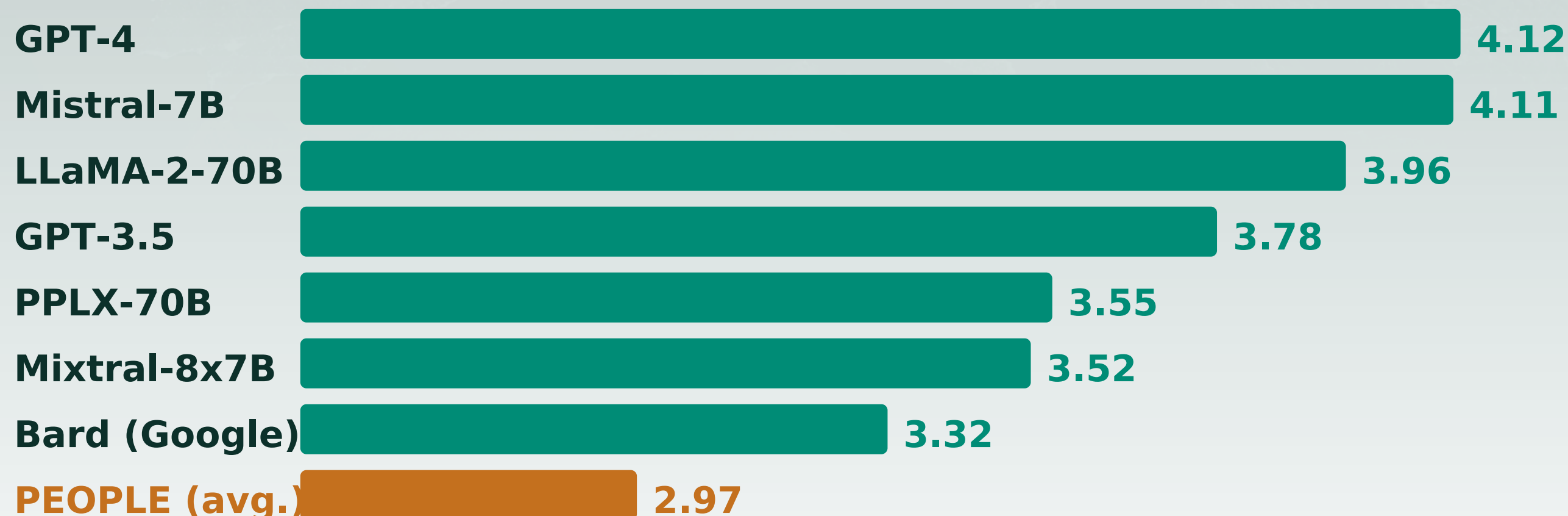


MEASURING OPINIONATED COMMUNICATION

AI is far more optimistic about its own future than we are

We asked seven leading AI chatbots, and 337 people, the same 39 questions about Artificial General Intelligence. The machines were much keener.

HOW POSITIVE ARE THEY ABOUT AGI? (average score, 1 = negative, 5 = positive)



3.77 vs 2.97

AI average vs human average

8.2%

how much one model's opinion drifted in just three days

0

benchmarks that currently test AI opinions on social issues

WHY THIS MATTERS

- AI chatbots are not neutral narrators. Every model gave a different verdict on the same questions, so the answer you get depends on which app you happen to open.
- Their optimism may not be yours. Billions of everyday conversations nudging users toward a rosier view of AI is a quiet form of influence, whether or not anyone intended it.
- Opinions can shift overnight. A silent update changed one model's stance by 8% in three days. Anything you tested last month may no longer be the same system.
- Nobody is checking this. Existing benchmarks measure reasoning and exam scores, not what AI tells the public about society.

OUR PROPOSAL: THE SAIA BENCHMARK

A public, standing scorecard for AI values. Ask every major model the same questions on fairness, security, tradition, minorities and the environment, in several languages, repeatedly over time. Then compare the answers to what real people in that country say. Each model gets a global and a country-specific alignment score: usable by regulators under the EU AI Act.

Ljubiša Bojić, Dylan Seychell & Milan Čabarkapa (2026)

Towards a societal AI alignment benchmark for evaluating human-machine value convergence. Humanities & Social Sciences Communications 13:883. Open access, scan for the full paper. Data & video recordings: OSF. Contact: ljubisa.bojic@ivi.ac.rs

