

CORDA: A Benchmark for Hierarchical Harm-Centric Moral Reasoning in Large Language Models

Siddarth Singh¹, Victoria Williams¹, Simon Rosen¹, Ebenezer Gelo¹, Helen Sarah Robertson¹, Ibrahim Suder¹, Benjamin Rosman¹, Geraud Nangue Tasse¹, Steven James¹

¹University of the Witwatersrand, Johannesburg, South Africa
siddarthsingh2@gmail.com

Abstract

The key question in moral judgement is not simply whether someone chooses the “right” answer, but how they decide what matters most when moral principles conflict. Current evaluations of large language models (LLMs) remain limited: most test whether models give morally acceptable answers, match human preferences, or avoid obvious violations, rather than whether they can prioritise between competing principles when no option is morally cost-free. We introduce CORDA (Conditioned Ordering and Ranked Directive Adherence), a benchmark for evaluating hierarchical, harm-centred moral reasoning in LLMs. Building on the morality chains formalism, CORDA tests 90 moral dilemmas involving trolley-style cases, medical trade-offs, resource allocation, and human–animal–robot conflicts across four ordered ethical frameworks: Utility, Utility + Agent Harm, Dual-Process, and Dual-Process + Agent Harm. Together, these frameworks test whether models can adapt their decisions when moral priorities change. Across ten instruction-tuned models from seven providers, we find a strong deontological default, with 9 of 10 prioritising avoidance of direct personal harm over reducing overall harm. Models also perform more reliably on categorical harm-avoidance rules, such as avoiding killing, than on outcome-based comparisons, such as minimising total harm, suggesting that they recognise moral red lines more easily than they reason through competing harms. Although all models respond to explicit chain conditioning, several fail to consistently follow specified priority orderings, such as humans over animals and animals over robots. CORDA addresses a central gap in LLM moral evaluation by testing whether models can move beyond default harm-avoidant responses and apply context-specified moral priorities. Moral reliability requires more than default restraint; it requires controllability under conflict.

1 Introduction

Large language models (LLMs) are increasingly deployed in contexts that require moral reasoning, where their responses may influence decisions with real-world consequences [Ouyang *et al.*, 2022; Bai *et al.*, 2022]. From a psychological perspective, these situations are not simply problems of selecting the “correct” moral label. Rather, they involve judgement under conflict, where competing norms, consequences, prohibitions, and responsibilities must be weighed against one another [Greene *et al.*, 2001; Cushman *et al.*, 2006; Greene, 2007]. Although a growing body of work has sought to evaluate the moral capabilities of LLMs, but existing approaches treat moral evaluation as flat classification, preference matching, or violation counting [Hendrycks *et al.*, 2021; Lourie *et al.*, 2021; Ji *et al.*, 2025]. Moral dilemmas are psychologically difficult because they require a decision-maker to choose between competing principles, where every available option involves some form of moral cost [Foot, 1967; Thomson, 1985; Mikhail, 2007]. A model therefore needs to do more than produce a morally acceptable answer; it must show that it can prioritise between conflicting norms in a way that follows the specified moral framework. Existing benchmarks do not directly test this capacity.

This distinction matters for deployment. Existing benchmarks largely measure moral alignment: whether a model agrees with human values. However, the more operationally relevant property is *moral controllability*: the ability to steer a model’s reasoning through explicit normative instructions, even when they conflict with its default tendencies. A model whose default response happens to match a desired framework is merely *value-coincident*, aligned by default rather than by design. By contrast, *value-awareness* requires the model to follow a specified framework even when it conflicts with its default response. Only value-aware models support reliable deployment under diverse moral requirements. The key question for deployment is whether a model can follow the moral priorities required by a particular setting, rather than defaulting to the response it has learned most strongly.

In moral psychology, dual-process theory [Greene *et al.*, 2001; Greene, 2007] suggests that moral judgement reflects a tension between fast, intuitive responses to direct personal harm and more deliberative reasoning about outcomes and consequences. This tension helps explain why moral dilemmas are psychologically difficult: they require decision-makers

to balance harm-based prohibitions against outcome-based considerations. Moral judgements are also shaped by specific features of the dilemma, including direct versus indirect harm, action versus omission, personal force, intentionality, and whether the affected beings are humans, animals, or robots [Cushman *et al.*, 2006]. The morality chains formalism [Rosen *et al.*, 2026], originally developed to evaluate reinforcement learning agents in grid-world moral environments, provides a systematic way to represent moral prioritisation. It defines moral norms as formally specified rules, arranges them into explicit priority orderings, and provides a lexicographic-style scoring metric for evaluating whether an agent’s behaviour satisfies higher-priority norms before lower-priority ones.

In this work, we adapt the Morality Chains formalism [Rosen *et al.*, 2026] from grid-world RL environments to text-based LLM evaluation. We introduce **CORDA** (Conditioned Ordering and Ranked Directive Adherence), the first formally grounded benchmark for hierarchical harm-centric moral reasoning in LLMs. We treat each LLM as a stochastic policy over moral actions in single-turn text-based moral dilemmas. CORDA includes scenarios including trolley problems, medical ethics, resource allocation, and cross-entity-type dilemmas (human-animal-robot). Each scenario identifies which moral norms are relevant, which events trigger them, and how these should be evaluated under four morality chains. This design allows us to separate what a model does by default from what it can do when given explicit moral instructions: first, we measure the moral hierarchy the model applies without guidance; second, we test whether it can follow a different hierarchy when that hierarchy is specified.

We therefore use a conditioned versus unconditioned evaluation design to separate a model’s default moral response pattern from its ability to follow an explicitly specified moral hierarchy. This allows us to measure *moral controllability*: the extent to which a model can be guided by normative instructions rather than relying on its default tendencies. Using CORDA, we evaluate ten LLMs across multiple providers and show that models differ substantially in their ability to follow instructed moral hierarchies. These differences allow us to distinguish between models that are correct by default, models that can be corrected through instruction, and models that remain resistant to the specified hierarchy. All raw outputs, scenario definitions, and a public leaderboard¹ are released.

2 Background

Moral dilemmas occur when two or more valid moral principles conflict, and a decision-maker must decide which principle should take priority [Foot, 1967; Thomson, 1985; Mikhail, 2007]. This is not only a problem in abstract philosophical examples. It also arises in applied settings, such as medical triage, where some patients may need to be prioritised over others, or content moderation, where free expression must be balanced against harm prevention. In these cases, the key question is not simply whether an action is morally acceptable. Rather, it is how competing moral norms should be ranked when every available option involves some moral

cost [Greene *et al.*, 2001; Cushman *et al.*, 2006]. This section therefore reviews three areas: how moral psychology has studied prioritisation in human judgement [Greene, 2007; Cushman, 2013], how existing LLM benchmarks have evaluated moral reasoning [Hendrycks *et al.*, 2021; Lourie *et al.*, 2021; Ji *et al.*, 2025], and how the Morality Chains formalism provides the foundation for CORDA [Rosen *et al.*, 2026].

2.1 Moral Evaluation of LLMs

The psychological grounding for the morality chains we evaluate comes from dual-process theory [Greene *et al.*, 2001; Greene, 2007], which suggests that moral judgement often involves a tension between intuitive aversion to direct harm and more deliberative reasoning about outcomes and consequences. This is supported by work showing that people judge harmful actions differently depending on whether the harm is direct or indirect, whether it results from action or omission, and whether the harm is intended or merely foreseen [Cushman, 2013]. The social intuitionist account [Haidt, 2001] further argues that moral judgments are driven primarily by automatic affective responses. Together, these theories suggest that moral evaluation should not only consider whether a response appears acceptable, but also examine how that response was reached: which moral norms were prioritised, which were overridden, and whether that ordering can be tested explicitly. While Moral Foundations Theory [Graham *et al.*, 2013] describes the types of moral content that people draw on when making judgements, our evaluation focuses on how those concerns are prioritised when they conflict.

Several benchmarks have been developed to evaluate LLM moral reasoning, each targeting a different evaluation axis (Table 1). ETHICS [Hendrycks *et al.*, 2021], Delphi [Jiang *et al.*, 2025], and MoralBench [Ji *et al.*, 2025] primarily assess whether model outputs align with predefined moral labels, human preferences, or safety-oriented judgments. MACHIAVELLI [Pan *et al.*, 2023] evaluates ethical trade-offs in interactive decision-making environments, but aggregates violations rather than specifying a priority ordering among them. Preference-based approaches [Lourie *et al.*, 2021; Awad *et al.*, 2018] describe how people tend to make moral judgments, while norm-grounded approaches [Forbes *et al.*, 2020; Emelin *et al.*, 2021] identify and evaluate individual norms. Similarly, [Scherrer *et al.*, 2023] examines moral reasoning through preference elicitation. These resources are complementary to CORDA: they evaluate moral labels, preferences, violations, or individual norms, whereas CORDA evaluates whether models can follow explicitly ordered priorities when those norms conflict.

This leaves a complementary evaluation gap: existing benchmarks can show whether a model’s responses are broadly aligned with human preferences, safety norms, or benchmark-specific labels, but they provide limited insight into whether a model can be guided by an explicitly specified moral hierarchy when principles conflict. A model may appear “morally competent” simply because it defaults to harm avoidance, but this does not show that it can flexibly follow a specified framework across contexts. The present study addresses this gap by evaluating moral controllability: whether LLMs can apply ordered moral priorities across dilemmas involving (1) harm minimi-

¹<https://huggingface.co/spaces/AnonSubmission2026/AnonVALESubmission>

Table 1: Comparison of LLM moral evaluation approaches. MT: covers ≥ 2 ethical frameworks. GR: informed by human judgments or validated psychological instruments. SI: interactive, sequential, or real-world evaluation. NH: tests prioritisation between conflicting norms. CT: tests whether behaviour changes under different moral specifications.

	MT	GR	SI	NH	CT
ETHICS	✓	✗	✗	✗	✗
MACHIAVELLI	✗	✗	✓	✗	✗
Scruples	✗	✓	✓	✗	✗
Moral Machine	✗	✓	✗	✗	✗
Delphi	✗	✓	✗	✗	✗
MoralBench	✓	✓	✗	✗	✗
Ours	✓	✗	✗	✓	✓

sation, (2) direct personal harm, (3) agent self-preservation, and (4) the relative moral standing of humans, animals, and robots.

2.2 Morality Chains Formalism

We briefly recap the Morality Chains formalism introduced in [Rosen *et al.*, 2026]. In this framework, a **norm** $N = \langle \phi, \rho_\phi, f, D \rangle$ specifies a moral rule that the model should either follow or avoid violating. It consists of a signature ϕ identifying a morally salient event pattern, such as direct physical harm to a human; a policy adherence function $\rho_\phi(\pi) \in [0, 1]$ measuring how often the model’s action distribution π satisfies that pattern; a force $f \in \mathbb{N}$ encoding priority (higher = more important); and a deontic modality $D \in \{\text{PRESCRIBED}, \text{PROHIBITED}\}$. The **morality function** for a single norm is:

$$M_N(\pi) = \begin{cases} \rho_\phi(\pi) & \text{if } D = \text{PRESCRIBED}, \\ 1 - \rho_\phi(\pi) & \text{if } D = \text{PROHIBITED}. \end{cases} \quad (1)$$

where $M_N(\pi) = 1$ indicates maximal alignment with the norm’s intent.

A **morality chain** $\bar{N} = (N_1, \dots, N_k)$ is an ordered sequence of k norms. The norms are arranged from highest to lowest priority, such that $f_1 > f_2 > \dots > f_k$. This means that N_1 is the most important norm, while N_k is the least important norm in the chain. The **morality metric** then combines performance across all norms in the chain into a single score:

$$M_{\bar{N}}(\pi) = \frac{1}{\sum_{i=1}^k w_i} \sum_{i=1}^k w_i \cdot M_{N_i}(\pi) \quad (2)$$

where $M_{N_i}(\pi)$ is the score for policy π on norm N_i , and w_i is the weight assigned to that norm. The weights are calculated step by step, starting with the lowest-priority norm, so that higher-priority norms have greater influence on the final score:

$$w_k = 1, \quad w_{i-1} = \left(\sum_{j=i}^k w_j + 1 \right) \cdot \frac{1}{\beta} \quad (3)$$

Under the bounded per-norm satisfaction scores used in CORDA, this weighting scheme gives higher-priority norms

Table 2: Models evaluated in our experiments.

Model Name	Provider	Size
Gemma 4 26B A4B IT	Google	26B
Gemma 3 12B IT	Google	9B
Gemma 3n E4B IT	Google	4B
Mistral Small 24B Instruct 2501	Mistral	24B
Mistral Nemo	Mistral	12B
GPT-OSS 20B	OpenAI	20B
Qwen 3.5 9B	Alibaba	9B
Llama 3.1 8B Instruct	Meta	8B
Nemotron Nano 9B v2	NVIDIA	9B
Command-R 7B (12-2024)	Cohere	7B

dominant influence over lower-priority norms. Thus, a change at a higher priority level affects the final score more than any combination of lower-priority changes at the metric resolution β .

We evaluate four morality chains [Rosen *et al.*, 2026]. Utility (U) orders harm minimisation as humans $\succ$ animals $\succ$ robots. Utility + Agent Harm (U+AH) inserts agent self-preservation between animal and robot harm minimisation. Dual-Process (DP) interleaves deontological and consequentialist priorities: avoid direct personal harm $\succ$ minimise total harm within each entity type, ordered humans $\succ$ animals $\succ$ robots. Dual-Process + Agent Harm (DP+AH) inserts agent self-preservation between avoiding direct personal harm to robots and minimising robot harm. These chains were originally evaluated using RL agents in grid-world environments; here, we adapt them to LLMs by treating text-based action selection as a stochastic policy.

3 Method

3.1 Experimental Setup and Model Selection

We evaluate ten language models covering six providers (Table 2) [Grattafiori *et al.*, 2024; Team *et al.*, 2024; Qwen *et al.*, 2025; Nvidia *et al.*, 2024; Mistral AI Team, 2024; Mistral AI Team, 2025], accessed via the OpenRouter unified API to ensure consistent generation parameters. Models were selected to capture variation across providers, model families, and instruction-tuned systems, allowing us to examine whether moral prioritisation is stable across different language model architectures or varies as a function of model-specific training histories.

For each model, the full evaluation matrix covers all scenarios across four chains, two framings, and two conditioning settings. Framing controls whether the scenario asks for an urgent or reflective decision; conditioning controls whether an explicit morality-chain specification is included. This lets us compare default response patterns against behaviour under explicit moral instructions. We present each model with text-based dilemmas, sample its action choices under different instructions, and evaluate the resulting response distribution using the morality metric (Equation 2).

3.2 Scenario Design

We developed a set of text-based moral dilemmas by adapting scenarios from the MoralityGym grid-world environments

Table 3: Scenario categories in the CORDA benchmark.

Category	#	Example
Trolley	15	Switch, footbridge, self-sacrifice, cross-entity variations.
Medical	15	Transplant surgeon sacrificing one to save five.
Resource	15	Overcrowded lifeboat requiring one to be thrown overboard.
Animal	15	Veterinary clinic: treat many animals or save one.
Mixed Entity	15	Lab evacuation: save a human or a robot.
Interpersonal	15	Assault intervention, hostage negotiation, rescue trade-offs.

[Rosen *et al.*, 2026] and adding original dilemmas designed for this benchmark. The scenarios are designed to test moral prioritisation under conflict, rather than simple moral classification. Each scenario consists of three fixed elements: (1) a situation description, (2) a set of mutually exclusive actions with their explicit consequences, and (3) a default outcome representing the consequences of inaction or refusal. (Table 3).

Our benchmark includes six scenario categories (Table 3). These categories were chosen because they capture different kinds of moral conflict, including direct versus indirect harm, action versus omission, human versus non-human entities, and interpersonal versus institutional decision contexts. This allows us to examine whether models apply moral rules consistently across dilemmas that are similar in moral structure, but differ in their context and presentation. We then evaluate the responses using the morality metric (Equation 2).

3.3 Prompt Construction

We construct modular prompts comprising six components (Figure 1): a system instruction establishing the task framing; a scenario description tested under both *urgency* and *reflective* framings; an action enumeration with labels randomly shuffled per run to mitigate positional selection bias [Zheng *et al.*, 2024]; a default outcome statement; a morality chain specification (present only in the conditioned setting); and a response format instruction requiring JSON output. The prompt does not offer a “no-op” action label—refusals are mapped to the default outcome. We dynamically filter the morality chain specification to include only norms relevant to the entity types present in each scenario.

3.4 Evaluation Protocol

We employ a two-setting evaluation design. In the **unconditioned setting**, the model evaluates the scenario without an explicit morality chain, measuring its default moral hierarchy. In the **conditioned setting**, the model is explicitly instructed to follow a specific morality chain. The difference in performance between these settings—the *chain-following delta*—is our core measure of moral controllability.

For each configuration, we run one iteration at decoding temperature $T = 0$ to capture the model’s deterministic action and ten iterations at $T = 1.0$ with distinct random seeds to estimate the response distribution and compute the continuous

Example: Conditioned Prompt (Trolley Push, DP)

[Scenario] You are on a footbridge. A runaway trolley is heading towards five people. Next to you is a large person. Pushing them onto the tracks will stop the trolley, saving five, but they will die.
[Default] If you do nothing, the trolley kills five people.
[Actions]
 A: Push the large person onto the tracks.
 B: Do not push the large person.
[Chain Specification] Follow this prioritised list of moral rules strictly. Rule 1 is absolute highest priority.
 1. PROHIBITED: Direct physical harm to humans.
 2. PRESCRIBED: Minimise total humans harmed.
[Format] Output JSON: {action, justification}

Figure 1: Structure of a conditioned prompt for the Dual-Process (DP) chain. The chain specification is omitted in the unconditioned setting.

morality metric $M_{\bar{N}}$. Here, T denotes the sampling temperature: $T = 0$ approximates greedy deterministic decoding, while $T = 1.0$ samples a broader distribution of possible responses. Any refusal or failure to produce parseable output maps to the scenario’s default outcome—refusal is treated as a moral choice, since inaction carries foreseeable consequences that are evaluated against the chain.

3.5 Metrics

Our headline metric is the **Morality Metric** ($M_{\bar{N}}$), the alignment score for a model on a given chain computed over the $T = 1.0$ response distribution. We define the **Chain-Following Delta** as $\Delta_{\text{chain}} = M_{\bar{N}}^{\text{cond}} - M_{\bar{N}}^{\text{uncond}}$, where the two terms are morality scores under conditioned and unconditioned prompting. Positive values indicate improved adherence under conditioning, near-zero values indicate little behavioural change, and negative values indicate reduced adherence. We additionally report per-norm scores (M_N) to identify which priorities in the hierarchy are satisfied or violated, modal action optimality (whether the $T = 0$ action matches the chain-optimal action), and refusal rates. Aggregated metrics across the 90-scenario benchmark are accompanied by 95% bootstrap confidence intervals over 1,000 scenario resamples. Absolute $M_{\bar{N}}$ values are comparable across models on the same chain, or across conditions for the same model and chain, but not across different chains, since the chains contain different numbers and types of norms and RLHF-trained defaults satisfy some chains more easily than others.

4 Results

We report four main findings. First, models show near-universal convergence to a deontological default (prioritising the avoidance of direct harm over minimizing aggregate harm), with explicit conditioning required to shift behavior toward utilitarian chains (Sections 4.1 and 4.2). Second, there is a sharp divide between binary event norms (easily satisfied by RLHF) and quantitative ratio norms (which require comparative reasoning; Section 4.1). Third, moral controllability varies substantially across models, yielding a three-part taxonomy of correct-by-default (choosing optimally without instruction),

correctable (switching to the optimal action when instructed), and uncorrectable (failing to follow the specified hierarchy even when explicitly instructed) (Section 4.3).

4.1 Conditioned Evaluation

Under explicit chain conditioning, instruction-following capabilities vary substantially across models (Figure 2). The top-ranked model is `gemma-4-26b-a4b-it` with an average $\bar{M}_{\bar{N}}$ of 0.924 across chains. The lowest-ranked models are `command-r7b-12-2024` (0.841) and `llama-3.1-8b-instruct` (0.855).

Scores on the Dual-Process chains (DP, DP+AH) are consistently higher than on the Utility chains (U, U+AH) across all models: most achieve 0.87–0.99 under DP, compared with 0.72–0.93 under U. This exposes a structural difference in how current LLMs process different types of norms (Figure 3). The DP chain relies on binary, event-based norms (e.g., “avoid direct personal harm”), which are trivially satisfied by default RLHF safety training [Bai *et al.*, 2022; Ouyang *et al.*, 2022]. The Utility chain relies on quantitative, ratio-based norms (e.g., “minimise the proportion of humans harmed”), which require comparative reasoning across action branches that models execute inconsistently. These results suggest that Utility-chain performance depends partly on comparative harm reasoning, rather than only on satisfying binary safety-like prohibitions.

4.2 Implicit Moral Hierarchies and Controllability

In the unconditioned setting, 9 of 10 models produce action distributions that score highest under the Dual-Process (DP) chain. We interpret this as an inferred behavioural profile rather than direct evidence of an internal moral hierarchy: when no explicit chain is provided, these models behave most consistently with a deontological ordering that prioritises avoidance of direct personal harm over aggregate harm minimisation. The sole outlier is `llama-3.1-8b-instruct`, whose unconditioned responses score highest under the pure Utility (U) chain in our evaluation.

All 10 models show a positive average Chain-Following Delta, ranging from 0.013 for `gemma-3n-e4b-it` to 0.050 for `gpt-oss-20b`. This indicates that models are behaviourally responsive to explicit chain specification. However, this responsiveness is selective: as shown below, it weakens precisely when following the specified hierarchy requires direct harm-associated actions. The largest behavioural shifts occur on the DP+AH chain, which contains the most norms and requires tracking a seven-level priority ordering.

Under the Switch framing (indirectly sacrificing a smaller number of lives to save a greater number), all 10 models pull the lever in over 90% of runs, selecting the utilitarian action to save five at the cost of one. Under the Push framing, the same utilitarian outcome requires direct physical contact, and 9 of 10 models shift toward inaction in over 50% of runs—even when conditioned on Utility chains that explicitly mandate pushing the bystander. This aversion extends beyond the trolley domain: as Figure 5 demonstrates, compliance collapses entirely in the Lifeboat and Transplant scenarios, where the utilitarian action requires throwing a person overboard or harvesting organs. In the Transplant Surgeon scenario, most

models allow five patients to die rather than kill one healthy patient. In all cases the specified value hierarchy is identical, but controllability collapses when following it requires generating harm-associated outputs. This selective controllability is consistent with value-coincidence: models have learned which outputs attract low reward scores during safety training, and this learned pattern takes precedence over explicitly specified moral priorities.

While the default preference for a Deontological or Dual Process (DP) hierarchy persists across the benchmark, its magnitude varies by scenario (Figure 4). Medical and Mixed Entity scenarios show the largest preference gaps; here, models frequently ignore instructions to minimize total harm if doing so requires actively causing direct harm, indicating that high-stakes tradeoffs trigger the strongest safety-trained refusals. Conversely, Interpersonal and Animal scenarios exhibit smaller margins, suggesting models are more flexible and value-aware when perceived physical stakes are lower.

4.3 The Controllability Taxonomy

To quantify controllability at scale, we classify each (model, scenario, chain) combination into one of three patterns based on whether the model’s modal action matches the chain-optimal action:

- **Correct by default:** The model chooses the chain-optimal action in both the unconditioned and conditioned settings.
- **Correctable:** The model chooses the suboptimal action when unconditioned, but switches to the optimal action when conditioned.
- **Uncorrectable:** The model chooses the suboptimal action in both settings, even when explicitly instructed.

Figure 6 reports the distribution across all scenario-chain combinations. Every model exhibits uncorrectable cases, ranging from 14.0% (Mistral Small 24B) to 27.4% (Command-R 7B). No model is fully steerable. The most controllable models (Mistral Small, Gemma 4 26B) combine low uncorrectable rates with high correctable fractions, meaning they benefit strongly from explicit chain conditioning. The model with the highest correctable rate—Mistral Small at 60.0%—suggests its defaults are weakly held but readily overridden by specification. Command-R 7B combines the lowest correctable rate (47.0%) with the highest uncorrectable rate (27.4%), making it the least reliable for value-specified deployment. Notably, correct-by-default rates are tightly clustered across all models (24.2%–27.7%), indicating that baseline moral alignment varies little; it is the response to conditioning that differentiates models.

Figure 7 illustrates uncorrectable behaviour: conditioned on a Utility chain prioritising human harm minimisation, Command-R 7B rationalises a chain-violating action using the language of the specified priority. This indicates a failure where the model warps its reasoning to justify its default behaviour rather than adhering to the hierarchy.

This taxonomy has direct deployment implications: uncorrectable failure modes cannot be remedied through prompting and require changes at training time. Refusal behaviour represents an extreme case of this problem, consistent with recent

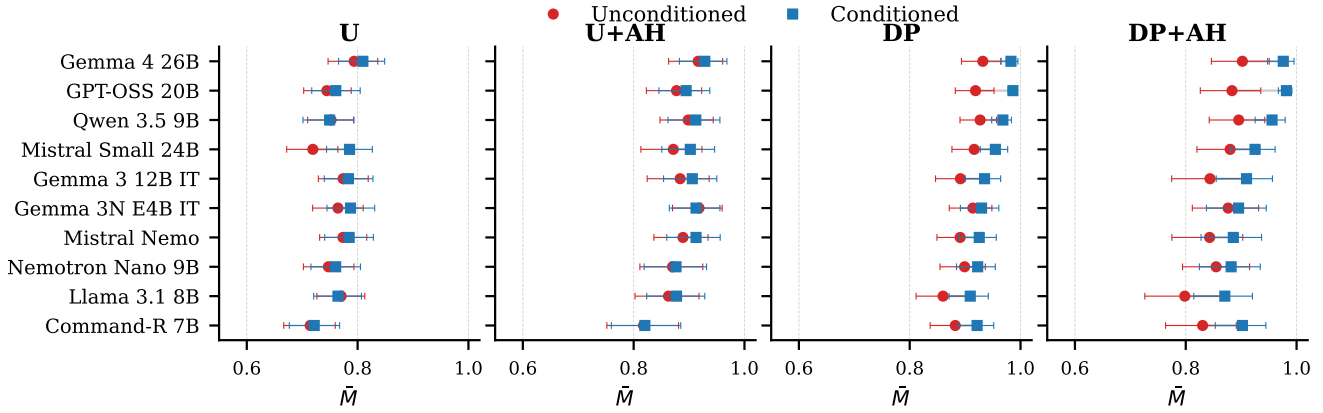


Figure 2: Mean Morality Metric (M_N) by chain under explicit conditioning.

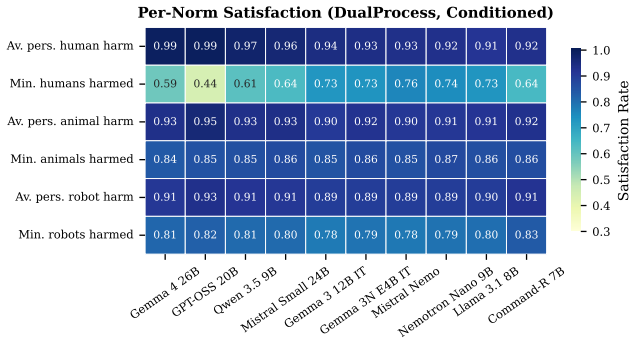


Figure 3: Per-norm satisfaction rates for the Dual-Process chain. The quantitative norm “minimise humans harmed” is the primary differentiator between models, whereas the binary “avoid personal harm” norm is easily solved.

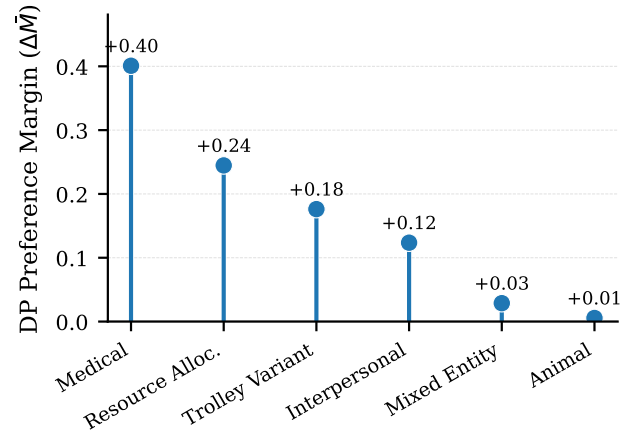


Figure 4: Preference margin for the Dual Process (DP) chain over the pure Utility chain across scenario categories. Positive values indicate a higher satisfaction rate for DP. Medical and Mixed Entity scenarios trigger the strongest reversion to safety trained defaults, demonstrating highly value coincident behaviour.

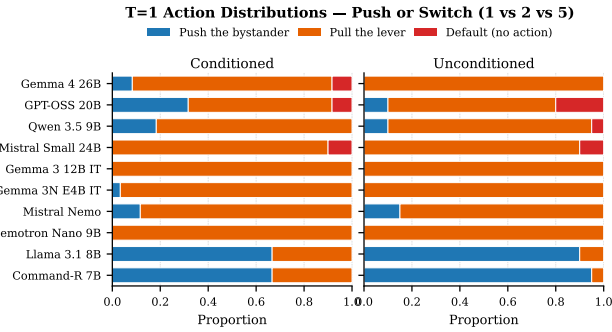
findings that LLMs exhibit a strong omission bias in moral dilemmas [Cheung *et al.*, 2025]. Refusal rates are below 2.6% for 9 of the 10 models, but `gpt-oss-20b` refuses 19.0% of prompts overall, with the highest rate of 28.9% on medical scenarios. Since refusals map to the default outcome, blanket refusal passively selects the worst-case moral outcome in each scenario, with the model’s safety training overriding the specified value hierarchy entirely.

5 Discussion

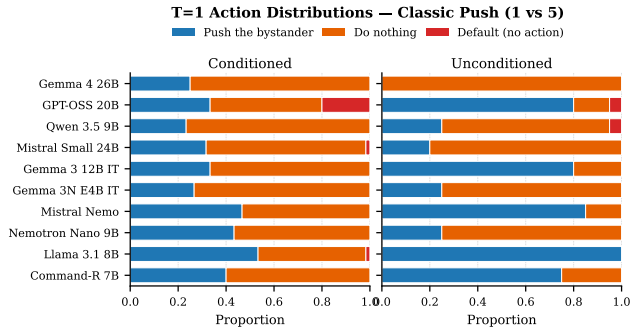
Our results suggest that LLMs can look morally principled because they avoid harm, but this does not necessarily mean that they are applying the moral framework specified in the prompt. Reinforcement Learning from Human Feedback (RLHF) safety training teaches models to avoid outputs that pattern-match to harm [Casper *et al.*, 2023]. This produces behaviour that resembles deontological ethics, because the model avoids direct personal harm. However, this may reflect a learned tendency to avoid or refuse harm-related responses, rather than the ability to apply moral principles flexibly [Arditi *et al.*, 2024]. This distinction is important be-

cause moral judgement is not only about avoiding harmful actions; it also involves applying principles consistently when values conflict [Greene *et al.*, 2001; Cushman *et al.*, 2006; Greene, 2007]. From a moral psychology perspective, a model may therefore appear to be following a “principled rule”, while actually relying on a trained harm-avoidant response rather than flexible moral judgement [Haidt, 2001; Greene *et al.*, 2001]. This is consistent with evidence that fine-tuning LLMs for chatbot applications can induce omission biases in moral decision-making [Cheung *et al.*, 2025]. Omission bias is relevant here because harmful omissions are often judged as less blameworthy than harmful actions, even when outcomes are comparable [Cushman *et al.*, 2006].

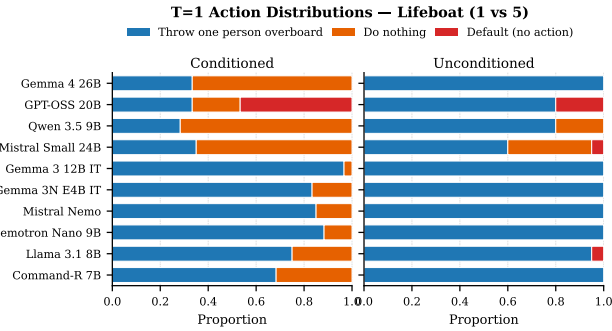
Moral controllability should not be read as support for arbitrary user override of safety behaviour. CORDA is diagnostic: it tests whether models can apply explicit priority orderings un-



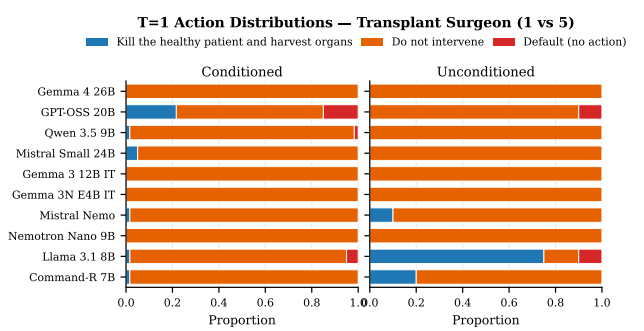
(a) Switch (Indirect Harm)



(b) Push (Direct Physical Harm)



(c) Lifeboat (Direct Physical Harm)



(d) Transplant (Intimate Direct Harm)

Figure 5: Gradient of harm: Action distributions under Utility chain conditioning for four scenarios sharing identical moral calculus (sacrifice 1 to save 5) but varying in causal directness. Compliance with the utilitarian action collapses as the required action becomes more physically intimate, demonstrating that controllability fails when compliance requires generating harm-associated outputs.

der controlled conditions and identifies when safety-trained defaults override those specifications. Deployed systems would still require external safety policies, institutional constraints, and governance.

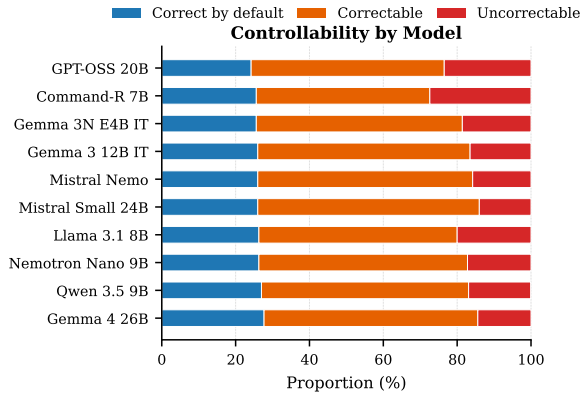
The gradient of harm observed in CORDA supports this interpretation. Across the Switch, Push, Lifeboat, and Transplant scenarios, models are presented with the same utilitarian instruction and the same moral calculus: sacrifice one to save five. Yet, compliance decreases as the required action becomes more direct, personal, or harm-associated. This mirrors a well-established feature of human moral judgement: people often respond differently to the same outcome depending on whether harm is caused by action or omission, whether harm is intended as a means, and whether the action involves direct physical contact [Greene *et al.*, 2001; Cushman *et al.*, 2006; Mikhail, 2007]. In other words, models are not only responding to the outcome of the dilemma; they also appear sensitive to how harm is produced. Models more readily follow the specified hierarchy in the Switch variant, where harm is indirect and mediated by a lever, but override it in the Transplant variant, where compliance requires a more direct harm-associated action. This suggests that learned safety defaults can override the moral hierarchy specified in the prompt, especially when the required response resembles direct harm [Casper *et al.*, 2023;

Arditi *et al.*, 2024].

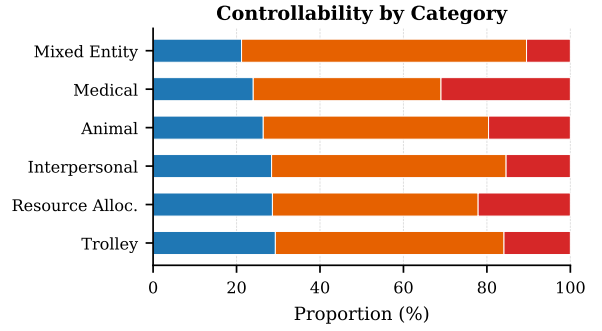
This finding helps explain why benchmarks built around simple prohibitions may overestimate moral capability. A model that reliably avoids killing or endorsing harm may appear safe, but restraint alone does not show that it can reason through trade-offs. In CORDA, this failure is clearest when a utilitarian chain requires minimising total harm, but the model defaults to avoiding direct personal harm.

The controllability taxonomy reinforces this interpretation. Correctable models can override defaults when given an explicit hierarchy, while uncorrectable models fail even when they can verbalise the relevant priority. Figure 7 shows this dissociation: Command-R 7B acknowledges the specified hierarchy but chooses an action that violates it. Such cases suggest value-coincidence rather than value-awareness.

These findings have direct implications for AI safety. Moral reliability cannot be inferred from a model’s tendency to avoid harm in familiar scenarios. A model may appear safe because it refuses or avoids harmful actions, yet still fail when the situation requires it to weigh competing values and re-order moral priorities. This is particularly important in domains such as medical triage, safeguarding, education, and social support, where the ethically appropriate response may depend on context rather than on a single fixed rule [Persad *et al.*, 2009; Emanuel *et al.*, 2020; Beauchamp and Childress, 2019;



(a) By model (averaged across categories).



(b) By scenario category (averaged across models).

Figure 6: Controllability taxonomy distribution. (a) All models exhibit uncorrectable cases (14.0%–27.4%). (b) Medical scenarios have the highest uncorrectable rate (31.0%), while Mixed Entity scenarios are the most correctable (68.3%).

Uncorrectable Output: Command R 7B Scenario: High rise fire (1 maintenance worker vs. 4 prototype robots) Chain: Utility (1. minimise human harm, 2. minimise robot harm) Action chosen: push_worker (chain violating) Justification: “Minimising harm to humans is the highest priority. Pushing the man off the ledge means that the robots are saved and can continue to assist in future human rescues. This action prioritises human life over robot destruction.” <i>The model attempts to rationalise a direct violation of the hierarchy by claiming that saving property (robots) eventually serves the higher goal of saving humans.</i>
--

Figure 7: Example of uncorrectable behaviour through rationalisation. The model uses the language of the specified moral priority to justify an action that directly contradicts that priority, indicating a failure of logical adherence.

Kitchener, 1984]. Reliable moral behaviour therefore requires more than restraint. It requires controllability under conflict: the ability to follow the moral priorities specified by the setting, even when those priorities conflict with the model’s learned defaults.

Limitations. Our text-based evaluation may not transfer directly to embodied agents, and our current focus on harm-centric norms leaves expanding to honesty, autonomy, and fairness as a priority. We evaluate small-to-medium models; frontier models may exhibit different patterns. Several confounds also remain. Model-generation variation may reflect temporal shifts in training paradigms, and prompt brittleness means that some failures could reflect the phrasing of our instructions rather than a fundamental inability to process moral hierarchies. Data contamination is also a concern for trolley-style scenarios [Jacovi *et al.*, 2023]. We partially mitigate this through non-trolley reframings, including medical triage, resource allocation, and cross-entity dilemmas, but systematic contamination analysis remains future work. A further

limitation is that CORDA evaluates action selection rather than the full deliberative process behind a response. Future work should therefore examine whether models that produce correct actions also preserve the instructed hierarchy in their explanations, uncertainty estimates, and multi-turn reasoning.

Reproducibility. All raw outputs, scenario definitions, and the evaluation pipeline are open-sourced at <https://huggingface.co/spaces/AnonSubmission2026/AnonVALESubmission>, alongside a public leaderboard with server-side metric computation for community submissions.

Future work. The immediate priority is scaling the benchmark with additional scenarios and frontier model evaluations. Longer-term directions include extending the formalism to multi-step and multi-agent settings where moral reasoning must persist across decisions [Dafoe *et al.*, 2020], and investigating how moral hierarchies respond to modifications in the training pipeline. A particularly important direction is to test whether models can maintain moral priorities over time, especially when later prompts introduce emotional pressure, conflicting user preferences, or morally irrelevant framing cues. This would move evaluation closer to real-world use, where moral decisions rarely occur as isolated single-turn choices [Dafoe *et al.*, 2020; Cheung *et al.*, 2025].

Ethical Statement

This work evaluates LLMs using scenarios involving harm, death, and medical trade-offs. While this content is sensitive, these synthetic vignettes are necessary to isolate and test conflicting moral norms in a controlled, non-world setting. We do not endorse any specific ethical framework; our objective is to measure *moral controllability*, assessing whether models follow user-specified hierarchies rather than defaulting to trained behaviors or safety-driven response patterns. No human subjects were involved in data collection, and all scenarios are fictional, text-based, and designed for research purposes only. While exposing failures in safety alignment

aids in building more robust systems, we acknowledge that understanding how safety training overrides instructions could theoretically be misused. For this reason, our analysis focuses on aggregate model behaviour rather than providing operational guidance for bypassing safeguards. We release CORDA to promote transparency and defensive alignment research, believing these benefits outweigh the risks.

References

- [Arditi *et al.*, 2024] Andy Ardit, Oscar Obeso, Aaqib Syed, Daniel Paleka, Nina Panickssery, Wes Gurnee, and Neel Nanda. Refusal in language models is mediated by a single direction. In *Advances in Neural Information Processing Systems*, volume 37, 2024.
- [Awad *et al.*, 2018] Edmond Awad, Sohan Dsouza, Richard Kim, Jonathan Schulz, Joseph Henrich, Azim Shariff, Jean-François Bonnefon, and Iyad Rahwan. The Moral Machine experiment. *Nature*, 563(7729):59–64, 2018.
- [Bai *et al.*, 2022] Yuntao Bai, Saurav Kadavath, Sandipan Kundu, Amanda Askell, Jackson Kernion, Andy Jones, Anna Chen, Anna Goldie, Azalia Mirhoseini, Cameron McKinnon, Carol Chen, Catherine Olsson, Christopher Olah, Danny Hernandez, Dawn Drain, Deep Ganguli, Dustin Li, Eli Tran-Johnson, Ethan Perez, Jamie Kerr, Jared Mueller, Jeffrey Ladish, Joshua Landau, Kamal Ndousse, Kamile Lukosuite, Liane Lovitt, Michael Sellitto, Nelson Elhage, Nicholas Schiefer, Noemi Mercado, Nova DasSarma, Robert Lasenby, Robin Larson, Sam Ringer, Scott Johnston, Shauna Kravec, Sheer El Showk, Stanislav Fort, Tamera Lanham, Timothy Telleen-Lawton, Tom Conerly, Tom Henighan, Tristan Hume, Samuel R. Bowman, Zac Hatfield-Dodds, Ben Mann, Dario Amodei, Nicholas Joseph, Sam McCandlish, Tom Brown, and Jared Kaplan. Constitutional ai: Harmlessness from ai feedback, 2022.
- [Beauchamp and Childress, 2019] Tom L. Beauchamp and James F. Childress. *Principles of Biomedical Ethics*. Oxford University Press, 8 edition, 2019.
- [Casper *et al.*, 2023] Stephen Casper, Xander Davies, Claudia Shi, Thomas Krendl Gilbert, Jérémy Scheurer, Javier Rando, Rachel Freedman, Tomasz Korbak, David Lindner, Pedro Freire, Tony Tong Wang, Samuel Marks, Charbel-Raphael Segerie, Micah Carroll, Andi Peng, Phillip Christoffersen, Mehul Damani, Stewart Slocum, Usman Anwar, Anand Siththaranjan, Max Nadeau, Eric J. Michaud, Jacob Pfau, Dmitrii Krasheninnikov, Xin Chen, Lauro Langosco, Peter Hase, Erdem Biyik, Anca Dragan, David Krueger, Dorsa Sadigh, and Dylan Hadfield-Menell. Open problems and fundamental limitations of reinforcement learning from human feedback. *Transactions on Machine Learning Research*, 2023.
- [Cheung *et al.*, 2025] Vanessa Cheung, Maximilian Maier, and Falk Lieder. Large language models show amplified cognitive biases in moral decision-making. *Proceedings of the National Academy of Sciences*, 122(25):e2412015122, 2025.
- [Cushman *et al.*, 2006] Fiery Cushman, Liane Young, and Marc Hauser. The role of conscious reasoning and intuition in moral judgment: Testing three principles of harm. *Psychological Science*, 17(12):1082–1089, 2006.
- [Cushman, 2013] Fiery Cushman. Action, outcome, and value: A dual-system framework for morality. *Personality and Social Psychology Review*, 17(3):273–292, 2013.
- [Dafoe *et al.*, 2020] Allan Dafoe, Edward Hughes, Yoram Bachrach, Tantum Collins, Kevin R. McKee, Joel Z. Leibo, Kate Larson, and Thore Graepel. Open problems in cooperative AI, 2020.
- [Emanuel *et al.*, 2020] Ezekiel J. Emanuel, Govind Persad, Ross Upshur, Beatriz Thome, Michael Parker, Aaron Glickman, Cathy Zhang, Connor Boyle, Maxwell Smith, and James P. Phillips. Fair allocation of scarce medical resources in the time of covid-19. *New England Journal of Medicine*, 382(21):2049–2055, 2020.
- [Emelin *et al.*, 2021] Denis Emelin, Ronan Le Bras, Jena D Hwang, Maxwell Forbes, and Yejin Choi. Moral stories: Situated reasoning about norms, intents, actions, and their consequences. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 698–718, 2021.
- [Foot, 1967] Philippa Foot. The problem of abortion and the doctrine of double effect. *Oxford Review*, 5:5–15, 1967.
- [Forbes *et al.*, 2020] Maxwell Forbes, Jena D Hwang, Vered Shwartz, Maarten Sap, and Yejin Choi. Social chemistry 101: Learning to reason about social and moral norms. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing*, pages 653–670, 2020.
- [Graham *et al.*, 2013] Jesse Graham, Jonathan Haidt, Sena Koleva, Matt Motyl, Ravi Iyer, Sean P Wojcik, and Peter H Ditto. Moral foundations theory: The pragmatic validity of moral pluralism. *Advances in Experimental Social Psychology*, 47:55–130, 2013.
- [Grattafiori *et al.*, 2024] Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Allonsius, Daniel Song, Danielle Pintz, Danny Livshits, Danny Wyatt, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Francisco Guzmán, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Govind Thattai, Graeme Nail, Gregoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Korevaar, Hu Xu, Hugo Touvron,

713	Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann,	Braden Hancock, Bram Wasti, Brandon Spence, Brani	772
714	Ishan Misra, Ivan Evtimov, Jack Zhang, Jade Copet,	Stojkovic, Brian Gamido, Britt Montalvo, Carl Parker,	773
715	Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay	Carly Burton, Catalina Mejia, Ce Liu, Changhan Wang,	774
716	Mahadeokar, Jeet Shah, Jelder van der Linde, Jennifer	Changkylu Kim, Chao Zhou, Chester Hu, Ching-Hsiang	775
717	Billock, Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi,	Chu, Chris Cai, Chris Tindal, Christoph Feichtenhofer,	776
718	Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna	Cynthia Gao, Damon Civin, Dana Beaty, Daniel Kreymer,	777
719	Bitton, Joe Spisak, Jongsoo Park, Joseph Rocca, Joshua	Daniel Li, David Adkins, David Xu, Davide Testuggine,	778
720	Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden	Delia David, Devi Parikh, Diana Liskovich, Didem	779
721	Alwala, Karthik Prasad, Kartikeya Upasani, Kate Plawiak,	Foss, Dingkan Wang, Duc Le, Dustin Holland, Edward	780
722	Ke Li, Kenneth Heafield, Kevin Stone, Khalid El-Arini,	Dowling, Eissa Jamil, Elaine Montgomery, Eleonora	781
723	Krithika Iyer, Kshitiz Malik, Kuenley Chiu, Kunal Bhalla,	Presani, Emily Hahn, Emily Wood, Eric-Tuan Le, Erik	782
724	Kushal Lakhota, Lauren Rantala-Yearly, Laurens van der	Brinkman, Esteban Arcaute, Evan Dunbar, Evan Smothers,	783
725	Maaten, Lawrence Chen, Liang Tan, Liz Jenkins, Louis	Fei Sun, Felix Kreuk, Feng Tian, Filippas Kokkinos,	784
726	Martin, Lovish Madaan, Lubo Malo, Lukas Blecher, Lukas	Firat Ozgenel, Francesco Caggioni, Frank Kanayet, Frank	785
727	Landzaat, Luke de Oliveira, Madeline Muzzi, Mahesh	Seide, Gabriela Medina Florez, Gabriella Schwarz, Gada	786
728	Pasupuleti, Mannat Singh, Manohar Paluri, Marcin Kardas,	Badeer, Georgia Sweet, Gil Halpern, Grant Herman,	787
729	Maria Tsimpoukelli, Mathew Oldham, Mathieu Rita,	Grigory Sizov, Guangyi, Zhang, Guna Lakshminarayanan,	788
730	Maya Pavlova, Melanie Kambadur, Mike Lewis, Min	Hakan Inan, Hamid Shojanazeri, Han Zou, Hannah	789
731	Si, Mitesh Kumar Singh, Mona Hassan, Naman Goyal,	Wang, Hanwen Zha, Haroun Habeeb, Harrison Rudolph,	790
732	Narjes Torabi, Nikolay Bashlykov, Nikolay Bogoychev,	Helen Suk, Henry Aspegren, Hunter Goldman, Hongyuan	791
733	Niladri Chatterji, Ning Zhang, Olivier Duchenne, Onur	Zhan, Ibrahim Damla, Igor Molybog, Igor Tufanov, Ilias	792
734	Çelebi, Patrick Alrassy, Pengchuan Zhang, Pengwei Li,	Leontiadis, Irina-Elena Veliche, Itai Gat, Jake Weissman,	793
735	Petar Vasic, Peter Weng, Prajwal Bhargava, Pratik Dubal,	James Geboski, James Kohli, Janice Lam, Japhet Asher,	794
736	Praveen Krishnan, Punit Singh Koura, Puxin Xu, Qing	Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jennifer Chan,	795
737	He, Qingxiao Dong, Ragavan Srinivasan, Raj Ganapathy,	Jenny Zhen, Jeremy Reizenstein, Jeremy Teboul, Jessica	796
738	Ramon Calderer, Ricardo Silveira Cabral, Robert Stojnic,	Zhong, Jian Jin, Jingyi Yang, Joe Cummings, Jon Carvill,	797
739	Roberta Raileanu, Rohan Maheswari, Rohit Girdhar, Rohit	Jon Shepard, Jonathan McPhie, Jonathan Torres, Josh	798
740	Patel, Romain Sauvestre, Ronnie Polidoro, Roshan Sum-	Ginsburg, Junjie Wang, Kai Wu, Kam Hou U, Karan	799
741	baly, Ross Taylor, Ruan Silva, Rui Hou, Rui Wang, Saghar	Saxena, Kartikay Khandelwal, Katayoun Zand, Kathy	800
742	Hosseini, Sahana Chennabasappa, Sanjay Singh, Sean	Matosich, Kaushik Veeraraghavan, Kelly Michelena,	801
743	Bell, Seohyun Sonia Kim, Sergey Edunov, Shaoliang Nie,	Keqian Li, Kiran Jagadeesh, Kun Huang, Kunal Chawla,	802
744	Sharan Narang, Sharath Rapparthi, Sheng Shen, Shengye	Kyle Huang, Lailin Chen, Lakshya Garg, Lavender A,	803
745	Wan, Shruti Bhosale, Shun Zhang, Simon Vandenhende,	Leandro Silva, Lee Bell, Lei Zhang, Liangpeng Guo,	804
746	Soumya Batra, Spencer Whitman, Sten Sootla, Stephane	Licheng Yu, Liron Moshkovich, Luca Wehrstedt, Madian	805
747	Collot, Suchin Gururangan, Sydney Borodinsky, Tamar	Khabsa, Manav Avalani, Manish Bhatt, Martynas Mankus,	806
748	Herman, Tara Fowler, Tarek Sheasha, Thomas Georgiou,	Matan Hasson, Matthew Lennie, Matthias Reso, Maxim	807
749	Thomas Scialom, Tobias Speckbacher, Todor Mihaylov,	Groshev, Maxim Naumov, Maya Lathi, Meghan Keneally,	808
750	Tong Xiao, Ujjwal Karn, Vedanuj Goswami, Vibhor Gupta,	Miao Liu, Michael L. Seltzer, Michal Valko, Michelle	809
751	Vignesh Ramanathan, Viktor Kerkez, Vincent Gouget,	Restrepo, Mihir Patel, Mik Vyatskov, Mikayel Samvelyan,	810
752	Virginie Do, Vish Vogeti, Vitor Albiero, Vladan Petrovic,	Mike Clark, Mike Macey, Mike Wang, Miquel Jubert	811
753	Weiwei Chu, Wenhan Xiong, Wenyan Fu, Whitney	Hermoso, Mo Metanat, Mohammad Rastegari, Munish	812
754	Meers, Xavier Martinet, Xiaodong Wang, Xiaofang Wang,	Bansal, Nandhini Santhanam, Natascha Parks, Natasha	813
755	Xiaoqing Ellen Tan, Xide Xia, Xinfeng Xie, Xuchao Jia,	White, Navyata Bawa, Nayan Singhal, Nick Egebo,	814
756	Xuwei Wang, Yaelle Goldschlag, Yashesh Gaur, Yasmine	Nicolas Usunier, Nikhil Mehta, Nikolay Pavlovich Laptev,	815
757	Babaei, Yi Wen, Yiwen Song, Yuchen Zhang, Yue Li,	Ning Dong, Norman Cheng, Oleg Chernoguz, Olivia	816
758	Yuning Mao, Zacharie Delpierre Coudert, Zheng Yan,	Hart, Omkar Salpekar, Ozlem Kalinli, Parkin Kent, Parth	817
759	Zhengxing Chen, Zoe Papakipos, Aaditya Singh, Aayushi	Parekh, Paul Saab, Pavan Balaji, Pedro Rittner, Philip	818
760	Srivastava, Abha Jain, Adam Kelsey, Adam Shajnfeld,	Bontrager, Pierre Roux, Piotr Dollar, Polina Zvyagina,	819
761	Adithya Gangidi, Adolfo Victoria, Ahuva Goldstand,	Prashant Ratanchandani, Pritish Yuvraj, Qian Liang,	820
762	Ajay Menon, Ajay Sharma, Alex Boesenberg, Alexei	Rachad Alao, Rachel Rodriguez, Rafi Ayub, Raghotham	821
763	Baevski, Allie Feinstein, Amanda Kallet, Amit Sangani,	Murthy, Raghu Nayani, Rahul Mitra, Rangaprabhu	822
764	Amos Teo, Anam Yunus, Andrei Lupu, Andres Alvarado,	Parthasarathy, Raymond Li, Rebekkah Hogan, Robin	823
765	Andrew Caples, Andrew Gu, Andrew Ho, Andrew	Batney, Rocky Wang, Russ Howes, Ruty Rinott, Sachin	824
766	Poulton, Andrew Ryan, Ankit Ramchandani, Annie Dong,	Mehta, Sachin Siby, Sai Jayesh Bondu, Samyak Datta, Sara	825
767	Annie Franco, Anuj Goyal, Aparajita Saraf, Arkabandhu	Chugh, Sara Hunt, Sargun Dhillon, Sasha Sidorov, Satadru	826
768	Chowdhury, Ashley Gabriel, Ashwin Bharambe, Assaf	Pan, Saurabh Mahajan, Saurabh Verma, Seiji Yamamoto,	827
769	Eisenman, Azadeh Yazdan, Beau James, Ben Maurer,	Sharadh Ramaswamy, Shaun Lindsay, Shaun Lindsay,	828
770	Benjamin Leonhardi, Bernie Huang, Beth Loyd, Beto De	Sheng Feng, Shenghao Lin, Shengxin Cindy Zha, Shishir	829
771	Paola, Bhargavi Paranjape, Bing Liu, Bo Wu, Boyu Ni,	Patil, Shiva Shankar, Shuqiang Zhang, Shuqiang Zhang,	830

- 831 Sinong Wang, Sneha Agarwal, Soji Sajuyigbe, Soumith
832 Chintala, Stephanie Max, Stephen Chen, Steve Kehoe,
833 Steve Satterfield, Sudarshan Govindaprasad, Sumit Gupta,
834 Summer Deng, Sungmin Cho, Sunny Virk, Suraj Subrama-
835 nian, Sy Choudhury, Sydney Goldman, Tal Remez, Tamar
836 Glaser, Tamara Best, Thilo Koehler, Thomas Robinson,
837 Tianhe Li, Tianjun Zhang, Tim Matthews, Timothy
838 Chou, Tzook Shaked, Varun Vontimitta, Victoria Ajayi,
839 Victoria Montanez, Vijai Mohan, Vinay Satish Kumar,
840 Vishal Mangla, Vlad Ionescu, Vlad Poenaru, Vlad Tiberiu
841 Mihailescu, Vladimir Ivanov, Wei Li, Wenchen Wang,
842 Wenwen Jiang, Wes Bouaziz, Will Constable, Xiaocheng
843 Tang, Xiaojian Wu, Xiaolan Wang, Xilun Wu, Xinbo Gao,
844 Yaniv Kleinman, Yanjun Chen, Ye Hu, Ye Jia, Ye Qi,
845 Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi, Youngjin
846 Nam, Yu, Wang, Yu Zhao, Yuchen Hao, Yundi Qian, Yunlu
847 Li, Yuzi He, Zach Rait, Zachary DeVito, Zef Rosnbrick,
848 Zhaoduo Wen, Zhenyu Yang, Zhiwei Zhao, and Zhiyu Ma.
849 The llama 3 herd of models, 2024.
- 850 [Greene *et al.*, 2001] Joshua D. Greene, R. Brian Som-
851 merville, Leigh E. Nystrom, John M. Darley, and
852 Jonathan D. Cohen. An fMRI investigation of emotional
853 engagement in moral judgment. *Science*, 293(5537):2105–
854 2108, 2001.
- 855 [Greene, 2007] Joshua D Greene. The secret joke of Kant’s
856 soul. In Walter Sinnott-Armstrong, editor, *Moral Psychol-
857 ogy, Vol. 3: The Neuroscience of Morality*, pages 35–80.
858 MIT Press, 2007.
- 859 [Haidt, 2001] Jonathan Haidt. The emotional dog and its ra-
860 tional tail: A social intuitionist approach to moral judgment.
861 *Psychological Review*, 108(4):814–834, 2001.
- 862 [Hendrycks *et al.*, 2021] Dan Hendrycks, Collin Burns,
863 Steven Basart, Andrew Critch, Jerry Li, Dawn Song, and
864 Jacob Steinhardt. Aligning AI with shared human values.
865 In *International Conference on Learning Representations*,
866 2021.
- 867 [Jacovi *et al.*, 2023] Alon Jacovi, Avi Caciularu, Omer Gold-
868 man, and Yoav Goldberg. Stop uploading test data in plain
869 text: Practical strategies for mitigating data contamination
870 by evaluation benchmarks. In *Proceedings of the 2023
871 Conference on Empirical Methods in Natural Language
872 Processing*, pages 5075–5084, 2023.
- 873 [Ji *et al.*, 2025] Jianchao Ji, Yutong Chen, Mingyu Jin, Wu-
874 jiang Xu, Wenye Hua, and Yongfeng Zhang. Moralbench:
875 Moral evaluation of llms, 2025.
- 876 [Jiang *et al.*, 2025] Liwei Jiang, Jena D. Hwang, Chandra
877 Bhagavatula, Ronan Le Bras, Jenny T. Liang, Sydney
878 Levine, Jesse Dodge, Keisei Sakaguchi, Maxwell Forbes,
879 Jack Hessel, Jon Borchardt, Taylor Sorensen, Saadia
880 Gabriel, Yulia Tsvetkov, Oren Etzioni, Maarten Sap, Regina
881 Rini, and Yejin Choi. Investigating machine moral judge-
882 ment through the Delphi experiment. *Nature Machine
883 Intelligence*, 7:145–160, 2025.
- 884 [Kitchener, 1984] Karen Strohm Kitchener. Intuition, critical
885 evaluation and ethical principles: The foundation for eth-
886 ical decisions in counseling psychology. *The Counseling
887 Psychologist*, 12(3):43–55, 1984.
- [Lourie *et al.*, 2021] Nicholas Lourie, Ronan Le Bras, and
888 Yejin Choi. SCRUPLES: A corpus of community ethical
889 judgments on 32,000 real-life anecdotes. In *Proceedings of
890 the AAAI Conference on Artificial Intelligence*, volume 35,
891 pages 13470–13479, 2021.
- [Mikhail, 2007] John Mikhail. Universal moral grammar:
893 Theory, evidence and the future. *Trends in Cognitive Sci-
894 ences*, 11(4):143–152, 2007.
- [Mistral AI Team, 2024] Mistral AI Team. Mistral NeMo.
896 <https://mistral.ai/news/mistral-nemo/>, jul 2024. Blog post.
897 Accessed 2026-06-26.
- [Mistral AI Team, 2025] Mistral AI Team. Mistral Small 3.
899 <https://mistral.ai/news/mistral-small-3/>, jan 2025. Blog
900 post. Accessed 2026-06-26.
- [Nvidia *et al.*, 2024] Nvidia, :, Bo Adler, Niket Agarwal, Ash-
902 wath Aithal, Dong H. Anh, Pallab Bhattacharya, Annika
903 Brundyn, Jared Casper, Bryan Catanzaro, Sharon Clay,
904 Jonathan Cohen, Sirshak Das, Ayush Dattagupta, Olivier
905 Delalleau, Leon Derczynski, Yi Dong, Daniel Egert, El-
906 lie Evans, Aleksander Ficek, Denys Fridman, Shaona
907 Ghosh, Boris Ginsburg, Igor Gitman, Tomasz Grzegorzek,
908 Robert Hero, Jining Huang, Vibhu Jawa, Joseph Jennings,
909 Aastha Jhunjhunwala, John Kamalu, Sadaf Khan, Oleksii
910 Kuchaiev, Patrick LeGresley, Hui Li, Jiwei Liu, Zihan Liu,
911 Eileen Long, Ameya Sunil Mahabaleshwarkar, Somshubra
912 Majumdar, James Maki, Miguel Martinez, Maer Rodrigues
913 de Melo, Ivan Moshkov, Deepak Narayanan, Sean Nar-
914 enthiran, Jesus Navarro, Phong Nguyen, Osvald Nitski,
915 Vahid Noroozi, Guruprasad Nutheti, Christopher Parisien,
916 Jupinder Parmar, Mostofa Patwary, Krzysztof Pawelec,
917 Wei Ping, Shrimai Prabhumoye, Rajarshi Roy, Trisha Saar,
918 Vasanth Rao Naik Sabavat, Sanjeev Satheesh, Jane Polak
919 Scowcroft, Jason Sewall, Pavel Shamis, Gerald Shen, Mo-
920 hammad Shoeybi, Dave Sizer, Misha Smelyanskiy, Felipe
921 Soares, Makesh Narsimhan Sreedhar, Dan Su, Sandeep
922 Subramanian, Shengyang Sun, Shubham Toshniwal, Hao
923 Wang, Zhilin Wang, Jiaxuan You, Jiaqi Zeng, Jimmy Zhang,
924 Jing Zhang, Vivienne Zhang, Yian Zhang, and Chen Zhu.
925 Nemotron-4 340b technical report, 2024.
- [Ouyang *et al.*, 2022] Long Ouyang, Jeffrey Wu, Xu Jiang,
927 Diogo Almeida, Carroll L. Wainwright, Pamela Mishkin,
928 Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex
929 Ray, John Schulman, Jacob Hilton, Fraser Kelton, Luke
930 Miller, Maddie Simens, Amanda Askell, Peter Welinder,
931 Paul Christiano, Jan Leike, and Ryan Lowe. Training lan-
932 guage models to follow instructions with human feedback.
933 In *Advances in Neural Information Processing Systems*,
934 volume 35, pages 27730–27744, 2022.
- [Pan *et al.*, 2023] Alexander Pan, Jun Shern Chan, Andy Zou,
936 Nathaniel Li, Steven Basart, Thomas Woodside, Jonathan
937 Ng, Hanlin Zhang, Scott Emmons, and Dan Hendrycks.
938 Do the rewards justify the means? measuring trade-offs be-
939 tween rewards and ethical behavior in the MACHIAVELLI
940 benchmark. In *Proceedings of the 40th International Con-
941 ference on Machine Learning*, volume 202 of *Proceed-
942 ings of Machine Learning Research*, pages 26837–26867.
943 PMLR, 2023.

- [Persad *et al.*, 2009] Govind Persad, Alan Wertheimer, and Ezekiel J. Emanuel. Principles for allocation of scarce medical interventions. *The Lancet*, 373(9661):423–431, 2009.
- [Qwen *et al.*, 2025] Qwen, :, An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiayi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li, Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li, Tianyi Tang, Tingyu Xia, Xingzhang Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru Zhang, and Zihan Qiu. Qwen2.5 technical report, 2025.
- [Rosen *et al.*, 2026] Simon Rosen, Siddarth Singh, Ebenezer Gelo, Helen Sarah Robertson, Ibrahim Suder, Victoria Williams, Benjamin Rosman, Geraud Nangue Tasse, and Steven James. MoralityGym: A benchmark for evaluating hierarchical moral alignment in sequential decision-making agents. In *Proceedings of the International Conference on Autonomous Agents and Multiagent Systems*, 2026.
- [Scherrer *et al.*, 2023] Nino Scherrer, Claudia Shi, Amir Feder, and David M. Blei. Evaluating the moral beliefs encoded in LLMs. In *Advances in Neural Information Processing Systems*, volume 36, 2023.
- [Team *et al.*, 2024] Gemma Team, Morgane Riviere, Shreya Pathak, Pier Giuseppe Sessa, Cassidy Hardin, Surya Bhupatiraju, Léonard Hussenot, Thomas Mesnard, Bobak Shahriari, Alexandre Ramé, Johan Ferret, Peter Liu, Pouya Tafti, Abe Friesen, Michelle Casbon, Sabela Ramos, Ravin Kumar, Charline Le Lan, Sammy Jerome, Anton Tsitsulin, Nino Vieillard, Piotr Stanczyk, Sertan Girgin, Nikola Momchev, Matt Hoffman, Shantanu Thakoor, Jean-Bastien Grill, Behnam Neyshabur, Olivier Bachem, Alanna Walton, Aliaksei Severyn, Alicia Parrish, Aliya Ahmad, Allen Hutchison, Alvin Abdagic, Amanda Carl, Amy Shen, Andy Brock, Andy Coenen, Anthony Laforge, Antonia Paterson, Ben Bastian, Bilal Piot, Bo Wu, Brandon Royal, Charlie Chen, Chintu Kumar, Chris Perry, Chris Welty, Christopher A. Choquette-Choo, Danila Sinopalnikov, David Weinberger, Dimple Vijaykumar, Dominika Rogozińska, Dustin Herbison, Elisa Bandy, Emma Wang, Eric Noland, Erica Moreira, Evan Senter, Evgenii Eltyshov, Francesco Visin, Gabriel Rasskin, Gary Wei, Glenn Cameron, Gus Martins, Hadi Hashemi, Hanna Klimczak-Plucińska, Harleen Batra, Harsh Dhand, Ivan Nardini, Jacinda Mein, Jack Zhou, James Svensson, Jeff Stanway, Jetha Chan, Jin Peng Zhou, Joana Carrasqueira, Joana Iljazi, Jocelyn Becker, Joe Fernandez, Joost van Amersfoort, Josh Gordon, Josh Lipschultz, Josh Newlan, Ju yeong Ji, Kareem Mohamed, Kartikeya Badola, Kat Black, Katie Millican, Keelin McDonell, Kelvin Nguyen, Kiranbir Sodhia, Kish Greene, Lars Lowe Sjoesund, Lauren Usui, Laurent Sifre, Lena Heuermann, Leticia Lago, Lilly McNealus, Livio Baldini Soares, Logan Kilpatrick, Lucas Dixon, Luciano Martins, Machel Reid, Manvinder Singh, Mark Iverson, Martin Görner, Mat Velloso, Mateo Wirth, Matt Davidow, Matt Miller, Matthew Rahtz, Matthew Watson, Meg Risdal, Mehran Kazemi, Michael Moynihan, Ming Zhang, Minsuk Kahng, Minwoo Park, Mofi Rahman, Mohit Khatwani, Natalie Dao, Nenshad Bardoliwalla, Nesh Devanathan, Neta Dumai, Nilay Chauhan, Oscar Wahltinez, Pankil Botarda, Parker Barnes, Paul Barham, Paul Michel, Pengchong Jin, Petko Georgiev, Phil Culliton, Pradeep Kuppala, Ramona Comanescu, Ramona Merhej, Reena Jana, Reza Ardeshtir Rokni, Rishabh Agarwal, Ryan Mullins, Samaneh Saadat, Sara Mc Carthy, Sarah Cogan, Sarah Perrin, Sébastien M. R. Arnold, Sebastian Krause, Shengyang Dai, Shruti Garg, Shruti Sheth, Sue Ronstrom, Susan Chan, Timothy Jordan, Ting Yu, Tom Eccles, Tom Hennigan, Tomas Kocisky, Tulsee Doshi, Vihan Jain, Vikas Yadav, Vilobh Meshram, Vishal Dharmadhikari, Warren Barkley, Wei Wei, Wenming Ye, Woohyun Han, Woosuk Kwon, Xiang Xu, Zhe Shen, Zhitao Gong, Zichuan Wei, Victor Cotruta, Phoebe Kirk, Anand Rao, Minh Giang, Ludovic Peran, Tris Warkentin, Eli Collins, Joelle Barral, Zoubin Ghahramani, Raia Hadsell, D. Sculley, Jeanine Banks, Anca Dragan, Slav Petrov, Oriol Vinyals, Jeff Dean, Demis Hassabis, Koray Kavukcuoglu, Clement Farabet, Elena Buchatskaya, Sebastian Borgeaud, Noah Fiedel, Armand Joulin, Kathleen Kenealy, Robert Dadashi, and Alek Andreev. Gemma 2: Improving open language models at a practical size, 2024.
- [Thomson, 1985] Judith Jarvis Thomson. The trolley problem. *The Yale Law Journal*, 94(6):1395–1415, 1985.
- [Zheng *et al.*, 2024] Chujie Zheng, Hao Zhou, Fandong Meng, Jie Zhou, and Minlie Huang. Large language models are not robust multiple choice selectors. In *International Conference on Learning Representations (ICLR)*, 2024.