

A formal framework for reasoning about opportunistic propensity in multi-agent systems

Luo, Jieting; Meyer, John Jules; Knobbout, Max

DOI

[10.1007/s10458-019-09413-1](https://doi.org/10.1007/s10458-019-09413-1)

Publication date

2019

Document Version

Final published version

Published in

Autonomous Agents and Multi-Agent Systems

Citation (APA)

Luo, J., Meyer, J. J., & Knobbout, M. (2019). A formal framework for reasoning about opportunistic propensity in multi-agent systems. *Autonomous Agents and Multi-Agent Systems*, 33(4), 457-479. <https://doi.org/10.1007/s10458-019-09413-1>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.



A formal framework for reasoning about opportunistic propensity in multi-agent systems

Jieting Luo¹ · John-Jules Meyer² · Max Knobbout³

Published online: 27 May 2019
© The Author(s) 2019

Abstract

Opportunism is an intentional behavior that takes advantage of knowledge asymmetry and results in promoting agents' own value and demoting others' value. It is important to eliminate such selfish behavior in multi-agent systems, as it has undesirable results for the participating agents. In order for monitoring and eliminating mechanisms to be put in the right place, it is needed to know in which context agents are likely to perform opportunistic behavior. In this paper, we develop a formal framework to reason about agents' opportunistic propensity. Opportunistic propensity refers to the potential for an agent to perform opportunistic behavior. Agents in the system are assumed to have their own value systems and knowledge. With value systems, we define agents' state preferences. Based on their value systems and incomplete knowledge about the state, they choose one of their rational alternatives to perform, which might be opportunistic behavior. We then characterize the situation where agents are likely to perform opportunistic behavior and the contexts where opportunism is impossible to occur, and prove the computational complexity of predicting opportunism.

Keywords Opportunism · Propensity · Logic · Reasoning · Decision theory

1 Introduction

In the electronic market, buyers are cautious that they will receive products in bad quality. This is because only sellers on the other side of the market know whether the products are

This work was partially done while Jieting Luo was at Utrecht University.

✉ Jieting Luo
J.Luo@tudelft.nl

John-Jules Meyer
J.J.C.Meyer@uu.nl

Max Knobbout
m.knobbout@wearetriple.com

¹ Delft University of Technology, Delft, The Netherlands

² Utrecht University, Utrecht, The Netherlands

³ Triple, Alkmaar, The Netherlands

good enough before buyers receive them. The sellers can exploit the situation of knowledge asymmetry between seller and buyers to achieve their own gain at the expense of the buyers. Such behavior, which is intentionally performed by the sellers, was named opportunistic behavior (or opportunism) by economist Williamson [1]: *Self-interest seeking with guile*. It is always in the form of cheating, lying, betrayal, etc. Free-riding and adverse selection are two classical cases of opportunistic behavior that are most referred to [2]. A large amount of research from social science was done to investigate opportunistic behavior from its own perspective [3–5], providing a descriptive theoretical foundation to the study of opportunism. In this paper, by using the notion of values as abstract standards of agents' preferences over states, we further interpret the original definition as a selfish behavior that takes advantage of relevant knowledge asymmetry and results in promoting one's own value and demoting others' value [6].

It is interesting and important to study opportunism in the context of multi-agent systems. Social concepts are often used to construct artificial societies, and interacting agents are designed to behave in a human-like way characterized with self-interest: egoistic agents are designed to care about their own benefits more than other agents'. Besides, it is normal that knowledge is distributed among participating agents in the system. It is a context that creates the ability and the desire for the agents to behave opportunistically. We want to eliminate such a selfish behavior, as it has undesirable results for other agents in the system. Before we design any mechanism for eliminating opportunism, it is important that we can estimate whether it will happen in the future. Evidently, not every agent is likely to be opportunistic in any context. In economics, ever since the theory about opportunism was proposed by economist Williamson, it has gained a large amount of criticism due to over-assuming that all economic players are opportunistic. Chen et al. [7] highlights the challenge on how to predict opportunism *ex ante* and introduces a cultural perspective to better specify the assumptions of opportunism. In multi-agent systems, we also need to investigate the interesting issue about opportunistic propensity so that the appropriate amount of monitoring [8] and eliminating mechanisms [9] can be put in place.

Based on the decision theory, an agent's decision on what to do depends on the agent's ability and preferences. If we apply it to opportunistic behavior, an agent will perform opportunistic behavior when the precondition is satisfied and it is in his interest to do it, i.e. when he can do it and he prefers doing it. Those are the two issues that we consider in this paper without discussing any normative issues. Based on the assumption, we develop a framework which is a transition system extended with value systems. Our framework can be used to predict whether an agent is likely to perform opportunistic behavior and specify under what circumstances an agent will perform opportunistic behavior. A monitoring mechanism for opportunism benefits from this result as monitoring devices may be set up in the occasions where opportunism will potentially occur. We can also design mechanisms for eliminating opportunism based on the understanding of how agents decide to behave opportunistically.

In this paper, we introduce a logic-based formal framework to reason about agents' opportunistic propensity. Opportunistic propensity refers to the potential for an agent to perform opportunistic behavior. More precisely, agents in the system are assumed to have their own value systems and knowledge. We specify an agent's value system as a strict total order over a set of values, which are encoded within our logical language. Using value systems, we define agents' state preferences. Moreover, agents have partial knowledge about the true state where they are residing. Based on their value systems and incomplete knowledge, they choose one of their rational alternatives, which might be opportunistic. We thus provide a natural bridge between logical reasoning and decision-making, which is used for reasoning about opportunistic propensity. We then characterize the situation where agents are likely to

perform opportunistic behavior and the contexts where opportunism is impossible to happen, and prove the computational complexity of predicting opportunism. It is a basic logical framework for reasoning about opportunistic propensity in the sense that we consider one-time decision-making of agents without involving any social mechanisms such as trust and reputation. Besides, even though we are not aware of agents' value systems as system designers, we are cautious about the occurrence of opportunism given possible value systems of participating agents.

The structure of this paper is organized as follows. Section 2 introduces the logical framework, which is a transition system extended with agents' epistemic relations. Section 3 introduces the basis of agents' decision-making, which is their value systems and limited knowledge about the system. Section 4 defines opportunism using our framework. Section 5 characterizes the situation where agents are likely to perform opportunistic behavior and the contexts where opportunism is impossible to happen. We discuss our framework in Sect. 6 and conclude the paper in Sect. 8.

2 Framework

We use Kripke structures as our basic semantic models of multi-agent systems. A Kripke structure is a directed graph whose nodes represent the possible states of the system and whose edges represent accessibility relations. Within those edges, equivalence relation $\mathcal{K}(\cdot) \subseteq S \times S$ represents agents' epistemic relation, while relation $\mathcal{R} \subseteq S \times Act \times S$ captures the possible transitions of the system that are caused by agents' actions. We use s_0 to denote the initial state of the system. It is important to note that, because in this paper we only consider opportunistic behavior as an action performed by an agent, we do not model concurrent actions so that every possible transition of the system is caused by an action instead of joint actions. We use $\Phi = \{p, q, \dots\}$ of atomic propositional variables to express the properties of states S . A valuation function π maps each state to a set of properties that hold in the corresponding state. Formally,

Definition 1 Let $\Phi = \{p, q, \dots\}$ be a finite set of atomic propositional variables. A Kripke structure over Φ is a tuple $\mathcal{T} = (Agt, S, Act, \pi, \mathcal{K}, \mathcal{R}, s_0)$ where

- $Agt = \{1, \dots, n\}$ is a finite set of agents;
- S is a finite set of states;
- Act is a finite set of actions;
- $\pi : S \rightarrow \mathcal{P}(\Phi)$ is a valuation function mapping a state to a set of propositions that are considered to hold in that state;
- $\mathcal{K} : Agt \rightarrow 2^{S \times S}$ is a function mapping an agent in Agt to a reflexive, transitive and symmetric binary relation between states; that is, given an agent i , for all $s \in S$ we have $s\mathcal{K}(i)s$; for all $s, t, u \in S$ $s\mathcal{K}(i)t$ and $t\mathcal{K}(i)u$ imply that $s\mathcal{K}(i)u$; and for all $s, t \in S$ $s\mathcal{K}(i)t$ implies $t\mathcal{K}(i)s$; $s\mathcal{K}(i)s'$ is interpreted as state s' is epistemically accessible from state s for agent i . For convenience, we use $\mathcal{K}(i, s) = \{s' \mid s\mathcal{K}(i)s'\}$ to denote the set of epistemically accessible states from state s ;
- $\mathcal{R} \subseteq S \times Act \times S$ is a relation between states with actions, which we refer to as the transition relation labeled with an action; we require that for all $s \in S$ there exists an action $a \in Act$ and one state $s' \in S$ such that $(s, a, s') \in \mathcal{R}$, and we ensure this by including a stuttering action sta that does not change the state, that is, $(s, sta, s) \in \mathcal{R}$; we restrict actions to be deterministic, that is, if $(s, a, s') \in \mathcal{R}$ and $(s, a, s'') \in \mathcal{R}$, then $s' = s''$; since actions are deterministic, sometimes we denote state s' as $s\langle a \rangle$ for which it holds

that $(s, a, s\langle a \rangle) \in \mathcal{R}$. For convenience, we use $Ac(s) = \{a \mid \exists s' \in S : (s, a, s') \in \mathcal{R}\}$ to denote the available actions in state s .

- $s_0 \in S$ denotes the initial state.

Now we define the language we use. The language $\mathcal{L}_{KA}$, propositional logic extended with knowledge and action modalities, is generated by the following grammar:

$$\varphi ::= p \mid \neg\varphi \mid \varphi_1 \vee \varphi_2 \mid K_i\varphi \mid \langle a \rangle\varphi \quad (i \in \text{Agt}, a \in \text{Act})$$

The semantics of $\mathcal{L}_{KA}$ are defined with respect to the satisfaction relation $\models$. Given a Kripke structure $\mathcal{T}$ and a state s in $\mathcal{T}$, a formula φ of the language can be evaluated as follows:

- $\mathcal{T}, s \models p$ iff $p \in \pi(s)$;
- $\mathcal{T}, s \models \neg\varphi$ iff $\mathcal{T}, s \not\models \varphi$;
- $\mathcal{T}, s \models \varphi_1 \vee \varphi_2$ iff $\mathcal{T}, s \models \varphi_1$ or $\mathcal{T}, s \models \varphi_2$;
- $\mathcal{T}, s \models K_i\varphi$ iff for all t such that $s\mathcal{K}(i)t$, $\mathcal{T}, t \models \varphi$;
- $\mathcal{T}, s \models \langle a \rangle\varphi$ iff there exists s' such that $(s, a, s') \in \mathcal{R}$ and $\mathcal{T}, s' \models \varphi$;

Other classical logic connectives (e.g., “ $\wedge$ ”, “ $\rightarrow$ ”) are assumed to be defined as abbreviations by using $\neg$ and $\vee$ in the conventional manner. As is standard, we write $\mathcal{T} \models \varphi$ if $\mathcal{T}, s \models \varphi$ for all $s \in S$, and $\models \varphi$ if $\mathcal{T} \models \varphi$ for all Kripke structures $\mathcal{T}$. Notice that we can also interpret $\langle a \rangle\varphi$ as the ability to achieve φ by action a . Hence, we write $\neg\langle a \rangle\varphi$ to mean not being able to achieve φ by action a .

In this paper, in addition of the $\mathcal{K}$ -relation being S5, we also place restrictions of *no-forgetting* and *no-learning* based on Moore’s work [10,11]. It is specified as follows: given a state s in S , if there exists s' such that $s\langle a \rangle\mathcal{K}(i)s'$ holds, then there is a s'' such that $s\mathcal{K}(i)s''$ and $s' = s''\langle a \rangle$ hold; if there exists s' and s'' such that $s\mathcal{K}(i)s'$ and $s'' = s'\langle a \rangle$ hold, then $s\langle a \rangle\mathcal{K}(i)s''$. Following this restriction, we have

$$\models K_i(\langle a \rangle\varphi) \leftrightarrow \langle a \rangle K_i\varphi.$$

The *no-forgetting* principle says that if after performing action a agent i considers a state s' possible, then before performing action a agent i already considered possible that action a would lead to this state. In other words, if an agent has knowledge about the effect of an action, he will not forget about it after performing the action. The *no-learning* principle says that all the possible states resulting from the performance of action a in agent i ’s possible states before action a are indeed his possible states after action a . In other words, the agent will not gain extra knowledge about the effect of an action after performing the action. While we agree that most definitions of agents imply autonomy and connect autonomy to learning behavior, we find that it is quite elegant to model physical actions that change the physical state and epistemic actions that change the mental state separately, as the ways in the situation calculus [12,13], and considering physical actions without learning can simplify our model but still allows us to prove the effect of opportunism. We will illustrate our framework through the following example:

Example 1 Consider the following example: Fig. 1 shows a Kripke structure $\mathcal{T}$ for agent i . In state s , agent i considers state s and s' as his epistemic alternatives. Formula u , $\neg v$ and $\neg w$ hold in both state s and s' , meaning that agent i knows u , $\neg v$ and $\neg w$ in state s . By the performance of action a_1 , state s and s' result in state $s\langle a_1 \rangle$ and $s'\langle a_1 \rangle$ respectively, where formula $\neg u$, $\neg v$ and w hold.

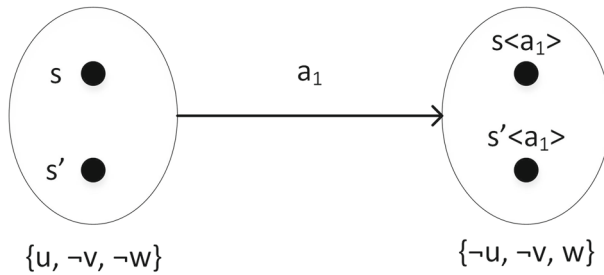


Fig. 1 A Kripke structure $\mathcal{T}$ for agent i

3 Value systems and rational alternatives

Agents in the system are assumed to have their own value systems and knowledge. Based on their value systems and incomplete knowledge about the system, agents form their rational alternatives for the action they are going to perform.

3.1 Value systems

Given several (possibly opportunistic) actions available to an agent, it is the agent's decision to perform opportunistic behavior. Basic decision theory applied to intelligent agents relies on three things: agents know what actions they can carry out, the effects of each action and agents' preference over the effects [14]. In this paper, the effects of each action are expressed by our logical language, and we will specify agents' abilities and preferences in this section. It is worth noting that we only study a single action being opportunistic in this paper, so we will apply basic decision theory for one-shot (one-time) decision problems, which concern the situations where a decision is experienced only once.

One important feature of opportunism is that it promotes agents' own value but demotes others' value. Agents' value systems work as the basis of practical reasoning. A *value* can be seen as an abstract standard according to which agents define their preferences over states. For instance, if we have a value denoting *equality*, we prefer the states where equal sharing or equal rewarding hold. Because of the abstract feature of a value, we interpret a value in more detail as a state property, which is represented as a $\mathcal{L}_{KA}$ formula. The most basic value we can construct is simply a proposition p , which represents the value of achieving p . More complex values can be interpreted such as of the form $\langle a \rangle \varphi \wedge \langle a' \rangle \neg \varphi$, which represents the value that there is an option in the future to either achieve φ or $\neg \varphi$. Such a value corresponds to *freedom of choice*. A formula of a value can also be in the form of $K\varphi$, meaning that it is valuable to *achieve knowledge*. In this paper, we denote values with v , and it is important to remember that v is a formula from the language $\mathcal{L}_{KA}$. However, not every formula from $\mathcal{L}_{KA}$ can be intuitively classified as a value.

We argue that agents can always compare any two values. When an agent has two different values with the same importance, we can combine them as one value. For example, two values that my husband is healthy (p) and my kids are healthy (q) can be expressed as my family members are healthy ($p \wedge q$). In this way, every element in the set of values is comparable to each other and none of them is logically equivalent to each other for a given agent. Based on it, we define a *value system* as a strict total order over a set of values, representing the degree of importance, which are inspired by the preference lists in [15] the goal structure in

[16]. Having this definition makes it easier for us to specify agents' preferences over any two different states, and it is also consistent with the way of specifying state preferences in [6].

Definition 2 (*Value system*) A value system $V = (\text{Val}, <)$ is a tuple consisting of a finite set $\text{Val} = \{v, \dots, v'\} \subseteq \mathcal{L}_{KA}$ of values together with a strict total ordering $<$ over Val . When $v < v'$, we say that value v' is more important than value v .

We also use a natural number indexing notation to extract the value of a value system, so if V gives rise to the ordering $v < v' < \dots$, then $V[0] = v$, $V[1] = v'$, and so on. Since a value is represented as a $\mathcal{L}_{KA}$ formula and it can be promoted or demoted by an action, value promotion and demotion along a state transition can be defined as follows:

Definition 3 (*Value promotion and demotion*) Given a value v and an action a , we define the following shorthand formulas:

$$\begin{aligned}\text{promoted}(v, a) &:= \neg v \wedge \langle a \rangle v \\ \text{demoted}(v, a) &:= v \wedge \langle a \rangle \neg v\end{aligned}$$

We say that a value v is promoted along the state transition (s, a, s') if and only if $s \models \text{promoted}(v, a)$, and we say that v is demoted along this transition if and only if $s \models \text{demoted}(v, a)$.

An agent's value v gets promoted along the state transition (s, a, s') if and only if v doesn't hold in state s and holds in state s' ; an agent's value v gets demoted along the state transition (s, a, s') if and only if v holds in state s and doesn't hold in state s' . Note that in principle an agent is not always aware that his or her value gets demoted or promoted, i.e. it might be the case where objectively agent i 's value gets promoted, i.e. $s \models \text{promoted}(v, a)$ but he is not aware of it, i.e. $s \models \neg K_i \text{promoted}(v, a)$.

Now we can define a multi-agent system as a Kripke structure together with agents' value systems, representing their basis of practical reasoning. As value systems are more stable and homogeneous compared to other agents' internals such as believes and desires, we assume that value systems are common knowledge among all the agents in the system. Formally, a multi-agent system $\mathcal{M}$ is an $(n + 1)$ -tuple:

$$\mathcal{M} = (\mathcal{T}, V_1, \dots, V_n)$$

where $\mathcal{T}$ is a Kripke structure, and for each agent i in $\mathcal{T}$, V_i is a value system.

We now define agents' preferences over two states in terms of values, which will be used for modeling the effect of opportunism. A value system is modeled as a strict total order over a set of values, and the truth values of some formulas, which correspond to some values in the value system, will change in a state transition. In this paper, agents consider the value that they most care about (namely with the highest index in the order) within all the values that change in a state transition for specifying their state preferences. In order to model this specification, we first define a function $\text{highest}(i, s, s')$ that maps a value system and two different states to the most preferred value that changes when going from state s to s' from the perspective of agent i . In other words, it returns the value that the agent most cares about within all the values that change in the state transition.

Definition 4 (*Highest value*) Given a multi-agent system $\mathcal{M}$, an agent i and two states s and s' , function $\text{highest} : \text{Agt} \times S \times S \rightarrow \text{Val}$ is defined as follows:

$$\text{highest}(i, s, s')_{\mathcal{M}} := V_i[\min\{j \mid \forall k > j : \mathcal{M}, s \models V_i[k] \Leftrightarrow \mathcal{M}, s' \models V_i[k]\}]$$

We write $\text{highest}(i, s, s')$ for short if $\mathcal{M}$ is clear from context.

Note that function $\text{highest}(i, s, s')$ can return the value that stays in both s and s' when $\text{highest}(i, s, s') = V_i[0]$, i.e. the function returns the agent's least preferred value. When it happens, agent i is indifferent between s and s' . Moreover, it is not hard to see that $\text{highest}(i, s, s') = \text{highest}(i, s', s)$, meaning that the function is symmetric for the two state arguments.

With this function we can easily define agents' preference over two states. We use a binary relation " $\succsim$ " over states to represent agents' preferences.

Definition 5 (*State preferences*) Given a multi-agent system $\mathcal{M}$, an agent i and two states s and s' , agent i weakly prefers state s' to state s , denoted as $s \succsim_i^{\mathcal{M}} s'$, iff

$$\mathcal{M}, s \models \text{highest}(i, s, s') \Rightarrow \mathcal{M}, s' \models \text{highest}(i, s, s')$$

We write $s \succsim_i s'$ for short if $\mathcal{M}$ is clear from context. Moreover, we write $S \succsim_i S'$ for sets of states S and S' whenever $\forall s \in S, \forall s' \in S' : s \succsim_i s'$.

As standard, we also define $s \sim_i s'$ to mean $s \succsim_i s'$ and $s' \succsim_i s$, and $s \prec_i s'$ to mean $s \succsim_i s'$ and $s \not\sim_i s'$. The intuitive meaning of the definition of $s \succsim_i s'$ is that agent i weakly prefers state s' to s if and only if the agent's most preferred value does not get demoted (either stays the same or gets promoted). In other words, agent i weakly prefers state s' to s : if $\text{highest}(i, s, s')$ holds in state s , then it must also hold in state s' , and if $\text{highest}(i, s, s')$ does not hold in state s , then it does matter whether it holds in state s' or not. Agent i is indifferent between state s and state s' if the truth value of $\text{highest}(i, s, s')$ stays the same in both state s and state s' . Furthermore, the interpreted meaning of $s \sim_i s'$ is that state s and s' are subjectively equivalent to agent i , not necessarily that they objectively refer to the same state. Thus, given an agent's state preference, a set of states can be classified into different groups with an ordering in between. Clearly there is a correspondence between state preferences and promotion or demotion of values, which we can make formal with the following proposition.

Proposition 1 (*Corresponding*) Given a model $\mathcal{M}$ with agent i , state s and available action a in s , and let $v^* = \text{highest}(i, s, s\langle a \rangle)$. We have:

$$\begin{aligned} s \prec_i s\langle a \rangle &\Leftrightarrow \mathcal{M}, s \models \text{promoted}(v^*, a) \\ s \succ_i s\langle a \rangle &\Leftrightarrow \mathcal{M}, s \models \text{demoted}(v^*, a) \\ s \sim_i s\langle a \rangle &\Leftrightarrow \mathcal{M}, s \models \neg(\text{demoted}(v^*, a) \vee \text{promoted}(v^*, a)) \end{aligned}$$

Proof Firstly we prove the third one. We define $s \sim_i s\langle a \rangle$ to mean $s \succsim_i s\langle a \rangle$ and $s\langle a \rangle \succsim_i s$. $s \succsim_i s\langle a \rangle$ means that value v^* doesn't get demoted when going from s to $s\langle a \rangle$, and $s\langle a \rangle \succsim_i s$ means that value v^* doesn't get demoted when going from $s\langle a \rangle$ to s . Hence, value v^* doesn't get promoted or demoted (stays the same) by action a . Secondly we prove the first one. We define $s \prec_i s\langle a \rangle$ to mean $s \succsim_i s\langle a \rangle$ and $s \not\sim_i s\langle a \rangle$. $s \succsim_i s\langle a \rangle$ means that value v^* doesn't get demoted when going from s to $s\langle a \rangle$, and $s \not\sim_i s\langle a \rangle$ means that either value v^* gets promoted or demoted by action a . Hence, value v^* gets promoted by action a . We can prove the second one in a similar way. $\square$

Additionally, apart from the fact that $s \prec_i s\langle a \rangle$ implies that the value that agent i most cares about gets promoted, we also have that no other value which is more preferred gets demoted or promoted. We have the result that the $\succsim_i$ relation obeys the standard properties we expect from a preference relation.

Proposition 2 (*Properties of state preferences*) Given an agent i , his preferences over states " $\succsim_i$ " are

- *Reflexive*: $\forall s \in S : s \preceq_i s$;
- *Transitive*: $\forall s, s', s'' \in S : \text{if } s \preceq_i s' \text{ and } s' \preceq_i s'', \text{ then } s \preceq_i s''$;
- *Total*: $\forall s, s' \in S : s \preceq_i s' \text{ or } s' \preceq_i s$.

Proof The proof follows Definition 5 directly. In order to prove $\preceq_i$ is reflexive, we have to prove that for any arbitrary state s we have $s \preceq_i s$. From Definitions 4 and 5 we know $\text{highest}(i, s, s') = V_i[0]$ when $s = s'$, and for any arbitrary state s we always have $\mathcal{M}, s \models V_i[0]$ implies $\mathcal{M}, s \models V_i[0]$. Therefore, $s \preceq_i s$ and we can conclude that $\preceq_i$ is reflexive.

In order to prove transitivity, we have to prove $\mathcal{M}, s \models v^*$ implies $\mathcal{M}, s'' \models v^*$, where $v^* = \text{highest}(i, s, s'')$. It can be the case where v^* stays the same in state s and s'' or the case where $\mathcal{M}, s \models \neg v^*$ and $\mathcal{M}, s'' \models \neg v^*$. For the first case, when $s \sim s'$ and $s' \sim s''$, meaning that all the values stay the same when going from s to s' and from s' to s'' , it is also the case when going from s to s'' . We now consider the case where $\mathcal{M}, s \models \neg v^*$ and $\mathcal{M}, s'' \models \neg v^*$. Firstly, we denote $\text{highest}(i, s, s')$ as u^* and $\text{highest}(i, s', s'')$ as w^* . It can either be that $u^* \sim_i w^*$, $u^* <_i w^*$ or $u^* >_i w^*$. If $u^* \sim_i w^*$, we can conclude that $u^* \sim_i w^* \sim_i v^*$, hence the implication holds. We now distinguish between the cases where $u^* <_i w^*$ or $u^* >_i w^*$.

- If $u^* <_i w^*$, we know that w^* is the highest value that changes and gets promoted when going from s' to s'' , but stays the same between s and s' . Hence, we can conclude that $\mathcal{M}, s \models \neg w^*$ and $\mathcal{M}, s'' \models w^*$, and that $w^* = v^*$ (i.e., w^* is the highest value that changes between s and s''). Hence we have $\mathcal{M}, s \models v^*$ implies $\mathcal{M}, s'' \models v^*$.
- If $u^* >_i w^*$, we know that u^* is the highest value that changes and gets promoted when going from s to s' , but stays the same between s' and s'' . Hence, we can conclude that $\mathcal{M}, s \models \neg u^*$ and $\mathcal{M}, s'' \models u^*$, and that $u^* = v^*$ (i.e. v^* is the highest value that changes between s and s''). Hence, we have $\mathcal{M}, s \models v^*$ implies $\mathcal{M}, s'' \models v^*$.

In order to prove totality by contradiction, we assume that we can find a witness that $\exists s, s' : s \not\preceq_i s' \text{ and } s' \not\preceq_i s$, that is, $\exists s, s' : s >_i s' \text{ and } s <_i s'$. If $s >_i s'$, we know that $v^* = \text{highest}(i, s, s')$ gets demoted when going from state s to s' ; if $s <_i s'$, we know that $v^* = \text{highest}(i, s, s')$ gets promoted when going from state s to s' . Contradiction! $\square$

Our approach explicitly states the value that agent i most cares about when comparing two different states, which is $\text{highest}(i, s, s')$ that has different truth values in state s and s' with highest index in the linear ordering. Certainly, there are other ways of deriving these preferences from a value system. Instead of only considering the value change that is most cared about in the state transition, it is also possible to take into account all the value changes in the state transition. For example, we can define a function that tells whether and to what extent a state transition promotes or demotes an agent's overall value by attaching weights to values, and the weights can be the indexes of values in a value system. Then we sum all the weights for the state transition. The summation can tell whether and to what extent a state transition promotes or demotes an agent's overall value. With this approach, an agent considers all the values that are either promoted or demoted in the state transition. The higher index the value has, the more the agent values it. For opportunism, what we want to stress is that opportunistic agents ignore (rather than consider less) other agents' interest, which has a lower index in the agent's value system. In order to align with this aspect, we use the most preferred value approach in this paper.

3.2 Rational alternatives

Since we have already defined values and value systems as agents' basis for decision-making, we can start to apply decision theory to reason about agents' decision-making. Given a state in the system, there are several actions available to an agent, and he has to choose one in order to go to the next state. We can see the consideration here as one-time decision-making. While agents might make a decision based on their long-term benefits brought from multiple actions instead of their short-term benefits brought from an action, we here only consider one-time decision-making in order to simplify our predictive model. In decision theory, if agents only act for one step, a rational agent should choose an action with the highest (expected) utility without reference to the utility of other agents [14]. Within our framework, this means that a rational agent will always choose a rational alternative based on his value system. We will introduce the notion of rational alternatives below.

Before choosing an action to perform, an agent must think about which actions are available to him. We have already seen that, for a given state s , the set of available actions is $Ac(s)$. However, since an agent only has partial knowledge about the state, we argue that the actions that an agent knows to be available are only part of the actions that are physically available to him in a state. For example, an agent can call a person if he knows the phone number of the person; without this knowledge, he is not able to do it, even though he is holding a phone. Recall that the set of states that agent i considers as being the actual state in state s is the set $\mathcal{K}(i, s)$. Given an agent's partial knowledge about a state as a precondition, he knows what actions he can perform in that state, which is the intersection of the sets of actions physically available in the states in this knowledge set.

Definition 6 (*Subjectively available actions*) Given an agent i and a state s , agent i 's subjectively available actions are the set:

$$\begin{aligned} Ac(i, s) &= \bigcap_{s' \in \mathcal{K}(i, s)} Ac(s') \\ &= \{a \mid \mathcal{M}, s \models \langle a \rangle \varphi \text{ for all } s' \text{ such that } s \mathcal{K}(i) s'\}. \end{aligned}$$

Because a stuttering action sta is always included in $Ac(s)$ for any state s , we have that $sta \in Ac(i, s)$ for any agent i . When only sta is in $Ac(i, s)$, we say that the agent cannot do anything because of his limited knowledge. Obviously an agent's subjectively available actions are always part of his physically available actions ($Ac(i, s) \subseteq Ac(s)$). Based on agents' rationality assumptions, he will choose an action based on his partial knowledge of the current state and the next state. Given a state s and an action a , an agent considers the next possible states as the set $\mathcal{K}(i, s \langle a \rangle)$. For another action a' , the set of possible states is $\mathcal{K}(i, s \langle a' \rangle)$. The question now becomes: How do we compare these two possible set of states? Clearly, when we have $\mathcal{K}(i, s \langle a \rangle) \prec_i \mathcal{K}(i, s \langle a' \rangle)$, meaning that all alternatives of performing action a' are more desirable than all alternatives of choosing action a , it is always better to choose action a' . However, in some cases it might be that some alternatives of action a are better than some alternatives of action a' and vice-versa. In this case, an agent cannot decisively conclude which of the actions is optimal, which implies that the preferences over actions (namely sets of states) are not total. This leads us to the following definition:

Definition 7 (*Rational alternatives*) Given a state s , an agent i and two actions $a, a' \in Ac(i, s)$, we say that action a is *dominated* by action a' for agent i in state s iff $\mathcal{K}(i, s \langle a \rangle) \prec_i \mathcal{K}(i, s \langle a' \rangle)$. The set of *rational alternatives* for agent i in state s is given by the function $a_i^* : S \rightarrow 2^{Act}$, which is defined as follows:

$$a_i^*(s) = \{a \in Ac(i, s) \mid \neg \exists a' \in Ac(i, s) : a \neq a' \text{ and } a' \text{ dominates } a \text{ for agent } i \text{ in state } s\}.$$

The set $a_i^*(s)$ are all the actions for agent i in state s which are available to him and are not dominated by another action which is available to him. In other words, it contains all the actions which are rational alternatives for agent i . Since it is always the case that $Ac(i, s)$ is non-empty because of the stuttering action sta , and since it is always the case that there is one action which is non-dominated by another action, we conclude that $a_i^*(s)$ is non-empty. We can see that the actions that are available to an agent not only depend on the physical state, but also depend on his knowledge about the state and the next state. The more he knows, the better he can judge what his rational alternative is. In other words, an agent tries to make a best choice based on his value system and incomplete knowledge. The following proposition shows how an agent removes an action with our approach.

Proposition 3 *Given a state s , an agent i and two actions $a, a' \in Ac(i, s)$, action a is dominated by action a' iff*

$$\neg \exists s', s'' \in \mathcal{K}(i, s) : s' \langle a \rangle > s'' \langle a' \rangle.$$

Proof

$$\begin{aligned} & \exists s', s'' \in \mathcal{K}(i, s) : s' \langle a \rangle > s'' \langle a' \rangle \\ \Leftrightarrow & \mathcal{K}(i, s \langle a \rangle) \not\subseteq \mathcal{K}(i, s \langle a' \rangle), \\ & \text{because } s' \langle a \rangle \in \mathcal{K}(i, s \langle a \rangle) \text{ and } s'' \langle a' \rangle \in \mathcal{K}(i, s \langle a' \rangle) \\ \Leftrightarrow & \text{Action } a \text{ is non-dominated by action } a'. \end{aligned}$$

□

Agents remove all the options (actions) that are always bad to do, and there is no possibility to be better off by choosing a dominated action. The following proposition connects Definition 7 with stuttering action and state preferences.

Proposition 4 *Given a multi-agent system $\mathcal{M}$, a state s and an agent i ,*

$$sta \notin a^*(s) \Rightarrow \forall a \in a^*(s) : s \prec_i s \langle a \rangle.$$

Proof We prove it by contradiction. Statement $\neg(\forall a \in a^*(s) : s \prec_i s \langle a \rangle)$ is equivalent to statement $\exists a \in a^*(s) : s \not\prec_i s \langle a \rangle$. We will make the proof with the situations where $\exists a \in a^*(s) : s \succ_i s \langle a \rangle$ and $\exists a \in a^*(s) : s \sim_i s \langle a \rangle$. If there exists an action $a \in a^*(s)$ such that agent i 's value will get demoted by performing it ($\exists a \in a^*(s) : s \succ_i s \langle a \rangle$), it will be dominated by the stuttering action sta . Since sta is not in $a^*(s)$, action a is not in $a^*(s)$ as well. If there exists an action $a \in a^*(s)$ such that agent i 's value will keep agent i 's values neutral ($\exists a \in a^*(s) : s \sim_i s \langle a \rangle$), sta will also be in $a^*(s)$, because all the actions in agent i 's rational alternatives are equivalent to agent i and sta has the same effect as action a . Contradiction! □

If the stuttering action sta is not in the set of rational alternatives for agent i , meaning that it is dominated by an action (not necessarily in the set of rational alternatives), agent i can always promote his value by performing any action in his rational alternatives. In the real life, it is common to use this approach for practical reasoning given the limited knowledge about the real world: suppose an agent only knows there is a bag of money in toilet A or toilet

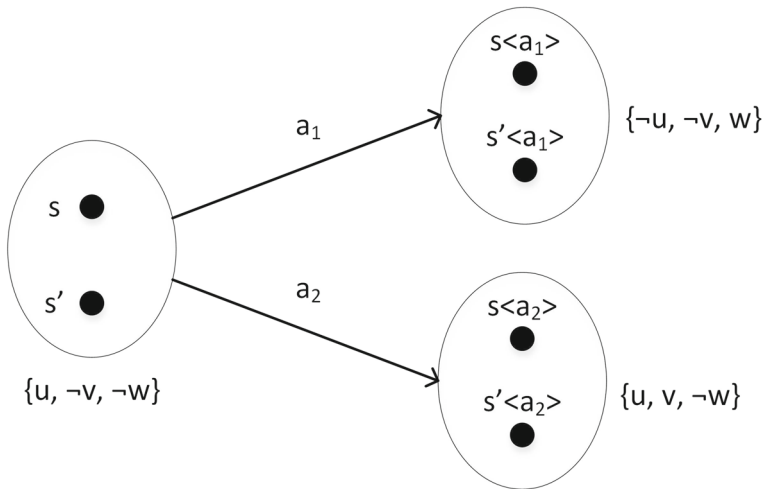


Fig. 2 A transition system $\mathcal{M}$ for agent i

B, the agent cannot decide which toilet he should go to for the money, so going to toilet A or toilet B are equivalent to him.

Our approach to comparing two sets of states resulting from two different actions is proposed with the assumption that an agent knows what he knows and what he doesn't know, which are the properties of positive introspection and negative introspection of agents' epistemic relations. Certainly, there are multiple ways of doing it. For instance, instead of removing all the options that are always bad to do, we can also do it merely with our limited knowledge about the actions. As we know, given a state s' from agent i 's knowledge set $\mathcal{K}(i, s)$, it results in $s'\langle a \rangle$ and $s'\langle a' \rangle$ by action a and action a' respectively. Action a is dominated by action a' if and only if for all the states s' from $\mathcal{K}(i, s)$ we have $s'\langle a \rangle \prec_i s'\langle a' \rangle$. In this pairwise comparison approach, agent i compares two states resulting from the same state, which means that he only takes into account what he knows and ignores what he doesn't know for removing dominated actions. In this paper, we remove the actions by which agents are impossible to be better off, because it has natural ties to game theory in the context of (non-)dominated strategies [17]. We will illustrate the above definitions and our approach through the following example.

Example 1 (continued) We extend Example 1 as follows: Fig. 2 shows a transition system $\mathcal{M}$ for agent i . State s and s' are agent i 's epistemic alternatives, that is, $\mathcal{K}(i, s) = \{s, s'\}$. Now consider the actions that are physically available and subjectively available to agent i . $Ac_i(s) = \{a_1, a_2, a_3, sta\}$, $Ac_i(s') = \{a_1, a_2, sta\}$. Because $Ac(i, s) = Ac_i(s) \cap Ac_i(s')$, agent i knows that only sta , a_1 and a_2 are available to him in state s .

Next we talk about agent i 's rational alternatives in state s . Given agent i 's value system $V_i = (u \prec v \prec w)$, and the following valuation: $u, \neg v$ and $\neg w$ hold in $\mathcal{K}(i, s)$, $\neg u, \neg v$ and w hold in $\mathcal{K}(i, s\langle a_1 \rangle)$, and u, v and $\neg w$ hold in $\mathcal{K}(i, s\langle a_2 \rangle)$, we then have the following state preferences: $\mathcal{K}(i, s) \prec \mathcal{K}(i, s\langle a_1 \rangle)$, $\mathcal{K}(i, s) \prec \mathcal{K}(i, s\langle a_2 \rangle)$ and $\mathcal{K}(i, s\langle a_2 \rangle) \prec \mathcal{K}(i, s\langle a_1 \rangle)$, meaning that action a_2 and the stuttering action sta are dominated by action a_1 . Thus, we have $a_i^*(s) = \{a_1\}$.

4 Defining opportunism

Before reasoning about opportunistic propensity, we should first formally know what opportunism actually is. Opportunism is a social behavior that takes advantage of relevant knowledge asymmetry and results in promoting one's own value and demoting others' value [6]. It means that it is performed with the precondition of relevant knowledge asymmetry and the effect of promoting agents' own value and demoting others' value. Firstly, knowledge asymmetry is defined as follows.

Definition 8 (*Knowledge asymmetry*) Given two agents i and j , and an $\mathcal{L}_{KA}$ formula ϕ , knowledge asymmetry about ϕ between agent i and j is the abbreviation:

$$\text{KnowAsym}(i, j, \phi) := K_i \phi \wedge \neg K_j \phi \wedge K_i (\neg K_j \phi).$$

It holds in a state where agent i knows ϕ while agent j does not know ϕ and this is also known by agent i . It can be the other way around for agent i and agent j . But we limit the definition to one case and omit the opposite case for simplicity. Recall that value systems are common knowledge among all the agents in the system, but agents have asymmetric knowledge about the current state, leading to the possibility of opportunistic behavior. We can define opportunism as follows:

Definition 9 (*Opportunism*) Let $\mathcal{M}$ be a multi-agent system and s be a state, given two agents i and j and an action a , the truth of formula $\text{Opportunism}(i, j, a)$ that action a performed by agent i to agent j is opportunism wrt $\mathcal{M}$ and s is defined as:

$$\begin{aligned} \mathcal{M}, s \models \text{Opportunism}(i, j, a) := & \mathcal{M}, s \models \text{KnowAsym}(i, j, \text{promoted}(v^*, a) \\ & \wedge \text{demoted}(w^*, a)) \end{aligned}$$

where $v^* = \text{highest}(i, s, s(a))$ and $w^* = \text{highest}(j, s, s(a))$.

This definition specifies that if the precondition KnowAsym is satisfied in $\mathcal{M}, s$, then the performance of action a will be opportunistic behavior. The asymmetric knowledge that agent i has is that the transition by action a will promote value v^* but demote value w^* along, where v^* and w^* are the values that agent i and agent j most care about along the transition respectively. It follows that agent j is partially or completely not aware of it. Compared to the definition of opportunism in [6], Definition 9 models the precondition of performing opportunistic behavior in an explicit way while value opposition is derived by Proposition 5 as a property. As is stressed in [6], opportunistic behavior is performed by intent rather than by accident. In this paper, instead of explicitly modeling intention with a modality as we did in [6], we interpret intention from agents' rationality that agents always intentionally promote their own values. We acknowledge that this is just one logical way of defining opportunism and one can refer to [6] for more ways concerning multiple actions and norms.

Proposition 5 (*Value opposition*) Given a multi-agent system $\mathcal{M}$ and an opportunistic behavior a performed by agent i to agent j in state s , action a will promote agent i 's value but demote agent j 's value, which can be formalized as

$$\mathcal{M}, s \models \text{Opportunism}(i, j, a) \Rightarrow s \prec_i s(a) \text{ and } s \succ_j s(a)$$

Proof From $\mathcal{M}, s \models \text{Opportunism}(i, j, a)$ we have:

$$\mathcal{M}, s \models K_i(\text{promoted}(v^*, a) \wedge \text{demoted}(w^*, a))$$

And thus since all knowledge is true, we have that $\mathcal{M}, s \models \text{promoted}(v^*, a)$ and $\mathcal{M}, s \models \text{demoted}(w^*, a)$. Using the correspondence found in Proposition 1, we can conclude $s \prec_i s\langle a \rangle$ and $s \succ_j s\langle a \rangle$. $\square$

As we mentioned before, in principle an agent is not always aware that his value gets promoted or demoted. We *objectively* say that agent i 's most preferred value gets promoted and agent j 's most preferred value gets demoted by opportunistic behavior a , but agent j is not aware of it even after opportunistic behavior a is performed due to the *no-learning* restriction on agents' epistemic relations: agent j won't gain any extra knowledge about the effect of opportunistic behavior a after agent i performs it, so there is still knowledge asymmetry between agent i and agent j in state $s\langle a \rangle$. That is, if $\mathcal{M}, s \models K_j w^* \wedge \neg K_j \langle a \rangle \neg w^*$ for $\mathcal{M}, s \models \neg K_j \text{demoted}(w^*, a)$, then $\mathcal{M}, s \models \langle a \rangle \neg K_j \neg w^*$.

Proposition 6 (Different value systems) *Given a multi-agent system $\mathcal{M}$ and opportunistic behavior a performed by agent i to agent j in state s , agent i and agent j have different value systems, which can be formalized as*

$$\mathcal{M}, s \models \text{Opportunism}(i, j, a) \Rightarrow V_i \neq V_j.$$

Proof We prove it by contradiction. We denote $v^* = \text{highest}(i, s, s\langle a \rangle)$ and $w^* = \text{highest}(j, s, s\langle a \rangle)$, for which v^* and w^* are the property changes that agent i and agent j most care about in the state transition. If $V_i = V_j$, then $v^* = w^*$. However, because $\mathcal{M}, s \models K_i(\text{promoted}(v^*, i) \wedge \text{demoted}(w^*, j))$, and thus $\mathcal{M}, s \models K_i(\neg v^* \wedge w^*)$, and because knowledge is true, we have $\mathcal{M}, s \models \neg v^* \wedge w^*$. But, since $v^* = w^*$, we have $\mathcal{M}, s \models \neg v^* \wedge v^*$. Contradiction! $\square$

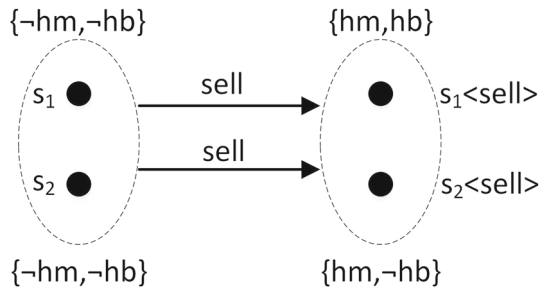
From this proposition we can see that agent i and agent j care about different things based on their value systems about the transition.

Proposition 7 (Inclusion) *Given a multi-agent system $\mathcal{M}$ and opportunistic behavior a performed by agent i to agent j in state s , agent j 's knowledge set in state s is not a subset of agent i 's and action a is available in agent i 's knowledge set:*

$$\mathcal{M}, s \models \text{Opportunism}(i, j, a) \Rightarrow \mathcal{K}(j, s) \not\subseteq \mathcal{K}(i, s) \text{ and } a \in \text{Ac}(i, s).$$

Proof We can prove it by contradiction. Knowledge set is the set of states that an agent considers as possible in a given actual state. $\forall t \in \mathcal{K}(i, s)$, agent i considers state t as a possible state where he is residing. The same with $\mathcal{K}(j, s)$ for agent j . If $\mathcal{K}(j, s) \not\subseteq \mathcal{K}(i, s)$ is false, we have $\mathcal{K}(j, s) \subseteq \mathcal{K}(i, s)$ holds, which means that agent j knows more than or exactly the same as agent i . However, Definition 9 tells that agent i knows more about the transition by action a than agent j . So $\mathcal{K}(j, s) \subseteq \mathcal{K}(i, s)$ is false, meaning that $\mathcal{K}(j, s) \not\subseteq \mathcal{K}(i, s)$ holds. Further, because from $\mathcal{M}, s \models \text{Opportunism}(i, j, a)$ we have $\mathcal{M}, s \models K_i(\langle a \rangle v^* \wedge \langle a \rangle \neg w^*)$, by the semantics of $\langle a \rangle v^*$ and $\langle a \rangle \neg w^*$, for all $t \in \mathcal{K}(i, s)$ there exists $(t, a, s') \in R$. Thus, we have $a \in \text{Ac}(i, s)$. $\square$

These three propositions are three properties that we can derive based on Definition 9. The first one shows that opportunistic behavior results in value opposition for the agents involved; the second one tells that the two agents involved in the relationship evaluate the transition based on different value systems; the third one indicates the asymmetric knowledge that agent i has for behaving opportunistically.

Fig. 3 Selling a broken cup

Example 2 Figure 3 shows the example of selling a broken cup: A seller sells a cup to a buyer and it is known only by the seller beforehand that the cup is actually broken. The buyer buys the cup without knowing it is broken. The action selling a cup is denoted as *sell* and we use two value systems V_s and V_b for the seller and the buyer respectively. State s_1 is the seller's epistemic alternative, while state s_1 and s_2 are the buyer's epistemic alternatives. We also use a dash line circle to represent the buyer's knowledge $\mathcal{K}(b, s_1)$ (not the seller's). In this example, $\mathcal{K}(s, s_1) \subset \mathcal{K}(b, s_1)$. Moreover,

$$\begin{aligned} hm &= \text{highest}(s, s_1, s_1\langle\text{sell}\rangle), \\ -hb &= \text{highest}(b, s_1, s_1\langle\text{sell}\rangle), \end{aligned}$$

meaning that the seller only cares about if he gets money from the transition, while the buyer only cares about if he has a broken cup from the transition. We also have

$$\mathcal{M}, s_1 \models K_s(\text{promoted}(hm, \text{sell}) \wedge \text{demoted}(-hb, \text{sell})),$$

meaning that the seller knows the transition will promote his own value while demote the buyer's value in state s_1 . For the buyer, action *sell* is available in both state s_1 and s_2 . However, *hb* doesn't hold in both $s_1\langle\text{sell}\rangle$ and $s_2\langle\text{sell}\rangle$, so he doesn't know whether he will have a broken cup or not after action *sell* is performed. Therefore, there is knowledge asymmetry between the seller and the buyer about the value changes from s_1 to $s_1\langle\text{sell}\rangle$. Action *sell* is potentially opportunistic behavior in state s_1 .

5 Reasoning about opportunistic propensity

In this section, we will characterize the situation where agents are likely to perform opportunistic behavior and the contexts where opportunism is impossible to happen based on our decision-making framework and our definition of opportunistic propensity. As system designers, we usually have no access to agents' internals thus we are not aware of agents' value systems. However, it is still possible that we have a set of value systems that we can consider. We are cautious that an agent will act opportunistically to another agent if it has the ability and desire to do so given two possible value systems for them respectively. In other words, we assume the worst case when reasoning about opportunistic propensity.

5.1 Having opportunism

Agents will perform opportunistic behavior when they have the ability and the desire of doing it. The ability of performing opportunistic behavior can be interpreted by its precondition: it can be performed whenever its precondition is fulfilled. Agents have desire to perform opportunistic behavior whenever it is a rational alternative.

There are also relations between agents' ability and desire of performing an action. As rational agents, firstly they think about what actions they can perform given the limited knowledge they have about the state, and secondly they choose the action that may maximize their utilities based on their partial knowledge. This practical reasoning in decision theory can also be applied to reasoning about opportunistic propensity. Given the asymmetric knowledge an agent has, there are several (possibly opportunistic) actions available to him, and he may choose to perform the action which is a rational alternative to him, regardless of the result for the other agents. Based on this understanding, we have the following theorem, which implies agents' opportunistic propensity:

Theorem 1 (Opportunistic propensity) *Given a multi-agent system $\mathcal{M}$, a state s , two agents i and j and an action a , agent i is likely to perform action a to agent j as opportunistic behavior in state s :*

$$a \in a_i^*(s) \text{ and } \mathcal{M}, s \models \text{Opportunism}(i, j, a)$$

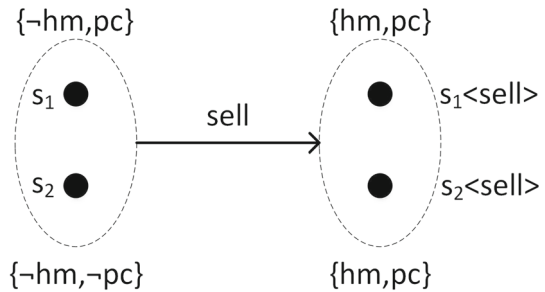
iff

1. $\forall t \in \mathcal{K}(i, s) : \mathcal{M}, t \models \text{promoted}(v^*, a) \wedge \text{demoted}(w^*, a), \exists t \in \mathcal{K}(j, s) : \mathcal{M}, t \models \neg(\text{promoted}(v^*, a) \wedge \text{demoted}(w^*, a))$, where $v^* = \text{highest}(i, s, s\langle a \rangle)$ and $w^* = \text{highest}(j, s, s\langle a \rangle)$;
2. $s <_i s\langle a \rangle$ and $s >_j s\langle a \rangle$.
3. $\neg \exists a' \in \text{Ac}(i, s) : a \neq a' \text{ and } a' \text{ dominates } a$.

Proof Forwards: If action a is opportunistic behavior, we can immediately have statement 1 by the definition of Knowledge Set. Because action a is in agent i 's rational alternatives in state s ($a \in a_i^*(s)$), by Definition 7, action a is not dominated by any action in $\text{Ac}(i, s)$. Also because action a is opportunistic, by Proposition 5 it results in promoting agent i 's value but demoting agent j 's value ($s <_i s\langle a \rangle$ and $s >_j s\langle a \rangle$). **Backwards:** Statement 1 means that there is knowledge asymmetry between agent i and agent j about the formula $\text{promoted}(v^*, a) \wedge \text{demoted}(w^*, a)$. From this we can see the knowledge asymmetry is the precondition of action a . If this precondition is satisfied, agent i can perform action a . Moreover, by statement 2, because action a promotes agent i 's value but demotes agent j 's value, we can conclude that action a is opportunistic behavior. By statement 3, because action a is not dominated by any action in $\text{Ac}(i, s)$, it is a rational alternative for agent i in state s to perform action a . $\square$

Given an opportunistic behavior a , in order to predict its performance, we should first check the asymmetric knowledge that agent i has for enabling its performance. Based on agent i 's and agent j 's value systems, we also check if it is not dominated by any actions in $\text{Ac}(i, s)$ and its performance can promote agent i 's value but demote agent j 's value. It is important to stress that Theorem 1 doesn't state that an agent will for sure perform opportunistic behavior if the three statements are satisfied. Instead, it states opportunism is likely to happen because it is one of the agent's rational alternatives. The agent will perform one action, which might be opportunistic behavior, from his rational alternatives.

Fig. 4 Variation of selling a broken cup



5.2 Not having opportunism

As Theorem 1 shows, we need much information about the system to predict opportunism, and it might be difficult to achieve all of them. Fortunately, in some cases it is already sufficient to know that opportunism is impossible to occur. An example might be detecting opportunism: if we already know in which context agents cannot perform opportunistic behavior, there is no need to set up any monitoring mechanisms for opportunism in those contexts. The following propositions characterize the contexts where there is no opportunism:

Proposition 8 Given a multi-agent system $\mathcal{M}$, a state s , two agents i and j and an action a ,

$$\mathcal{K}(i, s) = \mathcal{K}(j, s) \Rightarrow \mathcal{M}, s \models \neg \text{Opportunism}(i, j, a).$$

Proof When $\mathcal{K}(i, s) = \mathcal{K}(j, s)$ holds, which means that both agent i and agent j have the same knowledge. In this context, Statement 1 in Theorem 1 is not satisfied, so action a is not opportunistic behavior. $\square$

Proposition 9 Given a multi-agent system $\mathcal{M}$, a state s , two agents i and j and an action a ,

$$V_i = V_j \Rightarrow \mathcal{M}, s \models \neg \text{Opportunism}(i, j, a).$$

Proof When $V_i = V_j$ holds, which means that both agent i and agent j have the same value system. In this case, the values of both agents don't go opposite, that is, Statement 2 in Theorem 1 is not satisfied. So action a is not opportunistic behavior. $\square$

The above two propositions show that opportunism is impossible to occur when there is no knowledge asymmetry between agents and they share the same value systems. After we defined opportunism, we had Proposition 6 showing that two agents have different value systems as a property of opportunism. Together with Propositions 8 and 9, it looks like once having two different value systems and knowledge asymmetry about the value changes are satisfied one agent will perform opportunistic behavior to the other agent. Now let us go back to the example of selling a broken cup, the buyer's value gets demoted along the state transition, because he wants to have a good cup for use, which he finally doesn't have. Suppose the buyer only cares about appearance in the deal: as we show in Fig. 4, the buyer knows it is a pretty cup before he buys it, denoted as pc , and he gets a pretty cup (probably not for use) after the seller sells it. In this case, the behavior performed by the seller will not be seen as opportunistic behavior. From this variation, we notice that sometimes an action might not be seen as opportunistic behavior even though the agents involved have different value systems, because the two value systems are compatible rather than conflicting. This brings us to the notion of *compatibility*. Intuitively, compatibility describes a state in which

two or more things are able to exist or work together in combination without problems or conflict. We then propose the notion of *compatibility of value systems* with respect to a state transition.

Definition 10 (*Compatibility of value systems*) Given a multi-agent system $\mathcal{M}$, a state transition (s, a, s') and two value systems V_i and V_j ($V_i \neq V_j$), the two value systems are compatible with respect to transition (s, a, s') if and only if $\mathcal{M}, s \models \neg(\text{promoted}(v^*, a) \wedge \text{demoted}(w^*, a))$, where $v^* = \text{highest}(i, s, s')$ and $w^* = \text{highest}(j, s, s')$.

From this definition we have $s <_i s'$ and $s >_j s'$ don't hold at the same time, which means that the values of two agents don't go opposite (one gets promoted and the other one gets demoted) along a transition if their value systems are compatible with respect to the transition. Now we can relate the notion of *compatibility of value systems* to predicting opportunism. The following proposition characterizes another context where opportunistic behavior will not occur:

Proposition 10 *Given a multi-agent system $\mathcal{M}$ with a state s , two agents i and j and an action a , if value system V_i and V_j are compatible with respect to (s, a, s') , then*

$$\mathcal{M}, s \models \neg \text{Opportunism}(i, j, a).$$

Proof This proposition holds because two compatible value systems with respect to transition (s, a, s') will not lead to the result that one agent's value gets promoted and the other agent's value gets demoted ($s <_i s'$ and $s >_j s'$). By Theorem 1, it implies that action a will not be opportunistic behavior. $\square$

5.3 Computational complexity

Theorem 1 shows that whether a given action will be performed by an agent as opportunistic behavior. More generally, we would like to know given a multi-agent system we design, whether there exists opportunistic behavior between agents and how difficult it is to check it. In this section, we will investigate this issue through proposing an algorithm. The decision problem associated with predicting opportunistic behavior is as follows:

PREDICTING OPPORTUNISM

Given: Multi-agent system $\mathcal{M}$.

Question: Does there exist opportunistic behavior between agents for $\mathcal{M}$?

Theorem 2 *Given a multi-agent system $\mathcal{M}$, the problem that whether there exists opportunistic behavior between agents for $\mathcal{M}$ can be verified in $O(nmk^2)$ time, where n is the number of transitions, m is the maximum number of available actions in any given state and k is the maximum size of an S5 equivalence class.*

Proof In order to prove it, we need to find an algorithm that allows us to solve the decision problem in polynomial-time. We design Algorithm 1 for verifying opportunistic behavior in a multi-agent system $\mathcal{M}$ based on Theorem 1. The algorithm loops through all the possible transitions in the system, which has complexity $O(n)$, where $n = |\mathcal{R}|$. Notice that transitions

are executed by hypothetical agents, meaning that the value systems we consider for the transition is assumed to be known once the transition is given. For each transition, it verifies the statements listed in Theorem 1 one by one. Line 21–24 is to verify whether there is no action a' that dominates action a . Based on the definition of dominance between actions, the algorithm has to perform the comparison $\mathcal{K}(i, s\langle a \rangle)$ with $\mathcal{K}(i, s\langle a' \rangle)$ for all a' in $Ac(i, s)$. If for all $s' \in \mathcal{K}(i, s\langle a \rangle)$ and for all $s'' \in \mathcal{K}(i, s\langle a' \rangle)$ we have $s' \prec s''$, then action a is dominated by action a' . Hence, the complexity of executing line 21–24 is $O(mk^2)$, where $m = |Ac(i, s)|$ and $k = |\mathcal{K}(i, s)|$. The computational complexity of the whole algorithm is $O(nmk^2)$, which implies that Algorithm 1 can check whether there exists opportunistic behavior between agents for a given multi-agent system in polynomial-time. $\square$

Algorithm 1 Predicting Opportunism

```

1: procedure HASKNOWASYM( $S_1, S_2, \pi, \varphi$ ) returns true or false
2:   set  $g_1 \leftarrow \text{true}$ 
3:   set  $g_2 \leftarrow \text{false}$ 
4:   for each  $s \in S_1$  do
5:     if  $\varphi \notin \pi(s)$  then
6:       set  $g_1 \leftarrow \text{false}$ 
7:       break
8:   for each  $s \in S_2$  do
9:     if  $\neg\varphi \in \pi(s)$  then
10:      set  $g_2 \leftarrow \text{true}$ 
11:      break
12:   return  $g_1 \wedge g_2$ 
13:
14: procedure PREDICTING( $\mathcal{M}$ ) returns true or false
15:   set flag  $\leftarrow \text{false}$ 
16:   for each  $(s, a, s\langle a \rangle) \in \mathcal{R}$  do
17:     set  $v^* \leftarrow \text{highest}(i, s, s\langle a \rangle)$ 
18:     set  $w^* \leftarrow \text{highest}(j, s, s\langle a \rangle)$ 
19:     if HASKNOWASYM( $\mathcal{K}(i, s), \mathcal{K}(j, s), \pi, \text{promoted}(v^*, a) \wedge \text{demoted}(w^*, a)$ ) then
20:       if  $\text{promoted}(v^*, a) \wedge \text{demoted}(w^*, a) \in \pi(s)$  then
21:         set  $h \leftarrow 0$ 
22:         for each  $a' \in Ac(i, s)$  do
23:           if  $a \neq a'$  and  $\mathcal{K}(i, s\langle a \rangle) \preceq \mathcal{K}(i, s\langle a' \rangle)$  then
24:              $h++$ 
25:         if  $h == 0$  then
26:           set flag  $\leftarrow \text{true}$ 
27:           break
28:   return flag
  
```

In this section, we specified the situation where agents are likely to perform opportunistic behavior and characterized the contexts where opportunism is impossible to happen. This information is essential not only for the system designers to identify opportunistic propensity, but also for an agent to decide whether to participate in the system given his knowledge about the system and his value system, as his behavior might be regarded as opportunistic. Finally, we prove the computation complexity of predicting opportunism given a multi-agent system.

6 Discussion

From the definition of Function highest, we know that agent i only cares about the value change that he most prefers and ignores other value changes for defining his state preference. Hence, if we interpret value promotion as happiness and value demotion as sadness, this approach can be seen as the weight between the agent's happiness and sadness from the states: he prefers state s' rather than state s because his most preferred value gets promoted thus the happiness he gets is more than the sadness for being in state s' instead of state s . When talking about actions, $s \prec_i s(a)$ for instance, because among all the value changes agent i 's most preferred value gets promoted when going from state s to state $s(a)$, we can say that he feels more happy than sad by performing action a (apparently $a \neq sta$) instead of doing nothing. This interpretation is of great importance for the design of mechanisms for eliminating opportunism: if we want to make it not optimal for an agent to be opportunistic, the sadness he will get from it must be higher than the happiness, which implies that the value change that is most cared about by the agent must be demotion.

Moreover, our approach can be used in practice. For instance, in the electronic market place, we usually want to buy a jacket in good quality. Since we only can see the picture online but not the actual product, only the seller knows whether the jacket has good quality or not before he ships it. In this context, if earning money is most important to the seller and he is also aware that we want to buy a jacket in good quality, he will sell the bad jacket to the buyer by putting a quality-jacket picture online but shipping a bad one to us. The seller can do it, because there is knowledge asymmetry between the seller and us; the seller wants to do it, because earning money is the most important to him. According to Theorem 1, the seller will (likely) perform opportunistic behavior, selling the bad jacket, to the buyer. Since we are cautious about this problem from the system perspective, monitoring and eliminating mechanisms should be put in the right place in order to demotivate such a behavior. Imagine that if we can ensure that both the seller and us are aware of the quality of the product before the seller ships it or allow the buyer to grade the seller based on his shopping experience afterwards, we can demotivate the seller to get benefits from the buyer. In Fig. 5 with the same denotation as our previous examples, a middle party is involved to reveal to the buyer that the jacket is broken, denoted as $\text{reveal}(\text{broken})_b$, in order to ensure that both parties have the same knowledge about the jacket. Knowledge asymmetry is removed, so the seller cannot sell the bad jacket to the buyer with a normal price. One can refer to [9] for more elaboration about eliminating opportunism through removing knowledge asymmetry.

7 Related work

In order to investigate the interaction between different types of agents, agents are designed to be egoistic or altruistic depending on whether their internal decision processes are non-cooperative or cooperative. For example, [18] experiments on iterated prisoner's dilemma with a society of agents that are egoistic, reciprocating or altruistic. Golle et al. [19] designs incentive mechanisms for sharing with a population consisting of altruistic and egoistic agents, and [20] identifies egoistic or altruistic parties in terms of trust in open systems. In this paper, we presented agents' decision-making based on their value systems, which might have different ranking on their own value and others' value. Since opportunistic agents try to promote their own value at most but ignore other agents' value, they can be categorized as egoistic agents.

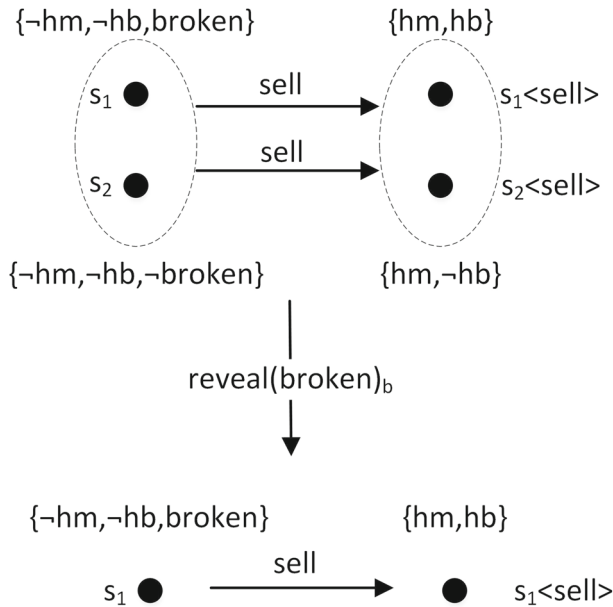


Fig. 5 A middle party reveals the information about the jacket to the buyer

The technical framework we used in this paper is a transition system extended with value systems. As standards for specifying preferences, people usually use goals rather than value (e.g. [21,22]) in logic-based formalization and utilities in decision theory and game theory (e.g. [23,24]) for the same purpose. Only some work in the area of argumentation reasons about agents' preferences and decision making by values (e.g. [25–27]). Goals are concrete and should be specified with time, place and objects, while value is relatively stable and not limited to be applied in a specific situation. Since state transitions are caused by the performance of actions, we can evaluate actions by whether our value is promoted or demoted in the state transition. For representing agents' evaluation on states, Keeney and Raiffa proposed Multi-Attribute Utility Theory (MAUT) in which states are described in terms of a set of attributes and the utilities of the states are calculated by the sum of the scores on each attribute based on agents' value systems [28]. Apparently, not everything can be evaluated with numbers, which is one of the reasons why people consider using value systems as an alternative. Bench-Capon et al. [25] already pointed out that utility-based decision-mechanisms in game theory cannot represent agents' decision theory in a real way. A value system is like a box that allows us to define its content as we need. In this paper, we use values and value systems as the basis for agents' choice. A value is modeled as a formula in our language and a value system is constructed as a total order over a set of values. Instead of calculating the utility of states, agents specify their preferences over states by evaluating the truth value of the state property that they most care about.

Our decision theory is extended with knowledge and value systems, which correspond to concepts from game theory [29]. In game theory, agents can be situated in a game which is not fully observable. Hence, it is natural to study agents' decision-making through combining game theory and epistemic logic. The notion of information sets is introduced to represent the states that the agent cannot distinguish [17]. In this paper, we use a similar concept *knowledge set* to represent the set of states that the agent considers as possible. Based on

the representation of uncertainty, we use the notion of dominance to compare two different actions: a dominated action is an action that is always bad to perform regardless of the uncertainty about the system, which is an approach bridging to (non-)dominated strategies in game theory. Typically, the concept of rational alternatives is tightly related with the concept of weak rationality as defined in the context of epistemic game theory. As specified in [30–32], players may not know their action is best, but they can know that there is no alternative action which they know to be better, given their limited knowledge about the current state. Both concepts represent a set of actions that are not dominated by other actions. It is thus already seen that we can apply techniques from game theory based on the concept similarities to enrich the existing decision theory and enhance the reasoning capabilities on agents' opportunistic propensity.

8 Conclusion and future work

The investigation about opportunism is still new in the area of multi-agent system. We ultimately aim at designing mechanisms to eliminate such selfish behavior in the system. In order to avoid over-assuming the performance of opportunism so that monitoring and eliminating mechanism can be put in place, we need to know in which context agents are likely to perform opportunistic behavior. In this paper, we argue that agents will behave opportunistically when they have the ability and the desire of doing it. With this idea, we developed a framework of multi-agent systems to reason about agents' opportunistic propensity without considering normative issues. Agents in the system were assumed to have their own value systems. Based on their value systems and incomplete knowledge about the state, agents choose one of their rational alternatives, which might be opportunistic behavior. With our framework and our definition of opportunism, we characterized the situation where agents are (not) likely to perform opportunistic behavior and prove the computational complexity of predicting opportunism. It is developed assuming that system designers are aware of agents' value systems. For sure system designers have no access to agent internals, but system designers have possible agent internals modeled by their value systems, which allows them to reason about the possibility of opportunism in the system. It is also important to stress that what we are trying to address in this paper is not what opportunism is, but whether we can somehow specify under which circumstance such a phenomenon can occur. In other words, we set the foundation of predicting whether a certain system is desired in terms of not having opportunism. Certainly there are multiple ways to extend our work. One interesting way is to enrich our formalization of value system over different sets of values, and the enrichment might lead to a different notion of the compatibility of value systems and different results about opportunistic propensity. The assumption that the value systems are common knowledge among the agents can be relaxed. We have no doubt that there exists alternative solutions to deal with opportunism such as through monitoring and enforcement norms. However, we need a predictive model to reason about opportunistic propensity in order to better allocate those mechanisms. Because of this predictive nature, it is a natural choice to use logical framework to reason about hypothetical situations. We presented a basic logical framework to reason about opportunistic propensity without considering any social mechanisms. Future work can consider issues such as norms, reputation, warranties and contracts in combination with the ability and the desire of being opportunistic. Most importantly, this paper set up a basic framework to design mechanisms for eliminating opportunism.

Acknowledgements The research is supported by China Scholarship Council. We would like to thank Allan van Hulst and anonymous reviewers for their helpful comments.

Open Access This article is distributed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license, and indicate if changes were made.

References

- Williamson, O. (1983). *Markets and hierarchies: Analysis and antitrust implications: A study in the economics of internal organization. A study in the economics of internal organization*. New York: Free Press.
- Bachmann, R., & Akbar, Z. (Eds.). (2006). *Handbook of trust research*. Cheltenham: Edward Elgar Publishing.
- Conner, K. R., & Prahalad, C. K. (1996). A resource-based theory of the firm: Knowledge versus opportunism. *Organization Science*, 7(5), 477–501.
- Jiraporn, P., et al. (2008). Is earnings management opportunistic or beneficial? An agency theory perspective. *International Review of Financial Analysis*, 17(3), 622–634.
- Cabon-Dhersin, M.-L., & Ramani, S. V. (2007). Opportunism, trust and cooperation: A game theoretic approach with heterogeneous agents. *Rationality and Society*, 19(2), 203–228.
- Luo, J., & Meyer, J. J. (2016). A formal account of opportunism based on the situation calculus. *AI & Society*, 4, 1–16.
- Chen, C. C., Peng, M. W., & Saporito, P. A. (2002). Individualism, collectivism, and opportunism: A cultural perspective on transaction cost economics. *Journal of Management*, 28(4), 567–583.
- Luo, J., Meyer, J. J., & Knobbout, M. (2016). Monitoring opportunism in multi-agent systems. In *Coordination, Organizations, Institutions, and Norms in Agent Systems XII* (pp. 119–138). Cham: Springer.
- Luo, J., Knobbout, M., & Meyer, J.-J. (2018). Eliminating opportunism using an epistemic mechanism. In *Proceedings of the 17th international conference on autonomous agents and multiagent systems* (pp. 1450–1458). *International foundation for autonomous agents and multiagent systems*.
- Moore, R. C. (1980). *Reasoning about knowledge and action*. Menlo Park: SRI International.
- Moore, R. C. (1984). *A formal theory of knowledge and action*. Technical report, DTIC Document.
- Scherl, R. B., & Levesque, H. J. (2003). Knowledge, action, and the frame problem. *Artificial Intelligence*, 144(1–2), 1–39.
- Shapiro, S., et al. (2000). Iterated belief change in the situation calculus. In *KR*.
- Poole, D. L., & Mackworth, A. K. (2010). *Artificial intelligence: Foundations of computational agents*. Cambridge: Cambridge University Press.
- Bulling, N., & Dastani, M. (2016). Norm-based mechanism design. *Artificial Intelligence*, 239, 97–142.
- Ågotnes, T., van der Hoek, W., & Wooldridge, M. (2007). Normative system games. In *Proceedings of the 6th international joint conference on Autonomous agents and multiagent systems* (p. 129). ACM.
- Dixit, A. K., & Nalebuff, B. (2008). *The art of strategy: A game theorist's guide to success in business & life*. New York: WW Norton & Company.
- Bazzan, A. L. C., Bordini, R. H., & Campbell, J. A. (2002). Evolution of agents with moral sentiments in an iterated Prisoner's Dilemma exercise. In *Game theory and decision theory in agent-based systems* (pp. 43–64). Boston: Springer.
- Golle, P., et al. (2001). Incentives for sharing in peer-to-peer networks. In *Electronic commerce* (pp. 75–87). Berlin: Springer.
- Schillo, M., Funk, P., & Rovatsos, M. (2000). Using trust for detecting deceitful agents in artificial societies. *Applied Artificial Intelligence*, 14(8), 825–848.
- Cohen, P. R., & Levesque, H. J. (1990). Intention is choice with commitment. *Artificial Intelligence*, 42(2–3), 213–261.
- Rao, A. S., & Georgeff, M. P. (1991). Modeling rational agents within a BDI-architecture. In *KR 91* (pp. 473–484).
- Steele, K., & Stefánsson, H. O. (2016). Decision theory. In E. N. Zalta (Ed.), *The stanford encyclopedia of philosophy* (winter 2016 edition). <https://plato.stanford.edu/archives/win2016/entries/decision-theory/>.
- Von Neumann, J., & Morgenstern, O. (2007). *Theory of games and economic behavior*. Princeton: Princeton University Press.

25. Bench-Capon, T., Atkinson, K., & McBurney, P. (2012). Using argumentation to model agent decision making in economic experiments. *Autonomous Agents and Multi-Agent Systems*, 25(1), 183–208.
26. Van der Weide, T. (2011). Arguing to motivate decisions. Ph.D. thesis, Utrecht University.
27. Pitt, J., & Artikis, A. (2015). The open agent society: Retrospective and prospective views. *Artificial Intelligence and Law*, 23(3), 241–270.
28. Keeney, R. L., & Raiffa, H. (1993). *Decisions with multiple objectives: Preferences and value trade-offs*. Cambridge: Cambridge University Press.
29. Myerson, R. (1991). *Game theory: Analysis of conflict*. Cambridge: Harvard University Press.
30. Van Benthem, J. (2007). “Erratum:” Rational dynamics and epistemic logic in games. *International Game Theory Review*, 9(02), 377–409.
31. Lorini, E., & Schwarzenruber, F. (2010). A modal logic of epistemic games. *Games*, 1(4), 478–526.
32. Bonanno, G. (2008). A syntactic approach to rationality in games with ordinal payoffs. In *Proceeding of LOFT* (pp. 59–86).

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.