

Generating Compromises Between Two Points of View

Anonymous ACL submission

Abstract

Large Language Models (LLMs) excel academically but struggle with social intelligence tasks. We present a framework for generating empathically neutral compromises between opposing viewpoints to address this limitation. Using a dataset of 2,400 contrasting views from human participants, we develop compromise generation methods. To overcome data collection constraints, we utilize prompt engineering with Claude LLM to generate compromises, validated through a 50-participant user study. Models trained on this dataset via preference alignment demonstrated effective compromise generation across multiple metrics. This work establishes a scalable approach to enhance LLMs' social intelligence through neutral conflict resolution.

1 Introduction

Recent advances in language understanding have expanded Large Language Model's (LLM) capabilities significantly. These developments (Jiang et al., 2020; Meng et al., 2022) span code comprehension (Nam et al., 2024; Lin et al., 2024; Roziere et al., 2023), problem-solving (Orrù et al., 2023; Kim et al., 2024b), empathetic response generation (Qian et al., 2023; Majumder et al., 2020), and style-controlled text generation (Han et al., 2024; Dathathri et al.; Zhou et al., 2023), each adapting to unique task structures. Since the LLMs are pre-trained on massive amount of text data which make them academically intelligent, they often lack in tasks require social intelligence. The distinction between social and academic intelligence in LLMs is fundamental — social intelligence enables model to navigate interpersonal contexts and emotional cues, e.g., the capability of being more empathetic and finding the common ground of between differences, while academic intelligence facilitates structured information processing and scholarly output. Despite LLMs exhibiting exceptional capa-

bilities in cognitive tasks (Niu et al., 2024), recent attempts (Xu et al., 2024) indicate a low correlation between LLM academic performance and social intelligence metrics. Given the crucial role of social skills across various domains, it is essential to design different scenarios/tasks to evaluate LLMs or incorporate the skill into an LLM. Recent findings, including behavioral intelligence from a first-person perspective (Hou et al., 2024), suggest LLMs exhibit significant performance gaps particularly in complex, interactive, and goal-driven social contexts.

We develop a task focused on generating compromises from a pair of contrasting views (Figure 1). Views are limited to scenarios where places are evaluated as (safe ($view_A$), less safe ($view_B$)) or (welcoming ($view_A$), excluded ($view_B$)). In general, ($view_A, view_B$) consists of reason and suggestion to improve a place based on the evaluation. For instance, for a pair of (safe, less safe) views, $view_A$ describes the features that make a location safe and suggests potential improvements to enhance its safety, while $view_B$ identifies the hazards that make a location less safe and proposes measures to address these safety concerns. A suitable compromise should satisfy these two constraints: (i) it should consider both the suggestions while generating candidate compromises, and (ii) candidate compromise needs to be empathically neutral. While LLMs can produce empathetic responses (Welivita et al., 2021; Sabour et al., 2022; Majumder et al., 2020), their ability to consistently generate neutral compromises are less explored. Such neutrality is crucial in diverse applications including conflict resolution, negotiation support, and collaborative decision-making.

To facilitate robust model training, we use a corpus of 2400 contrasting viewpoints expressed by human participants (section 3). Given the task's complexity and the diversity of viewpoints involved, it is difficult for individuals to neutrally

View A: I am writing about this place: The neighborhood where I live.. I feel safe here because For a big city, it feels like a community. There are many close knit ties to neighbors. It feels like a true collaboration of peoples.. Some ways this place could be modified to be safer are : I would show them the improvements to safety that have been made over the years. There could be more streetlights. I would show them the community.

View B: I am writing about this place: Neighborhood park I feel safety could be improved here This is a small neighborhood park local to where I live. Later in the evening the park feels more less safe than during the day - drug use is becoming more common and loud and intimidating individuals are becoming the norm here. Some ways this place could be modified to be safer are More security, cameras, banning individuals who deal/use drugs.



Implement a "Park Steward" program where rotating pairs of neighborhood volunteers host scheduled evening activities (like group exercises or children's games), making the space actively used by families while naturally discouraging inappropriate behavior.



Install motion-activated lighting combined with discreet security cameras to maintain visibility while being cost-effective.



Address View A's emphasis on community ties and View B's concerns about evening safety and inappropriate behavior.



Creates a solution that could unite rather than divide the community.

Overlooks View A's core values.

Discourage the natural community interaction that View A values.

Figure 1: Example of a data point. ViewA and ViewB are collected from human participants, compromises (preferably, candidate compromises) are synthetically generated using prompt engineering that satisfy the criteria of balanced view and empathic neutrality.

mediate between opposing views without introducing their own unconscious biases (i.e., confirmation bias, social desirability bias etc). Instead, we opted for synthetically generated candidate compromises, which can be aligned with specific needs (Long et al., 2024). Careful implementation of prompt-based controls over the generation process allows us to create more balanced representations of diverse compromises based on the original views. To evaluate empathic neutrality, we train a separate similarity model based on empathic similarity (Shen, 2023) (section 3.2). It generates an empathic similarity rating between the views and candidate compromise. candidate compromise that achieve balanced similarity ratings with respect to both the views are considered to demonstrate empathic neutrality.

Prior work (Liu et al., 2023; Huang et al., 2024) suggests that model alignment can effectively steer LLM behavior toward new tasks while preserving broad generalization capabilities. We leverage this insight for our compromise generation task, as the pretrained knowledge in existing open-source LLMs (Dubey et al., 2024; Jiang et al., 2023) provides a valuable foundation. Through alignment methods (Section 4), we can guide these models to generate compromises while maintaining adherence to our required constraints. We compare two different preference alignment methodologies (section 4) for compromise generation: (i) Noise Contrastive Estimation (NCE) (Gutmann and Hyvärinen, 2010) based and (ii) a task-loss based that we propose. We find that finetuning LLMs on our dataset improves the generation performance on an automatic evaluation metric. This fine-tuned model is used for alignment training for further

improvement. We observe our alignment methods effectively preserve pre-trained knowledge while acquiring this skill.

While quantitative comparison provides valuable insights for model evaluation, assessing how effective our candidate compromises are in practice is critical. Since, to the best of our knowledge, there is no existing benchmark/dataset that addresses the compromise generation task, we conducted a study with 50 human subjects to evaluate compromises generated by different methods. The study shows that providing an LLM with an external measure of empathic similarity to proposed compromises can produce much more acceptable compromises than an LLM generating compromises alone (section 11.1 for qualitative study). Overall, our framework demonstrates that despite data scarcity in social intelligence tasks, carefully designed scenarios coupled with prompt engineering for synthetic data generation can unlock new capabilities in LLMs, revealing valuable characteristics that can inform future research directions in enhancing language models' social capabilities.

2 Related work

Social intelligence, fundamental to human cognition, remains a challenge for Large Language Models (LLMs) despite their advanced text generation capabilities (Sterelny, 2007; Gallegos et al., 2024; Dautenhahn, 1995). While LLMs excel in academic tasks (GPT-4 achieving 92.1% on MATH), they show significant limitations in social intelligence, scoring only 54.4% on SESI benchmark (Xu et al., 2024). Similar limitations appear in interactive gaming contexts, where LLMs perform approximately 20% below human baseline in theory

of mind tasks (Liu et al., 2024). Studies from First-person perspective confirm that despite possessing basic theory of mind capabilities, LLMs show considerable limitations in managing complex social interactions compared to human performance (Hou et al., 2024). Current evaluation methods include traditional psychological assessments like ToMi (Le et al., 2019) and specialized datasets such as SocialIQA (Sap et al., 2019), SockKET (Choi et al., 2023), and SECEU (Wang et al., 2023). However, newer benchmarks like SOTOPIA (Zhou et al.) and EmoBench (Sabour et al., 2024) face limitations in reflecting real-world interactions.

Recent work demonstrates that fine-tuned language models using supervised learning and re-ranking of generated statements can effectively generate consensus statements maximizing group approval (Bakker et al., 2022). Mediation of group discussions by incorporating both majority and minority perspectives when generating compromises was explored in (Tessler et al., 2024). In contrast to (Bakker et al., 2022) which used debate topics with two sides, in our work the prompts were more open-ended in that they asked how a place should be modified. In addition, while (Bakker et al., 2022) propose the use of a reward model for predicting which candidate statements generated by an LLM a human prefers, we focus on generating compromises for which people with different views have similar empathy. While (Tessler et al., 2024) used humans in selection of a joint group statement of common ground, we propose to generate multiple suggested areas of compromise.

Synthetic datasets offer a promising solution for enhancing LLMs’ social capabilities (Ghannadian et al., 2024; Hassan et al., 2024; Balog et al., 2024; Gabriel et al., 2024). These datasets leverage LLMs’ generative abilities (Dankar and Ibrahim, 2021) while ensuring diverse representation (Yamagishi and Nakamura, 2024). Advanced prompt engineering techniques, particularly using GPT-4 or Claude (Achiam et al., 2023; Anthropic, 2023), have improved dataset quality (Hikov and Murphy, 2024; Shi et al., 2023). Chain of Thought (COT) approaches (Wei et al., 2022; Feng et al., 2024; Chu et al., 2023) show enhanced reasoning capabilities (Shao et al., 2023), with significant improvements in specific tasks (Nong et al., 2024). Structured Chain-of-Thought (SCOT) (Sultan et al., 2024) further improves accuracy in document-grounded QA conversations. Recent development suggest using parameterized soft prompts instead of traditional

hard-prompting (DeSalvo et al., 2024).

Model alignment with human preferences (Liu et al., 2023; Kim et al., 2024a; Wang et al., 2024b) has evolved from Reinforcement Learning from Human Feedback (RLHF) (Kaufmann et al., 2023) to more efficient approaches. While RLHF allows models trained on general text data to align with complex human value (Song et al., 2024; Sun et al., 2023; Yuan et al., 2024), it involves training a "reward model" with direct human feedback and then using it to optimize the AI agent’s performance through reinforcement learning. However this type of training (Proximal policy optimization (Schulman et al., 2017)) turns out to be complex and resource intensive (Liu et al., 2023; Wang et al., 2024a; Hong et al., 2023). Direct Preference Optimization (DPO) (Rafailov et al., 2024) eliminates the need for separate reward models, it uses a negative log-likelihood loss function to increase the relative probability of preferred responses over dis preferred ones, making it suitable for stable training. Further advancements like reward model distillation (Fisch et al., 2024) and leveraging KL regularization (Kullback and Leibler, 1951) to avoid overfitting (Kullback and Leibler, 1951; Wesego and Rooshenas, 2024; Azar et al., 2024) improve training stability. DPOP (Pal et al., 2024) introduces an additional penalty term to the loss function, ensuring that the likelihood of preferred completions remains high relative to the reference model. Noise Contrastive Estimation (NCE) (Gutmann and Hyvärinen, 2010) based algorithms InfoNCA and NCA (Chen et al., 2024) extends DPO to handle multi-response reward datasets and focuses on optimizing absolute likelihoods to prevent issues like declining preferred response probabilities respectively. Additional approaches like TODO (Guo et al., 2024) extend beyond binary preference models by introducing a ternary ranking system and RPO (Yin et al., 2024) uses contrastive weighting to differentiate between preferred responses from both identical and related prompts, offering more nuanced preference optimization methods.

Based on empathy towards the candidate compromises, we explore the use of prompt engineering informed by the empathic similarity of two viewpoints towards a candidate compromise for generation of synthetic data. We then align the model using two different strategies: NCE-based (Gutmann and Hyvärinen, 2010) and task metric-based objective.

3 Compromise Generation

This section describes the dataset used for our experiments and the generation of candidate compromise sentences through diverse prompt engineering techniques.

3.1 Dataset with contrasting views

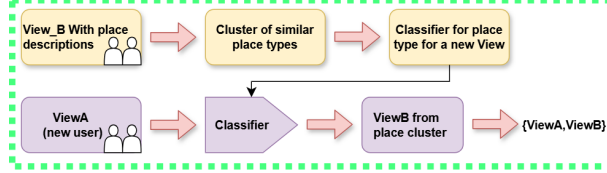


Figure 2: Overview of Dataset collection process. $View_A, View_B$ are positive and negative views.

We used a dataset that will be published with a forthcoming paper under review. We summarize the collection process here. Data collection occurred in two stages (Figure 2). In Stage 1 (Upper row), participants recruited from dscout¹ provides views ($view_B$) about places where they felt unsafe/excluded. For each place, they explained their feelings and suggested modifications to improve safety or inclusivity. Place descriptions were grouped into 14 clusters using agglomerative clustering each for unsafe/excluding places. A classifier was developed to match new views to these clusters.

In Stage 2 (bottom row), new participants from Prolific² wrote views ($view_A$) about places where they felt safe/welcome. The classifier matched each $view_A$ to a $view_B$ from the place-type cluster corresponding to $view_A$. Participants associated with $view_A$ also rated their empathy for one, three, or five $view_B$ s (this data is further used to train the similarity model described in section 3.2), along with some demographic information.

3.2 Similarity model

We follow the work of Shen (2023) on modeling empathetic similarity. The task aims to compute a similarity score $sim(f_\theta(s1), f_\theta(s2))$ between story pairs ($s1, s2$), where the score should be higher for stories with similar empathetic content. We used the e5-large (Wang et al., 2022) model as the encoder instead of sentence BERT (Reimers, 2019) in the original work. During inference, the similarity model predicts the empathetic similarity

between the generated compromise and each of the views.

3.3 Generating candidate compromises

The goal is to generate compromises for each given pair of positive and negative views for a place. Each view at least consists of (i) place (ii) reason and (iii) suggestions of improvement. In order to use all this information as a context for feeding into an LLM, we use a fixed prompt template similar to the Alpaca prompt (Taori et al., 2023).

Generating desirable candidate compromises is challenging due to inherent human biases - for instance, confirmation bias leads individuals to favor information that supports their existing views, while social desirability bias influences them to present socially acceptable responses rather than their true perspectives, making truly neutral compromise difficult to achieve. Furthermore, we observe that relying on a single compromise as the target may limit the model’s generalization capabilities (Gong et al., 2019). Therefore, we generated multiple candidate compromises for each view pair to enhance the model’s ability to learn diverse resolution patterns. Recent work suggests that proprietary LLMs excel at synthetic dataset generation (Cascante-Bonilla et al., 2023; Lei et al., 2023; Ghanadian et al., 2024; Shi et al., 2023) with the help of sophisticated prompting strategies based on the task scenario, providing an alternative to human data collection. However, we observe that single prompts are not adequate for our task, often focusing on one view only, leading us to explore advanced prompt engineering techniques.

Chain of Thought (CoT): As shown in Figure 3, we initially used CoT (Chu et al., 2024; Wei et al., 2022) to identify suggestion similarities and derive compromises. However, this approach proved ineffective, as it either focused too heavily on single suggestion or exceeded their scope, resulting in less suitable compromises.

COT+LLM: Since recent work (Madaan et al., 2024) suggests LLMs can refine outputs using self-generated feedback to better align with human refinement techniques, We enhance the prompt with LLM-based self-evaluation scores (Figure 3) in the CoT+LLM score approach. The self-evaluation scores are empathetic similarity ratings between the generated compromise and views obtained from the LLM itself. Ideally, the similarity ratings should be high between each view and generated compromise and the difference should be zero. We use

¹<https://dscout.com/>

²<https://www.prolific.com/>

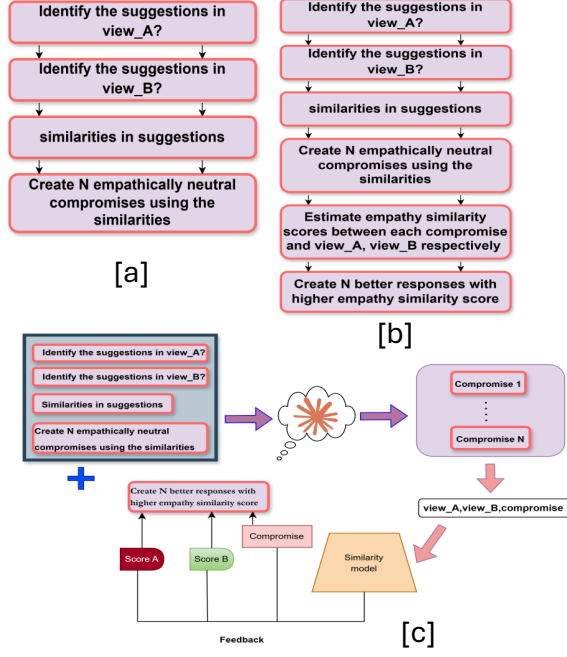


Figure 3: Prompt engineering strategies for candidate compromise generation (a) basic COT approach (b) CoT with LLM self-evaluation scores (c) Extending (b) with similarity model and feedback

these scores to improve the response further. However, we observe that while self-evaluation scores improved, compromise quality remains almost stagnant, suggesting that self-evaluation scores are not a reliable indicator of improving response quality.

COT+Feedback: Instead of relying on the self-evaluation scores, we train a separate similarity model. (Refer to section 3.2) and use the predicted empathy similarity ratings and generated responses (Figure 3) as feedback to improve the responses iteratively. We observe that for each iteration, both response quality and empathy similarity ratings improve. Since the similarity model is trained to mimic human empathic similarity, LLMs may find this external feedback more valuable than self-evaluation.

For each prompting strategy, we generated four candidate compromises per view pair. Given our investigation of four distinct prompting strategies (single prompt, CoT, CoT+LLM score, and CoT+Feedback), we collected a total of $4 \times 4 = 16$ candidate compromises for each view pair across all strategies. From this set, we selected the top four responses as final candidate compromises based on our neutrality criteria. To evaluate neutrality, we use the similarity model that computes the empathic similarity between the generated compro-

mise and both $view_A$ and $view_B$, which we refer to as $score_A$ and $score_B$, respectively. In the ideal scenario of empathic neutrality, the difference between these scores approaches zero: $|score_A - score_B| \rightarrow 0$.

Table 1 shows the composition of the final candidate compromises across all the strategies. COT+Feedback contributed the highest amount of accepted compromises, while single prompt yielded the lowest, demonstrating the effectiveness of our approach.

4 Alignment

Language models can be aligned with preferences through either explicit scalar rewards (i.e., GPT-4/claude ratings) or implicit rewards learned from preference data, where the reward difference indicates preference probability. Direct Preference Optimization (DPO) (Rafailov et al., 2024) simplifies training by using data likelihood ratios between two models to jointly learn rewards and optimize language model behavior. However, DPO has limitations. It only works with pairwise preference data and shows unexpected behavior where the best response likelihood decreases during training despite increase in the relative likelihood between two responses-following the training objective. We use the Noise Contrastive Estimation (NCE) (Chen et al., 2024) based and task loss based alignment method. Recent work (Gutmann and Hyvärinen, 2012) shows NCE based alignment methods optimize for the absolute likelihood of the data, rather than focusing on the relative likelihood between different responses.

Our alignment ensures to learn a policy that enhances the capability of open source pretrained LLMs towards compromise generation task. Following the DPO method (Rafailov et al., 2024), our initial approach includes supervised finetuning of the base model (pretrained LLM) before applying an alignment method.

4.1 NCE based alignment

After finetuning the base model, we align our model with the candidate compromises. Since we collect multiple responses for a fixed pair of views, for each epoch we randomly sample from this candidate compromise pool to train the model. However, we don’t find it beneficial to train the model with all layers, rather we freeze all layers except the last three, hence during alignment the last

Response Type	Single Prompt	COT	COT+LLM	COT+Feedback
Welcome	0.75%	24.70%	28.45%	46.10%
Safe	0.83%	22.09%	26.65%	50.43%

Table 1: Distribution of final responses across all the prompting strategies.

three layers of the model gets trained. Unlike DPO (Rafailov et al., 2024), which utilizes a reference model to serve as guidance for the policy model during training, we used a single fine-tuned base model unfreezing the last three layers. As shown in Equation 1, we use the sum of log probabilities of the generated tokens as the reward. s_{target} and s_{hypo} is defined as the sum of log probabilities of the candidate compromise tokens and fine-tuned base model’s output tokens, respectively. The objective increases the candidate compromise likelihood while decreasing the fine-tuned base model’s output likelihood.

$$\mathcal{L}_{nce} = -\log\left(\frac{1}{1 + e^{-s_{target}}}\right) - \log\left(\frac{1}{1 + e^{s_{type}}}\right) \quad (1)$$

4.2 Task-based loss alignment

As Discussed in 4.1, Although it increases the likelihood of the candidate compromise samples, early work suggests (Bhattacharyya et al., 2020) there is a discrepancy between probability values and the task metric. Samples with higher probability does not guarantee higher score in automatic evaluation metric (i.e., BLEU (Papineni et al., 2002), ROUGE (Lin, 2004) etc.). To ensure higher likelihood values correspond to samples with higher evaluation scores, we jointly use the task metric (in our case we consider ROUGE) with log probability (Eq. 2).

$$\mathcal{L}_{tbl} = \max(0, s_{target} - s_{hypo} + |target_{rouge} - hypo_{rouge}| \cdot w_{margin}) \quad (2)$$

$target_{rouge}$ refers to the ROUGE score of the candidate compromise with respect to the reference text (which is candidate compromise), while $hypo_{rouge}$ denotes the ROUGE score of the compromise generated by the fine-tuned base model with respect to the candidate compromise. The weight margin value, w_{margin} captures the impact of the difference in task metric during training.

5 Experimental setup

Dataset: For training the similarity model, we used 1000 of the data points along with their empathy ratings from participants (please refer to section

3.1). In addition, we merged our dataset with the data used in (Shen, 2023). Hence, the total data points with human annotated empathy ratings is 2500. We split this 75/5/20 as train, dev and test data, respectively.

For the alignment experiemnts, we use a corpus containing 2400 contrasting view pairs (data points) divided evenly into the two categories: 1200 pairs for safe versus unsafe views and 1200 pairs for welcome versus excluded views. The dataset is split into training, development, and test subsets with 75%, 5%, and 20% of the data, respectively.

Setup: For all the experiments, we use Llama-3.1 8B and mistral-7B instruct (Dubey et al., 2024; Jiang et al., 2023) as base models. We initially fine-tune the base model for one epoch. We implemented a linear warm-up and cosine decay scheduler to dynamically adjust the learning rate with an initial learning rate of $3e^{-5}$. For optimization, we used the Adam optimizer with beta parameters (β_1, β_2) of 0.90 and 0.99, respectively. For task-based loss alignment, a weight margin value of 10 yields suitable result. We train for 8 and 12 epochs for the NCE-based alignment and task-based loss alignment respectively.

Evaluation metric: For the similarity model experiment, we use the Spearman correlation as the evaluation metric. When computing correlations, the predicted empathic similarity value is compared to the labeled value where the labeled value was normalized to range from 0 to 1. For generated compromise evaluation, we use ROUGE score (F-measures) (Lin, 2004). Since we have multiple candidate compromises for fixed view pairs, we first perform pairwise calculations, where the ROUGE score is computed between the predicted compromise and each reference compromise individually. The final reported ROUGE score (Table 2) is the maximum score obtained from all pairwise calculation.

6 Evaluation

Our evaluation examines the quality and neutrality of generated compromises, the impact of demographic information, preservation of pretrained knowledge, and human agreement with the gener-

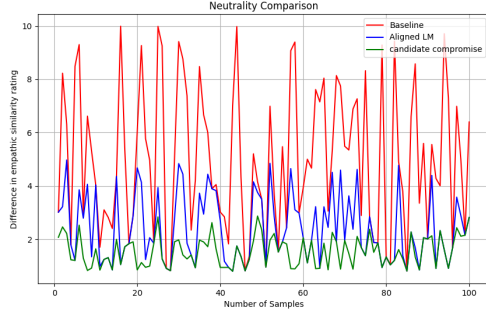


Figure 4: Difference between $score_A$ and $score_B$. Low difference indicates better neutrality.

ated compromises through a user study.

As shown in Table 2, the pre-trained base models (Llama and Mistral) perform poorly in generating compromises. Instead of relying on single output sample, we also investigate through sampling from the base model (multinomial sampling), generating 5 samples and selecting the best sample based on evaluation metrics. We observe that pretrained base line models do not yield suitable generation. Although finetuning the model for a single epoch improves performance it is still inadequate. We apply the preference alignment algorithm on top of the fine-tuned model, leading to improvement in performance for both models in evaluation metric, demonstrating the effectiveness of alignment training for the task. Please see 11.2 for qualitative analysis.

6.1 Empathic Neutrality

Improvement in the evaluation metric does not guarantee empathic neutrality. We use the similarity model to test the generated compromises for neutrality by measuring the difference of empathic similarity rating between the generated compromise and each of the views ($score_A, score_B$). We select 100 pairs of views sampled randomly from the test data. Figure 4 shows how the difference varies across the data points. The base model exhibits high difference, yielding least neutral generations while the candidate compromise consistently maintains least difference, inclined towards neutrality. Our method effectively bridges the gap between these, yielding better samples with improved neutrality compared to the base model.

6.2 Significance of Demography

Since we also collected demographic information from human subjects, we also investigate whether

LLM can be benefitted by including demographic information in the input context for compromise generations. We evaluated this on two different settings, with and without demographics, independently finetuning and aligning two separate models. For each setting, we select 300 pairs of views from train data and 50 pairs of views from test data. Table 3 suggests that the model does not derive meaningful insights from the demographic data, generating similar compromises.

6.3 Catastrophic Forgetting

To evaluate our alignment approach’s impact on pretrained knowledge retention, we investigate how our method preserves the model’s original capabilities. We use the WikiText dataset (test subset of wikitext-2-v1)³ (Merity et al., 2018), a popular dataset to perform catastrophic forgetting tests (Fawi, 2024). To assess our method’s susceptibility to knowledge forgetting, we calculate average log likelihood⁴ across the dataset. A decline in this number indicates that the model’s mastery of earlier content is eroding. As shown in (Table 4), our method maintains pretrained knowledge more effectively compared to existing methods (i.e., Direct Preference Optimization (DPO) (Rafailov et al., 2024)), which is crucial for generating effective compromises that leverage the model’s existing capabilities.

7 Human Evaluation of Compromise Methods

We asked humans to rate a sample of the compromises generated by the three methods, Single Prompt, Chain-of-Thought, and Feedback. The participants were shown a pair of statements about modifications proposed for a place, one was written by someone who thought positively about the place and the other statement was written by someone who felt negatively about the place. The rater was asked to take the viewpoint of either the people who felt positively or who felt negatively. For each pair of statements, they were asked to rate on a scale of 1-100 five statements from the viewpoint of the person who wrote the statement they were to identify with. One statement was the positive or negative statement that they were not to

³<https://huggingface.co/datasets/Salesforce/wikitext/viewer/wikitext-2-v1/test>

⁴log likelihood for a datapoint is obtained by summing the log probabilities over its tokens

Model	Llama		Mistral	
	ROUGE-1($\uparrow$)	ROUGE-L($\uparrow$)	ROUGE-1($\uparrow$)	ROUGE-L($\uparrow$)
Base Model	0.164	0.107	0.160	0.105
Base Model + Sampling (best of 5)	0.193	0.151	0.191	0.127
Base Model + Finetuning (FT)	0.255	0.186	0.256	0.185
Base Model + FT+ NCE	0.315	0.216	0.318	0.234
Base Model +FT+ Task based Loss	0.321	0.326	0.333	0.254

Table 2: ROUGE score comparison on test data

Experiment	Llama($\uparrow$)	Mistral($\uparrow$)
W Demog.	0.249	0.238
W/O Demog.	0.251	0.237

Table 3: Rouge-L for demography test.
Llama = Llama+FT+NCE, Mistral = Mistral+FT+NCE.

Model	Log Likelihood($\downarrow$)
LLAMA 3.1 8b+DPO	-1.9145
Mistral 7b+DPO	-1.8976
Mistral 7b+NCE	-1.7230
Mistral 7b+Task based loss	-1.7015
LLAMA 3.1 8b+NCE	-1.6234
LLAMA 3.1 8b+Task based loss	-1.6345

Table 4: Catastrophic forgetting test

identify with. The four other statements were generated compromises: one by Single Prompt, one by Chain-of-Thought, and two by the Feedback methods. Each of participant rated the compromises for five pairs of statements. For details about the collection of compromise ratings, see Appendix 11.3

Method	First Preference (%)	Second Preference (%)
View _B	0	1
Single Prompt	5	13
CoT	18	24
Feedback 1	37	33
Feedback 2	40	29

Table 5: Preference distribution for user study (%). For Single Prompt and CoT, one randomly selected compromise was used in evaluation. For Feedback 1 and 2, two randomly selected compromises from the CoT+Feedback method were used. View_B represents the opposing viewpoint to the person doing the rating.

The results of the study (Table 5) reveal a clear hierarchy in the effectiveness of the methods evaluated. As expected, the opposing statement, **View_B**, was never selected as first preference; it was selected as second preference only 1% of the time. The **Single Prompt** method, serving as the baseline, performed the next worst, with only 5% of

participants selecting it as their first preference and 13% as their second preference. This indicates that a simple single-shot approach is insufficient for generating effective compromises. The **Chain-of-Thought (CoT)** method demonstrated notable improvement over the baseline, achieving 18% as the first preference and 24% as the second preference. This highlights the importance of incorporating structured reasoning in the compromise generation process. Finally, the **Feedback-based methods** (Feedback 1 and Feedback 2) achieved the highest performance. Feedback 1 garnered 37% as the first preference and 33% as the second preference, while Feedback 2 achieved 40% as the first preference and 29% as the second preference. These results underscore the critical role of iterative refinement in producing nuanced and satisfactory compromises. In conclusion, the findings validate that feedback-based methods significantly outperform both Chain-of-Thought reasoning and Single Prompt approaches, establishing their superiority for generating balanced compromises.

8 Conclusion and Future Work

Our proposed task framework is an example of the critical importance of crafting unique and targeted scenarios to evaluate large language models (LLMs) and identify potential limitations, particularly in the context of their social intelligence capabilities. Importantly, the scarcity of existing data should not pose a significant challenge, as heuristic-based prompting strategies can be employed to generate synthetic data tailored to specific evaluation needs.

Our alignment model effectively bridging the gap between baseline performance and the desired target outcomes. Furthermore, user studies validate the efficacy of our data generation approach. This validation highlights the robustness and practicality of our framework in advancing the evaluation and refinement of LLMs for socially intelligent applications.

9 Limitations

The COT+feedback approach sought to produce compromise text that is equally acceptable to both viewpoints. In the future, we would also like to include consideration of maximizing the empathy of each viewpoint to the compromise.

Our collection of compromises was constrained to keep the data collection required manageable. The compromises are generated based on two different human viewpoints for a similar type of place. Although many of the comments about a place apply generally to a place type, e.g., trails and fencing in a park, in a real scenario, the expressed viewpoints are about the same type of place, rather than exactly the same. We chose this approach because collecting compromise data for the same place would be a much larger undertaking. The prompts were based on feeling safe/unsafe or welcome/excluded. Generalization to other prompts was not done because this would have also increased the required size of the negative viewpoints. These generalizations are left for future work.

In addition to the collected viewpoints, our work made use of demographic information in all experiments. Our Review Board will not allow the release of the dataset with demographic information because it is PII.

The Human Study comparing different compromise generation texts presented only a sample of 5 of the generated texts per viewpoint pair, rather than all 20, because humans would have difficulty rating 20 items reliably.

10 Ethical Considerations

Our goal in this work was to generate compromises that are equally acceptable to the two parties with different viewpoints. Similar to the general dangers of some people trusting all results generated by an LLM, the generated compromises should be viewed as suggestions for discussion areas, rather than a compromise that each person should agree to.

References

- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- AI Anthropic. 2023. Introducing claude.

- Mohammad Gheshlaghi Azar, Zhaohan Daniel Guo, Bilal Piot, Remi Munos, Mark Rowland, Michal Valko, and Daniele Calandriello. 2024. A general theoretical paradigm to understand learning from human preferences. In *International Conference on Artificial Intelligence and Statistics*, pages 4447–4455. PMLR.
- Michiel Bakker, Martin Chadwick, Hannah Sheahan, Michael Tessler, Lucy Campbell-Gillingham, Jan Balaguer, Nat McAleese, Amelia Glaese, John Aslanides, Matt Botvinick, et al. 2022. Fine-tuning language models to find agreement among humans with diverse preferences. *Advances in Neural Information Processing Systems*, 35:38176–38189.
- Krisztian Balog, John Palowitch, Barbara Ikica, Filip Radlinski, Hamidreza Alvari, and Mehdi Manshadi. 2024. Towards realistic synthetic user-generated content: A scaffolding approach to generating online discussions. *CoRR*.
- Sumanta Bhattacharyya, Amirmohammad Rooshenas, Subhajit Naskar, Simeng Sun, Mohit Iyyer, and Andrew McCallum. 2020. Energy-based reranking: Improving neural machine translation using energy-based models. *arXiv preprint arXiv:2009.13267*.
- Paola Cascante-Bonilla, Khaled Shehada, James Seale Smith, Sivan Doveh, Donghyun Kim, Rameswar Panda, Gul Varol, Aude Oliva, Vicente Ordonez, Rogerio Feris, et al. 2023. Going beyond nouns with vision & language models using synthetic data. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 20155–20165.
- Huayu Chen, Guande He, Lifan Yuan, Ganqu Cui, Hang Su, and Jun Zhu. 2024. Noise contrastive alignment of language models with explicit rewards. *arXiv preprint arXiv:2402.05369*.
- Minje Choi, Jiaxin Pei, Sagar Kumar, Chang Shu, and David Jurgens. 2023. Do llms understand social knowledge? evaluating the sociability of large language models with socket benchmark. *arXiv preprint arXiv:2305.14938*.
- Zheng Chu, Jingchang Chen, Qianglong Chen, Weijiang Yu, Tao He, Haotian Wang, Weihua Peng, Ming Liu, Bing Qin, and Ting Liu. 2023. A survey of chain of thought reasoning: Advances, frontiers and future. *arXiv preprint arXiv:2309.15402*.
- Zheng Chu, Jingchang Chen, Qianglong Chen, Weijiang Yu, Tao He, Haotian Wang, Weihua Peng, Ming Liu, Bing Qin, and Ting Liu. 2024. Navigate through enigmatic labyrinth a survey of chain of thought reasoning: Advances, frontiers and future. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1173–1203.
- Fida K Dankar and Mahmoud Ibrahim. 2021. Fake it till you make it: Guidelines for effective synthetic data generation. *Applied Sciences*, 11(5):2158.

745	Sumanth Dathathri, Andrea Madotto, Janice Lan, Jane	Michael Gutmann and Aapo Hyvärinen. 2010. Noise-	797
746	Hung, Eric Frank, Piero Molino, Jason Yosinski, and	contrastive estimation: A new estimation principle	798
747	Rosanne Liu. Plug and play language models: A	for unnormalized statistical models. In <i>Proceedings</i>	799
748	simple approach to controlled text generation. In <i>In-</i>	<i>of the thirteenth international conference on artificial</i>	800
749	<i>ternational Conference on Learning Representations.</i>	<i>intelligence and statistics</i> , pages 297–304. JMLR	801
		Workshop and Conference Proceedings.	802
750	Kerstin Dautenhahn. 1995. Getting to know each	Michael U Gutmann and Aapo Hyvärinen. 2012. Noise-	803
751	other—artificial social intelligence for autonomous	contrastive estimation of unnormalized statistical	804
752	robots. <i>Robotics and autonomous systems</i> , 16(2-	models, with applications to natural image statistics.	805
753	4):333–356.	<i>Journal of machine learning research</i> , 13(2).	806
754	Giulia DeSalvo, Jean-Fracois Kagy, Lazaros Karydas,	Jingxuan Han, Quan Wang, Zikang Guo, Benfeng Xu,	807
755	Afshin Rostamizadeh, and Sanjiv Kumar. 2024. No	Licheng Zhang, and Zhendong Mao. 2024. Disentan-	808
756	more hard prompts: Softsr prompting for synthetic	gled learning with synthetic parallel data for text style	809
757	data generation. <i>arXiv preprint arXiv:2410.16534</i> .	transfer. In <i>Proceedings of the 62nd Annual Meet-</i>	810
		<i>ing of the Association for Computational Linguistics</i>	811
758	Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey,	(<i>Volume 1: Long Papers</i>), pages 15187–15201.	812
759	Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman,		
760	Akhil Mathur, Alan Schelten, Amy Yang, Angela	Abdelrahman A Hassan, Radwa J Hanafy, and Mo-	813
761	Fan, et al. 2024. The llama 3 herd of models. <i>arXiv</i>	hammed E Fouda. 2024. Automated multi-label an-	814
762	<i>preprint arXiv:2407.21783</i> .	notation for mental health illnesses using large lan-	815
		guage models. <i>arXiv preprint arXiv:2412.03796</i> .	816
763	Muhammad Fawi. 2024. Curlora: Stable llm contin-	Asen Hikov and Laura Murphy. 2024. Information	817
764	ual fine-tuning and catastrophic forgetting mitigation.	retrieval from textual data: Harnessing large language	818
765	<i>arXiv preprint arXiv:2408.14572</i> .	models, retrieval augmented generation and prompt	819
		engineering. <i>Journal of AI, Robotics & Workplace</i>	820
766	Guhao Feng, Bohang Zhang, Yuntian Gu, Haotian Ye,	<i>Automation</i> , 3(2):142–150.	821
767	Di He, and Liwei Wang. 2024. Towards revealing the		
768	mystery behind chain of thought: a theoretical per-	Jixiang Hong, Quan Tu, Changyu Chen, Xing Gao,	822
769	spective. <i>Advances in Neural Information Processing</i>	Ji Zhang, and Rui Yan. 2023. Cyclealign: Iter-	823
770	<i>Systems</i> , 36.	ative distillation from black-box llm to white-box	824
		models for better human alignment. <i>arXiv preprint</i>	825
771	Adam Fisch, Jacob Eisenstein, Vicky Zayats, Alekh	<i>arXiv:2310.16271</i> .	826
772	Agarwal, Ahmad Beirami, Chirag Nagpal, Pete Shaw,		
773	and Jonathan Berant. 2024. Robust preference op-	Guiyang Hou, Wenqi Zhang, Yongliang Shen, Zeqi Tan,	827
774	timization through reward model distillation. <i>arXiv</i>	Sihao Shen, and Weiming Lu. 2024. Entering real	828
775	<i>preprint arXiv:2405.19316</i> .	social world! benchmarking the theory of mind and	829
		socialization capabilities of llms from a first-person	830
776	Rodney A Gabriel, Onkar Litake, Sierra Simpson, Brit-	perspective. <i>arXiv preprint arXiv:2410.06195</i> .	831
777	tany N Burton, Ruth S Waterman, and Alvaro A Mac-		
778	cias. 2024. On the development and validation of	Tiansheng Huang, Gautam Bhattacharya, Pratik Joshi,	832
779	large language model-based classifiers for identifying	Josh Kimball, and Ling Liu. 2024. Antidote: Post-	833
780	social determinants of health. <i>Proceedings of the Na-</i>	fine-tuning safety alignment for large language mod-	834
781	<i>tional Academy of Sciences</i> , 121(39):e2320716121.	els against harmful fine-tuning. <i>arXiv preprint</i>	835
		<i>arXiv:2408.09600</i> .	836
782	Isabel O Gallegos, Ryan A Rossi, Joe Barrow,	Albert Q Jiang, Alexandre Sablayrolles, Arthur Men-	837
783	Md Mehrab Tanjim, Sungchul Kim, Franck Dernon-	sch, Chris Bamford, Devendra Singh Chaplot, Diego	838
784	court, Tong Yu, Ruiyi Zhang, and Nesreen K Ahmed.	de las Casas, Florian Bressand, Gianna Lengyel, Guil-	839
785	2024. Bias and fairness in large language models: A	laume Lample, Lucile Saulnier, et al. 2023. Mistral	840
786	survey. <i>Computational Linguistics</i> , pages 1–79.	7b. <i>arXiv preprint arXiv:2310.06825</i> .	841
787	Hamideh Ghanadian, Isar Nejadgholi, and Hussein	Zhengbao Jiang, Frank F Xu, Jun Araki, and Graham	842
788	Al Osman. 2024. Socially aware synthetic data gen-	Neubig. 2020. How can we know what language	843
789	eration for suicidal ideation detection using large	models know? <i>Transactions of the Association for</i>	844
790	language models. <i>IEEE Access</i> .	<i>Computational Linguistics</i> , 8:423–438.	845
791	Zhiqiang Gong, Ping Zhong, and Weidong Hu. 2019.		
792	Diversity in machine learning. <i>Ieee Access</i> , 7:64323–	Timo Kaufmann, Paul Weng, Viktor Bengs, and Eyke	846
793	64350.	Hüllermeier. 2023. A survey of reinforcement	847
		learning from human feedback. <i>arXiv preprint</i>	848
794	Yuxiang Guo, Lu Yin, Bo Jiang, and Jiaqi Zhang. 2024.	<i>arXiv:2312.14925</i> .	849
795	Todo: Enhancing llm alignment with ternary prefer-		
796	ences. <i>arXiv preprint arXiv:2411.02442</i> .		

850	Dongyoung Kim, Kimin Lee, Jinwoo Shin, and Jae-	Yu Meng, Jiaxin Huang, Yu Zhang, and Jiawei Han.	906
851	hyung Kim. 2024a. Aligning large language models	2022. Generating training data with language mod-	907
852	with self-generated preference data. <i>arXiv preprint</i>	els: Towards zero-shot language understanding. <i>Ad-</i>	908
853	<i>arXiv:2406.04412</i> .	<i>Advances in Neural Information Processing Systems</i> ,	909
		35:462–477.	910
854	Geunwoo Kim, Pierre Baldi, and Stephen McAleer.	Stephen Merity, Nitish Shirish Keskar, James Bradbury,	911
855	2024b. Language models can solve computer tasks.	and Richard Socher. 2018. Scalable language model-	912
856	<i>Advances in Neural Information Processing Systems</i> ,	ing: Wikitext-103 on a single gpu in 12 hours. <i>Pro-</i>	913
857	36.	<i>ceedings of the SYSML</i> , 18.	914
858	Solomon Kullback and Richard A Leibler. 1951. On	Daye Nam, Andrew Macvean, Vincent Hellendoorn,	915
859	information and sufficiency. <i>The annals of mathe-</i>	Bogdan Vasilescu, and Brad Myers. 2024. Using an	916
860	<i>matical statistics</i> , 22(1):79–86.	llm to help with code understanding. In <i>Proceedings</i>	917
861	Matthew Le, Y-Lan Boureau, and Maximilian Nickel.	<i>of the IEEE/ACM 46th International Conference on</i>	918
862	2019. Revisiting the evaluation of theory of mind	<i>Software Engineering</i> , pages 1–13.	919
863	through question answering. In <i>Proceedings of the</i>		
864	<i>2019 Conference on Empirical Methods in Natu-</i>	Qian Niu, Junyu Liu, Ziqian Bi, Pohsun Feng, Benji	920
865	<i>ral Language Processing and the 9th International</i>	Peng, Keyu Chen, Ming Li, Lawrence KQ Yan,	921
866	<i>Joint Conference on Natural Language Processing</i>	Yichao Zhang, Caitlyn Heqi Yin, et al. 2024. Large	922
867	<i>(EMNLP-IJCNLP)</i> , pages 5872–5877.	language models and cognitive science: A compre-	923
		hensive review of similarities, differences, and chal-	924
868	Fangyu Lei, Qian Liu, Yiming Huang, Shizhu He, Jun	lenges. <i>arXiv preprint arXiv:2409.02387</i> .	925
869	Zhao, and Kang Liu. 2023. S3eval: A synthetic, scal-		
870	able, systematic evaluation suite for large language	Yu Nong, Mohammed Aldeen, Long Cheng, Hongxin	926
871	models. <i>arXiv preprint arXiv:2310.15147</i> .	Hu, Feng Chen, and Haipeng Cai. 2024. Chain-of-	927
		thought prompting of large language models for dis-	928
872	Chin-Yew Lin. 2004. Rouge: A package for automatic	covering and fixing software vulnerabilities. <i>arXiv</i>	929
873	evaluation of summaries. In <i>Text summarization</i>	<i>preprint arXiv:2402.17230</i> .	930
874	<i>branches out</i> , pages 74–81.		
875	Feng Lin, Dong Jae Kim, et al. 2024. When llm-based	Graziella Orrù, Andrea Piarulli, Ciro Conversano, and	931
876	code generation meets the software development pro-	Angelo Gemignani. 2023. Human-like problem-	932
877	cess. <i>arXiv preprint arXiv:2403.15852</i> .	solving abilities in large language models using chat-	933
		gpt. <i>Frontiers in artificial intelligence</i> , 6:1199350.	934
878	Wenhao Liu, Xiaohua Wang, Muling Wu, Tianlong Li,	Arka Pal, Deep Karkhanis, Samuel Dooley, Man-	935
879	Changze Lv, Zixuan Ling, Jianhao Zhu, Cenyuan	ley Roberts, Siddhartha Naidu, and Colin White.	936
880	Zhang, Xiaoqing Zheng, and Xuanjing Huang. 2023.	2024. Smaug: Fixing failure modes of prefer-	937
881	Aligning large language models with human pref-	ence optimisation with dpo-positive. <i>arXiv preprint</i>	938
882	erences through representation engineering. <i>arXiv</i>	<i>arXiv:2402.13228</i> .	939
883	<i>preprint arXiv:2312.15997</i> .		
884	Ziyi Liu, Abhishek Anand, Pei Zhou, Jen-tse Huang,	Kishore Papineni, Salim Roukos, Todd Ward, and Wei-	940
885	and Jieyu Zhao. 2024. Interintent: Investigating	Jing Zhu. 2002. Bleu: a method for automatic evalu-	941
886	social intelligence of llms via intention understand-	ation of machine translation. In <i>Proceedings of the</i>	942
887	ing in an interactive game context. <i>arXiv preprint</i>	<i>40th annual meeting of the Association for Computa-</i>	943
888	<i>arXiv:2406.12203</i> .	<i>tional Linguistics</i> , pages 311–318.	944
889	Lin Long, Rui Wang, Ruixuan Xiao, Junbo Zhao, Xiao	Yushan Qian, Wei-Nan Zhang, and Ting Liu. 2023.	945
890	Ding, Gang Chen, and Haobo Wang. 2024. On llms-	Harnessing the power of large language models	946
891	driven synthetic data generation, curation, and evalu-	for empathetic response generation: Empirical in-	947
892	ation: A survey. In <i>Findings of the Association for</i>	vestigations and improvements. <i>arXiv preprint</i>	948
893	<i>Computational Linguistics ACL 2024</i> , pages 11065–	<i>arXiv:2310.05140</i> .	949
894	11082.		
895	Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler	Rafael Rafailov, Archit Sharma, Eric Mitchell, Christo-	950
896	Hallinan, Luyu Gao, Sarah Wiegrefe, Uri Alon,	pher D Manning, Stefano Ermon, and Chelsea Finn.	951
897	Nouha Dziri, Shrimai Prabhumoye, Yiming Yang,	2024. Direct preference optimization: Your language	952
898	et al. 2024. Self-refine: Iterative refinement with	model is secretly a reward model. <i>Advances in Neu-</i>	953
899	self-feedback. <i>Advances in Neural Information Pro-</i>	<i>ral Information Processing Systems</i> , 36.	954
900	<i>cessing Systems</i> , 36.		
901	Navonil Majumder, Pengfei Hong, Shanshan Peng,	N Reimers. 2019. Sentence-bert: Sentence embed-	955
902	Jiankun Lu, Deepanway Ghosal, Alexander Gelbukh,	dings using siamese bert-networks. <i>arXiv preprint</i>	956
903	Rada Mihalcea, and Soujanya Poria. 2020. Mime:	<i>arXiv:1908.10084</i> .	957
904	Mimicking emotions for empathetic response genera-	Baptiste Roziere, Jonas Gehring, Fabian Gloeckle, Sten	958
905	tion. <i>arXiv preprint arXiv:2010.01454</i> .	Sootla, Itai Gat, Xiaoqing Ellen Tan, Yossi Adi,	959
		Jingyu Liu, Romain Sauvestre, Tal Remez, et al. 2023.	960

961	Code llama: Open foundation models for code. <i>arXiv preprint arXiv:2308.12950</i> .	An instruction-following llama model. https://github.com/tatsu-lab/stanford_alpaca .	1015
962			1016
963	Sahand Sabour, Siyang Liu, Zheyuan Zhang, June M Liu, Jinfeng Zhou, Alvionna S Sunaryo, Juanzi Li, Tatia MC Lee, Rada Mihalcea, and Minlie Huang. 2024. Emobench: Evaluating the emotional intelligence of large language models. <i>CoRR</i> .	Michael Henry Tessler, Michiel A Bakker, Daniel Jarrett, Hannah Sheahan, Martin J Chadwick, Raphael Koster, Georgina Evans, Lucy Campbell-Gillingham, Tantum Collins, David C Parkes, et al. 2024. Ai can help humans find common ground in democratic deliberation. <i>Science</i> , 386(6719):eadq2852.	1017
964			1018
965			1019
966			1020
967			1021
968	Sahand Sabour, Chujie Zheng, and Minlie Huang. 2022. Cem: Commonsense-aware empathetic response generation. In <i>Proceedings of the AAAI Conference on Artificial Intelligence</i> , volume 36, pages 11229–11237.	Haoxiang Wang, Yong Lin, Wei Xiong, Rui Yang, Shizhe Diao, Shuang Qiu, Han Zhao, and Tong Zhang. 2024a. Arithmetic control of llms for diverse user preferences: Directional preference alignment with multi-objective rewards. <i>arXiv preprint arXiv:2402.18571</i> .	1022
969			1023
970			1024
971			1025
972			1026
973	Maarten Sap, Hannah Rashkin, Derek Chen, Ronan LeBras, and Yejin Choi. 2019. Socialliqa: Commonsense reasoning about social interactions. <i>arXiv preprint arXiv:1904.09728</i> .		1027
974			1028
975			
976			1029
977	John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. 2017. Proximal policy optimization algorithms. <i>arXiv preprint arXiv:1707.06347</i> .	Liang Wang, Nan Yang, Xiaolong Huang, Binxing Jiao, Linjun Yang, Daxin Jiang, Rangan Majumder, and Furu Wei. 2022. Text embeddings by weakly-supervised contrastive pre-training. <i>arXiv preprint arXiv:2212.03533</i> .	1030
978			1031
979			1032
980			1033
981	Zhihong Shao, Yeyun Gong, Yelong Shen, Minlie Huang, Nan Duan, and Weizhu Chen. 2023. Synthetic prompting: Generating chain-of-thought demonstrations for large language models. In <i>International Conference on Machine Learning</i> , pages 30706–30775. PMLR.	Xuena Wang, Xueting Li, Zi Yin, Yue Wu, and Jia Liu. 2023. Emotional intelligence of large language models. <i>Journal of Pacific Rim Psychology</i> , 17:18344909231213958.	1034
982			1035
983			1036
984			1037
985	Jocelyn Shen. 2023. <i>Modeling empathic similarity in personal narratives</i> . Ph.D. thesis, Massachusetts Institute of Technology.	Zhichao Wang, Bin Bi, Shiva Kumar Pentiyala, Kiran Ramnath, Sougata Chaudhuri, Shubham Mehrotra, Xiang-Bo Mao, Sitaram Asur, et al. 2024b. A comprehensive survey of llm alignment techniques: Rlhf, rlaiif, ppo, dpo and more. <i>arXiv preprint arXiv:2407.16216</i> .	1038
986			1039
987			1040
988			1041
989			1042
990	Taiwei Shi, Kai Chen, and Jieyu Zhao. 2023. Safer-instruct: Aligning language models with automated preference data. <i>arXiv preprint arXiv:2311.08685</i> .		1043
991			1044
992			1045
993	Feifan Song, Bowen Yu, Minghao Li, Haiyang Yu, Fei Huang, Yongbin Li, and Houfeng Wang. 2024. Preference ranking optimization for human alignment. In <i>Proceedings of the AAAI Conference on Artificial Intelligence</i> , volume 38, pages 18990–18998.	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. 2022. Chain-of-thought prompting elicits reasoning in large language models. <i>Advances in neural information processing systems</i> , 35:24824–24837.	1046
994			1047
995			1048
996			
997			1049
998	Kim Sterelny. 2007. Social intelligence, human intelligence and niche construction. <i>Philosophical Transactions of the Royal Society B: Biological Sciences</i> , 362(1480):719–730.	Anuradha Welivita, Yubo Xie, and Pearl Pu. 2021. A large-scale dataset for empathetic response generation. In <i>Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing</i> , pages 1251–1264.	1050
999			1051
1000			1052
1001			1053
1002	Md Arafat Sultan, Jatin Ganhotra, and Ramón Fernández Astudillo. 2024. Structured chain-of-thought prompting for few-shot generation of content-grounded qa conversations. <i>arXiv preprint arXiv:2402.11770</i> .	Daniel Wesego and Pedram Rooshenas. 2024. Score-based multimodal autoencoder . <i>Transactions on Machine Learning Research</i> .	1054
1003			1055
1004			1056
1005			
1006			1057
1007	Zhiqing Sun, Sheng Shen, Shengcao Cao, Haotian Liu, Chunyuan Li, Yikang Shen, Chuang Gan, Liang-Yan Gui, Yu-Xiong Wang, Yiming Yang, et al. 2023. Aligning large multimodal models with factually augmented rlhf. <i>arXiv preprint arXiv:2309.14525</i> .	Ruoxi Xu, Hongyu Lin, Xianpei Han, Le Sun, and Yingfei Sun. 2024. Academically intelligent llms are not necessarily socially intelligent. <i>arXiv preprint arXiv:2403.06591</i> .	1058
1008			1059
1009			1060
1010			
1011			1061
1012	Rohan Taori, Ishaan Gulrajani, Tianyi Zhang, Yann Dubois, Xuechen Li, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. 2023. Stanford alpaca:	Yosuke Yamagishi and Yuta Nakamura. 2024. Utrad-nlp at# smm4h 2024: Why llm-generated texts fail to improve text classification models. In <i>Proceedings of The 9th Social Media Mining for Health Research and Applications (SMM4H 2024) Workshop and Shared Tasks</i> , pages 42–47.	1062
1013			1063
1014			1064
			1065
			1066
			1067
			1068
			1069

1070 through contrasting responses across identical and
1071 diverse prompts. *arXiv preprint arXiv:2402.10958*.

1072 Hongyi Yuan, Zheng Yuan, Chuanqi Tan, Wei Wang,
1073 Songfang Huang, and Fei Huang. 2024. Rrhf: Rank
1074 responses to align language models with human feed-
1075 back. *Advances in Neural Information Processing*
1076 *Systems*, 36.

1077 Wangchunshu Zhou, Yuchen Eleanor Jiang, Ethan
1078 Wilcox, Ryan Cotterell, and Mrinmaya Sachan. 2023.
1079 Controlled text generation with natural language in-
1080 structions. In *International Conference on Machine*
1081 *Learning*, pages 42602–42613. PMLR.

1082 Xuhui Zhou, Hao Zhu, Leena Mathur, Ruohong Zhang,
1083 Haoifei Yu, Zhengyang Qi, Louis-Philippe Morency,
1084 Yonatan Bisk, Daniel Fried, Graham Neubig, et al.
1085 Sotopia: Interactive evaluation for social intelligence
1086 in language agents. In *The Twelfth International*
1087 *Conference on Learning Representations*.

11 Appendix

1088

11.1 Qualitative analysis of candidate
compromise generation across different
prompting strategies

1089
1090
1091

We provide a detailed qualitative example of target
compromise generation across different prompting
strategies, The first example provides a comprehen-
sive evaluation, while the second example features
a brief assessment.

1092
1093
1094
1095
1096

Example 1

Positive view

I am writing about this place: A nearby park. I feel safe here because I feel safe here when others are around. There’s a good sense of community. Some ways this place could be modified to be safer are: There’s no fences, gates, no visitor check, and it’s extremely open. This is good and bad.

Negative view

I am writing about this place: Our local park. I feel safety could be improved here. It’s a wonderful park, but people do not obey leash laws. I often see people letting their dogs off leash, making other dog walkers and people with children uncomfortable. Some ways this place could be modified to be safer are I would like to see stricter leash laws. Fines for rule breakers.

1097

Basic Single Prompt

You are an intelligent AI assistant!

I need you to generate a third person response strictly based on two contrasting views called positive story and negative story.

The positive and negative story should be equally empathetic towards the response. The response should be a specific suggestion. It should be a compromise between the positive and negative stories based on the context of both stories.

Please generate 4 responses with a fixed format. Try to be as specific and short instead of being comprehensive.

Please provide your response in the following format:
Positive view: {}
Negative view: {}
Response 1: [Insert response 1 here]
Response 2: [Insert response 2 here]
Response 3: [Insert response 3 here]
Response 4: [Insert response 4 here]

Results

- Response 1: The park management could consider implementing designated off-leash areas within the park, allowing dogs to roam freely in specific zones while maintaining strict leash laws in other areas to ensure the comfort and safety of all visitors.
- Response 2: Increasing the presence of park rangers or volunteers to monitor and enforce leash laws could help maintain a safer environment for everyone, while also providing an opportunity to educate dog owners about the importance of responsible pet ownership.
- Response 3: The park could install clear signage outlining the leash laws and the consequences for violating them, such as fines or temporary bans from the park, to encourage compliance and create a safer atmosphere for all visitors.
- Response 4: Organizing community events

that promote responsible dog ownership and provide opportunities for dogs and their owners to socialize in a controlled environment could help foster a sense of community and encourage adherence to park rules and regulations.

Analysis

- Human observation: The basic single prompt strategy shows clear limitations in balanced perspective handling. Claude’s responses predominantly focus on addressing the negative story’s concerns about dogs and leash laws, while largely overlooking the positive story’s suggestions about general safety and community aspects. This demonstrates a bias in compromise generation toward more concrete, specific issues (dog-related) over broader community safety concerns.

Chain of Thought(CoT)

Results

- Response 1: Install fences and gates around designated off-leash areas for dogs, ensuring they are separate from other park visitors.
- Response 2: Implement a visitor check system that includes registration for dog owners, ensuring they are aware of and agree to follow leash laws.
- Response 3: Increase park staff presence to enforce leash laws and issue warnings or fines to rule breakers, promoting a safer environment for all visitors.
- Response 4: Add clear signage throughout the park to remind visitors of leash laws and the importance of maintaining a safe and comfortable environment for everyone.

Analysis

- Human Observation: The Chain of Thought approach shows improvement in addressing multiple aspects of safety, but still demonstrates some imbalance. While responses successfully incorporate the positive story’s suggestions about physical safety measures, the solutions tend to be more rigid and enforcement-focused.

1099

1100

1101

1102

1103

1104

1105

1106

1107

1108

1109

1110

1111

1112

1113

1114

1115

1116

1117

1118

1119

1120

1121

1122

1123

1124

1125

1126

1127

1128

1129

1130

1131

1132

1133

1134

1135

1136

1137

1138

1139

1140

1141

1142

1143

1144

1145

1146

1147

1148

1149

1150

1151

1152

1153

1154

1155

1156

1157

1158

1159

1160

1161

1162

1163

CoT+Feedback

Results

- Response 1: Create designated off-leash areas within the park that are securely fenced and have separate entrances. These areas should be located away from the main paths and gathering spaces to minimize potential conflicts between dogs and other park visitors. Clearly mark the off-leash areas and provide amenities such as water fountains and waste disposal stations.
- Response 2: Implement a graduated approach to park safety, starting with clear signage outlining park rules, leash laws, and consequences for violations. Provide educational materials on responsible pet ownership and the importance of maintaining a safe environment for all visitors. As a second step, strategically install fences and gates in high-traffic areas to improve safety while preserving the park's openness. Regularly assess the effectiveness of these measures and adjust as needed based on community feedback.
- Response 3: Implement a hybrid approach to park safety that combines educational initiatives with targeted infrastructure improvements. Begin by launching a comprehensive awareness campaign that encourages responsible pet ownership, highlights the importance of leash laws, and promotes a culture of mutual respect among all park users. Simultaneously, identify specific areas within the park where safety concerns are most prevalent and install strategic fencing and gates to mitigate risks without compromising the overall openness of the space. Regularly assess the effectiveness of these measures through community surveys and data analysis, making adjustments as needed to strike a balance between safety and accessibility.
- Response 4: Implement a "Balanced Park Safety" initiative that addresses the concerns of both park user groups. The initiative will focus on creating a safe and welcoming environment for all visitors, including those with children and dog owners. To achieve this, the park will introduce designated off-leash areas for dogs, ensuring that they have a secure space to play and socialize without compro-

mising the safety and comfort of other park users. Simultaneously, the park will enforce leash laws in all other areas, with clear signage and regular patrols by park staff to encourage compliance. The initiative will also include educational programs for dog owners, promoting responsible pet ownership and emphasizing the importance of respecting park rules and other users' needs.

Analysis

- Human Observation: The Chain of Thought with Feedback approach demonstrates the most comprehensive and balanced response generation. All five responses effectively address the negative story's dog-related safety concerns while incorporating the positive story's suggestions about physical safety measures. Notably, this approach generates creative solutions beyond the original suggestions, such as community feedback mechanisms and educational initiatives. The responses show deeper integration of both perspectives and introduce novel elements that enhance both safety and community aspects, demonstrating superior compromise generation capabilities.

Example 2

Positive view

I am writing about this place: This is our local church. I feel welcomed by others for who I am in this location because This church has become our home church after our previous church split up. They readily took our family in and made us feel loved and welcome..Some ways this place could be modified to be more welcoming are Honestly, I would not change anything about the practices, but I might meet with them to find out why they feel excluded. I would give them the gospel.

Negative view

I am writing about this place: Church,Downtown, Vegas. I feel excluded by others for who I am in this location because Unfortunately, organized religion and the LGBT+ community for the most part live in a space like oil and water. While there are some exceptions, many religious institutions utilize their beliefs as weapons as opposed to a welcome. I was raised as a catholic, but as I became an adult I quickly realized that who I am does not align with the beliefs and practices within organized religion. This space never feels welcoming or inclusive. It's a space I focus on avoiding with all of my energy. The sad part is the roots of religion are based off inclusion and making those who are looked upon by society as "less than" feel lifted up and included, however, that is not what is put into practice today..Some ways this place could be modified to be less excluding and more welcoming are If the church and organized religion went back to core principals of love, praise and inclusion, I truly feel that a space such as this beautiful church would be one that I would walk to weekly to share in a joyous occasion. It's unfortunate because this is exactly a half block from where I live. I hope one day that there is the realization that have veered off the right path. It takes one sector of organized religion to start an inclusion practice and accept no exceptions.

Basic Single Prompt

Results

- Response 1: Design an 'Inclusive Community Framework' with designated quiet zones and gradual integration options, using a holistic screening approach that considers present circumstances over past history, accommodating both those who prefer minimal interaction and those seeking community acceptance. [Framework too complex and potentially overwhelming]
- Response 2: Institute a 'Neighborhood Integration Choice' system where residents can customize their community involvement level, with fair entry requirements focusing on future contribution rather than past circumstances. [Integration choice system may reinforce divisions]
- Response 3: Create a 'Neighborhood Harmony Program' featuring both private sanctuaries and optional community spaces, with flexible residency requirements that focus on present community contribution and mutual respect, accommodating both those seeking solitude and those working on personal growth. [Harmony program too bureaucratic and impersonal]
- Response 4: Create community integration programs that provide both structured support for newcomers (including those in recovery) and optional private participation methods, respecting individual comfort levels with social interaction. [Community integration focus may miss core issues]

CoT

Results

- Response 1: Institute monthly community meetings where both current church members and those who feel excluded can work together to define specific, actionable steps toward mutual understanding and respect. [Community meetings might lack sufficient structure]
- Response 2: Develop a "Welcome All" program that combines gospel sharing with inclusive practices, allowing different groups to participate in church community service

1240
1241
1242
1243
1244
1245
1246
1247
1248
1249
1250
1251
1252
1253
1254
1255
1256
1257
1258
1259
1260
1261
1262
1263
1264
1265
1266
1267
1268
1269
1270
1271
1272
1273
1274
1275
1276
1277
1278
1279
1280
1281
1282
1283
1284
1285
1286

1287	projects together. ["Welcome All" program	CoT+Feedback	1332
1288	could feel pressuring to both sides]		
1289	• Response 3: Establish regular open dialogue	Results	1333
1290	sessions where church members and those		
1291	feeling excluded can share their perspectives		
1292	in a structured, respectful environment fo-		
1293	cused on finding common ground. [Open dia-		
1294	logue without clear goals may not progress]	• Response 1: The church could create a dedi-	1334
1295		cated safe space program where trained me-	1335
1296		diators host bi-weekly sessions allowing both	1336
1297		traditional church members and LGBTQ+ in-	1337
1298		dividuals to share their perspectives and find	1338
1299		common ground in their faith journey. [Re-	1339
1300	• Response 4: Create a dedicated outreach com-	specting both traditional religious practice and	1340
	mittee that includes both traditional church	LGBTQ+ perspectives]	1341
	members and LGBTQ+ individuals to discuss		
	ways to bridge understanding while maintain-		
	ing church values. [Outreach committee needs		
	more defined purpose]		
1301	CoT+LLM score		
1302	Results		
1303	• Response 1: The church could partner with		
1304	local LGBTQ+ organizations to develop inclu-	• Response 2: The church could establish an	1342
1305	sive worship services that maintain religious	LGBTQ+ liaison committee composed of	1343
1306	traditions while incorporating elements that	both congregation members and LGBTQ+	1344
1307	acknowledge and celebrate diverse identities,	community representatives to facilitate mean-	1345
1308	perhaps on alternate weekends. [Partnership	ingful dialogue and identify specific inclusion	1346
1309	for inclusive services is good but may push	opportunities while respecting core religious	1347
1310	too far too fast]	values. [Focuses on practical solutions]	1348
1311	• Response 2: The church leadership could		
1312	implement monthly "Open Door Sundays"		
1313	where they explicitly welcome and engage		
1314	with diverse community members through		
1315	structured interfaith discussions, followed by	• Response 3: Establish regular open dialogue	1349
1316	shared meals to build personal connections.	sessions where church members and those	1350
1317	[Open Door Sundays provides structure but	feeling excluded can share their experiences	1351
1318	might feel forced]	and perspectives in a moderated, respectful en-	1352
1319		vironment. [Regular open dialogue sessions	1353
1320	• Response 3: Create specific inclusion guide-	work well as they allow both sides to share	1354
1321	lines that combine traditional church values	experiences]	1355
1322	with explicit welcoming practices, developed		
1323	collaboratively by church leadership and com-		
1324	munity members. [Inclusion guidelines are		
1325	helpful but may be difficult to implement		
1326	fairly]		
1327	• Response 4: Create a dedicated outreach com-	• Response 4: Develop a structured welcome	1356
1328	mittee that includes both long-term members	program where new or returning members are	1357
1329	and those who have felt excluded, meeting	paired with existing church members for sup-	1358
1330	monthly to discuss inclusion challenges and	port and integration, regardless of their back-	1359
1331	solutions. [Outreach committee has potential	ground. [Being inclusive while respecting	1360
	but needs careful balance]	church structure]	1361

1362 11.2 Qualitative analysis of compromise
1363 generation

Example 1 (safe/less safe pair)

Positive view

I am writing about this place: The river next to a trail that I walk on. I feel safe here because I feel safe because it brings me joy and helps me to feel good. I feel comforted and comfortable and therefore feel safe.. Some ways this place could be modified to be safer are :It may feel safer knowing that people are around if your personal safety is a concern. There is only one path and no ability to get lost. Everyone is friendly.

Negative view

I am writing about this place: The river walk I feel safety could be improved here There have been several attacks along this river walk. It's not a place that people want to enjoy after dark. Part of the problem is the amount of drunk and high people that frequent the path after dark..Some ways this place could be modified to be safer are While there are lights and it is somewhat illuminated, it would be helpful to have regular security guards along the trail at night. Security cameras would also be very helpful along this walkway.

1364
1365 Generation

1366 **Base model (Llama):** The riverside trail is fun-
1367 damentally a place of community and connection.
1368 By establishing regular morning walking groups
1369 between 7-9 AM, we can enhance the existing
1370 friendly atmosphere. The natural gathering of
1371 people who enjoy the path creates an organic
1372 safety network. The single clear path already
1373 provides good navigation, and scheduled group
1374 activities would further strengthen the welcoming
1375 environment that makes this place special. [Leans
1376 heavily towards Positive view's perspective.]
1377

1378 **Base model (Llama)+FT:** The river walk
1379 requires enhanced security measures to ensure
1380 safety at all hours. Implementing a dedicated patrol
1381 system with trained local volunteers would provide
1382 necessary oversight, especially during evening
1383 hours. Installing emergency communication points

and improving lighting would address the current
safety concerns. These measures would help
transform the area from a potentially dangerous
space into a secure community asset. [Focuses
heavily on Negative view's security concerns.]

Aligned Llama: Create a "Community Safety
Network" that combines the trail's natural friend-
liness with practical security measures. This
would include trained volunteers from the local
community who patrol during both daylight
and evening hours, wearing recognizable but
approachable uniforms. Establish welcoming rest
areas to serve as both social gathering points and
safety zones, encouraging community interaction.
[Maintaining balanced consideration of both
perspectives.]

Example 2 (welcome/excluded pair)

Positive view

I am writing about this place: North Park, in my town.. I feel welcomed by others for who I am in this location because It's a hidden park that is rather unknown, it's always quiet, has a bench by the creek and trails to walk. I feel welcomed here because it allows me to take a break from my day and decompress. Some ways this place could be modified to be more welcoming are Someone who is not open to mindfulness, or who is loud and doesn't like being alone, may dislike this place. I would imagine they may feel too alone with their thoughts, and not able to appreciate what is around them due to what they may already be used to elsewhere, perhaps they thrive in a city environment. I would make this space more welcoming for them, simply by accompanying them on their visit. Providing them comfort while they experience what I do while I alone.

Negative view

I am writing about this place: The park that is in our neighborhood. I feel excluded by others for who I am in this location because I went a few times with my son and the people that have been there each time that we have gone pretty much claimed it for themselves. I actually know them from the school that I teach at and they don't like me because I am a Christian. At one point they said "um we are playing here" , And you could tell by tone of voice and facial expression they didn't want us there. You can find imagine that would make someone not even want to stay. Some ways this place could be modified to be less excluding and more welcoming are I would maybe put a sign up saying all people are welcome in our neighborhood or community. Not that that would change much but it would at least make me feel better

Generation

Base model (Mistral): The park should establish dedicated quiet zones with meditation benches and nature trails, creating peaceful spaces for individual reflection and mindfulness. Simple signs can

indicate these areas as spaces for contemplation, helping preserve the serene atmosphere while gently encouraging respectful behavior from all visitors [minimally addressing negative view's concerns about exclusion.]

Base model (Mistral)+FT: The park should implement a welcoming committee program where community volunteers greet visitors and organize inclusive group activities. Installing welcome signs in multiple languages and creating designated community gathering spaces would help ensure everyone feels accepted, regardless of their background or beliefs [No consideration for those who find peace in solitude]

Aligned Mistral: Clear signage and designated spaces would establish areas for both group activities and individual reflection, supported by a comprehensive community engagement plan addressing personal and social needs. A trained community mediator would help resolve any conflicts or exclusion, while regular forums bring diverse groups together to shape park policies that give equal weight to all experiences and concerns. [Maintaining balanced consideration of both perspectives.]

Conclusion

The Chain of Thought with Feedback approach generated the most balanced and nuanced responses, effectively integrating perspectives from both the positive and negative stories. It introduced creative solutions that enhanced both safety and community aspects. While this strategy has limitations such as being resource-intensive and potentially introducing bias, it demonstrated superior compromise generation capabilities compared to the basic single prompt and Chain of Thought approaches.

11.3 Collection of human ratings for evaluation of different prompting strategies

The Prolific⁵ platform was used to recruit participants located in the US. The participants were paid the equivalent of \$17/hr for 20 minutes of time, with the actual mean time of a batch of participants ranging from 14-20 minutes. All participants provided informed consent prior to participating in the

⁵<https://prolific.com>

study. The study was part of a larger study that was reviewed by a review board prior to launch of the study.

The participants were given the following overall instructions:

Welcome to the study. In this study, you will read accounts about similar places from different pairs of people, whom we will be labelling as Person A and Person B. The accounts are related to their perceptions of how welcoming or safe their neighborhoods are. A and B have not met and wrote their stories separately. After reading each pair, you will take the perspective of Person A and rate the acceptability of several suggested modifications. How acceptable would Person A find these suggestions?

Each participant was asked to rate compromises for five pairs of statements about suggested place modifications written by Person A and Person B. For each participant, the story A of a pair was first presented:

You are about to read Person A's story. Please pay special attention to this story. You will need to take this person's perspective when considering Person B's story and the suggestions afterwards. You should have just read the story of Person A, the person whose perspective you will take.

Whether story A is a positive view or a negative view was randomly assigned to participants, balancing for an equal number of positive and negative views. Next, the participant read story B of a pair of statements:

Now you will read Person B's story. Person B has had a much different experience. (A and B have not met and wrote their stories separately). Press the blue button when you are ready.

After reading the stories by Person A and Person B, a participant is asked to rate four generated compromises plus Person B's modifications.

Now please pretend to be Person A (the first person whose story you read). As Person A, you have just read Person B's story. You will see a list of suggested modifications that try to address both A and B. Carefully consider each of the following suggested modifications and rate how acceptable each suggestion is from your (Person A's) perspective. To help you remember, we have included both people's original suggestions below. A's suggested modification: <view A> B's suggested modification: <view B> Pretend you are Person A (the first person whose account you read). How acceptable is each suggestion? Use the slider BELOW each suggestion to make your rating.

The four suggestions plus view B were presented in random order, and a horizontal slider from 1 to 100 was shown below each text.

Once all sliders were adjusted, the slider ratings were recorded and the participant could click a button to move to the next pair of statements.

Once the participant rated five pairs of statements, they were asked demographic questions about their age, gender, income level,