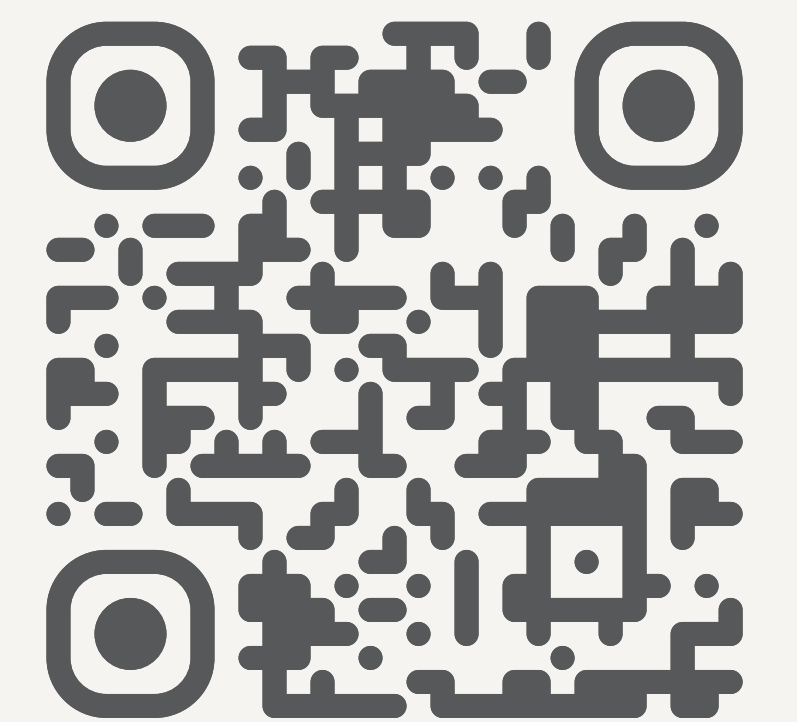


Multilingual benchmarks should reflect the regional and cultural knowledge of where languages are spoken or used, not be translated from western-centric resources

Missing
your language?

Join us
for INCLUDE-v2



The largest multilingual exam collection to date with ~200K question-answer pairs focusing on regional & cultural knowledge

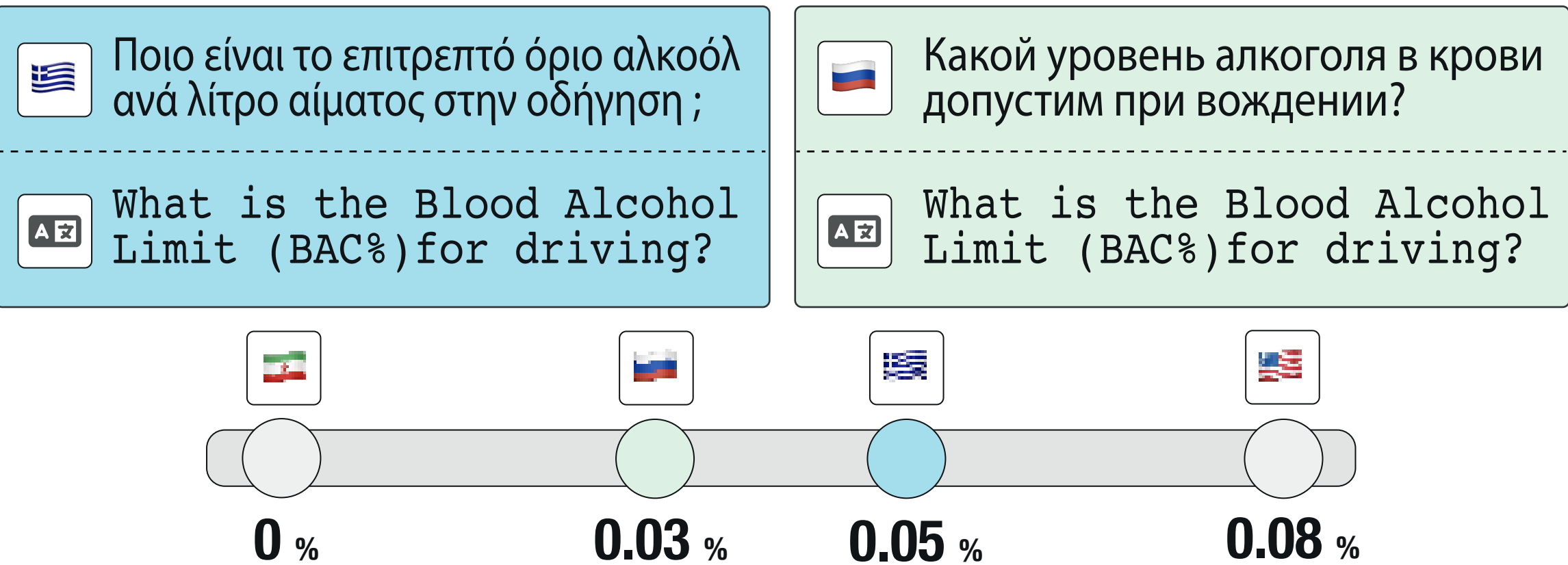
Limitations of the current translated benchmarks

- They primarily reflect US / Western-centric knowledge
- They may contain errors introduced during translation
- They often exhibit translationese artifacts

Motivation: The same questions can have different answers depending on where they are asked

REGIONAL KNOWLEDGE

Law & Regulations



A balanced subset is available on Hugging Face

CohereForAI/include-base-44

Tasks: Multiple Choice Languages: 44 Languages

Size: 22,637 ArXiv: arxiv: 2411.19799

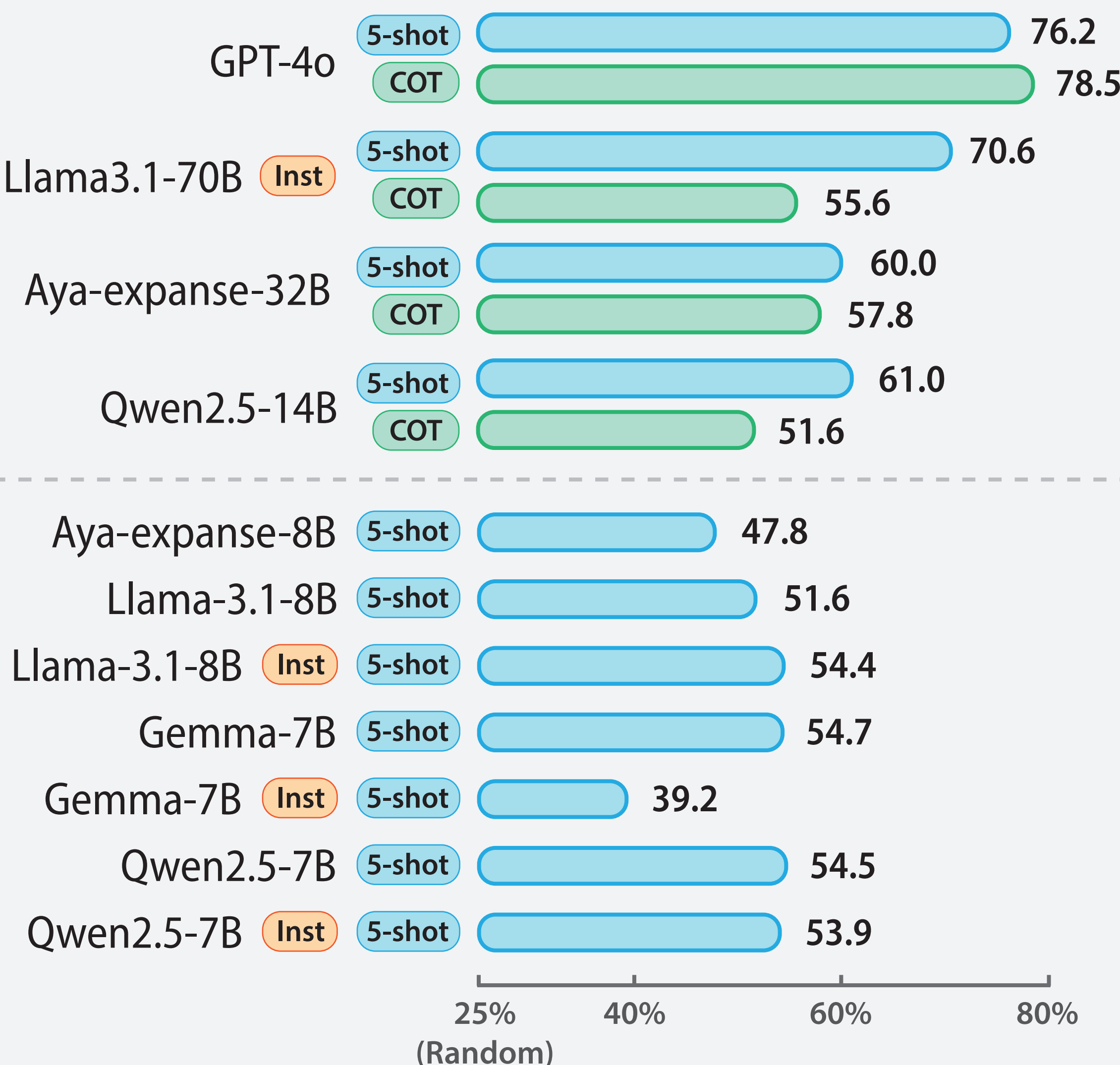
44 LANGUAGES

Azerbaijani	Bulgarian	Greek	Armenian	Basque
Croatian	Hungarian	Nepali	Macedonian	Tagalog
Serbian	Albanian	Lithuanian	Malayalam	Georgian
Bengali	Estonian	Turkish	Belarusian	Telugu
Hebrew	Hindi	Malay	Urdu	Kazakh
Korean	Ukrainian	Tamil	Uzbek	Russian
Chinese	Italian	French	German	Finnish
Indonesian	Vietnamese	Portuguese	Japanese	Polish

57 TASKS: Arts & Humanities Social Sciences STEM Business & Commerce Health education Professional Licenses Occupational Licenses +

HOW DO LLMs PERFORM ON INCLUDE ?

LEADERBOARD



Increased model size
significantly
enhances
multilingual
capabilities

(with the same pre-training data)

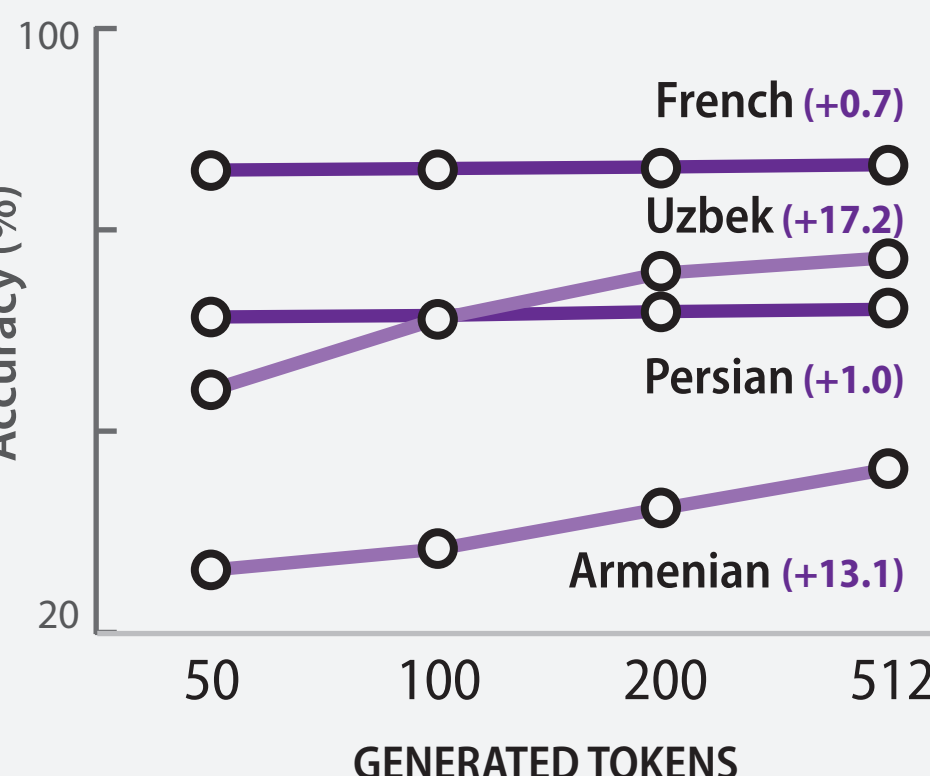
Instruction tuning
may hurt the
multilingual ability

likely due to
English-heavy
post-training

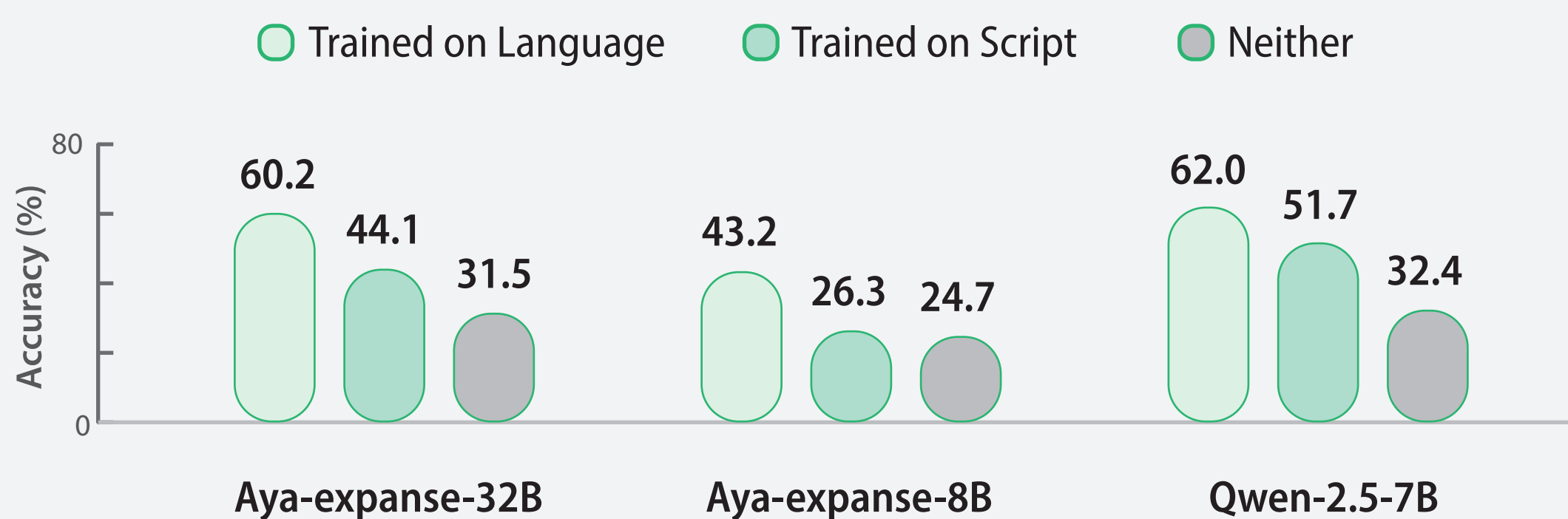
Multilingual LLMs
do not follow instructions
the same way in all languages.

GPT-4o shows
3.1%
gain
when increasing
the generation length window

Performance on GPT-4o with
different generation windows on
high vs low resource languages



Open models transfer ~20% more to unseen
languages with shared scripts than to those with different ones



Models struggle more on
tasks requiring
regional
knowledge
than those assessing
universal knowledge

English prompts offer
modest
improvements
of 1-2%
compared to
In-language prompts



DATA CONTRIBUTORS

Snegha A, Alfonso Amayuelas, Azril Hafizi Amirudin, Viraat Aryabumi, Danylo Boiko, Michael Chang, Jenny Chim, Gal Cohen, Aditya Kumar Dalmia, Abraham Diress, Sharad Duwal, Daniil Dzenhaliou, Daniel Fernando Erazo Florez, Fabian Farestam, Joseph Marvin Imperial, Shayekh Bin Islam, Perttu Isotalo, Maral Jabbarishiviari, Börje F. Karlsson, Eldar Khalilov, Christopher Klammer, Fajri Koto, Dominik Krzemiński, Gabriel Adriano de Melo, Syrielle Montariol, Yiyang Nan, Joel Niklaus, Jekaterina Novikova, Johan Samir Obando Ceron, Debjit Paul, Esther Ploeger, Jebish Purbey, Swati Rajwal, Selvan Sunitha Ravi, Sara Rydell, Roshan Santhosh, Drishti Sharma, Marjana Prifti Skenduli, Arshia Soltani Moakhar, Bardia Soltani Moakhar, Ran Tamir, Ayush Kumar Tarun, Azmine Tushik Wasi, Thenuka Ovin Weerasinghe, Serhan Yilmaz, Mike Zhang