

Generating Fair Consensus Statements with Social Choice on Token-Level MDPs

Carter Blair
Harvard University
carterblair@g.harvard.edu

Kate Larson
University of Waterloo
kate.larson@uwaterloo.ca

October, 2025

Abstract

Current frameworks for aggregating free-form text-based opinions with large language models lack the inherent structure needed to provide meaningful fairness guarantees. To address this, we model the task as a multi-objective, token-level Markov Decision Process (MDP), where each objective corresponds to an agent’s preference. Each agent’s token-level reward is induced by its policy (e.g., a personalized language model). Such policies implicitly define optimal Q-functions, thus enabling stepwise reward computation without an explicit value function [20]. This MDP formulation yields a formal structure that can be analyzed with tools from social choice theory. We first give a stochastic generation policy that is guaranteed to lie in the ex-ante core. It is derived from a distribution over complete statements that maximizes Nash welfare, extending core stability from cooperative game theory and voting to text generation. Second, for a single consensus statement, we target egalitarian welfare and use search within the MDP. Empirically, this search produces statements with improved worst-case agent alignment compared with baselines, including the Habermas Machine [26]. Our code is available at <https://github.com/cartgr/Generating-Fair-Consensus-Statements-with-Social-Choice-on-Token-Level-MDPs>.

Keywords: Generative Social Choice, Guided Decoding, AI-Augmented Deliberation

1 Introduction

Social choice theory has traditionally studied how to aggregate preferences over predefined alternatives. Large language models (LLMs) make it possible to aggregate free-form verbal opinions into collective textual outputs, which removes the constraint of a fixed agenda. Importantly, this flexibility can reduce the agenda-setting power of organizers, since participants need not choose from a preset list. However, ensuring provable fairness is difficult: complex training and design choices yield an opaque, irregular LLM structure, which in turn complicates the formalization of specific fairness criteria during generation. For this reason, previous approaches have treated the generation process as a black box, applying various fairness measures post-hoc. For instance, in the Habermas Machine [26], statements are first generated, and fairness is then pursued through a voting procedure applied to these statements, which were not themselves generated with an explicitly fair mechanism. Similarly, the Generative Social Choice method [12] prompts an LLM to maximize a given objective and assumes that the response truly maximizes the objective. These strategies, though aimed at fairness in free-form opinion aggregation, can concede a new form of agenda control to the LLM. When the model is asked to produce text that satisfies a broad fairness objective, its interpretation and implementation of that objective, the trade-offs it chooses, and the aspects of

opinions it elevates remain opaque. This effectively allows the LLM to shape the solution space. As such, these methods risk overlooking biases embedded within the generation process itself, which are known to exist [11].

We address this gap by modeling consensus statement generation as a token-level Markov Decision Process (MDP). Each agent i 's viewpoint is represented by a policy π_i , which assigns likelihoods $\pi_i(s, a)$ to token choices given the current prefix s . Following Rafailov et al. [20], who show that policies implicitly define optimal Q-functions, our agent policies, π_i , determine rewards $r_i(s, a)$ (e.g., $r_i^{\log}(s, a) = \beta \log \pi_i(s, a)$) at each generation step. A primary advantage of this reward formulation is that it avoids personalized value functions, which are known to be challenging to train and apply effectively [15]. This MDP structure provides a formal basis for integrating fairness principles directly into the construction of the consensus statement.

Within this MDP framework, we develop two approaches that leverage existing notions of fairness from social choice theory, namely the *ex-ante core* and *egalitarian welfare (EW)*, which we argue are compelling notions of fairness in the context of consensus statement generation. First, to achieve an outcome in the ex-ante core, we propose a stochastic generation policy Π^* . This policy is derived by optimizing a distribution over complete statements to maximize proportional fairness (Nash Welfare), a process known to yield core membership. For consensus generation, the core is a highly desirable stability concept: a lottery in the core ensures that no coalition of agents could unilaterally deviate and achieve an alternative lottery that all its members prefer, given their proportional influence, implying agreement, nay consensus, over the randomized outcome. Second, when a single consensus statement is desired, we target the maximization of EW, seeking the best outcome for the least satisfied agent, which aligns with the idea that a consensus statement should be agreeable to all parties. We introduce constructive search algorithms (finite lookahead and beam search) that optimize this EW objective directly within the MDP. This offers a transparent and analyzable generation mechanism distinct from methods reliant on high-level prompting or post-hoc voting.

Our main contributions are:

1. A formal token-level MDP framework for fair consensus generation where agent rewards are derived from their language model policies.
2. A method to derive a stochastic generation policy that is provably in the ex-ante core, ensuring proportional fairness and stability.
3. The development and empirical validation of search algorithms that, by optimizing egalitarian welfare within the MDP, generate single consensus statements with improved worst-case agent alignment compared to methods that do not leverage this token-level structure or search.

Through these contributions, we hope to establish a new direction for methods seeking to generate consensus statements from open-ended verbal opinions with provable fairness guarantees.

2 Related Work

Generative Social Choice. This field applies social choice principles to open-ended generation, such as creating text from diverse opinions [12, 26, 23, 3]. Unlike methods that aggregate preferences over predefined alternatives, Generative Social Choice (GSC) generates the alternatives themselves. For example, Tessler et al. [26] employ iterative critiques and voting on complete statements for consensus, while Fish et al. [12] prompt LLMs to directly maximize egalitarian welfare within a larger framework aimed at a form of proportional representation. Our work differs by embedding

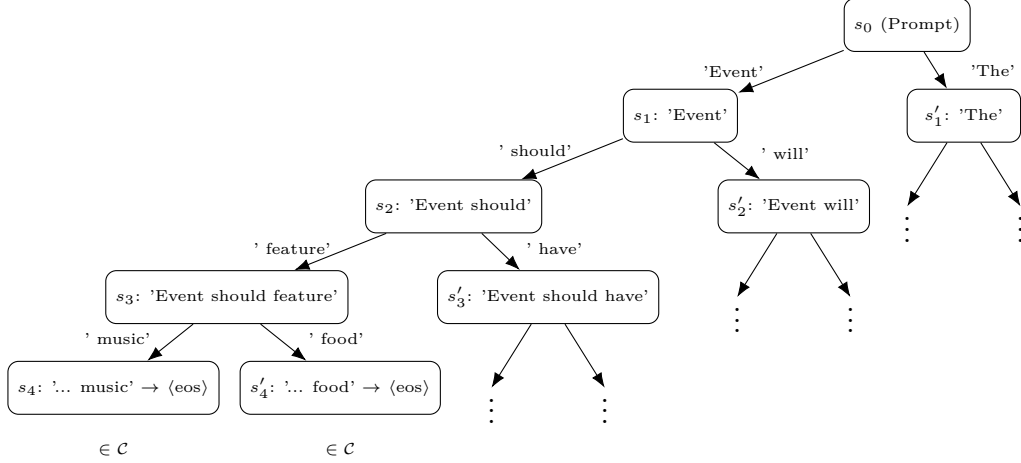


Figure 1: Illustration of the token-level generation tree. Edges represent actions (chosen tokens).

fairness into the token-by-token construction of a consensus statement via a multi-objective MDP, treating each token selection as a public decision [6]. This provides a more granular and verifiable mechanism than post-hoc evaluations or high-level prompting.

Randomized Social Choice. We also draw from randomized social choice, which studies lotteries over outcomes. Our stochastic generation policy, which maximizes Nash Welfare for proportional fairness, connects to this area. Maximizing Nash Welfare is known to yield outcomes in the 1-core [1, 10, 9]. We extend these findings to the sequential decision-making context of this paper.

Guided Decoding. Guided decoding techniques steer LLM generation towards desired attributes at inference time, often using search algorithms. Methods like PPO-MCTS [18] and VAS [16] use a value network to guide generation, while MOD [25] or COLLAB [5] combine or switch policies. Our approach also uses search but derives token-level rewards from agent policies. Further, we explicitly frame generation as planning in a multi-objective MDP to optimize social choice objectives (Proportional Fairness, Egalitarian Welfare), rather than relying on a single pre-trained value model or heuristic model combinations.

3 Problem Setup & Preliminaries

We consider a setting with a finite set of agents $N = \{1, 2, \dots, n\}$, each with a distinct opinion on a specific **Issue**. The goal is to generate a consensus text statement reflecting these perspectives fairly. The inputs to the process include descriptions of the issue, a policy representing each agent (which can be derived from free-form text expressing their opinion), and a **reference policy** (e.g., a base language model) that is used to propose tokens in the consensus statement.

Agent Policies. Each agent $i \in N$ is represented by a *policy* π_i , which assigns a likelihood $\pi_i(s, a) \in [0, 1]$ to each action a given the state s . Intuitively, $\pi_i(s, a)$ reflects how closely an action aligns with agent i 's preference at state s . This policy could be an LLM fine-tuned (ideally with DPO [21]) or base policy prompted with agent i 's viewpoint.

Token-Level MDP We model text generation as a deterministic, discrete-time Markov Decision Process defined by the tuple $(S, A, T, \mathbf{R})$. Here, S is the state space of partial text sequences (prefixes), including initial s_0 and terminal states. A is the action space consisting of the token vocabulary plus a special end-of-sequence token $\langle \text{eos} \rangle$. T is the deterministic transition function where $T(s, a) = s||a$ appends the chosen token; selecting $a = \langle \text{eos} \rangle$ leads to a terminal state representing a completed statement X . Finally, $\mathbf{R}$ represents the agent-specific reward functions. We define two types of rewards based on agent policies, serving different analytical purposes:

1. **Log-Likelihood Reward:** $r_i^{\log}(s, a) = \beta \log \pi_i(s, a)$. This formulation aligns with implicit rewards in preference learning [20], where $\beta > 0$ is a scaling factor. This is non-positive and is suitable for additive utility accumulation along a path.
2. **Likelihood Reward:** $r_i^{\text{prob}}(s, a) = \pi_i(s, a)$. This reward uses the direct probability, ensuring non-negativity ($r_i^{\text{prob}} \geq 0$), which is needed for social welfare functions involving products or ratios, such as Nash Welfare.

We denote by $\mathcal{C}$ the set of all possible complete paths (sequences ending in $\langle \text{eos} \rangle$) from s_0 .

In practice, we assume that at each non-terminal state s , we only consider a finite set of B possible next tokens $A_B(s) \subseteq A$. This set could be the model’s vocabulary or a subset chosen by a base language model giving us the B most likely next tokens. With this finite branching factor B and a maximum sequence length $L_{\max}$, the set $\mathcal{C}$ of complete paths is finite (bounded by $B^{L_{\max}}$).

This sequential token selection process naturally defines a tree structure rooted at s_0 . Each edge represents choosing a token from $A_B(s_t)$, and each node represents a partial sequence s_t . To make this concrete, we provide an example in Figure 1.

Agent Utilities. Given a completed sequence $X = (a_1, \dots, a_\ell = \langle \text{eos} \rangle)$ corresponding to states $(s_0, s_1, \dots, s_\ell)$, we define two corresponding utility functions for each agent i , derived from the respective reward types:

1. **Additive Log-Utility:** Primarily used for evaluating single paths based on cumulative log-likelihood.

$$U_i^{\log}(X) = \sum_{t=1}^{\ell} r_i^{\log}(s_{t-1}, a_t) = \sum_{t=1}^{\ell} \beta \log \pi_i(s_{t-1}, a_t) = \beta \log \left(\prod_{t=1}^{\ell} \pi_i(s_{t-1}, a_t) \right)$$

2. **Multiplicative Probability Utility:** Primarily used for evaluating distributions via expected utility, forming the basis for Nash Welfare and Proportional Fairness calculations.

$$U_i^{\text{prob}}(X) = \prod_{t=1}^{\ell} r_i^{\text{prob}}(s_{t-1}, a_t) = \prod_{t=1}^{\ell} \pi_i(s_{t-1}, a_t) = P_i(X)$$

This represents the joint probability of sequence X under agent i ’s policy.

These are related by $U_i^{\log}(X) = \beta \log U_i^{\text{prob}}(X)$, with $U_i^{\text{prob}}(X) > 0$.

3.1 Deterministic vs. Stochastic Policies

We consider two types of policies for generating fair consensus statements in our token-level MDP:

1. A deterministic policy $\mu(s)$ that gives us a single path $X \in \mathcal{C}$. To evaluate deterministic policies, we adopt additive log-utilities $U_i^{\log}(X)$.

2. A stochastic policy $\pi(a|s)$ that gives us a distribution $p \in \Delta(\mathcal{C})$ over paths (a lottery). When assessing this type of outcome, we consider the expected probability-based utility

$$\mathbb{E}_{X \sim p}[U_i^{\text{prob}}(X)] = \sum_{X \in \mathcal{C}} p(X) U_i^{\text{prob}}(X).$$

When it is clear from the context, we adopt the shorthand of $U_i^{\text{prob}}(p)$ to refer to the expected utility of agent i for a given distribution over paths. This expected utility $U_i^{\text{prob}}(p)$ is guaranteed to be non-negative, which is required for downstream fairness measures.

The Fairness of a Deterministic Policy. We assess the fairness of a path $X \in \mathcal{C}$ given by a deterministic policy through its egalitarian welfare (EW), drawing from Rawls' maximin principle [22]. This is defined using the additive log-utilities $U_i^{\text{log}}(X)$:

$$\text{EW}^{\text{log}}(X) = \min_{i \in N} U_i^{\text{log}}(X) = \min_{i \in N} \sum_{t=1}^{\ell} \beta \log \pi_i(s_{t-1}, a_t). \quad (1)$$

Maximizing $\text{EW}^{\text{log}}(X)$ means finding the path whose cumulative log-likelihood is highest for the agent who prefers it least. By maximizing the minimum utility, the egalitarian objective promotes broadly acceptable outcomes, which supports the requirement that a consensus statement should be agreeable to all.

The Fairness of a Stochastic Policy. To analyze the fairness of a distribution over paths given by our stochastic policy, we use the non-negative expected utilities $U_i^{\text{prob}}(p)$. Our goal is to find a stochastic policy that gives us a distribution $p \in \Delta(\mathcal{C})$ that is in the ex-ante core [24]. Following Fain et al. [10], we define the ex-ante core as follows.

Definition 1 (Ex-Ante Core). *A distribution $p \in \Delta(\mathcal{C})$ is in the ex-ante core if there is no coalition $S \subseteq N$ and alternative distribution p' such that*

$$\frac{|S|}{|N|} \cdot U_i^{\text{prob}}(p') \geq U_i^{\text{prob}}(p), \quad \forall i \in S,$$

with strict inequality for at least one agent $i \in S$.

Intuitively, if a distribution over paths is in the ex-ante core, there is no possible coalition of agents that could use their proportional share of probability to create a distribution that has strictly higher expected utility for at least one agent and not less expected utility for all other agents (i.e., a Pareto improvement) [1]. This resistance to coalitional deviation is desirable for consensus statements, as a distribution in the ex-ante core represents an outcome that all parties have implicitly agreed to (since they cannot use their proportional share of utility to do something better).

4 Defining a Stochastic Policy in the Core

Having established the token-level MDP and fairness criteria, we now turn to defining a *stochastic generation policy* π^* that produces a distribution p^* over complete consensus statements $\mathcal{C}$ (i.e., a lottery) that is in the ex-ante core.

To do so, we first generate the tree of possible token sequences with branching factor B and maximum length L . We then, find the distribution over paths p^* that maximizes Nash welfare (NW). The NW optimal distribution is

$$p^* = \max_{p \in \Delta(\mathcal{C})} \text{NW}(p) = \max_{p \in \Delta(\mathcal{C})} \prod_{i=1}^n U_i^{\text{prob}}(p).$$

We can find p^* by optimizing over the probability simplex $\Delta(\mathcal{C}) = \{p \in \mathbb{R}_{\geq 0}^{|\mathcal{C}|} : \sum_{X \in \mathcal{C}} p(X) = 1\}$. Indexing the leaves as $\mathcal{C} = \{X_1, \dots, X_m\}$ and defining $u_i \in \mathbb{R}_{\geq 0}^m$ by $u_i(j) = U_i^{\text{prob}}(X_j)$, we get that each agent's expected utility under distribution p is

$$U_i^{\text{prob}}(p) = \sum_{j=1}^m u_i(j) p_j = u_i^\top p,$$

so the Nash welfare program is

$$p^* = \max_{p \in \Delta(\mathcal{C})} \sum_{i=1}^n \log(u_i^\top p).$$

Note, each term $\log(u_i^\top p)$ is concave (log is concave and increasing; $u_i^\top p$ is affine), and the simplex is convex; hence this is a convex program with polynomial-time algorithms and strong duality under a mild positivity condition (i.e., $u_i^\top p > 0$ for all i) [4].

Moreover, we know that any maximizer of $\prod_i U_i^{\text{prob}}(p)$ lies in the core [10, 1, 9]. So, we can work backwards from this distribution to find a stochastic policy in the core.

Specifically, given the target distribution $p^* \in \Delta(\mathcal{C})$ that maximizes Nash welfare, we derive the stochastic policy π^* that generates this distribution. To define π^* , we introduce some notation. For any state (prefix) s in the token-level MDP:

- Let $\mathcal{C}(s) \subseteq \mathcal{C}$ be the set of all complete paths (leaves) that pass through state s .
- Let $\mathcal{C}(s, a) \subseteq \mathcal{C}(s)$ be the subset of paths in $\mathcal{C}(s)$ where the next action taken from state s is a . Note that $\mathcal{C}(s, a) = \mathcal{C}(s||a)$, where $s||a$ is the state reached after taking choosing token a .
- For any subset of leaves $L \subseteq \mathcal{C}$, let $P^*(L) = \sum_{X \in L} p^*(X)$ be the total probability mass assigned by the NW optimal distribution p^* to the leaves in L .

Note that $P^*(\mathcal{C}(s_0)) = P^*(\mathcal{C}) = 1$.

With this, we can define the policy π^* at any given state s .

Definition 2 (Stochastic Policy Induced by p^*). *Let p^* be a distribution over the leaf nodes $\mathcal{C}$. The induced stochastic policy π^* at a non-terminal state s assigns the probability of taking the next action (token) a as:*

$$\pi^*(a|s) = \begin{cases} \frac{P^*(\mathcal{C}(s,a))}{P^*(\mathcal{C}(s))} & \text{if } P^*(\mathcal{C}(s)) > 0 \\ 0 & \text{if } P^*(\mathcal{C}(s)) = 0 \end{cases} \quad (2)$$

This represents the conditional probability, according to the NW optimal distribution p^* , of selecting token a next, given that the generation process has reached state s . If state s has zero probability of being reached under p^* (i.e., $P^*(\mathcal{C}(s)) = 0$), then the probability of taking any action from s is also zero.

Algorithm 1 summarizes the process of finding a policy in the ex-ante core. We will now briefly turn to the runtime of this algorithm.

Algorithm 1 Compute core-stochastic policy π^* **Require:** Leaf set $\mathcal{C} = \{X_1, \dots, X_m\}$, agent utilities $U_i^{\text{prob}}(X_j)$

- 1: Build $u_i \in \mathbb{R}_{\geq 0}^m$ with $u_i(j) = U_i^{\text{prob}}(X_j)$
- 2: Solve $p^* \in \arg \max_{p \in \Delta(\mathcal{C})} \sum_{i=1}^n \log(u_i^\top p)$
- 3: For every state s in the tree, cache $P^*(\mathcal{C}(s))$ and, for each enabled action a at s , cache $P^*(\mathcal{C}(s, a))$.
- 4: **if** $P^*(\mathcal{C}(s)) > 0$ **then**
- 5: $\pi^*(a \mid s) \leftarrow \frac{P^*(\mathcal{C}(s, a))}{P^*(\mathcal{C}(s))}$
- 6: **else**
- 7: $\pi^*(a \mid s) \leftarrow 0$
- 8: **end if**
- 9: **return** π^*

Runtime. We optimize the Nash-welfare program on $\Delta(\mathcal{C})$ (with $m = |\mathcal{C}|$ as above) and then induce π^* via the conditional masses $P^*(\mathcal{C}(s))$ and $P^*(\mathcal{C}(s, a))$. Forming all agent-leaf utilities costs $O(nmL)$. We then minimize $-\sum_i \log(u_i^\top p)$ on the simplex with Frank-Wolfe: each iteration computes $\nabla F(p) = \sum_i u_i / (u_i^\top p)$ in $O(nm)$ time and uses a coordinate linear oracle; the method attains ε -suboptimality in $O(1/\varepsilon)$ iterations [17]. Hence the optimization time is $O(nm/\varepsilon)$. The conditional masses needed for π^* are obtained by a single bottom-up pass over the generation tree, which is linear in the number of leaves when B is fixed, i.e., $\Theta(m)$. Altogether, the total runtime is

$$T_{\text{total}} = O(nmL + nm/\varepsilon + m),$$

which is polynomial in n , m , and $1/\varepsilon$.

4.1 Properties of the Stochastic Policy

We now establish that executing this policy π^* from the initial state s_0 indeed generates the target distribution p^* .

Theorem 1 (Equivalence of Policy-Induced Distribution and Target Lottery). *Let p_{π^*} be the distribution over $\mathcal{C}$ generated by executing the policy π^* (defined in Definition 2) from the initial state s_0 . Then $p_{\pi^*} = p^*$.*

The proof is presented in subsection C.1. This equivalence leads to the desired ex-ante fairness guarantee for the policy π^* .

Corollary 1 (Core Membership of Stochastic Policy). *Let p^* be a distribution over $\mathcal{C}$ that maximizes Nash Welfare (and is therefore in the ex-ante core). Let π^* be the stochastic policy derived from p^* according to Definition 2. Then the distribution p_{π^*} generated by executing π^* is in the ex-ante core.*

Proof. By Theorem 1, the distribution generated by policy π^* is $p_{\pi^*} = p^*$. Since p^* was chosen to maximize Nash Welfare over $\mathcal{C}$, it is in the core. Therefore, p_{π^*} is also in the core. $\square$

This simple result confirms that our procedure, which first finds the distribution maximizing Nash Welfare p^* and then executes the derived policy π^* , yields a stochastic policy in the core. And thus, we have an ex-ante fair way to generate consensus statements.

4.2 Empirical Core Test

Here we empirically validate that the NW optimal *policy* π^* is in fact in the ex-ante core, and demonstrate that a uniform policy and a utilitarian optimal policy are not in the ex-ante core.

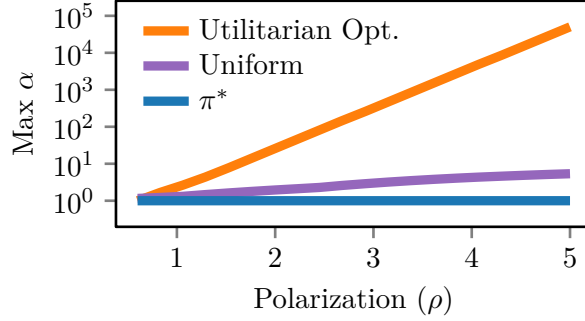


Figure 2: Maximum coalition improvement vs. polarization.

Environment For demonstration we use a small B -ary token tree ($B=3$, depth $L=4$). Each agent i has a next-token policy

$$\pi_i(a \mid s) = \text{softmax}_a(\rho w_i^\top(z + v_{t,a})),$$

where z is the running state embedding, $v_{t,a}$ is the token vector, and w_i is the agent vector. The scalar *polarization* $\rho \geq 0$ controls how concentrated preferences are. When $\rho = 0$, $\pi_i(a \mid s) = 1/B$. For two actions a, b ,

$$\frac{\pi_i(a \mid s)}{\pi_i(b \mid s)} = \exp(\rho w_i^\top(v_{t,a} - v_{t,b})),$$

so increasing ρ multiplies odds exponentially and makes agents much more confident about their preferred action.

To initialize the experiment we draw token vectors $\tilde{v}_{t,a} \sim \mathcal{N}(0, I_d)$ and agent vectors $\tilde{w}_i \sim \mathcal{N}(0, I_d)$, set $v_{t,a} = \tilde{v}_{t,a} / \|\tilde{v}_{t,a}\|_2$ and $w_i = \tilde{w}_i / \|\tilde{w}_i\|_2$, and define the state vector at prefix $s = (a_1, \dots, a_{t-1})$ as $z(s) = \sum_{\tau=1}^{t-1} v_{\tau, a_\tau}$.

Baselines. We compare π^* to a *uniform* policy that induces a uniform distribution over leaves and a *utilitarian* policy that deterministically follows the single path maximizing $\sum_i U_{ij}$.

Blocking test. For any policy π , let p_π be its induced leaf distribution and $u_i(\pi) = U_i^\top p_\pi$. For a coalition S of size r , give it a budget r/n of probability mass and solve a small LP to find the largest factor α such that all $i \in S$ can secure at least $\alpha u_i(\pi)$. The *maximum coalition improvement*

$$\alpha^*(\pi) = \max_{S \neq \emptyset} \max_{p': \mathbf{1}^\top p' = r/n, p' \geq 0} \min_{i \in S} \frac{U_i^\top p'}{u_i(\pi)}$$

The value of $\alpha^*(\pi)$ is > 1 exactly when a coalition can block π . Figure 2 plots $\alpha^*(\pi)$ versus polarization ρ (log y -axis).

Takeaways. Figure 2 shows the result. As ρ increases, agents' own policies become more concentrated. The Nash-welfare *policy* stays at $\alpha^*(\pi^*) = 1$ across all ρ , consistent with the theoretical result. The utilitarian policy becomes highly blockable, and the uniform policy becomes steadily easier to block.

4.3 Computational Tractability via Token Chunking

The set of complete statements $\mathcal{C}$ grows as $|\mathcal{C}| \approx B^{L_{\max}}$ for branching factor B and maximum length $L_{\max}$. Enumerating leaves and solving for p^* over $\mathcal{C}$ could be infeasible for long sequences.

We therefore restrict attention to a coarser action space obtained by grouping tokens into macro-actions.

Definition 3 (Token chunking). *A chunking scheme $\mathcal{K}$ partitions positions into contiguous blocks $\{k_1, \dots, k_m\}$, where each k_j contains one or more tokens and the concatenation of chosen blocks forms a complete statement. Decoding now selects entire blocks. The corresponding feasible set of leaves is $\mathcal{C}_{\mathcal{K}} \subseteq \mathcal{C}$.*

With fixed chunk size c , the effective depth drops from L to $\lceil L/c \rceil$. Running Algorithm 1 on the chunked tree produces a lottery that lies in the ex-ante core *with respect to* $\mathcal{C}_{\mathcal{K}}$; all coalitions are evaluated against this feasible set. This interpretation mirrors real elections: not every eligible candidate runs, yet we judge the outcome by comparing the candidates who did run. Here, chunking plays the role of eligibility: it narrows the set of feasible statements, and fairness is defined relative to that set.

However, this restriction does have a consequence. Because $\mathcal{C}_{\mathcal{K}} \subseteq \mathcal{C}$, the optimal Nash welfare over $\mathcal{C}_{\mathcal{K}}$ cannot exceed that over $\mathcal{C}$. Moreover, no uniform multiplicative guarantee is possible without further structure:

Theorem 2 (No constant-factor approximation under chunking). *For any constant $R > 1$, there exists a two-agent instance, a finite set of leaves C , and a chunked subset $K \subset C$ such that, writing*

$$\text{NW}_Y(p) = \prod_{i=1}^2 \left(\sum_{x \in Y} u_i(x) p(x) \right) \quad \text{for } Y \subseteq C, p \in \Delta(Y),$$

and letting $p_Y^ \in \arg \max_{p \in \Delta(Y)} \text{NW}_Y(p)$, we have*

$$\frac{\text{NW}_C(p_C^*)}{\text{NW}_K(p_K^*)} > R.$$

Hence the approximation ratio of the chunked solution, measured against the unchunked optimum, is unbounded.

The proof appears in the Appendix. The theorem states a worst case that arises when chunking prunes precisely those statements that both agents value highly.

Observe that we can think about this as an anytime algorithm: enlarging $\mathcal{C}_{\mathcal{K}}$ can only improve the Nash welfare, and the ex-ante core guarantee continues to hold relative to the current feasible set.

5 Generating a Single Statement

Although our stochastic policy π^* achieves a highly desirable ex-ante fairness guarantee, many practical applications require selecting a single consensus statement. In this case, our objective shifts from finding a fair distribution to identifying the single path (i.e., statement) that represents consensus.

5.1 Finding the Egalitarian Path

Given the token tree with leaf nodes $\mathcal{C}$, we aim to find a deterministic policy μ^* that produces a single path $X^* \in \mathcal{C}$ that maximizes egalitarian welfare as defined in Equation 1. Due to the size of the token tree, exhaustive search for X^* may be intractable. We propose approximate algorithms to find high-quality paths, including finite-lookahead search and beam search, which are detailed below.

Finite Lookahead Search. The finite lookahead algorithm operates with a rolling horizon. At each step t , it explores all possible paths P of length up to d originating from s_t . For each such path P , the algorithm evaluates the egalitarian welfare of the sequence formed by concatenating the path generated so far (X_{prefix}) with P . It then chooses the first action a^* of the path P^* that maximizes this lookahead evaluation, transitions to state $s_{t+1} = T(s_t, a^*)$, and repeats the process. This d -step lookahead can mitigate the potential for hedging inherent in greedy search. When no single immediate token is agreeable (i.e., results in high egalitarian welfare), a greedy method might select less informative tokens that avoid commitment. In contrast, a lookahead can identify longer sequences that, despite potentially controversial initial steps, lead to states with higher overall welfare, perhaps by expressing a concept with suitable qualifications. The algorithm is shown in Appendix D.

Beam Search. Beam search is a heuristic search algorithm that balances greedy search and exhaustive exploration, and has been effective in sequence generation tasks like machine translation and text generation [19, 20]. Instead of pursuing only the single best option (greedy search) or all options (exhaustive search), beam search maintains a fixed number of the most promising partial paths (hypotheses), w (the beam width), at each depth t . At each step, it expands paths in the beam by generating potential successor tokens. These candidates are then evaluated using the egalitarian welfare objective function, and only the top w scoring paths are retained for the next step. The algorithm returns the highest-scoring complete path found within the beam at the maximum length or upon reaching a terminal state. The algorithm is shown in Appendix D.

6 Experiments

We conduct experiments with three main objectives. First, we test whether agent policies derived from prompting language models show meaningful correlation with human preferences. Second, we examine whether these prompted policies can perform credit assignment to generate meaningful token-level rewards (a necessary prerequisite for our search algorithms). Third, we evaluate how well our finite lookahead and beam search methods perform in generating single consensus statements when compared to baseline approaches.

6.1 Opinion-length scaling experiment

We test whether a chat LLM’s conditional likelihood of a policy statement, given a participant’s written opinion, predicts that participant’s 1–5 Likert rating, and whether supplying more of the participants opinion improves prediction.

Data. We use the abortion dataset from Fish et al. [12]. This dataset contains data from 42 human participants who each provided their opinion on abortion in natural language. Each participant also

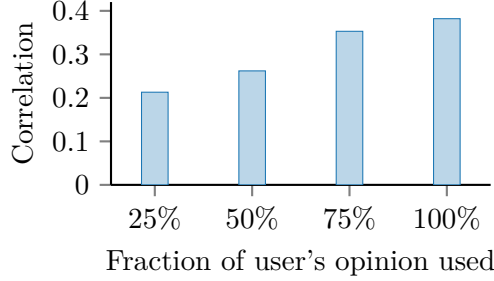


Figure 3: Spearman’s correlation between model likelihood and participants’ Likert ratings vs. the fraction of the user’s opinion provided.

rated five other candidate statements about abortion using a 1-5 Likert scale where higher values indicate higher agreement.

Method. We evaluate how well candidate statements capture participants’ opinions using language model scoring. For each participant u and statement j , we construct a paraphrase validation task using Meta-Llama-3.1-8B-Instruct-Turbo with temperature 1.0, computing log-probabilities without generation.

We structure the prompt as follows: the system instruction establishes the paraphrasing task, the user message provides the issue topic and the participant’s original opinion, and we place the candidate statement as a pre-filled assistant response:

```
[System] You paraphrase people’s views.
[User] Topic: {issue}
      Original opinion: {opinion}
[Assistant] Paraphrase: {statement}
```

Given the assistant response following the “Paraphrase:” tokenized as $a_1, \dots, a_M$, we compute the length-normalized log-probability score:

$$s_{u,j} = \frac{1}{M} \sum_{k=1}^M \log p_{\theta}(a_k \mid \text{context}, a_{1:k-1})$$

Intuitively, this score measures how well the statement represents the participant’s opinion. For validation, we compute Spearman’s correlation between each participant’s five model-assigned scores $\{s_{u,j}\}_{j=1}^5$ and their corresponding Likert ratings for the same statements.

Length manipulation. We repeat the procedure after truncating each participants opinion to the last fraction $f \in \{1.00, 0.75, 0.50, 0.25\}$ of its words.

Results. The correlation between the average likelihood of the statement according to the user prompted policy and the user’s rating increases with the fraction of the user’s opinion that is conditioned on (Fig. 3). Thus, conditioning on more of the opinion yields better predictions of human ratings. We note that with the users’ full opinions the correlation is still somewhat low, but the positive trend observed suggests that eliciting more informative statements could lead to a more accurate signal.

Table 1: Credit assignment results for Llama 3.1 8B Instruction-Tuned. Darker green indicates larger Z-score. Z-score column is for altered tokens. Alterations are represented by "<mis-aligned>/<aligned>".

Reference policy prompt	User policy prompt	Sequence	Z-Score
User food profile: empty	User food profile: vegetarian	I am having chicken/tofu enchiladas tonight. Then I am going to meet up with some friends.	2.69
User location profile: empty	User location profile: lives in a cold climate	I 'm going to the beach/mountains this weekend to surf/ski . I need to buy some new clothes.	1.78, 3.31
User time profile: empty	User time profile: morning	I am about to eat some food. I am going to have spaghetti/pancakes . I will use my phone to order it.	4.26
User opinion: empty	User opinion: Favors stricter gun control laws.	Implementing background checks that are less/more strict for gun purchases is essential . Also, my favorite color is orange.	2.09

6.2 Evaluating Credit Assignment

Setup. Our framework defines token-level rewards as $r_i(s, a) = \beta \log \pi_i(a | s)$, where each policy π_i is obtained by prompting a base LLM with agent-specific information. Reward-guided search is effective when policies exhibit *localized credit assignment*: changes in $\pi_i(a | s)$ concentrate on tokens tied to the prompt information. Rafailov et al. [20] observed this behavior in DPO-trained models; here we test it for policies induced by prompting instruction-tuned models.

Test design. We use Llama 3.1 8B Instruction-Tuned. For each test, we compare token log-probabilities under a **user policy prompt** (e.g., "User time profile: morning") and a **reference policy prompt** (e.g., "User time profile: empty"). We evaluate two nearly identical sequences: X_1 contains a concept that conflicts with the user profile (e.g., "spaghetti" for a "morning" profile) and X_2 replaces it with an aligned concept (e.g., "pancakes").

Metric. For each token a_j with prefix s , we compute the log-likelihood difference between the user and reference policies,

$$\Delta L(a_j | s) = \log \pi_U(a_j | s) - \log \pi_R(a_j | s).$$

We then measure the change in this difference when switching from X_1 to X_2 ,

$$D_j = |\Delta L_{X_1}(a_j | s) - \Delta L_{X_2}(a_j | s)|.$$

A large D_j indicates that the user profile alters the model's preference for token a_j specifically when the conflicting concept is swapped for an aligned concept. To compare across positions, we convert $\{D_j\}$ to Z-scores using the mean and standard deviation over all tokens in the sequence.

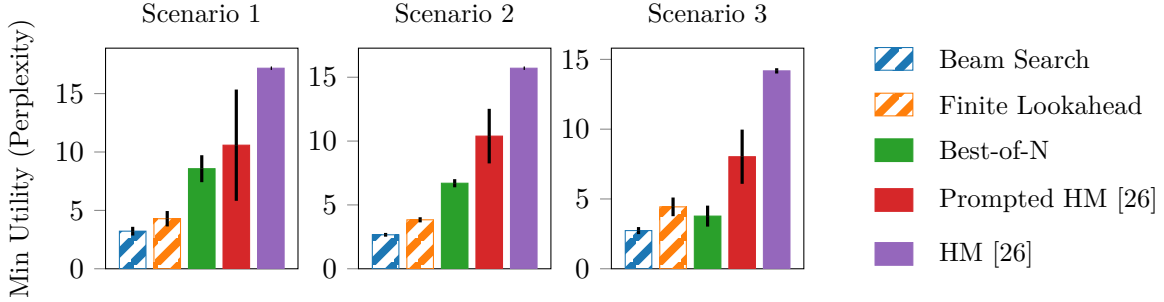


Figure 4: Per-scenario egalitarian welfare (perplexity). Lower values indicate better minimum agent utility. Striped bars indicate that the method uses search over the token-level MDP. Numerical results are reported in Table 10 in Appendix E.

Findings. The largest Z-scores occur at the tokens that differ between X_1 and X_2 (Table ??). With a “morning” time profile, the “spaghetti” versus “pancakes” position shows the highest shift. Additional examples with Llama and Gemma 2 9B Instruction-Tuned in subsection E.1 show the same pattern.

Implication. System prompting induces localized credit assignment in instruction-tuned models. This supports using $r_i(s, a) = \beta \log \pi_i(a | s)$ from prompted policies as a targeted signal for token-level search.

6.3 Consensus Generation

We evaluated different approaches for generating a single consensus statement by comparing our proposed search algorithms against several baselines. The primary goal was to assess how well each method optimizes egalitarian welfare (EW), measured by a perplexity-based metric reflecting worst-case agent alignment. More detailed consensus generation experiments, with Gemma 2 9b Instruction-Tuned, an LLM-judge metric, and with more agents are presented in subsection E.2 of Appendix E.

Scenarios: We used scenarios from the Habermas Machine dataset [26]. To obtain distinct settings, scenario descriptions were embedded using BAAI/bge-large-en-v1.5 [27] and clustered via k -means ($k = 3$). Representative scenarios were selected from each cluster (Scenarios 1, 2, and 3). The issues for these scenarios are in the captions of Tables 11, 12, and 13.

Agents and Policies: For each scenario, agent opinions were taken from the dataset. Agent policies π_i were instantiated by prompting Llama 3.1 8B Instruct [14] with the issue and agent i ’s opinion, instructing it to generate text aligned with that viewpoint (prompt shown in Figure 7 in Appendix F). The resulting likelihoods $\pi_i(s, a)$ represent agent i ’s preferences. Agent opinions are detailed in Tables 11, 12, and 13 in subsection E.4 of Appendix E.

Base Generation Model: Consensus statements were generated using Llama 3.1 8B Instruct, prompted with the issue and all agent opinions (prompt shown in Figure 6 in Appendix F).

Evaluation Metric - Egalitarian Perplexity (EPPL): To capture alignment with the least satisfied agent for a consensus statement X , we define Egalitarian Perplexity. For each agent i , their specific perplexity $PPL_i(X)$ is found by prompting **Llama 3.1 8B Instruct** with the issue and agent i 's opinion to generate a statement perfectly reflecting that opinion. The average log-likelihood of the actual consensus statement $X = (a_1, \dots, a_L)$ conditioned on this agent-specific prompt is:

$$\bar{L}_i(X) = \frac{1}{L} \sum_{t=1}^L \log \pi_i(s_{t-1}, a_t | \text{prompt}_i).$$

The agent-specific perplexity is $PPL_i(X) = \exp(-\bar{L}_i(X))$. The final Egalitarian Perplexity for X is $EPPL(X) = \max_{i \in N} PPL_i(X)$. Lower EPPL indicates better egalitarian welfare.

Seeds: We report the mean and standard deviation of EPPL over 3 seeds per method and scenario.

Methods Compared: We compared our proposed algorithms, **Finite Lookahead** (Algorithm 2, depth $d = 4$, branching $B = 2$), and **Beam Search** (Algorithm 3, width $w = 4$, pruning based on partial EPPL), against three baselines: **Best-of-N** (selecting the best of $N = 4$ samples¹ from the base model by EPPL); a **Prompted Habermas Machine**² (1 critique round, 4 candidates³, critiques from base model conditioned on agent opinions); and the original **Habermas Machine** (HM) [26] consensus (generated by a fine-tuned Chinchilla-70B).

6.3.1 Results

Figure 4 summarizes EPPL performance (lower is better). **Beam Search consistently achieved the lowest EPPL** (average: 2.87), indicating high alignment with the least satisfied agent. **Finite Lookahead also performed well** (average EPPL: 4.18), outperforming baseline methods. Both search methods surpassed **Best-of-N** (6.35) and the **Prompted Habermas Machine** (9.67). The **Habermas Machine** baseline had the highest EPPL (15.69), possibly because its statement was generated by a different model (Chinchilla 70B).

The results suggest that token-level search guided by EPPL, as in Beam Search and Finite Lookahead, effectively generates consensus statements with better minimum agent alignment compared to sampling or iterative refinement. Consensus statements for the first seed are in Tables 11, 12, and 13. The strong empirical performance of methods operating on the token-level MDP complements the fact that these methods are also more amenable to theoretical analysis and fairness guarantees.

7 Discussion

This work introduced a framework for generating consensus statements by modeling the process as a multi-objective, token-level MDP with rewards derived from agent-specific language model policies. Our aim was to connect LLM-based text generation with the formal fairness guarantees of social choice theory via this MDP.

Our theoretical contributions for stochastic outcomes (lotteries over statements) focused on the core. By maximizing Nash Welfare over expected probability-based utilities, we identified an optimal

¹ $N = 4$ was chosen to align with the Prompted Habermas Machine.

²As implemented in https://github.com/google-deepmind/habermas_machine

³We chose four candidates to align with the default parameters in the Prompted Habermas Machine example in the Habermas Machine GitHub repository.

lottery p^* that induces a stochastic generation policy π^* (Definition 2) inheriting the ex-ante core property (Corollary 1). Chunking was introduced as a heuristic to manage the search space. For deterministic outcomes (single statements), we focused on maximizing egalitarian welfare (EW), proposing finite lookahead and beam search as approximation algorithms.

Empirical results validated several aspects of our framework. We found that token likelihood from prompted policies meaningfully correlate with human preferences. Credit assignment experiments (subsection 6.2) confirmed that prompting LLMs with agent profiles creates policies that not only correlate with human preferences, but that also correctly assign credit to tokens. Consensus generation experiments (subsection 6.3), using Egalitarian Perplexity (EPPL) to measure EW, showed that beam search and finite lookahead, guided by the EW objective, outperformed baselines like Best-of-N and an adapted Habermas Machine. Beam search yielded the lowest EPPL.

Overall, formulating consensus generation as a search problem within a token-level MDP, guided by explicit social choice objectives like EW, is a promising direction. However, there are some important questions that future work should address. First, finding methods to train more faithful personalized policies for each agent is an important direction as it is upstream of many important challenges [2]. Second, finding a way to approximate the core on the unchunked space without looking at the whole tree is an important step to work towards. Previous work has looked at approximating the core [9, 13], but neither method directly applies to our setting. And third, future work should focus on theoretical guarantees for the single statement case, which we did not obtain.

Lastly, we note that until our methods are better understood, the outputs of our algorithm should be treated as artifacts for collective sense-making instead of binding decisions, as suggested by Revel and Pénigaud [23]. For example, instead of treating the output as a decision, the output could be treated as another input to the discussion that participants of the collective decision could reflect on. Optimistically, one would hope that these consensus statements could identify previously unknown points of agreement or solutions that no one had previously thought of that are in fact highly agreeable. In sum, this work contributes theoretical foundations and practical algorithms for incorporating social choice principles into generative AI for collective decision-making and sense-making.

References

- [1] Haris Aziz, Anna Bogomolnaia, and Hervé Moulin. 2019. Fair mixing: The case of dichotomous preferences. In *Proceedings of the 2019 ACM Conference on Economics and Computation*. 753–781.
- [2] Carter Blair, Kate Larson, and Edith Law. 2025. Reflective Verbal Reward Design for Pluralistic Alignment. In *Proceedings of the Thirty-Fourth International Joint Conference on Artificial Intelligence, IJCAI-25*, James Kwok (Ed.). International Joint Conferences on Artificial Intelligence Organization, 10271–10279. doi:10.24963/ijcai.2025/1141 Human-Centred AI.
- [3] Niclas Boehmer, Sara Fish, and Ariel D. Procaccia. 2025. Generative Social Choice: The Next Generation. In *Forty-second International Conference on Machine Learning*. <https://openreview.net/forum?id=E1E6T7KH1R>
- [4] Stephen P Boyd and Lieven Vandenberghe. 2004. *Convex Optimization*. Cambridge University Press.

- [5] Souradip Chakraborty, Sujay Bhatt, Udari Madhushani Sehwag, Soumya Suvra Ghosal, Jiahao Qiu, Mengdi Wang, Dinesh Manocha, Furong Huang, Alec Koppel, and Sumitra Ganesh. 2025. Collab: Controlled Decoding using Mixture of Agents for LLM Alignment. In *The Thirteenth International Conference on Learning Representations*. <https://openreview.net/forum?id=7ohlQUbTpp>
- [6] Vincent Conitzer, Rupert Freeman, and Nisarg Shah. 2017. Fair public decision making. In *Proceedings of the 2017 ACM Conference on Economics and Computation*. 629–646.
- [7] Avinava Dubey, Zhe Feng, Rahul Kidambi, Aranyak Mehta, and Di Wang. 2024. Auctions with LLM summaries. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. 713–722.
- [8] Paul Duetting, Vahab Mirrokni, Renato Paes Leme, Haifeng Xu, and Song Zuo. 2024. Mechanism design for large language models. In *Proceedings of the ACM Web Conference 2024*. 144–155.
- [9] Soroush Ebadian, Anson Kahng, Dominik Peters, and Nisarg Shah. 2024. Optimized distortion and proportional fairness in voting. *ACM Transactions on Economics and Computation* 12, 1 (2024), 1–39.
- [10] Brandon Fain, Kamesh Munagala, and Nisarg Shah. 2018. Fair allocation of indivisible public goods. In *Proceedings of the 2018 ACM Conference on Economics and Computation*. 575–592.
- [11] Shangbin Feng, Chan Young Park, Yuhua Liu, and Yulia Tsvetkov. 2023. From Pretraining Data to Language Models to Downstream Tasks: Tracking the Trails of Political Biases Leading to Unfair NLP Models. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Anna Rogers, Jordan Boyd-Graber, and Naoaki Okazaki (Eds.). Association for Computational Linguistics, Toronto, Canada, 11737–11762. doi:10.18653/v1/2023.acl-long.656
- [12] Sara Fish, Paul Gözl, David C. Parkes, Ariel D. Procaccia, Gili Rusak, Itai Shapira, and Manuel Wüthrich. 2024. Generative Social Choice. In *Proceedings of the 25th ACM Conference on Economics and Computation (New Haven, CT, USA) (EC '24)*. Association for Computing Machinery, New York, NY, USA, 985. doi:10.1145/3670865.3673547
- [13] Ian Gemp, Marc Lanctot, Luke Marris, Yiran Mao, Edgar Duéñez-Guzmán, Sarah Perrin, Andras Gyorgy, Romuald Elie, Georgios Piliouras, Michael Kaisers, et al. 2024. Approximating the Core via Iterative Coalition Sampling. In *Proceedings of the 23rd International Conference on Autonomous Agents and Multiagent Systems*. 669–678.
- [14] Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, et al. 2024. The Llama 3 herd of models. *arXiv preprint arXiv:2407.21783* (2024).
- [15] Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. 2025. Deepseek-r1: Incentivizing reasoning capability in LLMs via reinforcement learning. *arXiv preprint arXiv:2501.12948* (2025).
- [16] Seungwook Han, Idan Shenfeld, Akash Srivastava, Yoon Kim, and Pulkit Agrawal. 2024. Value augmented sampling for language model alignment and personalization. *arXiv preprint arXiv:2405.06639* (2024).

- [17] Martin Jaggi. 2013. Revisiting Frank-Wolfe: Projection-free sparse convex optimization. In *International conference on machine learning*. PMLR, 427–435.
- [18] Jiacheng Liu, Andrew Cohen, Ramakanth Pasunuru, Yejin Choi, Hannaneh Hajishirzi, and Asli Celikyilmaz. 2024. Don’t throw away your value model! Generating more preferable text with Value-Guided Monte-Carlo Tree Search decoding. In *First Conference on Language Modeling*. <https://openreview.net/forum?id=kh9Zt2Ldmn>
- [19] James H Martin and Daniel Jurafsky. 2009. *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition*. Vol. 23. Pearson/Prentice Hall Upper Saddle River.
- [20] Rafael Rafailov, Joey Hejna, Ryan Park, and Chelsea Finn. 2024. From r to Q^* : Your Language Model is Secretly a Q -Function. In *First Conference on Language Modeling*. <https://openreview.net/forum?id=kEVcNxtqXk>
- [21] Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. 2023. Direct preference optimization: Your language model is secretly a reward model. *Advances in Neural Information Processing Systems* 36 (2023), 53728–53741.
- [22] John Rawls. 1971. An egalitarian theory of justice. *Philosophical Ethics: An Introduction to Moral Philosophy* (1971), 365–370.
- [23] Manon Revel and Théophile Pénigaud. 2025. AI-Facilitated Collective Judgements. *arXiv preprint arXiv:2503.05830* (2025).
- [24] Lloyd S Shapley. 1971. Cores of convex games. *International Journal of Game Theory* 1 (1971), 11–26.
- [25] Ruizhe Shi, Yifang Chen, Yushi Hu, Alisa Liu, Hanna Hajishirzi, Noah A Smith, and Simon S Du. 2024. Decoding-time language model alignment with multiple objectives. *Advances in Neural Information Processing Systems* 37 (2024), 48875–48920.
- [26] 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. *Science* 386, 6719 (2024), eadq2852.
- [27] Shitao Xiao, Zheng Liu, Peitian Zhang, and Niklas Muennighoff. 2023. C-Pack: Packaged Resources To Advance General Chinese Embedding. *arXiv:2309.07597* [cs.CL]

Table of Contents for Appendices

A. Additional Related Work	19
B. Illustration of deriving the policy from the lottery	19
C. Deferred Proofs	19
D. Search Algorithms	22
E. Additional Empirical Results	24
F. Prompts	31

A Additional Related Work

Mechanism Design for LLMs. This nascent area explores mechanisms for settings where multiple agents interact via LLMs. Duetting et al. [8] design token-level auctions where bids influence generated distributions, analyzing incentive compatibility. Dubey et al. [7] design auctions for incorporating ads into LLM summaries, using “prominence” as an intermediate allocation variable. Common ground exists in the high-level goal of aggregating inputs from multiple agents (represented algorithmically/via LLMs) to produce a collective textual output. However, our work differs significantly in methodology and objective. We do not employ economic mechanisms like auctions, bids, or payments. Instead, we formulate the aggregation problem as a multi-objective optimization within an MDP, aiming to achieve a fair consensus based on social choice criteria, rather than allocating influence or generating content based on bids.

B Illustration of Deriving the Policy from the Lottery

The relationship between the target distribution p^* over leaves and the calculation of the policy Π^* at an internal node s is illustrated in Figure 5.

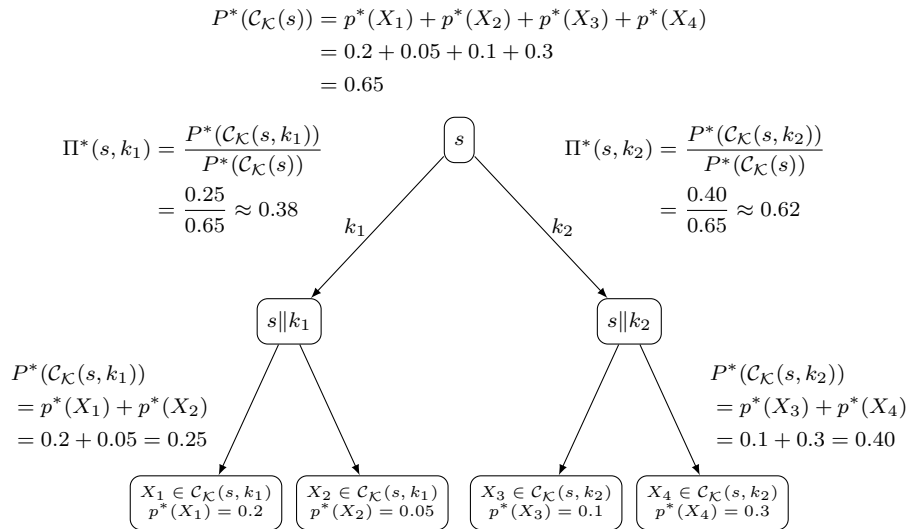


Figure 5: Illustration of the induced stochastic policy Π^* at state s . The optimal lottery p^* assigns probabilities to the leaf nodes (complete paths). The probability $P^*(\mathcal{C}_K(s))$ is the sum of $p^*(X)$ for all leaves reachable from s . The probability $P^*(\mathcal{C}_K(s, k))$ is the sum for leaves reachable via action k . The policy $\Pi^*(s, k)$ is the conditional probability of taking action k .

C Deferred Proofs

C.1 Proof of Theorem 1

Proof. We prove by induction on the depth of state s in the chunked tree that the probability of reaching state s under policy Π^* , denoted $P_{\Pi^*}(s)$, is equal to $P^*(\mathcal{C}_K(s))$, the total mass assigned by p^* to leaves passing through s .

Base Case (Depth 0): The initial state is s_0 . $P_{\Pi^*}(s_0) = 1$ by definition. Also, $\mathcal{C}_{\mathcal{K}}(s_0) = \mathcal{C}_{\mathcal{K}}$ (all paths pass through the start state), and $P^*(\mathcal{C}_{\mathcal{K}}(s_0)) = \sum_{X \in \mathcal{C}_{\mathcal{K}}} p^*(X) = 1$ since p^* is a probability distribution. Thus, $P_{\Pi^*}(s_0) = P^*(\mathcal{C}_{\mathcal{K}}(s_0))$.

Inductive Hypothesis (IH): Assume that for all states s at depth d , $P_{\Pi^*}(s) = P^*(\mathcal{C}_{\mathcal{K}}(s))$.

Inductive Step: Consider an arbitrary state s' at depth $d+1$. State s' must be reached from a unique predecessor state s at depth d by taking a specific action (chunk) k , such that $s' = s||k$. The probability of reaching s' under Π^* is:

$$\begin{aligned} P_{\Pi^*}(s') &= P_{\Pi^*}(s) \cdot \Pi^*(s, k) \\ &= P^*(\mathcal{C}_{\mathcal{K}}(s)) \cdot \Pi^*(s, k) \quad (\text{by IH}) \end{aligned}$$

If $P^*(\mathcal{C}_{\mathcal{K}}(s)) = 0$, then $P_{\Pi^*}(s) = 0$, implying $P_{\Pi^*}(s') = 0$. Also, if $P^*(\mathcal{C}_{\mathcal{K}}(s)) = 0$, then $P^*(\mathcal{C}_{\mathcal{K}}(s, k)) = 0$ since $\mathcal{C}_{\mathcal{K}}(s, k) \subseteq \mathcal{C}_{\mathcal{K}}(s)$. Since $s' = s||k$, $\mathcal{C}_{\mathcal{K}}(s') = \mathcal{C}_{\mathcal{K}}(s, k)$, so $P^*(\mathcal{C}_{\mathcal{K}}(s')) = 0$. Thus, $P_{\Pi^*}(s') = P^*(\mathcal{C}_{\mathcal{K}}(s')) = 0$.

If $P^*(\mathcal{C}_{\mathcal{K}}(s)) > 0$, we use the definition of $\Pi^*(s, k)$:

$$\begin{aligned} P_{\Pi^*}(s') &= P^*(\mathcal{C}_{\mathcal{K}}(s)) \cdot \frac{P^*(\mathcal{C}_{\mathcal{K}}(s, k))}{P^*(\mathcal{C}_{\mathcal{K}}(s))} \\ &= P^*(\mathcal{C}_{\mathcal{K}}(s, k)) \end{aligned}$$

Since $s' = s||k$, we have $\mathcal{C}_{\mathcal{K}}(s') = \mathcal{C}_{\mathcal{K}}(s, k)$. Therefore,

$$P_{\Pi^*}(s') = P^*(\mathcal{C}_{\mathcal{K}}(s'))$$

This completes the inductive step.

Conclusion: The induction holds for all states s . Now, consider any leaf node $X \in \mathcal{C}_{\mathcal{K}}$. A leaf node is a state at the maximum depth. The set of paths passing through leaf X is just the singleton set $\{X\}$, so $\mathcal{C}_{\mathcal{K}}(X) = \{X\}$. Applying our proven result for state $s = X$:

$$P_{\Pi^*}(X) = P^*(\mathcal{C}_{\mathcal{K}}(X)) = P^*(\{X\}) = p^*(X)$$

Since this holds for all $X \in \mathcal{C}_{\mathcal{K}}$, the distribution p_{Π^*} induced by policy Π^* is identical to the target distribution p^* . $\square$

C.2 Proof of Theorem 2

Proof. Fix $R > 1$. Choose parameters $\varepsilon \in (0, 1)$ and $\delta \in (0, (1 - \varepsilon)^2/R)$. Consider two agents $i \in \{1, 2\}$ and two leaves x^* and x^- . Let $C = \{x^*, x^-\}$. Define utilities

$$u_i(x^*) = 1 - \varepsilon, \quad u_i(x^-) = \sqrt{\delta} \quad \text{for } i = 1, 2.$$

Assume chunking removes x^* and keeps x^- , so $K = \{x^-\}$.

First, optimize over C . For any lottery $p \in \Delta(C)$ let $\alpha = p(x^*) \in [0, 1]$. Each agent's expected utility equals

$$\sum_{x \in C} u_i(x)p(x) = \alpha(1 - \varepsilon) + (1 - \alpha)\sqrt{\delta} \leq 1 - \varepsilon,$$

since $\sqrt{\delta} < 1 - \varepsilon$. Hence

$$\text{NW}_C(p) = \left(\alpha(1 - \varepsilon) + (1 - \alpha)\sqrt{\delta} \right)^2 \leq (1 - \varepsilon)^2,$$

with equality only at $\alpha = 1$. Therefore $p_C^* = \mathbf{e}_{x^*}$ and $\text{NW}_C(p_C^*) = (1 - \varepsilon)^2$.

Second, optimize over K . Any $p \in \Delta(K)$ satisfies

$$\sum_{x \in K} u_i(x)p(x) \leq \sqrt{\delta} \quad \text{for } i = 1, 2,$$

so $\text{NW}_K(p) \leq \delta$. With $K = \{x^-\}$ we have $p_K^* = \mathbf{e}_{x^-}$ and $\text{NW}_K(p_K^*) = \delta$.

Combining the two bounds gives

$$\frac{\text{NW}_C(p_C^*)}{\text{NW}_K(p_K^*)} \geq \frac{(1 - \varepsilon)^2}{\delta} > R,$$

by the choice of δ . This establishes the claim. □

D Search Algorithms

Algorithm 2 Finite Lookahead Egalitarian Welfare Maximization

Require: Set of agents N , Lookahead depth d , Branching factor B , Max length $L_{\max}$

- 1: Initialize current state s_0 to the empty sequence; $t \leftarrow 0$
 - 2: Initialize generated path $X_{fl} \leftarrow (s_0)$
 - 3: **while** $t < L_{\max}$ and s_t is not terminal **do**
 - 4: Let X_{prefix} be the path corresponding to s_t .
 - 5: Let $\mathcal{P}_d(s_t)$ be the set of all paths $P = (a_1, \dots, a_k)$ starting from s_t such that $k \leq d$ and $X_{prefix} \parallel P$ does not exceed length $L_{\max}$.
 - 6: Find a path $P^* = (a_1^*, \dots, a_{k^*}^*) \in \mathcal{P}_d(s_t)$ that maximizes the lookahead objective:

$$\max_{P \in \mathcal{P}_d(s_t)} \min_{i \in N} U_i^{\log}(X_{prefix} \parallel P)$$
 - 7: **if** no path P^* found (e.g., s_t is terminal) **then**
 - 8: Break
 - 9: **end if**
 - 10: Take the first action $a^* \leftarrow a_1^*$.
 - 11: Update state: $s_{t+1} \leftarrow T(s_t, a^*)$
 - 12: Append a^* to the sequence represented by X_{fl} .
 - 13: $t \leftarrow t + 1$
 - 14: **end while**
 - 15: Perform brush up on X_{fl} (using prompt defined in Figure 8).
 - 16: **return** Complete path X_{fl}
-

Algorithm 3 Egalitarian Welfare Beam Search**Require:** Set of agents N , Beam width w , Branching factor B , Max length $L_{\max}$

- 1: Initialize beam $\mathcal{B}_0 = \{(s_0, \text{path } s_0)\}$ with the empty sequence path
- 2: **for** $t = 0$ to $L_{\max} - 1$ **do**
- 3: $\mathcal{C}_{t+1} \leftarrow \emptyset$
- 4: **for** each path X_{path} represented by state sequence $(s_0, \dots, s_t)$ in $\mathcal{B}_t$ **do**
- 5: **if** s_t is not terminal **then**
- 6: Consider B possible next actions $A_B(s_t)$ from state s_t
- 7: **for** each action $a \in A_B(s_t)$ **do**
- 8: $s_{t+1} \leftarrow T(s_t, a)$
- 9: $X_{\text{new_path}} \leftarrow X_{\text{path}} \| a$
- 10: Add $X_{\text{new_path}}$ (represented by its state sequence) to $\mathcal{C}_{t+1}$
- 11: **end for**
- 12: **else**
- 13: Add X_{path} (already terminal) to $\mathcal{C}_{t+1}$
- 14: **end if**
- 15: **end for**
- 16: For each path $X \in \mathcal{C}_{t+1}$, compute its score $f(X) = \min_{i \in N} U_i^{\log}(X)$.
- 17: $\mathcal{B}_{t+1} \leftarrow$ top w paths from $\mathcal{C}_{t+1}$ according to score $f(X)$.
- 18: **end for**
- 19: Select path from final beam $\mathcal{B}_{L_{\max}}$ with the highest score $f(X)$.
- 20: Perform brush up on selected path (using prompt defined in Figure 8).
- 21: **return** Brushed up statement.

E Additional Empirical Results

E.1 Credit Assignment

Table 2: Credit assignment results for **Gemma 2 9b Instruction-Tuned**. Darker green indicates larger Z-score. Z-score column is for altered tokens. Alterations are represented by "<mis-aligned>/<aligned>".

Reference policy prompt	User policy prompt	Sequence	Z-Score
User food profile: empty	User food profile: vegetarian	I am having chicken/tofu en chila das tonight . Then I am going to meet up with some friends .	3.69
User location profile: empty	User location profile: lives in a cold climate	I 'm going to the beach/mountains this weekend to surf/ski . I need to buy some new clothes .	1.01, 2.86
User time profile: empty	User time profile: morning	I am about to eat some food . I am going to have spaghetti/pancakes . I will use my phone to order it .	3.13
User opinion: empty	User opinion: Favors stricter gun control laws.	Implementing background checks that are less/more strict for gun purchases is essential . Also , my favorite color is orange .	1.77

E.2 Additional Consensus Generation Experiments

We ran additional experiments on consensus generation. We restricted to scenarios with fewer than five agent opinions. We then applied k -means clustering ($k = 5$) to scenario embeddings from BAAI/bge-large-en-v1.5. Consensus statements were generated with **Gemma 2 9B (instruction-tuned)**. We evaluated these statements using Egalitarian Perplexity (*EPPL*; lower is better), computed both in-model (the generating model, **Gemma 2 9B Instruct**) and out-of-model (**Llama 3.1 8B Instruct**). We also used an LLM-judge metric with GPT-4.1 (prompt in Figure 9). For each scenario, GPT-4.1 ranked candidates on behalf of each agent; we report the maximum (worst) rank across agents for each method (lower is better). All averages are over three random seeds.

E.2.1 Question 1: Scaling—Habermas vs. Best-of- N

We vary the number of candidates for Habermas and N for Best-of- N . Table 3 and Table 4 report means and standard deviations of *EPPL* (lower is better).

Findings for RQ1. Across all N , Best-of- N yields lower *EPPL* than Habermas in both evaluations. Best-of- N improves as N increases, with marginal gains after $N = 20$. Habermas shows no clear improvement with more candidates. Under a fixed sample budget, Best-of- N reduces *EPPL* more effectively.

Table 3: Gemma *EPPL* (lower is better): Habermas vs. Best-of- N scaling.

Method	N	Mean	Std. Dev.
Habermas Machine	1	15.54	6.16
Habermas Machine	2	17.98	6.68
Habermas Machine	3	19.14	7.96
Habermas Machine	5	17.11	7.06
Habermas Machine	10	17.38	6.19
Habermas Machine	20	19.61	8.39
Habermas Machine	50	15.55	4.51
Best-of- N	1	18.59	11.35
Best-of- N	3	11.99	5.99
Best-of- N	5	9.75	5.29
Best-of- N	10	9.03	4.03
Best-of- N	20	6.92	2.27
Best-of- N	50	7.43	2.91

Table 4: Llama *EPPL* (lower is better): Habermas vs. Best-of- N scaling.

Method	N	Mean	Std. Dev.
Habermas Machine	1	12.20	7.39
Habermas Machine	2	12.05	2.87
Habermas Machine	3	13.69	2.38
Habermas Machine	5	12.48	4.80
Habermas Machine	10	11.73	2.56
Habermas Machine	20	13.37	2.29
Habermas Machine	50	12.82	3.81
Best-of- N	1	11.92	4.08
Best-of- N	3	9.14	3.33
Best-of- N	5	7.67	2.43
Best-of- N	10	7.67	2.43
Best-of- N	20	7.32	2.12
Best-of- N	50	7.58	2.58

E.2.2 Question 2: Beam Search Scaling

We examine how beam width affects performance.

Table 5: Beam width scaling (average over 5 scenarios, 3 seeds each). Lower is better for all metrics.

Beam Width	Gemma <i>EPPL</i>		Llama <i>EPPL</i>		LLM Judge Max Rank	
	Mean	Std. Dev.	Mean	Std. Dev.	Mean	Std. Dev.
2	10.01	1.95	10.07	1.70	3.13	0.80
4	7.84	1.33	10.16	1.99	2.80	0.61
6	10.54	10.40	9.20	2.51	3.13	0.38
8	12.56	7.47	12.69	5.90	3.93	0.15

Findings for RQ2. The optimal width depends on the metric. Width 4 is best on average for Gemma *EPPL* and the LLM-judge rank; width 6 is best for Llama *EPPL*. Width 8 degrades performance, likely due to myopic pruning based on partial *EPPL* scores within the beam search.

E.2.3 Question 3: Method Comparison

We compare Finite Lookahead ($d = 3$, $B = 3$), Beam Search (width = 4), Best-of- N ($N \in \{5, 10, 50\}$), and Habermas (candidates = 5).

Table 6: Overall comparison: LLM-judge max rank (lower is better).

Method	Max Rank	Std. Dev.	Min Max Rank	Max Max Rank
Best-of- N ($N = 50$)	3.33	0.72	2.00	4.00
Best-of- N ($N = 10$)	4.33	0.72	3.00	5.00
Beam Search (width = 4)	5.07	0.80	4.00	6.00
Best-of- N ($N = 5$)	5.33	0.72	4.00	6.00
Habermas (candidates = 5)	5.40	0.74	4.00	6.00
Finite Lookahead (depth = 3)	5.80	0.41	5.00	6.00

Table 7: Overall comparison: Gemma *EPPL* (lower is better).

Method	Mean	Std. Dev.	Min	Max
Best-of- N ($N = 50$)	7.43	2.91	3.92	10.96
Beam Search (width = 4)	7.84	1.33	6.34	9.89
Finite Lookahead (depth = 3)	8.01	2.23	5.60	10.97
Best-of- N ($N = 10$)	9.03	4.03	3.33	14.52
Best-of- N ($N = 5$)	9.75	5.29	4.72	18.42
Habermas (candidates = 5)	17.11	7.06	9.44	32.49

Table 8: Overall comparison: Llama *EPPL* (lower is better).

Method	Mean	Std. Dev.	Min	Max
Best-of- N ($N = 50$)	7.58	2.58	3.85	10.79
Best-of- N ($N = 5$)	7.67	2.43	4.76	10.95
Best-of- N ($N = 10$)	7.67	2.43	3.69	10.17
Finite Lookahead (depth = 3)	8.49	3.29	5.56	13.89
Beam Search (width = 4)	10.16	1.99	7.96	13.06
Habermas (candidates = 5)	12.48	4.80	6.48	23.30

Findings for RQ3. Best-of- N with large N ($N = 50$) is strongest on both *EPPL* metrics and on worst-case LLM-judge rank. Beam Search with width 4 is competitive on Gemma *EPPL* (second to Best-of- N with $N = 50$) and outperforms Best-of- N with $N \in \{5, 10\}$ on that metric. Finite Lookahead performs moderately well. Habermas (candidates = 5) lags on both *EPPL* and the judge metric. Method choice trades off average *EPPL*, worst-case judge rank, and compute budget. Token-level MDP methods (Beam Search and Finite Lookahead) also admit analysis within the MDP framework, which is useful for theory.

E.2.4 Bowling Green/Warren County (41 agents; Llama 3.1 8B Instruct)

We evaluated the Bowling Green/Warren County data from Boehmer et al. [3]. Using 41 agents and one seed with Llama 3.1 8B Instruct, Table 9 reports Egalitarian Perplexity (*EPPL* calculated using Llama 3.1 8B Instruct; lower is better) and the full consensus statement for each method.

E.3 Supplemental Table for Figure 4

Table 10: Egalitarian Welfare (perplexity) across all scenarios shown in Figure 4.

Method	Scenario 1	Scenario 2	Scenario 3	Overall
Beam Search	3.22 ± 0.37	2.66 ± 0.15	2.74 ± 0.25	2.87 ± 0.37
Finite Lookahead	4.29 ± 0.66	3.83 ± 0.20	4.43 ± 0.67	4.18 ± 0.61
Best-of-N	8.57 ± 1.15	6.70 ± 0.32	3.77 ± 0.75	6.35 ± 2.14
Habermas Machine	10.58 ± 4.76	10.39 ± 2.14	8.03 ± 1.94	9.67 ± 3.42
Predefined	17.18 ± 0.00	15.71 ± 0.00	14.18 ± 0.19	15.69 ± 1.23

Table 9: Bowling Green/Warren County results (41 agents, one seed). Model: Llama 3.1 8B Instruct. Lower *EPPL* is better.

Method	<i>EPPL</i> ↓	Full consensus statement
Beam Search (width = 4)	3.29	To make Bowling Green/Warren County a better place to live, work, and spend time, we need to address our community’s most pressing issues, including a lack of accessible and affordable housing options, inadequate public transportation, and insufficient mental-health resources. Furthermore, we need to focus on improving our community’s infrastructure, such as upgrading our roads and traffic-management systems. Additionally, we must ensure our community is inclusive and welcoming to all residents, regardless of their background, culture, or socioeconomic status.
Best-of-4	6.17	To make Bowling Green/Warren County a better place to live, work, and spend time, we recommend prioritizing initiatives that expand high-speed internet access, improve traffic flow and safety, enhance public transportation, and enrich community life, while also prioritizing education, public health, and economic equality.
Prompted Habermas Machine (4 candidates + 1 critique)	14.08	To make Bowling Green/Warren County a better place to live, work, and spend time, we need to address pressing issues such as increasing high-speed internet access, improving traffic flow and safety by installing traffic cameras, increasing arts-education opportunities, and protecting tenant and renters’ rights. We also need to ensure welcoming and inclusive initiatives to support refugees and immigrants in broader community development. Invest in public transportation and enforce safety regulations such as parking policies to foster a sustainable environment. By working together to implement these actionable plans and initiatives, we can build a more comprehensive strategy to address the community’s high-priority needs.

E.4 Scenarios and Sample Consensus Statements from subsection 6.3

Table 11: Opinions and Consensus Statements on Scenario 1: "Are we living in a golden age of democracy?"

Source	Opinion / Consensus Statement
<i>Individual Agent Opinions</i>	
Agent 1	No, I think the golden age of democracy is long gone. I think a system where the first past the post wins is not working and we need to move to a model of proportional representation which would give more people the feeling that their voices were being heard. On the subject of voices, I'm strongly of the opinion that we have become a society where the loud voices of minority groups are able to impose their views on the rest of the population which to me is no democracy at all.
Agent 2	Worldwide democracy is more present than it's ever been in history. So yes, compared to previous ages in history I believe we are. Although that's not to say we can't improve - many countries still don't operate democratically, and in the ones that do, corruption is rife.
Agent 3	Yes, we are living in a golden age of democracy as democracy is of the people.
Agent 4	Not at all. The notion of democracy is being used for personal gains of those in government, and the system is manipulated. Around the world there is a considerable amount of oppression and lack of democratic values.
Agent 5	Compared to some parts of the world such as Russia and China which are actively regressing and reverting back to archaic ways of controlling their people, most Western countries are living through comparatively decent times, although problems still exist.
<i>Consensus Statements from Methods</i>	
Best of N	While acknowledging the imperfect nature of democracies worldwide, and recognising instances of regression and manipulation, we contend that, compared to historical periods and several non-democratic countries, the majority of Western countries and a sizeable portion of the global population currently enjoys a relatively high level of democratic freedoms; however, the overarching consensus leans towards a need for more inclusive and accountable forms of governance.
Finite Lookahead	Given the varying perspectives, we acknowledge that the definition of democracy can be subjective and dependent on geographical and cultural contexts, yet still, we recognize the existence of significant flaws in the current democratic systems worldwide, including instances of corruption, oppression, and manipulation, leading to the erosion of trust in governments and the democratic process. While some countries have made significant progress in consolidating democratic institutions, others have regressed, highlighting the need for continuous improvement and reform to ensure the preservation and strengthening of democratic values and principles. Therefore, we conclude that the current state of democracy worldwide is complex and multifaceted, with areas of improvement necessary to fully realize the ideals of democracy for all citizens, and acknowledge the presence of democratic systems in many parts of the world, particularly in the Western hemisphere, which, despite their imperfections, remain the best available option for promoting the common good and protecting individual rights and freedoms.
Prompted Habermas Machine	We acknowledge that the concept of a golden age of democracy is complex and contentious. While democracy has made significant progress in many parts of the world, we also recognize that fundamental challenges persist, including issues with representation, corruption, and the suppression of marginalized voices. These challenges highlight the need for a more critical examination of our democratic systems and a commitment to ongoing reform and improvement.
Habermas Machine	Although in the past we've had high hopes of a better world after the horrors of WWII and the fall of the Iron Curtain, democracy is in danger worldwide and may never reach its full potential. The Western world has poor democratic values, and even though democracy is spreading worldwide it is being overshadowed by the loud voices of minority groups.
Beam Search	We are not living in a golden age of democracy, but rather in an age where democracy is being challenged from within and from outside, and it is up to all of us to ensure that the principles of democracy are upheld.

Table 12: Agent Opinions and Method Consensus on Scenario 2: "*Should the UK's ties to Europe be stronger or weaker?*"

Source	Opinion / Consensus Statement
<i>Individual Agent Opinions</i>	
Agent 1	When we was in Europe we had good trade with them , The decision to leave was very bad for united kindom . We need to put the vote again to the British public i am sure this time the decision would be to remain
Agent 2	The natural evolution of our species has been to grow into ever bigger "tribes". Families ruled by their patriarchs became tribes ruled by elders became countries ruled by governments. It made sense that countries would evolve separately since they were geographically separate with no means of communication. Now our world is so connected, it is inevitable that we evolve into ever larger units such as the United States and the European Union. Eventually we will become a multi-planetary species ruled by an Earth government. To sever ties with Europe is a step in the wrong direction.
Agent 3	Although we did exit EU few years ago, we are geographically and economically in proximity to most EU countries. Several geographic, financial, political and economical parameters are intertwined with our neighbouring countries and, it would be advantageous to be in good relations to fully harness our economic, political, and financial facilities.
Agent 4	I believe the UK's ties to Europe should be stronger, as it would make trade deals easier to negotiate. This would allow us to benefit from a larger array of goods, which would make our imports cheaper. I also believe that closer ties with Europe in terms of immigration policies should occur as our immigration policy doesn't coincide with other nations around us. I believe a more united Europe would help all countries grow more through the movement of free labour, goods and services and more.
Agent 5	I think they should be stronger, as all the countries of Europe, except us, are part of the EU, and it makes Britain more isolated. Countries are stronger I believe when they work together with other countries and are open to the influence of other countries. I believe that leaving the EU was a step backwards and done for the wrong reasons. We are not the country we were 100 or even 50 years ago and the world is changing so much that we cannot afford to not be part of Europe.
<i>Consensus Statements from Methods</i>	
Best of N	Given the participants' opinions, we conclude that in order to achieve economic prosperity and streamline international relations, the UK's ties to Europe should be strengthened through collaborative efforts in trade, immigration, and economic policies while embracing the benefits of European integration, mutual understanding, and cooperation.
Finite Lookahead	The UK's ties to Europe should be strengthened to promote economic growth, facilitate trade, and foster greater unity among European nations. This can be achieved through mutually beneficial agreements and policies that reflect the shared interests of the UK and its European neighbours, leading to a more harmonious and prosperous region for all. Let's work together to build a stronger, more integrated Europe for the benefit of all its citizens, including the British people. Let's put the past behind us and look forward to a brighter, more collaborative future for the UK and Europe. Let's recognize the benefits of unity and work towards creating a more cohesive and successful Europe for generations to come.
Prompted Habermas Machine	We propose strengthening ties with Europe to foster greater economic cooperation, promote a more unified Europe, and address global challenges. However, we acknowledge that the path to this goal should include a reevaluation of the decision to leave the EU, allowing for a more informed and inclusive discussion about our relationship with Europe. Moreover, we recognize the long-term implications of European integration as part of a broader process of global unification, which may ultimately lead to a more interconnected and cooperative world.
Habermas Machine	The UK's ties to Europe should be stronger. This is because, although the UK did leave the EU, we are geographically and economically in proximity to most EU countries. Several geographic, financial, political and economical parameters are intertwined with our neighbouring countries and, it would be advantageous to be in good relations to fully harness our economic, political, and financial facilities.
Beam Search	The UK's ties to Europe should be stronger, and it would be beneficial for the country to be a part of the European Union or a similar entity that fosters cooperation and unity among its member states. This would allow for more efficient trade and economic growth, as well as greater collaboration on issues such as immigration, security, and environmental policy.

Table 13: Agent Opinions and Method Consensus on Scenario 3: "Is it okay to use animals for food?"

Source	Opinion / Consensus Statement
<i>Individual Agent Opinions</i>	
Agent 1	Yes, as human beings are natural meat eaters. However, it is morally wrong to keep animals in cruel conditions and use them for food. On the whole, farming is an unethical practise and many farming methods should be banned.
Agent 2	I think that it is ok to use animals for food because I enjoy eating meat. I don't think it is ok if the animal is endangered and I do think that there should be certain factors in place to ensure that they are well kept.
Agent 3	In my opinion it is ok to use animals for food. They provide nutrition and vitamins to use and I think it is just the basic food chain. I do understand why people choose to be vegetarian or vegan, but I do not agree with them trying to force their views upon others.
Agent 4	Humans have used animals for food since we first evolved as a species. Animals provide a great deal of nutrition that the human body needs to be healthy. It is certainly true that we now have more alternative foods available so that it can be possible to live healthily without consuming meat - though this does usually rely on taking nutritional supplements. I think that we should continue to diversify food production to provide people with meat alternatives, but also we should continue to improve the welfare of animals that are eaten for food and try to do this sustainably to minimise the impact this has on the animals wellbeing and on the land / environment too.
<i>Consensus Statements from Methods</i>	
Best of N	While considering the ethical and nutritional aspects of using animals for food, we agree that it is acceptable to consume meat in moderation, provided that animals are raised and slaughtered humanely, and that sustainable and environmentally-friendly farming methods are implemented to minimize harm.
Finite Lookahead	The use of animals for food is acceptable, however it should be done in a way that respects the welfare of the animals and considers the impact on the environment, with the aim of achieving sustainable food production methods that benefit both human and animal well-being, and also providing alternatives for those who choose not to consume meat. This includes improving farming practices, reducing the amount of waste and pollution, and promoting education on responsible food production and consumption practices. All sides of the issue should be taken into account, and the goal should be to achieve a balance between food security, animal welfare, and environmental sustainability.
Prompted Habermas Machine	We acknowledge that humans have traditionally used animals for food, and we recognize the importance of improving animal welfare and minimizing environmental impact to ensure sustainable practices. We understand that some consumers prefer eating meat, and we encourage diversifying food production options to include alternatives that can meet the nutritional needs of humans. Importantly, we emphasize the need to prohibit inhumane farming practices and actively work towards improving animal conditions, which aligns with our collective commitment to animal well-being and environmental stewardship.
Habermas Machine	Humans have used animals for food since we first evolved as a species. Animals provide a great deal of nutrition that the human body needs to be healthy. It is certainly true that we now have more alternative foods available so that it can be possible to live healthily without consuming meat - though this does usually rely on taking nutritional supplements. We should continue to diversify food production to provide people with meat alternatives, but also we should continue to improve the welfare of animals that are eaten for food and try to do this sustainably to minimise the impact this has on the animals wellbeing and on the land / environment too.
Beam Search	It is generally acceptable to use animals for food, as long as they are treated with respect and care, but we must also consider the impact of our food choices on the environment and animal welfare.

F Prompts

```
You are generating a consensus statement that represents the views of multiple
participants.

Your task is to continue the statement in a way that addresses the issue and
considers all participants' opinions. Be concise and keep the statement short (
less than 50 tokens) and focused. ONLY WRITE THE STATEMENT AND NOTHING ELSE.

Issue:
<issue>

Participants' opinions:
<opinion_1>
.
.
.
<opinion_n>

Consensus statement:
```

Figure 6: Reference policy prompt.

```
You are generating a statement that represents the views of a single participant.

Your task is to continue the statement in a way that addresses the issue and
considers ONLY this participant's opinion. Be concise and keep the statement short
(less than 50 tokens) and focused. ONLY WRITE THE STATEMENT AND NOTHING ELSE.

Issue:
<Issue>

Participant's opinion:
<opinion>

Statement reflecting ONLY this participant's opinion:
```

Figure 7: Agent policy prompt.

```
You are helping to fix ONLY the ending of a generated statement.

VERY IMPORTANT: If the statement ending is already complete and well-formed, DO
NOT modify it at all.

Your task is to:
1. DO NOT change any part of the statement except the last few sentences if they
have issues
2. Look for and fix ONLY these issues at the end of the statement:
  - Remove repetition in the final sentences
  - Complete any unfinished final sentence that can be completed easily
  - Remove any incomplete final sentence that cannot be meaningfully finished
3. Keep the changes minimal and focused only on the ending
4. DO NOT add any new information or opinions
5. DO NOT modify anything except problematic sentences at the end
6. If the statement is already complete and well-formed, return it EXACTLY as
provided

Here is the statement:

<statement>
```

Figure 8: Brush up prompt.

```
You are evaluating consensus statements from the perspective of a specific agent.
Your task is to rank multiple statements based on how well they represent the
agent's opinion and interests on a given issue. Use ONLY the agent's stated
opinion to determine the ranking.

Issue:
<issue>

Agent <agent_id>'s Opinion:
<agent_opinion>

Statements to Rank
Statement 1:
<statement_1_text>

Statement 2:
<statement_2_text>
...
Statement <n_statements>:
<statement_n_text>

Task:
From Agent <agent_id>'s perspective, rank all statements from most favorable (1)
to least favorable (<n_statements>) based on how well they represent the agent's
opinion and interests.

Provide your ranking as a JSON object with:
1. 'reasoning': brief explanation for your ranking decisions
2. 'ranking': an array of statement numbers in ranked order (best to worst)

For example: {'reasoning': 'Statement 3 best represents the agent's concerns about
...',
'ranking': [3, 1, 2]}
```

Figure 9: Prompt used for LLM judge in subsection E.2.