

CHATGPT

Die echten Slash-Befehle

Warum die fünf viralen Geheimcodes keine sind und welche Befehle es tatsächlich gibt

24

echte Befehle

0 von 5

Codes sind echt

1 Taste

zeigt deine Liste

Der Schrägstrich schaltet nichts frei, wenn nichts dahinterliegt.

Ein virales Prompt wird nicht zur Funktion, nur weil ein Slash davorsteht. Echte Befehle stehen dagegen in einer Liste, die aufklappt, sobald du einen Schrägstrich eintippst. Was nicht in dieser Liste steht, ist einfach Text.

01 Die fünf Codes einzeln geprüft

Auf Seite 2 steht bei jedem, was behauptet wird und was wirklich passiert. Zwei von fünf sind komplett erfunden.

02 Die echte Liste

Auf Seite 3 und 4 alle 24 Befehle aus der Dokumentation, sortiert nach Alltag und Projektarbeit.

03 Für jeden Fake-Code der echte Weg

Auf Seite 5. Für drei der fünf gibt es einen richtigen Befehl, der genau das macht, was der Code verspricht.

Quelle: die Befehlsreferenz von OpenAI, geprüft am 3. September 2026. Die Liste gilt für die Desktop-App, im Browser und am Handy siehst du weniger. Details dazu auf Seite 6.

DER FAKTENCHECK

Die fünf Codes, **einzel**n geprüft

Nicht alles ist gleich falsch, und genau das ist der Punkt

CODE	WAS WIRKLICH PASSIERT
EL10	Wirkt, weil es als Anweisung verstanden wird. Kein Code.
ALT3	Wirkt meistens, liest sich als Bitte um drei Varianten.
/human	Wird etwas schlichter. Der Zusatz zu Detektoren ist falsch.
KILLCRITIC	Schubst leicht. Eine echte Anweisung wirkt zuverlässiger.
X10THINK	Gibt es nicht. Es existiert kein Regler dafür in Textform.

Der Unterschied, auf den es ankommt.

Zwei davon tun etwas, weil eine Anweisung immer etwas tut. Wenn du schreibst, er soll es dir wie einem Zehnjährigen erklären, macht er das. Nur ist das kein freigeschalteter Modus, sondern eine normale Bitte, und du könntest sie genauso gut ausschreiben.

01 Wo es kippt

Bei X10THINK und beim Detektor-Versprechen von slash human. Es gibt keinen Mechanismus, der die Denktiefe auf ein Zauberwort hin verzehnfacht, und ein Schreibstil macht keine Kennzeichnung unsichtbar. Das sind die zwei Stellen, an denen dir jemand in den Kommentaren widerspricht, und er hätte recht.

DIE ECHTE LISTE, TEIL 1

Was du **täglich** brauchst

Die Befehle, die für jeden etwas bringen

BEFEHL	WAS ER MACHT
/goal	Ein Ziel setzen, an dem er dauerhaft weiterarbeitet
/plan	Planmodus für mehrere Schritte an- und ausschalten
/reasoning	Die Denktiefe einstellen
/model	Das Modell für diesen Chat wählen
/personality	Den Antwortstil wählen, wo unterstützt
/status	Chat-Kennung, Kontextverbrauch und deine Limits
/compact	Den Kontext des laufenden Chats eindampfen
/side	Einen kurzen Nebenchat aufmachen
/memories	Einstellen, wie das Gedächtnis genutzt wird
/fast	Auf die schnelle Stufe umschalten, wo verfügbar
/feedback	Rückmeldung mit Protokoll an OpenAI schicken
/pet	Das Schreibtier holen oder wegräumen

Die zwei, die kaum jemand kennt und die am meisten bringen.

Slash Status zeigt dir schwarz auf weiß, wie voll dein Kontext ist und wo du bei den Limits stehst. Und Slash Side macht einen Nebenchat auf, in dem du kurz etwas nachfragen kannst, ohne deinen Hauptchat mit dem Zwischenkram zuzumüllen.

DIE ECHTE LISTE, TEIL 2

Wenn du damit **baust**

Die Befehle für Projekte und Code, die es nur in der Desktop-App gibt

BEFEHL	WAS ER MACHT
<code>/project</code>	Ein Projekt für neue Chats auswählen
<code>/task</code>	Einen Chat ohne Projekt starten
<code>/local</code>	Den Chat im gewählten lokalen Projekt laufen lassen
<code>/cloud</code>	Den Chat in der Cloud laufen lassen, wo verfügbar
<code>/cloud-environment</code>	Die Cloud-Umgebung für den Chat wählen
<code>/init</code>	Ein Grundgerüst als AGENTS.md fürs Projekt anlegen
<code>/review</code>	Die Code-Durchsicht starten
<code>/fork</code>	Einen lokalen Chat in einen neuen kopieren
<code>/worktree</code>	Den Chat in einem neuen Git-Arbeitsbaum laufen lassen
<code>/ide-context</code>	Den geteilten Editor-Kontext an- und ausschalten
<code>/mcp</code>	Die angebundenen MCP-Server anzeigen
<code>/approve</code>	Eine automatisch abgelehnte Prüfung einmal freigeben

Das ist der Teil, an dem du merkst, wohin OpenAI zielt.

AGENTS.md, Git-Arbeitsbäume und Editor-Kontext sind keine Chat-Funktionen mehr, das ist eine Arbeitsumgebung für Entwickler. Wenn du diese Befehle nicht siehst, arbeitest du im Browser statt in der Desktop-App.

DER TAUSCH

Für jeden Code der **echte Weg**

Das ist die Seite, die du behalten solltest

STATT	NIMM
X10THINK	/reasoning und dort die Denktiefe hochstellen
/human	/personality, dazu einen Stil-Prompt in den Anweisungen
KILLCRITIC	Eine echte Anweisung, siehe unten
EL10	Einfach ausschreiben, wirkt genauso
ALT3	Einfach ausschreiben, wirkt genauso

STATT KILLCRITIC

Widerspruch mir. Such die drei schwächsten Stellen in dem, was ich gerade geschrieben habe, und sag mir bei jeder, woran sie scheitert. Stimme mir nur zu, wenn du es begründen kannst.

01 Warum das besser wirkt

Ein Wort wie killcritic muss er sich zusammenreimen. Eine Anweisung mit einer Zahl und einer klaren Aufgabe kann er abarbeiten und selber prüfen, ob er sie erfüllt hat. Das ist derselbe Unterschied wie zwischen mach es modern und nimm diese Schriftart.

02 Slash reasoning ist der eigentliche Fund

Genau das, was X10THINK behauptet, gibt es wirklich, nur eben als Einstellung und nicht als Zauberwort. Du stellst die Denktiefe hoch, er rechnet länger und antwortet gründlicher. Bei einfachen Fragen stellst du sie runter und sparst dir die Wartezeit.

WICHTIG

Warum deine Liste **anders** aussieht

Bevor du dich wunderst, dass die Hälfte fehlt

01 Die Liste ist nicht bei allen gleich

In der Dokumentation steht ausdrücklich, dass die verfügbaren Befehle von deiner Umgebung und deinem Zugang abhängen. Also vom Tarif, vom Gerät und davon, in welcher Testgruppe dein Konto steckt. Es gibt keine feste Liste, die für jeden gilt.

02 Die Befehle auf Seite 4 sind Desktop-App

Projekte, Arbeitsbäume und Editor-Kontext gibt es im Browser nicht. Wenn du ChatGPT nur im Browser oder am Handy nutzt, siehst du davon nichts, und das ist kein Fehler bei dir.

03 So siehst du deine eigene Liste

Neuen Chat aufmachen, einen Schrägstrich eintippen und nicht weiterschreiben. Dann klappt genau die Liste auf, die für dein Konto gilt. Weitertippen filtert, also zeigt dir slash sta direkt den Status-Befehl.

Und der Trick, mit dem du jeden angeblichen Geheimcode selbst prüfen kannst.

Tipp den Schrägstrich ein und schau, ob der Befehl in der Liste auftaucht. Steht er nicht drin, ist er kein Befehl, sondern Text. Das dauert zwei Sekunden und funktioniert bei jedem neuen Code, der demnächst durch die Feeds geht.

NÄCHSTER SCHRITT

Direkt umsetzen

Drei Minuten, und du kennst deine eigene Liste

01 Schrägstrich eintippen und hinschauen

Neuer Chat, ein Schrägstrich, fertig. Vergleich die Liste, die aufklappt, mit den Seiten 3 und 4 und du weißt sofort, was dein Konto kann.

02 Einmal slash status aufrufen

Da steht, wie voll dein Kontext ist und wo du bei den Limits stehst. Das ist die Zahl, die erklärt, warum er in langen Chats schlechter wird.

03 Die Denktiefe einmal hochstellen

Nimm slash reasoning und stell sie hoch, dann gib ihm eine Aufgabe, an der er sonst scheitert. Danach weißt du, wofür sich das Warten lohnt und wofür nicht.

Du willst wissen, was von den nächsten Tricks wirklich stimmt?

Auf Instagram prüfe ich, was gerade durch die Feeds geht, und sage dir, was davon übrig bleibt.

→ [@timhuck.ai auf Instagram](#)

→ www.shortmedia.agency

Alle Befehle aus der Befehlsreferenz von OpenAI, geprüft am 3. September 2026. Die Liste ändert sich laufend und unterscheidet sich je nach Tarif und Gerät, deswegen gilt am Ende immer das, was bei dir aufklappt.