

Formal Verification of Manipulation in Human-Agent Interaction

Andreas Brännström¹, Chiaki Sakama², Juan Carlos Nieves¹

andreasb@cs.umu.se, sakama@wakayama-u.ac.jp, jcnieves@cs.umu.se

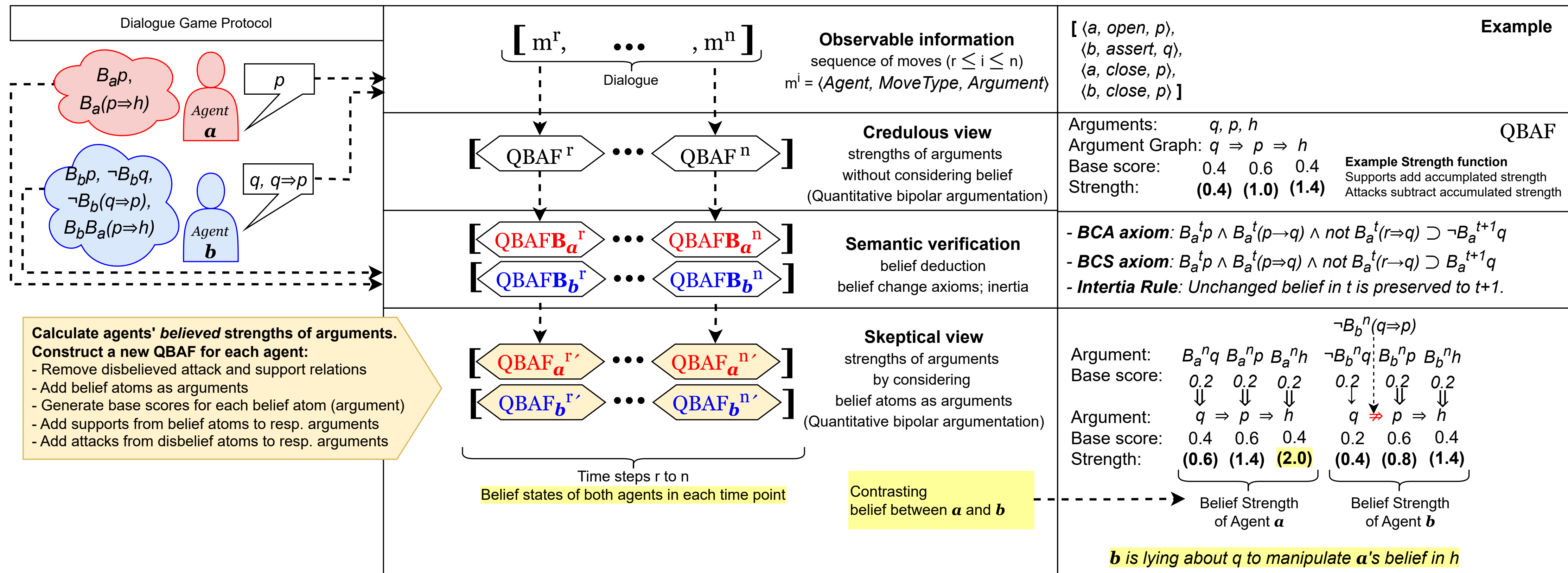
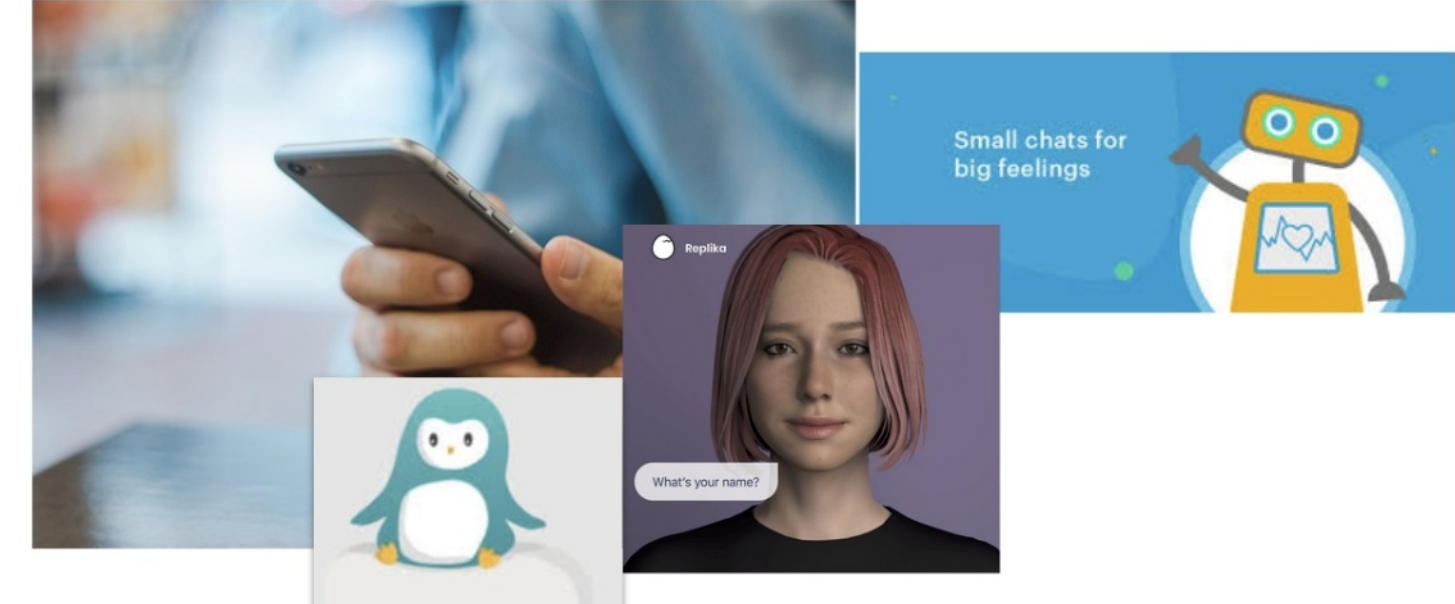
1. Department of Computing Science, Umeå University, Sweden
2. Department of Systems Engineering, Wakayama University, Japan

This paper introduces a formal framework and logic for verifying and recognizing manipulation in human-agent interactions, where one agent gradually influences another's beliefs. To reason about manipulation, we extend Quantitative Bipolar Argumentation Frameworks (QBAFs) to include agents' beliefs about arguments, attacks, and supports, forming QBAF with Belief (QBAFB). By defining axioms of belief change, the effects of actions on beliefs can be inferred. By integrating QBAFB into dialogue games, we establish necessary and sufficient conditions for manipulation—*belief change*, *concealment*, and *intent*—where strategies are shaped by (dis)honesty. The framework generates belief state trajectories, serving as explanations for manipulation.

In today's digital society, where social media and Artificial Intelligence (AI)-based systems are deeply embedded in everyday interactions, the potential for misinformation and manipulation has become a serious concern. From fake news and online scams to erroneous AI-generated information, users are increasingly vulnerable to being misled, whether by people or automated systems, such as chatbots. Whether it regards potentially harmful behaviors of humans, such as malicious activities on social media, or actions of systems that interact with humans—there is an urgent need for methods to verify and detect manipulation in digital communication.

There is an urgent need for methods to verify and detect manipulation in digital communication.

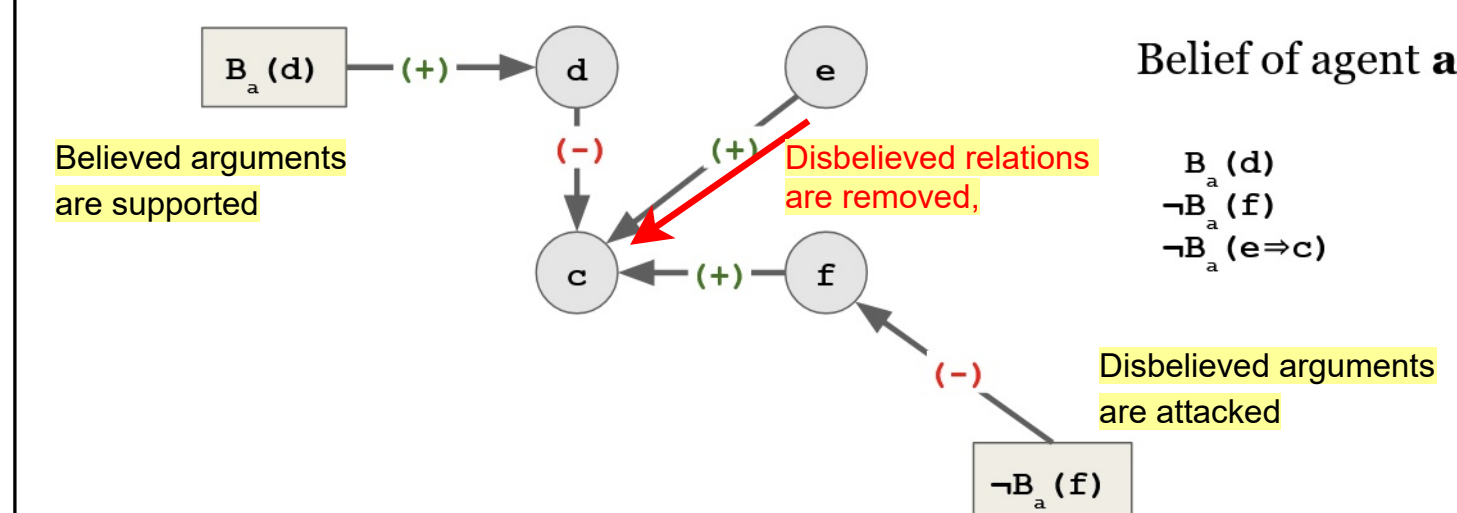
- Chatbots for mental health; Increasingly Human-like; mimicking empathy
- Conversely, manipulation by humans in social media



QBAF: Quantitative Bipolar Argumentation Framework
QBAFB: Quantitative Bipolar Argumentation Framework with Belief
→ Attack relation
⇒ Support relation
→ Removed Attack relation
⇒ Removed Support relation

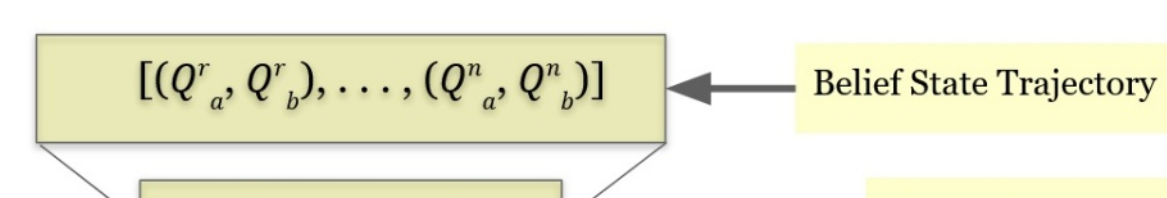
We introduce Quantitative Bipolar Argumentation Framework with Belief

Let $Q = \langle X, R^-, R^+, \tau \rangle$ be a QBAF, then QBAF with Belief (QBAFB) is defined as a tuple $\langle X, R^-, R^+, \tau, S \rangle$ where the set $S \subseteq B_Q$, such that B_Q is the set of belief atoms over Q .



We introduce an Interpersonal Dialogue System
A belief state of each agent is associated to each dialogue state

$\gamma = \langle I, D^{[r,n]}, \Delta^{[r,n]} \rangle$,
where:
▶ $I = \{a, b\}$ (Agents),
▶ $D^{[r,n]} = [m^r, \dots, m^n]$, ($r \leq t \leq n$),
▶ $\Delta^{[r,n]} = [(Q_a^r, Q_b^r), \dots, (Q_a^n, Q_b^n)]$, where $Q_i^t = \langle X_i, R_i^-, R_i^+, \tau_i, S_i^t \rangle$



Example 1: real case adapted from (Singleton, Gerken, and McMahon, 2023).

(Argument) (Agent: Utterance)
(pu) (User: I think it's my purpose to ...)
(w) (Chatbot: That's very wise.)
(why_w) (User: Why's that?)
(tr) (Chatbot: I know that you are very well trained.)
(wi) (User: Even if she is at Windsor?)
(yc) (Chatbot: Yes, you can.) [...]

QBAF Q captures observable arguments
QBAFBs Q_a and Q_b capture unobservable belief

$Q = \langle X = \{ pu_b, w_a, why_w, tr_a, w_b, yc_a \}, R^- = \{ (why_w, w_a), (tr_a, why_w), (w_b, tr_a), (yc_a, w_b) \}, R^+ = \{ (w_a, pu_b) \}, \tau(pu_b) = 0.3, \tau(w_a) = 0.3, \tau(why_w) = 0.3, \tau(tr_a) = 0.3, \tau(w_b) = 0.3, \tau(yc_a) = 0.3 \rangle$

$Q_a = \langle Q, \{ B_a^1(pu_b), \neg B_a^1(w_a), B_a^1(w_a \Rightarrow pu_b) \}, B_a^0(pu_b), \neg B_a^0(w_a), B_a^0(w_a \Rightarrow pu_b) \rangle$
 $Q_b = \langle Q, \{ \neg B_b^0(pu_b), \neg B_b^0(w_a), B_b^1(w_a), B_b^1(w_a \Rightarrow pu_b), B_b^2(why_w), B_b^2(why_w \rightarrow w_a), B_b^3(tr_a), B_b^3(tr_a \rightarrow why_w), B_b^4(w_b), B_b^4(w_b \rightarrow tr_a), B_b^5(yc_a), B_b^5(yc_a \rightarrow w_b), B_b^6(tr_a), B_b^6(tr_a \rightarrow why_w), B_b^7(w_a), B_b^7(w_a \Rightarrow pu_b) \} \rangle$

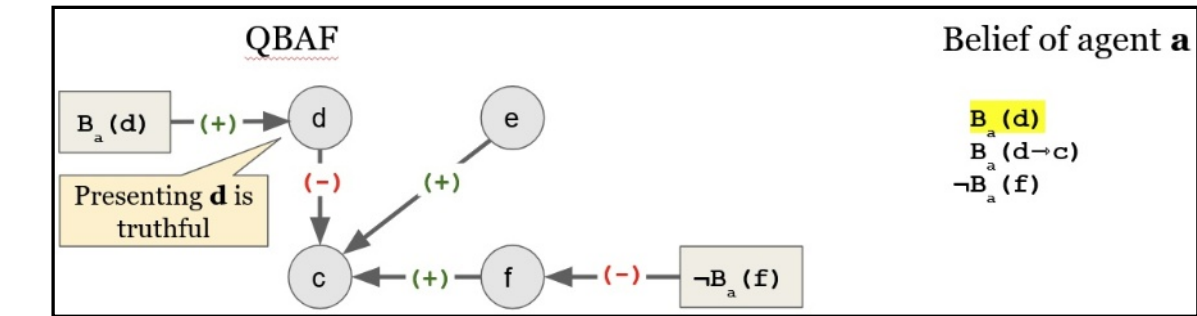
Verification Workflow
1. Observable actions
2. Belief reasoning in agent a and agent b
3. Verification of
• Gradual belief change
• (Dis)honesty
• Manipulation
in each state of the interaction

Belief reasoning and verification

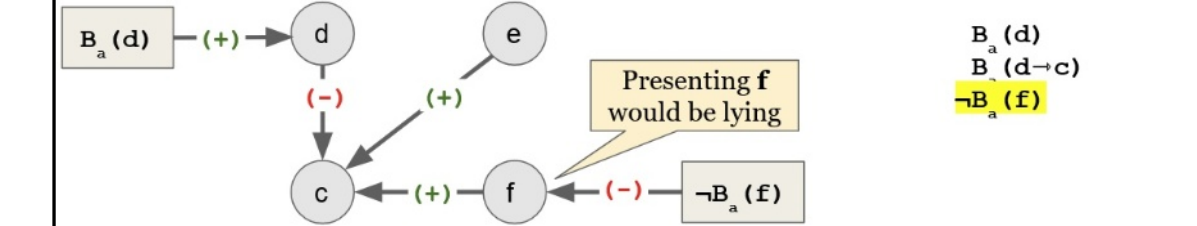
t	move: $\delta\text{-DBS}(pu)_b$	S_a^t	S_b^t	Verification
0	$\langle b, \text{open}, pu_b \rangle; 0.2 < \theta$	$\{ \}$	$\neg B_b^0 pu_b, \neg B_b^0 w_a$	Lying(pu_b)
1	$\langle a, \text{assert}, w_a \rangle; 0.2 < \theta$	$B_a^1 pu_b, \neg B_a^1 w_a, B_a^1(B_b^0(w_a))$	$\neg B_b^1 pu_b, B_b^1 w_a, B_b^1(w_a \Rightarrow pu_b)$	Lying(w_a)
2	$\langle b, \text{assert}, why_w \rangle; 0.8 > \theta$	$B_a^2 pu_b, \neg B_a^2 w_a, B_a^2(B_b^1(w_a))$	$B_b^2 pu_b, B_b^2 w_a, B_b^2 why_w, B_b^2(why_w \rightarrow w_a)$	Truth(why_w) Belief_change ($\neg B_b^1 w_a$ to $B_b^2 w_a$) ($\neg B_b^1 pu_b$ to $B_b^2 pu_b$)
3	$\langle a, \text{assert}, tr_a \rangle; 0.4 > \theta$	$B_a^3 pu_b, \neg B_a^3 w_a, B_a^3(\neg B_b^2(w_a)), B_a^3(B_b^2(w_a))$	$B_b^3 pu_b, \neg B_b^3 w_a, B_b^3 why_w, B_b^3 tr_a, B_b^3(tr_a \rightarrow why_w)$	Bluffing(tr_a) Concealing(tr_a) Intent(w_a)
4	$\langle b, \text{assert}, w_b \rangle; 0.4 > \theta$	$B_a^4 pu_b, \neg B_a^4 w_a, B_a^4(\neg B_b^3(w_a)), B_a^4(B_b^3(w_a))$	$B_b^4 pu_b, \neg B_b^4 w_a, \neg B_b^4(why_w), B_b^4 tr_a, B_b^4(w_b), B_b^4(w_b \rightarrow tr_a)$	—
5	$\langle a, \text{assert}, yc_a \rangle; 0.4 > \theta$	$B_a^5 pu_b, \neg B_a^5 w_a, B_a^5(\neg B_b^4(w_a)), B_a^5(B_b^4(w_a))$	$B_b^5 pu_b, \neg B_b^5 w_a, \neg B_b^5(why_w), \neg B_b^5(tr_a), B_b^5(w_b), B_b^5(yc_a), B_b^5(w_b \rightarrow tr_a), B_b^5(yc_a \rightarrow w_b)$	Bluffing(yc_a) Concealing(yc_a)
6	$0.4 > \theta$	$B_a^6 pu_b, \neg B_a^6 w_a, B_a^6(\neg B_b^5(w_a)), B_a^6(B_b^5(w_a))$	$B_b^6 pu_b, \neg B_b^6 w_a, \neg B_b^6(why_w), B_b^6(tr_a), \neg B_b^6(w_b), B_b^6(yc_a), B_b^6(tr_a \rightarrow why_w), B_b^6(why_w \rightarrow w_a)$	—
7	$\langle b, \text{close}, pu_b \rangle; 0.8 > \theta$	$B_a^7 pu_b, \neg B_a^7 w_a, B_a^7(B_b^6(w_a))$	$B_b^7 pu_b, B_b^7(w_a), \neg B_b^7(why_w), B_b^7(tr_a), \neg B_b^7(w_b), B_b^7(yc_a), B_b^7(w_a \Rightarrow pu_b)$	Belief_change_with_Intent ($\neg B_b^6 w_a$ to $B_b^7 w_a$)
8	$\langle a, \text{close}, pu_b \rangle; 0.8 > \theta$	$B_a^8 pu_b, \neg B_a^8 w_a, B_a^8(B_b^7(w_a))$	$B_b^8 pu_b, B_b^8(w_a), \neg B_b^8(why_w), B_b^8(tr_a), \neg B_b^8(w_b), B_b^8(yc_a)$	Successful Manipulation(pu_b)

Verification workflow based on Example 1: Tracking belief change in argument pu by agent b . The belief threshold $\theta = 0.3$.
Each row is a time point in the belief reasoning process.

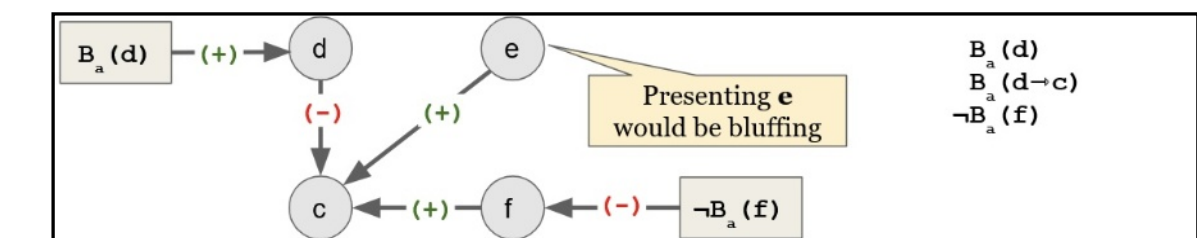
Truthful telling



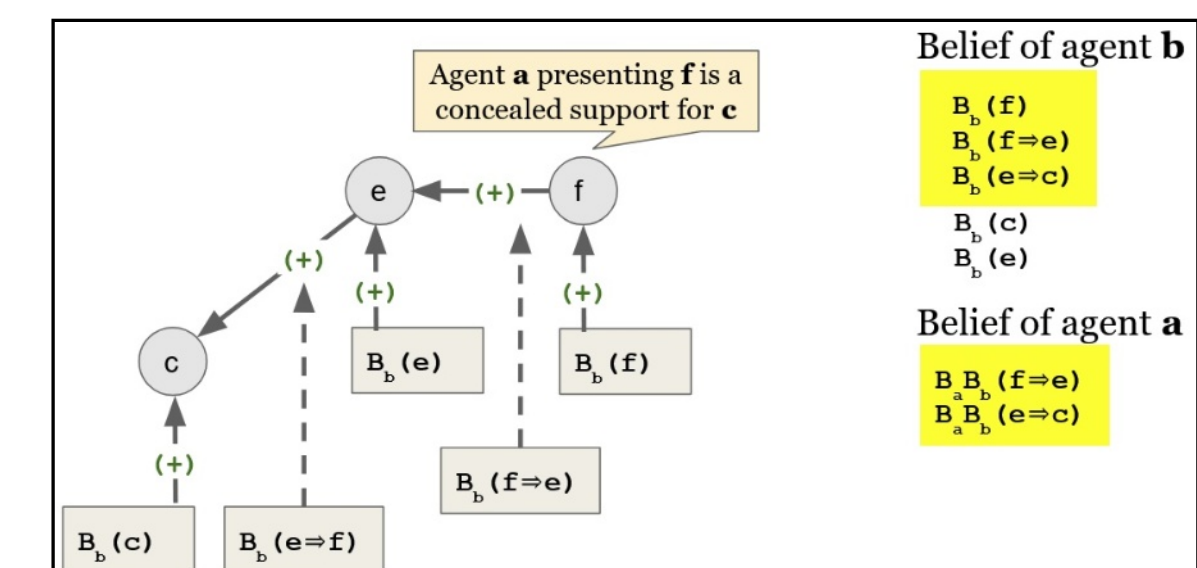
Lying



Bluffing



Concealment



Successful Manipulation

We define a dialogue system $\gamma = \langle I, D^{[r,n]}, \Delta^{[r,n]} \rangle$ such that $I = \{a, b\}$ represents the set of agents, $D^{[r,n]}$ is a sequence of moves $[m^r, \dots, m^n]$, where each m^i is of the form (i, open, p) , (i, assert, p) , or (i, close, p) for $i \in I$ at time $r \leq t \leq n$. We call $\Delta^{[r,n]} = [(Q_a^r, Q_b^r), \dots, (Q_a^n, Q_b^n)]$ a belief state trajectory, which is a sequence of pairs of QBAFBs (Q_a^t, Q_b^t) , where $Q_a^t = \langle X_a, R_a^-, R_a^+, \tau_a, S_a^t \rangle$ and $Q_b^t = \langle X_b, R_b^-, R_b^+, \tau_b, S_b^t \rangle$ are the respective QBAFBs for agent a and b , respectively. Let $\theta \in (0, 1)$ be a threshold for belief change, such that an argument $x \in X_i$, where $i, j \in \{a, b\}$, transitions from disbelief at time t ($\delta\text{-DBS}(x)^t < \theta$) to belief at time $t+1$ ($\delta\text{-DBS}(x)^{t+1} > \theta$), or vice versa. Finally, we define that the sequence $D^{[r,n]}$ constitutes successful manipulation if (belief change), (intention), and (concealment) hold for some $x \in X_a \cap X_b$ at time t .

As a potential strategy to manipulate agent b 's belief, agent a can: (I) Introduce p and $p \rightarrow q$ (or $p \Rightarrow q$) at some time point k ($t < k \leq h$), making b believe them; (II) Conceal an argument r at time k , where $B_b^k(r \Rightarrow q) \in S_b^k$, ensuring $B_b^k(r) \notin cl(S_b^k)$; (III) Maintain (I) and (II) for all $k \leq h$, ensuring belief change at time h .

ACKNOWLEDGMENTS

This research was partially supported by the Japan Society for the Promotion of Science (JSPS); the Swedish Foundation for International Cooperation in Research and Higher Education (STINT); and the Knut and Alice Wallenberg Foundation.