



Modèle dynamique pour la gestion du tour de parole d'humains virtuels

Mathieu Jégou

► To cite this version:

Mathieu Jégou. Modèle dynamique pour la gestion du tour de parole d'humains virtuels. Synthèse d'image et réalité virtuelle [cs.GR]. 2013. dumas-00854996

HAL Id: dumas-00854996

<https://dumas.ccsd.cnrs.fr/dumas-00854996>

Submitted on 28 Aug 2013

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.



STAGE DE MASTER RECHERCHE

RAPPORT DE STAGE

Modèle dynamique pour la gestion du tour de parole d'humains virtuels

Auteur :
Mathieu JÉGOU

Encadrant :
Pierre CHEVAILLIER
Lab-STICC
Centre Européen de Réalité
Virtuelle



Table des matières

1	Introduction	3
2	Tour de parole et modèles de prise de décision et d'action	4
2.1	Tour de parole dans une interaction dialogique	4
2.1.1	Caractéristiques du tour de parole	4
2.1.2	Signaux interprétés dans le cadre du tour de parole	5
2.1.3	Synthèse et hypothèses par rapport au tour de parole	6
2.2	Architectures d'agents pour la gestion du tour de parole	6
2.2.1	Contraintes sur les architectures d'agents conversationnels	6
2.2.2	Ymir (Thórisson, 1999)	7
2.2.3	Synthèse sur les architectures d'agents	8
2.3	Modélisation de la prise de décision et de l'action chez un humain	8
2.3.1	Dynamique comportementale	9
2.3.2	Choix mutuellement exclusifs et prise de décision	10
2.3.3	Positionnement et choix du modèle	12
3	Modèle conceptuel pour la gestion du tour de parole	14
3.1	Principes du modèle	14
3.1.1	Prise de décision	14
3.1.2	Contrôle de l'action	15
3.1.3	Mise en relation des deux modèles	15
3.2	Application du modèle au tour de parole	17
3.2.1	Prise de décision	17
3.2.2	Production de signaux	19
3.3	Couplage des modèles du locuteur et de l'auditeur	23
3.3.1	Premier scénario	23
3.3.2	Deuxième scénario	24
3.3.3	Troisième scénario	27
3.3.4	Synthèse	29
4	Architecture d'agent pour l'implémentation du modèle	29
4.1	Principes d'Ymir (Thórisson, 1999)	30
4.2	Solution pour l'intégration du modèle	31
4.2.1	Capteurs	31
4.2.2	Perception, décision et sélection d'action	32
4.2.3	Actionneurs	32
4.3	Place du modèle conceptuel dans l'architecture	33
4.4	Implémentation de l'architecture	33
5	Discussion	34
5.1	Vérification des hypothèses	34
5.2	Dynamique comportementale pour la prise de décision	35
5.3	Interaction entre plus de deux agents	35
5.4	Signaux interprétés	36
5.5	Validation du modèle	36
6	Conclusion	36

Table des figures

1	Canaux de communication non verbale	5
2	Architecture Ymir	7
3	Cycle de perception-action	9
4	Cycle d'hystérésis	11
5	Simulation de DDM	13
6	Simulation de prise de décision	13
7	Couplage entre interlocuteurs	16
8	Implémentation du modèle	18
9	Fonction $f(a_l, r_{int}, r_{conf})$ pour le contrôle du niveau sonore	21
10	Évolution temporelle des actions et de la prise de décision du locuteur et de l'auditeur pour le premier scénario	24
11	Niveaux sonores du premier scénario	25
12	Évolution temporelle des actions et de la prise de décision du locuteur et de l'auditeur pour le second scénario	26
13	Niveaux sonores du second scénario	27
14	Évolution temporelle des actions et de la prise de décision du locuteur et de l'auditeur pour le troisième scénario	28
15	Niveaux sonores du troisième scénario	29
16	Schéma général de l'architecture	31
17	Scène utilisée pour l'implémentation du modèle	34

Au cours d’une interaction dialogique, les participants ont deux rôles possibles, celui de locuteur et celui d’auditeur. Le participant ayant le tour de parole à un instant donné est celui qui a le rôle de locuteur. Au cours de la conversation, les rôles changent. Les transitions de tour de parole, ne sont pas contrôlées par un participant en particulier mais leur gestion est distribuée entre les différents protagonistes. Le tour de parole est ainsi émergent des interactions entre les participants et auto-organisé. La gestion du tour de parole dans la plupart des architectures d’agents conversationnels, se fonde sur des prises de décisions ponctuelles, sur la base de règles. Nous proposons, au contraire, un modèle continu de décision et de mise en action, rendant compte du caractère auto-organisé du tour de parole. Ce modèle conceptuel se fonde sur deux modèles de psychologie cognitive, le modèle d’accumulation-diffusion (*drift diffusion model* ou *DDM*) de Ratcliff (1978) et la dynamique comportementale de Warren (2006). Nous montrons que ce modèle rend compte des observations rapportées par différents auteurs sur l’occurrence du changement de tour de parole entre les locuteurs.

Mots clés : dynamique comportementale, modèle d’accumulation-diffusion, agents conversationnels, tour de parole

1 Introduction

Les environnements virtuels collaboratifs tels que les environnements collaboratifs pour l’apprentissage sont peuplés d’humains virtuels qui interagissent entre eux et avec l’utilisateur. Ceux-ci sont contrôlés soit par un utilisateur, soit par un agent autonome (Gerbaud et al. (2007) ; Swartout et al. (2006) ; Edward et al. (2008)).

Une interaction spontanée entre un utilisateur et des humains virtuels autonomes nécessite de la part des agents virtuels le respect des mêmes modes d’interactions que ceux observés lors des interactions humaines. La communication verbale est une des interactions possibles entre un utilisateur et un agent virtuel. Le langage naturel impose du côté de l’agent, l’interprétation de ce que dit l’utilisateur et la génération de réponse, mais aussi le respect des règles d’interaction entre humains dans une conversation.

Le dialogue définit deux rôles, celui de locuteur et d’auditeur, le participant ayant le tour de parole est celui qui, à un instant donné est dans le rôle de locuteur. La dynamique du changement de tour de parole est émergente, auto-organisée et distribuée. Chaque participant modifie ses signaux de prise et d’abandon de tour de parole en fonction des signaux produits par l’autre, de ce couplage entre les participants émerge le moment de transition de tour de parole.

Plusieurs architectures proposent un modèle de prise de décision dans le cadre du tour de parole (Rickel and Johnson, 1997 ; Nakano et al., 2003 ; Thórisson, 2002). Elles sont toutes fondées sur des prises de décision ponctuelles sur la base de signaux provenant de l’environnement et de règles de déclenchement d’action. Ces architectures ne rendent pas compte du caractère continu de la prise de décision dans le cadre de la coordination d’actions entre humains. L’objectif ici est de proposer un modèle continu permettant, sur la base d’une accumulation de signaux provenant d’un autre interlocuteur, et de la modulation en continu des signaux des agents, de rendre compte du caractère émergent du tour de parole.

Le modèle élaboré se fonde sur le couplage de deux modèles, le modèle d’accumulation-diffusion (DDM ou *drift diffusion model*, (Ratcliff, 1978)), simulant une prise de décision

continue par accumulation, pour l'interprétation du comportement de l'autre interlocuteur et la dynamique comportementale (Warren, 2006) pour la production de signaux de début et de fin de tour de parole.

Ce rapport se structure de la manière suivante. Les fondements théoriques nécessaires au développement de notre modèle, et de l'architecture sont tout d'abord exposés. Il s'agit de définir les caractéristiques du comportement humain dans une interaction dialogique, puis d'étudier les architectures d'agents conversationnels existantes, avant de présenter des modèles généraux du comportement humain permettant de rendre compte du caractère auto-organisé de la dynamique du tour de parole, ainsi que l'anticipation de la transition de tour de parole par les humains. Ensuite, le modèle réalisé est exposé en présentant d'abord les principes généraux du modèle, puis son application dans le cadre du tour de parole, avant de décrire les résultats observés lors de l'exécution du modèle. Dans la partie suivante, l'architecture d'agent réalisée pour l'implémentation du modèle est présentée. Une discussion sur le modèle réalisé conclut enfin ce mémoire.

2 Tour de parole et modèles de prise de décision et d'action

2.1 Tour de parole dans une interaction dialogique

2.1.1 Caractéristiques du tour de parole

La gestion du tour de parole entre les participants est cruciale pour le bon déroulement d'une conversation. Un tour de parole est défini comme un intervalle durant lequel un seul participant a le rôle de locuteur.

Au cours de la conversation le tour de parole change. Le passage du tour de parole d'un interlocuteur à un autre s'effectue à des instants précis, que Sacks et al. (1974) nomment des moments pertinents de transition (*transition relevant place* ou *TRP*). Au cours de ces transitions, on observe que le changement de locuteur s'effectue de trois manières (Sacks et al., 1974) :

1. si le locuteur courant sélectionne un interlocuteur comme futur locuteur, celui-ci prend le tour ;
2. si le locuteur ne sélectionne personne comme prochain locuteur, n'importe qui peut prendre le tour de parole ;
3. si personne ne prend le tour de parole, le locuteur courant reprend le tour.

Au cours des TRP, le tour de parole se transmet de manière fluide entre les participants. Les recouvrements de parole sont évités, perçus comme agressifs (ter Maat et al., 2010), et non-coopératifs (Grice, 1989, cité par Thórisson 2002), tout comme un moment de silence trop long entre deux tours de parole (ter Maat et al., 2010). Une transition de tour de parole sans recouvrement ni rupture trop longue constitue la transition observée le plus souvent selon Sacks et al.. Les moments de la conversation où deux participants prennent le tour en même temps ne sont pas rares mais sont brefs. Dans ce second cas de figure un des deux locuteurs s'arrêtera afin de ne pas laisser le recouvrement continuer (Sacks et al., 1974).

L'observation montre ainsi que les temps de transition entre deux tours (temps où un locuteur finit de parler et où un autre commence à parler) sont inférieurs à 100 millisecondes soit inférieurs au temps minimal nécessaire pour percevoir et agir chez un être

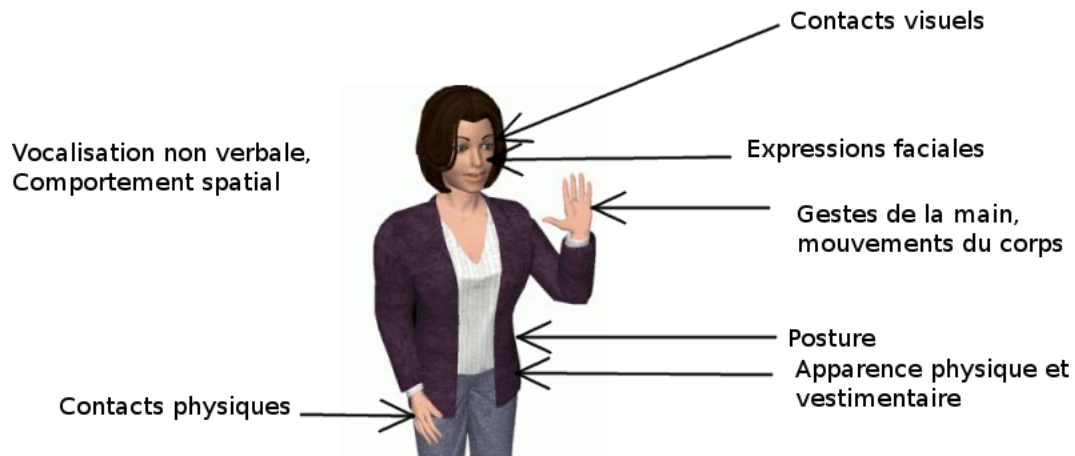


FIGURE 1 – Canaux de communication non verbale illustré ici sur l’agent GRETA (source : Mancini et al., 2007)

humain. Un interlocuteur décide donc d’agir avant que le moment de transition ait lieu (Goodwin, 1981 cité par Thórisson, 2002).

La décision de l’apparition d’un moment de changement de tour n’est de plus sous le contrôle d’aucun participant en particulier. Sacks et al. (1974), énonce trois propriétés concernant la dynamique du tour de parole :

1. c’est une propriété émergente,
2. de décisions locales,
3. qui s’effectuent sur la base de prédictions par les participants.

Le tour de parole est émergent dans le sens où la durée des tours, le prochain locuteur ainsi que les moments où s’effectuent les transitions dépendent des interactions complexes entre les décisions prises par chaque participant. La gestion du tour de parole est aussi vu comme géré localement, la distribution des tours, la durée d’un tour, et l’ordre des tours ne sont pas planifiés à l’avance. Au contraire, toutes les décisions prises par les participants sont dirigées vers la prochaine transition, sans prévision des tours succédant au prochain tour. Enfin, toujours selon Sacks et al. (1974), les auditeurs prédisent le moment où la fin de tour apparaît en intégrant de manière continue des indices provenant du discours du locuteur.

2.1.2 Signaux interprétés dans le cadre du tour de parole

Les signaux intégrés dans le cadre du tour de parole sont de nature multimodale. La communication se fait ainsi de manière verbale (relatif au contenu) et non verbale. Lewis (1998) cité par Allbeck and Badler (2001) définit neuf canaux de communication non verbale présents dans une conversation (cf figure 1) :

1. les expressions faciales (sourires, hochements de tête),
2. les gestes,
3. les mouvements du corps,
4. la posture,
5. l’orientation visuelle,

6. les contacts physiques entre interlocuteurs,
7. le comportement spatial,
8. l'apparence,
9. les vocalisations non verbales (l'intonation, le niveau sonore, des sons produits ne correspondant pas à des mots).

Ainsi dans le cadre du tour de parole, le contenu de la conversation autant que les canaux non verbaux sont des intermédiaires utilisés par un interlocuteur pour communiquer ses intentions concernant la prise ou l'abandon du tour de parole.

Certains auteurs ont observé qu'un certain nombre de signaux verbaux et non verbaux était utilisé pour l'identification du moment de transition. Sacks et al. (1974) introduit les unités de construction de tour (*turn constructional unit* ou *TCU*), des indices syntaxiques utilisés pour la prédiction de la fin de tour. Le locuteur les produit volontairement afin de construire le moment de changement de tour, la tâche des auditeurs est alors de les reconnaître afin de prédire la fin de tour du locuteur. D'autres auteurs ont observé l'utilisation de signaux non verbaux. Duncan (1972) introduit ainsi une intonation finale montante ou descendante et la terminaison de gestes de la main dans les signaux de fin de tour. Enfin, Gravano and Hirschberg (2011) déterminent six signaux et variations de signaux verbaux ou non verbaux utilisés pour l'anticipation :

1. une intonation montante ou descendante,
2. une vitesse de parole plus grande,
3. une intensité sonore qui diminue,
4. un point de complétion textuelle (savoir si le discours du locuteur est complet syntaxiquement ou sémantiquement),
5. une diminution du nombre de pauses entre deux unités de parole,
6. une voix plus bruitée.

2.1.3 Synthèse et hypothèses par rapport au tour de parole

Les sections 2.1.1 et 2.1.2 ont mis en avant le caractère auto-organisé de la dynamique des tours de parole. L'agent ne se décide de plus pas à agir lorsque le TRP apparaît, il agit avant que le moment de transition se produise, sur la base de signaux multimodaux provenant des autres interlocuteurs.

Nous émettons de plus l'hypothèse que le moment de transition n'est pas prédit par les participants mais que chaque agent accumule de manière continue des indices éventuellement contradictoires sur l'intention de l'autre. L'agent module alors ses propres signaux selon les indices reçus de la part des autres interlocuteurs. Cette modulation de signal rétroagit sur l'autre participant qui, en retour modifie ses propres signaux. Le moment de transition n'est alors ni complètement contrôlé par chaque participant, ni planifié à l'avance, puisque les actions d'un participant sont couplées aux actions des autres participants.

2.2 Architectures d'agents pour la gestion du tour de parole

2.2.1 Contraintes sur les architectures d'agents conversationnels

L'analyse du tour de parole intègre un certain nombre de contraintes pour la modélisation d'un agent capable de gérer de manière aussi fluide qu'un humain le tour de parole.

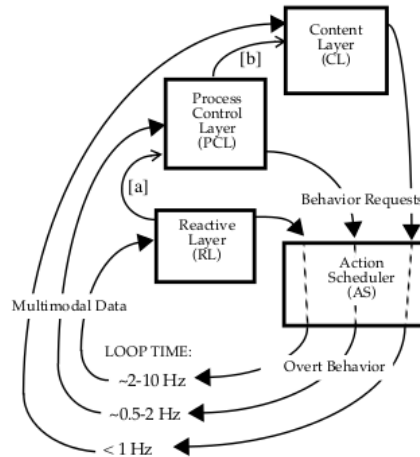


FIGURE 2 – Schéma représentant le fonctionnement de l’architecture Ymir (tiré de Thórisson, 2002). Les modules des trois couches reçoivent des signaux provenant des capteurs et prennent une décision sur l’action à produire sur la base de ces signaux. L’action est envoyée au planificateur d’action qui se charge de choisir le ou les mouvements associés à cette action ainsi que de gérer les priorités entre actions

- La multimodalité des signaux implique qu’un agent doit être capable d’interpréter et de produire en parallèle des signaux de nature différente.
- La modulation en continu des signaux des agents nécessite qu’une action puisse être paramétrable et, qu’une fois lancée, celle-ci puisse être modifiée.
- Les notions de moment où apparaît la transition et de longueur de transition impliquent une contrainte de temps réel forte du côté de l’agent. Pour que l’agent soit crédible concernant la gestion du tour de parole, il doit percevoir et agir aux mêmes échelles temporelles qu’un humain.

Certains agents conversationnels comme STEVE (Rickel and Johnson, 1997) utilisant l’architecture SOAR, planifient leurs actions à partir d’un modèle interne et selon un système de règles, la décision y est ponctuelle.

D’autres agents comme MACK (Nakano et al., 2003) interprètent avant tout le contenu de la conversation, peu de signaux non verbaux sont intégrés par l’agent, ces signaux sont des signaux de l’utilisateur lui indiquant sa compréhension ou son incompréhension de ce qu’a dit l’agent. Thórisson (2002) s’est spécifiquement intéressé à la gestion du tour de parole, son architecture, Ymir, a un certain nombre de caractéristiques intéressantes concernant les propriétés énoncées et les hypothèses faites dans la section 2.1.3. Cette architecture est décrite dans la section suivante.

2.2.2 Ymir (Thórisson, 1999)

Ymir répond à un besoin de modéliser les contraintes de temps réel et la prise en compte de signaux et d’actions multimodaux. L’architecture est composée de modules s’exécutant en parallèle, répartis en plusieurs couches et définissant plusieurs fréquences d’exécution des modules de l’architecture (figure 2). L’utilisation de couches d’exécution permet de définir des temps de réaction et de mise en action précis pour l’agent.

Deux types de modules s’exécutent dans l’architecture, les percepteurs se chargent de faire correspondre aux données provenant des capteurs des données perceptuelles utiles à

la prise de décision de l'agent.

Les décideurs reçoivent un ensemble de données provenant des percepteurs, et se chargent d'exécuter une action selon les données reçues.

Les données échangées entre les modules sont de nature booléenne. Ces données sont intégrées dans les conditions de déclenchement du module. Pour un percepteur, la donnée en sortie est associée à la valeur de vérité de cette condition. Pour les décideurs, les décisions associées sont déclenchées lorsque la condition est remplie.

Le processus de transformation de la décision en action s'effectue grâce à un lexique moteur. Le lexique moteur permet d'associer à une décision un ensemble d'actions réalisées par l'agent, ces actions pouvant être effectuées en parallèle sur des canaux différents.

Une action réalisée par le planificateur d'action a :

- un nom,
- une durée,
- un délai avant le début de son exécution,
- des paramètres dépendants de l'action réalisée.

2.2.3 Synthèse sur les architectures d'agents

Dans Ymir, les actions produites par les agents ne sont pas continuellement modulables : les actions sont déclenchées, durent un certain temps, puis s'arrêtent. Ceci est contradictoire avec les propriétés énoncées dans la partie 2.1.3. La continuité de la décision et de l'action implique un contrôle en continu de l'action entamée, puisqu'un interlocuteur ne planifie pas le moment où la transition se fera, il ne peut de même pas prévoir la durée de son action.

De plus, les actions physiques réalisées par l'agent sont certes contrôlées par des paramètres continus, néanmoins, la décision prise conduisant à l'action est, elle, discrète. Une décision prise par les décideurs n'est pas paramétrable.

Cependant l'architecture est intéressante dans la mesure où elle gère l'interprétation et la production de signaux multimodaux. Le temps réel y est de même inclus, en spécifiant les fréquences liées au cycle de perception-action.

L'architecture Ymir, moyennant certaines adaptations, notamment la transformation du mécanisme de perception, de décision et d'action pour intégrer un modèle continu, semble intéressante pour l'implémentation du modèle. Néanmoins, Ymir est un modèle computationnel d'agent, il ne spécifie pas les processus amenant un agent, à partir des signaux provenant de l'environnement, à moduler son action, d'où la nécessité de définir un tel modèle conceptuel.

2.3 Modélisation de la prise de décision et de l'action chez un humain

À partir de l'étude des caractéristiques du tour de parole, deux composantes à prendre en compte dans le modèle du tour de parole sont identifiées. La prise de décision sur accumulation de signaux nécessite de trouver un modèle déterminant comment varie l'accumulation de signaux en fonction des données provenant de l'environnement. La production de signal dans le cadre du tour de parole nécessite, lui, un second modèle. Ces deux composantes peuvent être simulées avec deux modèles continus de psychologie cognitive, la dynamique comportementale de Warren (2006) pour la production de signal et le modèle d'accumulation-diffusion (*drift-diffusion model* ou *DDM*) pour la prise de décision (Rat-

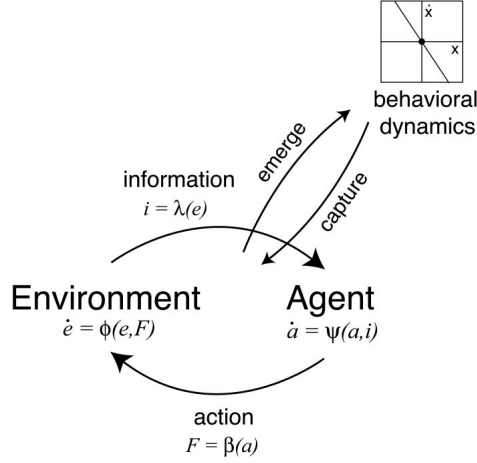


FIGURE 3 – Schéma du cycle de perception-action et de la dynamique comportementale proposé par Warren (2006)

cliff, 1978).

2.3.1 Dynamique comportementale

La dynamique comportementale s'inscrit dans le cadre de la théorie écologique de la perception formalisée par Gibson. Gibson considère que le processus de perception provient directement d'un type d'information de l'environnement, les invariants (Goldstein, 1981). Les stimuli reçus de l'environnement varient continuellement, néanmoins certains éléments de ces stimuli demeurent constants. Ces invariants représentent les éléments sur lesquels se base le contrôle de l'action de l'agent. Cette vision de la perception s'oppose aux théories incluant, entre l'information reçue de l'environnement et ce qui est perçu par l'agent, une étape de traitement intermédiaire, transformant les stimuli reçus en information symbolique. Plus précisément, le traitement de cette information inclut la mise à jour d'un modèle interne de l'environnement, ce modèle interne est ensuite utilisé comme base pour la prise de décision de l'agent (Loomis and Beall, 2004). Puisque, selon la théorie écologique, la décision est prise sur la base des invariants, Gibson considère que l'information est un régulateur direct du comportement (Warren, 2006). Cette régulation du comportement par l'information amène alors à dire que la structure de contrôle du comportement ne réside pas entièrement dans l'agent, ni dans l'environnement, mais est distribué entre l'agent et l'environnement (Warren, 2006).

L'idée d'un comportement s'établissant spontanément est aussi abordée dans la dynamique de la coordination élaborée par Kelso (Kelso, 2009). Kelso étudie les phénomènes de coordination entre êtres vivants. Dans certaines conditions, un système composé d'éléments différents devient une unité réalisant une fonction précise. Les éléments se spécialisent dans une fonction et se coordonnent afin d'assurer le comportement du système. Le comportement global du système apparaît spontanément des interactions entre les composants du système. Le comportement est stable, une fluctuation d'une partie de ce système est rapidement compensée par les autres éléments du système. La dynamique de la coordination de Kelso est une théorie formalisant ce type de système par des systèmes dynamiques. L'émergence d'un comportement stable du système est formalisée par les attracteurs. Le changement de comportement par des bifurcations. L'étude s'effectue à deux niveaux. Au niveau du système, on y décrit la dynamique générale, exprimée en terme de variables de contrôle (induisant un changement dans le comportement), de variables collectives (quan-

tités créées par la coopération entre les individus et rétroagissant sur eux), d'attracteurs et de bifurcations du système. Au niveau des individus, il s'agit de formaliser la manière dont ils interagissent entre eux, incluant le type d'information transmis entre les composants.

La dynamique comportementale s'inspire à la fois de la théorie écologique de la perception et de la dynamique de la coordination. Elle reprend les deux niveaux d'analyse de la dynamique de la coordination, décrivant le comportement global et les composants du système sous forme de système dynamique. Le système décrit est composé de deux entités, l'agent et l'environnement. De leur couplage apparaît spontanément le comportement, celui-ci rétroagissant sur les actions de l'agent. La figure 3 décrit ce fonctionnement. L'agent contrôle ses variables d'actions en fonction des informations ($i = \lambda(e)$) de l'environnement ($\dot{a} = \Psi(a, i)$), ces informations se traduisant en terme de variables de contrôle dans le système de l'agent. Les actions produites par l'agent se traduisent en forces exercées sur l'environnement ($F = \beta(a)$), représentant des variables de contrôle du système environnement ($\dot{e} = \Phi(e, F)$).

Le comportement est alors une trajectoire dans l'espace d'état du système agent-environnement ($\dot{x} = \Omega(x, r)$). La stabilité du comportement est due à la convergence de ce système vers un attracteur, sa flexibilité, c'est à dire la capacité du système à changer de type de comportement, est due aux bifurcations du système dynamique.

La dynamique comportementale a été appliquée avec succès, dans la modélisation de certaines tâches, notamment la locomotion (Fajen and Warren, 2003). Néanmoins la dynamique comportementale pose quelques problèmes de recherche actuels. Warren (2006) identifie quatre types d'activités dont la formulation actuelle de la dynamique comportementale ne rend pas compte :

1. les tâches séquentielles sont représentées par des buts se décomposant en sous-buts. Elles se décomposent en une série d'actions dont l'enchaînement demande des connaissances préalables sur la tâche ;
2. certaines tâches impliquent des comportements d'anticipation. Dans ce cas, un humain cherche à atteindre des buts qui ne sont pas immédiatement présents dans son environnement ;
3. dans les tâches prédictives, l'agent utilise des connaissances préalables sur les objets de l'environnement, non disponibles par l'information reçue. La manière d'aborder le problème du comportement prédictif est d'établir des relations entre les propriétés cachées des objets et les informations réellement disponibles sur l'objet. Cette relation entre les propriétés visibles et cachées d'un objet peuvent être apprises par l'expérience ;
4. enfin, les tâches stratégiques se basent sur un contexte dépassant la connaissance actuelle de l'environnement, un agent peut par exemple supposer qu'une personne se comportera à un instant donné d'une certaine manière.

2.3.2 Choix mutuellement exclusifs et prise de décision

Dans le cadre du tour de parole, la décision concernant la nature des intentions de l'autre s'effectue entre deux alternatives, pour un auditeur, le locuteur laisse ou garde la parole, pour un locuteur, l'auditeur prend ou pas la parole. Plusieurs modèles traitent du choix d'un agent entre deux alternatives. Le modèle de Frank et al. (2009) utilise le cadre de la dynamique comportementale, pour traiter de la prise de décision entre la saisie à

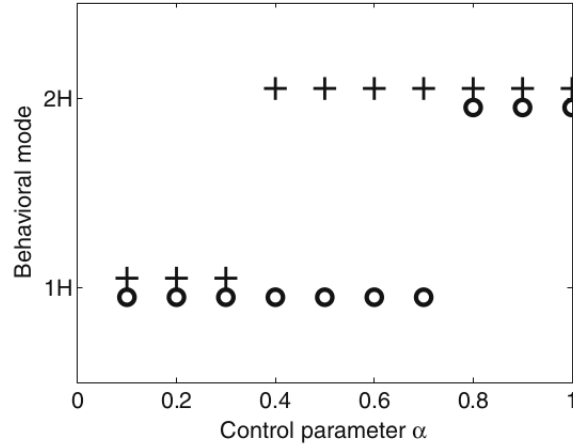


FIGURE 4 – Cycle d’hystérésis observé lors d’un comportement de saisie (Frank et al., 2009). Le paramètre α représente le rapport entre la taille de la main et la taille de l’objet. Le symbole o représente l’action pour une variation ascendante de la variable de contrôle, le symbole $+$ représente l’action pour une variation descendante de la variable de contrôle.

une main d’un objet et la saisie à deux mains (*grasping*). Le modèle élaboré se présente sous la forme des équations suivantes :

$$\begin{aligned}\dot{\xi}_1 &= \lambda_1 \xi_1 - B \xi_2^2 \xi_1 - C(\xi_1^2 + \xi_2^2) \xi_1 \\ \dot{\xi}_2 &= \lambda_2 \xi_2 - B \xi_1^2 \xi_2 - C(\xi_1^2 + \xi_2^2) \xi_2\end{aligned}\tag{1}$$

ξ_1 et ξ_2 représentent l’amplitude des mouvements à une main et à deux mains, ils sont, soit strictement positifs, soit nuls, selon que l’on observe l’action ou non. λ_1 et λ_2 représentent les taux d’accroissement de ξ_1 et ξ_2 lorsque ceux ci sont proches de 0. Ces paramètres sont fonctions d’une variable de contrôle α représentant le rapport entre la taille de l’objet et la taille de la main. En accord avec les observations expérimentales, le paramètre α montre deux seuils de transition d’action, l’un pour la transition de la saisie à une main à la saisie à deux mains, l’autre pour la transition de la saisie à deux mains à la saisie à une main, ces seuils de transition sont montrés sur la figure 4. La prise de décision dans le cadre de la saisie à deux mains montre donc un cycle d’hystérésis. Les modèles de tâches de choix forcé à deux alternatives (*two alternative forced choice tasks* ou *T AFC*), s’appliquent lorsque la prise de décision remplit ces trois conditions d’application :

1. la décision à prendre doit être effectuée entre deux alternatives ;
2. la décision doit être réalisée en temps limité ;
3. l’environnement de l’agent est bruité, il acquiert plus ou moins bien les informations provenant de l’environnement.

Les modèles sont de plus fondés sur trois hypothèses concernant la nature du processus de décision :

1. le processus de décision s’effectue en accumulant de manière continue des indices provenant de l’environnement ;
2. celui-ci est soumis à des fluctuations aléatoires ;

3. une décision est prise lorsque l'agent a récolté suffisamment d'indices favorisant l'une ou l'autre des alternatives.

Les modèles de TAFC sont tous bâtis de la manière suivante. La quantité d'indices favorisant chaque alternative varie au cours du temps. Le processus de décision n'est pas infini, l'agent prend à un moment donné une décision définitive sur le choix à réaliser. Cette décision est prise, soit lorsqu'une valeur seuil de quantité d'indice favorisant une alternative par rapport à l'autre est atteinte, soit une limite de temps est atteinte, l'alternative ayant récolté le plus d'indices est alors choisie.

Il existe plusieurs modèles de TAFC néanmoins sous certaines conditions d'optimalité, ces modèles peuvent se réduire au modèle de DDM (Bogacz et al., 2006). L'équation du DDM se présente sous la forme suivante :

$$dx = A dt + c dW \quad (2)$$

Cette équation est une équation différentielle stochastique représentant un processus aléatoire. Les indices ne sont pas récoltés séparément, on évalue la différence entre les indices accumulés. dx représente la variation d'accumulation d'indices, A représente le taux d'accroissement moyen de la différence dans les indices accumulés, $c dW$ est un terme aléatoire suivant une loi normale centrée, de variance c . Le DDM définit de plus deux seuils de décision, un seuil positif et un seuil négatif représentant les deux alternatives possibles. Si le taux d'accumulation est négatif, la quantité d'indices variera en moyenne vers le seuil négatif, si le taux d'accumulation est positif, la quantité d'indices variera en moyenne vers le seuil positif. Ces deux seuils peuvent être égaux en terme de valeur absolue, ou différents. Une différence représente un biais dans la prise de décision, l'agent est au début du processus de décision plutôt favorable envers l'alternative dont le seuil a la valeur absolue la plus petite. La quantité d'indice de départ de l'agent peut être soit nulle, soit non nulle. Une quantité d'indice non nulle représente, de la même manière qu'une différence dans les seuils, un biais dans la décision.

Selon les applications au modèle, le taux d'accumulation peut être soit constant pendant le processus de prise de décision (Ratcliff, 1978), soit variable (Ratcliff, 1980). Le choix d'une variable A constante, s'effectue lorsque l'environnement de l'agent ne varie pas, les informations reçues par l'agent ne changent pas. Au contraire dans un environnement dynamique, les indices reçus varient au cours du temps, changeant le taux d'accumulation. La figure 5 montre un exemple de processus de décision selon le DDM avec A constant et la figure 6 montre un exemple de processus de décision avec A variant au cours du temps. Dans le premier cas, A étant positif, la quantité d'indice varie en moyenne en faveur de l'alternative représentée par le seuil positif, la variance de 0,2 induit des fluctuations fortes dans cet accroissement. x atteint le seuil 1 au bout de 3,8 secondes. Dans le deuxième cas, on constate que la différence dans la quantité d'indices se situe autour de 0 jusqu'à 5 secondes, dû à un taux d'accumulation nul. Ensuite, le taux d'accumulation varie de 0 à 0,5 résultant dans le processus de décision en une évolution de l'accumulation d'indices vers l'alternative représentée par le seuil 1. La décision est prise en faveur de cette alternative au bout de 7 secondes.

2.3.3 Positionnement et choix du modèle

Nous avons analysé différents modèles pour la prise de décision et l'action dans le cadre du tour de parole. La dynamique comportementale de Warren (2006) possède les mêmes caractéristiques que le tour de parole, c'est à dire l'auto-organisation du comportement,



FIGURE 5 – Exemple de simulation du DDM. L’axe des abscisses représente le temps en secondes, l’axe des ordonnées représente la quantité x de preuves accumulées au cours du temps. Le taux d’accroissement A est ici égal à $0,2$, le paramètre c est égal à $0,2$, les seuils valent 1 et -1 et la quantité d’indices initiale est à 0 .

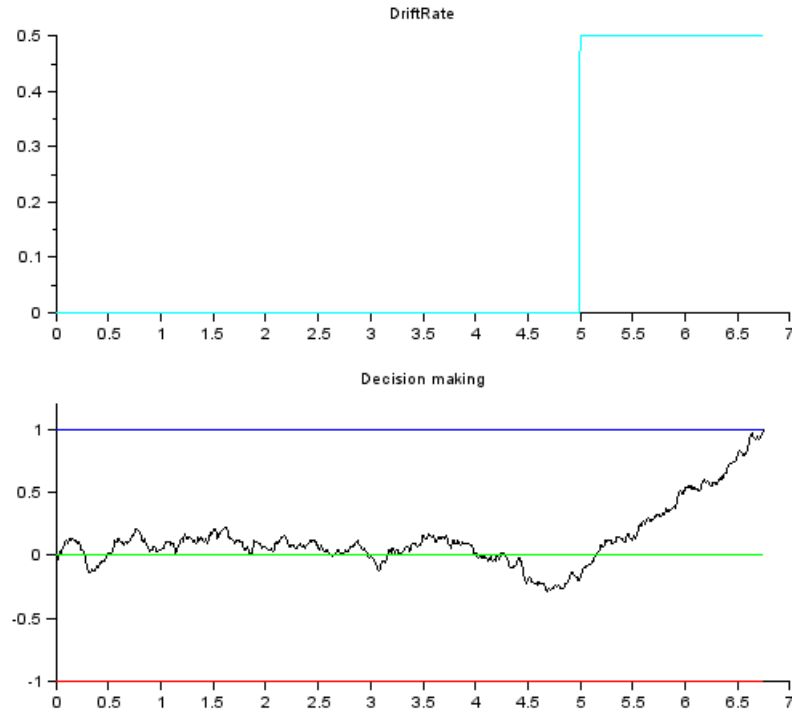


FIGURE 6 – Exemple de simulation de prise de décision avec un taux d’accumulation non constant. L’axe des abscisses représente le temps en secondes, l’axe des ordonnées représente la quantité d’indices x . Ici $c = 0,2$, les seuils valent 1 et -1 , la quantité d’indice initiale est égale à 0 .

ne dépendant pas seulement de l'agent mais aussi de l'environnement dans lequel l'agent évolue, dans notre cas, les autres participants à la conversation. Il semble ainsi pertinent d'utiliser le cadre de la théorie de Warren.

Un modèle fondé sur la dynamique comportementale traite d'une activité où un agent a deux actions possibles et mutuellement exclusives (Frank et al., 2009). Néanmoins, ce modèle se positionne dans le cas où un agent est certain de la nature des informations provenant de l'environnement, et ne modélise pas la notion d'accumulation de signal énoncée dans la section 2.1.3. De plus, la gestion du tour de parole est une activité d'anticipation (l'agent agit avant le moment de transition), et stratégique, types d'activités identifiés comme des limites au modèle actuel de dynamique comportementale (Warren, 2006).

La gestion du tour de parole entre au contraire parfaitement dans les conditions d'application des modèles de choix forcé à deux alternatives, et plus particulièrement du DDM. Nous proposons ainsi de coupler les deux modèles, la dynamique comportementale pour le contrôle de l'action de l'agent, et le DDM pour la prise de décision.

3 Modèle conceptuel pour la gestion du tour de parole

Nous allons maintenant présenter le modèle élaboré durant le stage. La section se structure comme suit. Nous allons d'abord présenter les principes généraux du modèle. Il s'agit d'appliquer tout d'abord les modèles de DDM et de dynamique comportementale à des tâches de coordination où l'agent accumule des indices concernant les signaux de l'autre et agit en fonction de cette accumulation. Ensuite, nous montrons comment la décision de l'agent est incluse dans les équations de mise en action. Nous présentons ensuite l'application de ce modèle au tour de parole, en précisant les différents signaux produits et interprétés par le locuteur et l'auditeur, et les équations associées d'interprétation et de contrôle. Puis nous illustrerons la mise en application du modèle avec trois scénarios de transition de tour de parole.

3.1 Principes du modèle

3.1.1 Prise de décision

Les signaux des agents variant au cours du temps, nous avons fait le choix d'un taux d'accumulation non constant. Celui-ci varie continuellement en fonction des signaux produits par l'autre agent.

Dans notre modèle, à chaque signal produit par celui-ci, est associée une valeur d'accumulation. Ce taux d'accumulation est pondéré selon l'importance de ce signal par rapport à d'autres signaux accumulés par l'agent. L'équation 3 représente la forme générale de modification du taux d'accumulation global en fonction des signaux produits par l'autre agent.

$$A = \sum_{j=0}^{n_s} c_{s_j} f_{Accj}(s_j) \quad (3)$$

Avec, n_s le nombre de signaux interprétés par l'agent, s_j un type de signal accumulé par l'agent, c_{s_j} le coefficient de pondération du signal s_j par rapport aux autres signaux, f_{Accj} la fonction attribuant une valeur d'accumulation au signal s_j .

3.1.2 Contrôle de l'action

Forme générale de l'équation La production de signal de l'agent est représentée par un ensemble d'équations différentielles. Ces équations différentielles s'inspirent de l'équation de locomotion élaborée par Fajen and Warren (2003). L'équation 4 est la formule générale utilisée pour la production de signal.

$$\ddot{x} = -b\dot{x} - f_{attr}(x) \quad (4)$$

La fonction $f_{attr}(x)$ dépend du type d'action que réalise l'agent. Elle capture la définition des points fixes du système, et les bifurcations possibles du système en fonction de variables de contrôle que nous définirons plus loin.

L'équation est composée de deux termes, le premier représente une inertie (*damping*) et le deuxième, la convergence de l'action vers un attracteur.

Il reste à ajouter dans l'équation les variables de contrôle modifiant l'action que réalise l'agent.

Prise en compte d'une valeur d'intention Dans le modèle de Fajen et Warren, le but de l'agent est implicite et ne change pas au cours du temps. Dans le cadre du tour de parole, le but d'un agent correspond au fait que l'agent veut parler ou écouter un autre participant. L'agent manifeste une intention variable à adopter l'un ou l'autre des rôles, locuteur ou auditeur. Ce qui gouverne cette intention n'entre pas dans le champ de notre étude. Dans notre modèle, l'intention de l'agent prend la forme d'une variable de contrôle de l'action de l'agent. L'équation 5 reprend l'équation 4 en incluant la variable d'intention comme variable de contrôle de la fonction f_{attr} .

$$\ddot{x} = -b\dot{x} - f_{attr}(x, r_{int}) \quad (5)$$

Avec r_{int} la valeur d'intention définie dans l'intervalle $[0; 1]$. La manière dont r_{int} modifie l'action n'est pas prédéfinie, son influence varie selon l'action réalisée.

3.1.3 Mise en relation des deux modèles

Prise en compte d'une variable de confiance Le modèle réalisé consiste en un couplage des modèles présentés au dessus. La liaison entre les deux modèles s'effectue par l'intermédiaire d'une valeur de confiance ajoutée au modèle de mise en action. Cette valeur de confiance représente la certitude ou non que les autres interlocuteurs produisent ou non une action particulière et provient du modèle de DDM. Cette notion de confiance influençant l'action modélise l'hésitation d'un interlocuteur à produire un signal, ou au contraire l'urgence à le produire, soit parce que l'agent est proche du seuil de décision, soit parce qu'il est incertain du type d'action de l'autre. Aussi, de la même manière que pour l'intention, la valeur de confiance est une variable de contrôle de l'équation de mouvement. Cette variable est égale à la quantité d'indices accumulée par l'équation de prise de décision de l'agent (équation 2). L'équation 6 reprend l'équation 5 en ajoutant comme variable de contrôle de f_{attr} la variable de confiance r_{conf} .

$$\ddot{x} = -b\dot{x} - f_{attr}(x, r_{int}, r_{conf}) \quad (6)$$

r_{conf} est défini, pour le modèle, dans l'intervalle $[-1; 1]$.

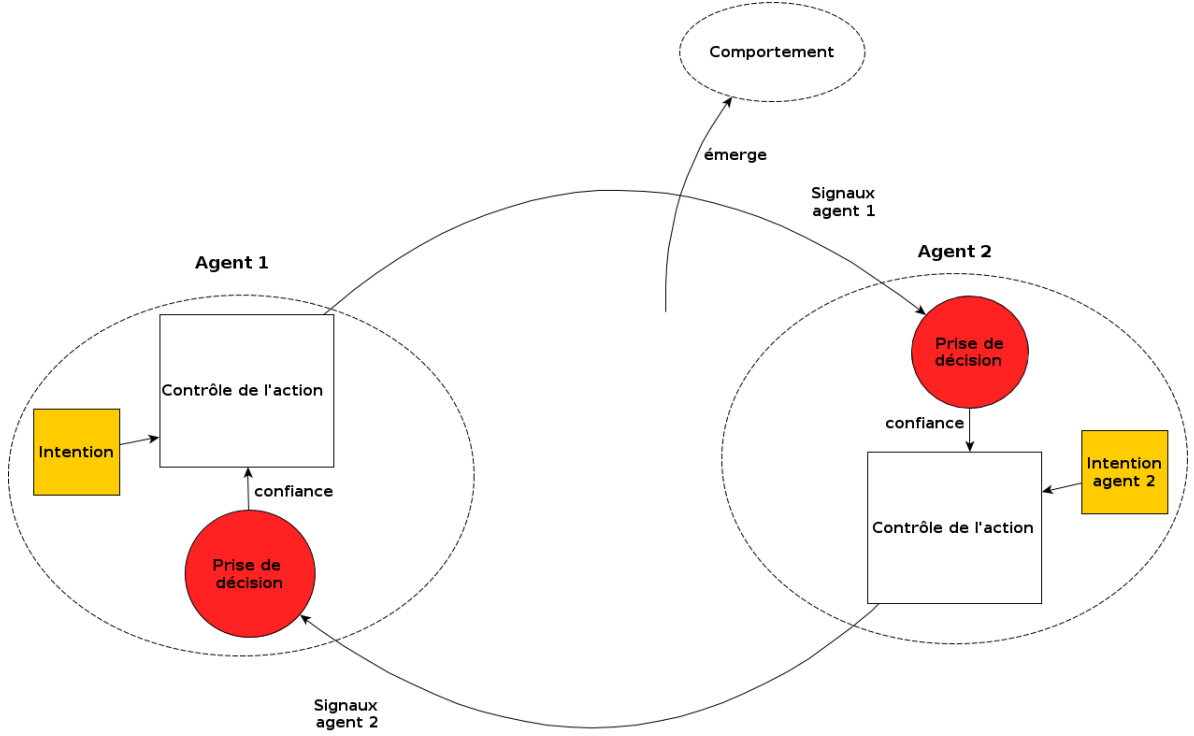


FIGURE 7 – Représentation du couplage entre agents généré par le modèle et émergence d'un comportement global au système formé par les deux agents

Prise en compte du franchissement de seuil L'équation de prise de décision franchit à un moment donné un des deux seuils de décision. La décision prise par l'agent est alors définitive. Lorsque ce seuil est franchi, l'agent exécute une autre action spécifiquement associée au franchissement de ce seuil. L'action peut être différente de l'action produite pendant le processus de décision.

Le choix effectué dans l'équation 7 est d'y ajouter deux variables discrètes p_{thre} et n_{thre} afin de modéliser une notion de déclenchement d'action lorsque l'agent franchit un seuil. p_{thre} représente le franchissement du seuil positif et n_{thre} représente le franchissement du seuil négatif. Si ces variables sont à 0, le seuil n'a pas encore été franchi par l'agent, si une de ces variables est à 1, le seuil correspondant a été franchi. L'équation 7 reprend l'équation 6 en y ajoutant les variables n_{thre} et p_{thre} .

$$\begin{aligned} \ddot{x} = & -b\dot{x} - ((1 - p_{thre})(1 - n_{thre})f_{attr}(x, r_{int}, r_{conf}) \\ & + n_{thre}k_2(x - x_{attr2}) \\ & + p_{thre}k_3(x - x_{attr3})) \end{aligned} \quad (7)$$

L'agent se retrouve avec trois actions possibles, une action produite lors du processus de décision (terme en $(1 - p_{thre})(1 - n_{thre})$, $r_{conf} \in]-1; 1[$), une action produite correspondant au franchissement du seuil positif (terme en p_{thre} , $r_{conf} = 1$), une action produite correspondant au franchissement du seuil négatif (terme en n_{thre} , $r_{conf} = -1$). Les paramètres k_2 et k_3 représentent des paramètres de raideurs, ils déterminent la rapidité de convergence vers les attracteurs x_{attr2} et x_{attr3} .

Ces trois termes sont non nuls à des moments différents. Les seuils de l'équation de prise de décision ne sont pas franchis en même temps, empêchant les termes en n_{thre} et p_{thre} d'être non nuls en même temps. Le franchissement d'un seuil a pour effet d'annuler le terme en $(1 - p_{thre})(1 - n_{thre})$, empêchant les premier et deuxième termes, et les premier

et troisième termes d'être non nuls au même moment.

L'intention et la confiance n'interviennent pas dans les actions de franchissement de seuil, lorsque l'agent a pris sa décision, il produit l'action associée sans tenir compte de sa réticence ou sa résolution à changer de rôle. Si un interlocuteur est au départ réticent à changer de rôle, le franchissement d'un seuil positif le force à changer de rôle.

3.2 Application du modèle au tour de parole

Une fois formalisé le modèle de gestion du tour de parole, il s'agit de définir les rôles possibles des interlocuteurs, la nature de la prise de décision faite par les interlocuteurs, et les types de signaux produits par les agents.

Les rôles possibles des interlocuteurs proviennent directement de l'analyse du tour de parole effectuée dans la section 2.1 : soit locuteur, soit auditeur. Chaque agent interprète les signaux provenant d'autres agents, et produit ses propres signaux. Les signaux produits et interprétés par les agents diffèrent selon les rôles. Les signaux interprétés par le locuteur sont les suivants :

- l'auditeur se tourne spécifiquement vers le locuteur,
- le niveau sonore et le débit de l'auditeur monte,
- l'auditeur bouge le bras.

Les signaux interprétés par l'auditeur sont les suivants :

- le locuteur se tourne vers lui,
- le niveau sonore et le débit du locuteur diminue.

Le but du modèle est avant tout de montrer que du couplage entre les participants peut émerger le moment de transition. Certains signaux tels que le niveau sonore proviennent ainsi de l'analyse des signaux relatifs au tour de parole (cf Gravano and Hirschberg, 2011). D'autres signaux tels que bouger le bras et se tourner vers l'autre ont été imaginés comme de possibles signaux de prise ou d'abandon de tour de parole (cf section 2.1.2). Le signal "se tourner" est pertinent dans le contexte d'un environnement constitué de plus de deux agents. Le locuteur n'est pas forcément tourné spécifiquement vers un auditeur mais s'il remarque des signaux de prise de parole de la part d'un auditeur, il peut se tourner spécifiquement vers lui. La figure 8 montre une implémentation complète du modèle en précisant les rôles, la nature de la prise de décision, et les signaux produits par les deux participants.

3.2.1 Prise de décision

Décisions à prendre par les participants La nature de la prise de décision faite par le locuteur concernant l'action de l'auditeur est relative à la prise de tour de parole ou non de l'auditeur. Le locuteur répond à cette question par "oui" ou "non", la quantité d'indices associée à cette décision est initialement à 0. Si cette valeur de décision franchit le seuil positif, l'agent répond "oui" à cette question. Au contraire, si elle franchit le seuil négatif, l'agent répond "non" à cette question.

L'auditeur doit, lui, répondre à la question "le locuteur me donne le tour?". Pour une valeur initiale à 0, si cette valeur franchit le seuil positif, il répond "oui" à la question. Si cette valeur franchit le seuil négatif, il répond "non" à la question.

Le paramètre de variance du DDM est fixé : $c = 0, 1$. Les signaux interprétés par les agents sont tous définis arbitrairement dans l'intervalle $[0; 1]$, une valeur de 0 correspondant à un signal non produit.

Pour le niveau sonore, des valeurs de cette variable comprises entre 0 et 0,3 représentent

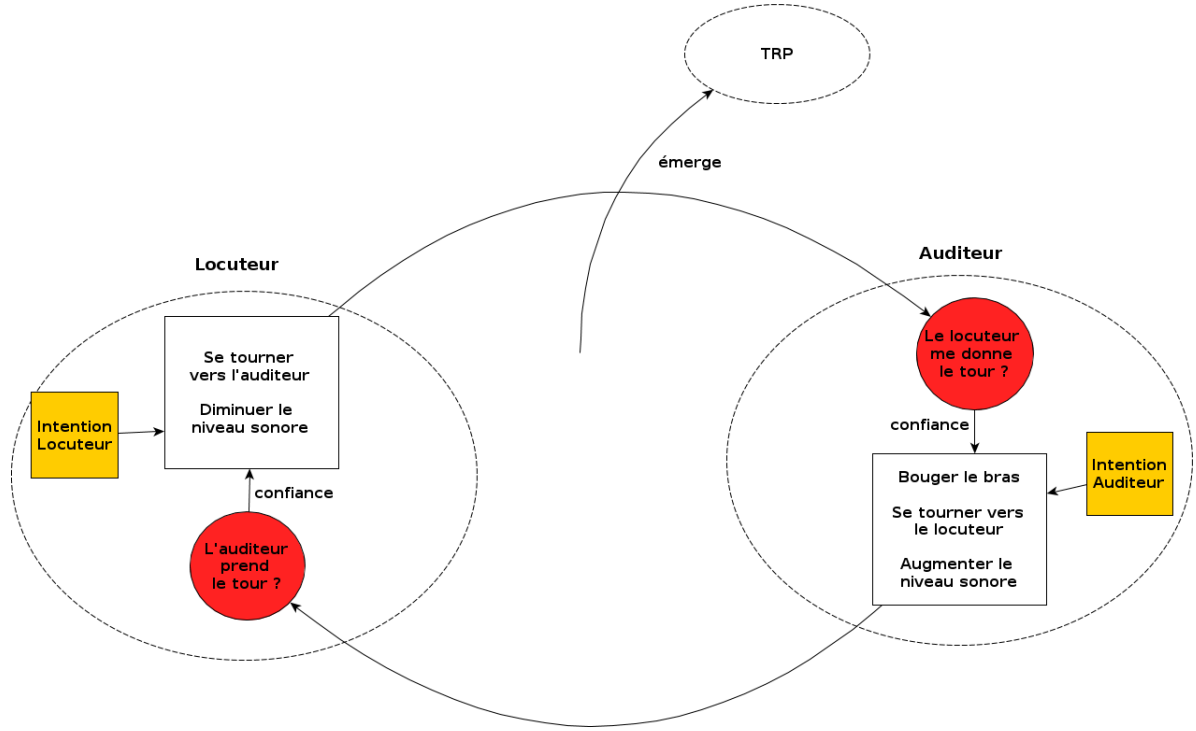


FIGURE 8 – Reprise du schéma de la figure 7 précisant la nature des rôles, des prises de décision et des actions réalisées

un niveau sonore faible, des valeurs comprises entre 0,3 et 0,6 représentent un niveau sonore moyen, des valeurs comprises entre 0,6 et 1 représentent un niveau sonore élevé. Les participants parlent le plus souvent avec un niveau sonore de 0,5, l'élévation du niveau sonore au-dessus de la valeur 0,5 correspond à un agent haussant la voix.

Nous allons maintenant détailler la manière dont les différents signaux sont pris en compte par le locuteur et l'auditeur, notamment la manière dont ils sont inclus dans l'équation de prise de décision.

Prise de décision du locuteur La valeur du taux d'accumulation A des équations de prise de décision varie en fonction des signaux reçus. Pour le locuteur la variation du taux d'accumulation est donnée par l'équation 8.

$$A_t = 0,8f_{audioL}(\dot{s}_a, s_a) + 0,6s_m + 0,3s_f \quad (8)$$

Les variables s_a , s_m et s_f représentent respectivement le niveau sonore de l'auditeur, le mouvement du bras, et une valeur associée au fait que l'auditeur fait face à au locuteur. Pour un seul signal, le niveau sonore, à la fois la valeur du signal en lui-même et sa dérivée sont pris en compte. Cette dérivée a une influence sur le taux d'accumulation, en effet un auditeur descendant son niveau sonore vers un niveau moyen ne traduit pas la même intention qu'un auditeur augmentant son niveau sonore vers cette même valeur. Pour les autres indices, le locuteur n'intègre pas la dérivée des signaux de l'agent. Leur dérivée a certes une influence sur le taux d'accumulation de l'agent, mais cette influence est plus négligeable que pour le niveau sonore. Afin de ne pas complexifier le modèle, la dérivée temporelle n'est pas représentée dans le modèle pour ces types de signaux.

Les valeurs de pondération 0,8, 0,6, représentent respectivement une prise en compte

forte du niveau sonore et du mouvement du bras du locuteur et la valeur 0,3, une prise en compte faible du fait que l'auditeur fait face à l'agent.

Lorsque le niveau sonore est faible le taux d'accumulation de l'agent est négatif. La convergence vers le seuil de décision négatif est alors inversement proportionnelle à la valeur du niveau sonore. Des niveaux sonores moyens et forts sont associés à une valeur positive du taux d'accumulation, la rapidité de convergence du niveau sonore est alors proportionnelle à la valeur du niveau sonore. Cette valeur est pondérée par la dérivée du niveau sonore. Pour un niveau sonore fixe, une dérivée positive a pour effet d'augmenter la valeur du taux d'accumulation, une dérivée négative produit l'effet inverse. Afin de modéliser ces contraintes, nous avons donné à f_{audioL} la forme indiquée par l'équation 9.

$$f_{audioL}(\dot{s}_a, s_a) = (0,25\dot{s}_a - 1,75)(1 - s_a) + 1 \quad (9)$$

Prise de décision de l'auditeur Pour l'auditeur le taux d'accumulation est déterminé par l'équation 10.

$$A_a = 0,6f_{audioA}(\dot{s}_a, s_a) + 0,1s_f \quad (10)$$

Les valeurs de pondération 0,6 et 0,1 représentent une prise en compte forte du niveau sonore du locuteur et faible de la valeur associée au fait que celui-ci fait face à l'agent.

Les variables s_a et $\dot{s}_a$, correspondent toujours au niveau sonore et sa dérivée, les valeurs qualitatives du niveau sonore liées aux valeurs de s_a sont les mêmes que pour le locuteur. La variable s_f représente la valeur associée au fait que le locuteur fait face à l'agent.

La fonction f_{audioA} se comporte à l'inverse de la fonction f_{audioL} . Un niveau sonore faible montre un locuteur en train de donner le tour, la valeur du taux d'accumulation est alors d'autant plus grande que le niveau sonore est faible. Au contraire, un niveau sonore élevé montre un locuteur ne souhaitant pas donner le tour, la valeur du taux d'accumulation est donc négative et devient d'autant plus petite que le niveau sonore est élevé. De la même manière que pour f_{audioL} , la valeur du taux d'accumulation est pondérée par la dérivée du niveau audio, l'action n'est en effet pas la même si le locuteur atteint une valeur de niveau sonore avec une dérivée positive, que si le locuteur atteint ce même niveau avec une dérivée négative. Le passage d'un taux d'accumulation négatif à un taux d'accumulation positif se situe à des valeurs faibles du niveau sonore, vers 0,3. La fonction f_{audioA} est donnée par l'équation 11.

$$f_{audioA}(\dot{s}_a, s_a) = (-\dot{s}_a - 2)(s_a - 0,5) - 0,3 \quad (11)$$

3.2.2 Production de signaux

Les actions implémentées dans le modèle correspondent aux signaux interprétés par les équations de prise de décision. Nous avons considéré deux types de signaux pour la prise de décision de l'auditeur (équation 10), le niveau sonore et le signal lié à l'angle relatif entre le locuteur et l'auditeur. Deux actions sont donc à considérer pour le locuteur, se tourner vers l'auditeur et moduler son niveau sonore. Pour la prise de décision du locuteur, trois signaux étaient pris en compte, le niveau sonore, l'angle relatif entre l'auditeur et le locuteur et le fait que l'auditeur lève le bras. L'auditeur exhibe donc trois types d'action, se tourner vers le locuteur, moduler son niveau sonore et bouger le bras. Nous avons fixé le paramètre b à la valeur 3.25 cette valeur est reprise de l'équation de locomotion de Fajen and Warren (2003).

Nous détaillons maintenant les équations de contrôle des signaux du locuteur et de l'auditeur. Pour chaque signal, nous présentons le résultat voulu, puis nous présentons les choix en termes d'attracteurs et de bifurcations. Nous présentons ensuite l'équation.

Actions du locuteur

Se tourner vers l'auditeur L'action de se tourner est un signal exhibé lorsque l'agent n'est pas encore certain des intentions de l'autre interlocuteur, l'agent se tourne vers l'autre interlocuteur pour inciter l'autre à clarifier son intention concernant la prise ou non de tour de parole. Au contraire, lorsque l'agent a des valeurs de confiance proches de 1 ou proches de -1, l'agent ressent moins le besoin de se tourner. Plus l'intention est grande dans la résolution ou la réticence à donner le tour plus, l'agent ressentira de même le besoin de se tourner. Lors du franchissement de seuil positif l'agent continue de se tourner vers l'autre interlocuteur, s'il ne lui fait pas complètement face. Le franchissement de seuil négatif a au contraire pour conséquence l'arrêt de cette action.

La variable d'état que l'équation fait varier est l'angle relatif Φ entre l'orientation du locuteur et la position de l'auditeur. L'axe de référence pris pour la mesure de l'angle est la direction actuelle du regard du locuteur. Cet angle se situe dans le repère du locuteur, aussi, l'action "se tourner" consiste pour le locuteur à ce que cet angle évolue vers 0. L'équation ne comporte ainsi qu'un seul attracteur situé en (0,0). Les variables d'intention et de confiance font varier le paramètre de raideur de l'équation, pour une intention fixe, le paramètre de raideur est maximal pour une valeur de confiance proche de 0.

Nous avons donné à l'équation de contrôle de l'action "se tourner" la forme suivante :

$$\ddot{\Phi} = -b\dot{\Phi} - k_g(1 - n_{thre})(\Phi)f_{turnL}(r_{int}, r_{conf}) \quad (12)$$

Avec $f_{turnL}(r_{int}, r_{conf})$ la fonction suivante :

$$f_{turnL}(r_{int}, r_{conf}) = -2(r_{int} - 0,5)^2(r_{conf} + 1)(r_{conf} - 2) \quad (13)$$

Le facteur $(r_{conf} + 1)(r_{conf} - 2)$ permet d'avoir un paramètre de raideur nul pour $r_{conf} = -1$, plus l'agent est proche du seuil -1 moins il se tournera vite, l'autre solution $r_{conf} = 2$, n'est jamais atteint ($r_{conf} \leq 1$), néanmoins à partir de la valeur $\frac{1}{2}$, le paramètre de raideur décroît plus la confiance augmente, ce qui correspond bien au fait que plus la confiance de l'agent augmente lorsque celle-ci est supérieure à $\frac{1}{2}$ moins il se tourne rapidement.

Le facteur $(r_{int} - 0,5)^2$, permet d'exhiber une action similaire selon que le locuteur est réticent ou résolu. La raideur est maximale pour $r_{int} = 0$ et $r_{int} = 1$, puis décroît moins l'agent est réticent ou résolu à donner le tour. La valeur du paramètre k_g est fixé à 0,75.

Modulation du niveau sonore Le locuteur a plusieurs actions possibles concernant la modulation de son niveau sonore. Si sa valeur de confiance est négative, il garde un niveau sonore normal. Si la valeur de confiance est positive, l'action dépend de la valeur d'intention. S'il est réticent à donner le tour alors s'il détecte que l'auditeur exhibe des signaux de prise de tour, il cherchera à montrer à l'auditeur qu'il ne souhaite pas céder le tour, il augmente alors son niveau sonore. S'il est résolu à donner le tour alors il baissera son niveau sonore pour laisser l'auditeur prendre le tour.

Lorsque l'équation de prise de décision franchit le seuil négatif, le locuteur reprend un

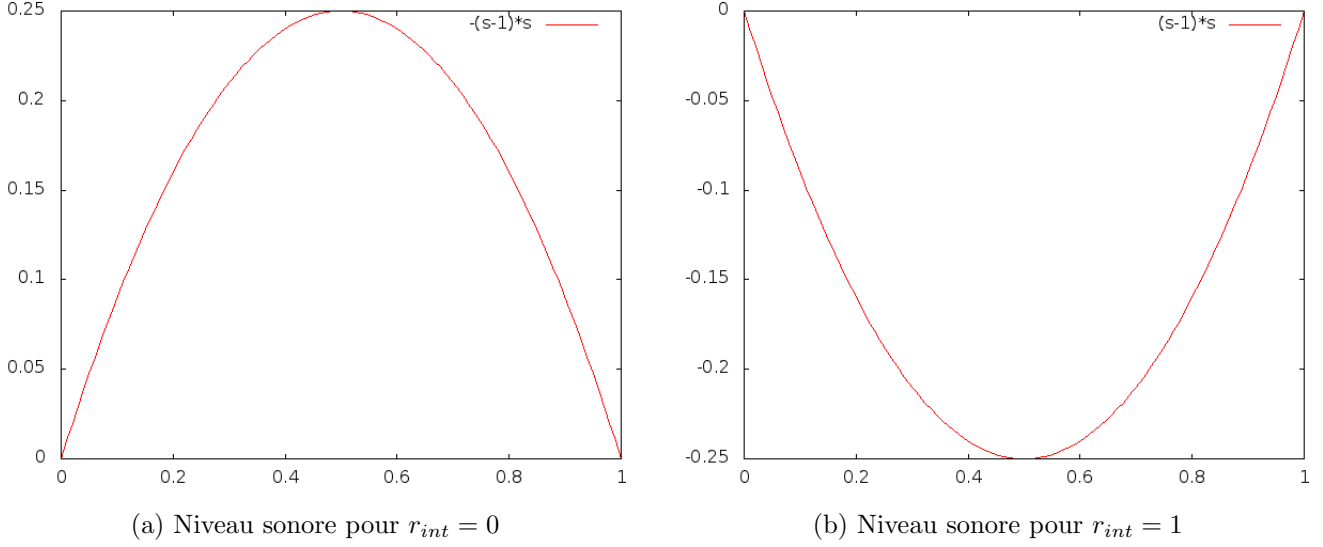


FIGURE 9 – Deux exemples d’évolution temporelle du niveau sonore pour le locuteur en faisant varier le paramètre d’intention. L’axe des ordonnées représente $\dot{a}_l$, l’axe des abscisses a_l . Le paramètre d’inertie de l’équation ne modifiant pas les attracteurs et bifurcations du système, l’évolution de $\ddot{a}_l$ n’est pas représenté. $r_{conf} = 1$ pour les deux courbes.

niveau sonore normal. Au contraire, si cette équation franchit le seuil positif de décision, alors le locuteur arrête de parler.

Ces actions sont maintenant formulées en terme d’attracteurs. Deux cas sont à traiter lorsque le processus de décision du locuteur est en cours. Si la confiance est positive alors deux points fixes sont présents, $(0; 0)$ et $(1; 0)$. La nature de ces points fixes dépend du paramètre d’intention. Si l’intention est supérieure à 0.5, le point en $(0; 0)$ est un attracteur tandis que le point $(1; 0)$ est un répulseur. Au contraire, si l’intention est inférieure à 0.5, le système a un attracteur en $(1; 0)$ et un répulseur en $(0; 0)$. La figure 9 montre ces deux cas de figure. Si la confiance est inférieure à 0 le système n’a alors qu’un seul attracteur situé en $(0, 5; 0)$. La confiance et l’intention induisent non seulement des bifurcations dans l’action, mais contrôle le paramètre de raideur. Ainsi plus la confiance du locuteur sera élevé plus il convergera rapidement vers son attracteur. Cette convergence est d’autant plus rapide que le locuteur est résolu ou réticent à donner le tour.

L’action de modulation du niveau sonore est représentée comme suit :

$$\begin{aligned}
 \ddot{a}_l &= -b\dot{a}_l \\
 &\quad - k_g((1 - p_{thre})(1 - n_{thre})f(a_l, r_{int}, r_{conf}) \\
 &\quad + n_{thre}(a_l - 0.5) \\
 &\quad + p_{thre}a_l) \\
 f_{levelL}(a_l, r_{int}, r_{conf}) &= (-2r_{int} + 1)r_{conf}u(r_{conf})(a_l - 1)a_l \\
 &\quad - (1 - u(r_{conf}))r_{conf}(a_l - 0, 5)
 \end{aligned} \tag{14}$$

La variable a_l représente le niveau sonore. Cette action est modélisée par la fonction $f_{levelL}(a_l, r_{int}, r_{conf})$. La fonction $u(r_{conf})$ est la fonction d’Heaviside, si $r_{conf} > 0$, $u(r_{conf}) = 1$ si $r_{conf} < 0$, $u(r_{conf}) = 0$.

Le premier terme de la fonction f_{audioL} est non nul lorsque la confiance est positive. Dans ce cas, le facteur $(-2r_{int} + 1)$ induit une bifurcation dans l’action. Si $r_{int} < 0, 5$ alors ce facteur est négatif. Le niveau sonore converge alors vers 1. Si $r_{int} > 0, 5$, le facteur est

positif et le niveau sonore converge vers 0. Le second terme est non nul lorsque $r_{conf} < 0$, pour cette plage de valeur du niveau sonore converge vers 0,5. La paramètre k_g est fixé à 0,75.

Actions de l'auditeur

Se tourner vers le locuteur L'équation contrôlant l'action "se tourner" est identique à l'action du locuteur (cf équation 12).

Lever le bras L'auditeur bouge le bras lorsqu'il est encore incertain envers les intentions du locuteur. Il permet de pousser le locuteur à clarifier ses intentions. En fonction de la confiance et de l'intention, l'auditeur varie l'amplitude du mouvement du bras. Aussi, lorsque la confiance est nulle, l'amplitude du mouvement est maximale. L'action de bouger le bras est produite lorsque l'auditeur est résolu à prendre le tour. Plus l'auditeur est résolu à prendre le tour, plus il se décidera rapidement à bouger le bras pour un niveau de confiance faible. Le franchissement d'un seuil, a pour effet l'arrêt de l'action.

Le système a donc un attracteur qui varie en fonction de la confiance. L'intention modifie le paramètre de raideur du système. L'équation contrôlant l'action "bouger le bras" est la suivante :

$$\begin{aligned}
 \ddot{m} &= -b\dot{m} - k_g((1 - p_{thre})(1 - n_{thre})f_{arm}(r_{int}, r_{conf}, m) \\
 &\quad + p_{thre}m \\
 &\quad + n_{thre}m) \\
 f_{arm}(m, r_{int}, r_{conf}) &= (m - r_{int}(1 - r_{conf}^2))(0,5r_{int}u(r_{int} - 0,5))
 \end{aligned} \tag{15}$$

La fonction $u(r_{int} - 0,5)$ représente toujours la fonction d'Heaviside. Si $r_{int} < 0,5$ alors $u(r_{int} - 0,5) = 0$, si $r_{int} > 0,5$ alors $u(r_{int} - 0,5) = 1$. Cette fonction permet de s'assurer que l'auditeur n'exhibera pas de mouvement de bras lorsque son intention est inférieure à 0,5. Le facteur $r_{int}(1 - r_{conf}^2)$ contrôle la position de l'attracteur. L'attracteur est maximal pour $r_{conf} = 0$ et $r_{int} = 1$, et vaut 1. Le facteur $(0,5r_{int}u(r_{int} - 0,5))$ contrôle la raideur. Pour une valeur d'intention supérieure à 0,5, cette raideur est proportionnelle à r_{int} . Le paramètre k_g , vaut comme pour les équations précédentes 0,75.

Modulation du niveau sonore Pendant le processus de prise de décision, pour une confiance positive, si l'auditeur est résolu à prendre le tour de parole, il augmentera son niveau sonore. Au contraire, si la confiance est négative, son niveau sonore baissera vers la valeur 0. Le franchissement d'un seuil négatif, a de même pour effet du côté de l'auditeur, de s'arrêter de parler, tandis que le franchissement d'un seuil positif a pour effet de prendre la parole.

Pour $r_{conf} > 0$ si $r_{int} < 0,5$ alors le système présentera un attracteur en $(0;0)$. Au contraire si $r_{int} > 0,5$ alors le système présentera un attracteur en $(0,5;0)$. Si $r_{conf} < 0$ alors l'attracteur de l'agent est en $(0;0)$ La figure ?? montre les différentes possibilités de l'agent lors du processus de décision lorsque la confiance est positive. L'action de

modulation du niveau sonore est représentée par l'équation suivante :

$$\begin{aligned}
\ddot{a}_a &= -b\dot{a}_a \\
&- k_g((1 - p_{thre})(1 - n_{thre})f_{levelA}(a_a, r_{int}, r_{conf}) \\
&+ n_{thre}a_a + p_{thre}(a_a - 0.5)) \\
f_{levelA}(a_a, r_{int}, r_{conf}) &= (2r_{int} - 1)u(r_{conf})r_{conf}(a_a - 0.5)a_a \\
&- (1 - u(r_{conf}))r_{conf}a_a
\end{aligned} \tag{16}$$

La fonction f_{levelA} comporte deux termes. Le premier terme est non nul pour $r_{conf} > 0$. Le système est composé de deux attracteurs, l'un en 0 l'autre en 0,5.

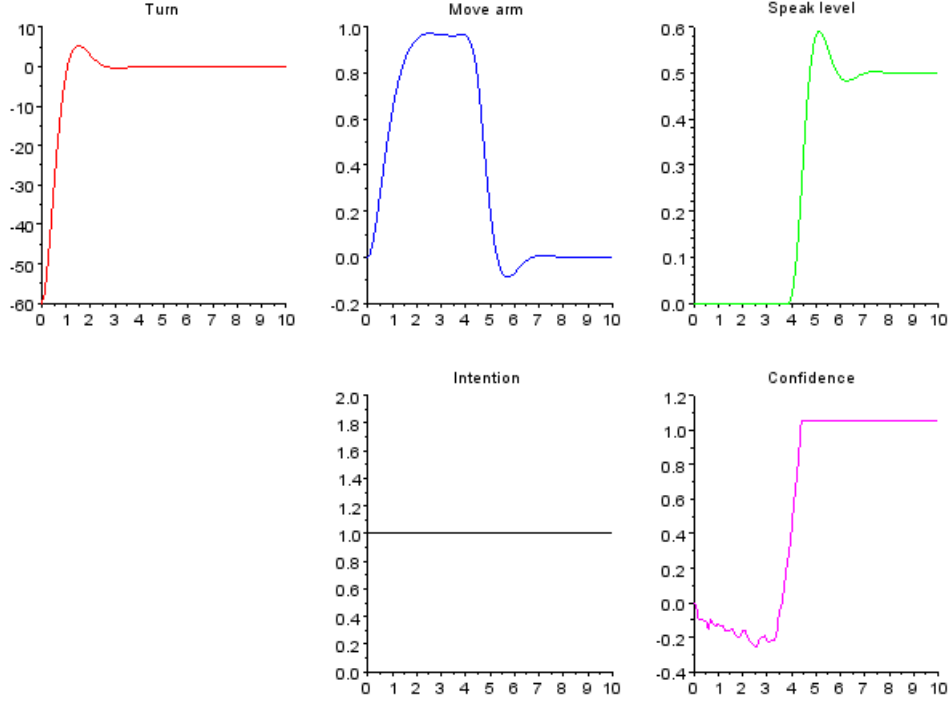
3.3 Couplage des modèles du locuteur et de l'auditeur

Une fois implémenté les modèles de contrôle des actions pour le locuteur et l'auditeur, il reste à coupler les deux implémentations afin d'analyser le comportement de tour de parole, en particulier le moment de changement de tour. Trois scénarios sont décrits dans la suite. Entre chaque scénario, le paramètre d'intention des agents varie, induisant une variation dans le comportement global du tour de parole.

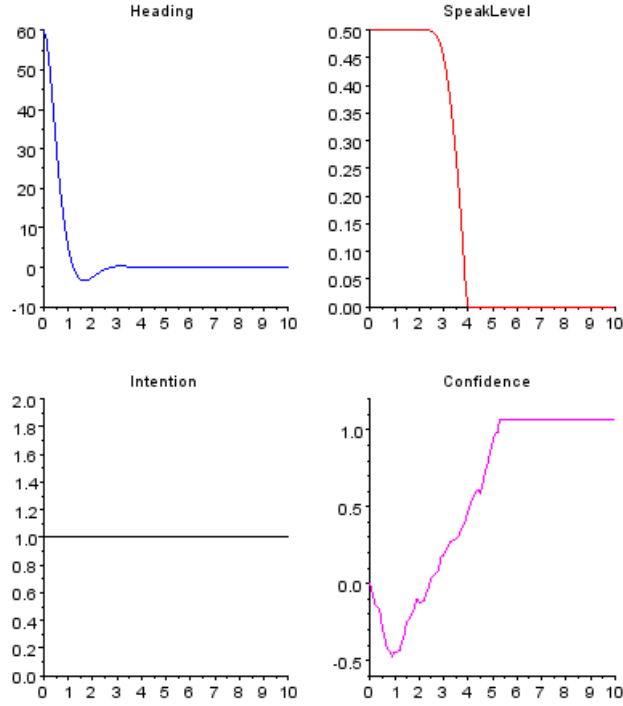
3.3.1 Premier scénario

Le premier scénario, dont l'évolution est décrite sur les figures 10 et 11, montre un locuteur et un auditeur résolus respectivement à donner et à prendre le tour de parole. Le locuteur ainsi que l'auditeur ont aussi leur variable d'intention à 1. Le locuteur et l'auditeur étant incertains sur l'intention de l'autre se tournent d'abord l'un vers l'autre. L'incertitude de l'auditeur le pousse de même à bouger le bras. En réaction à l'auditeur bougeant le bras, la confiance du locuteur diminuant au départ commence à remonter. Plus la confiance du locuteur remonte, plus son niveau sonore diminue. En réaction à la diminution du niveau sonore, la confiance de l'auditeur remonte. Plus la confiance de l'auditeur augmente, moins celui-ci ressent le besoin de lever le bras, celui-ci commence alors à se baisser. Au contraire le niveau sonore de l'auditeur augmente, à partir d'une valeur située entre 0,2 et 0,4, jusqu'à prendre complètement la parole. La troisième courbe montre l'évolution conjointe des niveaux sonores du locuteur et de l'auditeur.

On constate une transition de tour de parole très courte entre la fin de tour du locuteur et le début de tour du locuteur ainsi qu'un léger recouvrement de parole. Cette transition courte est bien due à l'anticipation de la fin du tour de parole, les deux participants se mettent en action avant d'avoir totalement pris la décision concernant la fin du tour de parole (située vers 4 secondes), l'un baisse son niveau sonore, l'autre augmente son niveau sonore.



(a) Évolution temporelle des actions et de la confiance de l'auditeur



(b) Évolution temporelle des actions et de la confiance du locuteur

FIGURE 10 – Évolution temporelle des actions et de la prise de décision du locuteur et de l'auditeur pour le premier scénario

3.3.2 Deuxième scénario

Le second scénario (figures 12 et 13) montre un auditeur réticent à prendre le tour de parole (intention à 0) et un locuteur résolu à lui donner (intention à 1). L'auditeur ne

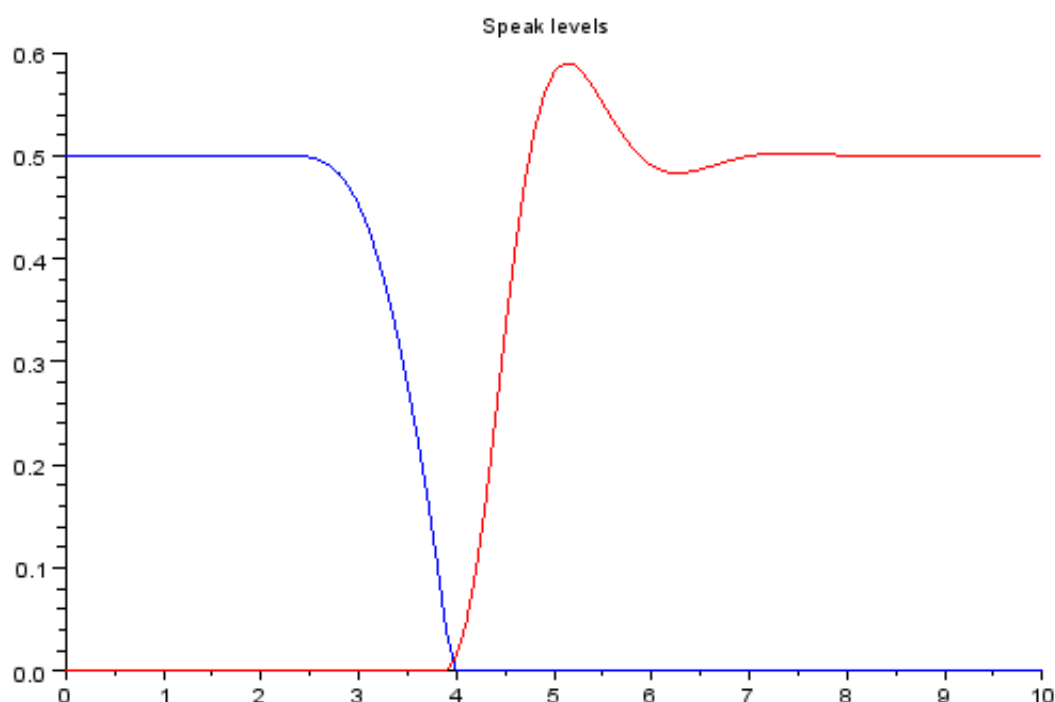
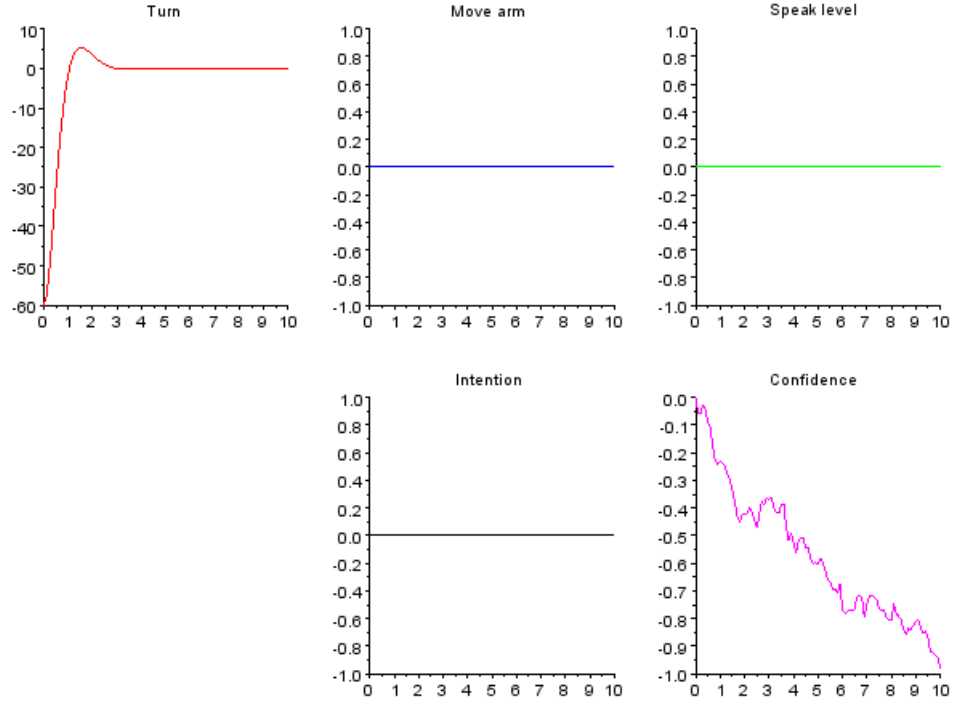
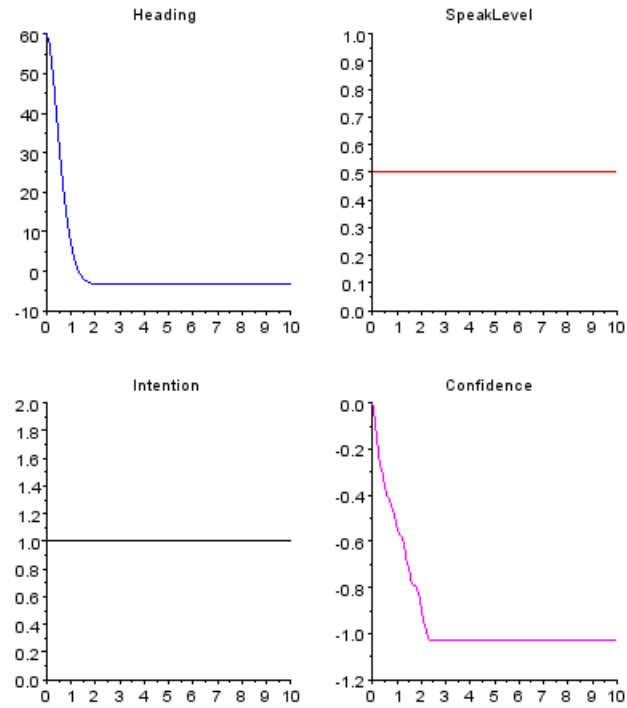


FIGURE 11 – Évolution temporelle des niveaux sonores des deux agents (en bleu le locuteur, en rouge l’auditeur) pour le premier scénario

montre ainsi aucun signal de prise de tour de parole. La valeur de confiance du locuteur étant nulle au départ celui-ci entame l’action de se tourner vers l’auditeur. Néanmoins, ne voyant aucune réaction du côté de l’auditeur, sa valeur de confiance continue à diminuer. Il n’entame pas de baisse de niveau sonore, ce qui a pour effet de baisser le niveau de confiance de l’auditeur. Le locuteur remarque alors rapidement que l’auditeur ne souhaite pas prendre le tour (la valeur de confiance franchit le seuil négatif au bout de 2, 5 secondes). La valeur de confiance de l’auditeur est plus lente à atteindre le seuil négatif, ce qui est dû à la confiance liée au fait que le locuteur s’est tourné vers l’auditeur. Au final, le moment de transition n’a pas eu lieu dans ce cas de figure.



(a) Évolution temporelle des actions et de la confiance de l'auditeur



(b) Évolution temporelle des actions et de la confiance du locuteur

FIGURE 12 – Évolution temporelle des actions et de la prise de décision du locuteur et de l'auditeur pour le second scénario

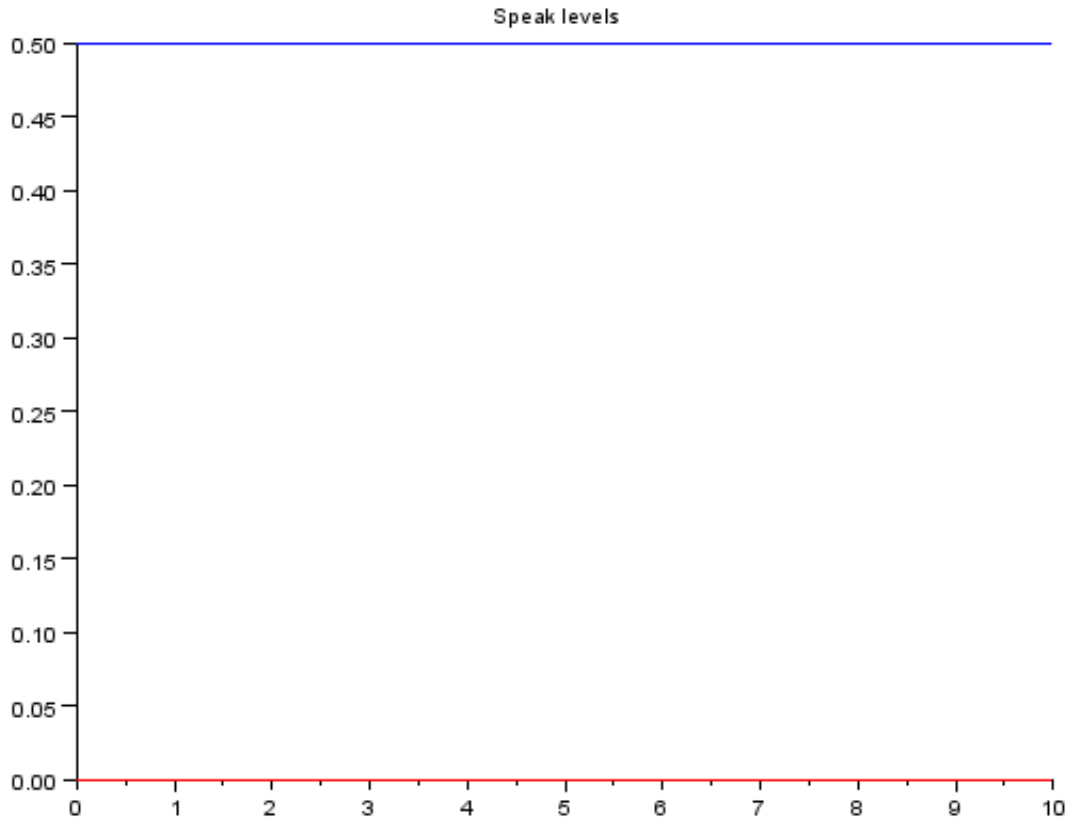
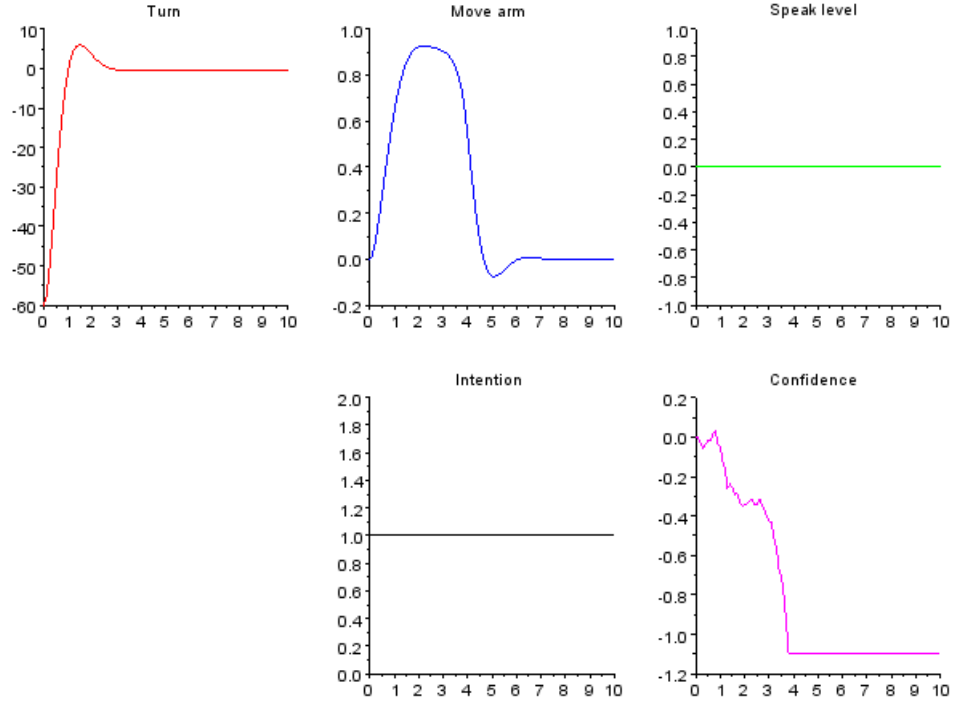


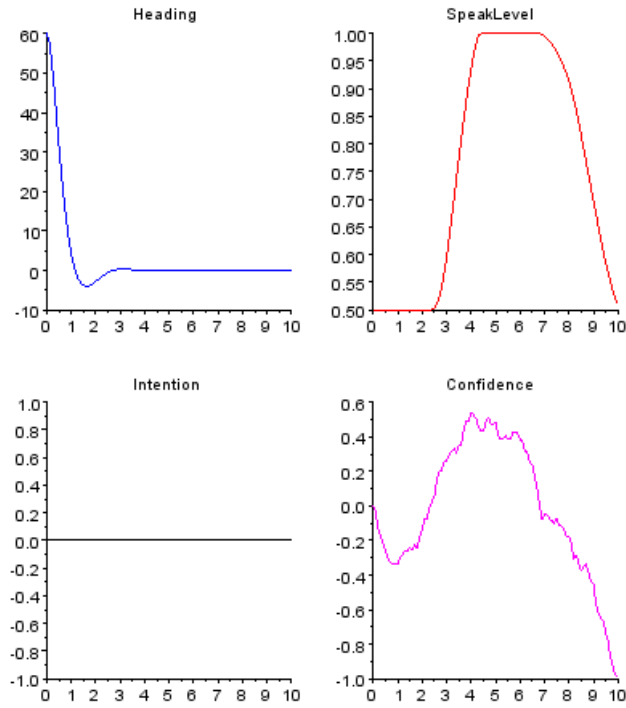
FIGURE 13 – Évolution temporelle des niveaux sonores des deux agents (en bleu le locuteur, en rouge l’auditeur) pour le deuxième scénario

3.3.3 Troisième scénario

Dans le troisième scénario (figures 14 et 15), le locuteur est plutôt réticent à donner le tour de parole (intention à 0,2), l’auditeur est résolu à prendre le tour (intention à 1). Le locuteur ne montre ainsi aucun signal d’abandon de tour de parole. La valeur de confiance du locuteur étant nulle, il se tourne vers l’auditeur, sa valeur d’intention étant néanmoins faible, il se tourne lentement vers ce dernier. La valeur de confiance du locuteur diminue au départ, du fait que l’auditeur n’a pas encore produit de signaux de prise de tour. Au bout d’une seconde, lorsque les amplitudes des actions ”se tourner” et ”bouger le bras”, sont suffisamment importantes, la valeur de confiance du locuteur augmente à nouveau. Lorsque la valeur de confiance devient positive du côté du locuteur, sa réticence le pousse à augmenter son niveau sonore pour dissuader l’auditeur à prendre le tour. Le moment de transition ne se fait alors pas.



(a) Évolution temporelle des actions et de la confiance de l'auditeur



(b) Évolution temporelle des actions et de la confiance du locuteur

FIGURE 14 – Évolution temporelle des actions et de la prise de décision du locuteur et de l'auditeur pour le troisième scénario

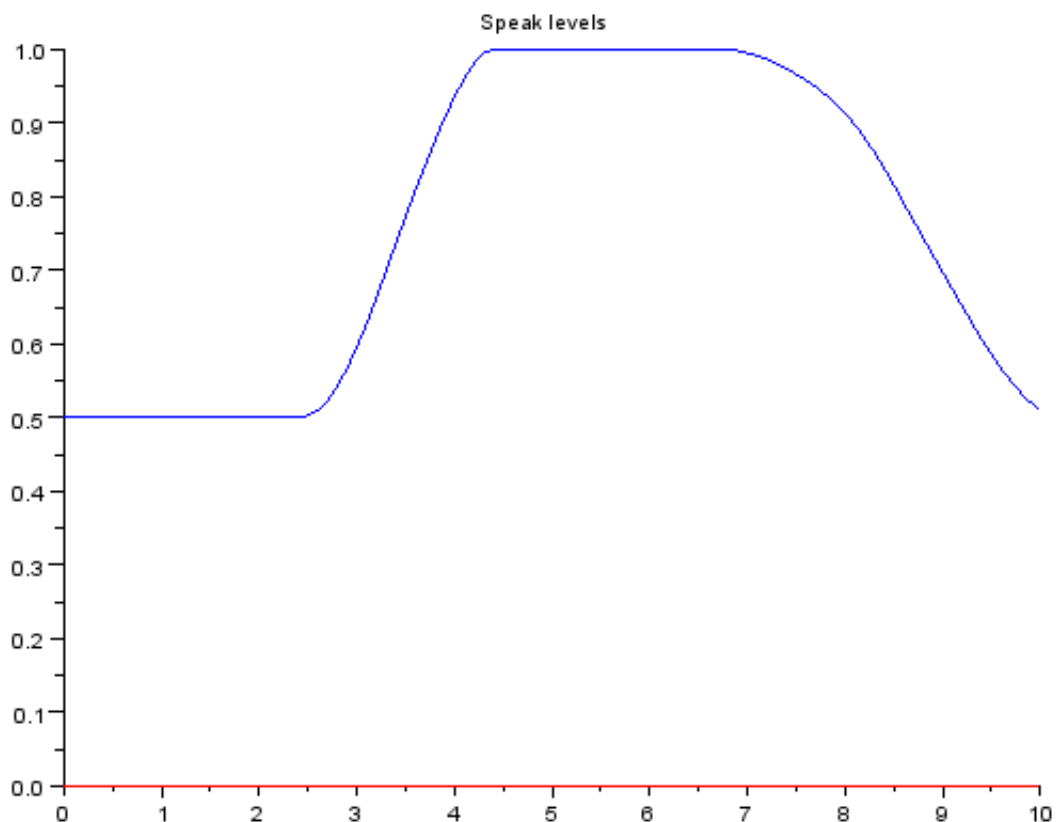


FIGURE 15 – Évolution temporelle des niveaux sonores des deux agents (en bleu le locuteur, en rouge l’auditeur) pour le troisième scénario

3.3.4 Synthèse

Les trois scénarios illustrent trois comportements différents relatifs aux transitions de tour de parole dont le modèle rend compte : selon les cas, la transition peut se faire ou ne pas se faire, et les actions du locuteur et de l’auditeur varient en fonction de leur intention propre et de l’intention de l’autre. Le modèle montre donc une variabilité dans les comportements globaux de transition. Le couplage entre agents y est de même illustré. Dans les premier et troisième scénarios, l’auditeur a une variable d’intention identique. Néanmoins, la variation de l’intention du locuteur produit des courbes d’action de l’auditeur différentes.

4 Architecture d’agent pour l’implémentation du modèle

Le modèle conceptuel a été élaboré, il reste maintenant à l’implémenter dans une architecture d’agent. De l’analyse des architectures d’agents conversationnels existantes de la section 2.2 nous avons retenu qu’Ymir semblait la plus adaptée pour l’implémentation d’un modèle du tour de parole. Néanmoins certaines caractéristiques d’Ymir ne permettent

pas l'intégration de modèles continus, d'où la nécessité de la création d'une architecture inspirée d'Ymir pour l'implémentation du modèle.

4.1 Principes d'Ymir (Thórisson, 1999)

Certaines caractéristiques d'Ymir ont été vues dans la section 2.2.2, dont la répartition des modules en plusieurs couches d'exécution, la distinction entre percepteurs et décideurs, la sélection d'action à partir d'une décision. Revenons maintenant plus en détail sur les différents éléments de l'architecture.

Il existe trois couches d'exécution dans Ymir, la couche réactive comprend les modules s'exécutant à la fréquence la plus élevée. La couche de supervision contient les modules dédiés au contrôle de la dynamique du dialogue (changements de tour, signaux de rétroaction), indépendamment du contenu du dialogue. Enfin les modules traitant du contenu sont regroupés dans la couche de gestion du contenu s'exécutant à la fréquence la moins élevée.

Les modules peuvent être actifs ou inactifs dans l'architecture. Ils sont de plus associé à un état (vrai ou faux) représentant, pour les percepteurs, la valeur de vérité de la donnée qu'ils publient, pour les décideurs l'état de la décision auxquels ils sont associés (déclenchée ou non).

Dans Ymir, les percepteurs sont classés en deux catégories, les percepteurs unimodaux et les percepteurs multimodaux. Les percepteurs unimodaux récupèrent en entrée une donnée brute provenant des capteurs de l'agent. Les percepteurs multimodaux récupèrent, eux, plusieurs données en entrée. Ces données peuvent être des données brutes provenant de capteurs ou des données provenant d'autres percepteurs. Le mécanisme des percepteurs multimodaux permet une mise en cascade des percepteurs, donc une chaîne de traitement des données provenant de l'environnement, amenant à la création de données perceptuelles de haut niveau, fondées sur l'intégration de données provenant de plusieurs canaux de communication.

Les décideurs sont de même classés en deux catégories : les décideurs déclarés et les décideurs non déclarés. Les décideurs non déclarés sont chargés de prendre une décision sur l'état cognitif de l'agent. Les décideurs déclarés contrôlent le comportement extérieur de l'agent, ils contrôlent le déclenchement d'une décision amenant au déclenchement d'une ou plusieurs actions de la part de l'agent.

Les modules communiquent entre eux par l'intermédiaire de plusieurs tableaux noirs (*blackboard*). Ces *blackboard* constituent des bases de données dans lesquelles les données publiées par les percepteurs sont stockées. Récupérer des données en entrée consiste pour les modules à effectuer une requête au *blackboard* contenant la donnée à aller chercher. Deux *blackboard* sont utilisés dans Ymir, un pour la communication entre les modules de la couche réactive et les modules de la couche de supervision, l'autre pour la communication entre les modules de la couche de supervision et les modules de la couche de gestion du contenu.

Une fois la décision prise par les décideurs déclarés d'exécuter une action, une requête d'action est envoyée à un planificateur d'action. Les tâches du planificateur à la réception de la requête sont multiples. Le planificateur attribue d'abord une priorité à l'action reçue. Cette priorité est liée à la couche à laquelle appartient le décideur. Plus la couche s'exécute à une fréquence élevée plus l'action est prioritaire. Le planificateur d'action est ensuite chargé d'attribuer une suite de "morphologies", c'est à dire d'actions de bas niveau réalisées dans l'environnement par l'agent. Il recherche alors dans un lexique com-

portemental la bonne action à réaliser. Le lexique comportemental est formé d'un arbre comportant l'ensemble des actions disponibles à l'agent. Deux types d'actions y sont stockées. Les comportements *act* correspondent à des comportements où plusieurs alternatives sont possibles. Ces alternatives sont choisies par le planificateur en fonction des ressources de l'agent disponibles. Chaque alternative comporte elle-même une suite d'actions à réaliser lorsque cette alternative est choisie. Les feuilles de l'arbre représentent au contraire des comportements n'ayant pas d'alternatives, ces comportements sont appelés des *motor schemes*. Ces *motor schemes* sont associés à une suite d'actions physiques exécutées par l'agent. La tâche du planificateur d'action est alors à partir d'un nœud représentant le comportement à exécuter, de parcourir l'arbre jusqu'à trouver un ou plusieurs *motor schemes* à exécuter par l'agent. Une fois le ou les *motor schemes* sélectionnés, le planificateur envoie les requêtes d'actions à exécuter aux actionneurs chargés d'exécuter le comportement.

4.2 Solution pour l'intégration du modèle

La figure 16 montre le schéma général de l'architecture élaborée. Les éléments avec un fond rouge sont des éléments n'existant pas dans Ymir, et créés pour les besoins de notre architecture. Les éléments entourés de rouge sont des éléments existants dans Ymir mais modifiés dans notre architecture, les éléments sur fond blanc sont des éléments extérieurs à l'architecture.

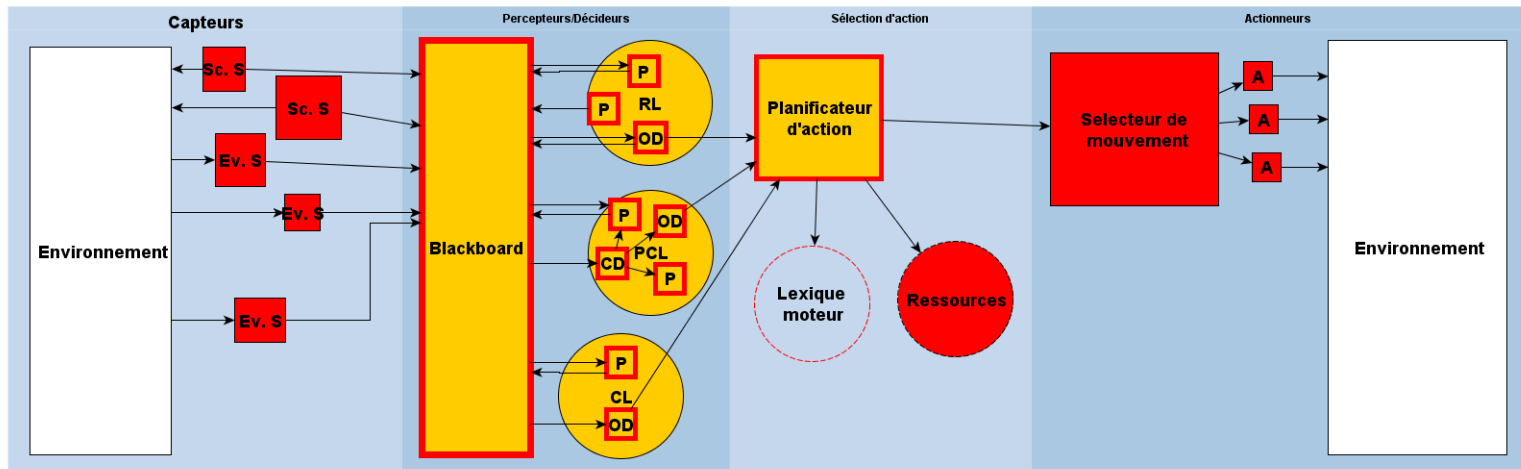


FIGURE 16 – Schéma général de l'architecture réalisée. Sc. S fait référence aux capteurs par scrutation, Ev. S fait référence aux capteurs événementiels, P fait référence aux percepteurs, OD aux décideurs déclarés, CD aux décideurs non déclarés, RL à la couche réactive, PCL à la couche de supervision, CL à la couche de contenu et A aux actionneurs de l'agent.

4.2.1 Capteurs

Thórisson (1999) ne décrivant pas comment les capteurs récupèrent des données de l'environnement, cette partie de l'architecture a entièrement été définie. Les capteurs sont liés à l'outil utilisé pour l'implémentation de l'environnement virtuel. Deux types de capteurs sont possibles dans l'architecture, les capteurs par scrutation, ayant une fréquence

d'envoi des données, et les capteurs par événement, envoyant une donnée sur occurrence d'un événement. Une fois la donnée récupérée celle-ci est mise en forme et envoyée aux percepteurs et décideurs. La fréquence d'exécution des capteurs peut être différente de la fréquence d'exécution des percepteurs et décideurs de l'architecture.

4.2.2 Perception, décision et sélection d'action

L'architecture réalisée a introduit un certain nombre de modifications au niveau des couches d'exécution. Il fallait tout d'abord changer la nature des données transmises entre les modules. Deux types de variables ont été définies :

- Les variables continues représentent des données prenant leurs valeurs dans l'ensemble des réels. Elles sont associées à une dimension, et à un domaine de validité ;
- Les variables événementielles sont liées aux données perceptuelles de nature événementielle. L'événement "un agent a parlé" est un exemple d'une donnée de type événementielle. Chaque variable événementielle est liée à une valeur de confiance sur l'occurrence réelle de l'événement et éventuellement de champs optionnels auxquels est aussi associé une valeur de confiance. Ces données événementielles ont été créées par souci de compatibilité avec certains outils utilisés, notamment la reconnaissance vocale, n'envoyant que des informations de type événementiel.

Dans Ymir, les modules ont un état, l'action associée au module n'est effectuée que lorsque l'état du module vaut *true*. Pour un percepateur, l'action réalisée est de poster son changement d'état dans le *blackboard*. Le positionnement pris dans notre architecture, est que les données continues sont publiées à chaque exécution du module, même si cette donnée ne varie pas. Les percepteurs traitant des données événementielles vérifient en premier lieu, par l'intermédiaire de l'horodatage liée à la donnée, si celle-ci a été mise à jour depuis la dernière fois que celui-ci a consulté la donnée. Il ne publie alors la donnée, que si celle-ci a été mise à jour.

Ymir utilise deux *blackboards*, dans l'architecture réalisée un seul *blackboard* est utilisé. Il est commun à tous les percepteurs et décideurs de l'architecture peu importe leur couche d'exécution. Les données provenant des capteurs sont de même insérées dans ce *blackboard*.

De même que pour la perception, la décision doit être modifiée pour intégrer un contrôle en continu de l'action réalisée. Pour cela un décideur déclaré a deux types de décisions à prendre. Une décision concernant le déclenchement de l'action, et, une fois l'action déclenchée le contrôle de l'action. Lorsque l'action est déclenchée, les décideurs déclarés envoient ainsi à chaque exécution l'action à réaliser. Afin de rendre l'action modulable des paramètres ont été ajoutés à la requête d'action envoyée au planificateur. Le mécanisme de planification a de même été modifié afin d'assurer le contrôle en continu de l'action. Le planificateur d'action garde dans une liste les actions actuellement exécutées dans l'environnement. Lors de la réception d'une requête d'action, le planificateur vérifie en premier lieu si l'action n'est pas actuellement en train d'être exécutée dans l'environnement. Dans ce cas, il n'arrête pas l'action en cours mais envoie aux actionneurs la modification à effectuer sur l'action.

4.2.3 Actionneurs

Comme pour les capteurs, le mécanisme de réception et de déclenchement de l'action au niveau des actionneurs n'est pas précisé par Thórisson. Les actionneurs sont liés

aux outils utilisés pour l'implémentation de l'environnement. Ils se chargent de réaliser l'action physiquement dans l'environnement. Un planificateur de mouvement est chargé de réceptionner les requêtes de mouvement envoyés par le planificateur d'action, puis d'activer le comportement si le mouvement n'était pas actuellement déclenché dans l'environnement, de modifier ses paramètres si celui-ci était déjà exécuté dans l'environnement ou de l'annuler.

4.3 Place du modèle conceptuel dans l'architecture

Il reste maintenant à définir où les différents éléments du modèle seront implémentés. Les équations de contrôle de l'action sont liées à des actions externes de l'agent, elles seront donc implémentées en utilisant les décideurs déclarés. L'équation de prise de décision représente, elle, un jugement sur une donnée perceptuelle de l'environnement, ces types d'équations seront donc implémentés dans un percepteur.

4.4 Implémentation de l'architecture

L'architecture a été implémentée en utilisant le langage C# et le framework .NET de Microsoft.

Le travail réalisé s'inscrit dans le cadre du projet CORVETTE (CollaboRative Virtual Environment Technical Training and Experiment), dont le but est de créer et évaluer sur des utilisateurs des environnements virtuels d'apprentissage humain. Dans le cadre de ce projet, une scène présentée sur la figure 17 a été réalisée sous Unity3d. Dans cette scène, l'utilisateur collabore avec deux agents (au second plan sur l'image) pour monter un meuble. Le troisième agent (au premier plan) supervise l'opération. Il est chargé d'expliquer la tâche à l'utilisateur au début du scénario. La collaboration entre l'utilisateur et les agents nécessite une coordination entre eux, le mode d'interaction principal permettant cette coordination étant la parole.

Une première implémentation d'un comportement de tour de parole a été réalisée, sans utiliser le modèle élaboré, en utilisant cette scène.

Dans l'objectif de l'élaboration de cette première version, certains modules ont été implémentés.

Au niveau des capteurs, des capteurs de distance et d'angle ont été implémentés, ils calculent et envoient aux percepteurs, les angles relatifs et les distances de tous les agents de l'environnement. D'autres capteurs sont liés à la reconnaissance vocale (utilisant la Speech API de Microsoft), et plus généralement aux événements vocaux de l'environnement de l'agent, provenant soit de l'utilisateur soit d'autres agents. Ils récupèrent des événements tels que "un agent a commencé à parler" et des informations relatives à des phrases reconnues. Ainsi le capteur "AudioLevel" récupère le niveau sonore de l'environnement de l'agent.

Les percepteurs unimodaux de l'architecture servent à publier dans le blackboard les données des capteurs. Ces données ne subissent pas de traitement de la part de ces percepteurs, ceux-ci permettent la publication d'une donnée à la fréquence de la couche d'exécution des percepteurs unimodaux (la couche réactive ayant une période d'exécution de 100 millisecondes). Plusieurs percepteurs multimodaux sont présents dans l'architecture. Le percepteur "Facing" correspond au signal s_f des modules de prise de décision. "LocutorGiveTurnTo" et "OtherTakesTurn" déterminent respectivement si le locuteur donne le tour à un agent de l'environnement si l'agent est un auditeur, et si quelqu'un est en train de prendre le tour dans l'environnement. Deux actions sont possibles pour



FIGURE 17 – Scène du projet CORVETTE utilisée pour l’implémentation du modèle

l’agent, il peut se tourner à une vitesse définie, et parler. L’actionneur "Speak" utilise la synthèse vocale de Microsoft, et permet la modulation du niveau sonore de l’agent. Le décideur déclaré TurnToLocutor, permet le contrôle de l’action "se tourner" afin que l’auditeur se tourne vers le locuteur, tandis que le décideur "Talk" permet à l’agent d’envoyer une phrase à l’actionneur. Un certain nombre de ces modules seront utilisés pour l’implémentation de notre modèle. Les capteurs d’angle et de niveau sonore permettent la récupération de tous les signaux nécessaires à l’implémentation du modèle. Le perceuteur multimodal "FacingMe" permet de même de récupérer l’information, "fait face à l’agent", nécessaire pour l’auditeur et le locuteur dans leur prise de décision.

Les équations de prise de décision seront implémentés dans les perceuteurs "OtherTakesTurn" et "LocutorGivesTurn".

Il manque néanmoins l’actionneur "RaiseArm" permettant à un agent de lever son bras. Le décideur déclaré "MoveArm" chargé de contrôler l’amplitude et la vitesse du lever de bras de l’agent devra de même être implémenté.

5 Discussion

5.1 Vérification des hypothèses

Le modèle comportemental de la gestion du tour de parole que nous avons proposé vérifie bien l’hypothèse énoncée dans la section 2.1.3 : un moment de transition peut émerger uniquement de l’interaction entre le locuteur et l’auditeur, sans planification du prochain locuteur, ni prédiction du moment où la transition apparaît. Les décisions des participants sont locales temporellement, un agent ne se soucie que de l’occurrence du prochain moment de transition, et locales à chaque agent, chaque interlocuteur prend ses propres décisions concernant le comportement de l’autre et ajuste ses propres variables d’action. Le TRP n’est pas prédit par les interlocuteurs, à aucun endroit dans les

équations n'a été spécifié de variable indiquant le moment où le tour de parole allait se produire. Ce moment est plutôt anticipé, les participants se mettent en action avant que le tour de parole ne soit réellement transmis. En suivant ces critères, le modèle est capable de faire émerger des moments de transition. On obtient de plus, comme observé par Goodwin (1981) et Sacks et al. (1974) des transmissions fluides, sans moment de silence trop important ni recouvrement. Le modèle permet de même de produire une variété de comportements de tour de parole, illustré par les trois scénarios de la section 3.3.

5.2 Dynamique comportementale pour la prise de décision

Le modèle choisi pour le tour de parole, couple deux modèles de psychologie cognitive, le modèle d'accumulation-diffusion (Bogacz et al., 2006) et le modèle de dynamique comportementale (Warren, 2006). Certains modèles utilisant uniquement la dynamique comportementale ont permis de modéliser des tâches où un agent devait choisir entre deux actions mutuellement exclusives (Frank et al., 2009).

Le modèle de Frank et al. (2009), et les modèles de dynamique comportementale en général, décrivent les actions observés en fonction des informations reçus de l'environnement. Ces modèles ne décrivent pas la manière dont la prise de décision est effectuée. L'application de ce type de modèle aurait permis la mise en avant des propriétés d'émergence et d'auto-organisation du tour de parole. Il n'aurait néanmoins pas permis de rendre compte de l'hypothèse de l'accumulation de signal de la part de l'agent. Le DDM introduit un effet mémoire de la décision prise auparavant par l'agent. Les signaux accumulés précédemment influent sur la décision de l'agent ensuite. Si un auditeur accumule d'abord des signaux défavorables concernant le fait que le locuteur lui donne le tour, il mettra plus de temps à prendre la décision inverse s'il reçoit des signaux favorables concernant l'intention de donner le tour, que s'il n'avait reçu que des signaux favorables. Le moment de transition observé ne sera alors pas le même.

5.3 Interaction entre plus de deux agents

Le modèle élaboré permet de rendre compte de la variabilité des transitions possibles de tour de parole dans le cadre d'une interaction entre deux agent, mais ne permet pas de rendre compte de transitions de tour de parole entre plusieurs agents.

Plusieurs solutions peuvent être mises en place pour une généralisation à n agents. La dynamique comportementale s'applique à un agent en relation avec son environnement. Les interlocuteurs sont donc dans ce cadre intégrés dans un seul système, l'environnement. Dans la logique de la dynamique comportementale, du côté de l'agent, les équations de prise de décision et de contrôle de l'action ne seraient donc pas spécifiquement dirigés vers un agent en particulier mais vers l'environnement de l'agent dans son ensemble. Pour le locuteur, la décision ne porterait alors pas sur "l'autre prend le tour" mais sur "quelqu'un prend le tour". Plusieurs alternatives concernant l'intention du locuteur sont possibles dans un environnement comportant plusieurs agents, il peut, soit donner le tour à un auditeur précis, soit finir son tour sans avoir désigné quelqu'un (Sacks et al., 1974). Le DDM étant un modèle de prise de décision entre deux alternatives, plusieurs équations de prise de décision seraient nécessaires du côté des auditeurs. Le problème réside alors, pour le locuteur dans l'identification de l'agent souhaitant prendre le tour. Si un auditeur souhaite prendre le tour de parole, le locuteur peut produire des signaux donnant à cet auditeur le tour de parole. Il faut alors qu'il soit capable, à partir de l'équation de prise

de décision, de produire des signaux spécifiques envers cet auditeur. Plutôt qu'utiliser une équation de prise de décision pour le locuteur, la solution serait alors d'en avoir une par agent. Pour l'auditeur, le locuteur changeant au cours d'une conversation, un auditeur doit être capable d'identifier qui est le locuteur courant, afin de récupérer les bons signaux de l'environnement.

5.4 Signaux interprétés

L'application du modèle décrite dans la section 3.2 a été avant tout réalisée dans le but de mettre en avant le fait qu'un modèle fondé sur le couplage entre deux agents permet de voir émerger un comportement global de tour de parole. Les signaux utilisés sont très limités. Cette limitation provient du fait que celui-ci doit fonctionner dans l'architecture mise en place pour la validation du modèle. Une prise en compte des travaux sur le tour de parole, notamment ceux de Gravano and Hirschberg (2011) permettrait d'avoir un modèle plus crédible de la décision et de la mise en action des agents. Néanmoins cette prise en compte de signaux nécessite un grand nombre de capteurs et d'actionneurs du point de vue de l'architecture, incluant notamment la capacité pour un agent de varier son intonation en continu, et une analyse sémantique de ce qui est dit, ce qui n'est actuellement pas implémenté dans l'architecture.

5.5 Validation du modèle

La validation entre deux agents dans un environnement virtuel est en cours. Une première implémentation d'un modèle de tour de parole a été effectuée dans l'architecture. Ce modèle se fonde sur une machine à états décrivant les différents états du dialogue, notamment l'état "locuteur" et "auditeur", ainsi que les règles amenant à une transition entre ces états. L'observation réalisée concernant le tour de parole, est avant tout que les transitions entre les tours sont brusques, avec un moment de silence entre les tours importants par rapport aux temps de transitions observés dans les conversations réelles. Ce modèle est de plus limité dans la variabilité des comportements de transition possibles. Les résultats que l'on cherche à observer, sont au contraire des temps de transitions courts voire inexistantes entre les tours, et des comportements de transition variés.

Pour des raisons de temps, la validation entre un agent et un utilisateur n'est pas envisagé d'ici la fin du stage. La validation par rapport à un utilisateur est plus compliquée due aux limitations liées aux capteurs de l'environnement, comme énoncé dans la partie précédente. Observer les actions d'un utilisateur dans le cadre du tour de parole implique en effet, la récupération de signaux tels que l'intonation à partir de la reconnaissance vocale, ce qui n'est pas réalisable avec la reconnaissance vocale actuellement utilisée, ou encore des signaux relatifs à la gestuelle, ce qui implique l'utilisation de reconnaissance gestuelle. Une validation avec un utilisateur reposerait sur le caractère spontané de l'interaction, et sur l'engagement de l'utilisateur dans l'interaction avec l'agent. Si l'on observe ces deux caractéristiques, alors nous aurons démontré que le modèle est crédible par rapport aux conversations réelles.

6 Conclusion

Nous avons proposé un modèle pour la gestion du tour de parole entre humains virtuels dans le cadre d'une interaction dialogique entre deux participants. Le modèle, contraire-

ment à beaucoup de modèles d’agents existants, est continu. Il est fondé sur une accumulation d’indices provenant de l’environnement, et une modulation des signaux de l’agent en fonction de la quantité d’indices intégrés par l’agent sur les intentions de l’autre. Le modèle utilise ainsi deux modèles issus de la psychologie cognitive, le DDM (Ratcliff, 1978) et la dynamique comportementale de Warren (2006). La simulation de plusieurs scénarios de transition de tour de parole a montré l’émergence de transitions de tour de parole, sans que celui-ci soit prédit ni planifié par les agents. Une variabilité dans les transitions possibles de tour de parole a de plus été observée. Le modèle n’a pas encore été appliqué au contrôle d’un humain virtuel. La prochaine étape sera la validation du modèle dans le cadre d’une simulation de conversation entre deux agents virtuels autonomes. En vue de l’implémentation du modèle pour le contrôle d’un humain virtuel, une architecture d’agent inspirée d’Ymir (Thórisson, 2002) a été créée.

Remerciements

Ce projet a été financé par l’ANR, Agence Nationale de la Recherche dans le cadre du projet CORVETTE (ANR-10-CORD-012).

Références

- Allbeck, J. M. and Badler, N. I. (2001). Towards behavioral consistency in animated agents. In *Proceedings of DEFORM/AVATARS’00*, pages 191–205.
- Bogacz, R., Brown, E., Moehlis, J., Holmes, P., and Cohen, J. D. (2006). The physics of optimal decision making : a formal analysis of models of performance in two-alternative forced choice tasks. *American Psychological Association*, 113(4) :700–765.
- Duncan, S. (1972). Some signals and rules for taking speaking turns in conversations. *Journal of Personality and Social Psychology*, 23(2) :283–292.
- Edward, L., Lourdeaux, D., Barthès, J.-P. A., Lenne, D., and Burkhardt, J.-M. (2008). Modelling autonomous virtual agent behaviours in a virtual environment for risk. *IJVR*, 7(3) :13–22.
- Fajen, B. R. and Warren, W. H. (2003). Behavioral dynamics of steering, obstacle avoidance, and route selection. *Journal of Experimental Psychology*, 29(2) :343–362.
- Frank, T. D., Richardson, M. J., Lopresti-Goodman, S. M., and Turvey, M. (2009). Order parameter dynamics of body-scaled hysteresis and mode transitions in grasping behavior. *Journal of Biological Physics*, 35 :127–147.
- Gerbaud, S., Mollet, N., and Arnaldi, B. (2007). Virtual environments for training : From individual learning to collaboration with humanoids. In *Technologies for E-Learning and Digital Entertainment*, volume 4469 of *Lecture Notes in Computer Science*, pages 116–127. Springer-Verlag.
- Goldstein, B. (1981). The ecology of J.J Gibson’s perception. *Leonardo*, 14(3) :191–195.
- Gravano, A. and Hirschberg, J. (2011). Turn-taking cues in task-oriented dialogue. In *Computer Speech and Language* 25, pages 601–634. Elsevier.
- Kelso, J. S. (2009). Coordination dynamics. *Encyclopedia of complexity and systems science*, 1 :1537–1564.

- Loomis, J. M. and Beall, A. C. (2004). Model based control of perception/action. In *Optic flow and beyond*, pages 421–441.
- Mancini, M., Castellano, G., Bevacqua, E., and Peters, C. (2007). Copying behaviour of expressive motion. In Gagalowicz, A. and Philips, W., editors, *MIRAGE*, volume 4418 of *LNCS*, pages 180–191.
- Nakano, Y. I., Reinstein, G., Stocky, T., and Cassell, J. (2003). Towards a model of face-to-face grounding. In *Proceedings of the Annual Meeting of the Association for Computational Linguistics*, pages 553–561.
- Ratcliff, R. (1978). A theory of memory retrieval. *Psychological Review*, 85(2) :59–107.
- Ratcliff, R. (1980). A note on modeling accumulation of information when the rate of accumulation change over time. *Journal of Mathematical Psychology*, 21 :178–184.
- Rickel, J. and Johnson, W. L. (1997). Steve : An animated pedagogical agent for procedural training in virtual environments (extended abstract). *SIGART Bulletin*, 8 :16–21.
- Sacks, H., Schegloff, E. A., and Jefferson, G. A. (1974). A simplest systematics for the organisation of turn-taking in conversation. *Language*, 50 :696–735.
- Swartout, W. R., Gratch, J., Jr., R. W. H., Hovy, E. H., Marsella, S., Rickel, J., and Traum, D. R. (2006). Toward virtual humans. *AI Magazine*, 27(2) :96–108.
- ter Maat, M., Truong, K., and Heylen, D. (2010). How turn-taking strategies influence users’ impressions of an agent. In Allbeck, J. et al., editor, *IVA*, volume 6356 of *LNCS*, pages 441–453.
- Thórisson, K. R. (2002). Natural turn-taking needs no manual : computational theory and model, from perception to action. In *Multimodality in Language and Speech Systems*, pages 173–207.
- Thórisson, K. R. (1999). A mind model for multimodal communicative creatures & humanoids. *INTERNATIONAL JOURNAL OF APPLIED ARTIFICIAL INTELLIGENCE*, 13(4) :519–538.
- Warren, W. H. (2006). The dynamics of perception and action. *Psychological Review*, 113(2) :358–389.