



Intention et principes de cooperation pour le traitement des requetes et des questions fermees au travers des assertifs

Andreas Herzig, Dominique Longin

► To cite this version:

Andreas Herzig, Dominique Longin. Intention et principes de cooperation pour le traitement des requetes et des questions fermees au travers des assertifs. 13ème Congrès Francophone AFRIF-AFIA de Reconnaissance des Formes et Intelligence Artificielle (RFIA 2002), Association Française d'Intelligence Artificielle; Association Française pour la Reconnaissance et l'Interprétation des Formes, Jan 2002, Angers, France. hal-03534110

HAL Id: hal-03534110

<https://hal.science/hal-03534110>

Submitted on 19 Jan 2022

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Intention et principes de coopération pour le traitement des requêtes et des questions fermées au travers des assertifs

Andreas Herzig

Dominique Longin

Institut de Recherche en Informatique de Toulouse (IRIT)

Équipe Logic, Interaction, Language and Computation (LILAC)

118 Route de Narbonne, F-31062 Toulouse cedex 4

{herzig,longin}@irit.fr

<http://www.irit.fr/ACTIVITES/LILaC/>

Résumé

La formalisation de l'évolution des connaissances d'un agent est un problème crucial dès lors que cet agent peut percevoir des informations sur le monde dans lequel il évolue. Ces mécanismes sont d'autant plus complexe à décrire que le langage est riche. Cette contribution propose un cadre d'actes de langage simplifiés qui permet tout de même de traiter les actes nécessaires à un dialogue orienté-tâche.

Pour cela, une caractérisation de la notion d'intention ainsi que de certaines lois coopératives est nécessaire, et permet d'inférer des requêtes et des questions fermées à partir d'assertions.

Nous présentons les justificatifs philosophiques de ce mécanisme d'inférence, qui est formalisé dans une logique modale dynamique épistémique. Nous montrons enfin comment l'appliquer à une architecture BDI d'agent rationnel. Tout ce travail est dirigé dans le but d'obtenir une logique minimale et donc simple à mettre en œuvre par des méthodes de preuve.

Mots Clef

Intention, actes de langage, logique modale, agent rationnel intelligent.

Abstract

Formalising the dynamics of an agent's knowledge is important as soon as that agent gets information about his environment. The more the language is rich, the more these mechanisms get complex. In this paper we propose a simple framework for communication allowing nevertheless to handle task-oriented dialogues.

To that end it is necessary to characterise the notion of intention as well as some laws of cooperation. These will allow to infer requests and yes-no-questions from assertions. We present the philosophical justifications of such an inference mechanism, which is formalised in an epistemic dynamic logic. We finally show how it can be applied to a BDI

architecture of a rational agent.

It is the aim of this work to obtain a minimal logic that can be mechanized in a simple way.

Keywords

Intention, speech acts, modal logic, intelligent rational agent.

1 Introduction

À partir des solutions au problème du décor (*frame problem*) proposées dans les années 90 [11, 22], des tentatives ont été faites ces dernières années pour le résoudre dans le cadre de l'évolution de connaissances [16, 22]. Nous employons ici le mot « connaissance » par opposition à « croyance », cette dernière étant vue comme une connaissance subjective (autrement dit, potentiellement non conforme au monde réel).

Modifier ces solutions dans le but de prendre en compte non plus l'évolution des connaissances mais celle des croyances n'est pas chose aisée. Nous avons montré dans [6] les difficultés rencontrées par l'approche élaborée dans [21].

Un des buts communs à toutes ces approches est alors de définir la logique d'action et de croyance la plus simple, tout en étant capable de rendre compte de l'évolution dans le temps des croyances d'agents rationnels.

Dans ce but, nous nous focalisons dans ce qui suit sur les actes communicatifs et partons du constat que la difficulté à décrire l'évolution des croyances d'un agent croît avec le nombre de types d'actes communicatifs que cet agent est capable de prendre en compte. Néanmoins, pour que cet agent puisse interagir convenablement avec son environnement, il doit au minimum être capable de faire des assertions, poser des questions, et donner des directives ; il doit bien sûr également être capable de les comprendre et d'y répondre.

Notre propos est donc de montrer dans ce qui suit que l'on peut doter un agent de ces capacités d'interaction par les

seuls assertifs, associés à une notion d'intention adéquate et des principes de coopération, dont nous donnons une caractérisation formelle. Nous illustrons l'adéquation de ce formalisme au traitement d'énoncés (d'un type autre que celui des assertifs) par deux exemples.

Nous montrons tout d'abord que cette démarche est justifiée par des travaux en philosophie des actes de langage et des états mentaux (Sect. 2). Puis nous présentons notre langage logique (Sect. 3) afin de pouvoir présenter formellement notre notion d'intention (Sect. 4) et les actes d'assertion (Sect. 5). Nous pouvons alors formuler les principes de coopération nécessaires à l'interprétation d'actes communicatifs et à la génération de réponses appropriées (Sect. 6). Enfin, nous appliquons ces résultats à une architecture BDI d'agent rationnel (Sect. 7).

2 Actes indirects

Dans ce qui suit, nous montrons comment les travaux issus de la théorie des actes de langage peuvent être utilisés pour faire le lien entre les assertifs d'une part, et des principes de coopération d'autre part.

L'idée selon laquelle « en disant quelque chose, on *peut* vouloir dire quelque chose d'autre » trouve sa source dans les fondements philosophiques de la théorie des actes de langage et de celle des états mentaux [1, 18, 17], sous la dénomination de : *actes de langage indirects*.

Un acte de langage indirect est l'acte de langage produit (indirectement) par le biais de l'accomplissement d'un autre acte de langage (« direct », par opposition au premier). Ainsi, l'énonciation de « peux-tu me passer le sel ? » réalise l'acte direct correspondant à une question fermée (*i.e.* une question dont la réponse attendue est « oui » ou « non ») sur la capacité physique ou morale de l'interlocuteur à passer le sel au locuteur. Dans certains contextes, cette énonciation *peut*¹ réaliser l'acte indirect correspondant à une requête du locuteur à destination de son interlocuteur, où le premier demande au second de lui passer le sel.

On peut légitimement douter de la réalité de cette dichotomie entre actes direct et indirect. Nous avons nous-même montré comment des résultats en psycholinguistique remettaient en cause le fait que, pour comprendre le sens indirect d'un énoncé, on comprenait d'abord son sens direct : il s'agit plutôt d'associer différents effets à un même acte selon son contexte d'accomplissement [4].

Néanmoins, si on laisse de côté les préoccupations relatives au « réalisme cognitif » du modèle, on peut avantageusement tirer parti du fait qu'une assertion interprétée systématiquement de façon directe, nous permettra *via* des principes de coopération, d'extraire également son interprétation indirecte.

Il est important de souligner ici que notre préoccupation

1. Nous insistons ici sur le fait qu'il s'agit d'une possibilité, et non d'une certitude : si un énoncé est une forme d'indirection, il peut ne pas être utilisé comme tel. Ceci est relatif à la propriété selon laquelle « tout acte de langage indirect est *annulable* » [23].

première n'est pas de chercher à traiter des actes indirects : simplement, nous souhaitons exploiter les mécanismes qui les régissent, dans le but de coder des types d'actes illocutoires (autres que les assertifs) *via* des actes de type assertif. Partant de ces considérations générales, et étendant en cela les travaux de Searle [17], Virbel a montré qu'un acte de langage « peut être réalisé par des assertions ou des questions portant sur trois types principaux d'arguments » [24, 3] :

- (a) des conditions de succès de l'acte direct correspondant (*i.e.* celui que l'on souhaite indirectement accomplir) ;
- (b) des raisons de (ne pas) faire l'action indiquée ;
- (c) des éléments relatifs à la planification de l'action indiquée.

Par exemple, soit α l'acte réalisé par l'énoncé « Peux-tu me passer le sel ? ». Dans le cas où cet énoncé est à interpréter de façon indirecte, l'accomplissement de α produit un acte (indirect) α' correspondant à l'énoncé « Passe-moi le sel ». La condition préparatoire de α' étant que le locuteur pense que l'auditeur peut lui passer le sel, α est bien une question fermée sur la condition préparatoire de α' (cas (a)).

De même, la condition de sincérité d'une requête telle que « Exécute l'action β » est que le locuteur souhaite que l'auditeur accomplisse l'action β . Ainsi, une assertion de cette condition (« Je veux que tu exécutes β », ou encore « Je voudrais que tu fasses β ») constitue une forme d'indirection pour signifier « Exécute l'action β ».

Enfin, une raison pour demander à quelqu'un si p est vrai ou non, est de vouloir savoir si p est vrai ou non. Ainsi, une requête telle que « Dis-moi si p est vrai ou non » peut être accomplie en faisant une assertion sur cette raison, *c.-à-d.* *via* un énoncé du type « Je veux savoir si p est vrai ou non ». Nous venons donc d'exhiber une façon d'accomplir une requête et une façon de poser une question par le biais d'une assertion. C'est ce que nous exploitons formellement dans la suite de cet article.

Il est utile de préciser que dans [2], on trouve des schémas d'indirection (du type : « si tel acte est accompli, cela peut signifier que tel acte a été accompli »). Mais il n'est pas précisé comment ces schémas sont construits, ni quand ils sont appliqués ou non. L'approche sur laquelle nous nous fondons, quant à elle, a le mérite de pouvoir construire des règles d'indirection de façon systématique et elle est en cela explicative de l'indirection.

3 Le langage logique

Notre langage est celui de la logique multimodale du premier ordre sans égalité ni symbole de fonction. La sémantique associée est définie en termes de mondes possibles, et de relations d'accessibilité ou de fonctions de voisinage (pour l'intention). Les formules atomiques sont notées p , q , ... ou $P(t_1, \dots, t_n)$, et ATM est l'ensemble de toutes les formules atomiques. Les formules sont notées A , B , ... et $FORM$ est l'ensemble de toutes les formules.

Soit $\mathcal{AGT}$ l'ensemble des agents. Pour $i \in \mathcal{AGT}$, $Bel_i A$ est lu « l'agent i croit que A » et $Int_i A$ est lu « l'agent i a l'intention que A ». Nous soulignons que les opérateurs d'intention Int_i sont définis dans une logique non normale. En particulier, la règle de nécessité (RN) et l'axiome (K) ne sont pas valables dans cette logique (cf. [10] pour plus de détails). Enfin, $Bel_{i,j} A$ est lu « les agents i et j croient mutuellement que A ».

Par exemple, $Bel_i Dest(\text{Paris})$ exprime que l'agent i croit que la destination est Paris. $BelIf_i A$ est une abréviation pour $Bel_i A \vee Bel_i \neg A$ qui est lue « l'agent i sait si A est vrai ou non ». Nous utilisons ici *savoir* par pure commodité linguistique. Dans ce qui suit, toute expression du type *savoir si fait* en réalité référence à une disjonction de croyances. En accord avec la théorie des actes de langage, un acte est représenté par une force illocutoire appliquée sur un contenu propositionnel [18]. Comme nous ne manipulons que des actes d'assertion, nous n'explicitons pas la force illocutoire. Ainsi, un acte de langage de type assertif est un triplet de la forme :

$$\langle i, j, A \rangle$$

où i est l'agent auteur de l'acte, j l'agent destinataire, et A une formule logique représentant le contenu propositionnel de l'acte. Par exemple, $\langle u, s, \text{Bleu}(\text{ciel}) \rangle$ représente l'acte d'assertion réalisé par l'énoncé de l'agent u à destination de l'agent s : « Le ciel est bleu ».

À chaque acte α est associé un opérateur modal que l'on note $Done_\alpha$. $Done_\alpha A$ est lu « α vient juste d'être accompli, avant quoi A était vrai »². $Done_\alpha \top$ est lu « α vient juste d'être accompli ».

On associe également à chaque acte un opérateur modal $Feasible_\alpha$, et $Feasible_\alpha A$ est lu « α est faisable, juste après quoi A sera vrai »³.

Enfin, on note $After_\alpha A$ la formule $\neg Feasible_\alpha \neg A$, et doit donc être lu « après toute exécution de α , A est vrai ». D'une manière similaire on note $Before_\alpha A$ la formule $\neg Done_\alpha \neg A$. $Before_\alpha A$ est lu « avant toute exécution de α , A est vrai ».

4 La notion d'intention

Nous ne pouvons définir des principes de coopérations relatifs à la croyance et l'intention des agents (cf. Sect. 6) sans définir ces notions et les propriétés que nous leurs attribuons.

En particulier, nous adoptons KD45 comme logique modale pour la croyance, ce qui est relativement standard. Les axiomes que nous allons définir sont bien-entendu dépendants de cette logique, et nous signalerons explicitement les cas où nous nous appuyons sur ses propriétés.

En ce qui concerne la notion d'intention que nous utilisons, celle-ci est très proche de celle de Cohen & Levesque [5] et de celle de Sadek [12, 13].

Néanmoins, alors que dans leurs travaux elle est le résultat d'une construction complexe d'opérateurs de croyance et

de but, la nôtre est une primitive du langage et un certain nombre d'axiomes en restituent assez fidèlement les propriétés. Comme cela a été signalé précédemment (Sect. 3), l'intention est définie ici dans une logique non normale restreinte à la seule règle d'équivalence (RE) :

$$\frac{A \leftrightarrow B}{Int_i A \leftrightarrow Int_i B} \quad (RE_{Int})$$

plus des axiomes rendant compte des propriétés de l'intention par rapport à la croyance ou à l'action. Du fait de $(Rel_{IntBel1})$ (v.i.) et de la logique KD45 pour les opérateurs de croyance, les opérateurs d'intention vérifient également l'axiome (D) de la logique modale.

Une propriété très importante de l'intention, et commune à toutes les approches issues de ces travaux (y compris la nôtre), est que l'on ne peut avoir l'intention qu'une propriété soit vraie s'il est possible qu'elle le soit déjà. En d'autres termes :

$$Int_i A \rightarrow Bel_i \neg A \quad (Rel_{IntBel1})$$

est un axiome de la logique.

THÉORÈME 1 $Bel_i A \rightarrow \neg Int_i A$.

PREUVE 1 Par $(Rel_{IntBel1})$ et l'axiome (D) de la logique modale pour la croyance on a que $Int_i A \rightarrow \neg Bel_i A$ à partir de quoi nous obtenons le théorème 1 par contraposition.

La propriété $(Rel_{IntBel1})$ a fait l'objet de beaucoup de critiques dans la littérature car elle rend compte d'une notion très forte d'intention : par exemple, si un agent ne sait pas si la lumière est éteinte dans une pièce, il ne pourra adopter l'intention d'aller l'éteindre. Formellement, si p représente « la lumière est éteinte », cela s'illustre par le fait que $\neg BelIf_i p \wedge Int_i p$ est une contradiction (par définition de $BelIf_i$ et par $(Rel_{IntBel1})$).

Le comportement qui paraît généralement le plus raisonnable, est de considérer que l'agent devrait aller dans la pièce concernée, et que si la lumière est déjà éteinte, il abandonne son intention d'éteindre la lumière puisque cette intention est satisfaite. Ainsi, on serait tenté d'affaiblir $(Rel_{IntBel1})$ en

$$Int_i A \rightarrow \neg Bel_i A$$

ce qui autoriserait un agent à avoir l'intention que p tout en ne sachant pas si p est vrai ou non (c'est ce qui est fait dans [19] par exemple).

Notre point de vue est qu'en fait il y a ici deux intentions et non pas une seule :

- 1° dans un premier temps il y a l'intention de savoir si la lumière est allumée ou éteinte ;
- 2° puis, dans un deuxième temps (seulement !), il y a l'intention de l'éteindre (dans le cas où elle est allumée).

2. $Done_\alpha A$ est similaire à $\langle \alpha^{-1} \rangle A$ en logique dynamique [8].

3. $Feasible_\alpha A$ est similaire à $\langle \alpha \rangle A$ en logique dynamique.

En accord avec l'argument donné par Sadek [12, p. 120], il faut conserver à l'esprit que l'intention est une attitude mentale par laquelle on s'engage (de façon persistante) à *réaliser* un but, et qu'en général⁴ il semble donc irrationnel de chercher à réaliser ce but si on pense qu'il l'est peut-être déjà. Ainsi, en général, avant de chercher à éteindre la lumière, il semble rationnel de chercher tout d'abord à savoir si elle est allumée.

En fait, l'idée sous-jacente à la notion d'intention formalisée par (Rel_{IntBel}1), est que chaque fois où il y a un doute sur la nécessité d'engendrer une intention (comme dans l'exemple précédent), l'agent doit d'abord adopter l'intention de connaître l'état du monde. Et ce n'est que si l'état du monde ne satisfait pas la propriété en question, que l'on va adopter l'intention de la réaliser.

À notre connaissance, [12] est le seul travail évoquant ce problème (Sadek y propose de le résoudre formellement au travers d'une attitude mentale de *besoin*, également appelée *intention potentielle*). C'est pourtant un problème majeur dès que l'on souhaite qu'un agent maintienne en permanence un équilibre rationnel entre ses différentes attitudes mentales (ici, croyance et intention).

Notre point de vue est donc que l'axiome (Rel_{IntBel}1) n'est pas trop fort mais que, simplement, il ne décrit pas intégralement l'équilibre rationnel entre croyance et intention. Contrairement à Sadek, nous ne pensons pas utile de définir une nouvelle attitude mentale. Tout d'abord, parce que cette attitude mentale est définie dans [15] comme une abréviation de $Bel_i A \vee Int_i Bel_i A$, donc en termes de croyance et d'intention. D'autre part, cette définition introduit certains effets de bord qui ne sont pas très intuitifs : par exemple $Bel_i A \rightarrow (Bel_i A \vee Int_i Bel_i A)$ est valide et signifie que si un agent croit A , alors il a besoin que A soit vrai. Enfin, intégrer la notion de besoin dans la notion même d'intention et dans ses propriétés par rapport à la croyance évitera de complexifier la logique ce qui est, rappelons-le, un de nos objectifs.

Formellement, cette *intention de connaître l'état du monde* correspond à l'intention de croire qu'une certaine propriété est vraie. Ici, avoir « l'intention de croire » fait référence à un mécanisme introspectif. Une instance du théorème 1 est :

$$Bel_i Bel_i A \rightarrow \neg Int_i Bel_i A$$

Autrement dit, un agent ne peut chercher à croire A s'il a conscience de le croire déjà.

Ainsi, nous proposons de compléter (Rel_{IntBel}1) par un second principe que nous formalisons par l'axiome suivant :

$$(Int_i Bel_i A \wedge Bel_i \neg A) \rightarrow Int_i A \quad (\text{Rel}_{\text{IntBel}2})$$

Cet axiome signifie que si un agent i a l'intention de croire A et qu'il croit actuellement $\neg A$, il adoptera l'intention que

A soit vrai. Cet axiome traduit une volonté, de la part de l'agent i de changer d'avis par rapport à A : il a l'intention de croire que A , or il pense que A est actuellement faux, i doit donc agir afin de faire évoluer le monde (ce qui justifie, dans ce cas, l'intention que A soit vrai). En revanche, si i ne savait pas si A était vrai, son intention sous-jacente pourrait se limiter à croire que A est vrai (sans entreprendre des actions pour changer l'état de A).

On peut ainsi démontrer le théorème suivant :

$$\text{THÉORÈME 2 } (Int_i Bel_i A \wedge \neg Int_i A) \rightarrow \neg Bel_i A$$

Ce théorème signifie que si i a l'intention de croire que A sans avoir l'intention que A , alors il ignore si A .

PREUVE 2 On peut prouver le théorème 2 par (Rel_{IntBel}2) d'où découle $(Int_i Bel_i A \wedge \neg Int_i A) \rightarrow \neg Bel_i \neg A$, et par (Rel_{IntBel}1) d'où découle $Int_i Bel_i A \rightarrow Bel_i \neg Bel_i A$, cette dernière formule étant équivalente à $\neg Bel_i A$ dans KD45.

Enfin, une troisième propriété importante constitutive de l'équilibre rationnel entre intention et croyance (on la trouve également dans [12]) est formalisée par l'axiome suivant :

$$Int_i A \rightarrow Int_i Bel_i A \quad (\text{Rel}_{\text{IntBel}3})$$

qui signifie qu'un agent ne peut avoir l'intention de réaliser une propriété sans avoir l'intention de croire que cette propriété est vraie.

Sa réciproque n'est pas valide (logiquement et intuitivement) : on peut avoir l'intention de croire qu'une certaine propriété est vraie sans pour autant avoir l'intention qu'elle le devienne.

Par exemple, un agent ne sachant pas si la lumière d'une pièce est allumée adopte d'abord l'intention de croire que la lumière est allumée avant, éventuellement, d'adopter l'intention qu'elle soit éteinte. Mais dans ce cas, le fait qu'il ait l'intention de croire la lumière allumée ne constitue pas une raison pour qu'il ait l'intention que la lumière soit allumée (sinon, dans le cas où la lumière est éteinte, il l'allumerait pour pouvoir ensuite... l'éteindre !).

Il est utile de remarquer que (Rel_{IntBel}1) et (Rel_{IntBel}3) font de (Rel_{IntBel}2) une équivalence.

Enfin, il y a deux propriétés essentielles (que l'on trouve par ailleurs dans [12]) en relation avec la capacité d'introspection d'un agent :

$$Bel_i Int_i A \leftrightarrow Int_i A \quad (\text{Rel}_{\text{IntBel}4})$$

$$Bel_i \neg Int_i A \leftrightarrow \neg Int_i A \quad (\text{Rel}_{\text{IntBel}5})$$

Ces deux axiomes signifient que les intentions (respectivement, les non-intentions) d'un agent sont correctes et complètes par rapport aux intentions (respectivement, aux non-intentions) qu'il croit avoir.

4. On peut exhiber des cas où, par principe de précaution ou pour des contraintes temporelles, on exécute une action dont le but est peut-être déjà atteint.

5 Effets d'un acte d'assertion

Selon la force illocutoire de l'acte d'assertion considéré, les effets peuvent être plus ou moins importants. Dans [14], Sadek associe aux actes communicatifs des effets de différentes natures :

- l'*effet rationnel* correspondant à l'effet attendu par le locuteur sur le destinataire de l'acte ;
- l'*effet intentionnel* (selon la vision gricéenne de la communication) décrivant le fait que le locuteur a l'intention de produire un « certain effet » (en fait, celui attendu) sur son auditeur ;
- l'*effet indirect*⁵ correspondant à la persistance au travers de l'accomplissement de l'acte, de ses préconditions de faisabilité.

L'*effet rationnel* n'est pas directement produit chez l'auditeur, mais sera considéré comme tel à partir du moment où il est dérivable de ses croyances à l'issue de l'incorporation des autres effets. Par exemple, si $\alpha = \langle i, j, p \rangle$ est l'acte assertif venant d'être accompli, son effet rationnel est $Bel_j p$. Il est utile de préciser que dans [14], Sadek formalise les effets de façon plus fine (mais plus complexe) que ce que nous évoquons. Nous pensons disposer d'un bon compromis entre adéquation de représentation et complexité de celle-ci en simplifiant son modèle (comme l'illustrent nos exemples).

L'*effet intentionnel* est construit autour de l'effet rationnel comme l'intention qu'a le locuteur que son auditeur croie qu'il a l'intention de produire un certain effet. Ainsi, l'effet intentionnel produit par l'accomplissement de $\alpha = \langle i, j, p \rangle$ est : $Int_i Bel_j Int_i Bel_j p$.

L'*effet indirect* est lié à la préservation de la précondition de capacité et de celle de pertinence au contexte. Par exemple, pour $\alpha = \langle i, j, p \rangle$, la précondition de capacité est $Bel_i p$; celle de pertinence au contexte est $\neg Bel_i Bel_j p$. Sadek décrit des actes communicatifs, qui ne sont pas forcément des actes linguistiques. Mais si on se restreint à ces derniers, la capacité peut être assimilée à la sincérité.

Nous supposons que les lois d'actions sont du domaine de la croyance commune : tout observateur de l'acte croit que les effets indirect et intentionnel se sont produits.

Ainsi, si k est un observateur de l'accomplissement de α , nous avons que : $Bel_k Int_i Bel_j Int_i Bel_j p \wedge Bel_k Bel_i p \wedge Bel_k \neg Bel_i Bel_j p$ est vrai. Si l'agent j (qui, en tant que destinataire de l'acte est un observateur particulier de ce dernier) en vient à croire p (produit *via* les effets précédents, ses – autres – croyances et ses règles de rationalité et de coopération), alors c'est que l'effet rationnel a également été produit, et que l'acte a atteint son but.

6 Principes de coopération

Notre notion d'intention étant ainsi définie, nous sommes en mesure de présenter des principes de coopération qui, à partir d'assertifs particuliers, nous permettront de dériver

5. « Indirect » est à prendre ici dans le sens de « effet de bord », et non comme l'effet d'un quelconque acte indirect.

deux types d'actes communicatifs particulièrement utiles dans le cadre d'un dialogue homme-machine. Nous ferons ainsi l'économie de ces deux types d'actes en tant que primitives de notre langage, et l'axiomatique relative à l'évolution des croyances de l'agent en sera simplifiée.

Comme nous l'avons signalé précédemment (Sect. 2), les deux formes d'assertion qui nous intéressent sont les suivantes :

- $\langle i, j, Int_i Done_\beta \top \rangle$ (où j doit être l'auteur de β) est l'acte réalisé par l'énoncé : « i veut que j fasse β » ;
- $\langle i, j, Int_i Bel_j A \rangle$ est l'acte réalisé par l'énoncé : « i souhaite savoir si A (est vrai ou non) » ;

Le premier réalise une requête qui correspond à un acte du type $\langle Request_{i,j} A \rangle$, tandis que le second réalise une question fermée correspondant à un acte du type $\langle QueryYN_{i,j} A \rangle$ (et dont la réponse attendue est « oui » ou « non »).

Tous les actes de ces types seront traités comme des assertions. Pour cela, nous devons définir des principes (de coopération) destinés à ce qu'un agent puisse exploiter l'indirection sous-jacente à ces types d'acte. Ce sont ces principes que nous exposons dans la présente section.

Globalement, on peut considérer qu'être coopératif, c'est contribuer à satisfaire les buts d'autrui. C'est une définition assez connue et répandue aujourd'hui, bien que superficielle.

Idéalement, cette contribution devrait tenir compte des règles sociales et des capacités cognitives de l'individu que l'on est supposé aider. Par exemple, écouter quelqu'un quand il parle et ne pas l'interrompre est une règle de coopération sociale (on l'aide ainsi à *mieux* satisfaire un des ses buts – celui de parler) ; ne pas fournir plus de réponse à sa question que ce qu'il pourra retenir est une règle de coopération cognitive (ceci est à mettre en relation avec la maxime de quantité de Grice [7]).

Être coopératif, c'est également tenter de comprendre et de satisfaire les buts ultimes d'un agent (cf. [12] pour cet aspect). Comme nous l'avons montré dans [3], la gestion de la communication non littérale peut être vue comme une forme de coopération ; l'adoption de croyances ou d'intentions de son interlocuteur également ; ainsi que la génération d'intention destinées à (indirectement) permettre de satisfaire les intentions de ce dernier.

Toujours dans le but de décrire une logique minimale d'agent rationnel, nous ne prenons pas en compte les capacités cognitives ou les règles sociales, les axiomes établis n'étant *a priori* pas contradictoires avec des propriétés plus fines. Ainsi, la coopération est décrite au travers de deux principes :

- l'adoption de croyance (par laquelle un agent adopte les croyances d'un autre agent) ;
- la génération d'intention (par laquelle un agent génère des intentions, notamment afin de répondre aux questions qui lui sont posées et de corriger des croyances erronées de son interlocuteur).

Ces principes sont complétés par des principes de préservation des croyances que nous avons étudiés précédemment [9, 6] et sur lesquels nous ne revenons pas.

6.1 L'adoption de croyance

Adopter une croyance, c'est se mettre à croire ce qu'on pense qu'un autre agent croit. L'adoption doit être contrainte par une condition pour éviter d'adopter toutes les croyances des autres agents.

Ici, cette condition est liée à la notion de compétence : un agent adoptera la croyance d'un autre agent s'il pense que cet agent est compétent à propos de cette croyance. Cette notion est également utilisée chez Cohen & Levesque [5] et chez Sadek [12].

Formellement, pour décrire la compétence d'un agent à propos d'une formule, nous utilisons une *relation de dépendance* telle qu'elle est décrite dans [9] : un agent i est compétent à propos de p si et seulement si il existe une relation de compétence entre i et p . Nous notons alors : $i \rightsquigarrow p$.

$$\begin{array}{l} Bel_i A \rightarrow A \\ \text{si } i \rightsquigarrow A \text{ et } A \text{ est objective} \end{array} \quad (\text{Adopt}_{\text{Bel}1})$$

où une formule est objective si elle ne contient aucun opérateur modal Bel ou Int . Notons que $\rightsquigarrow$ est une notion métalinguistique.

REMARQUE 1 La logique modale pour la croyance que nous utilisons est KD45. De plus, nous supposons que la compétence d'un agent est du domaine de la connaissance commune. Ainsi, l'axiome ($\text{Adopt}_{\text{Bel}1}$) ci-dessus entraîne logiquement que si $i \rightsquigarrow A$ alors $Bel_k Bel_i A \rightarrow Bel_k A$ est vrai pour tout agent k .

Un autre principe d'adoption de croyance est que si un agent dit quelque chose qui n'entre pas en contradiction avec les croyances d'un autre agent, alors ce dernier adopte la croyance du premier.

Autrement dit : si l'agent k ne croit pas que A est faux, alors après que l'agent i a dit que A est vrai, l'agent k croit que A est vrai. Soit, formellement :

$$\neg Bel_k \neg A \rightarrow \text{After}_{\langle i, j, A \rangle} Bel_k A \quad (\text{Adopt}_{\text{Bel}2})$$

Nous soulignons que dans l'axiome précédent, i peut être incompétent sur A . On trouve un principe similaire dans [20].

6.2 La génération d'intention

La génération d'intention complète nos principes de coopération. Elle concerne :

- la génération d'*intention intermédiaire* propre à « fixer l'état du monde » ;
- la génération d'*intention adoptée*, qui consiste à générer une intention identique à celle d'un autre agent (mécanisme appelé couramment *adoption d'intention*).

Nous formalisons ces mécanismes ci-dessous.

Supposons que $Bel_i Int_j A$. Intuitivement, la difficulté de cette étape consiste à tenir compte de l'axiome ($\text{Rel}_{\text{IntBel}1}$) précédent : l'intention de réaliser une certaine propriété ne doit être générée que si l'agent considéré croit que cette propriété est actuellement fausse.

Ainsi, dans le cas où l'agent i ne croit pas que A est actuellement vrai, il n'a pas forcément l'intention que A soit vrai⁶, mais seulement l'intention de croire que A est vrai (cf. la discussion Sect. 4 à propos de ($\text{Rel}_{\text{IntBel}1}$)). Nous obtenons ainsi notre axiome central :

$$\begin{array}{l} (Bel_i Int_j A \wedge \neg Bel_i A \wedge \\ \neg Int_i Bel_i \neg A) \rightarrow Int_i Bel_i A \end{array} \quad (\text{Gen}_{\text{Int}1})$$

REMARQUE 2 Notons ici que nous négligeons une propriété de la notion d'intention de [5], à savoir qu'un agent i ne peut entretenir l'intention que A s'il croit que A n'est jamais réalisable : d'une part, « l'intention que A » doit être abandonnée dès lors que i croit que A sera toujours faux, et d'autre part « l'intention que A » ne peut être générée si i croit que A sera toujours faux.

On pourrait tenir compte de cet aspect dans l'axiome ($\text{Gen}_{\text{Int}1}$). Nous ne l'avons pas incorporé ici afin de ne pas alourdir notre formalisme par un opérateur temporel.

REMARQUE 3 Notons que ($\text{Gen}_{\text{Int}1}$) est trop fort dans le cas où il y a plus de deux agents. En effet, supposons que i coopère avec deux agents j et k , et que i pense que j et k ont des intentions contradictoires : $Bel_i Int_j A \wedge Bel_i Int_k \neg A$. Supposons en plus que $\neg Bel_i Int_j A \wedge \neg Int_i \neg A \wedge \neg Int_i A$ (i.e. i « se moque totalement de A »), alors avec ($\text{Gen}_{\text{Int}1}$) i génère les intentions $Int_i Bel_i A$ et $Int_i Bel_i \neg A$, ce qui est inconsistant par le réalisme fort ($\text{Rel}_{\text{IntBel}1}$).

Une manière de tenir compte de telles inconsistances éventuelles est d'affaiblir ($\text{Gen}_{\text{Int}1}$) en rajoutant une condition dans les prémisses stipulant que l'intention de j doit être consistante avec les intentions attribuées par i aux autres agents :

$$(Bel_i Int_j A \wedge \neg Bel_i A \wedge C) \rightarrow Int_i Bel_i A$$

où C est une formule de la forme $\neg Bel_i Int_{k_1} \neg A \dots \wedge \neg Bel_i Int_{k_n} \neg A$. Une telle condition pourra également tenir compte de priorités et préférences de i relatives aux intentions des autres agents.

Une autre manière d'affaiblir ($\text{Gen}_{\text{Int}1}$) est de stipuler que i ne peut pas rester sans intention dès qu'il apprend que j et k ne sont pas d'accord. Ceci peut être obtenu par le principe

$$Bel_i Int_j A \rightarrow (Int_i A \vee Int_i \neg A)$$

6. En effet, dans le cas où, en plus, $\neg Bel_i \neg A$, on ne peut avoir l'intention que A puisque $\neg Bel_i \neg A$ implique $\neg Int_i A$.

6.3 La génération d'intention : théorèmes et principes dérivés

Dans la suite nous discutons deux autres principes importants, et nous montrons qu'elles peuvent être dérivés de notre axiome central.

D'abord, notre axiome central (Gen_{Int1}) assure dans une certaine mesure que l'agent ne générera l'intention de croire A que s'il n'a pas déjà une intention contradictoire :

THÉORÈME 3 $\neg Bel_i A \rightarrow \neg Int_i \neg Bel_i A$

PREUVE 3 Ceci peut être montré à partir de l'axiome (5) de la logique pour la croyance, et du théorème 1 : une instance de ce théorème est $Bel_i \neg Bel_i A \rightarrow \neg Int_i \neg Bel_i A$. L'axiome (5) de la logique pour la croyance étant $\neg Bel_i A \rightarrow Bel_i \neg Bel_i A$, on obtient donc le théorème 3.

En accord avec ($Rel_{IntBel2}$) et (Gen_{Int1}), si un agent i croit qu'un agent j a l'intention que A soit vrai et qu'il n'a pas l'intention que A soit faux, alors l'agent i adopte l'intention que A soit vrai. Par le théorème 1, si l'agent i croit que A est faux alors il n'a pas l'intention que A soit faux. Soit le premier principe suivant :

PRINCIPE 1

$$(Bel_i Int_j A \wedge Bel_i \neg A) \rightarrow Int_i A \quad (Gen_{Int2})$$

Dans le cas où i croit qu'il faut nécessairement agir pour changer le monde, l'intention de j est donc directement adoptée en tant que telle.

THÉORÈME 4 L'axiome (Gen_{Int1}) implique (Gen_{Int2}).

PREUVE 4 L'hypothèse est $Bel_i Int_j A \wedge Bel_i \neg A$. D'une part, $Bel_i \neg A \rightarrow \neg Bel_i A$ avec l'axiome (D), et nous obtenons ainsi la deuxième hypothèse de (Gen_{Int1}).

Pour établir la troisième hypothèse de (Gen_{Int1}) : d'abord, $Int_i Bel_i \neg A \rightarrow Bel_i \neg Bel_i \neg A$ avec ($Rel_{IntBel1}$). Ensuite, $Bel_i \neg Bel_i \neg A \rightarrow \neg Bel_i \neg A$ avec l'axiome (5). Nous obtenons donc $Int_i Bel_i \neg A \rightarrow \neg Bel_i \neg A$, et par contraposition $Bel_i \neg A \rightarrow \neg Int_i Bel_i \neg A$.

Ainsi les hypothèses de l'axiome (Gen_{Int2}) impliquent celles de l'axiome (Gen_{Int1}). Ce dernier nous permet donc d'obtenir $Int_i Bel_i A$:

$$(Bel_i Int_j A \wedge Bel_i \neg A) \rightarrow Int_i Bel_i A$$

Or $(Int_i Bel_i A \wedge Bel_i \neg A) \rightarrow Int_i A$ d'après ($Rel_{IntBel2}$), ce qui établit (Gen_{Int2}).

Le second principe de génération d'intention stipule que si l'agent i croit que l'agent j a l'intention que A soit vrai, et qu'il croit que A est actuellement vrai, alors il va générer l'intention que j croie A . Soit :

PRINCIPE 2

$$(Bel_i Int_j A \wedge Bel_i A) \rightarrow Int_i Bel_j A \quad (Gen_{Int3})$$

REMARQUE 4 Si (Gen_{Int1}) est appliqué, et que l'intention ainsi générée est satisfaite, alors (Gen_{Int3}) sera applicable.

L'axiome que nous avons formulé nous donne de bonnes propriétés, comme nous venons de le montrer, pour l'intention et la coopération. Dans la section suivante, nous montrons que c'est également le cas pour la compréhension de directifs et de questions via des actes d'assertion.

REMARQUE 5 Les conditions $Bel_i Int_j A$ et $Bel_i A$ de (Gen_{Int3}) ne peuvent être vraies simultanément que si j n'est pas compétent sur A . En effet, $Bel_i Int_j A \rightarrow Bel_i Bel_j \neg A$ par ($Rel_{IntBel1}$), et si j était compétent sur A alors on aurait $Bel_i Bel_j \neg A \rightarrow Bel_i \neg A$, ce qui ne peut être le cas puisque $Bel_i A$ et $Bel_i \neg A$ sont inconsistantes.

THÉORÈME 5 L'axiome (Gen_{Int1}) implique (Gen_{Int3}).

PREUVE 5 (Gen_{Int1}) permet de dériver (Gen_{Int2}) dont $(Bel_i Int_j Bel_j A \wedge Bel_i \neg Bel_j A) \rightarrow Int_i Bel_j A$ en est une instance.

Comme $Int_j Bel_j A$ implique $Bel_j \neg Bel_j A$ par ($Rel_{IntBel1}$), nous avons $Bel_i Int_j Bel_j A \rightarrow Bel_i Bel_j \neg Bel_j A$ par les principes de la logique modale K, et comme par ailleurs $Bel_j \neg Bel_j A$ est équivalent à $\neg Bel_j A$ dans KD45, nous obtenons $Bel_i Int_j Bel_j A \rightarrow Bel_i \neg Bel_j A$.

Ainsi, nous obtenons $Bel_i Int_j Bel_j A \rightarrow Int_i Bel_j A$ à partir de (Gen_{Int2}).

D'autre part, $Int_j A$ impliquant $Int_j Bel_j A$ par ($Rel_{IntBel3}$), nous avons $Bel_i Int_j A \rightarrow Bel_i Int_j Bel_j A$ par les principes de la logique modale K.

La transitivité de $\rightarrow$ nous permet alors de conclure que $Bel_i Int_j A \rightarrow Int_i Bel_j A$.

7 Application à une architecture BDI

Nous allons maintenant montrer comment des questions fermées et des requêtes peuvent être inférées à partir d'assertions.

7.1 Traitement d'une question fermée

Soit $\alpha = \langle u, s, Int_u Bel If_u A \rangle$ l'acte venant d'être accompli et correspondant à l'énoncé « u affirme qu'il aimerait savoir si A est vrai ou non ». Comme nous l'avons montré précédemment (Sect. 2), cet acte constitue une forme d'indirection qui doit être interprétée comme une question fermée.

Les effets produits par l'acte α sur l'agent s sont les suivants (cf. Sect. 5) :

1. $Bel_s Int_u Bel_s Int_u Bel_s Int_u Bel If_u A$
2. $Bel_s Bel_u Int_u Bel If_u A$
3. $Bel_s \neg Bel_u Bel If_s Int_u Bel If_u A$

et correspondent respectivement à l'effet intentionnel (1) et l'effet indirect (à savoir la sincérité (2) et la pertinence au contexte (3)).

En supposant que tout agent est compétent sur ses propres attitudes mentales (comme l'illustrent les axiomes (Rel_{IntBel}4) et (Rel_{IntBel}5)) :

$$Bel_s Int_u BelIf_u A$$

est une conséquence de (2) via (Rel_{IntBel}4). Le principe (Rel_{IntBel}1) (avec les principes de la logique KD45 pour les croyances) implique que

$$Bel_s Bel_u \neg BelIf_u A$$

De nouveau par les principes de KD45, $Bel_u \neg BelIf_u A$ est équivalent à $\neg BelIf_u A$,⁷ et nous obtenons ainsi

$$Bel_s \neg BelIf_u A$$

Enfin, étant données les prémisses $Bel_s Int_u BelIf_u A$ et $Bel_s \neg BelIf_u A$, le principe (Gen_{Int}2) permet de conclure

$$Int_s BelIf_u A$$

Ce faisant, l'agent s satisfait l'intention initiale de l'agent u qui était que s adopte l'intention que u sache si A est vrai ou non.

7.2 Traitement d'une requête

Soit $\alpha = \langle u, s, Int_u Done_\beta \top \rangle$ l'acte venant d'être accompli, où s est l'auteur de β . Cet acte constitue une forme d'indirection qui doit être interprétée comme une requête (cf. Sect. 2).

Les effets produits par l'acte α sur l'agent s sont les suivants (cf. Sect. 5) :

1. $Bel_s Int_u Bel_s Int_u Done_\beta \top$
2. $Bel_s Bel_u Int_u Done_\beta \top$
3. $Bel_s \neg Bel_u BelIf_s Int_u Done_\beta \top$

et correspondent respectivement à l'effet intentionnel (1) et l'effet indirect (resp. (2) et (3)).

En supposant que tout agent est compétent sur ses propres attitudes mentales,

$$Bel_s Int_u Done_\beta \top$$

est une conséquence de (2). Si on suppose que l'agent s croit qu'il ne vient pas d'accomplir l'acte β , i.e.

$$Bel_s \neg Done_\beta \top$$

alors il va adopter l'intention d'exécuter β (via (Gen_{Int}2)). Ce faisant, l'agent s satisfait l'intention initiale de l'agent u qui était que s accomplisse β .

REMARQUE 6 Si β est une action devant être exécutée non pas par s , mais par u lui-même, l'énoncé de u correspondant serait du type « [Je veux / j'ai l'intention] de faire β ».

7. $Bel_i \neg BelIf_i A$ est par définition $Bel_i (\neg Bel_i A \wedge \neg Bel_i \neg A)$, qui est équivalent à $Bel_i \neg Bel_i A \wedge Bel_i \neg Bel_i \neg A$ par les principes modaux standards. Par les principes d'introspection de KD45, cette dernière formule est équivalente à $\neg Bel_i A \wedge \neg Bel_i \neg A$, qui n'est rien d'autre que $\neg BelIf_i A$.

Il est tout à fait envisageable d'interpréter cet énoncé de façon indirecte, c.-à-d. de considérer que par cet énoncé, u demande (sur le mode allusif) à s de faire l'action à sa place. Pour traiter ce cas, il faudrait alors appliquer un principe supplémentaire du type « si un agent i croit qu'un agent j a l'intention d'exécuter une action, alors l'agent aura l'intention d'exécuter cette action ». Ainsi, nous pouvons toujours ajouter des axiomes pour prendre en compte des phénomènes langagiers plus fins.

Pour les besoins de notre exemple, nous avons supposé que l'agent s savait qu'il ne venait pas déjà d'accomplir β . Si maintenant on suppose au contraire que l'agent s croit qu'il a déjà accompli β , il générera l'intention que l'agent u le sache (via (Gen_{Int}3)⁸). Selon la réaction de u (« Je n'ai pas entendu | pas compris | pas retenu | ... »), l'agent s pourra éventuellement exécuter β de nouveau (ce qui nécessitera la génération d'une nouvelle intention puisque la précédente a été satisfaite en exécutant β la première fois).

De même, on pourrait supposer que l'agent ne sait plus s'il a ou non exécuté l'action β . L'intention générée par (Gen_{Int}1) pourrait alors être reliée à une recherche approfondie de l'agent s dans sa mémoire afin de savoir s'il a déjà exécuté β ou pas. Selon la réponse, il générera alors de ce fait l'une des intentions via (Gen_{Int}2) ou (Gen_{Int}3).

Ce dernier cas illustre formellement notre point de vue précédent (cf. Sect. 4) sur le problème de l'extinction de la lumière dans une pièce où on ne sait pas si la lumière est allumée ou déjà éteinte. Il vient en cela à l'appui de notre thèse selon laquelle, contrairement à ce qui a été dit, l'axiome (Rel_{IntBel}1) ne constitue pas un lien trop fort entre croyance et intention (bien que, par ailleurs, il nécessite d'être complété par d'autres principes).

8 Conclusion

Nous avons présenté une logique minimale pour l'interaction coopérative. Elle est basée sur une notion d'intention primitive satisfaisant le principe selon lequel *avoir l'intention que A implique croire que actuellement A est faux*, et elle contient seulement des assertions comme type d'acte de langage.

Dans ce cadre, nous avons montré comment des requêtes et questions fermées peuvent être inférées à partir d'assertions d'une manière similaire aux actes de langage indirects. L'inférence se fait moyennant des principes de coopération dont les plus importants sont originaux. Nous avons ainsi montré que notre logique minimale est néanmoins capable de raisonner sur la communication dans un environnement coopératif.

L'intérêt d'une telle démarche est qu'en réduisant le nombre de types d'actes illocutoires traités par la logique, d'une part on diminue les risques d'erreur sur la force illocutoire de l'énoncé lors du traitement de ce dernier au cours

8. Du fait de notre sémantique associée à l'opérateur $Done_\beta$, $Bel_s Done_\beta \top$ ne pourra être vrai, le dernier acte accompli étant forcément α (l'acte d'assertion de l'agent u).

des phases de reconnaissance / compréhension de parole, et d'autre part on limite la prolifération d'axiomes propres à chaque type d'illocution qui ne fait qu'augmenter la complexité de la logique (axiomatique et sémantique), en particulier la gestion de l'évolution des croyances des agents. Ainsi, assertifs, requêtes, questions fermées, évolution des croyances et des intentions des agents, sont traités à l'aide de seulement huit axiomes (en plus de l'axiomatique liée à chacune des attitudes mentales des agents, et constituant les propriétés communes à tous les agents, en dehors de toute autre hypothèse de coopération, de sincérité, etc.).

Remerciements

Nous tenons à remercier les rapporteurs anonymes de RFIA ainsi que l'auteur du rapport de lecture complémentaire pour leurs lectures attentives et leurs commentaires de qualité, qui nous ont permis de clarifier certains points et corriger quelques coquilles.

Merci également à Jérôme Lang, Fabio Massacci et Jacques Virbel pour leurs commentaires.

Références

- [1] John L. AUSTIN. *How To Do Things With Words*. Oxford University Press, 1962.
- [2] Philippe BRETIER. « *La communication orale coopérative : contribution à la modélisation logique et à la mise en œuvre d'un agent rationnel dialoguant* ». Thèse de doctorat, Université Paris Nord, Paris, France, 1995.
- [3] Maud CHAMPAGNE, Rémi FAURE, Andreas HERZIG, Dominique LONGIN, Christophe LUC, Jean-Luc NESPOULOUS, et Jacques VIRBEL. « Formalisation logique de la communication non littérale à la lumière d'aperçus pragmatiques et neuropsycholinguistiques ». Dans *Proc. Journées Francophones Modèles Formels de l'Interaction (MFI'01)*, 2001. 17 pages.
- [4] Maud CHAMPAGNE et Dominique LONGIN. « Non literal communication: From pragmatic to logical and psycholinguistical aspects ». Dans *Int. Colloc. on Cognitive Sciences (ICCS'01)*, 2001. 23 pages.
- [5] Philip R. COHEN et Hector J. LEVESQUE. « Intention is Choice with Commitment ». *Artificial Intelligence Journal*, 42(2–3), 1990.
- [6] Olivier GASQUET, Andreas HERZIG, et Dominique LONGIN. « An analysis of communication in a logic of belief, intention and action ». Rapport Technique IRIT/2001-07-R, Institut de Recherche en Informatique de Toulouse, mars 2001. 22 pages. Available in <http://www.irit.fr/ACTIVITES/LILaC>.
- [7] H. Paul GRICE. *Logic and Conversation*. Dans J.P. COLE et J.L. MORGAN, éditeurs, *Syntaxe and Semantics: Speech acts*, volume 3: *Speech Acts*. Academic Press, 1975.
- [8] David HAREL. *Dynamic Logic*. Dans D. GABBAY et F. GUENTHNER, éditeurs, *Handbook of Philosophical Logic*, volume II. D. Reidel Publishing Company, 1984.
- [9] Andreas HERZIG et Dominique LONGIN. « Belief Dynamics in Cooperative Dialogues ». *Journal of Semantics*, 17(2), 2000. 20 pages.
- [10] Andreas HERZIG, Dominique LONGIN, et Jacques VIRBEL. « Towards an Analysis of Dialogue Acts and Indirect Speech Acts in a BDI Framework ». Dans Massimo POESIO et David TRAUM, éditeurs, *Proc. Fourth Int. Workshop on the Semantics and Pragmatics of Dialogue (GötaLog-2000)*, Göteborg University, Sweden, 2000.
- [11] Ray REITER. « The frame problem in the situation calculus: A simple solution (sometimes) and a completeness result for goal regression ». *Artificial Intelligence and Mathematical Theory of Computation*, Papers in Honor of John McCarthy:359–380, 1991.
- [12] M. D. SADEK. « *Attitudes mentales et interaction rationnelle : vers une théorie formelle de la communication* ». Thèse de doctorat, Université de Rennes I, Rennes, France, 1991.
- [13] M. D. SADEK. « A Study in the Logic of Intention ». Dans Bernhard NEBEL, Charles RICH, et William SWARTOUT, éditeurs, *Proc. Third Int. Conf. on Principles of Knowledge Representation and Reasoning (KR'92)*, pages 462–473. Morgan Kaufmann Publishers, 1992.
- [14] M. D. SADEK. « Dialogue Acts are Rational Plans ». Dans M.M. TAYLOR, F. NÉEL, et D.G. BOUWHUIS, éditeurs, *The structure of multimodal dialogue*, pages 167–188, Philadelphia/Amsterdam, 2000. John Benjamins publishing company. From ESCA/ETRW, Workshop on The Structure of Multimodal Dialogue (Venaco II), 1991.
- [15] M. D. SADEK, P. BRETIER, et F. PANAGET. « ARTIMIS: Natural dialogue meets rational agency ». Dans Martha E. POLLACK, éditeur, *Proc. 15th Int. Joint Conf. on Artificial Intelligence (IJCAI'97)*, pages 1030–1035. Morgan Kaufmann Publishers, 1997.
- [16] Richard SCHERL et Hector J. LEVESQUE. « The frame problem and knowledge producing actions ». Dans *Proc. Nat. Conf. on AI (AAAI'93)*, pages 689–695. AAAI Press, 1993.
- [17] J. R. SEARLE. *Expression and Meaning*. Cambridge University Press, 1979.
- [18] John R. SEARLE. *Speech acts: An essay in the philosophy of language*. Cambridge University Press, 1969.
- [19] Christel SEGUIN. « *De l'action à l'intention : vers une caractérisation formelle des agents* ». Thèse de

doctorat, Université Paul Sabatier, Toulouse, France, mars 1992.

- [20] S. SHAPIRO, Yves LESPÉRANCE, et Hector J. LEVESQUE. « Specifying communicative multi-agent systems with ConGolog ». Dans *Working notes of the AAAI fall symposium on Communicative Action in Humans and Machines*, pages 75–82. AAAI Press, 1997.
- [21] S. SHAPIRO, M. PAGNUCCO, Y. LESPÉRANCE, et H. J. LEVESQUE. « Iterated Belief Change in the Situation Calculus ». Dans *Proc. IJCAI'99*, pages 527–538, 2000.
- [22] Michael THIELSCHER. « Representing the Knowledge of a Robot ». Dans A. COHN, F. GIUNCHIGLIA, et B. SELMAN, éditeurs, *Proc. KR'00*, pages 109–120. Morgan Kaufmann, 2000.
- [23] Daniel VANDERVEKEN. « Formal Pragmatics of Non Literal Meaning ». *Linguistische Berichte*, 1997.
- [24] Jacques VIRBEL. Contributions de la théorie des actes de langage à une taxinomie des consignes. Dans J. VIRBEL, J-M. CELLIER, et J-L. NESPOULOUS, éditeurs, *Cognition, Discours procédural, Action*, volume II. PRESCOT, 1999.