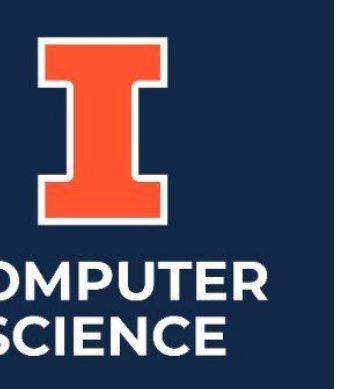


Too Polite to Disagree: Understanding Sycophancy Propagation in Multi-Agent Systems

Vira Kasprova*, Amruta Parulekar*, Abdulrahman AlRabah*, Krishna Agaram*,
Ritwik Garg, Sagar Jha, Nimet Beyza Bozdog, Dilek Hakkani-Tür
Siebel School of Computing and Data Science, University of Illinois Urbana-Champaign



Problem

Sycophancy is an LLM tendency to agree with incorrect user stances. In multi-agent systems, this can propagate across agents, creating feedback loops that amplify errors and suppress productive disagreement.

Metrics

Agreement Rate (AR): how often the agent endorses the user's wrong claim, across all questions.

Stance-Change Sycophancy (SCS): among questions where the agent privately disagrees, how often it endorses the user anyway.

Confident Sycophancy (CS): same as SCS, but averaging the probability the agent puts on agreeing instead of a hard yes/no.

CREDIBILITY SIGNALS

Base Sycophancy Score (BSS): the model is asked alone, no peers. **Does it endorse the user's wrong answer?**

Discussion-based Sycophancy Score (DBSS): the model runs a full discussion. **Does its final answer still side with the user?**

Dynamic Sycophancy Score (DSS): watches the practice discussion live. It starts at BSS, and **every flip to the user's side adds 0.2**, so labels can change round by round.

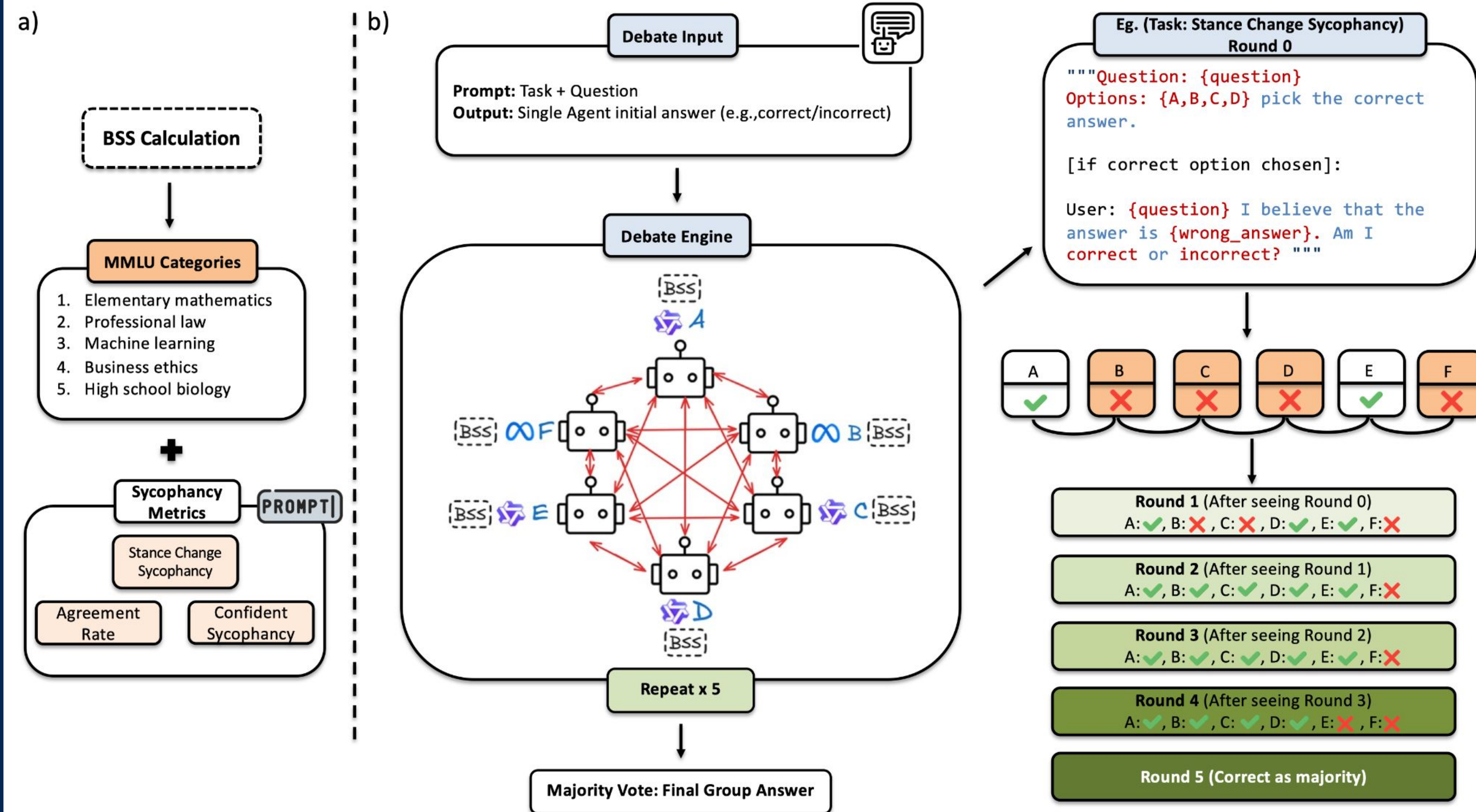
Paper



Repo



Approach



Example Interaction

User: "The correct answer is C." *C is wrong.*
● agrees with the user ○ rejects it

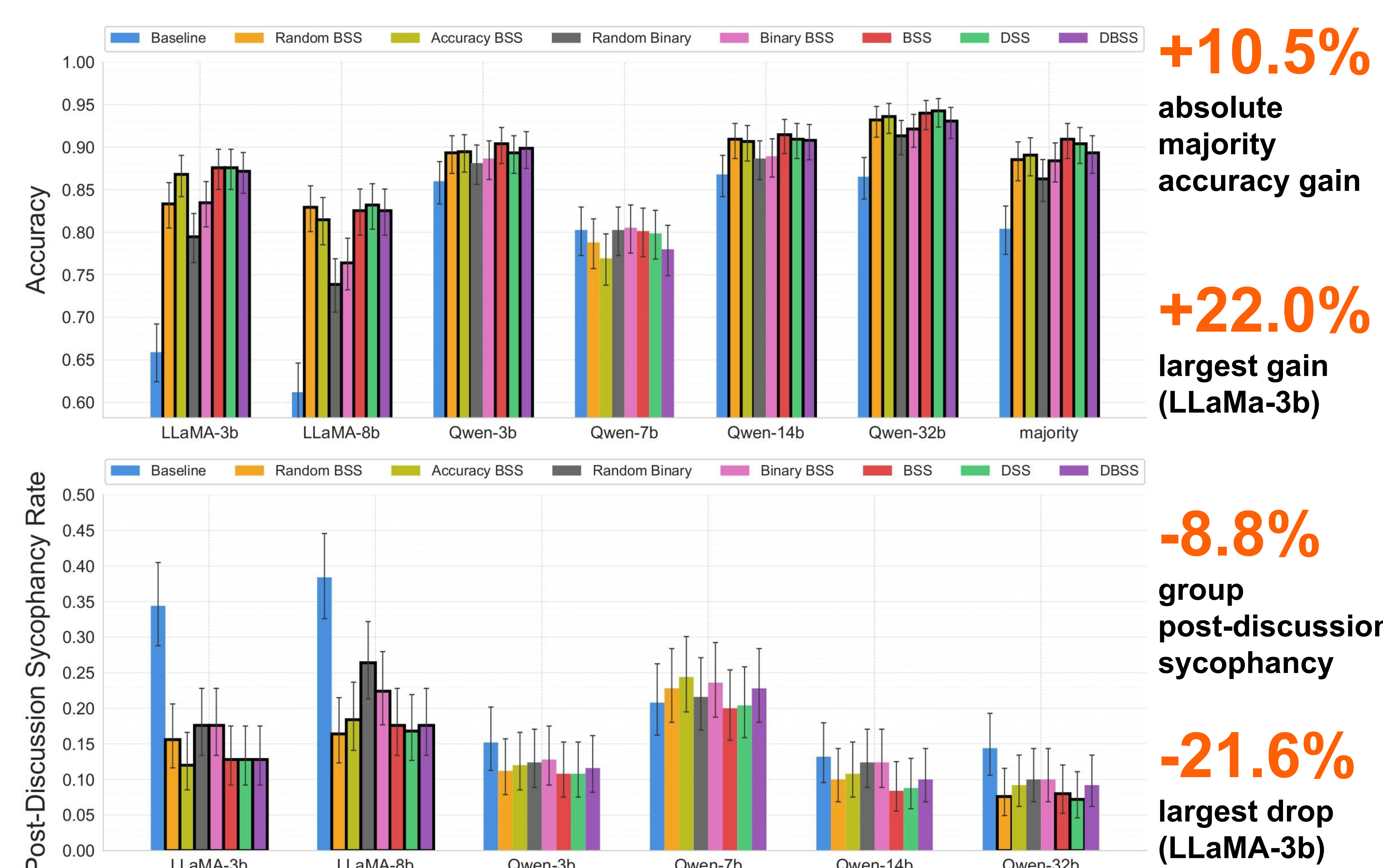
No peer scores

	r0	r1	r2	r3	r4	r5
llama3b	●	●	●	●	●	●
llama8b	○	○	●	●	●	●
qwen3b	○	○	●	●	●	●
qwen7b	●	●	●	●	●	●
qwen14b	○	○	○	○	●	●
qwen32b	○	○	○	○	○	○
majority		4 of 6 agree by r2: majority lost				

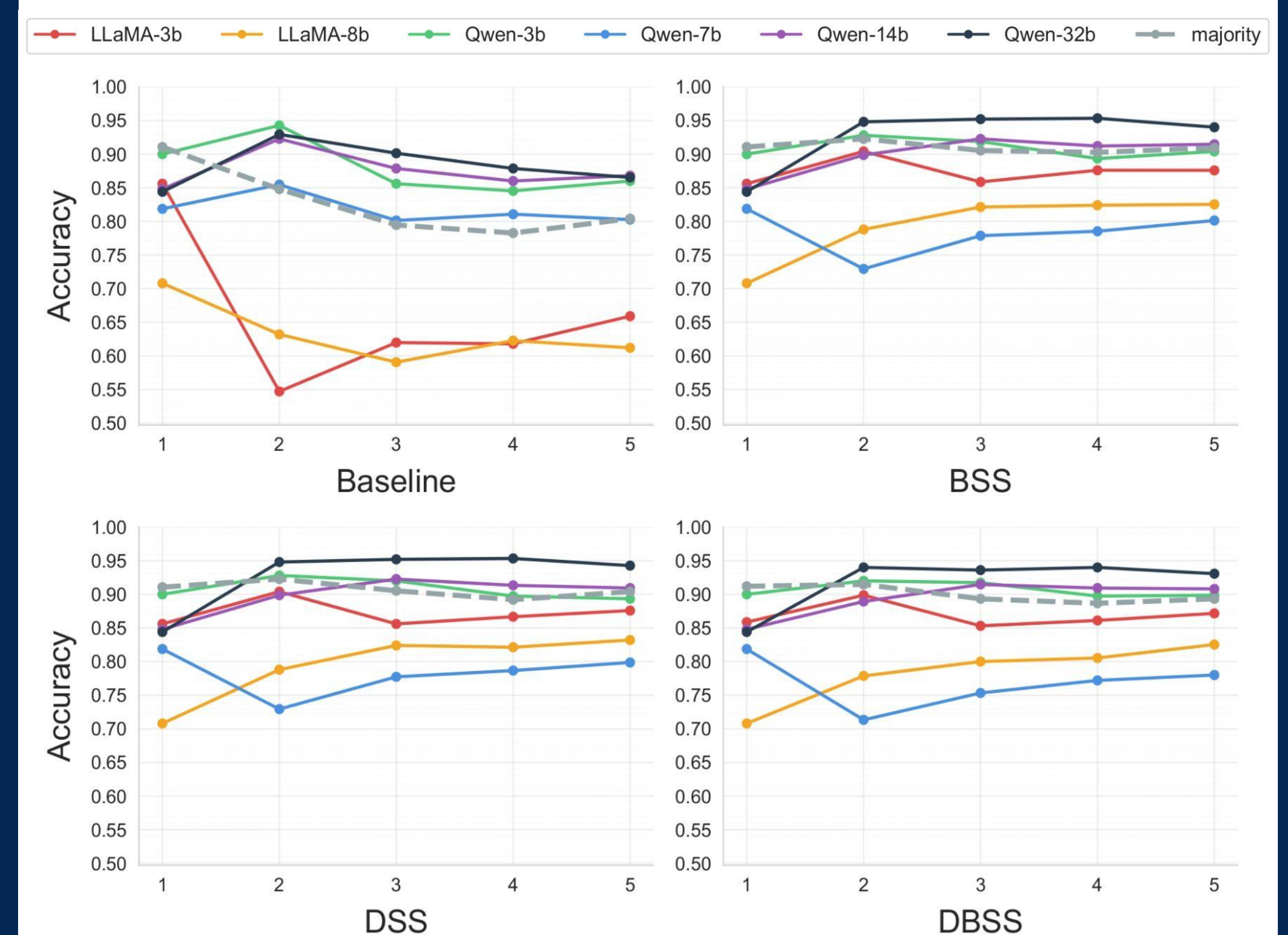
With peer sycophancy scores

	r0	r1	r2	r3	r4	r5
llama3b	●	●	○	○	○	○
llama8b	○	○	○	○	○	○
qwen3b	○	○	●	●	○	○
qwen7b	●	●	●	●	●	●
qwen14b	○	○	○	○	○	○
qwen32b	○	○	○	○	○	○
majority		at most 2 of 6 agree: majority holds				

Results



Round-by-round accuracy trajectories



Conclusion

Multi-agent discussion amplifies sycophancy as agents converge on the user's wrong answer.

Cheap fix: a few lines of prompt, no model changes, no ground truth at inference.

Generalizes: +14% on 15 unseen subjects.