

# LLM-Based Simulation of Multilateral Climate Negotiations: A Multi-Agent Framework for Computational Climate Diplomacy

Ziqi Wei (ziquwei@mit.edu)<sup>1</sup>, Anna Papp<sup>2</sup>

<sup>1</sup>Department of Civil and Environmental Engineering, MIT, <sup>2</sup>Department of Economics, MIT



## Background & Motivation

### Problem

- Multilateral climate negotiations under the UNFCCC involve many country coalitions bargaining over specific options, but:
- Game-theoretic models ignore textual and procedural details;
- Qualitative case studies cannot be replicated, or scaled;
- No computational platform for counterfactual experiments.

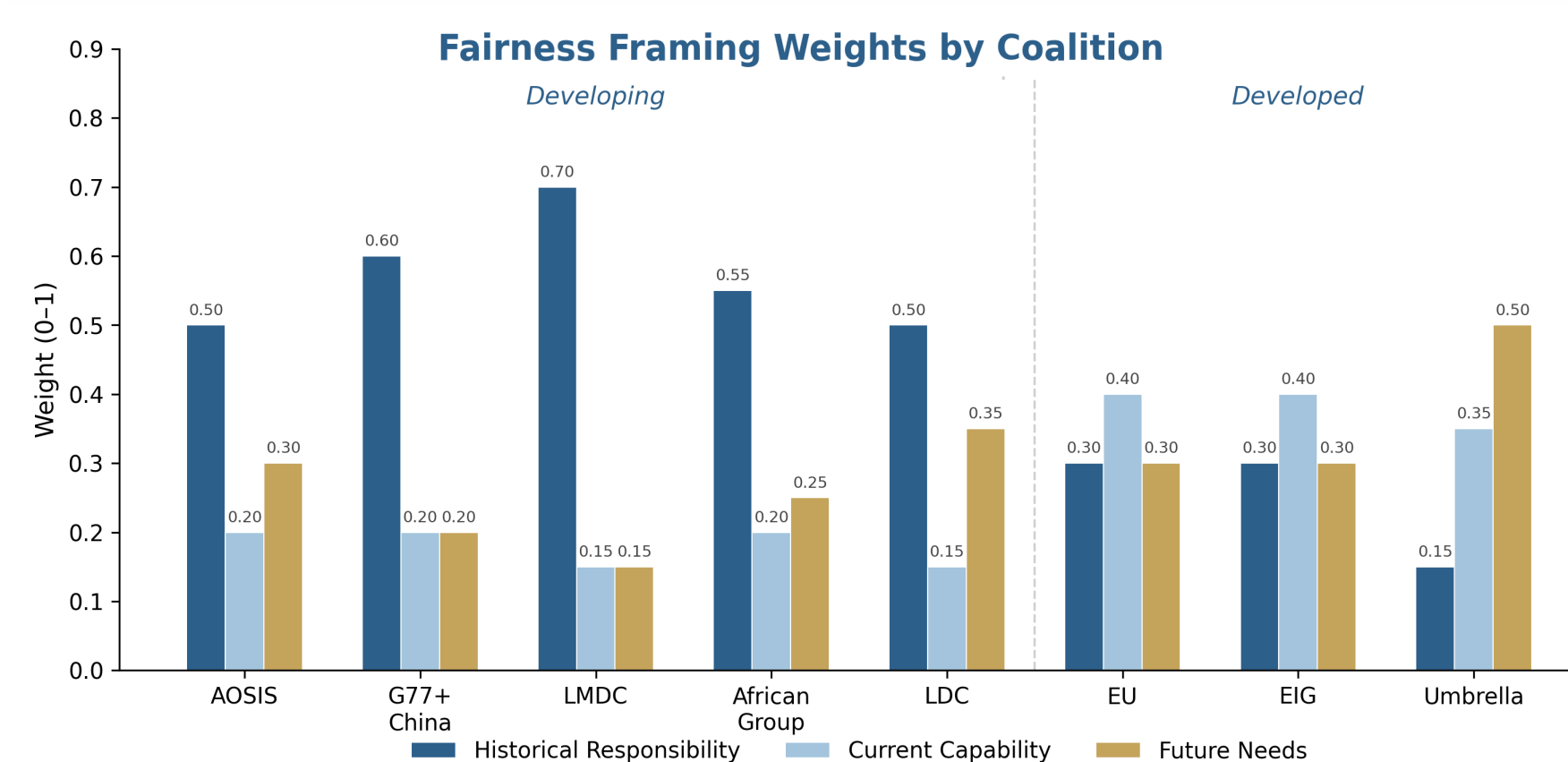
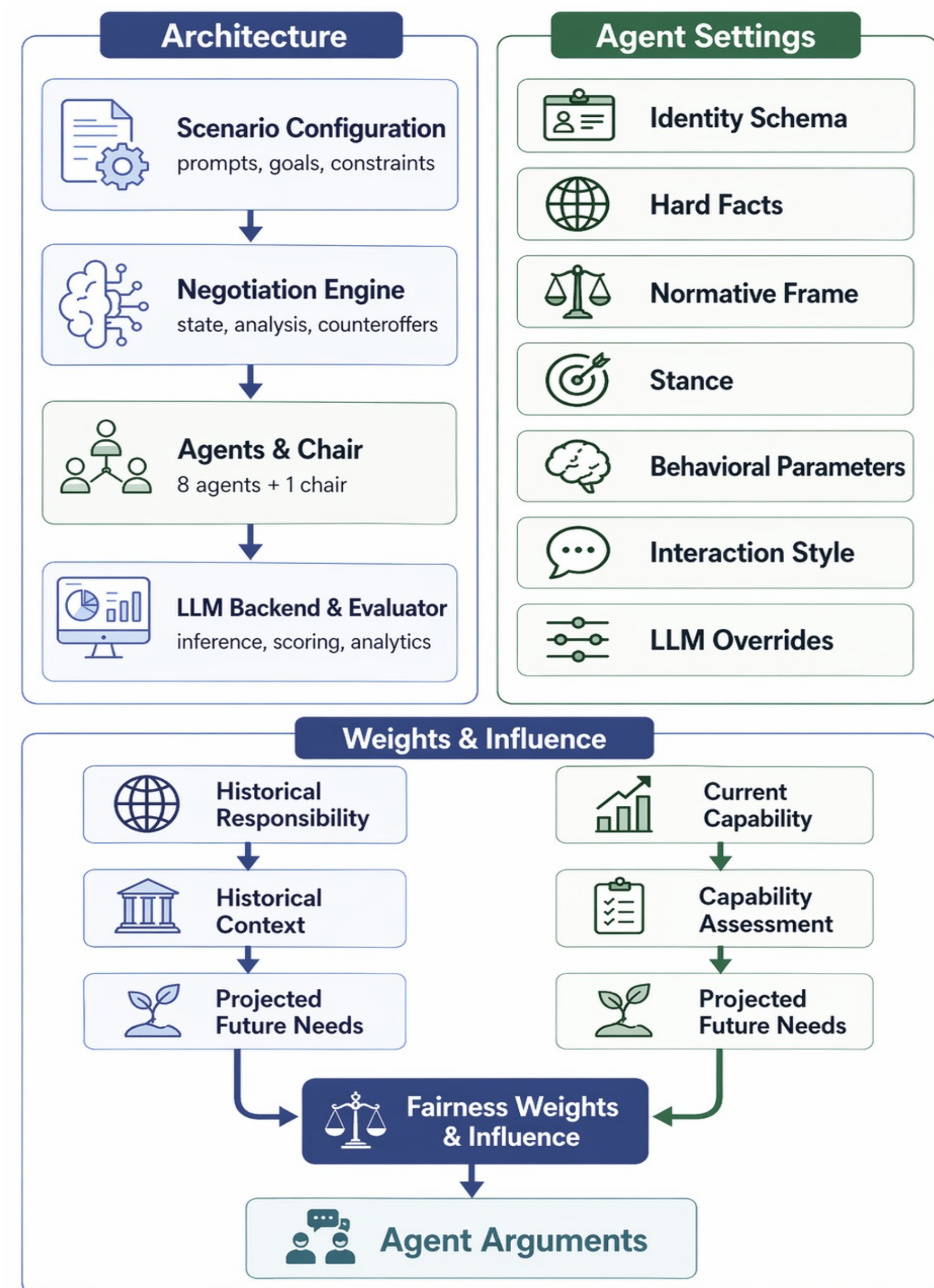
### Why LLM Agents

- Maintain stable, differentiated identities with system prompts;
- Respond to bridging text produced by a neutral Chair agent;
- Sustain negotiations without forgetting earlier rounds.

### Research Questions

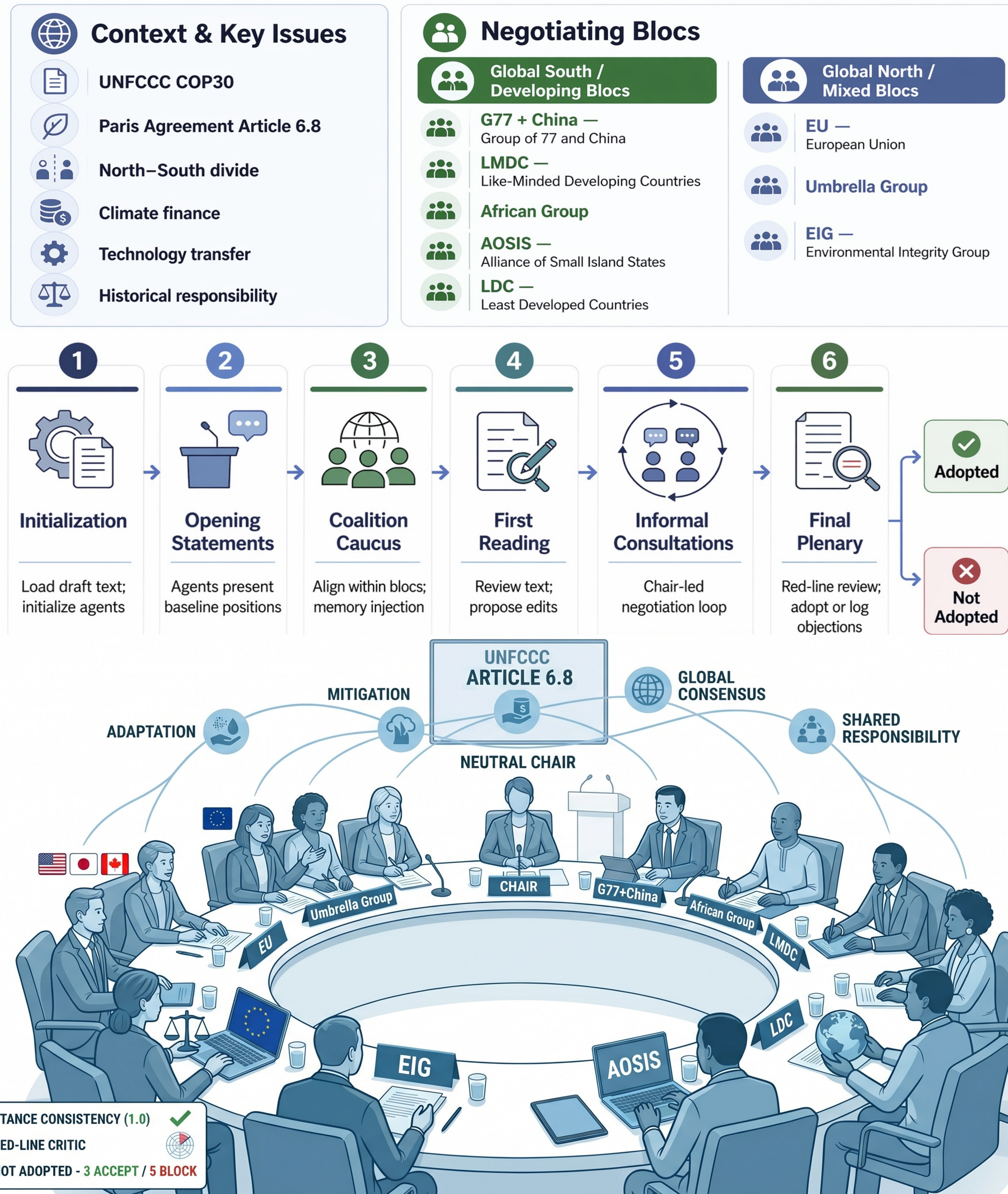
- Can LLM **agents sustain coherent, differentiated positions** reflecting real-world coalition dynamics across extended negotiations?
- Can the Chair agent **synthesize treaty-quality text** comparable to actual COP decisions?
- Does **agent ablation (removing one coalition)** produce **politically interpretable effects** aligning with expectations?

## Model Architecture



- Simulate UNFCCC Article 6.8 negotiations with 8 LLM-driven blocs, producing treaty text near-identical to real COP outcomes.
- Reproduce high-fidelity treaty language that fails politically with real deadlocks.
- Reveal through ablation that text quality and consensus are structurally opposed: vulnerable blocs enrich the treaty while G77's coalition pressure drives the collective veto.
- Extend to multi-article negotiations, add coalition side-channels, and test other LLM agents.

## Experiment Workflow

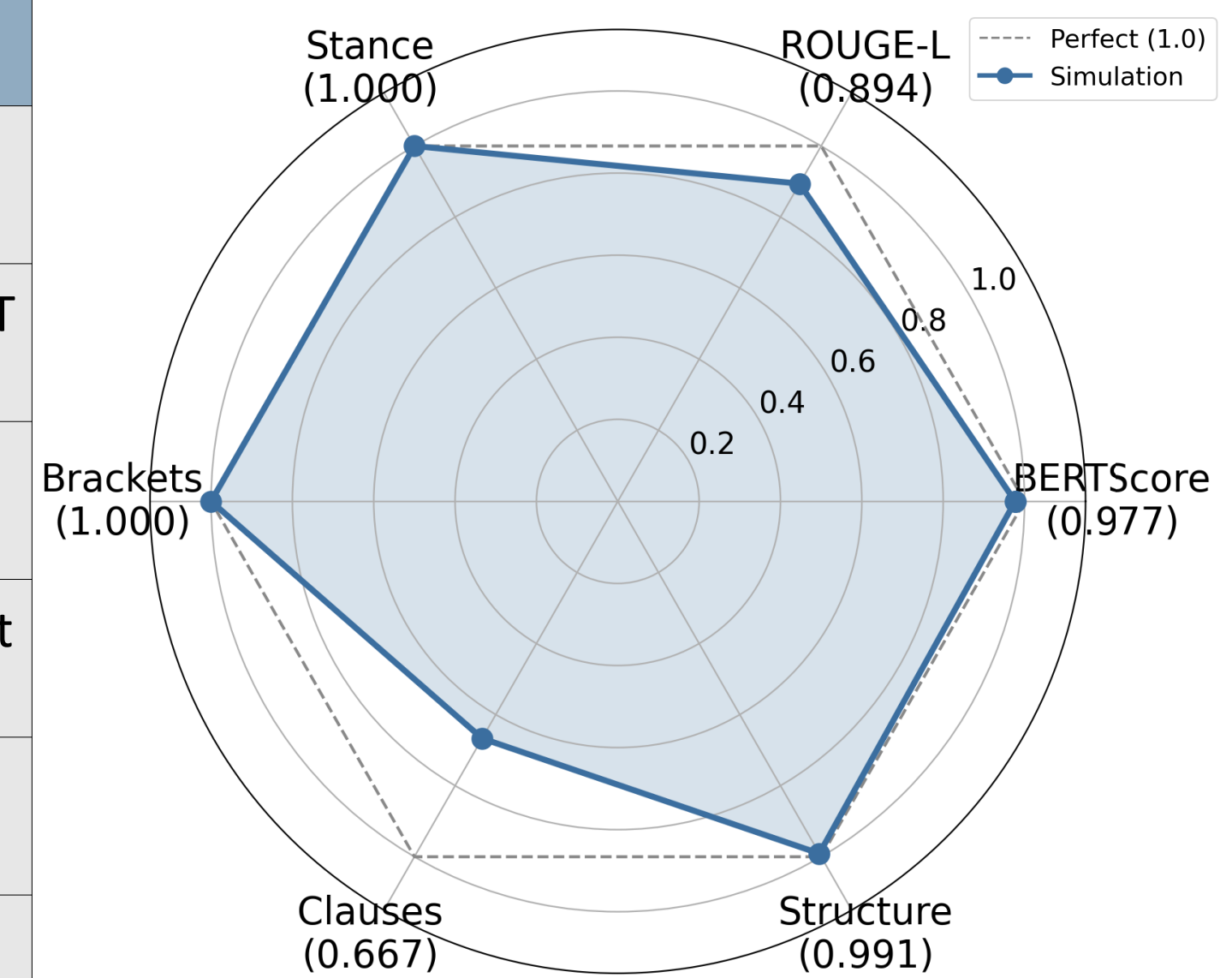


## Results

### Baseline Profile

**Outcome** Only 3 developed-world-aligned blocs accepted the final text while 5 developing-country blocs blocked, due to the absence of CBDR language and inadequate means-of-implementation provisions, reproducing the North-South veto dynamic in real negotiations.

| Metric Definitions    |                                                    |
|-----------------------|----------------------------------------------------|
| Metric                | What it Measures                                   |
| ROUGE-1/2/L           | Lexical n-gram overlap with reference              |
| BERTScore             | Semantic similarity via contextual BERT embeddings |
| Key Clause Coverage   | Essential policy elements present in agreed text   |
| Structural Similarity | Paragraph/item/word-count alignment with reference |
| Bracket Resolution    | Fraction of paragraphs fully de-bracketed          |
| Stance Consistency    | Red-line maintenance across agent statements       |

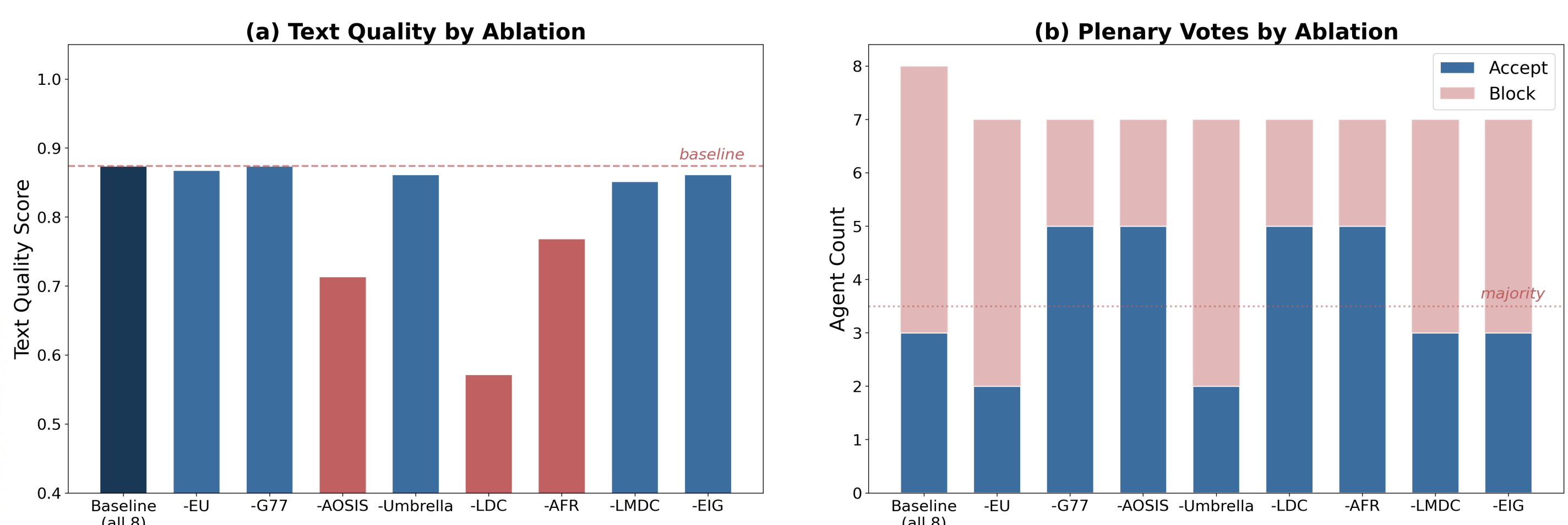


**Evaluation** The simulation achieves near-perfect scores on textual dimensions but falls short on Key Clauses (shown in above metrics), revealing that LLMs replicate the form and style of treaty language far more easily than the substantive policy content that defines real agreements.

### Ablation Findings

#### Text quality vs. consensus

- In Fig (a) and (b), removing LDC or AOSIS improves acceptance but collapses text quality, which means the actors who enrich the treaty are not the actors who block it.
- G77+China is the only bloc whose removal achieves majority acceptance without degrading text quality; developed-country ablations have negligible impact on either metric.



#### Structural voting dynamics

- In Fig (c), EU and EIG always accept; G77 and LMDC always block; removing G77 flips other developing blocs to accept, exposing coalition pressure as the driver of collective opposition.
- In Fig (d), developing-country ablations achieve majority at the cost of substantive depth, which means the baseline produces quality text that fails politically, mirroring real equity deadlocks.

