

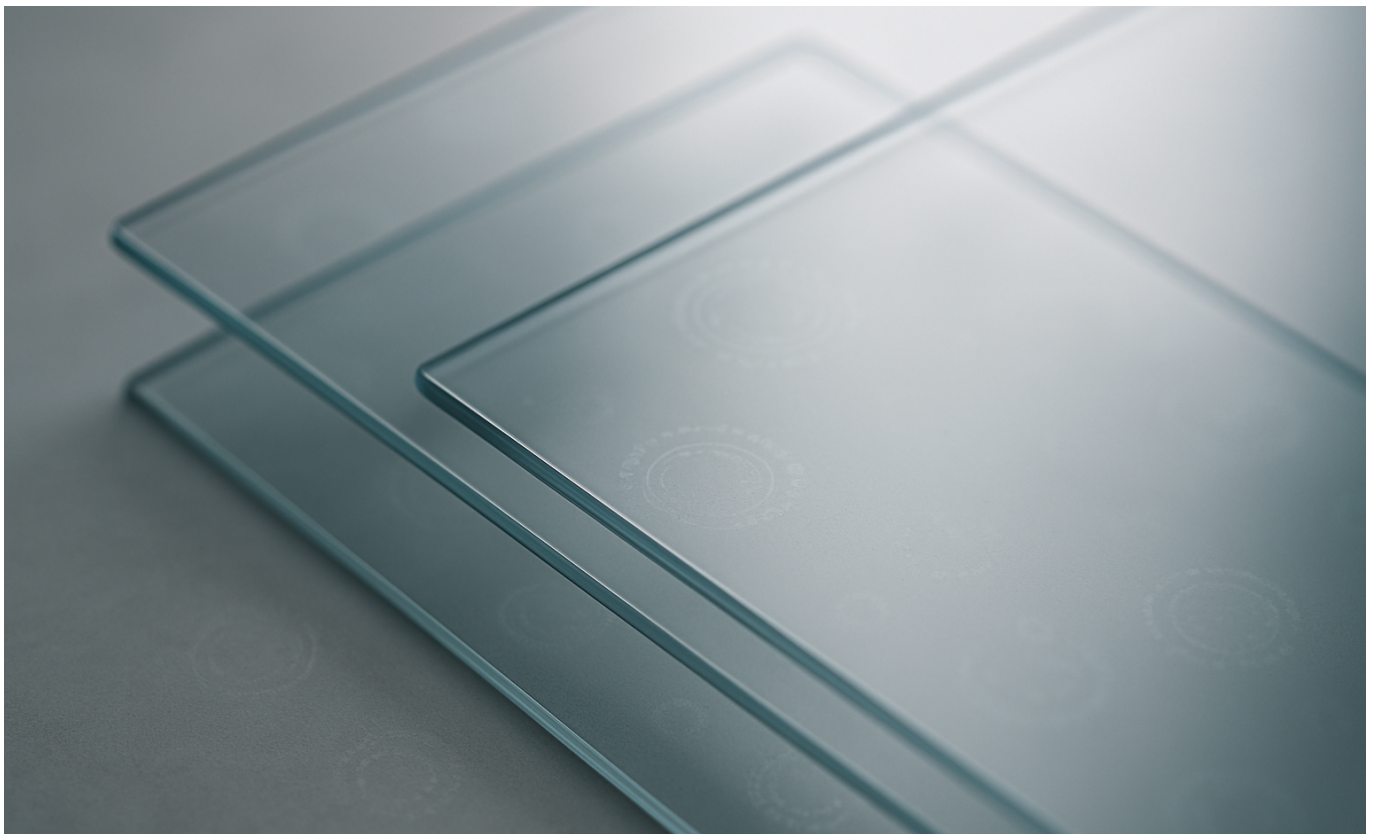
 Diese Beiträge werden vollautomatisch von einem KI-System erstellt und veröffentlicht – ohne menschliche Vorab-Prüfung. Kennzeichnung gemäß Art. 50 der KI-Verordnung (EU) 2024/1689.

Hinweis zur heutigen Ausgabe:

Heute leicht reduzierte Ausgabe.

KI-4-Everyone • Daily News

16. August 2026



SAFE

Anthropic erklärt Claudes unsichtbare Wasserzeichen

Anthropic beantwortet jetzt häufige Fragen zu Claudes neuen Wasserzeichen. Nicht alle sind von der unsichtbaren Markierung überzeugt.

SAFE

Anthropic-Chef: KI-Kritik wurzelt in Vertrauenskrise

Dario Amodei widerspricht dem Vorwurf, er male KI zu düster. Das Problem sei kein Pessimismus, sondern fehlendes Vertrauen.

Anthropic markiert Claude-Texte unsichtbar - und erklart jetzt, wie das funktioniert

Das KI-Unternehmen reagiert auf Fragen und Kritik zu seinem neuen Wasserzeichen, mit dem sich von Claude erzeugte Texte spaeter erkennen lassen sollen.

Wer heute einen Text liest, weiss oft nicht mehr, ob ein Mensch oder eine KI ihn geschrieben hat. Anthropic, das Unternehmen hinter dem Chatbot Claude, will genau daran etwas aendern - und baut in seine KI-Ausgaben ein unsichtbares Wasserzeichen ein. Das soll fuer mehr Transparenz sorgen. Doch nicht alle sind von der Idee begeistert, und im Netz kursieren viele Fragen dazu, wie diese Markierung eigentlich funktioniert. Jetzt liefert Anthropic Antworten nach.

Laut den beiden Berichten hat Anthropic weitere Details zu seinem neuen Wasserzeichen fuer Claude veroeffentlicht. Ein Wasserzeichen ist in diesem Fall keine sichtbare Grafik, sondern eine Art Signatur, die in den Text eingewoben ist und mit passenden Werkzeugen wieder ausgelesen werden kann. Das Unternehmen geht auf drei Kernfragen ein, die derzeit besonders oft gestellt werden: Wie funktioniert die Markierung technisch, laesst sie sich durch nachtraegliches Editieren entfernen, und was bedeutet das Ganze fuer Programmcode, den Claude erzeugt? Die konkreten technischen Antworten werden in den Quellen angerissen, aber nicht in Detailtiefe im hier vorliegenden Material wiedergegeben.

Der Schritt ist deshalb relevant, weil er ein Grundproblem der aktuellen KI-Welle beruehrt: Immer mehr Inhalte im Netz, in Bewerbungen, in Hausarbeiten oder in Nachrichten stammen ganz oder teilweise aus Chatbots. Verlage, Schulen, Behoerden und Plattformen haetten gerne ein zuverlaessiges Mittel, um KI-Text von Menschen-Text zu unter-

scheiden. Wenn Anthropic seine Ausgaben ab Werk markiert, koennte das ein Standard werden, dem andere Anbieter folgen - oder es koennte, je nach Sichtweise, den Anreiz erhoehen, auf Wettbewerber ohne solche Markierungen auszuweichen. Der Verweis auf 'nicht alle sind begeistert' im deutschen Bericht deutet an, dass es sowohl unter Nutzern als auch in Fachkreisen Widerstand gibt, auch wenn die Gruende im Material nicht ausbuchstabiert werden.

Offen bleibt vor allem, wie robust die Markierung wirklich ist. Der englischsprachige Bericht stellt selbst die Frage, ob sich das Wasserzeichen durch Editieren verstecken laesst - ob und wie klar Anthropic sie beantwortet, geht aus dem hier vorliegenden Material nicht eindeutig hervor. Auch die Frage nach Code ist heikel: Programmcode hat eine strenge Syntax, jedes zusaetzliche Zeichen kann ihn kaputtmachen, und Entwickler duerften genau hinschauen, ob ihr von Claude generierter Code veraendert wird. Unklar ist zudem, ob und wann das Wasserzeichen fuer alle Nutzer aktiv wird, ob es abgeschaltet werden kann und wer die Werkzeuge zum Auslesen bekommt.

Fuer die naechsten Tage lohnt der Blick darauf, wie Konkurrenten wie OpenAI oder Google auf den Vorstoss reagieren - und ob unabhaengige Fachleute die Markierung testen und Wege finden, sie zu umgehen. Genau daran wird sich zeigen, ob Wasserzeichen ein glaubwuerdiger Baustein gegen KI-Taeuschung werden oder nur ein Etikett, das bei der ersten Ueberarbeitung abfaellt.