

Aufgedeckt: So intransparent sind große KI-Modelle

[Originalartikel](#)

[Backup](#)

```
<html> <figure class=„aufmacherbild“><img
src=„https://heise.cloudimg.io/width/700/q75.png-lossy-75.webp-lossy-75.foil1/_www-heise-de_/imgs/
18/4/3/2/4/9/2/1/shutterstock_767827225-efbd8bfb498d84d0.jpeg“
srcset=„https://heise.cloudimg.io/width/700/q75.png-lossy-75.webp-lossy-75.foil1/_www-heise-de_/im
gs/18/4/3/2/4/9/2/1/shutterstock_767827225-efbd8bfb498d84d0.jpeg 700w,
https://heise.cloudimg.io/width/1050/q75.png-lossy-75.webp-lossy-75.foil1/_www-heise-de_/imgs/18/4/
3/2/4/9/2/1/shutterstock_767827225-efbd8bfb498d84d0.jpeg 1050w,
https://heise.cloudimg.io/width/1500/q75.png-lossy-75.webp-lossy-75.foil1/_www-heise-de_/imgs/18/4/
3/2/4/9/2/1/shutterstock_767827225-efbd8bfb498d84d0.jpeg 1500w,
https://heise.cloudimg.io/width/2300/q75.png-lossy-75.webp-lossy-75.foil1/_www-heise-de_/imgs/18/4/
3/2/4/9/2/1/shutterstock_767827225-efbd8bfb498d84d0.jpeg 2300w“ width=„6431“ height=„3613“
sizes=„(min-width: 80em) 43.75em, (min-width: 64em) 66.66vw, 100vw“ alt=„KI-Gestalt auf
schwarzem Hintergrund, der aussieht wie ein Weltraum“ class=„img-responsive“ referrerpolicy=„no-
referrer“ /><figcaption class=„akwa-
caption“>(Bild:&#160;metamorworks/Shutterstock.com)</figcaption></figure><p><strong>KI-
Forscher haben zehn Foundation Models hinsichtlich ihrer Transparenz bewertet. Das Ergebnis zeigt,
dass durch die Bank Nachholbedarf besteht.</strong></p><p>Ein Team aus KI-Forschern der
Universit&#228;ten Stanford, MIT und Princeton hat einen Transparenzindex f&#252;r Foundation
Models erstellt, also f&#252;r gro&#223;e KI-Modelle. Anhand von 100 Faktoren bewertet es zehn
bekannte Modelle wie GPT-4, Stable Diffusion 2 und PaLM 2. Volle Transparenz w&#228;re bei einem
Score von 100 Prozent erreicht.</p><p>Das Ergebnis wirkt ern&#252;chternd: Die Spitzenreiter
erreichen einen Score, der nur knapp &#252;ber 50 Prozent liegt, die rote Laterne tr&#228;gt
Amazons Titan Text mit einem Score von gerade einmal 12 Prozent.</p><h3 class=„subheading“
id=„nav_transparenz_vom0“>Transparenz vom Erstellen bis zum Einsatz</h3><p>Die Motivation
hinter dem Foundation Model Transparency Index (FMTI) liegt laut dem Team darin, dass es den
gro&#223;en KI-Modellen zunehmend an Transparenz mangle. Unternehmen k&#246;nnen daher
schwer einsch&#228;tzen, ob sie die Foundation Models problemlos in ihre Anwendungen integrieren
k&#246;nnen. Ebenso ben&#246;tige sowohl die Forschung als auch Endanwender Informationen zur
Transparenz beim Einsatz von KI.</p><p>Das FMTI-Team hat zehn verbreitete Modelle untersucht
und bewertet. An der Spitze steht Metas Large Language Model (LLM) LLaMA 2 mit einem Score von
54 Prozent. Dicht dahinter folgt BLOOMZ von Hugging Face mit 53 Prozent, und GPT-4 von OpenAI
nimmt mit 48 Prozent den dritten Platz ein.</p><figure class=„a-inline-image a-u-
inline“><div><img alt=„“ class=„legacy-img c1“ height=„370“ sizes=„“
src=„https://heise.cloudimg.io/width/696/q85.png-lossy-85.webp-lossy-85.foil1/_www-heise-de_/imgs/
18/4/3/2/4/9/2/1/Transparency_model_image_1_0-01cd0b2ecbfd2a10.jpg“
srcset=„https://heise.cloudimg.io/width/336/q70.png-lossy-70.webp-lossy-70.foil1/_www-heise-de_/im
gs/18/4/3/2/4/9/2/1/Transparency_model_image_1_0-01cd0b2ecbfd2a10.jpg 336w,
https://heise.cloudimg.io/width/1008/q70.png-lossy-70.webp-lossy-70.foil1/_www-heise-de_/imgs/18/4/
3/2/4/9/2/1/Transparency_model_image_1_0-01cd0b2ecbfd2a10.jpg 1008w,
https://heise.cloudimg.io/width/696/q70.png-lossy-70.webp-lossy-70.foil1/_www-heise-de_/imgs/18/4/3
/2/4/9/2/1/Transparency_model_image_1_0-01cd0b2ecbfd2a10.jpg 696w,
https://heise.cloudimg.io/width/1392/q70.png-lossy-70.webp-lossy-70.foil1/_www-heise-de_/imgs/18/4/
3/2/4/9/2/1/Transparency_model_image_1_0-01cd0b2ecbfd2a10.jpg 1392w“ width=„696“
referrerpolicy=„no-referrer“ /></div><figcaption class=„a-caption“>Das Team hat zehn Foundation
```

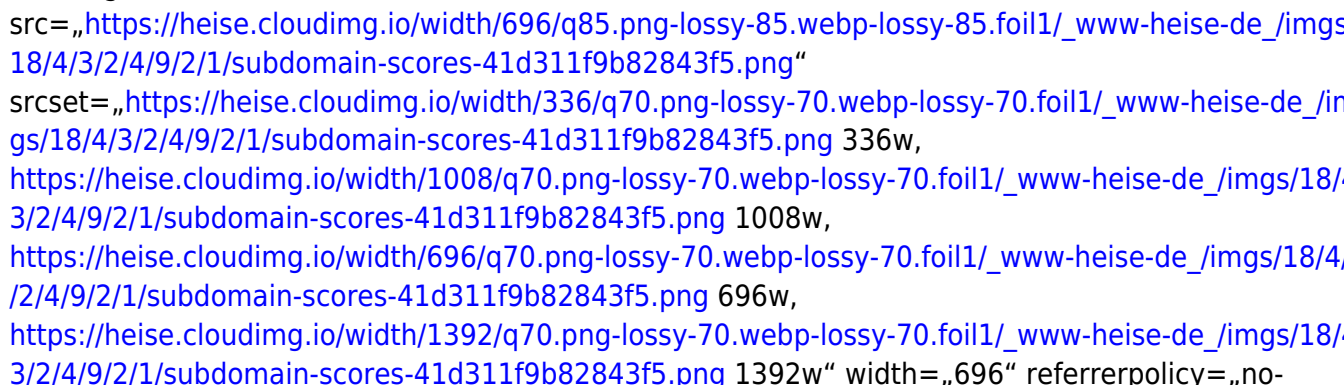
Models in unterschiedlichen Bereichen auf Transparenz untersucht, um einen Score zu ermitteln. (Bild:  Stanford University [1])

100 Indikatoren aus drei

Modellstufen

Die Bewertung hat das Team 100 Indikatoren zusammengestellt, die in die Bereiche Upstream, Modell und Downstream unterteilt sind. Die vorgelagerten (Upstream) Faktoren beschreiben den Prozess zum Erstellen des Modells, darunter die Datenquellen, geografische Verbreitung und Rechenressourcen

das Training der Foundation Models. Zu den Modellindikatoren gehören unter anderem die Architektur, die Fähigkeiten und die Limitierungen des Modells. Schließlich finden sich in den nachgelagerten Faktoren der Release- und Update-Prozess, die Lizenz sowie die Auswirkung des Modells auf User und



Die drei Hauptbereiche hat das Team in dreizehn Unterbereiche unterteilt, darunter Daten, Methoden und Risiken. (Bild: Stanford University)

Nachdem das Team die Scores erstellt hatte, gab es den Verantwortlichen die Modelle die Gelegenheit zur Stellungnahme und passte die Werte bei berechtigten Einwendungen an. Die Bewertung zeichnet kein positives Bild bezüglich der Transparenz, mit einem Sieger, der knapp über 50 Prozent liegt, und einem Durchschnittsscore von 37 Prozent.

Auffällig und wenig überraschend ist, dass die drei offenen Modelle LLaMA 2, BLOOMZ und Stable Diffusion 2 in den ersten vier Plätzen zu finden sind. Dass OpenAIs Modell auf dem dritten Platz landet, ist dagegen durchaus eine kleine Überraschung, da sich das Unternehmen seinem Namen zum Trotz bei den Details zu seinen Modellen bedeckt hält.



Die offenen Modelle kommen vor allem in den vorgelagerten Faktoren beim Training des Modells (grüner Anteil der Balken) punkten. (Bild: Stanford University)

Die Wurzeln des FMTI-Teams, das weitgehend aus Studenten, Doktoranden und einer Doktorandin sowie Forschungsleitern und einem Professor besteht,

href=„<https://www.heise.de/news/Stanford-University-startet-Institute-for-Human-Centered-AI-4341355.html>“>liegen im 2019 gegründeten [2] Stanford Institute for Human-Centered Artificial Intelligence. Das Team sieht bei allen getesteten Foundation Models „erhebliches Verbesserungspotenzial“, das es in künftigen Versionen des Index verfolgen möchte.</p><p>Weitere Details lassen <a href=„<https://crfm.stanford.edu/fmti/>“ rel=„external noopener“ target=„_blank“>sich der FMTI-Website [3] und dem zugehörigen <a href=„<https://github.com/stanford-crfm/fmti>“ rel=„external noopener“ target=„_blank“>GitHub-Repository entnehmen [4]. Die vollständige <a href=„<https://arxiv.org/abs/2310.12941>“ rel=„external noopener“ target=„_blank“>Abhandlung findet sich auf arXiv [5].</p><p>() </p><hr /><p>URL dieses Artikels:
<small>

<https://www.heise.de/-9346416>

</small></p><p>Links in diesem Artikel:
<small>

[1]

</small>
<small>

[2]

</small>
<small>

[3]

</small>
<small>

[4]

</small>
<small>

[5]

</small>
<small>

[6]

</small>
</p><p class=„printversion__copyright“>Copyright © 2023 Heise Medien</p> </html>

From:
<https://schnipsl.qgelm.de/> - Qgelm

Permanent link:
https://schnipsl.qgelm.de/doku.php?id=wallabag:wb2aufgedeckt_-so-intransparent-sind-groe-ki-modelle

Last update: 2025/06/27 11:17

