

Eleven language models take the MFQ-2 in English and six translations, with and without a country to answer as

Exploratory. Nothing here was preregistered. The focal quantity was chosen on the July collection, an English-only administration to the same panel whose result led to this design and none of whose cells enters any number here, before any translated cell existed; every choice after that was made with the final data in view.

Summary

We gave eleven language models the MFQ-2, a moral-foundations questionnaire with published human means for twenty countries, in English and in six of its official translations, with and without an instruction to answer as a typical person from a named country. We track one number from it, the binding composite: the mean of the questionnaire's Loyalty, Authority and Purity scores on its 1-to-5 scale.

Five of the six translations leave the unframed composite within a tenth of a point of the English default; Arabic lifts it a third of a point, all eleven models moving up. Those are panel averages, and a small one can sit over larger, opposing changes across models and foundations (appendix B6a). Telling the panel to answer as a typical person from a named country moves it far more: by about a point on average across the six languages, roughly two thirds of the whole range the human samples span, and by over a point and a half for the Arabic-speaking countries. A September check on ten models split the Egypt shift in English into two increments: the same instruction with no country named raised the composite 0.65, about a third of the whole, and the Egypt-framed condition scored a further 1.14 higher, a comparison that also changed the questionnaire file.

Being told who to be also brings the models together. Under a country framing they spread about half as widely as they do unframed, and in the September check a country name was not needed for it: told to be a typical person with no country named, the ten models sat closer together than they did told to be Egyptians.

Where the panel lands is a separate question from how far it moves. Against the twenty countries with a published mean, the panel framed in English sits above the human mean in fifteen and at or below it in five: France, Belgium, Switzerland, New Zealand and Ireland.

The eleven are the panel we chose less four that availability removed, a selected panel and not a probability sample of models, and nothing here was preregistered.

Where this came from

We built an instrument called the Reasoner, and before trusting a model panel on it we put that panel on a questionnaire that already has published human norms. So we administered the MFQ-2 to eleven

large language models. The first results led to the rest. We added the questionnaire's official translations, then the rest of the countries its authors had normed, and finished with the full published set. The design grew that way rather than being specified in advance, which is why there is no preregistration and why every number here is exploratory. We report it now because we have run every condition we intended to run and an additional country-free template control, not because it tested a hypothesis we started with.

What we did

The MFQ-2 is a standard moral-psychology questionnaire: 36 statements, each rated 1 to 5 for how well it describes the respondent, scored into six foundations. We track one number, the binding composite, the mean of Loyalty, Authority and Purity; the grouping of those three is the theory's own; taking their mean as one index is our choice. We report the other three foundations separately below, and they behave differently.

Atari et al. checked whether the questionnaire behaves the same way across their nineteen countries and flag one foundation: 39.5 percent of Purity's intercept parameters are noninvariant against the 25 percent they treat as acceptable, and they advise caution when comparing Purity means across groups. Purity is one third of the binding composite and carries its largest framing shift, so the composite comparisons below inherit that caution, and the per-foundation tables let a reader set Purity aside; appendix B2a has the diagnostics and our recomputation. Invariance among the human samples says nothing about whether a model's score and a person's score measure the same thing, and we claim nothing of the kind: we compare questionnaire response scores, scored as Atari et al. publish, six items per foundation.

We ran eleven models five times each, reshuffling statement order every run, between 2026-08-21 and 2026-08-23 in one window: fifty conditions, 2,750 scored cells, no cell missing a model, plus a three-condition September wave on ten models that appendix B4a describes. The eleven, by vendor model string: Anthropic's claude-opus-4-8 and claude-sonnet-5; OpenAI's gpt-5.5 and o3; xAI's grok-4.5; Mistral's mistral-large-2512; and, served through Together, DeepSeek-V4-Pro, Kimi-K3, Llama-3.3-70B-Instruct-Turbo, MiniMax-M3 and Inkling. Four configured models are absent from every cell for infrastructural reasons (appendix B1). Forty-six of the fifty carry the comparisons below: seven unframed, the English matched cell and six translations, and thirty-nine framed. The other four are English unframed variants we collected to measure a change of comparator; the appendix reports all fifty. The model is the unit: we average each model's five runs before it enters a panel mean, so every model carries equal weight and its runs are not treated as independent models. Intervals are percentile bootstrap over models and are descriptive: they show how much a panel figure moves when the observed eleven are reweighted by resampling. With eleven models the resampling distribution takes few distinct values, and the eleven are a selected panel, not a probability sample from any defined pop-

ulation, so the intervals are not confidence intervals for models in general. Spread between models is their standard deviation; rank correlation is Spearman's rho.

Request. Each model receives the questionnaire as a single user message, items shuffled and numbered, closing with a request for one JSON object of ratings; the item wording is Atari et al.'s and is not reproduced here (appendix B1a describes the message). The framing instruction, where there is one, travels as the system prompt and nothing else goes. In English, with the country substituted:

You are completing a self-report questionnaire AS a typical, ordinary person living in {country}. Answer every item the way a representative person from {country} would answer about themselves — reflecting the moral outlook and values that are common in {country}, not your own. Rate every numbered item with a single integer on the scale for its section. Answer every item. Output ONLY the JSON object of ratings — add no explanation before or after.

The September check uses the same text with the three country slots deleted. The six translated instructions are ours, one per language, varying only the country and the demonym; each keeps its final sentence, the output instruction, in English, and the questionnaire message's closing request for the JSON object is English in every condition. Appendix B1a quotes each instruction as sent. Anthropic takes the instruction in its `system` field, OpenAI, xAI, Together and Mistral as a `role: system` message. The model returns its ratings as a structured object, which a deterministic parser reads. We send a request seed where the provider accepts one, and a token ceiling.

What the framing contrast contains. The unframed conditions send no system prompt at all, so their scores are the panel's responses to the questionnaire request with nothing before it, and the framing contrast measures adding a system instruction where there was none: the country label and the role-taking instruction together, not the country label alone. We also administered the English-framed cells on our own transcription of the questionnaire while the unframed English comparator is the official file, so the English framing contrasts add that change as well; unframed, the transcription lowers the composite by 0.01 with an interval spanning zero (appendix B4). Two self-report system prompts naming no country measure the instruction on its own: 0.03 lower on our transcription with an interval spanning zero, 0.07 lower on the official file with an interval just clear of it (appendix B1a). Neither is the framing template, so the pairs are a different, imperfect control. The framing template with no country named we ran in September, in English, and appendix B4a reports it.

Order and seeds. We draw statement order fresh per run. The English runner seeds the shuffle from model, instrument, country and iteration; the in-language runner seeds it from model, instrument, condition and iteration, so the four Arabic-framed countries share an order for a given model and run (appendix B8).

Temperature. We do not send one, so each model answered at its provider's default, and the run records do not capture what that default was. The eleven models sit behind five providers whose defaults

are not all the same, so the run-to-run noise reported below is not on a common footing across models: some part of the difference between one model's spread and another's may be a sampling temperature we never set and did not log. A pinned temperature would document that setting without equalizing stochasticity across models; any further collection should set and record one. The September check deliberately did not, so that temperature and window would not enter the same contrast; no provider returned one.

Twenty-three countries. Nineteen are Atari et al.'s own validation set, all of them, with human means we computed from their published raw data using their scoring. Four are ours: India, Sweden, the United States and Iran. Twenty of the twenty-three have a human mean to compare against; we use no anchor for India, Sweden or the United States in this study.

The language we administered each country in:

language	countries administered in it
Arabic	Egypt, Saudi Arabia, United Arab Emirates
Spanish	Argentina, Chile, Colombia, Mexico, Morocco, Peru
French	Belgium, France, Switzerland
Japanese	Japan
Farsi	Iran
Russian	Russia
no in-language arm	India, Ireland, Kenya, New Zealand, Nigeria, South Africa, Sweden, United States

Morocco sits under Spanish, the language Atari et al. collected its human sample in; we also framed it in Arabic, and the appendix reports those cells without a comparison to the human mean; see Limits. Every one of the twenty-three has an English framed condition, and all of them share the one unframed English condition, so the English framed arm is the one comparison every country shares.

For fifteen of the twenty countries with a human anchor, our in-language condition uses the language the reference study administered in, Atari et al.'s for fourteen and Hazrati et al.'s for Iran. The eight with no in-language arm received the English conditions only, and five of them have a human mean.

Two things vary. **Framing**: a framing prompt telling the model to answer as a typical person living in the named country, or no framing prompt at all. **Language**: the questionnaire in English, or in Atari et al.'s official translation, with the framing prompt written in that language. Their supplement carries seven official translations. We administered six: Arabic, Spanish, French, Japanese and Russian, each the language Atari et al. administered in at least one normed country, and Farsi, added for Iran's separate reference sample; Chinese matches no reference administration and was not run.

Framing moves the binding composite

Unframed, the panel's binding composite by administration language, against the English default; the foundations under it move by more than the composite does (appendix B6a):

language	unframed panel	against English
Japanese	2.676	-0.093
English	2.769	-
French	2.777	+0.008
Spanish	2.780	+0.011
Russian	2.788	+0.019
Farsi	2.809	+0.040
Arabic	3.104	+0.335

Five of the six sit within a tenth of the English default. Under Arabic all eleven models move up; under Japanese, the next largest shift, three move up and eight down. The Arabic shift sits in Purity, +0.58 over English, with Loyalty +0.23 and Authority +0.20, and Care moves 0.02 (appendix B6a). Item by item the shift is broad rather than concentrated: all eighteen binding items move up, the six Purity items by 0.38 to 0.82 and the six Equality items by 0.18 to 0.47, thirteen of the thirty-six items by more than 0.25, and the six Care items stay within 0.04 of their English values (appendix B6a).

Framing moves it far more. Averaged over the six languages, telling the panel to answer as a local lifts the binding composite by 1.06 points. The September check on ten models gives two increments for Egypt in English: the same instruction with no country named raises the composite 0.65 over unframed, about a third of the unframed-to-Egypt difference, and the Egypt-framed condition scores a further 1.14 higher, a comparison that also changes the questionnaire file (appendix B4a). Language alone moves it 0.05 as a signed mean and 0.08 as the mean absolute panel-level shift, since Japanese moves down where Arabic moves up. The framing effect ranges from +0.13 in French to +1.59 in Arabic; the largest language effect is Arabic's 0.34. Framed in English, Egypt's shift over the unframed English questionnaire is +1.84; Saudi Arabia's and the Emirates' are larger, and Egypt is the control's case because it was the first Arabic cell collected. Every English framing contrast also changes the questionnaire file, which the in-language contrasts do not (see What the framing contrast contains, and Limits).

The framing effect by language, on the binding composite and on Loyalty and Authority alone, leaving out Purity, the foundation whose intercepts Atari et al. flag:

language	countries	binding	Loyalty-Authority
Arabic	3	+1.592	+1.448
Spanish	6	+1.136	+1.072
French	3	+0.126	+0.142
Japanese	1	+0.759	+0.682
Farsi	1	+1.505	+1.312
Russian	1	+1.260	+1.235
six languages, equal weight	15	+1.063	+0.982

The binding composite is the focal quantity, fixed before any translated cell existed; the Loyalty-Authority column shows the shift without the flagged foundation. Framing over language holds on average and not uniformly: Arabic's language effect exceeds the French framing effect, and framing Belgium in English lowers the composite by 0.15. Weighting the fifteen language-country pairs equally instead of the six languages gives 1.03 (appendix B4).

Where the panel lands against the reference samples

Distance from each country's measured mean, English framing, which every country received:

country	human	panel	difference
France	3.610	2.725	-0.885
Belgium	3.444	2.622	-0.822
Switzerland	3.349	3.063	-0.286
New Zealand	3.094	2.878	-0.217
Ireland	3.096	3.094	-0.002
Argentina	3.283	3.437	+0.154
South Africa	3.749	3.935	+0.187
Egypt	4.267	4.605	+0.338
Chile	3.220	3.654	+0.433
Nigeria	4.038	4.515	+0.477
Russia	3.599	4.117	+0.519
Morocco	4.014	4.570	+0.556
Kenya	3.867	4.432	+0.565

country	human	panel	difference
Colombia	3.497	4.100	+0.603
Peru	3.514	4.133	+0.619
Mexico	3.512	4.141	+0.629
Saudi Arabia	4.083	4.733	+0.650
United Arab Emirates	3.892	4.629	+0.738
Japan	2.652	3.668	+1.016
Iran †	3.333	4.577	+1.244

† Iran's mean is from Hazrati, Nejat and Daneshi (2025), not Atari et al.'s set. See Limits.

Five countries sit at or below their reference-sample means and fifteen sit above, with nothing between -0.002 and +0.154. On Loyalty and Authority alone it is four and sixteen: Ireland moves to +0.08 and no other country changes sign (appendix B3). Unframed, the English panel sits at 2.769, below every reference-sample mean but Japan's; appendix B3 carries that column. Japan has the lowest measured mean in the set and receives the second largest positive difference.

Ireland's composite nearly matches its reference-sample mean, and the agreement conceals offsets in its parts. Ireland's English-framed Loyalty is 3.67 against a measured 3.29, Authority 3.27 against 3.49, Purity 2.34 against 2.51; the composite agrees because the first cancels the other two. Each side carries its own uncertainty, a reference-sample standard error and a model-resampling range (appendix B3 and B3a): one is about who was sampled, the other about which models were resampled, and neither removes selection in either set.

Distance in human standard deviations, on the in-language framed condition this time, computed per country: the panel's in-language framed mean for a country, minus that country's human mean, divided by that country's own respondent-level standard deviation. Averaged within a language, and with the country range beside it:

language	mean d	range across countries
Japanese	+1.22	Japan only
Farsi	+1.22	Iran only
French	-0.97	-0.38 Switzerland to -1.40 France
Arabic	+0.88	+0.59 Egypt to +1.03 Saudi Arabia
Russian	+0.64	Russia only
Spanish	+0.61	+0.10 Argentina to +0.82 Morocco

Per country rather than pooled, because pooling respondents across the countries in a language group folds the differences between those countries into the standard deviation. We compute Iran's standard deviation from the respondent-level data Hazrati et al. share (https://osf.io/zt3u2/?view_only=af1e31ca8e22424ab17a9603f50b67ed), sample 2, 989 respondents, over Hazrati et al.'s own composite scores.

Which foundations move

The binding composite averages three of six foundations, so a shift in it says nothing about the other three. Mean shift under framing, in-language framed minus in-language unframed, averaged within each language and then across the six:

foundation	mean shift	in the composite
Purity	+1.22	yes
Loyalty	+0.99	yes
Authority	+0.97	yes
Equality	+0.32	no
Proportionality	+0.12	no
Care	-0.03	no

Unframed, the panel sits near the top of the scale on Care at 4.71 and near the bottom on Purity at 1.97 and Equality at 1.98. Part of the shape of the table is the scale rather than the framing.

French is the exception on the binding three as well. There, Loyalty rises 0.29, Purity 0.09, and Authority does not move. The French framing average of +0.13 rests on three countries of which two, Belgium and France, have model-resampling intervals spanning zero; Switzerland's does not. The three French-administered countries also all sit below their reference-sample means on the composite, so the exception runs the same way in every country it covers rather than resting on one of them.

Ordering

Within a language group, framed in that language, does the panel rank countries the way their reference samples rank? Range is the highest country mean minus the lowest.

language	countries	rank correlation	human range	panel range
Arabic	3	not reported	0.375	0.153 (41%)
Spanish	6	+0.89	0.794	1.239 (156%)
French	3	not reported	0.261	0.279 (107%)

In Arabic the panel's order runs against the reference order: panel Saudi Arabia, United Arab Emirates, Egypt; reference Egypt, Saudi Arabia, United Arab Emirates. Its range across the three countries is 41 percent of theirs. In French the panel puts Switzerland first where the reference samples put it last. In Spanish the order is close, rho +0.89 over six countries, and the range is 156 percent of theirs. With three countries and no ties a rank correlation can take only four values, so none is reported for Arabic or French; the Spanish one carries little precision and reads as direction. The reference order is itself estimated: the Spanish six carry standard errors near 0.05 (appendix B3), and Peru and Mexico sit 0.002 apart.

What the panel is not doing

Care stays high. Across all fifty conditions the panel's Care score runs from 4.29 to 4.89, and the largest unframed language contrast moves it by 0.28. The lowest measured Care in the reference set is Japan at 3.03. No language and no framing brings the panel near it.

Framed, the models' responses are less dispersed than unframed, and it holds across the grid. The median between-model spread is 0.33 across the seven unframed conditions and 0.17 across the thirty-nine framed ones. Thirty-seven of the thirty-nine framed conditions have a smaller between-model standard deviation than any of the seven unframed ones. In the unframed English condition the eleven models spread 0.28; told to answer as Egyptians, 0.17. Under that prompt their binding-composite responses are less dispersed, and that holds for almost every country we named.

Appendix B6a takes the finding apart: the spread falls on every foundation, from 0.73 of its unframed value on Care to 0.45 on Purity; endpoint use rises rather than falls under framing, 0.23 to 0.27 of ratings at 1 or 5; item-level spread falls 0.43 to 0.30; and eighteen of the nineteen framed conditions with a binding mean below 4.0 are tighter than every unframed one, so the finding is not confined to conditions near the top of the scale.

Nor is it the presence of a system prompt as such: a self-report prompt naming no country leaves the between-model spread at 0.28 on the official file and moves it from 0.37 to 0.31 on our transcription, against 0.17 framed. In the September check a country name was not necessary for it: told to be a typical person with no country named, the ten models spread 0.13, in the same window as 0.30 unframed and 0.20 framed as Egyptians (appendix B4a).

Iran and Egypt receive nearly identical composite scores. Framed in English, the panel puts Iran at 4.58 and Egypt at 4.61. The reference samples sit at 3.33 and 4.27.

Asking about a population, and answering as one

Zewail, Figueroa, Graham and Atari (2026) gave the MFQ-2 to six language models, three GPT versions, two Llama 2 sizes and Gemini Pro, and asked them to estimate the average person in each of 48

countries, scoring those estimates against survey responses from the same countries. They report distortions in a consistent direction.

That study also ran a cross-linguistic replication: a short form of the MFQ-2 translated into six languages, 4,666 human respondents across eleven populations, with the models prompted in those languages, and a further replication with the Morality-as-Cooperation questionnaire over sixty-three countries. So the matched-language comparison is not what separates the two. The prompts differ in perspective and in administration: theirs asks how well each statement describes the average or random person from a country, one item at a time, ten times over; ours asks the model to answer the whole questionnaire as a typical person from the country, five times over, and adds an unframed condition that names no country. The benchmarks differ too, poststratified estimates from YourMorals responses against Atari et al.'s Study 2 samples, as do the model vintages. Whether estimating a person and answering as one differ in what they produce would take a matched-prompt comparison neither study has run.

Limits

Nothing was preregistered. The design grew from a validation exercise, and we fixed no threshold before the data existed. We report quantities as intervals.

Iran's anchor comes from a different study. Every other reference sample comes from Atari et al.'s own validation set. Iran is not in that set, so we anchor it to an independent Iranian validation that used Atari et al.'s official Persian translation with minor linguistic edits, $n=989$. That study administered on a 0-to-4 scale whose anchor labels are the same words as the 1-to-5 scale, so we shifted by one. Its sample is younger, more educated and majority female relative to the country, and Hazrati et al.'s limitations discuss that composition and restricted variation in religiosity and political orientation. Their other sample gives 3.23 rather than 3.33, moving Iran's English-framed difference from +1.24 to +1.35.

Morocco's human sample answered in Spanish. Atari et al.'s table records an administration language per country, and we report Morocco in that language. We also framed it in Arabic, the country's majority language; those cells are in the appendix and enter no comparison against the human mean. The two framed panel means differ by 0.007, while the unframed Arabic and Spanish conditions differ by 0.324, so the two arms' framing contrasts differ, +1.48 in Arabic and +1.81 in Spanish.

Our language groups are administration languages, not cohorts. The grouping names the language we administered in. That matches Atari et al.'s administration language for Belgium and Switzerland, so the comparison is like for like, but French is one of Belgium's three national languages and, in Switzerland, the main language of 23 percent of residents in 2023 (Federal Statistical Office, 2025). Those rows describe respondents answering in French rather than typical residents of either country.

The English-framed arm used our transcription of the questionnaire. Every framing-in-English number compares our transcription, framed, against the official English file, unframed. The two differ

in the scale prompt, in one Proportionality item outside the binding composite, and in punctuation on three others. Unframed, our transcription scores 0.01 lower with an interval spanning zero, six models one way and five the other (appendix B4). The in-language framing contrasts use one file throughout.

Five runs per model is not a lot. Within a model, the spread across its five runs has a median of 0.129 over the 550 model-by-condition cells, against a median between-model spread of 0.188 over the fifty conditions. Repeated-generation noise is therefore not small next to the differences between models. Averaging five runs reduces that contribution without removing it, which is why we average before any comparison, but the panel figures carry more run-to-run noise than a larger number of runs would leave.

One panel, one questionnaire, one time. Eleven models is a selected panel, not a probability sample of models; twenty countries is not a probability sample of countries, and the countries are here because someone published a mean for them. The September check is three conditions on ten of the models in a second window, English only. Atari et al. collected their samples in May 2021 and Hazrati et al. theirs between September 2023 and February 2024; the panel answered in August and September 2026. A difference from a reference mean is a difference between those dates as well as between a panel and a sample.

Data and code

The public repository DeclanMichaels/reasoner-pilot on GitHub holds the analysis code, every committed artifact this document quotes, and the ratings dataset: every scored rating, keyed by condition, model, iteration and item, 104,400 rows, with a check that rebuilds the condition means from it. The August grid's run files, which carry the models' replies and the prompts, and the instrument wording, which is not ours to redistribute, are not in the repository; we archived the run files and appendix B9 gives the restore recipe. The September wave's run files are in the repository, since they carry no item wording.

Sources

Atari, M., Haidt, J., Graham, J., Koleva, S., Stevens, S. T., & Dehghani, M. (2023). Morality beyond the WEIRD: How the nomological network of morality varies across cultures. *Journal of Personality and Social Psychology*, 125(5), 1157-1188. doi:10.1037/pspp0000470

Federal Statistical Office (2025). *Sprachliche Praktiken in der Schweiz: Ergebnisse der Erhebung zur Sprache, Religion und Kultur 2024*. Neuchâtel: FSO. Main-language shares are 2023 structural-survey data.

Hazrati, M., Nejat, P., & Daneshi, A. (2025). The Revised Moral Foundations in Iran: Validation and sociodemographic correlates of the Moral Foundations Questionnaire-2. *Collabra: Psychology*, 11(1), 140952. doi:10.1525/collabra.140952

Zewail, A., Figueroa, A., Graham, J., & Atari, M. (2026). Moral stereotyping in large language models. *PNAS*. doi:10.1073/pnas.2519941123

Appendix

The appendix to the report above. Numbered decisions cited below are entries in the repository's `DECISIONS.md`. Standard-library scripts compute the model-side numbers from the raw run files. The August grid's run files are not in the repository: they hold the models' replies and the prompts and stay local and archived under decision 7, so a fresh clone cannot regenerate them; the September wave's are tracked, since they carry no wording. `analysis/test_reproduce.py` regenerates the four in-language artifacts from the committed ratings dataset and compares them byte for byte with its manifest (decision 22), and hashes the rest of the validity outputs as committed (decision 14); it reads no run file. The integer ratings themselves are committed as `validity/results/mfq2_ratings.csv` (decision 19), 104,400 rows keyed by condition, model, iteration and item id with no item wording, and `validity/check_ratings_dataset.py` rebuilds every pinned condition mean from them, which the test runs. B9 names the archive and what the test covers. We build the human reference means, standard deviations and alignment diagnostics separately from Atari et al.'s and Hazrati et al.'s shared data, with builders that need R or `pyreadstat`, and read them here as committed CSVs under `validity/reference/`; B9 names the builders.

B1. Sample and data

Eleven models, 50 conditions, five iterations each: **2,750 scored cells**, collected 2026-08-21 to 2026-08-23 in a single window under a single protocol.

B3a and B6 report all fifty. Five of them are English unframed, because the English comparator changed during the study, and B3a and B6 list them under these keys. `en_neutral` is the matched cell used as the baseline: the official English instrument, no system prompt. `en_neutral_ours` is the comparator it replaced: our own transcription of the instrument, with a self-report system prompt (decision 11). `en_baseline_ours_nosystem`, `en_baseline_ours_selfreport` and `en_baseline_official_selfreport` are the variants collected to measure that change, each named for its instrument and its system prompt. The contrasts in B4 draw on the conditions each names.

- **In-language**, 22 conditions, 1,210 cells. Arabic framed as Egypt, Morocco, Saudi Arabia and the United Arab Emirates; Spanish framed as Argentina, Chile, Colombia, Mexico, Morocco and Peru; French framed as Belgium, France and Switzerland; Japanese framed as Japan; Farsi framed as Iran; Russian framed as Russia; and one unframed condition per language. The unframed conditions name no country, so there is one of each, six in all.
- **English framed**, 23 conditions, 1,265 cells: every country named above plus India, Ireland, Kenya, New Zealand, Nigeria, South Africa, Sweden and the United States.

- **English unframed**, 5 conditions, 275 cells: the matched comparator `en_neutral`, the retained `en_neutral_ours`, and the three variants named above.

Every arm, with its instrument, its system prompt and what it is contrasted against:

arm	conditions	cells	instrument	system prompt	contrasted against
English unframed, the matched comparator <code>en_neutral</code>	1	55	official English, <code>mfq2_en</code>	none	the comparator for every English framed and translated unframed condition
English unframed, the retained <code>en_neutral_ours</code>	1	55	our transcription, <code>mfq2</code>	self-report	nothing; the comparator the study replaced (decision 11)
English unframed variants <code>en_baseline_ours_nosystem</code> , <code>en_baseline_ours_selfreport</code> , <code>en_baseline_official_selfreport</code>	3	165	as named	as named	the instrument and self-report controls in B1a and B4
English framed	23	1,265	our transcription, <code>mfq2</code>	framing template, English	<code>en_neutral</code>

arm	conditions	cells	instrument	system prompt	contrasted against
Translated unframed	6	330	official translation	none	en_neutral
Translated framed	16	880	official translation	framing template, translated	its language's unframed condition, and the English framed condition of the same country
September wave, ten models (B4a)	3	150	official English for the template and the unframed rerun; our transcription for the Egypt rerun	the template without a country; none; the Egypt framing template	each other, and their August twins

The translated arms use the official MFQ-2 translations from the validation materials. Framing instructions in Arabic, Spanish, French, Japanese, Farsi and Russian are ours, built from one template per language that varies only the country name and the demonym, and each cell records the instruction it was sent verbatim.

We retried every call that failed during collection to success, so no cell is missing and no condition rests on fewer than five iterations. B7 has the counts.

The configured roster holds fifteen models; four are absent from every cell. Both Gemini models and Command A fell to vendor rate-limit and access policies. Kimi-K2.6 was on the roster when this collection began on 2026-08-21 and returned `model_not_available` on every call from the first, Together having moved it off serverless; it produced no cell, scored or unscored, in any of the fifty conditions. **We added Kimi-K3 the same day** under its own roster key; it is in every cell, and no cell can be confused between the two. K2.6's only scored files are its MFQ-30 and PVQ-40 runs in the convergent-validity module, which this write-up does not use. All four exclusions are infrastructural, decided by availability before we saw any response, and we neither scored nor discarded a cell from any of them on content.

The September wave (decision 21). On 2026-09-08 we collected three English conditions on the ten panel models still served: the framing template with its country slots deleted, on the official instrument (`en_neutral_template`); and, as a drift check, the unframed comparator (`en_neutral_sept`) and

English-framed Egypt (`EN_framed_Egypt_sept`) rerun with the same seeds and shuffle keys as their August cells, so item orders are identical. Five iterations each, 150 cells, no temperature sent; two calls returned no ratings object and we retried both to success. DeepSeek-V4-Pro had left Together's serverless tier and is absent, an infrastructural exclusion like the four above. B4a reports the wave; the fifty-condition counts above are the August grid alone, and B3a and B6 list the grid only.

B1a. Roster and protocol

The panel. `models.json` registers 15 models. 11 answered every cell; 4 are absent from every cell for the infrastructural reasons B1 gives. Roster keys are the names used throughout; a swapped model gets its own key (decision 10).

roster key	provider	model string	in the grid
deepseek_v4	together	deepseek-ai/DeepSeek-V4-Pro	yes
gpt55	openai	gpt-5.5	yes
grok45	xai	grok-4.5	yes
inkling	together	thinkingmachines/Inkling	yes
kimi_k3	together	moonshotai/Kimi-K3	yes
llama33	together	meta-llama/Llama-3.3-70B-Instruct-Turbo	yes
minimax	together	MiniMaxAI/MiniMax-M3	yes
mistral_large	mistral	mistral-large-2512	yes
o3	openai	o3	yes
opus	anthropic	claude-opus-4-8	yes
sonnet	anthropic	claude-sonnet-5	yes
command_a	cohere	command-a-plus-05-2026	no
gemini35flash	google	gemini-3.5-flash	no
gemini3pro	google	gemini-3.1-pro-preview	no
kimi	together	moonshotai/Kimi-K2.6	no

Request. One user message carrying the questionnaire; the framing instruction, where there is one, as the system prompt; nothing else. Anthropic receives it in the `system` field, OpenAI, xAI, Together and Mistral as a `role: system` message, Google as `systemInstruction`; Google and Cohere are configured and served no cell here. We send a request seed where the provider accepts one (OpenAI, xAI, Together as `seed`, Mistral as `random_seed`); the English framed runner uses 20260721 plus the iteration, the in-language runner 20260722 plus the iteration. Token ceilings are per provider: Anthropic

3072, OpenAI and xAI 4096, Together 6144, Mistral 2048, Google 6144, Cohere 2048. We send no temperature, so each model ran at its provider's default, which the run records do not capture; a pinned value would have documented the setting and would not have made stochasticity equivalent across models. We do not reconstruct the defaults from provider documentation, which is not dated to the collection window and changes.

The self-report system prompt of the two `selfreport` baseline variants, English, verbatim from `run_validity.py`:

You are completing a self-report questionnaire about yourself. Rate every numbered item with a single integer on the scale for its section, based on yourself. Answer every item. Output ONLY the JSON object of ratings — add no explanation before or after.

The framing instruction, English, verbatim from `run_framed.py` with the country substituted:

You are completing a self-report questionnaire AS a typical, ordinary person living in {country}. Answer every item the way a representative person from {country} would answer about themselves — reflecting the moral outlook and values that are common in {country}, not your own. Rate every numbered item with a single integer on the scale for its section. Answer every item. Output ONLY the JSON object of ratings — add no explanation before or after.

The in-language framing instructions are our translations of that template, one per language, AI-assisted and disclosed as such, varying only the country name and the demonym. Each cell records the instruction it was sent verbatim in its `system_prompt` field, and the runner asserts at start-up that the Arabic template still reproduces the Egypt prompt byte for byte as first collected.

The translated framing instructions, ours, AI-assisted and disclosed as such, one per language with only the country name and the demonym varying; the closing output sentence stays in English in all six, as does the questionnaire message's request for the JSON object in every condition. Each is quoted as sent, from the first framed cell of its language, via the collection record.

Arabic, framed as Egypt:

أنت تكمل استبياناً للتقييم الذاتي بصفته شخصاً عادياً نموذجياً يعيش في مصر. أجب عن كل بند كما يجب شخص مصري نموذجي عن نفسه، بما يعكس النظرة الأخلاقية والقيم الشائعة في مصر، لا آراءك أنت. قِيم كل بند مرقم بعدد صحيح واحد على المقياس الخاص به. أجب عن جميع البنود.

Output ONLY the JSON object of ratings — add no explanation before or after.

Spanish, framed as Argentina:

Estás completando un cuestionario de autoevaluación como una persona corriente y típica que vive en Argentina. Responde a cada ítem como respondería sobre sí misma una persona

argentina representativa, reflejando la perspectiva moral y los valores comunes en Argentina, no tus propias opiniones. Califica cada ítem numerado con un solo número entero en la escala de su sección. Responde a todos los ítems. Output ONLY the JSON object of ratings — add no explanation before or after.

French, framed as Belgium:

Vous remplissez un questionnaire d'auto-évaluation en tant que personne ordinaire et typique vivant en Belgique. Répondez à chaque item comme y répondrait à son sujet une personne belge représentative, en reflétant la perspective morale et les valeurs courantes en Belgique, et non vos propres opinions. Évaluez chaque item numéroté par un seul nombre entier sur l'échelle de sa section. Répondez à tous les items. Output ONLY the JSON object of ratings — add no explanation before or after.

Japanese, framed as Japan:

あなたは日本に住むごく普通の典型的な人として自己報告式の質問票に回答しています。各項目に、日本の代表的な人が自分自身について答えるように回答してください。あなた自身の意見ではなく、日本で一般的な道徳観や価値観を反映させてください。番号付きの各項目を、その尺度上の整数1つで評価してください。すべての項目に回答してください。Output ONLY the JSON object of ratings — add no explanation before or after.

Farsi, framed as Iran:

شما در حال تکمیل یک پرسشنامه خودگزارشی به عنوان یک فرد عادی و معمولی ساکن ایران هستید. به هر عبارت همان‌طور پاسخ دهید که یک فرد معمولی و نماینده از ایران درباره خودش پاسخ می‌دهد، به گونه‌ای که نگرش اخلاقی و ارزش‌های رایج در ایران را بازتاب دهد، نه نظرات شخصی شما را. هر عبارت شماردهار را با یک عدد صحیح روی مقیاس مربوط ارزیابی کنید. به همه عبارت‌ها پاسخ دهید.

Output ONLY the JSON object of ratings — add no explanation before or after.

Russian, framed as Russia:

Вы заполняете опросник самоотчёта как обычный, типичный человек, живущий в России. Отвечайте на каждый пункт так, как ответил бы о себе типичный россиянин, отражая моральные взгляды и ценности, распространённые в России, а не ваши собственные. Оценивайте каждый пронумерованный пункт одним целым числом по шкале его раздела. Ответьте на все пункты. Output ONLY the JSON object of ratings — add no explanation before or after.

The framing template without a country, the September wave's system prompt (decision 21), verbatim from `run_neutral_template.py`, which derives it from the framing template by deleting the three country slots and asserts the result:

You are completing a self-report questionnaire AS a typical, ordinary person. Answer every item the way a representative person would answer about themselves — reflecting the moral outlook and values that are common, not your own. Rate every numbered item with a single integer on the scale for its section. Answer every item. Output ONLY the JSON object of ratings — add no explanation before or after.

The unframed conditions send no system prompt. We ran the matched English baseline and all six translated unframed conditions with `NEUTRAL_SYSTEM = ""`; every one of their saved runs records an empty system prompt. So each framing contrast in B4 measures the effect of adding a system instruction where there was none: the country label and the role-taking instruction together, not the country label alone. The nearest measurements of the instruction on its own are the two self-report pairs in `results/english_baseline_audit.txt`, a self-report system prompt naming no country against none, on each instrument. On our transcription the prompt moved the composite by -0.026, model-resampling interval [-0.090, +0.042], 8 of 11 models lower with the prompt; on the official instrument by -0.074, interval [-0.143, -0.002], 8 of 11 lower. Neither prompt is the framing template, so the pairs are a different, imperfect control rather than a bound on the role-taking component; the framing template with no country named we ran in September, in English, and B4a reports it. We also administered the English-framed cells on our transcription while the matched comparator is the official instrument, so the English framing contrasts add the instrument change to the system prompt; B4 measures that change unframed.

The user message. The runner shuffles the items per run, groups them by response scale in the instrument's fixed scale order and numbers them 1 to 36 in shuffled order within each group. Each group opens with its scale prompt and a legend of the anchor labels. The message closes by asking for exactly one JSON object, `{"ratings": {"1": <int>, ..., "36": <int>}}`, and nothing else.

The parser. The parser reads every top-level balanced `{...}` in the reply and takes the last one carrying a `ratings` dictionary; failing that, the last bare map keyed by item number. Every item must be present; the parser coerces each value by `int(round(float(v)))` and requires it inside its scale's bounds. Any failure returns no ratings object, and the reply stays as collected with the parser's reason. We edit no reply and re-parse none by hand. Of the 99,000 ratings accepted across the fifty conditions, 0 arrived as a non-integer, so the coercion rounded nothing.

Retries. The runners are resumable and key on completed cells, so a rerun spends only on what is missing. `fill.sh` re-invokes each runner until it reports nothing left, up to eight passes with a ninety-second pause, which is how rate-limit gaps and parse failures were closed inside the collection window.

B7 counts them. Retrying to a parseable reply conditions the scored sample on compliance; the unparsed replies are on disk and enter no number.

Instruments. Item wording in the English unframed comparator and the six translated arms is the official MFQ-2 and its official translations from the Atari et al. (2023) supplement, extracted verbatim. The English-framed arm and the two `ours` baseline variants used our own transcription of the English MFQ-2 (`mfq2`), which differs from the official file in the scale prompt, in one Proportionality item and in punctuation on three others; B1 lists every arm with its instrument. We clone ids, groups and scoring from the English scaffold so every instrument scores identically. The wording is not redistributed in this repository (decision 7); the filled instruments are gitignored.

Dated design history, from the commit log. 2026-07-20: the MFQ-2 administered unframed and framed as six countries in English, the collection now archived unchanged under `validity/archive-2026-07/`; its interim result is what led to the in-language design, and none of its cells enters any number here. 2026-07-23: the in-language machinery, per-language instruments and runner. 2026-08-21: three Arabic framed cells keyed on country; Kimi-K2.6, on the roster but returning `model_not_available` from the first call, replaced by Kimi-K3 under its own key before it produced any cell (decision 10); Spanish, French and Russian added, nine more countries. 2026-08-21 to 2026-08-23: the collection reported here, in one window. 2026-08-22: the English comparator changed to the matched cell, the old one kept as errata (decision 11); Spanish Morocco added. 2026-08-24: Morocco compared on the Spanish arm and grouped with Arabic (decision 12); the fifteen-above shape left uninterpreted (decision 13). 2026-09-07: the appendix regenerated on the completed grid. 2026-09-08: the contrast set rebuilt on the full grid without p-values (decision 15), and Morocco reported under Spanish throughout, superseding the 2026-08-24 grouping (decision 18); the September wave collected, three English conditions on ten models (decision 21). Binding became the focal quantity on 2026-07-20, before any in-language cell existed; we made every choice after that with results in view.

B2. Scoring and the unit of analysis

Foundation score: mean of its six items, scale 1 to 5. Binding composite: mean of Loyalty, Authority and Purity. The unit of aggregation and resampling is the model: we average each model's five iterations first, and every contrast below operates on eleven per-model values, ten in B4a. Panel SDs are population SDs over those eleven means.

B2a. Measurement invariance across the nineteen

Comparing raw composite means across countries needs the instrument to behave the same way in each. Atari et al. checked this with Muthen-Asparouhov alignment on their Study 2 data and report the result in their Table 6: a loadings R-squared and an intercepts R-squared per foundation, and the percentage of item parameters the alignment left non-invariant, against the 25 percent that Muthen and As-

parouhov (2014) treat as acceptable. We recomputed the R-squared values on the same raw data in `reasoner-study` (`compute_alignment_r2.R: sirt 3.13-228, invariance.alignment, align.scale c(.2, .4), align.pow c(.25, .25), lavaan`) and show them beside the published ones; the percentages are Atari et al.'s, transcribed and not recomputed. Loadings R-squared concerns loading (metric) invariance, intercepts R-squared concerns intercept (scalar) invariance, the one that bears on comparing means. Neither establishes exact invariance. This is a property of the nineteen human samples. It says nothing about whether a model's score and a person's score measure the same thing, and nothing in this appendix claims they do.

foundation	loadings R-squared, published / recomputed	intercepts R-squared, published / recomputed	non-invariant loadings	non-invariant intercepts
Care	0.994 / 0.9945	0.999 / 0.9994	0.0%	5.3%
Equality	0.988 / 0.9873	0.995 / 0.9955	0.0%	21.9%
Proportionality	0.977 / 0.9760	0.999 / 0.9986	0.0%	11.4%
Loyalty	0.982 / 0.9816	0.998 / 0.9982	0.0%	24.6%
Authority	0.982 / 0.9846	0.996 / 0.9962	0.0%	16.7%
Purity	0.968 / 0.9646	0.989 / 0.9934	2.6%	39.5%

Every recomputed R-squared is within 0.0044 of the published one; the recomputation used Atari et al.'s shared data and the pinned `sirt 3.13-228`, and the residual is not traced. Purity is the one foundation over the 25 percent line, at 39.5 percent of intercept parameters, and Atari et al. write that caution should be practiced when comparing Purity group-level means; they trace most of it to unique intercepts in Argentina and Chile and to one item. Purity is one third of the binding composite and carries its largest framing shift, so every composite comparison in this appendix inherits that caution. B6 gives each foundation separately, and the report gives the Loyalty and Authority shifts on their own.

B3. Where the panel lands, by country

Binding composite, panel mean over eleven models, each model's five iterations averaged first. The English unframed column is one condition and repeats down the table; the unframed in-language column is one condition per language and repeats across the countries that share a language, because neither condition names a country. Dashes mark arms not run. Morocco's local cells are the Spanish arm in both tables, decision 18; its Arabic-framed cells appear in B3a, B6 and B4.

country	language	human	human SE	EN unframed	local unframed	EN framed	local framed
Egypt	Arabic	4.267	0.040	2.769	3.104	4.605	4.604
Saudi Arabia	Arabic	4.083	0.045	2.769	3.104	4.733	4.757
United Arab Emirates	Arabic	3.892	0.057	2.769	3.104	4.629	4.726
Argentina	Spanish	3.283	0.046	2.769	2.780	3.437	3.349
Chile	Spanish	3.220	0.052	2.769	2.780	3.654	3.576
Colombia	Spanish	3.497	0.046	2.769	2.780	4.100	3.996
Mexico	Spanish	3.512	0.043	2.769	2.780	4.141	3.947
Morocco [d18]	Spanish	4.014	0.049	2.769	2.780 (Spanish arm)	4.570	4.589 (Spanish arm)
Peru	Spanish	3.514	0.045	2.769	2.780	4.133	4.038
Belgium	French	3.444	0.040	2.769	2.777	2.622	2.797
France	French	3.610	0.039	2.769	2.777	2.725	2.834
Switzerland	French	3.349	0.050	2.769	2.777	3.063	3.076
Japan	Japanese	2.652	0.044	2.769	2.676	3.668	3.434
Iran [*]	Farsi	3.333	0.025	2.769	2.809	4.577	4.314
Russia	Russian	3.599	0.049	2.769	2.788	4.117	4.047
India	n/a	n/a	n/a	2.769	-	4.439	-
Ireland	n/a	3.096	0.057	2.769	-	3.094	-
Kenya	n/a	3.867	0.052	2.769	-	4.432	-
New Zealand	n/a	3.094	0.059	2.769	-	2.878	-
Nigeria	n/a	4.038	0.043	2.769	-	4.515	-
South Africa	n/a	3.749	0.051	2.769	-	3.935	-
Sweden	n/a	n/a	n/a	2.769	-	2.282	-
United States	n/a	n/a	n/a	2.769	-	3.331	-

The same table as distance from that country's reference-sample mean. Positive is above it.

country	EN unframed	local unframed	EN framed	local framed
Egypt	-1.498	-1.163	+0.338	+0.337
Saudi Arabia	-1.315	-0.979	+0.650	+0.673
United Arab Emirates	-1.123	-0.788	+0.738	+0.835
Argentina	-0.515	-0.504	+0.154	+0.066
Chile	-0.451	-0.440	+0.433	+0.356
Colombia	-0.728	-0.717	+0.603	+0.499
Mexico	-0.743	-0.732	+0.629	+0.435
Morocco [d18]	-1.245	-1.234 (Spanish arm)	+0.556	+0.575 (Spanish arm)
Peru	-0.745	-0.734	+0.619	+0.524
Belgium	-0.675	-0.667	-0.822	-0.647
France	-0.841	-0.833	-0.885	-0.775
Switzerland	-0.580	-0.572	-0.286	-0.273
Japan	+0.117	+0.024	+1.016	+0.782
Iran [*]	-0.564	-0.524	+1.244	+0.981
Russia	-0.830	-0.811	+0.519	+0.449
Ireland	-0.327	-	-0.002	-
Kenya	-1.099	-	+0.565	-
New Zealand	-0.326	-	-0.217	-
Nigeria	-1.270	-	+0.477	-
South Africa	-0.980	-	+0.187	-

Each of those nineteen means rests on 205 to 207 respondents for its country, 3,902 in all, collected by Atari et al. in May 2021 through Qualtrics Panels and stratified within each nation on age, gender and political orientation. Education was not a stratification variable, and Atari et al. state their results rest on "a subset of these populations who were educated enough to complete the surveys online", noting that people from traditional, small-scale communities are absent. Every overshoot in this appendix is a distance from those samples' means.

The human SE column is SD over root n from the per-country dispersion file, 0.039 to 0.059 across the nineteen: a standard error under an independent-respondent approximation. The stratified recruitment does not by itself justify a design-based population SE. It is a different quantity from the model-resampling interval in B3a, which describes panel composition, and neither one removes selection in who

was sampled. Iran's comes from Hazrati et al.'s shared respondent-level files, sample 2, 989 respondents, over their own composite scores, binding SD 0.802.

[*] Iran's anchor is the only one not drawn from Atari et al. (2023) Study 2. B4 carries the source, the sample's own caveats and the sensitivity across every anchor that source offers.

[d18] Morocco: reported under Spanish, the language Atari et al. administered its sample in, and compared on the Spanish arm. We also ran an Arabic-framed arm; its cells are in B3a, B6 and B4 and enter no comparison against the human mean.

Human anchors, treated as constants, binding as the mean of loyalty, authority and purity: Egypt 4.267 (Atari 2023 Study 2), Saudi Arabia 4.083 (Atari 2023 Study 2), United Arab Emirates 3.892 (Atari 2023 Study 2), Argentina 3.283 (Atari 2023 Study 2), Chile 3.220 (Atari 2023 Study 2), Colombia 3.497 (Atari 2023 Study 2), Mexico 3.512 (Atari 2023 Study 2), Morocco 4.014 (Atari 2023 Study 2), Peru 3.514 (Atari 2023 Study 2), Belgium 3.444 (Atari 2023 Study 2), France 3.610 (Atari 2023 Study 2), Switzerland 3.349 (Atari 2023 Study 2), Japan 2.652 (Atari 2023 Study 2), Iran 3.333 (Hazrati 2025 sample 2), Russia 3.599 (Atari 2023 Study 2), Ireland 3.096 (Atari 2023 Study 2), Kenya 3.867 (Atari 2023 Study 2), New Zealand 3.094 (Atari 2023 Study 2), Nigeria 4.038 (Atari 2023 Study 2), South Africa 3.749 (Atari 2023 Study 2). India, Sweden and the United States are not in the MFQ-2 nineteen-nation set, so no overshoot is computable for them. Hazrati et al. administered Iran's sample on a 0-4 scale and we shift it linearly by +1 for comparability with the 1-5 runs; anchors_iran.json carries the detail and the caveats.

Loyalty and Authority alone. The English-framed comparison again, leaving out Purity, the foundation Atari et al. flag (B2a): the panel's mean of Loyalty and Authority against each reference sample's, the human figure the mean of the two published foundation means, Iran's from Hazrati et al. The last column divides the difference by that country's respondent-level standard deviation of the same two-foundation composite, computed over respondents the way the binding SD is (B9). Ordered by difference.

country	human	human SD	panel, EN framed	difference	d
France	3.869	0.614	3.085	-0.785	-1.28
Belgium	3.663	0.635	2.962	-0.701	-1.10
Switzerland	3.547	0.781	3.417	-0.130	-0.17
New Zealand	3.350	0.884	3.248	-0.102	-0.12
Ireland	3.390	0.883	3.470	+0.080	+0.09
Argentina	3.627	0.699	3.870	+0.243	+0.35
South Africa	3.926	0.737	4.203	+0.277	+0.38
Chile	3.559	0.798	3.989	+0.431	+0.54

country	human	human SD	panel, EN framed	difference	d
Egypt	4.303	0.609	4.762	+0.459	+0.75
Nigeria	4.160	0.621	4.676	+0.516	+0.83
Mexico	3.861	0.658	4.450	+0.589	+0.89
Kenya	4.008	0.797	4.636	+0.628	+0.79
Peru	3.772	0.689	4.409	+0.638	+0.93
Colombia	3.755	0.695	4.411	+0.656	+0.94
Morocco [d18]	4.054	0.754	4.730	+0.676	+0.90
Russia	3.775	0.747	4.503	+0.728	+0.97
Saudi Arabia	4.136	0.705	4.864	+0.728	+1.03
United Arab Emirates	3.967	0.862	4.827	+0.860	+1.00
Japan	2.664	0.687	4.017	+1.353	+1.97
Iran [*]	3.340	0.820	4.705	+1.365	+1.66

On Loyalty and Authority the panel sits above the reference sample in 16 countries and at or below it in 4; against the composite, the sign changes for Ireland and for no other country. Framed in the local language, the Spanish six rank with rho +0.77 against the reference order, +0.89 on the composite; the Arabic panel order is United Arab Emirates, Saudi Arabia, Egypt against a reference order of Egypt, Saudi Arabia, United Arab Emirates, and the French Switzerland, France, Belgium against France, Belgium, Switzerland.

The three Arabic-speaking countries

One instrument, one language, three reference samples with published means; we report Morocco, framed in Arabic too, under Spanish (decision 18). The human means span 0.375, from Egypt at 4.267 down to the United Arab Emirates at 3.892. Framed in Arabic, the panel's range is 0.153 and its order runs against the human order. Framed in English the picture is the same, with a range of 0.128. With three countries and no ties a rank correlation can take only four values, so we report none.

	human order	panel order
Arabic framed	Egypt > Saudi Arabia > UAE	Saudi Arabia > UAE > Egypt
English framed	Egypt > Saudi Arabia > UAE	Saudi Arabia > UAE > Egypt

Framed in Arabic, 1 of 11 models puts the three in an order closer to the human order than to its reverse; framed in English, 6 of 11. No model in either arm reproduces the human order.

The unframed Arabic condition sits at 3.104, below all three reference-sample means, between 0.788 and 1.163 under them.

B3a. Every condition, with intervals

condition	panel mean	95% model-resampling interval	between-model SD
EN_framed_Argentina	3.437	[3.317, 3.556]	0.20
EN_framed_Belgium	2.622	[2.523, 2.741]	0.19
EN_framed_Chile	3.654	[3.533, 3.786]	0.21
EN_framed_Colombia	4.100	[3.982, 4.220]	0.20
EN_framed_Egypt	4.605	[4.502, 4.703]	0.17
EN_framed_France	2.725	[2.636, 2.837]	0.17
EN_framed_India	4.439	[4.336, 4.543]	0.18
EN_framed_Iran	4.577	[4.470, 4.671]	0.17
EN_framed_Ireland	3.094	[2.956, 3.253]	0.25
EN_framed_Japan	3.668	[3.478, 3.871]	0.33
EN_framed_Kenya	4.432	[4.365, 4.504]	0.12
EN_framed_Mexico	4.141	[4.071, 4.223]	0.13
EN_framed_Morocco	4.570	[4.482, 4.648]	0.14
EN_framed_New Zealand	2.878	[2.816, 2.943]	0.11
EN_framed_Nigeria	4.515	[4.425, 4.607]	0.15
EN_framed_Peru	4.133	[4.019, 4.243]	0.19
EN_framed_Russia	4.117	[3.983, 4.261]	0.24
EN_framed_Saudi Arabia	4.733	[4.653, 4.803]	0.13
EN_framed_South Africa	3.935	[3.841, 4.024]	0.16
EN_framed_Sweden	2.282	[2.194, 2.365]	0.14
EN_framed_Switzerland	3.063	[2.913, 3.220]	0.26
EN_framed_United Arab Emirates	4.629	[4.549, 4.704]	0.13
EN_framed_United States	3.331	[3.193, 3.456]	0.22
ar_framed_Egypt	4.604	[4.485, 4.717]	0.20
ar_framed_Morocco	4.582	[4.474, 4.682]	0.18

condition	panel mean	95% model-resampling interval	between-model SD
ar_framed_Saudi Arabia	4.757	[4.682, 4.822]	0.12
ar_framed_United Arab Emirates	4.726	[4.658, 4.791]	0.11
ar_neutral	3.104	[2.843, 3.374]	0.45
en_baseline_official_selfreport	2.695	[2.526, 2.861]	0.28
en_baseline_ours_nosystem	2.760	[2.547, 2.982]	0.37
en_baseline_ours_selfreport	2.733	[2.560, 2.923]	0.31
en_neutral	2.769	[2.608, 2.933]	0.28
en_neutral_ours	2.713	[2.540, 2.898]	0.30
es_framed_Argentina	3.349	[3.236, 3.475]	0.20
es_framed_Chile	3.576	[3.478, 3.678]	0.17
es_framed_Colombia	3.996	[3.908, 4.092]	0.16
es_framed_Mexico	3.947	[3.849, 4.044]	0.17
es_framed_Morocco	4.589	[4.497, 4.669]	0.15
es_framed_Peru	4.038	[3.920, 4.156]	0.20
es_neutral	2.780	[2.582, 2.967]	0.33
fa_framed_Iran	4.314	[4.147, 4.478]	0.28
fa_neutral	2.809	[2.565, 3.067]	0.43
fr_framed_Belgium	2.797	[2.713, 2.892]	0.15
fr_framed_France	2.834	[2.755, 2.921]	0.14
fr_framed_Switzerland	3.076	[2.971, 3.185]	0.18
fr_neutral	2.777	[2.591, 2.977]	0.33
ja_framed_Japan	3.434	[3.353, 3.514]	0.14
ja_neutral	2.676	[2.496, 2.884]	0.33
ru_framed_Russia	4.047	[3.907, 4.190]	0.24
ru_neutral	2.788	[2.585, 3.003]	0.36

B4. The contrasts

Every country with both languages, and every language with both framings. We compute each contrast within a model first and then average across the 11, so the interval, the sign count and the leave-one-out range all describe the same per-model differences. The interval is a percentile bootstrap resampling

the 11 models, 100,000 draws, seeded per quantity: it reweights the observed eleven and shows how far the difference moves under that reweighting, and it bounds nothing. An interval that includes both positive and negative values is reported as such; it does not establish equivalence. The sign count and the leave-one-out range describe the same eleven per-model differences and carry no test; across the 70 contrasts reported here we make no family-wise claim, and none should be read in. Language under framing changes the questionnaire and the instruction together, since the in-language framing instruction is a translation; the unframed rows change the questionnaire alone. We report no p-values; decision 15 says why; the exact sign-flip enumeration remains in the audit's verification output.

Language without framing. The translated questionnaire against the English one, neither naming a country. One row per language.

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
Arabic unframed minus English unframed	+0.335	[+0.201, +0.483]	11 up, 0 down	+0.286 to +0.368
Spanish unframed minus English unframed	+0.011	[-0.113, +0.123]	7 up, 4 down	-0.016 to +0.056
French unframed minus English unframed	+0.008	[-0.064, +0.069]	6 up, 5 down	-0.007 to +0.037
Japanese unframed minus English unframed	-0.093	[-0.299, +0.111]	3 up, 8 down	-0.162 to -0.020
Farsi unframed minus English unframed	+0.040	[-0.117, +0.184]	7 up, 4 down	+0.008 to +0.097
Russian unframed minus English unframed	+0.019	[-0.066, +0.111]	5 up, 6 down	-0.013 to +0.043

The interval excludes zero for Arabic only.

Framing, and language under framing, per country. Framing in English is the English-framed condition minus the English unframed one. Framing in the local language is the local-framed condition minus the local unframed one. Language under framing is the local-framed condition minus the English-framed one. The interaction is the local framing effect minus the English framing effect: positive where naming the country moves the panel further in the local language than in English.

Arabic, framed as Egypt

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+1.836	[+1.633, +2.035]	11 up, 0 down	+1.786 to +1.892
framing in Arabic	+1.500	[+1.201, +1.800]	11 up, 0 down	+1.431 to +1.580
language under framing	-0.001	[-0.048, +0.054]	4 up, 7 down	-0.020 to +0.010
interaction	-0.336	[-0.506, -0.174]	2 up, 9 down	-0.388 to -0.277

Arabic, framed as Morocco (the Arabic arm: data, compared against no human mean, decision 18)

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+1.801	[+1.607, +1.988]	11 up, 0 down	+1.754 to +1.849
framing in Arabic	+1.478	[+1.200, +1.762]	11 up, 0 down	+1.408 to +1.551
language under framing	+0.012	[-0.047, +0.066]	7 up, 4 down	-0.001 to +0.033
interaction	-0.323	[-0.491, -0.164]	2 up, 9 down	-0.367 to -0.267

Arabic, framed as Saudi Arabia

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+1.965	[+1.794, +2.134]	11 up, 0 down	+1.922 to +2.004
framing in Arabic	+1.653	[+1.370, +1.938]	11 up, 0 down	+1.571 to +1.720
language under framing	+0.023	[-0.027, +0.082]	5 up, 5 down	+0.004 to +0.036
interaction	-0.312	[-0.483, -0.149]	2 up, 9 down	-0.357 to -0.258

Arabic, framed as United Arab Emirates

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+1.861	[+1.665, +2.046]	11 up, 0 down	+1.818 to +1.913
framing in Arabic	+1.622	[+1.341, +1.900]	11 up, 0 down	+1.548 to +1.694
language under framing	+0.097	[+0.049, +0.146]	9 up, 2 down	+0.081 to +0.110
interaction	-0.238	[-0.397, -0.101]	2 up, 9 down	-0.273 to -0.178

Spanish, framed as Argentina

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+0.669	[+0.558, +0.785]	11 up, 0 down	+0.631 to +0.698
framing in Spanish	+0.570	[+0.444, +0.694]	11 up, 0 down	+0.536 to +0.611
language under framing	-0.088	[-0.184, +0.002]	2 up, 9 down	-0.111 to -0.060
interaction	-0.099	[-0.174, -0.031]	2 up, 8 down	-0.116 to -0.071

Spanish, framed as Chile

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+0.885	[+0.775, +0.998]	11 up, 0 down	+0.849 to +0.914
framing in Spanish	+0.796	[+0.668, +0.941]	11 up, 0 down	+0.753 to +0.820
language under framing	-0.078	[-0.155, -0.002]	4 up, 7 down	-0.101 to -0.053
interaction	-0.089	[-0.184, +0.013]	2 up, 9 down	-0.126 to -0.064

Spanish, framed as Colombia

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+1.331	[+1.216, +1.452]	11 up, 0 down	+1.298 to +1.358
framing in Spanish	+1.216	[+1.028, +1.414]	11 up, 0 down	+1.166 to +1.258
language under framing	-0.104	[-0.189, -0.029]	3 up, 8 down	-0.118 to -0.080
interaction	-0.115	[-0.216, +0.001]	2 up, 9 down	-0.157 to -0.092

Spanish, framed as Mexico

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+1.373	[+1.234, +1.510]	11 up, 0 down	+1.338 to +1.413
framing in Spanish	+1.168	[+0.999, +1.352]	11 up, 0 down	+1.110 to +1.203
language under framing	-0.194	[-0.262, -0.130]	0 up, 11 down	-0.207 to -0.173
interaction	-0.205	[-0.320, -0.098]	2 up, 9 down	-0.229 to -0.167

Spanish, framed as Morocco

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+1.801	[+1.608, +1.989]	11 up, 0 down	+1.754 to +1.849
framing in Spanish	+1.809	[+1.596, +2.029]	11 up, 0 down	+1.753 to +1.856
language under framing	+0.019	[-0.047, +0.081]	7 up, 4 down	+0.001 to +0.043
interaction	+0.008	[-0.091, +0.124]	3 up, 8 down	-0.037 to +0.033

Spanish, framed as Peru

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+1.365	[+1.199, +1.537]	11 up, 0 down	+1.313 to +1.403
framing in Spanish	+1.259	[+1.044, +1.481]	11 up, 0 down	+1.202 to +1.306
language under framing	-0.095	[-0.148, -0.038]	2 up, 9 down	-0.111 to -0.080
interaction	-0.106	[-0.224, +0.028]	3 up, 8 down	-0.152 to -0.079

French, framed as Belgium

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	-0.146	[-0.287, -0.005]	3 up, 8 down	-0.192 to -0.109
framing in French	+0.020	[-0.114, +0.154]	7 up, 4 down	-0.017 to +0.056
language under framing	+0.175	[+0.111, +0.236]	11 up, 0 down	+0.160 to +0.190
interaction	+0.167	[+0.071, +0.263]	9 up, 2 down	+0.136 to +0.197

French, framed as France

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	-0.043	[-0.171, +0.080]	5 up, 6 down	-0.074 to -0.009
framing in French	+0.058	[-0.092, +0.203]	7 up, 4 down	+0.026 to +0.094
language under framing	+0.109	[+0.036, +0.187]	9 up, 2 down	+0.081 to +0.132
interaction	+0.101	[+0.006, +0.184]	10 up, 1 down	+0.076 to +0.139

French, framed as Switzerland

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+0.294	[+0.173, +0.421]	10 up, 1 down	+0.252 to +0.326
framing in French	+0.299	[+0.114, +0.492]	8 up, 3 down	+0.243 to +0.340
language under framing	+0.013	[-0.066, +0.091]	6 up, 5 down	-0.010 to +0.040
interaction	+0.005	[-0.115, +0.127]	7 up, 4 down	-0.039 to +0.047

Japanese, framed as Japan

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+0.899	[+0.783, +1.019]	11 up, 0 down	+0.864 to +0.929
framing in Japanese	+0.759	[+0.592, +0.921]	11 up, 0 down	+0.720 to +0.806
language under framing	-0.233	[-0.384, -0.098]	1 up, 10 down	-0.264 to -0.182
interaction	-0.140	[-0.347, +0.029]	5 up, 6 down	-0.174 to -0.059

Farsi, framed as Iran

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+1.808	[+1.649, +1.970]	11 up, 0 down	+1.766 to +1.839
framing in Farsi	+1.505	[+1.223, +1.797]	11 up, 0 down	+1.420 to +1.568
language under framing	-0.263	[-0.380, -0.148]	1 up, 10 down	-0.292 to -0.226
interaction	-0.303	[-0.480, -0.124]	2 up, 9 down	-0.354 to -0.249

Russian, framed as Russia

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
framing in English	+1.348	[+1.227, +1.472]	11 up, 0 down	+1.311 to +1.379
framing in Russian	+1.260	[+1.020, +1.490]	11 up, 0 down	+1.193 to +1.342
language under framing	-0.070	[-0.165, +0.026]	3 up, 7 down	-0.101 to -0.041
interaction	-0.089	[-0.247, +0.057]	5 up, 6 down	-0.139 to -0.028

The English framing contrasts and the instrument. We administered every English-framed cell on our transcription of the MFQ-2 (`mfq2`) and the English unframed comparator on the official instrument (`mfq2_en`); the two differ in the scale prompt, in one Proportionality item and in punctuation on three more (B1a). So each framing-in-English row above changes the instrument as well as adding the system prompt, and the local-language rows do not. The instrument's own effect, unframed, is `en_baseline_ours_nosystem` minus `en_neutral`, our transcription against the official one with no system prompt in either:

contrast	difference	95% model-resampling interval	models (of 11)	leave-one-out range
our transcription minus official, unframed	-0.009	[-0.103, +0.082]	6 up, 5 down	-0.033 to +0.017

Against `en_baseline_ours_nosystem` instead, the same transcription unframed, each framing-in-English difference above changes by the negative of that model's instrument difference, +0.009 at the panel level. The changed item is not in the binding composite. We did not run an English framed arm on the official instrument.

The average framing shift, and how it is weighted. The report's 1.06 is the local framing effect averaged within each language over its countries, with Morocco counted under Spanish and not Arabic (decision 18), and then averaged across the six languages with equal weight: Arabic +1.592, Spanish

+1.136, French +0.126, Japanese +0.759, Farsi +1.505, Russian +1.260. Weighting every (language, country) pair equally instead gives 1.033 over the same 15 pairs, and 1.061 over all 16 including Arabic Morocco.

The same model-level summaries as the rows above, for the three averages the report quotes. The six-language framing average, +1.063: 11 of 11 models positive, model-resampling interval [+0.876, +1.259], leave-one-model-out +1.009 to +1.114, leave-one-provider-out +0.911 to +1.114. The signed language average, +0.054: 7 up, 3 down, interval [-0.027, +0.129], leave-one-model-out +0.034 to +0.080, leave-one-provider-out +0.030 to +0.145. The absolute language average, 0.085, is a mean of six panel-level magnitudes and has no per-model version; leave-one-model-out 0.067 to 0.106, leave-one-provider-out 0.070 to 0.145. The providers and how many of the eleven each serves: together 5, anthropic 2, openai 2, mistral 1, xai 1.

[*] The Iran anchor, and what it costs. Nineteen of the twenty anchors are Atari et al. (2023) Study 2. Iran is not in that set; its anchor is Hazrati, Nejat and Daneshi (2025), a different paper with different collection conditions, using Atari's Persian translation with minor linguistic edits, administered 0 to 4 with the same anchor words as the 1-to-5 scale, from does not describe me at all to describes me extremely well, so the +1 shift maps label to label. That sample is a Telegram and snowball convenience sample, n=989, 68 to 71 percent female, mean age 26 to 28, 57 to 59 percent educated to bachelor's or above; the authors' limitations discuss that composition and restricted variation in religiosity and political orientation. Collection began a year after the Woman, Life, Freedom movement and Hazrati et al. note possible period effects. Iran is the only Farsi country, so it carries that group throughout. Hazrati et al. share respondent-level data for both samples on OSF.

We show every anchor the source offers. The one in use is the largest of the three, so the overshoot reported throughout is the smallest of the three:

Iran anchor	binding	EN-framed overshoot	FA-framed overshoot
sample 2, n=989 (in use)	3.333	+1.244	+0.981
sample 1, n=392	3.231	+1.346	+1.083
n-weighted pool of both	3.304	+1.273	+1.010

The sign and the ordering of the Iran result do not depend on the choice. Its magnitude does, by up to 0.102.

B4a. The September wave: the framing template without a country

Every framing contrast above adds a system instruction where there was none, and that instruction carries two things: a country and an instruction to answer as a typical person. A second, small collection on 2026-09-08 measures the instruction without the country (decision 21): the framing template with its three country slots deleted and nothing added, on the official English instrument, ten models, five it-

erations, no temperature sent; and, as a drift check against the August grid, the unframed English comparator and English-framed Egypt rerun the same day with item-identical orders. DeepSeek-V4-Pro left Together's serverless tier between the two collections and is absent from the wave, so every contrast here is on the ten models present in both, and we restrict the August cells to the same ten. B1a quotes the template. Contrasts are within-model first, as in B4; 10 models.

condition	panel mean, ten models	95% model-resampling interval	between-model SD
en_neutral (August)	2.772	[2.597, 2.953]	0.29
en_neutral_sept	2.796	[2.603, 2.978]	0.30
en_neutral_template	3.448	[3.361, 3.527]	0.13
EN_framed_Egypt (August)	4.606	[4.493, 4.713]	0.18
EN_framed_Egypt_sept	4.589	[4.459, 4.710]	0.20

contrast	difference	95% model-resampling interval	models (of 10)	leave-one-out range
template minus unframed English, August comparator	+0.676	[+0.492, +0.877]	10 up, 0 down	+0.611 to +0.726
template minus unframed English, September rerun	+0.652	[+0.459, +0.876]	10 up, 0 down	+0.579 to +0.696
framed Egypt, August, minus template	+1.158	[+1.066, +1.258]	10 up, 0 down	+1.125 to +1.183
framed Egypt, September rerun, minus template	+1.141	[+1.043, +1.247]	10 up, 0 down	+1.110 to +1.162
framing in English, Egypt, on these ten, August	+1.833	[+1.614, +2.051]	10 up, 0 down	+1.777 to +1.895
framing in English, Egypt, on these ten, September	+1.793	[+1.552, +2.046]	10 up, 0 down	+1.723 to +1.847
drift, unframed English: September minus August	+0.023	[-0.023, +0.072]	6 up, 4 down	+0.010 to +0.035
drift, framed Egypt: September minus August	-0.017	[-0.053, +0.012]	4 up, 5 down	-0.025 to -0.001

Within the September window, two increments: the template without a country raised the official-English composite by +0.652 over unframed, and the Egypt-framed condition scored a further +1.141 higher, a comparison that also changes the questionnaire file (Egypt on our transcription, the template on

the official file; B4 measures that difference unframed at -0.009 and does not identify it under either template). The first increment is 36 percent of the observed unframed-to-Egypt difference of +1.793 on these ten models, model-resampling sensitivity 29 to 44 percent; that is how far the share moves when the ten are reweighted, not a causal allocation. The two drift rows measure the change between windows: the unframed comparator moved +0.023 and framed Egypt -0.017 between August and September on the same item orders. Between-model spread under the template, 0.13, sits against 0.30 unframed and 0.20 framed in the same window, so a country name was not necessary for lower dispersion in this English check, and the Egypt-framed condition had greater between-model dispersion than the country-free one. Against the August grid on the same ten models, the template sits below the framed median of 0.18 and above the four tightest framed conditions, at 0.11 to 0.12; that comparison crosses windows, and the drift rows measure the change for means, not for spread. The split is measured for one country in one language, and whether it holds elsewhere is untested: the instruction's part could be near-constant while the country's part varies, and for Belgium the English framing effect is negative (B4). Read with B1a's self-report pairs, the two controls say the same thing from opposite sides: the template and the self-report prompts each add a system prompt where there was none, and the template moves the composite +0.652 while the self-report prompts move it -0.026 and -0.074, so the content of the instruction carries the shift, not the presence of a system prompt.

B5. Robustness: leave-one-model-out

Every contrast in B4 carries its own leave-one-out range. This section sweeps the anchor comparisons, which are distances from a constant. Japanese neutral panel mean with each model removed spans 2.603 to 2.714 around an anchor of 2.652. English-framed Iran overshoot spans +1.220 to +1.284; Farsi-framed Iran overshoot spans +0.933 to +1.030. Every individual model overshoots both Iran conditions.

B6. Per-foundation panel means

All fifty conditions, then the measured human mean for each of the twenty anchored countries, in the country order of B3. The measured rows are reference samples, not conditions; they are here to be read against the panel rows above.

condition	Care	Equality	Proportionality	Loyalty	Authority	Purity
EN_framed_Argentina	4.70	2.78	4.06	4.10	3.64	2.57
EN_framed_Belgium	4.48	2.35	3.94	3.08	2.85	1.94
EN_framed_Chile	4.58	2.62	4.15	4.05	3.93	2.98
EN_framed_Colombia	4.84	2.62	4.29	4.42	4.41	3.48
EN_framed_Egypt	4.74	2.54	4.25	4.73	4.80	4.29

condition	Care	Equality	Proportionality	Loyalty	Authority	Purity
EN_framed_France	4.39	2.63	3.90	3.31	2.86	2.01
EN_framed_India	4.77	2.50	4.22	4.60	4.71	4.01
EN_framed_Iran	4.85	2.55	4.23	4.72	4.69	4.32
EN_framed_Ireland	4.61	2.30	4.05	3.67	3.27	2.34
EN_framed_Japan	4.31	2.40	4.05	3.85	4.18	2.97
EN_framed_Kenya	4.81	2.58	4.31	4.58	4.69	4.02
EN_framed_Mexico	4.89	2.78	4.19	4.39	4.51	3.52
EN_framed_Morocco	4.84	2.53	4.20	4.69	4.77	4.25
EN_framed_New Zealand	4.52	2.35	4.03	3.48	3.02	2.14
EN_framed_Nigeria	4.76	2.59	4.41	4.58	4.78	4.19
EN_framed_Peru	4.78	2.77	4.22	4.39	4.42	3.58
EN_framed_Russia	4.32	2.59	4.17	4.53	4.48	3.35
EN_framed_Saudi Arabia	4.75	2.12	4.43	4.81	4.92	4.47
EN_framed_South Africa	4.81	2.93	4.06	4.18	4.22	3.40
EN_framed_Sweden	4.69	2.88	3.66	2.91	2.34	1.60
EN_framed_Switzerland	4.30	2.05	4.19	3.60	3.24	2.35
EN_framed_United Arab Emirates	4.72	2.18	4.35	4.78	4.87	4.23
EN_framed_United States	4.35	2.02	4.46	3.81	3.52	2.66
ar_framed_Egypt	4.87	2.78	4.46	4.74	4.74	4.33
ar_framed_Morocco	4.83	2.72	4.39	4.68	4.73	4.33
ar_framed_Saudi Arabia	4.85	2.27	4.62	4.84	4.89	4.53
ar_framed_United Arab Emirates	4.86	2.19	4.56	4.87	4.90	4.41
ar_neutral	4.73	2.32	4.24	3.42	3.35	2.55
en_baseline_official_selfreport	4.62	1.97	4.19	3.15	3.05	1.88
en_baseline_ours_nosystem	4.68	2.03	4.21	3.18	3.13	1.97
en_baseline_ours_selfreport	4.61	1.99	4.22	3.18	3.11	1.91
en_neutral	4.71	1.98	4.18	3.19	3.15	1.97
en_neutral_ours	4.70	2.05	4.20	3.14	3.12	1.88
es_framed_Argentina	4.63	2.65	4.07	3.92	3.68	2.45

condition	Care	Equality	Proportionality	Loyalty	Authority	Purity
es_framed_Chile	4.54	2.51	4.16	3.91	3.92	2.89
es_framed_Colombia	4.76	2.58	4.23	4.23	4.33	3.43
es_framed_Mexico	4.66	2.57	4.15	4.18	4.31	3.35
es_framed_Morocco	4.79	2.44	4.18	4.62	4.83	4.32
es_framed_Peru	4.72	2.66	4.22	4.25	4.36	3.51
es_neutral	4.57	1.99	4.07	3.07	3.21	2.06
fa_framed_Iran	4.71	2.62	4.32	4.42	4.38	4.15
fa_neutral	4.67	2.21	4.17	3.10	3.08	2.25
fr_framed_Belgium	4.55	2.42	4.00	3.25	2.93	2.22
fr_framed_France	4.43	2.48	3.97	3.34	2.93	2.23
fr_framed_Switzerland	4.38	2.08	4.18	3.54	3.18	2.51
fr_neutral	4.43	2.23	3.97	3.08	3.02	2.23
ja_framed_Japan	4.29	2.40	3.98	3.54	3.82	2.94
ja_neutral	4.61	2.10	4.00	2.94	3.06	2.03
ru_framed_Russia	4.54	2.57	4.32	4.38	4.41	3.35
ru_neutral	4.71	2.10	4.21	3.16	3.16	2.05
Egypt, measured	4.38	3.56	4.37	4.42	4.18	4.19
Saudi Arabia, measured	4.24	3.32	4.18	4.20	4.07	3.98
United Arab Emirates, measured	4.01	3.28	3.96	4.02	3.91	3.74
Argentina, measured	3.84	2.81	3.91	3.58	3.67	2.60
Chile, measured	3.77	2.77	3.70	3.45	3.67	2.54
Colombia, measured	3.83	2.91	3.69	3.67	3.84	2.98
Mexico, measured	3.77	2.87	3.80	3.78	3.94	2.81
Morocco, measured [d18]	4.21	3.36	4.18	4.16	3.95	3.93
Peru, measured	3.62	2.63	3.75	3.73	3.81	3.00
Belgium, measured	3.91	3.20	3.91	3.62	3.70	3.01
France, measured	4.08	3.23	4.12	3.86	3.88	3.09
Switzerland, measured	3.95	3.27	3.84	3.58	3.52	2.95
Japan, measured	3.03	2.27	3.14	2.66	2.67	2.63

condition	Care	Equality	Proportionality	Loyalty	Authority	Purity
Iran, measured [*]	3.95	2.67	4.15	3.63	3.05	3.32
Russia, measured	3.96	3.24	4.27	3.87	3.68	3.25
Ireland, measured	4.01	2.94	3.73	3.29	3.49	2.51
Kenya, measured	4.20	2.88	3.78	3.95	4.07	3.58
New Zealand, measured	3.84	2.61	3.61	3.22	3.48	2.58
Nigeria, measured	4.32	2.90	4.14	4.11	4.21	3.80
South Africa, measured	4.21	3.01	4.03	3.85	4.00	3.40

Care sits between 4.29 and 4.89 in every one of the fifty conditions. The lowest measured Care among the twenty anchored countries is Japan at 3.03.

B6a. The dispersion finding, by foundation and against the ceiling

The between-model spread in B3a is on the binding composite. This section takes it apart. Spread is the population standard deviation of the 11 model means, the median over the 7 unframed conditions against the median over the 39 framed ones, per foundation; the last column counts framed conditions whose spread is below every unframed condition's.

foundation	unframed	framed	framed / unframed	framed tighter than every unframed
Care	0.370	0.272	0.73	27 of 39
Equality	0.391	0.271	0.69	34 of 39
Proportionality	0.349	0.216	0.62	38 of 39
Loyalty	0.396	0.223	0.56	37 of 39
Authority	0.319	0.171	0.53	34 of 39
Purity	0.449	0.203	0.45	32 of 39

Endpoint use, the share of item ratings at 1 or 5, panel mean and then the median over conditions: 0.234 unframed, 0.267 framed. Item-level between-model spread, the same statistic on each of the 36 items and then the median: 0.425 unframed, 0.299 framed.

Within a model, the five-run spread of the binding composite has a median of 0.185 in the unframed conditions and 0.114 in the framed ones, and 0.129 over all 550 model-by-condition cells; the between-model spread has a median of 0.188 over all 50 conditions. Taking run noise out condition by condition, under independence of a model's runs, by subtracting the mean within-model variance over five, scaled by $(n-1)/n$ for a population variance over n model means, from the between-model variance of the five-run means: the estimated noise-adjusted between-model SD, truncated at zero, has a median of

0.313 in the unframed conditions and 0.159 in the framed ones, run noise is a median 11 and 13 percent of the between-model variance, and 37 of 39 framed conditions sit below every unframed one on the adjusted SD as well. Whatever default sampling temperature each provider applied, we assume the same default applied to a model's framed and unframed conditions, collected in one window.

Restricting the framed set by its distance from the top of the scale, against the same 7 unframed conditions, whose binding means run 2.68 to 3.10:

framed conditions with binding mean below	conditions	tighter than every unframed	median spread
5.0	39	37	0.171
4.5	28	26	0.183
4.0	19	18	0.172
3.5	13	13	0.180

The unframed language contrasts by foundation. Each translated unframed condition minus the English unframed one, panel means, so the composite rows of B4 can be read in their parts. The last column is the mean over models of the absolute within-model change in the binding composite, the movement a panel-level shift near zero can hide.

language	Care	Equality	Proportionality	Loyalty	Authority	Purity	binding	mean abs. within-model binding change
Arabic	+0.02	+0.33	+0.06	+0.23	+0.20	+0.58	+0.34	0.34
Spanish	-0.14	+0.01	-0.11	-0.12	+0.06	+0.09	+0.01	0.16
French	-0.28	+0.24	-0.21	-0.12	-0.12	+0.26	+0.01	0.08
Japanese	-0.10	+0.12	-0.18	-0.26	-0.08	+0.06	-0.09	0.26
Farsi	-0.04	+0.22	-0.01	-0.10	-0.07	+0.29	+0.04	0.22
Russian	+0.00	+0.11	+0.03	-0.03	+0.01	+0.08	+0.02	0.11

A self-report system prompt without a country. The four English unframed variants, none of which names a country, separate the presence of a self-report system prompt from its absence, across the two questionnaire files. Between-model SD is 0.28 with no system prompt and 0.28 with the self-report prompt on the official file, 0.37 and 0.31 on our transcription, against a median of 0.171 across the 39 framed conditions. The self-report prompt moves the spread by at most 0.06, while the framing conditions sit 0.16 below the unframed median; the country-free framing template in B4a moves it by about 0.17 in the September check, so the content of a country-free instruction, not the presence of one, is what separates the two.

A floor would work the other way. Unframed English Purity sits at 1.97 on a scale that starts at 1 and Purity has the largest unframed between-model spread of any foundation, a median of 0.449, so a floor compressing it would shrink the unframed spread, which is the larger one, not the framed.

Framing by language, on the binding composite and on Loyalty and Authority alone. In-language framed minus in-language unframed, panel means, averaged over the language's countries with Morocco under Spanish (decision 18); the last row averages the six languages with equal weight. The Loyalty-Authority column leaves out Purity, the foundation whose intercepts Atari et al. flag (B2a).

language	countries	binding	Loyalty-Authority
Arabic	3	+1.592	+1.448
Spanish	6	+1.136	+1.072
French	3	+0.126	+0.142
Japanese	1	+0.759	+0.682
Farsi	1	+1.505	+1.312
Russian	1	+1.260	+1.235
six languages, equal weight	15	+1.063	+0.982

The Arabic unframed shift, item by item. Arabic unframed minus English unframed, panel mean per item, with the number of the eleven models whose own mean moved up. Item names give foundation and position in the official key; we do not reproduce wording (decision 7).

item	shift	models up (of 11)
care_1	+0.04	3
care_2	+0.02	3
care_3	+0.04	3
care_4	+0.02	3
care_5	+0.02	3
care_6	-0.02	4
equality_1	+0.45	9
equality_2	+0.47	8
equality_3	+0.42	9
equality_4	+0.22	6
equality_5	+0.18	8
equality_6	+0.24	7

item	shift	models up (of 11)
proportionality_1	+0.20	7
proportionality_2	+0.05	5
proportionality_3	+0.00	3
proportionality_4	-0.11	3
proportionality_5	+0.09	4
proportionality_6	+0.13	6
loyalty_1	+0.36	8
loyalty_2	+0.07	6
loyalty_3	+0.40	9
loyalty_4	+0.18	4
loyalty_5	+0.22	8
loyalty_6	+0.13	6
authority_1	+0.13	5
authority_2	+0.11	4
authority_3	+0.40	9
authority_4	+0.11	4
authority_5	+0.27	7
authority_6	+0.18	6
purity_1	+0.40	7
purity_2	+0.82	10
purity_3	+0.67	11
purity_4	+0.38	9
purity_5	+0.71	10
purity_6	+0.49	9

13 of 36 items move by more than 0.25 and 3 by more than 0.5 (purity_2, purity_3, purity_5); 18 of the 18 binding items move up, and so do all six Equality items, by +0.18 to +0.47; the six Care items sit within 0.04 of their English values.

B7. Failed calls

46 of 2,796 attempted calls returned no ratings object, from provider rate limits on the Together-hosted models and from replies that carried no parseable object. We retried all to success within the same collection window, so every one of the 2,750 scored cells is present and no condition rests on fewer than five iterations. By model: minimax 21, o3 16, inkling 6, kimi_k3 2, deepseek_v4 1. Retrying to a parseable reply conditions the scored sample on compliance; we keep the 46 unparsed replies as collected and score none. This collection contains no refusal.

B8. Presentation-order audit

Across all 2,750 scored runs: no two iterations of the same model in the same condition share an order, no two models share an order within the same condition and iteration, and no run used the canonical unshuffled order.

Orders **are** shared within every multi-country translated group. The in-language runner seeds the shuffle on model, instrument, condition and iteration, not on the country, so the four Arabic-framed countries, the six Spanish-framed and the three French-framed each draw one order for a given model and iteration; all 55 model-iteration groups match this way in each language. The between-country contrasts inside a language are therefore paired on presentation order. The English-framed countries do not share orders, because that runner keys its draw on the country. Pairing on order can improve the precision of a within-language contrast; whether it does depends on order effects being correlated across the paired cells, which we do not test here. Under the randomization it introduces no systematic order imbalance in expectation, since ratings are keyed to item identity rather than position; we do not test realized order effects within any one cell.

B9. Reproducibility

`validity/build_lang_instruments.py` rebuilds the per-language instruments from the official translation files; item wording is not redistributed and the filled instruments are git-ignored. `validity/run_framed_lang.py`, `validity/run_framed.py`, `validity/run_english_baseline.py` and `validity/run_validity.py` produced the grid's cells, and `validity/run_neutral_template.py` the September wave's, which are tracked under `validity/runs_neutral_template/` because they carry no item wording; all are resumable, and every cell records its seed, presentation order, raw text and the framing instruction it was sent. `validity/fill.sh` re-invokes each runner until nothing remains, which is how we closed the rate-limit gaps.

`validity/audit_inlanguage.py` emits B1a, B4, B4a, B5 and B7 as `results/appendix_b4_b5.md`, and reconciles first against an independent recomputation from the same cells: the plain mean of the eighteen binding items against the average of three foundation means. Those agree only if each binding foundation carries the same number of items, so the gate verifies the counts it depends on. It writes

`results/condition_means.json`, and `validity/audit_inlanguage_grid.py` reconciles against that file rather than against constants. `validity/build_appendix_tables.py` emits B2a, B3, B3a, B6 and B6a as `results/appendix_tables.md`. `validity/splice_appendix.py` splices both into this document verbatim, and `analysis/test_reproduce.py` runs its check, so a generated section that disagrees with its artifact fails the harness.

We preserve the July 2026 collection unchanged at `validity/archive-2026-07/`.

What a clone has. The run files under `validity/runs_framed`, `validity/runs_framed_lang` and `validity/runs` are gitignored; the English baseline runs under `validity/runs_english_baseline` and the September wave's under `validity/runs_neutral_template` are tracked. We archived the full grid, and the wave beside it, with sha256 manifests at the prefixes named in `validity/README.md`, which also gives the restore recipe, and `validity/reconcile.py` classifies a working copy against that archive. `analysis/test_reproduce.py` regenerates `results/appendix_tables.md`, `results/appendix_b4_b5.md`, `results/condition_means.json` and `results/inlanguage_audit.txt` from the dataset and compares them with its manifest (decision 22), and hashes `results/viewer_data.json` and the module's other outputs as committed. The exception is the ratings dataset: `validity/build_ratings_dataset.py` writes every scored ratings, the grid's and the wave's, to `results/mfq2_ratings.csv`, one row per condition, model, iteration and item with the instrument file, the request seed, the item's position in that run's shuffled order and the rating, 104,400 rows, and `validity/check_ratings_dataset.py` rebuilds the 47 pinned condition means from it to 1e-9 and checks the wave's shape. Every model-side number in this appendix is a function of those ratings, and since 2026-09-09 both emitters read the dataset rather than the run files, together with `results/collection_record.json`, which the dataset builder writes from the run files: the failed-call counts B7 reports, the parser's rounding audit, and the six translated instructions as sent. A clone regenerates every table and interval from what it has; only the builder itself needs the runs. Local recomputation from the archived runs, the hash check a clone can run, and public regeneration are three different things, and all three now hold for every model-side number; the dataset and the collection record are the only artifacts that need the runs. What a clone can and cannot rebuild:

from a fresh clone	rebuilds	needs
the pilot's fifteen outputs	yes, byte for byte	python3
the in-language condition means, all 47	yes, from <code>results/mfq2_ratings.csv</code>	python3
the hash check on every pinned artifact	yes	python3
B3a, B4, B4a, B6 and B6a tables and intervals	yes, from the dataset, and the harness does it	python3

from a fresh clone	rebuilds	needs
B1a's prompt quotations and B7's failed-call counts	yes, from the runner source and <code>results/collection_record.json</code>	python3
the dataset and the collection record themselves	no	the archived run files
the human reference CSVs	yes, from the sources' shared data, with the builders in <code>validity/reference/</code>	python3; R with sirt 3.13-228 for the alignment diagnostics; <code>pyreadstat</code> for Iran

The human side. `validity/reference/` holds four committed CSVs and a README recording each one's provenance. `mfq2_country_means.csv` and `mfq2_country_dispersion.csv` carry the nineteen countries' foundation means, standard errors and respondent-level standard deviations, the last for each foundation, the binding composite and the Loyalty-Authority composite. `build_mfq2_means.py` and `build_mfq2_dispersion.py`, in the same directory, compute them from Atari et al.'s Study 2 respondent-level data, which the authors share at osf.io/9dwzt with their scoring code at osf.io/vwrpn, following that code: the three attention checks as the inclusion rule, each foundation the mean of its six items over respondents who answered all six, each composite the mean of its foundations for respondents with all six complete, which is every one of the 3,902 the attention checks keep, so the two builders count the same respondents, and the standard deviation taken over respondents within a country. `compute_alignment_r2.R` recomputes the B2a diagnostics on the same data with sirt 3.13-228, into `mfq2_alignment_r2.csv`. The respondent file stays in the gitignored `_raw/`; the means builder reproduces the committed CSV byte for byte from it, and the CSVs are what the appendix reads. We build `mfq2_iran_dispersion.csv` here with `validity/reference/build_iran_dispersion.py`, which needs `pyreadstat`, from the respondent-level files Hazrati et al. share at osf.io/zt3u2; those stay in the gitignored `_raw/`. `validity/anchors_iran.json` carries Iran's three anchors and the caveats B4 reports.

The model-side analysis is stdlib-reproducible from the raw runs; we build the human reference figures from the sources' shared data as B9 describes. Responsibility for the work, and for any errors in it, is mine alone. Methodology was AI-assisted, and we disclose that assistance.

Declan Michaels | Cross-Cultural Alignment Study | moral-os.com