidden costs.

Fits anywhere — drop it into your existing code.

It's a rare SDK that scales with your needs.

💻 Built for Developers

There's confidence in it.

It doesn't try to own your architecture.

That's the hallmark of a good SDK.

📦 Growing — Without Compromise

Even as OpenAI adds new features, the core remains consistent.

Each layer adds power without adding complexity.

You don't have to learn a new paradigm; it just fits naturally into what you already know.

That's rare — and it shows real design maturity.

❤️ It Deserves More Love

Maybe because it's not flashy, the OpenAI SDK often flies under the radar.

But it's elegant. It's dependable. And it's built with respect for developers who care about clarity and control.

If you're building serious GenAI systems — not just demos — it's absolutely worth a closer look.

🔜 Coming Up Next

In my next post, I'll dive deeper into the Agent SDK and Agent Builder, and explore how OpenAI is



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

redefining what “agentic” can mean — without the bloat.

Sometimes the best tools aren’t the loudest — they’re the ones that just fit the way you think.

□ Curious what others think — have you tried the OpenAI SDK yet? What’s your take compared to other GenAI frameworks?

26 · 1 Comment

Like

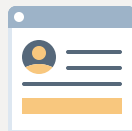
Comment

Share

To view or add a comment, [sign in](#)

Saurabh Suri

10mo



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn’s [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

openai just did an ama

here's some the most i

> just shipped gpt-5-co

diverse coding tasks

>model performs signi
optimization

>team dogfoods heavi
manual coding next ye

team usage & product

n focused training on

due to specialized

with goal of zero

> engineers can prototype large features with just ~5 turns of prompting, building 3 different feature versions in single day

> designers with no coding background now directly merge prs using codex, eliminating traditional developer bottleneck

> massive productivity gains reported: "no more writing crud endpoints or stream helpers" - focus shifts to higher-level architecture

> enables rapid experimentation where suboptimal prototypes can be scrapped without emotional attachment since they're "cheap to build"

performance trade-offs & complaints

> new gpt-5-codex significantly slower than previous models (15-20 minutes vs 3-4 minutes for same tasks)

> more intelligent but speed decrease breaks iterative development workflows for many users

technical capabilities & roadmap

> cli supports web search with --search flag, coming to ide extension soon with prompt caching fixes

> planning mode in development after internal debates about user control vs model autonomy

> exploring browser automation and voice integration (inspired by open source community demos)

> considering sub-agents and workflow automation but nothing actively in development

platform & integration challenges

> currently focused on timeline

> no free tier planned, windows experience problems

token economics & rat

> teams working with "getting started" for ma pro users (\$200/month)

> confusing limit descr minutes

> no usage tracking ui,

future vision & philoso

> tools should become

> abstraction level risin functions

> optimistic scenario: e design/product/engine

ustrated with no

ershell command

I but we are just

ecifically

le session can run 25

ut they said its coming

e background"

ng individual

eralist skills across



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

3

Like

Comment

Share

To view or add a comment, [sign in](#)

Haoyi Yin

10mo

🖥️ Here are 4 zero-cost methods for AI-powered coding in VSCode using Kilo Code ([Kilo Code](#)) a flexible alternative to GitHub Copilot that supports custom OpenAI-compatible APIs.

Method 1: Qwen-Code CLI + Free OAuth

Use Alibaba's Qwen-Code CLI to access the Qwen3-Coder model. You get 2,000 free tokens per day, no token caps.

Setup for Qwen-Code:

1. Install via npm: `npm install -g qwen-code`
2. Run `qwen` and authenticate with your Alibaba account
3. In Kilo Code settings, set the provider to `qwen3-coder` or `qwen3-coder-plus`.

Method 2: ModelScope

ModelScope, Alibaba's open-source platform, provides access to models like Qwen3-Coder after linking your Alibaba account.

Setup for ModelScope:

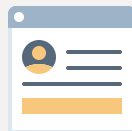
1. Register on [modelscope.cn](#)
2. In Kilo Code:
 - Provider: OpenAI Compatible
 - Base URL: [api.modelscope.cn/v1](#)
 - API Key: Your ModelScope API Key
 - Model: `qwen3-coder`

Method 3: OpenRouter Free Models

OpenRouter aggregates models like DeepSeek V3 and Gemini 2.5 Flash. A one-time \$10 top-up (your balance isn't used for free models) unlocks 1,000 free calls daily (20/min limit). Without it, you get 50.

Setup for OpenRouter:

1. Register on [OpenRouter](#) add \$10 credit, and generate an API key.
2. In Kilo Code:
 - Provider: OpenRouter



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

- API Key: Your OpenRouter Key
- Model: A free model like deepseek-v3:free or gemini-2.5-flash:free.

Method 4: Docker Self-hosted Gemini Balance

If you have multiple free Gemini API keys from Google AI Studio, you can self-host a load balancer to avoid rate limits. Gemini Balance deploys easily via Docker and is optimized for Gemini 2.5 Pro.

Setup for Gemini Balance:

1. Pull the Docker image: <https://lnkd.in/gWX9eP4d>
2. Configure a .env file with your Gemini keys and access tokens.
3. In Kilo Code:
 - Provider: OpenAI Com
 - Base URL: http://<you
 - API Key: Your access t

By load-balancing multiple Gemini API keys, you can avoid rate limits and use Gemini 2.5 Pro for high-frequency AI coding assistance, tailored for VSCode.

These four methods give you a powerful AI coding assistant inside VSCode.

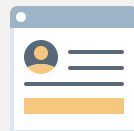
#VSCode #FreeAI #Qv

3

Like

Share

Patrick O'Sha
9mo · Edited



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

I've been having relatively good success using agentic AI CLI tools like Claude Code (Anthropic) and Codex CLI (OpenAI) for coding tasks. The fact that they have access to the full power of the shell, my full project context, and integrate with my preferred and familiar JetBrains IDEs (Rider, Webstorm,

etc), makes for a much better workflow than cut & paste code snippets. There is a continuity in conversation around the code base, and it does start to feel much more like a coding partner.

Of course, it is critical to use good software engineering practices with these new tools just as one would with any project - i.e. always review code diffs before committing to your git repo, do proper testing, etc. Also, their hallucinations and gaslighting over-exuberance about half-baked or even completely broken code can be a grating nuisance. But even with these limitations, I do find myself able to tackle more in less time. YMMV.

As I started spending more time in this mode, however, it became very jarring to abruptly hit up against API limits. My coding partner would suddenly take a multi-hours-long coffee break in the middle of a productive session!

So, I had an idea - what if I could have a coding session where I could use any of the open source models, and I could run them on my machine (TL;DR - no, I can't)

The result is this open source project: <https://lnkd.in/g5jkQ>

Currently supported CLI tools:

- * Claude Code
- * Codex CLI

Remote APIs:

- * OpenRouter
- * DeepSeek
- * Google Gemini
- * **Z.AI** GLM-4.5 (Claude 3.5 Sonnet)

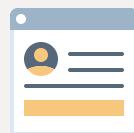
Local models:

Any local MLX or LM Studio models

You can add connections to your local models (I can add it for you)

For this project, I purposely was liberal with allowing the AI's to write nearly all of the code in the repo, including the git commit comments and README, to see how far it could get. I even used the tool itself while in development to switch between models to continue development on parts of it. Through the process, I learned a good deal about MLX, LiteLLM, LM Studio, the limitations in how LLM's access local tools, and the various coding LLM's that currently exist.

I hope it can serve as a useful resource for others that want to kickstart their own exploration of what is possible, and find the tools that work for them in the process. If you give it a try, I'd love to hear what you think!



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

(Btw, I wrote this post myself)

39 · 8 Comments

Like

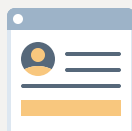
Comment

Share

To view or add a comment, [sign in](#)

Arnab Paul

10mo



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

My Learning Journey w

Over the past few days
to a very practical use

<-- What I Built -->

I implemented a pipeli
instead of feeding the
most relevant code sni

**Here's a breakdown

-- Text Splitting -> Leve
custom separators like

-- Embedding & Retrieval -> Used sfr_embedding_code_400m_r_model from Salesforce (Open Model)
for generating embeddings and retrieving top-k relevant chunks.

Prompt Engineering

- Created a custom PromptTemplate that ensures the model sticks to the given code context and
avoids hallucination.

Runnables & Chain

- Built a RunnableLambda to dynamically pass retrieved code snippets into the prompt, then

understanding --->

ration) and applying it
help of LLMs.

e, a Python file), and
sm to surface only the
that retrieved context.

code logically using

pipelined it with an OpenAI LLM and a StrOutputParser for clean outputs.

****Final Execution****

- Queried the system with questions like “Merging of File and explain the code” and received contextually accurate answers extracted from the code.

<--- Key Learnings --->

****Smart Retrieval > Dumping Context ****

- Feeding the entire codebase is inefficient and often exceeds context limits. Retrieval ensures only relevant pieces are used.

****Custom Splitting is Crucial****

- Default text splitting (e.g., by line) can break code logic. Custom splitting (e.g., by function/class, def) made retrieval more accurate.

****Hallucination Control****

- Explicitly instructing the LLM to only use retrieved context significantly improves reliability.

****Composable Chains****

- Using RunnableLambda and other LCEL tools allows for flexible chaining of retrieval and generation steps to extend functionality.

<--- Why It Matters --->

Firstly I am thankful to the community for their support and feedback. This experiment was a learning experience to create a practical solution for a common problem.

As codebases grow, understanding and managing code becomes a challenge. A RAG-powered approach can act like a personal assistant, providing quick answers with code references, explanations, and insights.

This experiment is just the beginning, and I’m excited to explore more advanced retrieval strategies and integrations for real-world developer workflows.

Github - <https://lnkd.in/gZK9BCDH>

Curious to know how others are approaching RAG for code. Let’s discuss!

#RAG #LangChain #LLM #CodeUnderstanding #AI #LearningInPublic

Like

Comment

Share

To view or add a comment, [sign in](#)

Omri Lahav

10mo

Claude Bug Postmortem

A few weeks ago, during a live demo, Claude's screen.

Claude is against me. The night. I'm tired too."

Did you also feel it?

We weren't hallucinating.

Anthropic just published a postmortem on Claude to act weird. Here's what we learned:

1) Requests were misinterpreted.

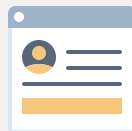
Some prompts went to the wrong model half the time the message was too long or needed.

→ Impact: outputs looked inconsistent, sometimes strangely off-topic, sometimes lower quality than usual.

2) Corrupted outputs

A config bug injected garbage characters into responses by mistakenly assigning higher probability to rare tokens, causing broken code, random tokens, nonsense in the middle of otherwise fine text.

→ Impact: developers got code that wouldn't even compile, or text that looked like Claude had a glitch mid-sentence.



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

3) Token-selection bug in the TPU compiler

An optimization bug sometimes caused Claude to literally skip over the right answer.

→ Impact: you'd get a confident but totally wrong response.

Each of these alone would be painful.

Overlapping together, they made it feel like Claude suddenly got stubborn and dumb (or both).

Now, what does it mean for your AI agents?

Your agents don't just rely on prompts and logic, they also rely on the LLM and its vendor.

These third-party LLMs have bugs that crash loudly. If you know it's wrong, you just

After testing and evaluating, they can start to behave in ways you're screwed.

When the LLM has this

- The output doesn't crash
- Guardrails don't save you
- "allowed but wrong" things
- Your careful pre-production

So remember: agentic

17 · 1 Comment

Like

Share



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

To view or add a comment, [sign in](#)

FireStrike

47 followers

10mo

Software eats the world of coding intelligence, GPT can now refactor your code. Novice? Doesn't matter. Translate it into code.

But, there's a twist, the OpenAI's move isn't just SWE-bench Verified, let's We're not looking at in productivity.

But Codex isn't just for regardless of whether non-coders and season

Industry-wide, these are prompt engines into s Competitors, brace yo double down on AI inte

However, if this is the r code in production? W race between Claude a robust adaptation to keep pace.

The bottom line: AI code agents are no longer just helpers; they're teammates at the keyboard. Tech leaders should start prepping for upskilled teams tasked with validating and securing auto-generated software, while non-tech business units can begin leveraging AI to prototype faster.

Whether you're a veteran coder or a business leader eyeing digital acceleration, keep close tabs on this next wave; the boundaries of who creates, maintains, and secures software just changed.

More details and benchmarks in the original coverage:

<https://lnkd.in/g5DXfbCP> (OpenAI's new GPT-5 Codex model takes on Claude Code)

Unleashes their latest in writes boilerplate, but out breaking a sweat. watch Codex

is Claude Code. impressive 74.5% on om 33.9% to 51.3%. it's possible for dev

ne can pilot it, transformative for on builder.

from monolithic e new normal. IDEs will need to

do we trust generated unch an app? As the pipelines will require



Sign in to view more content

Create your free account or sign in to continue your search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to LinkedIn's [User Agreement](#), [Privacy Policy](#), and [Cookie Policy](#).

Don't just take our word for it, read the full story here: <https://lnkd.in/g5DXfbCP>

[#SECURITYOPERATIONS](#) [#BLUETEAM](#) [#CYBERSECURITY](#) [#SOC](#) [#DIRECTOROFAI](#)

Like

Comment

Share

To view or add a comment, [sign in](#)

Scott Farrell

9mo

Your operating system
A traditional OS execut
execution.



Sign in to view more
content

Create your free account or sign in to continue your
search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to
LinkedIn's [User Agreement](#), [Privacy Policy](#), and
[Cookie Policy](#).

wn.

edit a binary mid-

[...more](#)

Like

Comment

Share

To view or add a comment, [sign in](#)



1,391 followers

599 Posts

View Profile

Explore content

Career

Productivity

Project Management

Show more

© 2026

Accessibility

Privacy Policy

Cookie Policy

Brand Policy

Community Guidelines



Sign in to view more
content

Create your free account or sign in to continue your
search

or

New to LinkedIn? [Join now](#)

By clicking Continue to join or sign in, you agree to
LinkedIn's [User Agreement](#), [Privacy Policy](#), and
[Cookie Policy](#).

Guest Controls

Language