



Some Remarks on the Semantics of FIPA's Agent Communication Language

JEREMY PITT AND ABE MAMDANI

{j.pitt,a.mamdani}@ic.ac.uk

Intelligent & Interactive Systems Group, Department of Electrical & Electronic Engineering, Imperial College of Science Technology and Medicine, London SW7 2BT, England

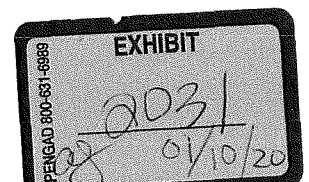
Abstract. The Foundation for Intelligent Physical Agents (FIPA) standardisation body has produced a set of specifications outlining a generic model for the architecture and operation of agent-based systems. The FIPA'97 Specification Part 2 is the normative specification of an Agent Communication Language (ACL) which agents use to 'talk' to each other. The FIPA ACL is based on speech act theory. Its syntax is defined by performatives parameterised by attribute value pairs, while its semantics is given in terms of the mental states of the communicating agents (i.e. intentionality). However, it is not clear if the formal semantics is meant as a normative or informative specification. The primary purpose of this paper is then to give an expository analysis of the FIPA ACL semantics to clarify this situation. We also offer some guidelines motivated from our own analysis, experience and understanding of how the semantic definitions and logical axioms should be interpreted and applied. However, our conclusion is that while the FIPA ACL specification offers significant potential to a developer using it for guidance, there are limitations on using an agent's mental state to specify the meaning of a performative as part of a normative standard. We consider some possibilities for making improvements in this direction.

Keywords: agent communication languages, semantics, standards, FIPA

1. Introduction

The Foundation for Intelligent Physical Agents (FIPA) standardisation body has produced a set of specifications outlining a generic model for the architecture and operation of agent-based systems. The FIPA'97 Specification Part 2 is the normative specification of an Agent Communication Language (ACL) which agents use to 'talk' to each other. This language is based on speech act theory and defined in terms of a set of performatives or communicative acts. A communicative act occurs whenever one agent sends a message to another.

Each communicative act is given a meaning informally via English descriptions, and a formal semantics in a first order modal logic of mental attitudes and actions, which define feasibility preconditions and rational effects (intended outcomes) of communicative acts, and various agent properties. The logical definitions generally characterise the belief state of the sending agent before performing the communicative act, and the intended belief state of the receiving agent after the act has been performed. Underlying these definitions are a number of properties which an agent can use either to plan a communicative act (i.e. send a message), or to make inferences from 'hearing'/'observing' a communicative act (i.e. receiving a message). These properties are formalised as axioms of the logic. Details of the underlying theory can be found in [20, 21, 22].



The FIPA ACL is a substantial technical achievement, but its key components are its performatives and its protocols. The intention of the semantics is to provide guidance for developers to produce systems that will be able to inter-operate co-operatively, and which can be formally checked for compliance. However, the semantics is very subtle in the way that its various properties interact, and the rationale behind some of the technical formulation is sometimes obscured by implicit assumptions and arcane notation. Clearly, the FIPA authors were very concerned to provide the right balance between a standard that was mathematically well-founded but also useful to non-mathematically inclined agent developers. In our experience on the MARINER project [17], the performatives and protocols were a useful starting point in developing a FIPA-compliant multi-agent system. We encountered, however, a number of difficulties with the interpretation and application of the formal semantics to implement co-operative behaviour through communicative acts (notwithstanding any issues in verifying compliance [28]).

Our primary criticism of the FIPA ACL standard is that it is not clear whether the ACL semantics should be interpreted as an informative specification providing guidance to developers, or as a normative specification providing conditions that the agents themselves are responsible for satisfying. The primary function of this paper is to offer some guidelines motivated from our own experience and understanding of how the semantic definitions and logical axioms should be interpreted and applied. We feel this is important and necessary because, as is well known from natural language pragmatics, any formalisation of communication based on the mental states of the participants must deal with what is fundamentally a process of revision and updating of those states. It is unlikely that a single set of axioms will cover all eventualities because communication is inherently context-dependent. As a result, what works well in one application may be inappropriate in another.

The authors of the FIPA specification have addressed this issue by specifying only the preconditions and intended outcomes of the given performatives and a limited number of axioms. The FIPA specification is then minimal with respect to what two agents can infer about each other's mental states after the performance of some communicative act. For example, a FIPA agent that 'observes' a communicative act may infer that the (persistent) preconditions for doing that action hold, and that the intended outcome of the act was indeed intended. However, what is not part of the specified semantics is that, after doing an action, the performing agent believes that the intended effect of (i.e. goal of, or reason for doing) that action has resulted. Moreover, the actual outcome of a communicative act, i.e. the change in the belief state of the receiving agent, is also not part of the semantics.

This 'under-specification' is not always sufficient if the standard formal semantics is to have practical use. The standard is effectively giving designers a partial axiomatisation of FIPA's performatives: saying, in effect, that the given axioms apply, it is then up to the designers to find a model that is consistent, and their own frame axioms. An agent developer may need further axioms to derive other logical consequences resulting from communicative acts, and may need variations on the specified semantics for given performatives with the same intuitive meaning. This paper, after a précis of the FIPA ACL semantics in Section 2, is concerned with exploring and explaining the ramifications of this, particularly with regard to some primitive

FIPA performatives. Section 3 undertakes an abstract analysis of the inform performative and certain axiomatic properties, and offers some guidelines for using the specification. Section 4 reports on practical experience from the MARINER project, and some variations associated with the semantics of the confirm and disconfirm performatives. Again we give guidelines on applying the specification for practical system development.

In conclusion, we would argue that the problems caused by variations in agent motivation and intent, communicative context, and application domain expose the limitations of intentionality as a basis for standardizing the semantics of an agent communication language. Intentionality is concerned with agent internals and a communicative act is an external phenomenon: mental attitudes can then only give a reason for and not the meaning of the performatives. In this way it can give guidance to developers but may be too strong a constraint on the agent. The protocols specified by FIPA ACL, on the other hand, can give purpose and meaning to individual performatives but only in the context of a conversation that has an identifiable objective. Section 5 provides further discussion on our understanding of the FIPA ACL specification and its application, but is also concerned with what steps might be taken to produce an ACL semantics that might resolve some of the criticisms of intentionality-based agent communication, both herein and elsewhere (e.g. [26]).

2. A précis of the FIPA ACL formal semantics

We are concerned with exploring the formal, (proposed) standard semantics of the FIPA ACL inform, confirm and disconfirm performatives. In this section, we give a summary of the formal semantics, discuss the motivation for the semantics being formulated as it is, and try to explain the difference in perception of these semantics from the developer's viewpoint as opposed to the agent's viewpoint.

The formal semantics of the FIPA ACL inform, confirm and disconfirm performatives is summarised in Table 1. An intuition of the semantics and explanation of the notation follows.

Intuitively, the inform performative is used by some agent s (say) to tell another agent r (say) that some proposition ϕ (say) is true. We will refer to the agent performing the communicative act, in this case s , as the sending agent, and the agent towards whom the act is performed, in this case r , as the receiving agent. This reflects the reality that the agents in a multi-agent system are in fact sending

Table 1. FIPA ACL performatives: formal semantics

Performative	Feasibility preconditions		Rational effect
	FP1	FP2	RE
$\langle s, \text{inform}(r, \phi) \rangle$	$\mathcal{B}_s \phi$	$\neg \mathcal{B}_s (\mathcal{B}if_r \phi \vee \mathcal{U}if_r \phi)$	$\mathcal{B}_r \phi$
$\langle s, \text{confirm}(r, \phi) \rangle$	$\mathcal{B}_s \phi$	$\mathcal{B}_s \mathcal{U}_r \phi$	$\mathcal{B}_r \phi$
$\langle s, \text{disconfirm}(r, \phi) \rangle$	$\mathcal{B}_s \neg \phi$	$\mathcal{B}_s (\mathcal{B}_r \phi \vee \mathcal{U}_r \phi)$	$\mathcal{B}_r \neg \phi$

and receiving electronic messages. The **confirm** performative has much the same intuition as **inform**, but different conditions on using it (see below). The **disconfirm** performative is used by some agent to tell another agent that some proposition is not (or is no longer) true.

In FIPA, agents either believe propositions or are uncertain about propositions. Belief and uncertainty are indicated by relativised modalities that, for agent s , would be $\mathcal{B}_s$ and $\mathcal{U}_s$ respectively. Belief and uncertainty are mutually exclusive: knowledge in FIPA's logic is an abbreviation for 'belief or uncertainty.' The abbreviation $\mathcal{B}if_s \phi$ is defined as $\mathcal{B}_s \phi \vee \mathcal{B}_s \neg \phi$ and $\mathcal{U}if_s \phi$ as $\mathcal{U}_s \phi \vee \mathcal{U}_s \neg \phi$.

The feasibility preconditions (*FP*) of a performative are those conditions that need to be true in order for an agent to (plan to) execute a communicative act. The first feasibility precondition for **inform** (*inform FP1*) therefore states that the agent performing the communicative act, the sending agent s , believes that the proposition ϕ is true. This precondition aims to ensure the so-called sincerity condition, i.e. that agents believe what they say and only say what they believe. The second feasibility precondition for **inform** states that the sending agent, s , does not believe that the intended recipient, r , either believes or is uncertain that either proposition ϕ or its negation is true. Thus *inform FP2* is intended to be used in goal formulation, so that if $\mathcal{B}_s \mathcal{B}_r \phi$, then agent s will not try to inform r of something, i.e. ϕ , that s believes that r already believes.

This means that if agent s believes that agent r either believes ϕ or is uncertain about ϕ , then no further **inform** speech act is required and so will not be planned. Otherwise, if agent s believes that proposition ϕ holds and believes that agent r is uncertain about ϕ , then s can plan to use the **confirm** performative to tell r that ϕ . Under other conditions, an agent may plan to use the **disconfirm** performative to refute propositions. Thus *disconfirm FP1* states that the sending agent s believes that proposition ϕ holds, while *disconfirm FP2* states that the sending agent s believes the intended recipient r believes or is uncertain that ϕ holds (i.e. has some prior knowledge about ϕ).

In addition to the conditions which must hold in order to plan a speech act, there are reasons for doing that act. These reasons are called Rational Effects (*RE*). The *RE* of a performative are, according to the FIPA specification, the reasons for which the act is selected, i.e. the intention that motivates the communicative act. From this perspective, it is a kind of pre-condition used in goal formulation. However, an agent "will select acts based on the relevance of the act's expected outcome or rational effect to its goals" ([8], Part 2v2.0, p46), and from this description, the rational effect can be identified with an intended actual outcome and so could be interpreted as a specification for a kind of post-condition. It is important to note that the latter reading is not intended and applies to all three performatives.

For example, the reason for an agent s informing another agent r of proposition ϕ is that s desires that (wants that) r should believe that ϕ is true. The rational effect is an intended outcome, i.e. a specification of a goal for the sending agent, but it is not necessarily an actual outcome, i.e. a specification of a post-condition for the receiving agent. Whether or not the *RE* is the actual change in the mental state of the receiving agent is therefore not part of the formal semantics. The FIPA specification therefore deliberately (and rightly) under-specifies the actual outcome:

FIPA is not trying to mandate how an agent should change its mental state according to the content of the ACL messages that it receives. Presumably because of this, the sending agent also “cannot assume that the rational effect will necessarily result from sending the message[s]” (*ibid*, p46).

However, the FIPA ACL does specify five properties that are given as part of the underlying semantic model and specified as axioms of an agent. Three are related to the reasoning processes involved in planning communicative acts and two to reasoning about communicative acts that are (in some sense) ‘observed’ (i.e. by receiving a message). The latter two do provide some guidance for inference, but the inferences that can be made are about the act itself, rather than any content.

FIPA Properties 1, 2 and 3 are concerned with planning communicative acts and need not concern us here. They are really concerned with agent internals (intra-agent properties) and so, in our opinion, superfluous to a standard for specifying semantics for inter-agent communications. They could provide some guidelines for implementers employing a BDI (Beliefs-Desires-Intentions) architecture, but, on the other hand, there are many different styles of planning, which need not be related or even reducible to BDI-style planning and reasoning. The significance of Properties 1–3 emerge when considering Properties 4 and 5, which define inferences an agent can make from observing a communicative act.

Property 4 and Property 5 are defined as:

Property 4: $\models \mathcal{B}_i(DONE(A) \wedge Agent(j, A) \rightarrow \mathcal{I}_j RE(A))$

Property 5: $\models \mathcal{B}_i(DONE(A) \rightarrow FP(A))$

Property 4 states that if an agent i ‘sees’ that action A has taken place and agent j did action A , then i can infer that j had the intention to bring about the rational effect of A . Property 5 states that, from the same observation, agent i can infer that the feasibility preconditions of action A (that are unrelated to time) hold.

The combined effect of the FP , RE and the latter two properties are, for the inform performative, as follows. Suppose agent s believes a proposition ϕ and does not believe that agent r either believes ϕ or is uncertain about ϕ , or its negation. If s generates the intention for r to believe ϕ , then s will plan to use the inform performative to (try to) make r believe ϕ . However, agent s can only try to get an agent r to believe a proposition ϕ by sending an inform, but cannot infer or assume that this will be the result. Furthermore, agent r is not mandated by the semantics to change its belief state just because it receives an inform. As the FIPA specification states: “Whether or not the receiver does, indeed, adopt belief in the proposition will be a function of the receiver’s trust in the sincerity and reliability of the sender” (*ibid*, p25). Agent r can, on the other hand, make inferences about the act by virtue of the defined properties.

Therefore the FIPA specification is minimal with respect to what one agent can infer about the mental state of another agent after having communicated with it, and with respect to the mental state of the agent receiving a communication. Instead, the performance of a communicative act allows the receiver to derive inferences about the mental state of the sender using Properties 4 and 5. Speech act theory has been criticized before for being too ‘speaker-centric’, i.e. concentrating only on

the mental state of the speaker [4], and it is right that the FIPA semantics should consider the mental state of the receiver.

However, the consequences of this minimality are two-fold:

- It is not clear how the semantics and axioms should be used: when they are merely providing guidance for the developers and when they are requirements that the agents are responsible for satisfying; and
- We are dealing with a process of revision and updating of belief states of both sender and receiver, and, according to context, more axioms are needed to infer that some condition holds in either of the participants after a communicative act has been either performed or observed.

We try to illustrate these consequences in Figure 1 below. This shows the developer's perspective on the state Σ of a multi-agent system with agents s and r . Agent s is behaving as if Properties 1, 2 and 3 were implemented, the *FPs* were satisfied, and the *RE* its goal. All this is internal to the agent: what the developer 'sees' is the FIPA ACL message being passed from s to r . In the formal semantics, this corresponds to a modality, and the whole system changes state to Σ' . Properties 4 and 5 come into effect in deciding the new state of the receiving agent, r' . However, there are unspecified frame axioms which may also apply, and the sending agent is changing state too (to s').

Thus, from the developer's viewpoint, there are two important points to note. Firstly, that the semantics is largely *external* and the agents are behaving as if they were implemented in this way, irrespective of how they are actually implemented (cf. [6]). Secondly, that on performing a communicative act the whole system is changing state, including both sender and receiver. The FIPA axioms apply and it is up to the developers to find a consistent model and their own frame axioms for their application.

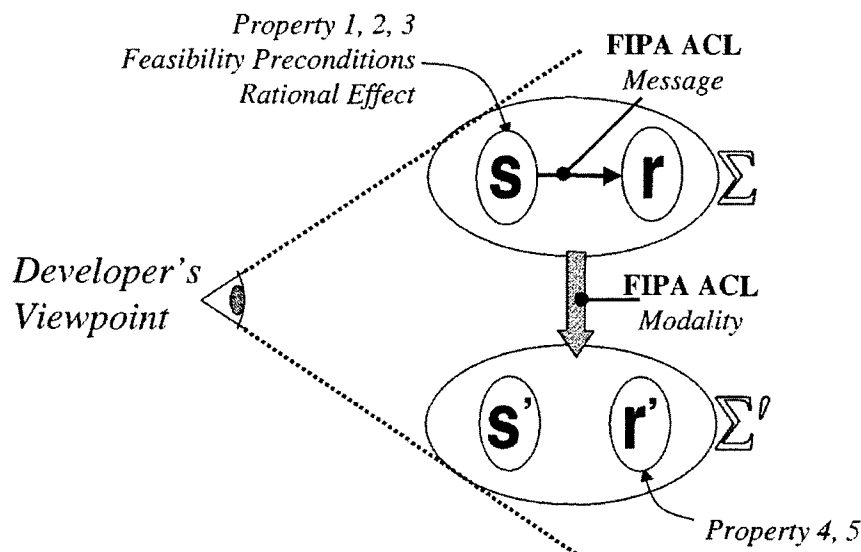


Figure 1. Developer's perspective on FIPA semantics.

The next two sections expand on the motivation for and consequences of this model in order to derive some guidelines to assist developers of agent-based systems using the FIPA ACL. In Section 3, we are concerned with two issues: firstly, the inclusion of the sincerity condition to constrain the *inform* action, and secondly, how an agent acquires information about the results of its *inform* actions. Section 4 briefly introduces the MARINER project, analyses the semantics of the *confirm* and *disconfirm* performatives in this context, and considers what to do if the semantics of the application don't quite fit the FIPA semantics.

3. The *inform* performative

3.1. *Inform FPI and the sincerity condition*

Recall the feasibility precondition *inform FPI*: $\mathcal{B}_s\phi$. This enforces the so-called sincerity condition, requiring that an agent believe whatever it informs another agent. However, there is a type of agent, nowadays called a self-interested agent, and a class of applications (involving competitive agents programmed by separate parties each with their own goals first and foremost [19]) where sincerity is not obligatory and can even be a decided disadvantage. It might therefore be necessary on occasions, for example some auctions, haggling, etc., to make provision for an intelligent agent to act sincerely but to be "economical with the truth." However, there are techniques (game theory and economic modelling, for example), for organizing the agents such that it can be shown that the system as a whole can not be manipulated even by agents that do act insincerely [23]. This is a property of the system, though, and not a feature of the standard semantics of the communication.

Furthermore, whether or not an agent abides by the sincerity condition is unlikely to be verifiable in practice. For example, an unscrupulous agent implementer could work-around the precondition in the following way. Suppose agent s does not believe ϕ but wants to inform r that ϕ . Then s adds ϕ to its database (*inform FPI* is satisfied), s performs $\langle \text{inform}(r, \phi) \rangle$ (s tells its lie), and then removes ϕ from its database.

Of course this is a trick, but an agent doing this is still technically complying with the semantic definitions. It is also not possible to distinguish between this behaviour and that of any agent with a mental state that is non-monotonic over time, even if all information is time stamped in some way (that is not explicit in the FIPA specification). It is even impossible to distinguish between a sincere agent that communicated something it believed to be true (that was actually false) and an insincere agent that communicated something it believed to be false (but was actually true).

The motivation for *inform FPI* is a pragmatic requirement for co-operative behaviour stemming from the original application of ARCOL (the progenitor of FIPA ACL), which was answering database queries [2]. The formalisation of intentional communicative behaviour for an intelligent agent often involves the axiomatisation of the felicity conditions of Searle and the conversational maxims of Grice. The sincerity condition underpinned Grice's analysis of conversational implicatures to

support the so-called co-operativity principle. Therefore, from the viewpoint of an agent developer who wants his/her agents to inter-operate co-operatively, the following facts are significant:

- that the sincerity condition is predicated on the notion of co-operation; and
- that all speech acts involve the receiver recognizing an intention on the part of the sender.

Therefore, when an agent uses an *inform* performative in the context of a negotiation protocol for some joint co-operation, the performing agent has to be sincere with respect to its beliefs, otherwise the receiving agent cannot reason (e.g. via Property 4 and 5) in order to co-operate. So long as the FIPA ACL semantics is aimed at developers of *co-operating* agents, then they will inter-operate as desired.

The problem with relaxing the sincerity condition is that it pulls down Property 5 with it. However, there is an intrinsic asymmetry in the FIPA specification. A FIPA agent is an imperfect witness of perfect integrity: it may communicate incorrect information but at least it believes it to be true. The receiver doesn't have to believe the content, but it is axiomatic that the sender must have believed it. If, however, whether or not the receiver actually believes the content is a function of trust or reliability of sender, surely then whether or not the receiver actually believes the sender's belief is a similar function of trustworthiness. Put another way: if you don't trust someone to tell you the truth, you don't trust them to believe what they tell you. It follows that Property 5 is not an axiom: it is just a reasonable assumption that follows from an initial assumption of an intention to co-operate.

What we see here is the danger of trying to standardise the meaning of a performative in isolation through an axiomatisation that implicitly expects the performative to be used in conversation. In isolation, for example, we could motivate the behaviour of a sincere agent by the behavioural axiom which includes *inform FPI*:

$$\models \mathcal{B}_s \phi \wedge \mathcal{D}_s \mathcal{B}_r \phi \rightarrow \mathcal{I}_s \text{DONE}(\langle s, \text{inform}(r, \phi) \rangle)$$

Here $\mathcal{B}$, $\mathcal{D}$ and $\mathcal{I}$ are respectively the beliefs, desires (goals) and intends modalities. *DONE* is an operator on actions, defined in the FIPA specification as *DONE(A)* being true after action *A* has taken place. This axiom therefore states that if the agent *s* believed ϕ and wanted another agent *r* to believe ϕ , then *s* would generate the intention to inform *r* of ϕ . Alternatively we could define the behaviour of a 'rapacious' agent by the axiom:

$$\models \mathcal{D}_s \psi \wedge \mathcal{B}_s (\text{DONE}(A) \rightarrow \psi) \rightarrow \mathcal{I}_s \text{DONE}(A)$$

This axiom states that if *s* desires (wants) ψ and believes that doing *A* will achieve ψ then *s* will intend to do *A*. If $\psi = \mathcal{B}_r \phi$, the same communication occurs (*A* is $\langle s, \text{inform}(r, \phi) \rangle$), without *s* having an explicitly held belief in ϕ . However, it is important to note that with the first axiom the agent is being sincere with respect to its beliefs, and with the second axiom it is being sincere with respect to its desires.

This use of axioms mixing belief and desire modalities goes back to the original work of [5], who explored different models of rationality using different axioms

combining modalities for beliefs, desires and intentions. This is not a particular feature of the FIPA semantics, although it is clear that the effects stipulated by the action schemas are not by themselves adequate to express the full semantics. Instead, both the action schemas and the axioms are required, an outcome perhaps to have been expected given Cohen and Levesque's earlier work.

3.2. Rational effects and actual outcomes

It is clear then what conclusions can be drawn from an inform communicative act via FIPA Properties 4 and 5: firstly, receiver's beliefs about the sender's intention to achieve the rational effect; and secondly, receiver's beliefs about the sender's belief of the content. We have seen in the previous section that the second conclusion is dependent on an implicit assumption of co-operativity. In this section, we address the question of how does an agent get to know the effect of its speech acts on the belief states of receiving agents, i.e. how does it know that its rational effect has been achieved?

We first note that the act of informing is very often, in agent-based systems, non-monotonic. Any agent that is embedded in an environment and taking actions in that environment, should be able to revise its beliefs, and intentions, based on changes in the environment or communications from other agents [10]. If that information conflicts with previously held beliefs, then the agent has to be able to select which information to believe.

As a result Property 5 also has some potentially curious effects if the information communicated is structured rather than propositional, i.e. more realistically, we are using first-order logic for knowledge representation. Suppose agent s believes that the colour of some object $obj1$ is red. It can inform agent r that $colour(obj1, red)$, and agent r will infer $\mathcal{B}_s colour(obj1, red)$ by Property 5. Now suppose that s comes to believe that the colour of this object is blue. The inform preconditions:

$$\mathcal{B}_s colour(obj1, blue) \wedge \neg \mathcal{B}_s (\mathcal{B}if_r colour(obj1, blue) \vee \mathcal{B}if_r \neg colour(obj1, blue))$$

are satisfied, so it can inform r . According to Property 5, r could now infer that agent s believes that $obj1$ is both red and blue. However, there is nothing in the FIPA semantics that says that the receiving agent does make this inference. If an object cannot be both red and blue (i.e., there is some frame axiom to this effect), with a KD45 belief modality the receiver will note that the sender believes the object to be blue. However, this frame axiom is *not* part of the FIPA semantics.

By the same token, the axiom that we call Property X is not part of the FIPA semantics and is defined by:

$$\text{Property } X: \models \mathcal{B}_i (DONE(A) \wedge Agent(i, A) \rightarrow RE(A))$$

This states that if an agent i believes that action A has occurred and s did A itself, then it can infer that the rational effect of that action holds. (Note that the axiom

$$\models \mathcal{B}_i (DONE(A) \rightarrow RE(A))$$

would force belief in the *RE* if agent *i* were either the sending or receiving agent, which is not the effect we want. If *j* were an agent and *A* were *inform*(*j*, ϕ), then *j*, as the recipient, could infer $\mathcal{B}_j\phi$, when it need not believe ϕ at all.) The purpose of this axiom is to give the sending agent a belief in the achievement of the rational effect of its actions. If it turns out that the belief is wrong (the *RE* does not hold), then the agent can do belief revision [10]. If it is the responsibility of agent implementers to implement the belief revision of the receiving agent (e.g. by defining frame axioms as above), it could of course equally be their responsibility to implement belief revision for the sending agent.

The rational effect of Property *X* gives the sending agent some idea of the intended resultant mental state of the receiving agent based on inference. Effectively, we are inferring from FIPA specification statement that an agent “cannot assume that the rational effect will necessarily result from sending the message[s]” (*ibid*, p46) the conclusion that an agent can assume that it may result.

It turns out that there even appears to be some requirement for Property *X*. This is because sending agents need some idea of the mental states of receiving agents, or otherwise *inform FP2*, $\neg\mathcal{B}_s(\mathcal{B}if_r\phi \vee \mathcal{U}if_r\phi)$, is vacuous. Using Property *X*, once an agent infers the *RE* of the *inform*, *FP2* is no longer satisfied. Without Property *X* and unless *s* believes $\mathcal{B}_r\phi$ (or even $\mathcal{U}_r\phi$) after performing an *inform*, *FP2* will not preclude formulating the goal again, and again, and so on. This is contrary to the intention of only communicating necessary information. Therefore an agent developer has to try to get the same result using the available semantic specification, or it implies that some aspects of message passing are pragmatic (e.g. an agent could maintain a message list and only repeats messages if they are definitely lost, etc.). Alternatively, the developer has to find his/her own solution to the under-specification.

Surprisingly, perhaps, having the receiving agent simply echo the content (via another *inform*) does not work because of Property 5: after $\langle s, \text{inform}(r, \phi) \rangle$ we get (for *r*) $\mathcal{B}_r\mathcal{B}_s\phi$, therefore *inform FP2* is not satisfied. An alternative is for *s* to ask *r* if it believes that ϕ is true. For this, *s* needs to use the performative $\langle s, \text{query_if}(r, \phi) \rangle$. *query_if* is defined as a macro action (i.e. defined in terms of other performatives), whereby the sending agent requests the receiver to inform it whether or not some proposition ϕ is true. The FIPA specifications are:

$$\begin{aligned} \langle s, \text{query_if}(r, \phi) \rangle &\equiv \langle s, \text{request}(r, \langle r, \text{inform_if}(s, \phi) \rangle) \rangle \\ \langle r, \text{inform_if}(s, \phi) \rangle &\equiv \langle r, \text{inform}(s, \phi) \rangle \mid \langle r, \text{inform}(s, \neg\phi) \rangle \end{aligned}$$

Here $\mid$ is the non-deterministic choice operator for actions.

It can be shown [20, 6] that the *FP* can be derived from the *FP* of the embedded actions. Presumably, then, from the embedded disjunctive *inform* the *FP* on the *inform_if* are (ignoring, for presentational simplicity, the uncertainty operator):

$$\begin{aligned} &\mathcal{B}_r\phi \vee \mathcal{B}_r\neg\phi \\ &\neg\mathcal{B}_r(\mathcal{B}_s\phi \vee \mathcal{B}_s\neg\phi) \vee \neg\mathcal{B}_r(\mathcal{B}_s\neg\phi \vee \mathcal{B}_s\neg\neg\phi) \end{aligned}$$

Now recall after $\langle s, \text{inform}(r, \phi) \rangle$ we get (for r) $\mathcal{B}_r \mathcal{B}_s \phi$. Consequently neither disjunct in the second precondition of `inform.if` is satisfied (assuming $\mathcal{B}_s \phi \leftrightarrow \mathcal{B}_s \neg \neg \phi$).

An alternative solution, rather than using Property X , and similar to that proposed by Smith et al. [27], is for r to enact the performative $\langle r, \text{inform}(s, \mathcal{B}_r \phi) \rangle$. However, the agent implementer must prevent cycles (e.g. $\langle s, \text{inform}(r, \mathcal{B}_s \mathcal{B}_r \phi) \rangle$ and so on) by operational means, and with Property 5, s will be able to infer formulas like $\mathcal{B}_s \neg \mathcal{B}_r (\mathcal{B}if_s \mathcal{B}_r \phi \vee \mathcal{U}if_s \mathcal{B}_r \phi)$, whose utility is not immediately obvious. There are therefore messaging, implementation and reasoning overheads associated with this solution, although one of the intended use of conversation policies [27] is to provide exactly this sort of operational control.

Yet another solution would be to define a new performative, say `assert`, with a semantics distinct from `inform` and `confirm` as shown in Table 2, which could be used in this context. Recall that the *Bif* operator is defined as $\mathcal{B}if_s \phi \leftrightarrow \mathcal{B}_s \phi \vee \mathcal{B}_s \neg \phi$. An agent would then use an `assert` instead of an `inform` if it wanted the intended receiver to make the weaker inference from Property 5.

We also note, *en passant*, another new performative `ask` with semantics as shown:

Feasibility precondition(s)	Rational effect	Performative
$\mathcal{B}if_s \phi$	$\mathcal{B}_r \phi$	$\langle s, \text{assert}(r, \phi) \rangle$
$\mathcal{U}_s \phi \wedge \mathcal{B}_s \mathcal{B}if_r \phi$	$\mathcal{I}_r \text{ DONE}(\langle r, \text{confirm}(s, \phi) \rangle $ $\langle r, \text{disconfirm}(s, \neg \phi) \rangle)$	$\langle s, \text{ask}(r, \phi) \rangle$

After an `ask`, the receiving agent can infer $\mathcal{B}_r \mathcal{U}_s \phi$, so the feasibility preconditions for `confirm` ($\mathcal{B}_r \phi \wedge \mathcal{B}_r \mathcal{U}_s \phi$) are satisfied if r believes ϕ ; or the feasibility preconditions for `disconfirm` ($\mathcal{B}_r \neg \phi \wedge \mathcal{B}_r (\mathcal{U}_s \phi \vee \mathcal{B}_s \phi)$) are satisfied if r believes $\neg \phi$. Furthermore, in each case the respective REs ($\mathcal{B}_s \phi$ and $\mathcal{B}_s \neg \phi$) satisfy the presumed goal behind the original `ask`, i.e. to resolve an uncertainty in $\mathcal{U}_s \phi$ one way or the other. Unfortunately this still does not enable an agent s to follow up a `inform` with an `ask`, although it does provide a neat alternative formulation of so-called yes-no queries and replies which could easily be standardised in the FIPA framework.

3.3. Guidelines for interpreting the `inform` performative semantics

The FIPA specification endeavours to provide a semantics in which communicative acts are motivated by changing the belief state of a receiving agent but constrained by the belief state of the performing agent. Inferences can only be drawn about the motivation and the constraints, and not the consequences. We say this is minimal because while one might intuitively expect a sending agent to hold a belief about the consequences of its actions, and for a receiving agent to hold new beliefs based on intended rational effects, the corresponding properties are not part of the semantics.

The previous two sections have shown that there is an implicit assumption of co-operation, a need to define additional axioms, and a subtle inter-play of actions and axioms which need to be considered when developing a FIPA-based multi-agent system using the ACL semantics.

To this end, we conclude this section with a (non-exclusive) list of guidelines for interpreting the inform performative semantics:

- Co-operation: that in order to co-operate an agent should behave as if it were attempting to satisfy the sincerity condition, although it does not require a specific representation or implementation of beliefs, desires or intentions.
- Choice of performative: that if the designer wants an agent to communicate new information that it (in some sense) believes to be true, then it should use an inform performative, as opposed to any other performative.
- Desired outcome: that in the absence of any counter-indication it is safe to assume that the desired outcome has been achieved (e.g. using Property *X*), but only if the agent can revise the assumption should it turn out to be incorrect.
- Recipient's perception: that among the possible inferences that can be made are that the content is true; that the sender believed it; and that the sender wanted the receiver to believe it.
- Context-sensitivity: that the inform performative is intended to be used in the context of protocols which deliver co-operative behaviour, and that additional axioms or actions may be required according to context, application, or 'one-off' communications.

4. The confirm and disconfirm performatives

4.1. *MARINER scenario 1*

The MARINER project is concerned with developing a FIPA compliant agent systems for load control in intelligent networks (IN) [1]. In MARINER scenario 1, there are agents associated with various network elements. One type of network element agents (the quantifiers) is concerned with computing the cost of an IN service to another type of network element agents (the allocators). Each quantifier communicates this cost to each allocator every 30 seconds or so. The allocators then know which network element to route service requests to. Intuitively, each quantifier should send each allocator an update message.

Considering MARINER scenario 1 with respect to the FIPA performatives, and following our own guidelines in Section 3, it is straightforward to implement this scenario to be compliant with the FIPA ACL semantics. For example, quantifiers can send informs to allocators, without knowing (or needing to know) about allocators' belief states or what they do with the content of the message. Therefore quantifier agents don't need Property *X*. Similarly, service information is updated every 30 seconds: therefore allocator agents don't need (or can't use) Property 5, because the FP are not unrelated to time.

In fact, we don't need intentional behaviour, FIPA semantics or even intelligent agents to implement MARINER scenario 1. However, for realistic agent-based load control in Intelligent Networks, we need a higher level of reasoning, decision making and communication normally associated with agents [17]. This requires quantifiers and allocators having (more or less) accurate or consistent models of each other's belief states. We then need to make two assumptions. Firstly, that the system is closed and quantifiers and allocators trust each other (but note that IN often span network operators so this might not always be a valid assumption). Secondly, that the agents' clocks are perfectly synchronised, so that they do not have to worry about time or errors.

Note that our aim is to design message passing as much as possible within ideal operating conditions, and then add stronger consequences to the semantics, or additional processing when the assumptions do not hold or something goes wrong. In these circumstances, then, Property 5 is safe: allocators can infer that quantifiers believe what they say. Furthermore, Property X is safe: in a closed system of trusted and trusting agents, quantifiers can infer that allocators believe what they are told.

In Sections 4.2 and 4.3 we investigate some of the consequences of working with adding stronger logical reasoning to MARINER agents and the implications this has for the FIPA semantics specified for the confirm and disconfirm performatives.

4.2. The confirm performative

The confirm performative is denoted by $\langle s, \text{confirm}(r, \phi) \rangle$. From inspection of Table 1 it has two feasibility preconditions. *Confirm FP1* is that sending agent s believes that proposition ϕ holds, and *confirm FP2* is that agent s believes intended recipient r is uncertain if ϕ holds. Note that the *confirm FP1* and *confirm RE* are identical to those for the inform performative. When an agent formulates a goal $\mathcal{B}_r\phi$, which performative will actually be planned is therefore dependent on the second feasibility precondition for each performative (which are also mutually exclusive).

However, a quantifier s and an allocator r cannot confirm information to each other after the communicative act $\langle s, \text{confirm}(r, \phi) \rangle$ for the following reasons:

- An allocator r with Property 5 has $\mathcal{B}_r\mathcal{B}_s\phi$ and so fails *confirm FP2*.
- An allocator r without Property 5 does not have $\mathcal{B}_r\mathcal{U}_s\phi$ and so fails *confirm FP2*.
- A quantifier s with Property X has $\mathcal{B}_s\mathcal{B}_r\phi$ and so fails *confirm FP2*.
- A quantifier s without Property X does not have $\mathcal{B}_s\mathcal{U}_r\phi$ and so fails *confirm FP2*.

In part this argument follows for the same reasons that one agent cannot reply to an inform from another agent with another inform of the same content. However, in this case neither agent can use a confirm performative. The confirm performative semantics therefore seem to be counter-intuitive: if one person arranges a meeting

with another in a week's time, either one may ring up a week later prior to the meeting, to confirm that the meeting is taking place.

In systems where the environment is dynamic and not completely known by any one agent, the need for periodic updating of information via a confirm speech act may be essential. For example, we have also been investigating the use of the BDI architecture in air traffic control (ATC) with FIPA ACL messages (cf. [18]). The idea is that a 'plane agent' negotiates with the 'ATC agent' to get a slot on the holding (circling) stack before joining the flight path for landing. Our problem is that having achieved a successful outcome to the negotiation, neither ATC agent nor plane agent can send a confirm to ensure that matters are proceeding as negotiated.

One possible solution is again to define a new performative, for example acknowledge (or even assert as discussed above), with a different semantics to inform, and use it similar to the way ack is used in network communication protocols.

4.3. The disconfirm performative

The disconfirm performative is denoted by $\langle s, \text{disconfirm}(r, \phi) \rangle$. From inspection of Table 1 it also has two feasibility preconditions. *Disconfirm FP1* is that sending agent s believes that proposition $\neg\phi$ holds, and *disconfirm FP2* is that agent s believes intended recipient r believes or is uncertain that ϕ holds. The *disconfirm RE* is similar to that of the inform performative, i.e. to share a belief, except that the sender wants the receiver also to believe that $\neg\phi$ holds.

In MARINER we encountered two problems with this semantic formulation: firstly, if we want to selectively disconfirm (deny) a proposition, and secondly if we wanted to withdraw a proposition without asserting that its negation holds.

An example of the first case is illustrated in Figure 2. Suppose there is a quantifier agent Q providing service c to allocator agents $A1$ and $A2$. Agent Q believes service c is available. Now suppose the traffic/load $T1$ from $A1$ using service c plus the traffic/load $T2$ from $A2$ using service c exceeds a certain threshold value. Agent Q seeks to throttle traffic: it does this by withdrawing the service from $A1$ but seeks to keep throughput up by continuing to offer it to $A2$. Thus Q believes the service is still available, except not to $A1$. A resolution to this problem could be found in refining the granularity of the content language, but it is arguable whether developers should be forced to fit the application language to the standard semantics or that the standard semantics is too restrictive.

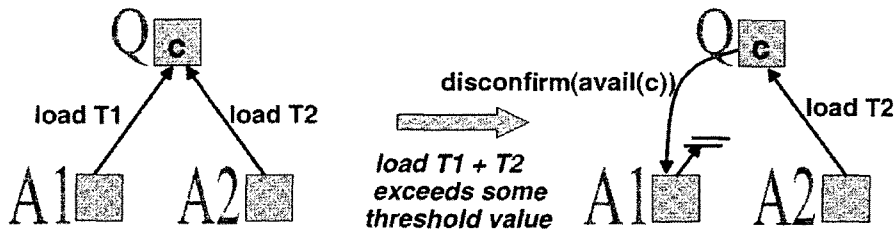


Figure 2. Withdrawing an IN service.

An abstract example of the second case is as follows. Suppose there is an agent s with the beliefs (which includes a meta-belief):

$$\{\mathcal{B}_r(\neg full \rightarrow empty), full, \neg empty\}$$

and an agent r with the beliefs:

$$\{\mathcal{B}_s(\neg full \rightarrow empty), \neg full \rightarrow empty\}$$

Now suppose that s wants to withdraw *full* without (yet) adding $\neg full$, and wants to disconfirm *full* to agent r . If the rational effect of the disconfirm is achieved, s can infer that r will believe something that s does not, namely $\mathcal{B}_r empty$, while disconfirming *empty* will, of course, allow r to infer *full* again. Also, from Property 5 and *disconfirm FPI*, r will believe $\mathcal{B}_s \neg full$ and conclude $\mathcal{B}_s empty$, which again, s does not.

The point of this is that there are applications where the closed world assumption does not apply, wherein one agent may want another one not to believe a proposition, without necessarily believing its negation. Its goal (*RE*) is actually $\neg \mathcal{B}_r \phi$; and $\neg \mathcal{B}_r \phi$ and $\mathcal{B}_r \neg \phi$ (the *RE* of disconfirm) are very different modal formulas.

4.4. Further guidelines

We conclude this section with three further guidelines on the interpretation and application of the FIPA ACL semantics.

- Use the standard protocols. The FIPA ACL specification make available a library of protocols for some commonly encountered interactions in agent-based systems. While not exhaustive, it does provide a useful starting point. Certainly this was our experience in designing the MARINER inter-agent communications.
- New performatives. As a last resort, the FIPA ACL make provision for the definition of new performatives, provided their semantics does not clash with the semantics specified for the standard performatives. Agent developers should freely take advantage of this according to the requirements of their application. Examples include a performative for acknowledgement or assertion similar to that described above, to be used after an *inform*; and a performative for updating, as a complement to *subscribe*, where the sender notifies the receiver of the *sender's* intention to inform the receiver whenever the value of some object changes or after a certain time. Table 2 provides a list of possible new primitive performatives for variants on the disconfirm performative.

The downside of defining new performatives is that inter-operability may be adversely affected, unless agents can learn or discover new semantic definitions. An alternative is to always submit new performatives to FIPA, with the different semantics clearly identified, and the circumstances giving the justification for the new performative.

Table 2. Variations on the disconfirm performative

Feasibility precondition(s)	Rational effect	Performative
$\mathcal{B}_s \neg \phi$	$\mathcal{B}_r \neg \phi$	$\langle s, \text{disconfirm}(r, \phi) \rangle$
$\neg \mathcal{B}_s \phi$	$\neg \mathcal{B}_r \phi$	$\langle s, \text{withdraw}(r, \phi) \rangle$
$\mathcal{B}_s \neg \phi$	$\mathcal{B}_r \neg \phi$	$\langle s, \text{deny}(r, \phi) \rangle$
$\text{Bel}_s \neg \phi$	$\neg \mathcal{B}_r \phi$	$\langle s, \text{retract}(r, \phi) \rangle$

- Be aware of time. Many of the performative definitions are too simple in that they do not identify the time at which the sender is supposed to have satisfied the preconditions. Indeed, Property 5 is only applicable to those *FPS* which “do not refer to time.” The system developer must carefully consider how the logical system is perturbed by temporal considerations.

In fact, the issue of time has presented a considerable problem to the design of a real-time agent-based system as required in MARINER [17], although MARINER has sought to contribute to FIPA'97 specifications by considering extensions for real-time agents.

While real-time is an issue to be addressed by FIPA'99, its impact on the ACL is likely to complicate the semantics substantially. This presents a considerable dilemma for standards creation: a faithful semantics will be more complicated and probably less accessible to developers, while a simpler semantics is likely to be more accessible but less general. However, the FIPA'97 approach is likely to prove useful again, wherein one exposition was made informally in terms of English descriptions, while a more formal exposition was made in logic, and each exposition was designed to serve a separate purpose.

5. Designing agent communications

The intentional stance has proved to be a highly profitable paradigm for exploring concepts like rationality, decision making, planning, attitude and of course agency (cf. [5, 25]). It is particularly effective for exploring the behaviour of individual agents and provides a basis for implementation, although we have found the FIPA ACL semantics is difficult to use with our operational model based on the practical BDI architecture of Kinny et al. [12].

However, we believe that there are fundamental limitations of the intentional stance for trying to articulate the meaning of communication between agents. This is because any interaction between intelligent entities is likely to be heavily context-dependent. An attempt to standardise communication is likely to work for some, but not all, contexts. We have seen numerous examples of this in the preceding: an agent that needs to be not strictly sincere, an agent that wants to confirm a previous arrangement, an agent that is continuously updating information, an agent that wants to withdraw specific information to selected agents, and so on.

This would indicate that the FIPA ACL semantics is not a good place to start when the designing the interactions between agents in a multi-agent system. By this we mean that is perhaps inadvisable to build an agent that implements the formal ACL semantics and then to examine the communicative requirements of the application. Instead, the FIPA ACL performatives (with intuitive specifications) and protocols do offer a sound starting point by focusing attention on the communicative requirements first. Indeed, the protocols should arguably be the focal point for standardisation of the ACL. In the following sub-sections, we offer some guidelines for applying the FIPA ACL in multiagent system design, namely:

- Design for error;
- Separation of concerns;
- Understand the granularity;
- Concentrate on conversations.

We draw some final conclusions in Section 6.

5.1. *Design for error*

In MARINER scenario 1, we made a number of simplifying assumptions about the operation of the system being controlled by the agents. It was then straightforward to identify the types of message to be communicated between the agents. On top of this, we needed to consider then timing arrangements. From this basis, assuming that everything was working correctly, we could generalise our design to cater for certain types of anticipated errors.

It is entirely possible that the communication fabric can malfunction, errors can occur in the functioning of the middleware (FIPA's Agent Management System), and so on. Messages can get delayed, lost, bottlenecked within a busy agent, even misunderstood by one 'stupid' agent in the midst of all the other good ones. There are some interesting implications here for theoretical models of computing because of the inherent autonomy, concurrency, non-determinism, asynchrony, and possible non-termination of distributed, emergent algorithms. Agent developers will need to consider the likely errors that may arise within the communicating agent community.

The issue can be approached by remembering that speech acts are not unrelated events but occur in the context of larger conversational structures [9]. Even silence can be interpreted as an event, and a 'thoughtful' agent will interpret it in useful ways. For example, if an agent has not responded to a message within a certain time, it is up to the sending agent to determine:

- if the email did not arrive, or
- if the email did arrive, and was not read, or
- if the email did arrive, was read, and was not understood, or
- if the email did arrive, was read, understood, and not responded to, or
- if the email did arrive, was read, understood, responded to, but the reply lost, etc.

The reasoning will determine the action taken to recover from the error.

Larger agent communities will inevitably display emergent behaviour that is not going to be easy to predict in advance. Such emergent behaviour may be serendipitous and beneficial, but it may equally be dysfunctional not to say pathological. More generally, problematic or erroneous behaviour is likely to vary from system to system and from application to application. A standard is only able to offer some guidance for when the system is operating as expected. It is important for the system designers to complement the standard with mechanisms for handling certain classes of exceptions for ensuring robustness in their particular application.

5.2. *Separation of concerns*

KQML (Knowledge Query & Manipulation Language) is an ACL developed by Finin and his associates [7], which was also motivated by speech act theory and the definition of a set of performatives. Before FIPA, KQML was emerging as a kind of de facto standard for multi-agent systems, but because there was no definitive standard and no semantics agent developers would each create their own ad hoc version of KQML, none of which could interoperate.

There has been work on a semantics for KQML, although Cohen & Levesque [6] identified a number of serious difficulties, and Labrou and Finin [13] followed the intentional route that has been found wanting as a basis for a normative semantics.

However, there was one aspect of KQML that was particularly appealing: the fact that the language description identified KQML as conceptually layered, into:

- the communication layer, which dealt with the mechanics of communication, and encoded a set of features in a message which described lower level communication facilities, such as the identify of sender and receiver, a unique identifier for the message, a tag for the reply, and so on;
- the message layer, which dealt with the logic of communication, which identified the performative or speech act of a ‘communicative action,’ the protocol being used, and parameters to describe the content, such as the language and ontology being used;
- the content layer, which dealt with the content of communication, and was the actual content of a message, in the agent’s own language, which could be anything from programming language constructs (e.g. Prolog queries), to ASCII strings, to binary-encoded compressed data.

The FIPA ACL specification notes the correspondence between the transport syntax and the logical representation, in general a formula $\langle s, \text{performative}(r, C) \rangle$ translates to:

```
( performative
  :sender    s
  :receiver  r
  :content   C
)
```

The FIPA ACL specification notes: “this also illustrates that some aspects of the operational use of the FIPA ACL fall outside the scope of this formal semantics but are still part of the specification” (*ibid*, p49). We find that there are three problems with this:

- the semantics has taken elements from each conceptual layer (the sender and receiver from the communication layer; the performative from the message layer, and the content from the content layer) and tried to give them a semantics in one go;
- the semantics has omitted to consider other attributes, which are equally semantically significant, such as the language, the ontology, the protocol, temporal aspects, etc., when it comes to specifying the actual meaning of a performative;
- the semantics has included elements which are necessarily concerned with agent internals.

We believe that the separation of concerns afforded by KQML's conceptual layering was one of the reasons that gave KQML its intuitive appeal and practical utility, even if the separation was not entirely ‘clean’. For example, performatives could have other performatives embedded in the content. Therefore, a clean separation would need to stipulate certain restrictions on the content: for example, don't embed performatives unless it can be shown how both the speech act performative and the embedded performative can be composed into a new action. Otherwise there may be intrinsic logical conflicts of the type identified by Cohen and Levesque [6].

However, we believe that the semantics should endeavour to capture the difference between and relation between the object level semantics, that of the content, its language and ontology, and the meta-level communication semantics, as given by the performative, protocol, and the identities of sender and receiver etc. The belief state of either party in a speech act is orthogonal to the meaning of the performative at any of these levels. The KQML semantics described by [14] is accordingly much more effective.

Therefore, when designing agents to interoperate in heterogeneous systems, the designers need to agree on:

- the observable behaviour, i.e. the meta-level communication semantics, which is nicely captured by FIPA ACL's performatives, protocols, message identifiers, conversation identifiers, etc., and are generic across applications and amenable to standardisation;
- the observable content, i.e. the object level content semantics, which is defined by a language and an ontology and is specific to an application, but still amenable to standardisation.

The feasibility preconditions and rational effects are then only an informative part of the standard, in the sense that agent developers may follow them if they want sending agents to be understood and co-operated with, and that receiving agents are entitled to (but may not) make the specified inferences (e.g. from Property 4 and 5). Similarly, while receiving agents may make the specified inferences, they should be aware that the *FPs* and *REs* are not binding.

5.3. *Understand the granularity*

FIPA ACL ostensibly appears to have three types of performative, for example:

inform, which could be a classifier for declarative sentences,
confirm, disconfirm, etc., which are verbs used in declarative sentences, and
accept-proposal, reject-proposal, etc., which are verbs and their complements.

However, closer inspection of the semantics reveals that only a few of the performatives are primitive and that the majority are macro-acts, whose semantics is generally given in terms of inform.

What FIPA ACL semantics then offers is a disciplined method for defining new performatives, thus providing a genuine standard for inter-operability. This is in marked contrast to the 'free-for-all' that developed with KQML, where in the absence of a 'standard' KQML and a formal semantics any number of dialects emerged.

It is important then to understand the granularity at which FIPA97 has chosen to standardise. This is because the openness of the standard does not exclude pathological behaviour on the part of the developers. For example, at one extreme, we could just say every agent communication is an inform and interpret the content according to the application. The agents would be FIPA-compliant even if the messages exchanged are of the form:

```
( inform
  :sender    s
  :receiver  r
  :language  KQML
  :content   (:tell ... )
)
```

At the other extreme, since we can define new performatives, we could keep on extending the set of performatives with every verb in the English language, or, even every verb and verb-complement demanded by an application.

The FIPA semantics helps to impose some discipline on the use of given performatives and the creation of new ones. Some new acts can be defined from the composition of constituent acts, while others are primitive and defined by distinct feasibility preconditions and rational effects. We have seen several instances throughout this paper where this appeared to be necessary. The FIPA ACL semantics gives all developers a consistent reference point, which is a necessary property of a good standard, and a systematic method for working beyond the standard.

Therefore, it is important to understand the granularity defined by FIPA ACL and the distinction between primitive and macro acts. This then has to be related to the granularity demanded by the application, which can be a non-trivial task [17]. In addition, Searle's own speech act typology had representative, commissive, directive, declarative and expressive categories. The FIPA performatives are largely declarative, and there is a potentially rich development in considering the other categories of Searle's typology, especially in relation to permissives and commissives [26].

5.4. Conversation analysis

The FIPA ACL Specification describes the inform performative as “one of the most interesting assertives regarding the core of mental attitudes it encapsulates” (*ibid*, p49). This appears to be the motivation behind the FIPA ACL semantics: to characterise the meaning of a performative in terms of the intentional stance by defining the locutionary conditions and perlocutionary effects on the ‘mental attitudes’ of speaker and hearer, or in this case, the sending and receiving agents. Locutionary conditions refer to the formulation of an utterance, defined by the preconditions. In Searle’s own terms [24] it is the specification of the relevant intentional content that plays a causal role in the production of behaviour. Perlocutionary effects refer to the updating of mental attitudes caused by the act itself.

We have encountered two problems with this:

- the specification of relevant intentional content that plays a causal role in the production of behaviour is only defined for a single speech act. What it fails to specify is the intentional content that plays a causal role in the production of behaviour in the context of a conversation between two or more agents;
- the specification of the perlocutionary effects for an inform are somewhat restrictive, and if strengthened logically, lead to several of the problems described in the previous section.

Speech act theory has been criticised before for being basically unidirectional, i.e. there is a lack of analysis of the wider interaction between the speaker and hearer [4]. The FIPA ACL semantics has addressed this issue to some extent through its Property 5, but some of the earliest work in inter-agent communication took the idea of a speech act as an event but formalised the interaction via protocols [3, 16, 11]. An agent was then expected to conform to a protocol in their interactions with other agents.

This can be seen as a restriction of Conversation Analysis [15], which has been used to inform the design of human-computer dialogues. Ideas from conversation analysis, such as opening and closing sequences, adjacency pairs, even correction and repair when interactions go wrong, might also be applicable to the design of agent-agent dialogues. The meaning of a speech act then becomes less centred on the motivation for performing the act and more centred on the context in which it is being performed—i.e. the conversation and protocol. Indeed, there appears to be an emerging consensus that conversations, dialogues and protocols are a more promising approach to specifying an ACL semantics, which is reflected in the FIPA’99 Call for Proposals (<http://www.fipa.org>).

6. Concluding remarks

With growing investment in open, agent-based system design, development and deployment, there is a pressing concern for the establishment of common standards for the software that comprises the interface between agents. This interface is partly

defined by the Agent Communication Language, and FIPA has endeavoured to standardise both the syntax and the semantics of this language. The potential advantages for developers of open, heterogeneous, and interoperable agent systems are plain to see.

However, it is important to understand what semantics is for. It has to be accessible to those who need to use it, and it is quite clear that the FIPA ACL semantics is not. The benefits of denotational semantics for programming languages, for example, were realised in improved compilers, consistency of compilation across different platforms, and new programming paradigms. The ACL semantics should be promising the same kind of benefits to the agent community. This paper should be perceived at least as an attempt to clarify the intentions and intuitions behind the FIPA ACL semantics, and provide some guidance and explanation for developers interested in engineering FIPA-compliant systems.

We believe this paper has shown that there are a number of limitations associated with the specification of the FIPA ACL semantics. We further believe that these problems stem from trying to standardise the semantics of an ACL based on the intentional stance. It is possible for the intuitions behind the original intentional semantics to be interpreted informally and accommodated if required. Interactions in a multi-agent system can then be designed on the basis of the standard performatives and protocols, and the ACL specification allows for alignment with new protocols and performatives should those prove to be necessary.

Standards, naturally, have their critics. There are those who argue that standards will not ensure consistent design, may be difficult for designers and implementers to use, and that generally stated standards can be wrong for specific applications. In particular, a standard that is difficult to interpret only gives ammunition to such critics, while those who aim to be compliant with the standard will find ways of subverting, circumventing or sidelining it. One suspects that this will be the fate of the FIPA ACL semantics, unless a more tractable solution can be found.

On the other hand, FIPA standards are not mandatory or binding, and any organization can apply the standards to build FIPA-compliant systems. Of course, this is a technically demanding engineering task, and the FIPA specifications can be correspondingly complex. While FIPA may be an open-to-all, not-for-profit organization, it is not a charity either. FIPA has made public a large body of knowledge contained in its seven normative and informative specifications. The understanding, interpretation and use of those specifications is another matter entirely. We hope that 'the agent community' will benefit from this expository analysis of just one part of just one specification, but our best advice to maximise the benefit of FIPA's output in the long term is to join FIPA itself.

Acknowledgments

This work owes much to discussion with numerous colleagues, in particular Jim Cunningham (Department of Computing, Imperial College of Science, Technology and Medicine) and we are very grateful for their various contributions. However, that they have discussed it does not necessarily imply that they endorse it, and the

authors alone are responsible for any errors in this document. We are also deeply indebted to two anonymous referees, whose comments and criticisms on a previous version of this paper were especially constructive. This work has been undertaken in the context of the EU-funded MARINER Project (ACTS 333) and the UK-EPSRC/Nortel Networks funded Project CASBAh (GR/L34440), and support from these funding bodies is gratefully acknowledged.

References

1. R. Brennan and B. Jennings, "MARINER Project," OMG IN/CORBA Workshop, Manchester, 1998.
2. P. Bretier and M. Sadek, "A rational agent as a kernel of a co-operative dialogue system: Implementing a logical theory of interaction," in *Proceedings ECAI'96 Workshop on Agent Theories, Architectures and Languages*, Springer-Verlag: Berlin, 1996, pp. 261–276.
3. B. Burmeister, A. Haddadi and K. Sundermeyer, "Generic configurable cooperation protocols for multi-agent systems," in *MAAMAW'93, Neuchatel*, 1993.
4. M. Chang and C. Woo, "A speech-act-based negotiation protocol: Design, implementation, and test use," *ACM Transactions on Information Systems*, vol. 12(4), pp. 360–382, 1994.
5. P. Cohen and H. Levesque, "Intention is choice with commitment," *Artificial Intelligence*, vol. 42(2-3), pp. 213–261, 1990.
6. P. Cohen and H. Levesque, "Communicative actions for artificial agents," in V. Lesser (ed.), *Proceedings ICMAS95*, AAAI Press: Menlo Park, CA, 1995.
7. T. Finin, Y. Labrou and J. Mayfield, "KQML as an agent communication language," in J. Bradshaw (ed.), *Software Agents*, MIT Press: Cambridge, MA, 1995.
8. FIPA, "FIPA'97 Specification Part 2: Agent Communication Language," Foundation for Intelligent Physical Agents, <http://drogo.cse.it.stet.it/fipa/>, 1997.
9. C. Flores and J. Ludlow, "Doing and speaking in the office," in G. Fisk and R. Sprague (eds.), *DSS: Issues and Challenges*, Pergamon Press: Elmsford, NY, pp. 95–118.
10. J. Galliers, "Autonomous belief revision and communication," in P. Gardenfors (ed.), *Belief Revision*, Cambridge University Press: Cambridge, UK, 1992.
11. R. Gustavsson, "Multi agent systems as open societies," in M. Singh, A. Rao and M. Wooldridge, (eds.), *Intelligent Agents IV, LNAI*, vol. 1365, Springer-Verlag: Berlin, 1998.
12. D. Kinny, M. Georgeff, J. Bailey, D. Kemp and K. Ramamohanarao, "Active databases and agent systems: A comparison," in *Proceedings Second International Rules in Database Systems Workshop*, 1995.
13. Y. Labrou and T. Finin, "Semantics for an agent communication language," in M. Singh, A. Rao and M. Wooldridge (eds.), *Intelligent Agents IV, LNAI*, vol. 1365, Springer-Verlag: Berlin, 1998.
14. Y. Labrou and T. Finin, "Semantics and conversations for an agent communication language," in M. Huhns and M. Singh, (eds.), *Readings in Agents*, Morgan Kaufmann: Los Altos, CA, 1998, pp. 235–242.
15. P. Luff, N. Gilbert and D. Frohlich (eds), *Computers and Conversation*, Academic Press: San Diego, 1990.
16. J. Pitt, M. Anderton, and J. Cunningham, "Normalized interactions between autonomous agents: A case study in inter-organizational project management," *Computer-Supported Cooperative Work*, vol. 5, pp. 201–222, 1996.
17. J. Pitt and K. Prouskas, "Initial specification of multi-agent system for realisation of load control and overload protection strategy," MARINER Project (EU ACTS AC333) Deliverable D3, <http://www.teltec.dcu.ie/mariner>, 1998.
18. A. Rao and M. Georgeff, "BDI agents: from theory to practice," in V. Lesser (ed.), *Proceedings ICMAS95*, AAAI Press: Menlo Park, CA, 1995.
19. J. Rosenschein and G. Zlotkin, *Rules of Encounter*, MIT Press: Cambridge, MA, 1998.

20. D. Sadek, "Dialogue acts are rational plans," in *Proceedings ESCA/ETRW Workshop on the Structure of Multimodal Dialogue*, Maratea, Italy, 1991, pp. 1–29.
21. D. Sadek, "A study in the logic of intention," in *Proceedings 3rd Conference on Principles of Knowledge Representation and Reasoning (KR'92)*, Cambridge, MA, 1992, pp. 462–473.
22. D. Sadek, "Communication theory = rationality principles + communicative act models," in *Proceedings AAAI94 Workshop on Planning for Interagent Communication*, Seattle, 1994, pp. 78–82.
23. T. Sandholm, "Agents in electronic commerce: Component technologies for automated negotiation and coalition formation," in M. Klusch (ed.), *Co-operative Information Agents*, LNAI, vol. 1435, Springer-Verlag: Berlin, 1998, pp. 113–134.
24. J. Searle, "Conversation," in H. Parret and J. Vershueren (eds.), (on) *Searle on Conversation*, Benjamins: Amsterdam, 1992, pp. 7–29.
25. M. Singh, "A logic of intentions and beliefs," *Journal of Philosophical Logic*, vol. 22, pp. 513–544, 1993.
26. M. Singh, "Agent communication languages: Rethinking the principles," *IEEE Computer*, pp. 40–47, 1993.
27. I. Smith, P. Cohen, J. Bradshaw, M. Greaves and H. Holmback, "Designing conversation policies using joint intention theory," in Y. Demazeau (ed.), *Proceedings ICSMAS98*, IEEE Press: New York, 1998.
28. M. Wooldridge, "Verifiable semantics for agent communication languages," in Y. Demazeau (ed.), *Proceedings ICMAS98*, IEEE Press: New York, 1998.