



PERSISTENCE, INTENTION, AND COMMITMENT

Technical Note 415

February 19, 1987

By: Philip R. Cohen, Sr. Computer Scientist
Artificial Intelligence Center
Computer and Information Sciences Division

Center for the Study of Language and Information
Stanford University

Hector J. Levesque*
Department of Computer Science
University of Toronto

APPROVED FOR PUBLIC RELEASE: DISTRIBUTION UNLIMITED

This research was made possible in part by a gift from the Systems Development Foundation, in part by support from the Natural Sciences and Engineering Research Council of Canada, and in part by support from the Defense Advanced Research Projects Agency under Contract N00039-84-K-0078 with the Naval Electronic Systems Command. The views and conclusions contained in this document are those of the authors and should not be interpreted as representative of the official policies, either expressed or implied, of the Defense Advanced Research Projects Agency, the United States Government, or the Canadian Government. Parts of this paper have appeared in *Reasoning about Actions and Plans: Proceedings of the 1986 Workshop at Timberline Lodge*, Morgan Kaufman Publishers, Inc. Those parts were reprinted there with the permission of the MIT Press.

*Fellow of the Canadian Institute for Advanced Research.

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 19 FEB 1987		2. REPORT TYPE		3. DATES COVERED 00-02-1987 to 00-02-1987	
4. TITLE AND SUBTITLE Persistence, Intention, and Commitment				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) SRI International,333 Ravenswood Avenue,Menlo Park,CA,94025				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT	18. NUMBER OF PAGES 46	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Abstract

This paper explores principles governing the rational balance among an agent's beliefs, goals, actions, and intentions. Such principles provide specifications for artificial agents, and approximate a theory of human action (as philosophers use the term). By making explicit the conditions under which an agent can drop his goals, i.e., by specifying how the agent is *committed* to his goals, the formalism captures a number of important properties of intention. Specifically, the formalism provides analyses for Bratman's three characteristic functional roles played by intentions [7,8], and shows how agents can avoid intending all the foreseen side-effects of what they actually intend. Finally, the analysis shows how intentions can be adopted relative to a background of relevant beliefs and other intentions or goals. By relativizing one agent's intentions in terms of beliefs about another agent's intentions (or beliefs), we derive a preliminary account of interpersonal commitments.

Contents

1	Introduction	1
1.1	Rational Balance	1
1.2	Artificial Intelligence Research on Planning Systems	2
1.3	Philosophical Theories of Intention	3
1.4	Desiderata for a Theory of Intention	4
1.5	The "Little Nell" Problem: Not giving up too soon.	5
1.6	Intention as a Composite Concept	6
2	Methodology	7
2.1	Strategy: A Tiered Formalism	7
2.2	Successive Approximations	7
2.3	Idealizations	7
2.4	Map of the Paper	8
3	Elements of A Formal Theory of Rational Action	8
3.1	Syntax	9
3.2	Semantics	10
3.2.1	Model Theory	10
3.2.2	Definitions	11
3.2.3	Satisfaction	11
3.2.4	Abbreviations	12
3.2.5	Constraints on the Model	13
3.3	Properties of the Model	13
3.3.1	Events and Action Expressions	13
3.3.2	Properties of Acts/Events under HAPPENS	14
3.3.3	Temporal Modalities: DONE, $\Diamond$, and $\Box$	15
3.3.4	Constraining Courses of Events	16
3.4	The Attitudes	16
3.5	Belief	16
3.6	Goals	17
3.6.1	Achievement Goals	18
3.6.2	No Persistence/Deferral Forever	18
3.6.3	Goals and their Consequences	19
4	Persistent goals	20
4.1	The Logic of P-GOAL	21
4.1.1	Conjunction, Disjunction, and Negation	21
4.1.2	No Consequential Closure of P-GOAL	21
4.2	Persistent Goals Constrain Future Beliefs and Actions	22
4.2.1	Acting on Persistent Goals	23

5	Intention as a Kind of Persistent Goal	24
5.1	Belief and Action	24
5.2	INTEND ₁	27
5.2.1	Intending actions	27
5.3	INTEND ₂	29
5.4	Meeting the Desiderata	30
5.4.1	Solving the "Little Nell" Problem: When not to give up goals	32
6	An End to Fanaticism	34
7	Conclusion	36
8	Acknowledgments	36

1 Introduction

Sometime in the not-so-distant future, you are having trouble with your new household robot.¹ You say "Willie, bring me a beer." The robot replies "OK, boss." Twenty minutes later, you screech "Willie, why didn't you bring that beer?" It answers "Well, I intended to get you the beer, but I decided to do something else." Miffed, you send the wise guy back to the manufacturer, complaining about a lack of commitment. After retrofitting, Willie is returned, marked "Model C: The Committed Assistant". Again, you ask Willie to bring a beer. Again, it accedes, replying "Sure thing". Then you ask: "What kind did you buy?" It answers: "Genessee". You say "Never mind." One minute later, Willie trundles over with a Genessee in its gripper. This time, you angrily return Willie for overcommitment. After still more tinkering, the manufacturer sends Willie back, promising no more problems with its commitments. So, being a somewhat trusting consumer, you accept the rascal back into your household, but as a test, you ask it to bring you your last beer. Willie again accedes, saying "Yes, Sir." (Its attitude problem seems to have been fixed). The robot gets the beer and starts toward you. As it approaches, it lifts its arm, wheels around, deliberately smashes the bottle, and trundles off. Back at the plant, when interrogated by customer service as to why it had abandoned its commitments, the robot replies that according to its specifications, it kept its commitments as long as required — commitments must be dropped when fulfilled or impossible to achieve. By smashing the last bottle, the commitment became unachievable.

Despite the impeccable logic, and the correct implementation, Willie is dismantled.

1.1 Rational Balance

This paper is concerned with specifying the "rational balance"² needed among the beliefs, goals, plans, intentions, commitments, and actions of autonomous agents. For example, it would be reasonable to specify that agents should act to achieve their intentions, and that agents should adopt only intentions that they believe to be achievable. Another constraint might be that once an agent has an intention, it believes it will do the intended act. Furthermore, we might wish that agents keep (or commit to) their intentions over time and not drop them precipitously. However, if an agent's beliefs change, it may well need to alter its intentions. Intention revision may also be called for when an agent tries and fails to fulfill an intention, or even when it succeeds. Thus, it is not enough to characterize what it means for an agent to have an intention; one also needs to describe how that intention affects the agent's beliefs, commitments to future actions, and ability to adopt still other intentions during plan formation.

Because autonomous agents will have to exist in *our* world, making commitments to us and obeying our orders, a good place to begin a normative study of rational balance is to examine various commonsense relationships among people's beliefs, intentions, and commitments that

¹This problematic robot is very loosely based on Willie, the robot in Philip K. Dick's novel *The Galactic Pot-Healer*.

²We thank Nils Nilsson for this apt phrase.

seem to justify our attribution of the term “rational”. However, rather than just characterizing agents in isolation, we propose a logic suitable for describing and reasoning about these mental states in a world in which agents will have to interact with others. Not only will a theorist have to reason about the kinds of interactions agents can have, agents may themselves need to reason about the beliefs, intentions, and commitments of other agents. The need for agents to reason about others is particularly acute in circumstances requiring communication. In this vein, the formalism serves as a foundation for a theory of speech acts [13,12], and applies more generally to situations of rational interaction in which communication may take place in a formal language.

In its emphasis on formally specifying constraints on the design of autonomous agents, this paper is intended to contribute to Artificial Intelligence research. To the extent that our analysis captures the ordinary concept of intention, this paper may contribute to the philosophy of mind. We discuss both areas below.

1.2 Artificial Intelligence Research on Planning Systems

AI research has concentrated on algorithms for finding plans to achieve given goals, on monitoring plan execution [17], and on replanning. Recently, planning in dynamic, multiagent domains has become a topic of interest, especially the planning of communication acts needed for one agent to affect the mental state and behavior of another [1,3,4,15,14,13,19,20,27,39,38]. Typically, this research has ignored the issues of rational balance — of precisely how an agent’s beliefs, goals, and intentions should be related to its actions.³ In such systems, the theory of intentional action embodied by the agent is expressed only as code, with the relationships among the agent’s beliefs, goals, plans, and actions left implicit in the agent’s architecture. If asked, the designer of a planning system may say that the notion of intention is defined operationally: A planning system’s intentions are no more than the contents of its plans. As such, intentions are representations of possible actions the system may take to achieve its goal(s). This much is reasonable; there surely is a strong relationship between plans and intentions [36]. However, what constitutes a plan for most planning systems is itself often a murky topic.⁴ Thus, saying that the system’s intentions are the contents of its plans lacks needed precision. Moreover, operational definitions are usually quite difficult to reason with and about. If the program changes, then so may the definitions, in which case there would not be a fixed set of specifications that the program implements. This paper can be seen as providing both a logic in which to write specifications for autonomous agents, and an initial theory cast in that logic.

³Exceptions include the work of Moore [34] who analyzed the relationship of knowledge to action, and that of Appelt [4], and Konolige [26,25]. However, none of these works addressed the issue of goals and intention.

⁴Rosenschein [40] discusses some of the difficulties of hierarchical planners, and presents a formal theory of plans in terms of dynamic logic.

1.3 Philosophical Theories of Intention

Philosophers have long been concerned with the concept of intention, often trying to reduce it to some combination of belief and desire. We shall explore their territory here, but cannot possibly do justice to the immense body of literature on the subject. Our strategy is to make connection with some of the more recent work, and hope our efforts are not yet another failed attempt, amply documented in *The Big Book of Classical Mistakes*.

Philosophers have drawn a distinction between future-directed intentions and present-directed ones [9,8,41]. The former guide agents' planning and constrain their adoption of other intentions [8], whereas the latter function *causally* in producing behavior [41]. For example, one's future-directed intentions may include cooking dinner tomorrow, and one's present-directed intentions may include moving an arm now. Most philosophical analysis has examined the relationship between an agent's doing something intentionally and that agent's having a present-directed intention. Recently, Bratman [9] has argued that intending to do something (or having an intention) and doing something intentionally are not the same phenomenon, and that the former is more concerned with the coordination of an agent's plans. We agree, and in this paper, we concentrate primarily on future-directed intentions. Hereafter, the term "intention" will be used in that sense only.

Intention has often been analyzed differently from other mental states such as belief and knowledge. First, whereas the content of beliefs and knowledge is usually considered to be in the form of propositions, the content of an intention is typically regarded as an action. For example, Castañeda [10] treats the content of an intention as a "practition", akin (in computer science terms) to an action description. It is claimed that by doing so, and by strictly separating the logic of propositions from the logic of practitioners, one avoids undesirable properties in the logic of intention, such as the fact that if one intends to do an action a one must also intend to do a or b. However, it has also been argued that needed connections between propositions and practitioners may not be derivable [7].

Searle [41] claims that the content of an intention is a causally self-referential representation of its conditions of satisfaction (and see also [24]). That is, for an agent to intend to go to the store, the conditions of satisfaction would be that the intention should cause the agent to go to the store. Our analysis is incomplete in that it does not deal with this causal self-reference. Nevertheless, the present analysis will characterize many important properties of intention discussed in the philosophical literature.

A second difference among kinds of propositional attitudes is that some, such as belief, can be analyzed in isolation — one axiomatizes the properties of belief apart from those of other attitudes. However, intention is intimately connected with other attitudes, especially belief, as well as with time and action. Thus, any formal analysis of intention must explicate these relationships. In the next sections, we explore what it is that theories of intention should handle.

1.4 Desiderata for a Theory of Intention

Bratman [8] argues that rational behavior cannot just be analyzed in terms of beliefs and desires (as many philosophers have held). A third mental state, intention, which is related in many interesting ways to beliefs and desires but is not reducible to them, is necessary. There are two justifications for this claim. First, noting that agents are resource-bounded, Bratman suggests that no agent can continually weigh his⁵ competing desires, and concomitant beliefs, in deciding what to do next. At some point, the agent must just *settle on* one state-of-affairs for which to aim. Deciding what to do establishes a limited form of *commitment*. We shall explore the consequences of such commitments.

A second reason is the need to coordinate one's future actions. Once a future act is settled on, that is, intended, one typically decides on other future actions to take with that action as given. This ability to plan to do some act *A* in the future, and to base decisions on what to do subsequent to *A*, requires that a rational agent *not* simultaneously believe he will *not* do *A*. If he did, the rational agent would not be able to plan past *A* since he believes it will not be done. Without some notion of commitment, deciding what else to do would be a hopeless task.

Bratman argues that unlike mere desires, intentions play the following three functional roles:

1. *Intentions normally pose problems for the agent; the agent needs to determine a way to achieve them.* For example, if an agent intends to fly to New York on a certain date, and takes no actions to enable him to do so, then the intention did not affect the agent in the right way.
2. *Intentions provide a "screen of admissibility" for adopting other intentions.* Whereas desires can be inconsistent, agents do not normally adopt intentions that they believe conflict with their present and future-directed intentions. For example, if an agent intends to hardboil an egg, and knows he has only one egg (and cannot get any more in time), he should not simultaneously intend to make an omelette.
3. *Agents "track" the success of their attempts to achieve their intentions.* Not only do agents care whether their attempts succeed, but they are disposed to replan to achieve the intended effects if earlier attempts fail.

In addition to the above functional roles, it has been argued that intending should satisfy the following properties. If an agent intends to achieve *p*, then:

4. *The agent believes p is possible.*
5. *The agent does not believe he will not bring about p.*⁶
6. *Under certain conditions, the agent believes he will bring about p.*

⁵or her. We use the masculine version here throughout.

⁶The rationale for this property was discussed above.

7. *Agents need not intend all the expected side-effects of their intentions.*⁷ For example, imagine a situation not too long ago in which an agent has a toothache. Although dreading the process, the agent decides that he needs desperately to get his tooth filled. Being uninformed about anaesthetics, the agent believes that the process of having his tooth filled will necessarily cause him much pain. Although the agent intends to ask the dentist to fill his tooth, and, believing what he does, he is willing to put up with pain, the agent could surely deny that he thereby *intends* to be in pain.

Bratman argues that what one intends is, loosely speaking, a subset of what one chooses. Consider an agent as choosing one desire to pursue from among his competing desires, and in so doing, choosing to achieve some state of affairs. If the agent believes his action(s) will have certain effects, the agent has chosen those effects as well. That is, one chooses a "scenario" or a possible world. However, one does not intend everything in that scenario, for example, one need not intend harmful expected side-effects of one's actions (though if one knowingly brings them about as a consequence of one's intended action, they have been brought about *intentionally*.) Bratman argues that side-effects do not play the same roles in the agent's planning as true intentions do. In particular, they are not goals whose achievement the agent will track; if the agent does not achieve them, he will not go back and try again.

We will develop a theory in which expected side-effects are *chosen*, but not intended.

These properties are our desiderata for a treatment of intention. However, motivated by AI research, we add one other, as described below:

1.5 The "Little Nell" Problem: Not giving up too soon.

McDermott [33] points out the following difficulty with a naively designed planning system:

Say a problem solver is confronted with the classic situation of a heroine, called Nell, having been tied to the tracks while a train approaches. The problem solver, called Dudley, knows that "If Nell is going to be mashed, I must rescue him." (He probably knows a more general rule, but let that pass.) When Dudley deduces that he must do something, he looks for, and eventually executes, a plan for doing it. This will involve finding out where Nell is, and making a navigation plan to get to his location. Assume that he knows where he is, and he is not too far away; then the fact that the plan will be carried out will be added to Dudley's world model. Dudley must have some kind of data-base-consistency maintainer to make sure that the plan is deleted if it is no longer necessary. Unfortunately, as soon as an apparently successful plan is added to the world model, the consistency maintainer will notice that "Nell is going to be mashed" is no longer true. But that removes any justification for the plan, so it goes too. But that means "Nell

⁷Many theories of intention are committed to the undesirable view that expected side-effects to ones intentions are intended as well.

is going to be mashed" is no longer contradictory, so it comes back in. And so forth. (p. 102).

The agent continually plans to save Nell, and abandons its plan because it believes it will be successful. McDermott attributes the problem to the inability of various planning systems to express "Nell is going to be mashed *unless* I save her", and to reason about the concept of prevention. Haas [21] blames the problem on a failure to distinguish between actual and possible events. The planner should be trying to save Nell based on a belief that it is possible that he will be mashed, rather than on the belief that he in fact will be mashed. Although reasoning about prevention, expressing "unless", and distinguishing between possible and actual events are important aspects of the original formulation of the problem, the essence of the Little Nell problem is the more general problem of an agent's giving up an intention too soon. We shall show how to avoid it.

As should be clear from the previous discussion, much rides on an analysis of intention and commitment. In the next section, we indicate how these concepts can be approximated.

1.6 Intention as a Composite Concept

Intention will be modeled as a composite concept specifying what the agent has chosen and how the agent is committed to that choice. First, consider the desire that the agent has chosen to pursue as put into a new category. Call this chosen desire, loosely, a goal.⁸ By construction chosen desires are consistent. We will give them a possible-world semantics, and hence the agent will have chosen a set of worlds in which the goal/desire holds.

Next, consider an agent to have a *persistent goal* if he has a goal (i.e., a chosen set of possible worlds) that will be kept at least as long as certain conditions hold. For example, for a fanatic these conditions might be that his goal has not been achieved but is still achievable. If either of those circumstances fail, even the fanatical agent must drop his commitment to achieving the goal. Persistence involves an agent's *internal* commitment to a course of events over time.⁹ Although a persistent goal is a composite concept, it models a distinctive state of mind in which agents have both chosen and committed to a state of affairs.

We will model intention as a kind of persistent goal. This concept, and especially its variations allowing for subgoals, interpersonal subgoals, and commitments relative to certain other conditions, is interesting for its ability to model much of Bratman's analysis. For example, the analysis shows that agents need not intend the expected side-effects of their intentions because agents need not be committed to the expected consequences of those intentions. To preview the analysis, persistence need not hold for expected side-effects because the agent's *beliefs* about the linkage of the act and those effects could change.

Strictly speaking, the formalism predicts that agents only intend the logical equivalences of their intentions, and in some cases intend their logical consequences. Thus, even using a possible-worlds approach, one can get a fine-grained modal operator that satisfies many desirable properties of a model of intention.

⁸Such desires are ones that speech act theorists claim to be conveyed by illocutionary acts such as requests.

⁹This is not a *social* commitment. It remains to be seen if the latter can be built out of the former.

2 Methodology

2.1 Strategy: A Tiered Formalism

The formalism will be developed in two layers: atomic and molecular. The foundational atomic layer provides the primitives for the theory of rational action. At this level can be found the analysis of beliefs, goals, and actions. Most of the work here is to sort out the relationships among the basic modal operators. Although the primitives chosen are motivated by the phenomena to be explained, few commitments are made at this level to details of theories of rational action. In fact, many theories could be developed with the same set of primitive concepts. Thus, at the foundational level, we provide a framework in which to express such theories.

The second layer provides new concepts defined out of the primitives. Upon these concepts, we develop a partial theory of rational action. Defined concepts provide economy of expression, and may themselves be of theoretical significance because the theorist has chosen to form some definitions and not others. The use of defined concepts elucidates the origin of their important properties. For example, in modeling intention with persistent goals, one can see how various properties depend on particular primitive concepts.

Finally, although we do not do so in this paper (but see [12]), one can erect theories of rational interaction and communication on this foundation. By doing so, properties of communicative acts can be derived from the embedding logic of rational interaction, whose properties are themselves grounded in rational action.

2.2 Successive Approximations

The approach to be followed in this paper is to approximate the needed concepts with sufficient precision to enable us to explore their interactions. We do not take as our goal the development of an exceptionless theory, but rather will be content to give plausible analyses that cover the important and frequent cases. Marginal cases (and arguments based on them) will be ignored when developing the first version of the theory.

2.3 Idealizations

The research presented here is founded on various idealizations of rational behavior. Just as initial progress in the study of mechanics was made by assuming frictionless planes, so too can progress be made in the study of rational action with the right idealizations. Such assumptions should approximate reality—for example, beliefs can be wrong and revised, goals not achieved and dropped—but not so closely as to overwhelm. Ultimately, choosing the right initial idealizations is a matter of research strategy and taste.

A key idealization we make is that no agent will attempt to achieve something forever—everyone has limited persistence. Similarly, agents will be assumed not to procrastinate forever. Although agents may adopt commitments that can only be given up when certain conditions (C) hold, the assumption of limited persistence requires that the agent eventually drop each commitment. Hence, it can be concluded that eventually conditions C hold.

Only because of this assumption are we able to draw conclusions from an agent's adopting a persistent goal. Our strategy will be first to explore the consequences of fanatical persistence—commitment to a goal until it is believed to be achieved or unachievable. Then, we will weaken the persistence conditions to something more reasonable.

2.4 Map of the Paper

In the next sections of the paper we develop elements of a formal theory of rational action, leading up to a discussion of persistent goals and the consequences that can be drawn from them with the assumption of limited persistence. Then, we demonstrate the extent to which the analysis satisfies the above-mentioned desiderata for intention, and show how the analysis of intention solves various classical problems. Finally, we extend the underlying concept of a persistent goal to a more general one, and briefly illustrate the utility of that more general concept for rational interaction and communication. In particular, we show how agents can have interlocking commitments.

3 Elements of A Formal Theory of Rational Action

The basis of our approach is a carefully worked out theory of rational action. The theory is expressed in a logic whose model theory is based on a possible-worlds semantics. We propose a logic with four primary modal operators — BELief, GOAL, HAPPENS (what event happens next), and DONE (which event has just occurred). With these operators, we shall characterize what agents need to know to perform actions that are intended to achieve their goals. The world will be modeled as a linear sequence of events (similar to linear time temporal models [28,29]).¹⁰ By adding GOAL, we can model an agent's intentions.

Intuitively, a model for these operators includes courses of events, which consist of sequences of primitive events, that characterize what has happened and will happen in each possible world.¹¹ Possible worlds can also be related to one another via accessibility relations that partake in the semantics of BEL and GOAL. Although there are no simultaneous primitive events in this model, an agent is not guaranteed to execute a sequence of events without events performed by other agents intervening.

As a general strategy, the formalism will be too strong. First, we have the usual consequential closure problems that plague possible-worlds models for belief. These, however, will be accepted for the time being, and we welcome attempts to develop finer-grained semantics (e.g., [6,16]). Second, the formalism will describe agents as satisfying certain properties that

¹⁰This is unlike the integration of similar operators by Moore [34], who analyses how an agent's knowledge affects and is affected by his actions. That research meshed a possible-worlds model of knowledge with a situation-calculus-style, branching-time, model of action [32]. Our earlier work [13] used a similar branching time/dynamic logic model. However, the model's inability to express beliefs about what was in fact about to happen in the future led to many difficulties.

¹¹For this paper, the only events that will be considered are those performed by an agent. These events may be thought of as event types, in that they do not specify the time of occurrence, but do include all the other arguments. Thus John's hitting Mary would be such an event type.

might generally be true, but for which there might be exceptions. Perhaps a process of non-monotonic reasoning could smooth over the exceptions, but we will not attempt to specify such reasoning here (but see [35]). Instead, we assemble a set of basic principles and examine their consequences for rational interaction. Finally, the formalism should be regarded as a description or specification of an agent, rather than one that any agent could or should use.

Most of the advantage of the formalism stems from the assumption that agents have a limited tolerance for frustration; they will not work forever to achieve their goals. Yet, because agents are (often) persistent in achieving their goals, they will work to achieve them. Hence, although all goals will be dropped, they will not be dropped too soon.

3.1 Syntax

For simplicity, we adopt a logic with no singular terms, using instead predicates and existential quantifiers. However, for readability, we will often use constants. The interested reader can expand these out into the full predicative form if desired.

$\langle \text{Action-var} \rangle ::= a, b, a_1, a_2, \dots, b_1, b_2, \dots, e, e_1, e_2, \dots,$

$\langle \text{Agent-var} \rangle ::= x, y, x_1, x_2, \dots, y_1, y_2, \dots$

$\langle \text{Regular-var} \rangle ::= i, j, i_1, i_2, \dots, j_1, j_2, \dots$

$\langle \text{Variable} \rangle ::= \langle \text{Agent-var} \rangle \mid \langle \text{Action-var} \rangle \mid \langle \text{Regular-var} \rangle$

$\langle \text{Pred} \rangle ::= (\langle \text{Pred-symbol} \rangle \langle \text{Variable} \rangle_1, \dots, \langle \text{Variable} \rangle_n)$

$\langle \text{Wff} \rangle ::= \langle \text{Pred} \rangle \mid \sim \langle \text{Wff} \rangle \mid \langle \text{Wff} \rangle \vee \langle \text{Wff} \rangle \mid \exists \langle \text{Variable} \rangle \langle \text{Wff} \rangle \mid \text{one of the following:}$

$\langle \text{Variable} \rangle = \langle \text{Variable} \rangle$

$(\text{HAPPENS } \langle \text{Action-expression} \rangle) - \langle \text{Action-expression} \rangle \text{ happens next.}$

$(\text{DONE } \langle \text{Action-expression} \rangle) - \langle \text{Action-expression} \rangle \text{ has just happened.}$

$(\text{AGT } \langle \text{Agent-var} \rangle \langle \text{Action-var} \rangle) - \langle \text{Agent-var} \rangle \text{ is the only agent of } \langle \text{Action-var} \rangle.$

$(\text{BEL } \langle \text{Agent-var} \rangle \langle \text{Wff} \rangle) - \langle \text{Wff} \rangle \text{ follows from } \langle \text{Agent-var} \rangle \text{'s beliefs.}$

$(\text{GOAL } \langle \text{Agent-var} \rangle \langle \text{Wff} \rangle) - \langle \text{Wff} \rangle \text{ follows from } \langle \text{Agent-var} \rangle \text{'s goals.}$

$\langle \text{Time-proposition} \rangle$

$\langle \text{Action-var} \rangle \leq \langle \text{Action-var} \rangle$

$\langle \text{Time-proposition} \rangle ::= \langle \text{Numeral} \rangle - (\text{see below})$

$\langle \text{Action-expression} \rangle ::= \langle \text{Action-var} \rangle \mid \text{one of the following:}$

$\langle \text{Action-expression} \rangle; \langle \text{Action-expression} \rangle - \text{sequential action,}$

$\langle \text{Action-expression} \rangle \mid \langle \text{Action-expression} \rangle - \text{non-deterministic choice action.}$

$\langle \text{Wff} \rangle? - \text{test action}$

$\langle \text{Action-expression} \rangle^* - \text{iterative action.}$

Time propositions are currently just numerals. However, for ease of exposition, we shall write them as if they were time-date expressions such as 2:30PM/3/6/85. These will be true or false in a course of events at a given index iff the index is the same as that denoted by the time proposition (i.e., numeral). Depending on the problem at hand, we may use timeless propositions, such as (At Robot NY). Other problems are more accurately modeled by conjoining a time proposition, such as (At Robot NY) $\wedge$ 2:30PM/3/6/85. Thus, if the above conjunction were a goal, both conjuncts would have to be true simultaneously.

3.2 Semantics

We shall adapt the usual possible-worlds model for belief to deal with goals. Assume there is a set of possible worlds T , each one consisting of a sequence (or course) of events, temporally extended infinitely in past and future. Each possible world characterizes possible ways the world could have been, and could be. Thus, each world specifies what happens in the future. Agents usually do not know precisely which world they are in. Instead, some of the worlds in T are consistent with the agents beliefs, and some with his goals, where the consistency is specified in the usual way, by means of an accessibility relation on tuples of worlds, agents, and an index, n , into the course of events defining the world (from which one can compute a time point, if need be).

To consider what an agent believes (or has as a goal), one needs to supply a world and an index into the course of events defining that world. As the world evolves, agents' beliefs and goals change. When an agent does an action in some world, he does not bring about a new world, though he can alter the facts of that world at that time. Instead, after an event has happened, we shall say the world is in a new "state" in which new facts hold and the set of accessible worlds has been altered. That is, the agent changes the possible ways the world could be, as well as ways he *thinks* the world could be or chooses the world to be.

3.2.1 Model Theory

A model M is a structure $\langle \Theta, P, E, Agt, T, B, G, \Phi \rangle$, where Θ is a set, P is a set of people, E is a set of primitive event types, $Agt \in [E \rightarrow P]$ specifies the agent of an event, $T \subseteq [Z \rightarrow E]$ is a set of possible courses of events (or worlds) specified as a function from the integers to elements of E , $B \subseteq T \times P \times Z \times T$ is the belief accessibility relation, $G \subseteq T \times P \times Z \times T$ is the goal accessibility relation, and Φ interprets predicates. Formulas will be evaluated with respect to some possible course of events, hereafter some *possible world*, and an "index" into that possible world, that is, at a particular point in the course of events.¹²

¹²For those readers accustomed to specifying possible worlds as real-numbered times and events as denoting intervals over them (e.g., [2]), we remark that we shall not be concerned in this paper about parallel execution of events over the same time interval. Hence, we model possible worlds as courses (i.e., sequences) of events.

3.2.2 Definitions

- $D = \Theta \cup P \cup E^*$, specifying the domain of quantification. Note that quantification over sequences of primitive events is allowed. Given this, $\Phi \subseteq [Pred^k \times T \times Z \times D^k]$.
- $AGT \subseteq E^* \times P$, where $x \in AGT[e_1, \dots, e_n]$ iff there is an i such that $x = Agt(e_i)$. That is, AGT specifies the partial agents of a sequence of events.

3.2.3 Satisfaction

Assume M is a model, σ a sequence of events, n an integer, v a set of bindings of variables to objects in D , and if $v \in [Vars \rightarrow D]$, then v_d^x is that function which yields d for x and is the same as v everywhere else. We now specify what it means for M, σ, v, n to *satisfy* a wff α , which we write as $M, \sigma, v, n \models \alpha$. Because of formulas involving actions, this definition depends on what it means for an expression a to *occur* between index points n and m . This, we write as $M, \sigma, v, n \llbracket a \rrbracket m$, and is itself defined in terms of satisfaction. The definitions are as follows:

1. $M, \sigma, v, n \models P(x_1, \dots, x_k)$ iff $\langle v(x_1) \dots v(x_k) \rangle \in \Phi[P, \sigma, n]$. Notice that the interpretation of predicates depends on the world σ and the event index n .
2. $M, \sigma, v, n \models \neg \alpha$ iff $M, \sigma, v, n \not\models \alpha$.
3. $M, \sigma, v, n \models (\alpha \vee \beta)$ iff $M, \sigma, v, n \models \alpha$ or $M, \sigma, v, n \models \beta$.
4. $M, \sigma, v, n \models \exists x \alpha$ iff $M, \sigma, v_d^x, n \models \alpha$ for some d in D .
5. $M, \sigma, v, n \models (x_1 = x_2)$ iff $v(x_1) = v(x_2)$.
6. $M, \sigma, v, n \models (e_1 \leq e_2)$ iff $v(e_1)$ is a subsequence of $v(e_2)$.
7. $M, \sigma, v, n \models (\text{HAPPENS } a)$ iff $\exists m, m \geq n$, such that $M, \sigma, v, n \llbracket a \rrbracket m$. That is, a describes a sequence of events that happens "next" (after n).
8. $M, \sigma, v, n \models (\text{DONE } a)$ iff $\exists m, m \leq n$, such that $M, \sigma, v, m \llbracket a \rrbracket n$. That is, a describes a sequence of events that *just* happened (before n).
9. $M, \sigma, v, n \models (\text{AGT } x \text{ e})$ iff $AGT[v(e)] = \{v(x)\}$. AGT thus specifies the *only* agent e .
10. $M, \sigma, v, n \models (\text{BEL } x \alpha)$ iff for all σ^* such that $\langle \sigma, n \rangle B[v(x)] \sigma^*$, $M, \sigma^*, v, n \models \alpha$. That is, α *follows from* the agents beliefs iff α is true in all possible worlds accessible via B , at index n .
11. $M, \sigma, v, n \models (\text{GOAL } x \alpha)$ iff for all σ^* such that $\langle \sigma, n \rangle G[v(x)] \sigma^*$, $M, \sigma^*, v, n \models \alpha$. Similarly, α *follows from* the agents goals iff α is true in all possible worlds accessible via G , at index n .
12. $M, \sigma, v, n \models \langle \text{Time-Proposition} \rangle$ iff $v(\langle \text{Time-Proposition} \rangle) = n$.

Turning now to the occurrence of actions, we have the following:

1. $M, \sigma, v, n[e]n+m$ (where e is an event variable) iff $v(e) = e_1 e_2 \dots e_m$ and $\sigma(n+i) = e_i, 1 \leq i \leq m$. Intuitively, e denotes some sequence of events of length m which appears next after n in the world σ .
2. $M, \sigma, v, n[a \mid b]m$ iff $M, \sigma, v, n[a]m$ or $M, \sigma, v, n[b]m$. Either the action described by a or that described by b occurs within the interval.
3. $M, \sigma, v, n[a; b]m$ iff $\exists k, n \leq k \leq m$, such that $M, \sigma, v, n[a]k$ and $M, \sigma, v, k[b]m$. The action described by a and then that described by b occurs.
4. $M, \sigma, v, n[\alpha?]n$ iff $M, \sigma, v, n \models \alpha$. The test action, $\alpha?$, involves no events at all, but occurs if α holds, or "blocks" (fails), when α is false.
5. $M, \sigma, v, n[a^*]m$ iff $\exists n_1, \dots, n_k$ where $n_1 = n$ and $n_k = m$ and for every i such that $1 \leq i \leq m$, $M, \sigma, v, n_i[a]n_{i+1}$. The iterative action a^* occurs between n and m provided only a sequence of what is described by a occurs within the interval.

A wff α is *satisfiable* if there is at least one model M , world σ , index n , and value assignment v such that $M, \sigma, v, n \models \alpha$. A wff α is *valid*, iff for every model M , world σ , event index n , and assignment of variables v , $M, \sigma, v, n \models \alpha$.

3.2.4 Abbreviations

It will be convenient to adopt the following:

Empty sequence: $\text{nil} \stackrel{\text{def}}{=} (\forall x (x=x))?$. $a=\text{NIL} \stackrel{\text{def}}{=} \forall b (a \leq b)$.

As a test action, NIL always succeeds; as an event sequence, it is a subsequence of every other one.

Conditional action: $[\text{IF } \alpha \text{ THEN } a \text{ ELSE } b] \stackrel{\text{def}}{=} \alpha?;a \mid \sim \alpha?;b$

That is, as in dynamic logic, an if-then-else action is a disjunctive action of doing action a at a time at which α is true or doing action b at a time at which α is false.

While-loops: $[\text{WHILE } \alpha \text{ DO } a] \stackrel{\text{def}}{=} (\alpha?;a)^*; \sim \alpha?$

While-loops are a sequence of doing action a a zero or more times, prior to each of which α is true. After the iterated action stops, α is false.

Eventually: $\Diamond \alpha \stackrel{\text{def}}{=} \exists x (\text{HAPPENS } x; \alpha?)$.

In other words, $\Diamond \alpha$ is true (in a given possible world) if there is something that happens after which α holds, that is, if α is true at some point in the future.

Always: $\Box \alpha \stackrel{\text{def}}{=} \sim \Diamond \sim \alpha$.

$\Box \alpha$ means that α is true throughout the course of events.

A useful application of $\Box$ is $\Box(p \supset q)$, in which no matter what happens, p still implies q . We can now distinguish between $p \supset q$'s being logically valid, its being true in all courses of events, and its merely being true after some event happens.

3.2.5 Constraints on the Model

1. *Consistency*: B is Euclidean, transitive and serial, G is serial. B 's being Euclidean essentially means that the worlds the agent thinks are possible (given what is believed) form an equivalence relation but do not necessarily include the real world [22]. Seriality implies that beliefs and goals are (separately) consistent. This is enforced by there always being a world that is either B - or G -related to a given world.
2. *Realism*: $\forall \sigma, \sigma^*$, if $\langle \sigma, n \rangle G[p] \sigma^*$, then $\langle \sigma, n \rangle B[p] \sigma^*$. In other words, $B \supseteq G$. That is, the worlds that are consistent with what the agent has chosen are not ruled out by his beliefs. Without this constraint, the agent could choose worlds involving (for example) future events that he believes will never happen. We believe this condition to be so strong, and its model theoretical statement so simple, that it deserves to be imposed as a constraint. It ensures that an agent does not want the opposite of what he believes to be unchangeable. For example, assume an agent knows that he will die in two months (and he does not believe in life after death). One would not expect that agent, if still rational, to buy a plane ticket to Miami in order to play golf three months hence. Simply, an agent cannot choose such worlds since they are not compatible with what he believes.

3.3 Properties of the Model

We begin by exploring the temporal and action-related aspects of the model, describing properties of our modal operators HAPPENS, DONE, and $\Diamond$. Next, we discuss belief and relate it to the temporal modalities. Then, we explore the relationships among all these and GOAL. Finally, we characterize an agent's persistence in achieving a goal.

Valid properties of the model are termed "Propositions". Properties that constitute our theory of the interrelationships among agent's beliefs, goals, and actions will be stated as "Assumptions". These are essentially nonlogical axioms that constrain the models that we consider.¹³ The term "Theorem" is reserved for major results.

3.3.1 Events and Action Expressions

The framework proposed here separates primitive events from action expressions. Examples of primitive events might include moving an arm, grasping, exerting force, and uttering a word or sentence. Action expressions denote sequences of primitive events that satisfy certain properties. For example, a movement of a finger may result in a circuit being closed, which may result in a light coming on. We will say that one primitive event happened, but one which can be characterized by various complex action expressions. This distinction between primitive events and complex action descriptions must be kept in mind when characterizing real world phenomena or natural language expressions.

¹³One also needs to show that there is at least one model that satisfies these assumptions. This consistency proof as well as proofs of the propositions will appear in a forthcoming report.

For example, to say that an action a occurs, we use $(\text{HAPPENS } a)$. To characterize world states that are brought about, we use $(\text{HAPPENS } \sim p?; a: p?)$, saying that event a brings about p . To be a bit more concrete, one would not typically have a primitive event type for closing a circuit. So, to say that John closed the circuit one would say that John did something (perhaps a sequence of primitive events) causing the circuit to be closed — $\exists e (\text{DONE } \sim (\text{Closed } c)?; e: (\text{Closed } c)?)$.

Another way to characterize actions and events is to have predicates be true of them. For example, one could have $(\text{Walk } e)$ to say that a given event (type) is a walking event. This way of describing events has the advantage of allowing complex properties (such as running a race) to hold for an undetermined (and unnamed) sequence of events. However, because the predication is made about the events, not the attendant circumstances, this method does not allow us to describe events performed only in certain circumstances. We will need to use both methods for describing actions.

3.3.2 Properties of Acts/Events under HAPPENS

We adopt the usual axioms characterizing how complex action expressions behave under HAPPENS, as treated in a dynamic logic (e.g., [23,34,37]), including the following:

Proposition 1 *Properties of complex acts —*

- $\models (\text{HAPPENS } a:b) \equiv (\text{HAPPENS } a; (\text{HAPPENS } b)?)$
- $\models (\text{HAPPENS } a|b) \equiv (\text{HAPPENS } a) \vee (\text{HAPPENS } b)$
- $\models (\text{HAPPENS } p?: q?) \equiv p \wedge q$
- $\models (\text{HAPPENS } a^*) \equiv (\text{HAPPENS NIL} \mid a:a^*)$

That is, action $a:b$ happens next iff a happens next producing a world state in which b then happens next. The “non-deterministic choice” action $a|b$ (read “|” as “or”) happens next iff a happens next or b does. The test action $p?$ happens next iff p is currently true. Finally, the iterative action a^* happens next iff nothing happens (which always does) or one step of the iteration has been taken followed by the a^* again.

Among many additional properties, note that after doing action a , a would have just been done:

Proposition 2 $\models \forall a (\text{HAPPENS } a) \equiv (\text{HAPPENS } a; (\text{DONE } a)?)$

Also, if a has just been done, then just prior to its occurrence, it was going to happen next.

Proposition 3 $\models \forall a (\text{DONE } a) \equiv (\text{DONE } (\text{HAPPENS } a)?; a)$

Although this may seem to say that the unfolding of the world is determined only by what has just happened, and is not random, this determinacy is entirely moot for our purposes. Agents need never know what possible world they are in and hence what will happen next. More serious would be a claim that agents have no “free will” — what happens next is determined without regard to their intentions. However, as we shall see, this is not a property

of agents; their intentions constrain their future actions. Next, observe that a test action is done whenever the condition hold:

Proposition 4 $\models p \equiv (\text{DONE } p?)$

That is, the test action filters out courses of events in which the proposition tested is false. The truth of Proposition 4 follows immediately from the definition of “?”.

For convenience, let us define versions of DONE and HAPPENS that specify the agent of the act.

Definition 1 $(\text{DONE } x \ a) \stackrel{\text{def}}{=} (\text{DONE } a) \wedge (\text{AGT } x \ a)$

Definition 2 $(\text{HAPPENS } x \ a) \stackrel{\text{def}}{=} (\text{HAPPENS } a) \wedge (\text{AGT } x \ a)$

Finally, one distinction is worth pointing out. When action variables are bound by quantifiers, they range over sequences of events (more precisely, event types). When they are left free in a formula, they are intended as schematic and can be instantiated with complex action expressions.

3.3.3 Temporal Modalities: DONE, $\Diamond$, and $\Box$

Temporal concepts are introduced with DONE (for past happenings) and $\Diamond$ (read “eventually”). To say that p was true at some point in the past, we use $\exists a (\text{DONE } p? : a)$. $\Diamond$ is to be regarded in the “linear time” sense and is defined above. Essentially, $\Diamond p$ is true iff somewhere in the future, p becomes true. $\Diamond p$ and $\Diamond \sim p$ are jointly satisfiable. Since $\Diamond p$ starts “now”, the following property is also true,

Proposition 5 $\models p \supset \Diamond p$

The following are also trivial consequences:

Proposition 6 $\models \Diamond(p \vee q) \wedge \Box \sim q \supset \Diamond p$

Proposition 7 $\models \Box(p \supset q) \wedge \Diamond p \supset \Diamond q$

To talk about propositions that are not true now, but will become true, we define:

Definition 3 $(\text{LATER } p) \stackrel{\text{def}}{=} \sim p \wedge \Diamond p$

A property of this definition that follows from the equivalence $\Diamond p$ and $\Diamond \Diamond p$ is:

Proposition 8 $\models \sim(\text{LATER } \Diamond p)$

3.3.4 Constraining Courses of Events

We will have occasion to state constraints on courses of events. To do so, we define the following:

Definition 4 $(\text{BEFORE } p \ q) \stackrel{\text{def}}{=} \forall c (\text{HAPPENS } c; q?) \supset \exists a (a \leq c) \wedge (\text{HAPPENS } a; p?)$

This definition states that p comes before q (starting at index n in the course of events) if, whenever q is true in a course of events, p has been true (after the index n). Obviously,

Proposition 9 $\models \Diamond q \wedge (\text{BEFORE } p \ q) \supset \Diamond p$

That is, if q is eventually true, and q 's being true requires that p has been true, then eventually p holds. Furthermore, we have

Proposition 10 $\models \sim p \supset (\text{BEFORE } (\exists a (\text{DONE } \sim p?; a; p?)) \ p)$

This basically says that worlds are consistent—no proposition changes truth-value without some event happening. In particular, there is no notion in this model for the simple passage of time (without any intervening events) affecting anyone's beliefs or goals. One would like to adopt the view that some event must *cause* that change, but as yet, there is no primitive relation of causality.

3.4 The Attitudes

BEL and GOAL characterize what is *implicit* in an agent's beliefs and goals (chosen desires), rather than what an agent actively or explicitly believes, or has as a goal.¹⁴ That is, these operators characterize what *the world would be like* if the agent's beliefs and goals were true. Importantly, we do not include an operator for wanting, since desires need not be consistent. Although desires certainly play an important role in determining goals and intentions, we assume that once an agent has sorted out his possibly inconsistent desires in deciding what he wishes to achieve, the worlds he will be striving for are consistent.

3.5 Belief

For simplicity, we assume the usual Hintikka-style axiom schemata for BEL [22] (corresponding to a "Weak S5" modal logic)

Proposition 11 *Belief Axioms:*

- a). $\models \forall x (\text{BEL } x \ p) \wedge (\text{BEL } x \ (p \supset q)) \supset (\text{BEL } x \ q)$
- b). $\models \forall x (\text{BEL } x \ p) \supset (\text{BEL } x \ (\text{BEL } x \ p))$
- c). $\models \forall x \sim(\text{BEL } x \ p) \supset (\text{BEL } x \ \sim(\text{BEL } x \ p))$
- d). $\models \forall x (\text{BEL } x \ p) \supset \sim(\text{BEL } x \ \sim p)$

¹⁴For an exploration of the issues involved in explicit vs. implicit belief, see [31].

And, we have the usual "necessitation" rule:

Proposition 12 *If $\models \alpha$ then $\models (\text{BEL } x \ \alpha)$*

If α is a theorem (i.e., is valid), then it follows from the agent's beliefs at all times. For example, all tautologies follow from the agent's beliefs. Clearly, we also have

Proposition 13 *If $\models \alpha$ then $\models (\text{BEL } x \ \Box \alpha)$*

That is, theorems are believed to be always true. Also, we introduce KNOW by definition:

Definition 5 $(\text{KNOW } x \ p) \stackrel{\text{def}}{=} p \wedge (\text{BEL } x \ p)$

Of course, this characterization of knowledge has many known difficulties, but will suffice for present purposes. Next, we will say an agent is COMPETENT with respect to p if he is correct whenever he thinks p is true.

Definition 6 $(\text{COMPETENT } x \ p) \stackrel{\text{def}}{=} (\text{BEL } x \ p) \supset (\text{KNOW } x \ p)$

Agents competent with respect to some proposition p adopt only beliefs about that proposition for which they have good evidence. For the purposes of this paper, we assume that agents are competent with respect to the primitive actions they have done:

Assumption 1 $\models \forall x.e \ (\text{AGT } x \ e) \supset [(\text{DONE } e) \equiv (\text{BEL } x \ (\text{DONE } e))]$

Note that this assumption does *not* hold when e is replaced by an arbitrary action expression, even if x is the agent. For example, if the agent does not know the truth value of p after just doing a , the agent may have done the action $a:p?$ without realizing it was done. But the assumption rules out unknowing execution of *primitive actions by an agent*. In section 5.1, we will make additional assumptions about the actions an agent is about to perform.

3.6 Goals

At a given point in a course of events, agents choose worlds they would like (most) to be in—ones in which their *goals* are true. $(\text{GOAL } x \ p)$ is meant to be read as p is true in all worlds, accessible from the current world, that are compatible with the agent's goals. Roughly, p follows from the agent's goals. Since agents choose entire worlds, they choose the (logically and physically) necessary consequences of their goals. At first glance, this appears troublesome if we interpret the facts that are true in all worlds compatible with an agent's goals as intended. However, intention will involve a form of commitment that will rule out such consequences as being intended, although they are chosen.

GOAL has the following properties:

Proposition 14 *Consistency: $\models \forall x \ (\text{GOAL } x \ p) \supset \sim (\text{GOAL } x \ \sim p)$*

What is implicit in someone's goals is closed under consequence:

Proposition 15 $\models (\text{GOAL } x \ p) \wedge (\text{GOAL } x \ p \supset q) \supset (\text{GOAL } x \ q)$

Again, we have a necessitation property:

Proposition 16 *If* $\models \alpha$ *then* $\models (\text{GOAL } x \ \alpha)$

That is, if α is a theorem, it is true in all chosen worlds. However, agents can distinguish such “trivial” goals from others, as explained below.

3.6.1 Achievement Goals

Agents can distinguish between achievement goals and maintenance goals. Achievement goals are those the agent believes to be false; maintenance goals are those the agent already believes to be true. We shall not be concerned in this paper with maintenance goals. But, to characterize achievement goals, we use:

Definition 7 $(\text{A-GOAL } x \ p) \stackrel{\text{def}}{=} (\text{GOAL } x \ (\text{LATER } p)) \wedge (\text{BEL } x \ \sim p)$

That is, x believes (and therefore accepts) that p is currently false, but in his chosen worlds, p is eventually true. In other words, this is the more standard notion of goal, where what is desired for the future is something that is believed to be currently false.

3.6.2 No Persistence/Deferral Forever

Agents are limited in both their persistence and their procrastination. They cannot try forever to achieve their goals; eventually they give up. On the other hand, agents do not forever defer working on their goals. The assumption below captures both of these desiderata.

Assumption 2 $\models \Diamond \sim (\text{GOAL } x \ (\text{LATER } p))$

Thus, agents eventually drop all achievement goals. Because one cannot conclude that agents always act on their goals, one needs to guard against infinite procrastination. However, one could have an agent who forever fails to achieve his goals, but believes success is still achievable. The limiting case here is an agent who executes an infinite loop. Another case is that of a compulsive gambler who continually thinks success is just around the corner. Our assumption rules out these pathological cases from consideration, but still allows agents to try hard. Finally, since no one ever said the world is fair (in the computer science sense), an agent who is ready to act in what he believes to be the correct circumstance may never get a chance to execute his action because the world keeps changing. We only require that if faced with such monumental unfairness, the agent reach the conclusion that the act is impossible.

One might object that there are still achievement goals that agents could keep forever. For example, one might argue that the goal expressed by “I always want more money than I have” is kept forever (or at least as long as the agent is alive);¹⁵ but consider a plausible logical representation of that sentence in our formal language:

¹⁵However, we have assumed immortal agents.

$$\Box[\text{GOAL } I \exists x, y (\text{HAVE } I x) \wedge y > x \wedge (\text{LATER } (\text{HAVE } I y))]$$

This sentence may be true, but it does not express an achievement goal since at some points the existential may be believed to be true (and the goal is merely to maintain that truth). To express the achievement aspect, it is necessary to quantify into the GOAL clause as in

$$\Box[\forall x (\text{KNOW } I (\text{HAVE } I x)) \supset (\text{A-GOAL } I \exists y (y > x \wedge (\text{HAVE } I y)))]$$

But here, there is no single sentence that the agent always has as a goal; the goal changes because of the quantified variables. Hence, one cannot argue he keeps anything as an achievement goal forever. Instead, the agent forever gets new achievement goals.

Important consequences will follow from Assumption 2 when combined with an agent's commitments. First, we need to examine what, in general, are the consequences of having goals.

3.6.3 Goals and their Consequences

Unlike BEL, GOAL needs to be characterized in terms of all the other modalities. In particular, we need to specify how goals interact with an agent's beliefs about the future.

The semantics of GOAL specifies that worlds compatible with an agent's goals must be included in those compatible with his beliefs. This is reflected in the following property:

Proposition 17 $\models (\text{BEL } x p) \supset (\text{GOAL } x p)$

From the semantics of BEL and GOAL, one sees that p will be evaluated at the same point in the B - and G -accessible worlds. So, if an agent believes p is true now, he cannot not now want it to be currently false; agents do not choose what they cannot change. Conversely, if p is now true in all the agent's chosen worlds, then the agent does not believe it is currently false. For example, if an agent believes he has not just done event e , then he cannot have $(\text{DONE } x e)$ as a goal. Of course, he *can* have $(\text{LATER } (\text{DONE } x e))$ as a goal.

This relationship between BEL and GOAL makes more sense when one considers the future. Let p be a proposition of the form $\Box q$. So, if an agent believes q is forever true (an example would be a tautology), Proposition 17 says that any worlds that the agent chooses must have q 's being true as well. Conversely, let p be of the form $\Diamond q$. From Proposition 17, we derive that if the agent wants q to be true sometime in the future, he does not believe it will be forever false.

Notice that although an agent may have to put up with what he believes is inevitable, he may do so reluctantly, knowing that if he should change his mind about the inevitability of that state-of-affairs, his choices would change. For example, the following is satisfiable:

$$(\text{BEL } x \Diamond p \wedge \Box[\neg(\text{BEL } x \Diamond p) \supset (\text{GOAL } x \Box \neg p)])$$

That is, the agent can believe p is inevitable (and hence will eventually be true in all the agent's chosen worlds), but at the same time believe that if he ever stops believing it is inevitable, he will choose worlds in which it is never true.

Notice also that, as a corollary of Proposition 17, agent's beliefs and goals "line up" with respect to their own primitive actions that happen next.

Proposition 18 $\models \forall x. a \text{ (BEL } x \text{ (HAPPENS } x \text{ a))} \supset (\text{GOAL } x \text{ (HAPPENS } x \text{ a)))}$

That is, if an agent believes he is about to do something that is next, then its happening next is true in all his chosen worlds. Of course, "successful" agents are ones who choose what they are going to do before believing they are going to do it; they come to believe they are going to do something because they have made certain choices. We discuss this further in our treatment of intention.

Next, as another simple subcase, consider the *consequences* of facts the agent believes hold in all of that agent's chosen worlds.

Proposition 19 *Expected consequences:* $\models (\text{GOAL } x \text{ p}) \wedge (\text{BEL } x \text{ p} \supset \text{q}) \supset (\text{GOAL } x \text{ q})$

By Proposition 17, if an agent believes $p \supset q$ is true, $p \supset q$ is true in all his chosen worlds. Hence by Proposition 15, q follows from his goals as well.

At this point, we are finished with the foundational level, having described agents' beliefs and goals, events, and time. In so doing, we have characterized agents as not striving for the unachievable, and eventually foregoing the contingent. What is missing is *commitment*, to ensure that none of these goals are given up too easily.

4 Persistent goals

To capture *one* grade of commitment (fanatical) that an agent might have towards his goals, we define a persistent goal, P-GOAL, to be one that the agent will not give up until he thinks it has been satisfied, or until he thinks it will never be true. The latter case could arise easily if the proposition p is one that specifically mentions a time. Once the agent believes that time is past, he believes the proposition is impossible to achieve. Specifically, we have:

Definition 8 $(\text{P-GOAL } x \text{ p}) \stackrel{\text{def}}{=} (\text{GOAL } x \text{ (LATER } p)) \wedge (\text{BEL } x \sim p) \wedge$
 $[\text{BEFORE } ((\text{BEL } x \text{ p}) \vee (\text{BEL } x \square \sim p))$
 $\sim(\text{GOAL } x \text{ (LATER } p))]$

Notice the use of LATER, and hence $\Diamond$, above. Clearly, P-GOALS are achievement goals; the agent's goal is that p be true in the future, and he believes it is not currently true. As soon as the agent believes it will never be true, we know the agent must drop his goal (by Proposition 17), and hence his persistent goal. Moreover, as soon as an agent believes p is true, the belief conjunct of P-GOAL requires that he drop the persistent goal that p be true. Thus, these conditions are necessary and sufficient for dropping a persistent goal. However, the BEFORE conjunct does *not* say that an agent *must* give up his *simple* goal when he thinks it is satisfied, since agents may have goals of maintenance. Thus, achieving one's persistent goals may convert them into maintenance goals.

4.1 The Logic of P-GOAL

The logic of P-GOAL is weaker than one might expect. Unlike GOAL, P-GOAL does not distribute over conjunction or disjunction, and is closed only under logical equivalence. First, we examine conjunction and disjunction. Then, we turn to implication.

4.1.1 Conjunction, Disjunction, and Negation

Proposition 20 *The Logic of P-GOAL:*

- a). $\not\models (P\text{-GOAL } x \ p \wedge q) \supset (P\text{-GOAL } x \ p) \wedge (P\text{-GOAL } x \ q)$
- b). $\not\models (P\text{-GOAL } x \ p \vee q) \supset (P\text{-GOAL } x \ p) \vee (P\text{-GOAL } x \ q)$
- c). $\models (P\text{-GOAL } x \ \sim p) \supset \sim (P\text{-GOAL } x \ p)$

First, $(P\text{-GOAL } x \ p \wedge q)$ does not imply $(P\text{-GOAL } x \ p) \wedge (P\text{-GOAL } x \ q)$ because, although the antecedent is true, the agent might believe q is already true, and thus cannot have q as a P-GOAL.¹⁶ Conversely, $(P\text{-GOAL } x \ p) \wedge (P\text{-GOAL } x \ q)$ does not imply $(P\text{-GOAL } x \ p \wedge q)$, because $(\text{GOAL } x \ (\text{LATER } p)) \wedge (\text{GOAL } x \ (\text{LATER } q))$ does not imply $(\text{GOAL } x \ (\text{LATER } p \wedge q))$; p and q could be true at different times.

Similarly, $(P\text{-GOAL } x \ p \vee q)$ does not imply $(P\text{-GOAL } x \ p) \vee (P\text{-GOAL } x \ q)$ because $(\text{GOAL } x \ (\text{LATER } p \vee q))$ does not imply $(\text{GOAL } x \ (\text{LATER } p)) \vee (\text{GOAL } x \ (\text{LATER } q))$; p could hold in some possible worlds compatible with the agent's goals, and q in others. But, neither p nor q is forced to hold in all G -accessible worlds. Moreover, the implication does not hold in the other direction either, because of the belief conjunct of P-GOAL; although the agent may believe $\sim p$ or he may believe $\sim q$, that does not guarantee he believes $\sim(p \vee q)$ (i.e., $\sim p \wedge \sim q$).

With respect to the last property, note that while it is impossible to be committed to achieving both p and $\sim p$ (since one of them is not believed to be false), it is quite possible to be committed to achieving $(p \wedge \Diamond \sim p)$. However, because of Proposition 8, $(P\text{-GOAL } x \ \Diamond p)$ is always false.

4.1.2 No Consequential Closure of P-GOAL

We demonstrate that P-GOAL is closed only under logical equivalence. Below are listed the possible relationships between a proposition p and a consequence q , which we term a "side-effect". Assume in all cases that $(P\text{-GOAL } x \ p)$. Then, depending on the relationship of p to q , we have the cases shown in the figure below. We will say a "case" fails, indicated by an "N" in the third column, if $(P\text{-GOAL } x \ q)$ does not hold.

Case 1 fails for a number of reasons, most importantly because the agent's persistent goals depends on his beliefs, not on the facts. However, consider Case 2. Even though the agent may believe $p \supset q$ holds, Case 2 fails because that implication cannot affect the agent's

¹⁶For example, I may be committed to your knowing q , but not to achieving q itself.

Figure 1: P-GOAL and progressively stronger relationships between p and q

Case	Relationship of p to q	(P-GOAL $\times$ q)?
1.	$p \supset q$	N
2.	$(\text{BEL} \times p \supset q)$	N
3.	$(\text{BEL} \times \Box(p \supset q))$	N
4.	$\Box(\text{BEL} \times \Box(p \supset q))$	N
5.	$\models p \supset q$	N
6.	$\models p \equiv q$	Y

persistent goals, which refer to p 's being true *later*. That is, the agent believes p is false and does not have the goal of it currently being true.

Consider Case 3, where the agent believes the implication *always* holds. Although Proposition 17 tells us that the agent has q as a goal, we show that the agent does not have q as a *persistent* goal. Recall that P-GOAL was defined so that the only reason an agent could give up a persistent goal was if it were believed to be satisfied or believed to be forever false. However, side-effects are goals only because of a belief. If the belief changes, the agent need no longer choose worlds in which $p \supset q$ holds, and thus need no longer have q as a goal. However, the agent would have dropped the goal for reasons other than those stipulated by the definition of persistent goal, and so does not have it as a persistent goal.

Now, consider Case 4, in which the agent *always* believes the implication. Again, q need not be a persistent goal, but for a different reason. Here, an agent could believe the side-effect already held. Hence, by the second clause in the definition of P-GOAL, the agent would not have a persistent goal. This reason also blocks Case 5, closure under logical consequence. However, instances of Case 4 and Case 5 in which the agent does not believe the side-effect already holds *would* require the agent to have the side-effect as a persistent goal. Finally, in Case 6, where q is logically equivalent to p the agent has q as a persistent goal. Having shown what cannot be deduced from P-GOAL, we now turn to its major consequences.

4.2 Persistent Goals Constrain Future Beliefs and Actions

An important property of agents is that they eventually give up their achievement goals (Assumption 2). Hence, if an agent takes on a P-GOAL, he must give it up subject to the constraints imposed by P-GOAL.

Proposition 21 $\models (\text{P-GOAL} \times q) \supset \Diamond[(\text{BEL} \times q) \vee (\text{BEL} \times \Box \sim q)]$

This proposition is a direct consequence of Assumption 2, the definition of P-GOAL, and Proposition 6. In other words, because agents eventually give up their achievement goals, and because the agent has adopted a persistent goal to bring about such a proposition q ,

eventually the agent must believe q or believe q will never come true. A simple consequence of Proposition 21 is:

Proposition 22 $\models (P\text{-GOAL } x \text{ (DONE } x \text{ } a)) \supset \Diamond[(\text{DONE } x \text{ } a) \vee (\text{BEL } x \Box \sim(\text{DONE } x \text{ } a))]$,

By Proposition 21, the agent eventually believes that he has done the act or that he will never do it. By Assumption 1, if the agent believes he has just done the act, then he has. We now give a crucial theorem:

Theorem 1 *From persistence to eventualities:*

If someone has a persistent goal of bringing about p , p is within his area of competence, and, before dropping his goal, the agent will not believe p will never occur, then eventually p becomes true:

$$\begin{aligned} &\models (P\text{-GOAL } y \text{ } p) \wedge \Box(\text{COMPETENT } y \text{ } p) \\ &\quad \wedge \sim(\text{BEFORE } (\text{BEL } y \Box \sim p) \sim(\text{GOAL } y \text{ (LATER } p))) \supset \Diamond p \end{aligned}$$

Proof: By Proposition 21, the agent eventually believes either that p is true, or that p is unachievable. If he eventually thinks p is true, since he is always competent with respect to p , he is correct. The other alternative sanctioned by Proposition 21, that the agent believes p is unachievable, is ruled out by the assumption that (it so happens to be the case that) any belief of the agent that the goal is unachievable can come only *after* the agent drops his goal. Hence, by Proposition 6, the goal comes about. ■

If an agent who is not competent with respect to p adopts p as a persistent goal, we cannot conclude that eventually p will be true, since the agent could forever create incorrect plans. If the goal is not persistent, we also cannot conclude $\Diamond p$ since the agent could give up the goal without achieving it. If the goal actually is unachievable, but the agent does not know this and commits to achieving it, then we know that eventually, perhaps after trying hard to achieve it, the agent will come to believe it is forever false and give up.

4.2.1 Acting on Persistent Goals

As mentioned earlier, one cannot conclude that, merely by committing to a chosen proposition (set of possible worlds), the agent will act; someone else could bring about the desired state of affairs. However, if the agent knows that he is the only one who could bring it about, then, under certain circumstances, we can conclude the agent will act. For example, propositions of the form $(\text{DONE } x \text{ } a)$ can only be brought about by the agent x . So, if an agent always believes the act a can be done (or at least believes it for as long as he keeps the persistent goal), the agent will act.

A simple instance of Proposition 21 is one where q is $(\text{HAPPENS } x \text{ } a)$. Such a goal is one in which the agent's goal is that eventually the next thing that happens is his doing action a . Eventually, the agent believes either the next action is his, or the agent eventually comes to believe he will never get the chance. We cannot guarantee that the agent will actually do

the action next, for someone else could act before him. If the agent never believes his act will never be done, then by Proposition 21, the agent will eventually believe (HAPPENS x a). By Proposition 18, we know that (GOAL x (HAPPENS x a)). If the agent acts just when he believes the next act is his, we know that he did so believing it would happen next and having its happening next as his goal. One could say, loosely, that the agent acted "intentionally".

Bratman [9] argues that one applies the term "intentionally" to foreseen consequences as well as to truly intended ones. That is, one intends a subset of what is done intentionally. Proposition 18 requires only that agents have expected effects as goals, but not as persistent goals. Hence, the agent would in fact bring about intentionally all those foreseen consequences of his goal that actually obtain from his doing the act. However, he would not be committed to bringing about the side-effects, and thus did not intend to do so.

If agents adopt *time-limited* goals, such as (BEFORE (DONE x e) 2:30pm/6/24/86), one *cannot* conclude the agent definitely will act *in time*, even if he believes it is possible to act. Simply, the agent might wait too long. However, one *can* conclude (see below) that the agent will not adopt another persistent goal to do a non-NIL act he believes would make the persistent goal unachievable. Still, the agent could unknowingly (and hence, by Proposition 18, accidentally) make his persistent goal forever false. If one makes the further assumption that agents always know what they are going to do just before doing it, then one can conclude agents will not in fact do anything to make their persistent goals unachievable.

All these conclusions are, we believe, reasonable. However, they do not indicate what the "normal" case is. Instead, we have characterized the possibilities, and await a theory of default reasoning to further describe the situation.

One final complication worth noting is that even if we assume that agents are perfectly competent about their beliefs and goals, it is unreasonable to assume that they are competent about their persistent goals. They may have incorrect beliefs about the BEFORE clause and misjudge the conditions under which they give up their achievement goals. A simple case is an agent that makes a promise (perhaps hastily), thinks he is committed, and then, finding out more about the situation, changes his mind and drops his goal without believing that it is satisfied or unachievable. Given that P-GOAL is based on whether an agent really is committed, the question remains as to the role of beliefs in one's commitments in a theory of this type.

5 Intention as a Kind of Persistent Goal

Before defining intention, we will need certain assumptions regarding the beliefs an agent has about what he is about to perform.

5.1 Belief and Action

Assumption 1 characterizes retrospective beliefs about the past performance of primitive actions. Also of interest are beliefs regarding primitive actions that are about to be done next. We then examine the case of actions characterized by action expressions.

First, consider an agent's beliefs about sequences of actions he is about to do.¹⁷

Assumption 3 $\models (\text{BEL } x (\text{HAPPENS } x a_1 a_2)) \supset$
 $(\text{BEL } x (\text{HAPPENS } x a_1: (\text{BEL } x (\text{HAPPENS } x a_2))?: a_2))$

In other words, if an agent now believes he is about to do a_1 followed by a_2 , then he now believes that he is about to do a_1 bringing about a state of the world in which he believes he is about to do a_2 . Of course, in actuality, things could go awry, and he could change his beliefs after doing the first action. But, what this assumption says is that he *now* believes that after doing the first action he will believe he is about to do the second.

Next, we assume that if an agent thinks he is about to do *something*, then there must be some initial sequence of which he believes he is going to do next.

Assumption 4 $\models (\text{BEL } x \exists e \neq \text{NIL} (\text{HAPPENS } x e)) \supset$
 $(\text{BEL } x \exists e \neq \text{NIL} [(\text{HAPPENS } x e) \wedge \exists e' \neq \text{NIL } e' \leq e \wedge (\text{BEL } x (\text{HAPPENS } x e'))])$

The antecedent would be true if the agent believes he is about to do a complex action (e.g., one containing a disjunction, or an iteration until a condition is satisfied). So, there may be uncertainty in his mind about what he is about to do. But for anything to happen at all, we assume that there must be some initial sequence that he has settled on and thinks he is going to do next.¹⁸

We make two additional assumptions to complete our statement that agents are "in control" of their primitive actions. First, agents are not undecided about which, if any, primitive actions they believe they are about to do that are next:

Assumption 5 $\models \forall e (\text{AGT } x e) \supset (\text{BEL } x (\text{HAPPENS } x e)) \vee (\text{BEL } x \sim(\text{HAPPENS } x e))$

In other words, for each primitive event of which x is the agent, either he believes the next event to happen is his doing that primitive event, or he believes it is not the next event to happen.

Lastly, the only way an agent's primitive action can happen next is if he believes it is about to happen next.

Assumption 6 $\models \forall e (\text{HAPPENS } x e) \supset (\text{BEL } x (\text{HAPPENS } x e))$

First, note that as with Assumption 1, this assumption does not hold when e is replaced by arbitrary action expressions. But again the assumption rules out accidental or unknowing execution of *primitive actions by an agent*. It does not rule out, say, one agent lifting up the

¹⁷Actually, what counts is that an agent is about to do something *that is next* in the course of events describing the world. This limitation occurs because we are not considering simultaneous actions. Future work should loosen this restriction.

¹⁸A fine point here is that this has been expressed with a belief that there is something about which there is another belief (that is to say, a sequence of the form $\text{BEL } \exists \text{ BEL}$). For our particular model of belief, this nested belief is always correct.

arm of another, or accidental actions. In the former case, the first agent is doing the action, even though it is the arm of the second that is going up. In the second case, accidental actions are possible, but happen *to* an agent. Notice that the converse of this proposition does not hold. Although an agent may believe he is about to do something that is next, he could be wrong, and another agent acts before he does (and hence acts next). Of course, there are long-standing philosophical problems lurking here about what is an action (e.g., is sneezing an action agents do, or is it something that happens to them.) We do not claim to solve such foundational problems, or even to address them seriously. Rather, we are characterizing what follows from being an agent of an action in our scheme.

With these assumptions, and given the expansion of complex action expressions in terms of the primitives, we can now complete the description of the consequences of an agent's believing he is about to do a complex action. First, consider disjunctive actions:

Proposition 23 *Agents are not non-deterministic:*

$$\begin{aligned} \models \forall e_1 \neq \text{NIL}. e_2 \neq \text{NIL} \ (\text{BEL } x \ (\text{HAPPENS } x \ e_1 | e_2) \supset \\ (\text{BEL } x \ (\text{HAPPENS } x \ e_1)) \vee (\text{BEL } x \ (\text{HAPPENS } x \ e_2))) \vee \\ \exists z \neq \text{NIL} \ (\text{BEL } x \ [z \leq e_1 \wedge z \leq e_2 \wedge (\text{HAPPENS } x \ z)]) \end{aligned}$$

That is, if an agent believes he is about to do a disjunction of (sequences of) primitive events, then he must believe he is about to do one, or believe he is about to do the other, or believe he is about to do their non-empty common subpart. For example, if an agent believes the next action in the world is his lifting his arm or his moving his foot, then the agent has an opinion on which act he will do. This Proposition follows from Assumption 4.

As a possible counterexample, imagine two agents, *A* and *B* having a fight. *A* believes he is about to block *B*'s punch by either lifting his right or lifting his left arm. However, in our model, *A* does not believe that his blocking action is the next action; the next action is *B*'s swinging. Once *B* swings, whichever act *A* does next will follow from *A*'s beliefs (albeit quickly, and perhaps unconsciously). If, in fact, there is no other intervening action (as with the example of the donkey placed between two bales of hay at equal distances) then nothing can change, so no decision will be made, and no action will take place.¹⁹

From Assumption 3 and Proposition 23, we can now show that an agent who is about to do an if-then-else, must believe the condition of if-then-else to be true, or believe it to be false. More generally, if he believes he will do an if-then-else in the future, he believes he will have an opinion on the truth of the condition. Formally, one can show that:

Proposition 24 *If-then-else:*

$$\begin{aligned} \models \forall e_1, e_2 \neq \text{NIL}. e_3 \neq \text{NIL} \ [\forall e_4 \ (e_4 \leq e_2 \wedge e_4 \leq e_3 \supset e_4 = \text{NIL}) \wedge \\ (\text{BEL } x \ (\text{HAPPENS } x \ e_1 : [\text{IF } p \ \text{THEN } e_2 \ \text{ELSE } e_3])) \supset \\ (\text{BEL } x \ (\text{HAPPENS } x \ e_1 : [(\text{BEL } x \ p \wedge (\text{HAPPENS } x \ e_2)) \vee (\text{BEL } x \ \sim p \wedge (\text{HAPPENS } x \ e_3))]))] \end{aligned}$$

This proposition says that an agent cannot believe he is about to do an if-then-else without coming to a belief about the condition. This holds true *provided* that the *then* part and the

¹⁹This is unrealistic, of course. Ultimately, the passage of time is sufficient to change beliefs. Perhaps one way to accomodate this in future models is to treat this as a natural event.

else part are disjoint (and hence, have no non-empty common subsequence). A case in which that would not be the case would be that of a physical “test”, such as testing that a liquid is acidic or basic.²⁰ An agent who believes he is about to do:

ACID?;Dip;(RED Paper)? | $\sim$ ACID?;Dip;(BLUE Paper)?

should not have to know in advance whether the liquid is acidic or basic. The third disjunct in Proposition 23 allows for this possibility.

At this point, we have characterized the dependencies among an agent’s beliefs, the actions he has taken, and the actions he is about to take (that are next). These dependencies will be vital to an understanding of intention. With our foundation laid, we are now in a position to define this concept. There will be two defining forms for INTEND, depending on whether the argument is an action or a proposition.

5.2 INTEND₁

Typically, one intends to do actions. Accordingly, we define INTEND₁ to take an action expression as its argument.

Definition 9 $(\text{INTEND}_1 x a) \stackrel{\text{def}}{=} (P\text{-GOAL } x [\text{DONE } x (\text{BEL } x (\text{HAPPENS } a))];a)],$
where *a* is any action expression.

Let us examine what this says. First of all, (fanatically) intending to do an action *a* is a special kind of commitment (i.e., persistent goal) to have done *a*. However, it is not a commitment just to doing *a*, for that would allow the agent to be committed to doing something accidentally or unknowingly. It seems reasonable to require that the agent be committed to believing he is about to do the intended action, and then doing it. Thus, intentions are future-directed, but here directed toward something happening *next*. This is as close as we can come to present-directed intention.

Secondly, it is a commitment to success — to having done the action. As a contrast, consider the following inadequate definition of INTEND₁:

$(\text{INTEND}_1 x a) \stackrel{\text{def}}{=} (P\text{-GOAL } x \exists e (\text{HAPPENS } x e;(\text{DONE } x a)))$

This would say that an intention is a commitment to being *on the verge* of doing some event *e*, after which *x* would have just done *a*.²¹ Of course, being on the verge of doing something is not the same as doing it; any unforeseen obstacle could permanently derail the agent from ever performing the intended act. This would not be much of a commitment.

5.2.1 Intending actions

Let us apply INTEND₁ to each kind of action expression. First, consider intentions to “test” *p*. $(\text{INTEND}_1 x p?)$ expands into

²⁰See [34] for another analysis of such tests.

²¹Notice that *e* could be the last step of *a*.

$(P\text{-GOAL } x \text{ (DONE } x \text{ [BEL } x \text{ (HAPPENS } x \text{ p?)]}; p?))$.

By Proposition 1, this is equivalent to $(P\text{-GOAL } x \text{ (DONE } x \text{ (KNOW } x \text{ p?))})$, which reduces to $(P\text{-GOAL } x \text{ (KNOW } x \text{ p}))$. That is, the agent is committed to coming to know p (and he does not know it now). However, the agent is not committed to bringing about p himself.

Second, consider action expressions of the form $e;p?$. An example would be knocking down a tree:

$\exists e \text{ (Chopping } e \text{ T) } \wedge \text{ (Tree T) } \wedge \text{ (INTEND}_1 x \text{ e; (Down T)?)}$

That is, there is a chopping event (type) e , such that the agent is committed to felling the tree by doing e , and he believes just prior to doing it that it will indeed knock down the tree. Notice that e is quantified outside of the INTEND_1 . This type of intention is appropriate when there is a fixed event (or event sequence) that an agent is willing to commit to. For example, with a small tree and a large axe, an agent may be very confident that the chopping event will do the trick.

However, not all trees are like this. Fortunately, chopping events can be repeated, although it need not be obvious how many times. Thus, certain intentions cannot be characterized in terms of a fixed sequence of events—an agent may never come to believe of any given event sequence that it will achieve the intention. In this case, the intention might be expressed by

$\exists e \text{ (Chopping } e \text{ T) } \wedge \text{ (Tree T) } \wedge \text{ (INTEND}_1 x \text{ [WHILE } \sim \text{(Down T) DO } e])}$.

That is, the agent intends to do e repeatedly until the tree is down. It is important to notice that at no time does the agent need to know precisely which chopping event will finally knock down the tree. Instead, the agent is committed solely to executing the chopping event until the tree is down. To give up the commitment (i.e., the persistent goal) constituting the intention, the agent must eventually come to believe he has done the iterative action believing it was about to happen. Also, by virtue of the definition of iterative actions, we know that when the agent believes he has done the iterative action, he will believe the condition is false (i.e., here, he will believe that the tree is down).

Lastly, consider intending a conditional action. $(\text{INTEND}_1 x \text{ [IF } p \text{ THEN } a \text{ ELSE } b])$ expands into

$(P\text{-GOAL } x \text{ (DONE } x \text{ [BEL } x \text{ (HAPPENS } x \text{ [IF } p \text{ THEN } a \text{ ELSE } b])]);$
 $\text{[IF } p \text{ THEN } a \text{ ELSE } b]))$

So, we know that eventually (unless, of course, he comes to believe the conditional is forever false), he will believe he has done the conditional in a state in which he believed he was just about to do it. As we discussed in Assumption 24 the agent cannot believe he is about to do a conditional (more generally, a disjunction) without either believing the condition is true or believing the condition is false. So, if one intends to do a conditional action, one expects (with the usual caveats) not to be forever ignorant about the condition. This seems just right.

In summary, we have defined intending to do an action in way that captures many reasonable properties, some of which are inherited from the commitments involved in adopting a persistent goal. However, it is often thought that one can intend to achieve states of affairs

in addition to just actions. Some cases of this are discussed above. But INTEND_1 cannot express an agent's intending to do *something* himself to achieve a state of affairs, since the event variables are quantified outside INTEND_1 . To allow for this case, we define another kind of intention, INTEND_2 .

5.3 INTEND_2

One might intend to become rich, become happy, or (perhaps controversially) to kill one's uncle,²² without having any idea how to achieve that state-of-affairs, not even having an enormous disjunction of possible options. In these cases, we shall say the agent x is committed merely to doing something himself to bring about a world state in which ($\text{RICH } x$) or ($\text{HAPPY } x$) or ($\text{DEAD } u$) hold. Notice that because of the constraints that come along with adopting such a commitment, this is stronger than having only a desire or a simple goal.

Definition 10 $(\text{INTEND}_2 x p) \stackrel{\text{def}}{=} (P\text{-GOAL } x$
 $\exists e (\text{DONE } x [(\text{BEL } x \exists e' (\text{HAPPENS } x e'; p?)) \wedge$
 $\sim(\text{GOAL } x \sim(\text{HAPPENS } x e; p?))] ?; e; p?))$

We shall explain this definition in a number of steps. First, notice that to INTEND_2 to bring about p , an agent is committed to doing some sequence of events e himself, after which p holds. However, as earlier, to avoid allowing an agent to intend to make p true by committing himself to doing something accidentally or unknowingly, we require the agent to think he is about to do *something* (event sequence e') bringing about p .²³ From Assumption 4 we know that even though the agent believes only that he will do *some* sequence of events achieving p , the agent will know which initial step he is about to take.

Now, it seems to us that the only way, short of truly wishful thinking, that an agent can believe he is about to do something to bring about p is if the agent in fact has a *plan* (good, bad, or ugly) for bringing it about. In general, it is quite difficult to define what a plan is, or define what it means for an agent to have a plan.²⁴ The best we can do, and that is not too far off, is to say that agent must believe he is about to do something (called e' here) that will bring about p . What is left for us to specify is under what conditions this belief is *justified*, ensuring for instance, that the agent never has such a belief when he has absolutely no idea of how to proceed.²⁵

Finally, we require that prior to doing e to bring about p , the agent not have as a goal e 's not bringing about p . In other words, while there may be uncertainty in the agent's mind as

²²We are not trying to be morbid here; just setting up a classic example.

²³The definition does not use e instead of e' because that would quantify e into the agent's beliefs, requiring that he (eventually) have picked out a precise sequence of events that he thinks will bring about p . If we wanted to do that, we could use INTEND_1 .

²⁴See [36] for a discussion of these issues.

²⁵One possibility is to make sure this belief *only* arises by existential generalization from a belief involving a particular action description (that is, the plan) achieving p . However, one cannot express this constraint in our logic since one cannot quantify over action expressions.

to which action will ultimately bring about p (for example, he may have a conditional plan), what does in fact happen had better be compatible with the agent's goals. This condition is required to handle the following example, due to Chisholm [11] and discussed in Searle's book [41]. An agent intends to kill his uncle. On the way to his uncle's house, this intention causes him to become so agitated that he loses control of his car, and runs over a pedestrian, who happens to be his uncle. Although the uncle is dead, we would surely say that the action that the agent did was not what was intended.

Let us cast this problem in terms of INTEND_2 , but without the condition stating that the agent should not want e not to bring about p . Call this INTEND_2' . So, assume the following is true: $(\text{INTEND}_2' \times (\text{DEAD } u))$. The agent thus has a commitment to doing some sequence of events resulting in his uncle's death, and immediately prior to doing it, he has to believe there would be some sequence (e') that he was about to do that would result in the uncle's death. However, the example satisfies these conditions, but the event he in fact does that kills his uncle may not be the one foreseen to do so. A jury requiring only the weakened INTEND_2 to convict for first-degree murder would find the agent to be guilty. Yet, we clearly have the intuition that the death was accidental.

Searle argues that a prior intention should cause an "intention-in-action" that presents the killing of the uncle as an "intentional object", and this causation is self-referential. To explain Searle's analysis would take us too far afield. However, we can handle this case by adding the second condition to the agent's mental state that just prior to doing the action that achieves p , the agent not only believes he is about to do some sequence of events to bring about p , he also does not want what he in fact does, e , *not* to bring about p . In the case in question (intuitively), the agent's plan is to get to his uncle's house and take it from there. Driving onto the sidewalk (and killing someone) is not one of the possible outcomes of this plan and so is ruled out by the agent's beliefs and goals. So, in swerving off the road, the agent may still have believed he was about to do something e' that would kill his uncle (and, by Proposition 17, he wanted e' to kill his uncle), but even allowing for indeterminacy in his plan, in none of his chosen worlds is his swerving off the road what kills his uncle.

Hence, our analysis predicts that the agent did not do what he intended, even though the end state was achieved, and resulted from his adopting an intention. Let us now see how this analysis stacks up against the problems and related desiderata.

5.4 Meeting the Desiderata

In this section we show how various properties of the commonsense concept of intention are captured by our analysis based on P-GOAL. In what follows, we shall use INTEND_1 or INTEND_2 as best fits the example. Similar results hold for analogous problems posed with the other form of intention.

We reiterate Bratman's [7,8] analysis of the roles that intentions typically play in the mental life of agents:

1. *Intentions normally pose problems for the agent; the agent needs to determine a way to achieve them.* If the agent intends an action as described by an action expression, then

the agent knows in general terms what to do. However, the action expression may have disjunctions and conditionals in it. Hence, the agent would not know at the time of forming the intention just what will be done. However, eventually, the agent will know which actions should be taken. In the case of nonspecific intentions, such as $(\text{INTEND}_2 x p)$, we can derive that, under the normal circumstances where the agent does not learn that p is unachievable, the agent eventually believes he is about to achieve p . Hence, our analysis shows the problem that is posed by adopting a non-specific intention, but does not encode the solution — that the agent will form a plan.

2. *Intentions provide a "screen of admissibility" for adopting other intentions.* If an agent has an intention to do b , and the agent (always) believes that doing a prevents the achievement of b , then the agent cannot have the intention to do $a;b$, or even the intention to do a before doing b . Thus, the following holds:

Theorem 2 Screen of Admissibility:

$\models \forall x (\text{INTEND}_1 x b) \wedge \Box (\text{BEL } x [(\text{DONE } x a) \supset \Box \sim (\text{DONE } x b)]) \supset \sim (\text{INTEND}_1 x a;b)$,
where a and b are arbitrary action expressions, and their free variables have been bound outside.

The proof is simply that there are no possible worlds in which the two intentions and the belief could all hold; in the agent's chosen worlds, if a has just been done, b will never be done. Hence, the agent cannot intend to do a before doing b . Similarly, if the agent first intends to do a , and believes the above relationship between a and b , then the agent cannot also adopt the intention to do b .²⁶

Notice that our agents cannot knowingly (and hence, by Proposition 18, deliberately) act against their own best interests. That is, they cannot intentionally act in order to make their persistent goals unachievable. Moreover, if they have adopted a time-limited intention, they cannot intend to do some other act knowing it would make achieving that time-limited intention forever false.

3. *Agents "track" the success of their attempts to achieve intentions.* In other words, agents keep their intentions after failure. Assume an agent has an intention to do a , and then does something, e , thinking it would bring about the doing of a , but he then comes to believe it did not. If the agent does not think that a can never be done, does the agent still have the intention to do a ? Yes.

Theorem 3 $\models \forall x \exists e (\text{DONE } x [(\text{INTEND}_1 x a) \wedge (\text{BEL } x (\text{HAPPENS } x a))] ?; e) \wedge$
 $(\text{BEL } x \sim (\text{DONE } x a)) \wedge \sim (\text{BEL } x \Box \sim (\text{DONE } x a)) \supset (\text{INTEND}_1 x a)$

The proof of this follows immediately from the definition of INTEND_1 , which is based on P-GOAL, which states that the intention cannot be given up until it is believed to have

²⁶Notice that the theorem does not require quantification over primitive acts, but allows a and b to be arbitrary action expressions.

been achieved or to be unachievable. Here, the agent believes it has not been achieved and does not believe it to be unachievable. Hence, the agent keeps the intention.

Other writers have proposed that if an agent intends to do a , then

4. *The agent believes that a can be done.* We do not have a modal operator for possibility. But we can state, via Proposition 17, that the agent does not believe a will never be done. This is not precisely the same as the desired property, but surely is close enough for current purposes.
5. *Sometimes, the agent believes he will in fact do a .* This is a consequence of Theorem 1, which states the conditions (call them C) under which $\Diamond(\text{DONE } x \ a)$ holds, given the intention to do a . So, if the agent believes he has the intention, and believes C holds, $\Diamond(\text{DONE } x \ a)$ follows from his beliefs as well.
6. *The agent does not believe he will never do a .* This principle is embodied directly in Proposition 17, which is validated by the simple model theoretical constraint that worlds that are consistent with one's choices are included in worlds that are consistent with one's beliefs (worlds one thinks one might be in).
7. *Agents need not intend all the expected side-effects of their intentions.* Recall that in an earlier problem, an agent intended to have his teeth filled. Not knowing about anaesthetics (one could assume this took place just as they were being first used in dentistry), he believed that it was always the case that if one's teeth are filled, one will feel pain. One could even say that surely the agent *chose* to undergo pain. Nonetheless, one would not like to say that he intended to undergo pain.

This problem is easily handled in our scheme: Let x be the patient. Assume p is (Filled-teeth x), and q is (In-pain x). Now, we know that the agent has surely chosen pain (by Proposition 19). Given all this, the following holds (see section 4.1.2, Case 3):

$$\models (\text{INTEND}_2 \ x \ p) \wedge (\text{BEL } x \ \Box(p \supset q)) \supset (\text{INTEND}_2 \ x \ q).$$

Thus, agents need not intend the expected side-effects of their intentions.

At this point, the formalism captures each of Bratman's principles. Let us now see how it avoids the "Little Nell" problem.

5.4.1 Solving the "Little Nell" Problem: When not to give up goals

Recall that in the "Little Nell" problem, Dudley never saves Nell because he believes he will be successful. This problem will occur (at least) in logics of intention that have the following initially plausible principles:

1. Give up the intention that p when you believe p holds.
2. Under at least some circumstances, an agent's intending that p entails the agent's believing that p will eventually be true.

These principles sound reasonable. One would think that a robot who forms the intention to bring David a beer should drop that intention as soon as he brings David a beer. Not dropping it constrains the adoption of other intentions, and may lead to David's receiving a year's supply of beer. The second principle says that at least in some (perhaps the normal) circumstances, one believes one's intentions will be fulfilled. Both of these principles can be found in our analysis.

One might say these two principles already entail a problem, because if p will be true, the agent might think it need not act (since p is already "in the cards"). However, this is not what Principle 1 says. Rather, the agent will drop his intentions when they are currently satisfied, not when they will be satisfied.

Now, the problem comes when one adopts a temporal or dynamic logic (modal or not, branching or not) that expresses " p will be true" as $(\text{FUTURE } p)$ or $\Diamond p$. Apply Principle 2 to a proposition p of the form $\Diamond q$. For example, let p represent "Nell is out of danger" by $\Diamond(\text{Saved Nell})$. Hence, if the agent has the intention to bring about $\Diamond(\text{Saved Nell})$, under the right circumstances ("all other things being equal"), the agent believes $\Diamond\Diamond(\text{Saved Nell})$. But, in most temporal logics, $\Diamond\Diamond(\text{Saved Nell})$ entails $\Diamond(\text{Saved Nell})$. So, it is likely that the agent believes that $\Diamond(\text{Saved Nell})$ holds as well. Now, apply Principle 1. Here, the agent had the intention to achieve $\Diamond(\text{Saved Nell})$ and the agent believes it is already true! So, the agent drops the intention, and Nell gets mashed.

Our theory of intention based on P-GOAL avoids this problem because an agent's having a P-GOAL requires that the goal be true later and that the agent not believe it is currently true. The cases of interest are those in which the agent purportedly adopts the intention and believes he will succeed. Thus, $(\text{BEL } x \Diamond\Diamond q)$, and so $(\text{BEL } x \Diamond q)$. However, while it is certainly possible to intend to achieve q , an agent *never* forms the intention to achieve anything like $\Diamond q$ since, as already noted, $(\text{P-GOAL } x \Diamond q)$ is always false.

One might argue that this analysis prevents agents from dropping their intentions when they think another agent will achieve the end goal. For example, one might want it to be possible for Dudley to drop the intention to save Nell (himself) because he thinks someone else, Dirk, is going to save him. There are two cases to consider. The first case involves goals that can only be achieved once. The second case concerns goals that can be re-achieved. We treat the second case in the next section. Regarding the first, one can easily show the following:

Theorem 4 Dropping Futile Intentions:

$$\models \forall x y \neq x \wedge (\text{BEL } x \Diamond \exists e (\text{DONE } y \sim p?:e:p?)) \supset \\ \sim(\text{INTEND}_2 x p) \vee \sim(\text{BEL } x [\exists e (\text{DONE } y \sim p?:e:p?) \supset \Box \sim \exists e (\text{DONE } x \sim p?:e:p?)])$$

That is, if an agent believes someone else truly is going to achieve p , then either the agent does not intend to achieve p himself, or he does not believe p can be achieved only once. Contrapositively, if an agent intends to achieve p , and always believes p can only be achieved once, the agent cannot simultaneously believe someone else is definitely going to achieve p . Intuitively, the reason is simple. If the agent believes someone else is definitely going to achieve p , then, because the agent believes that after doing so no one else could do so, the

agent cannot have the persistent goal of achieving p himself; he cannot consistently believe he will achieve p first, nor can he achieve p later. A more rigorous proof is left to the determined reader.

At this point, we have met the desiderata. Thus, the analysis so far has merit; but we are not finished. The definition of P-GOAL can be extended to make explicit what is only implicit in the commonsense concept of intention — the background of other justifying beliefs and intentions. Doing so will make our agents more reasonable.

6 An End to Fanaticism

As the formalism stands now, once an agent has adopted a persistent goal, he will not be deterred. For example, if agent A receives a request from agent B , and decides to cooperate by adopting a persistent goal to do the requested act, B cannot “turn A off”. This is clearly a defect that needs to be remedied. The remedy depends on the following definition.

Definition 11 *Persistent, Relativized Goal:*

$$\begin{aligned} (\text{P-R-GOAL } x \text{ } p \text{ } q) \stackrel{\text{def}}{=} & (\text{GOAL } x \text{ } (\text{LATER } p)) \wedge (\text{BEL } x \text{ } \sim p) \wedge \\ & (\text{BEFORE } [(\text{BEL } x \text{ } p) \vee (\text{BEL } x \text{ } \Box \sim p) \vee (\text{BEL } x \text{ } \sim q)] \\ & \sim (\text{GOAL } x \text{ } (\text{LATER } p))) \end{aligned}$$

That is, a necessary condition for giving up a P-R-GOAL is that the agent x believes it is satisfied, or believes it is unachievable, or believes $\sim q$. Such propositions q form a background that justifies the agent’s intentions. In many cases, such propositions constitute the agent’s *reasons* for adopting the intention. For example, x could adopt the persistent goal to buy an umbrella relative to his belief that it will rain. He could then consider dropping his persistent goal should he come to believe that the forecast has changed.

Our analysis supports the observation that intentions can (loosely speaking) be viewed as the contents of plans (e.g., [8,15,36]). Although we have not given a formal analysis of plans here (see [36] for such an analysis), the commitments one undertakes with respect to an action in a plan depend on the other planned actions, as well as the pre- and post-conditions brought about by those actions. If x adopts a persistent goal p relative to $(\text{GOAL } x \text{ } q)$, then necessary conditions for x ’s dropping his goal include his believing that he no longer has q as a goal. Thus, $(\text{P-R-GOAL } x \text{ } p \text{ } (\text{GOAL } x \text{ } q))$ characterizes an agent’s having a persistent *subgoal* p relative to the *supergoal* q . An agent’s dropping a supergoal is now a necessary (but not sufficient) prerequisite for his dropping a subgoal.²⁷ Thus, with the change to relativized persistent goals, we open up the possibility of having a complex web of interdependencies among the agent’s goals, intentions, and beliefs. We always had the possibility of conditional P-GOALS. Now, we have added background conditions that could lead to a revision of one’s persistent goals. The definitions of intention given earlier can now be recast in terms of P-R-GOAL.

²⁷Also, notice that $(\text{P-GOAL } x \text{ } p)$ is now subsumed by $(\text{P-R-GOAL } x \text{ } p \text{ } \sim p)$.

Definition 12 $(\text{INTEND}_1 x a q) \stackrel{\text{def}}{=} (\text{P-R-GOAL } x$
 $\quad \quad \quad [(\text{DONE } x$
 $\quad \quad \quad (\text{BEL } x (\text{HAPPENS } x a:p?))?:a:p?))]$
 $\quad \quad \quad q)$

Definition 13 $(\text{INTEND}_2 x p q) \stackrel{\text{def}}{=} (\text{P-R-GOAL } x$
 $\quad \quad \quad \exists e (\text{DONE } x [(\text{BEL } x \exists e' (\text{HAPPENS } x e':p?)) \wedge$
 $\quad \quad \quad \sim(\text{GOAL } x \sim(\text{HAPPENS } x e:p?))]:e:p?))$
 $\quad \quad \quad q)$

With these changes, the dependencies of an agent's intentions on his beliefs, other goals, intentions, and so on, become explicit. For example, we can express an agent's intending to take an umbrella relative to believing it will rain on March 5, 1986 as:

$\exists e.u (\text{Take } u e) \wedge [\text{INTEND}_1 x e:3/5/86? \Diamond (\text{Raining} \wedge 3/5/86)]$

One can now describe agents whose primary concern is with the end result of their intentions, not so much with achieving those results themselves. An agent may first adopt a persistent goal to achieve p , and then (perhaps because he does not know any other agent who will, or can, do so), subsequently decides to achieve p himself, relative to that persistent goal. So, the following is true of the agent: $(\text{P-GOAL } x p) \wedge (\text{INTEND}_2 x p (\text{P-GOAL } x p))$. If someone else achieves p (and the agent comes to believe is true), the agent must drop $(\text{P-GOAL } x p)$, and is therefore free to drop the commitment to achieving p himself. Notice, however, that for goals that can be reached, the agent is *not forced* to drop the intention, as the agent may truly be committed to achieving p himself.

Matters get more interesting still when we allow the relativization conditions q to include propositions about other agents. For example, if q is $(\text{GOAL } y s)$, then y 's goal is an *interpersonal supergoal* for x . The kind of intention that is engendered by a request seems to be a P-R-GOAL. Namely, the speaker tries to bring it about that $(\text{P-R-GOAL } \text{hearer } (\text{DONE } \text{hearer } a) (\text{GOAL } \text{speaker } (\text{DONE } \text{hearer } a)))$. The hearer can get "off the hook" if he learns the speaker does not want him to do the act after all.

Notice also that given this partial analysis of requesting, a hearer who merely says "OK" and thereby accedes to a request has (made it mutually believed that he has) adopted a commitment relative to the speaker's desires. In other words, he is committed to the speaker to do the requested action. This helps to explain how social commitments can arise out of communication. However, this is not the place to analyze speech acts (but see [12]).

Finally, interlocking commitments are obtained when two agents are in the following states: $(\text{P-R-GOAL } x p (\text{GOAL } y p))$, and $(\text{P-R-GOAL } y p (\text{GOAL } x p))$. Each agent will keep his intention at least as long as the other keeps it. For example, each might have the intention to lift a table. But each would not bother to try unless the other also had the same intention. This goes partway towards realizing a notion of "joint agency" espoused by Searle [42].²⁸

²⁸Ultimately, one can envision circular interlinkages in which one agent adopts a persistent goal provided another agent has adopted it relative to the first agent having adopted it relative to the second having adopted it, etc. For an analysis of circular propositions that might make such concepts expressible, see [5].

In summary, persistent relativized goals provide a useful analysis of intention, and extend the commonsense concept by making explicit the conditions under which an agent will revise his intentions.

7 Conclusion

This paper establishes basic principles governing the rational balance among an agent's beliefs, actions, and intentions. Such principles provide specifications for artificial agents, and approximate a theory of human action (as philosophers use the term). By making explicit the conditions under which an agent can drop his goals, that is, by specifying how the agent is *committed* to his goals, the formalism captures a number of important properties of intention. Specifically, the formalism provides analyses for Bratman's three characteristic functional roles played by intentions [7,8], and shows how agents can avoid intending all the foreseen side-effects of what they actually intend. Finally, the analysis shows how intentions can be adopted relative to a background of relevant beliefs and other intentions or goals. By relativizing one agent's intentions in terms of beliefs about another agent's intentions (or beliefs), we derive a preliminary account of interpersonal commitments.

The utility of the theory for describing people or artificial agents will depend on the fidelity of the assumptions. It does not seem unreasonable to require that a robot not procrastinate forever. Moreover, we surely would want a robot to be persistent in pursuing its goals, but not fanatically so. Furthermore, we would want a robot to drop goals given to it by other agents when it determines the goals need not be achieved. So, as a coarse description of an artificial agent, the theory seems workable.

The theory is not only useful for describing single agents in dynamic multiagent worlds, it is also useful for describing their interactions, especially via the use of communicative acts. In a companion paper [12], we present a theory of speech acts that builds on the foundations laid here.

Much work remains. The action theory only allowed for possible worlds consisting of single courses of events. Further developments should include basing the analysis on partial worlds/situations [6], and on temporal logics that allow for simultaneous actions [2,30,18]. Undoubtedly, the theory would be strengthened by the use of default and non-monotonic reasoning.

Lastly, we can now allay the reader's fears about the mental state of the rationally unbalanced robot, Willie. If manufactured according to our principles, it is guaranteed that the problems described will not arise again; Willie will act on its intentions, not in spite of them, and will give them up when you say so. Of course, Willie is not yet very smart (as we did not say *how* agents should form plans), but he is determined.

8 Acknowledgments

Michael Bratman, Joe Halpern, David Israel, Ray Perrault, and Martha Pollack provided many valuable suggestions. Discussions with Doug Appelt, Jim des Rivières, Michael

Georgeff, Georgia Green, Kurt Konolige, Amy Lansky, Joe Nunes, Calvin Ostrum, Fernando Pereira, Stan Rosenschein, and Moshe Vardi have also been quite helpful. Thanks to you all.

References

- [1] J. F. Allen. *A Plan-Based Approach to Speech Act Recognition*. Technical Report 131, Department of Computer Science, University of Toronto, Toronto, Canada, January 1979.
- [2] J. F. Allen. Towards a general theory of action and time. *Artificial Intelligence*, 23(2):123-154, July 1984.
- [3] J. F. Allen and C. R. Perrault. Analyzing intention in dialogues. *Artificial Intelligence*, 15(3):143-178, 1980.
- [4] D. Appelt. *Planning Natural Language Utterances to Satisfy Multiple Goals*. Ph. D. thesis, Stanford University, Stanford, California, December 1981.
- [5] J. Barwise and J. Etchemendy. Truth, the liar, and circular propositions. In preparation.
- [6] J. Barwise and J. Perry. *Situations and Attitudes*. MIT Press, Cambridge, Massachusetts, 1983.
- [7] M. Bratman. Casteñada's theory of thought and action. In J. Toberlin, editor, *Agent, Language, and the Structure of the World: Essays presented to Hector-Neri Casteñada with his Replies*, pages 149-169, Hackett, Indianapolis, Indiana, 1983.
- [8] M. Bratman. Intentions, plans, and practical reason. 1986. Harvard University Press, in preparation.
- [9] M. Bratman. Two faces of intention. *The Philosophical Review*, XCIII(3):375-405, 1984.
- [10] H. N. Casteñada. *Thinking and Doing*. Reidel, Dordrecht, Holland, 1975.
- [11] R. M. Chisolm. *Freedom and Action*. Random House, New York, 1966.
- [12] P. R. Cohen and H. J. Levesque. Rational interaction as the basis for communication. Technical Report 89, Center for the Study of Language and Information, Stanford University, 1987.
- [13] P. R. Cohen and H. J. Levesque. Speech acts and rationality. In *Proceedings of the Twenty-third Annual Meeting*, pages 49-59, Association for Computational Linguistics, Chicago, Illinois, July 1985.
- [14] P. R. Cohen and H. J. Levesque. Speech acts and the recognition of shared plans. In *Proceedings of the Third Biennial Conference*, pages 263-271, Canadian Society for Computational Studies of Intelligence, Victoria, B. C., May 1980.

- [15] P. R. Cohen and C. R. Perrault. Elements of a plan-based theory of speech acts. *Cognitive Science*, 3(3):177-212, 1979. Reprinted in *Readings in Artificial Intelligence*, Morgan Kaufman Publishing Co., Los Altos, California, B. Webber and N. Nilsson (eds.), pp. 478-495., 1981.
- [16] R. Fagin and J. Y. Halpern. Belief, awareness, and limited reasoning: preliminary report. In *Proceedings of the Ninth International Joint Conference on Artificial Intelligence*, pages 491-501, Morgan Kaufman Publishers, Inc., Los Altos, California, August 1985.
- [17] R. Fikes and N. J. Nilsson. Strips: a new approach to the application of theorem proving to problem solving. *Artificial Intelligence*, 2:189-208, 1971.
- [18] M. P. Georgeff. Actions, processes, and causality. In *Proceedings of the Timberline Workshop on Planning and Reasoning about Action*, Morgan Kaufman Publishers, Inc., Los Alto, California, 1987.
- [19] M. P. Georgeff. Communication and interaction in multi-agent planning. In *Proceedings of the National Conference on Artificial Intelligence*, pages 125-129, Morgan Kaufman Publishers, Inc., Los Altos, California, 1983.
- [20] M. P. Georgeff and A. L. Lansky. A BDI semantics for the procedural reasoning system. 1986. Technical Note in preparation, Artificial Intelligence Center, SRI International.
- [21] A. Haas. Possible events, actual events, and robots. *Computational Intelligence*, 1(2):59-70, 1985.
- [22] J. Y. Halpern and Y. O. Moses. A guide to the modal logics of knowledge and belief. In *Proceedings of the Ninth International Joint Conference on Artificial Intelligence, IJCAI*, Los Angeles, California, August 1985.
- [23] D. Harel. *First-Order Dynamic Logic*. Springer-Verlag, New York City, New York, 1979.
- [24] G. Harman. *Change in View*. Bradford Books, MIT Press, Cambridge, Massachusetts, 1986.
- [25] K. Konolige. *Experimental Robot Psychology*. Technical Note 363, Artificial Intelligence Center, SRI International, Menlo Park, California, November 1985.
- [26] K. Konolige. *A first-order formalization of knowledge and action for a multiagent planning system*. Technical Note 232, Artificial Intelligence Center, SRI International, Menlo Park, California, December 1980. (Appears in *Machine Intelligence*, Volume 10).
- [27] K. Konolige and N. J. Nilsson. Multiple-agent planning systems. In *Proceedings of the First Annual National Conference on Artificial Intelligence*, August 1980.

- [28] L. Lamport. "Sometimes" is sometimes better than "not never". In *Proceedings of the Seventh Annual ACM Symposium on Principles of Programming Languages*, pages 174-185, Association for Computing Machinery, 1980.
- [29] A. L. Lansky. *Behavioral Specification and Planning for Multiagent Domains*. Technical Note 360, Artificial Intelligence Center, SRI International, 1985.
- [30] A. L. Lansky. A representation of parallel activity based on events, structure, and causality. In *Proceedings of the Timberline Workshop on Planning and Reasoning about Action*, Morgan Kaufman Publishers, Inc., Los Alto, California, 1987.
- [31] H. J. Levesque. A logic of implicit and explicit belief. In *Proceedings of the National Conference of the American Association for Artificial Intelligence*, Austin, Texas, 1984.
- [32] J. McCarthy and P. J. Hayes. Some philosophical problems from the standpoint of artificial intelligence. In *Machine intelligence 4*, American Elsevier, New York, 1969.
- [33] D. McDermott. A temporal logic for reasoning about processes and plans. *Cognitive Science*, 6(2):101-155, April-June 1982.
- [34] R. C. Moore. *Reasoning about Knowledge and Action*. Technical Note 191, Artificial Intelligence Center, SRI International, Menlo Park, California, October 1980.
- [35] C. R. Perrault. An application of default logic to speech act theory, In preparation.
- [36] M. E. Pollack. *Inferring Domain Plans in Question Answering*. Ph. D. thesis, Department of Computer Science, University of Pennsylvania, 1986.
- [37] V. R. Pratt. *Six Lectures on Dynamic Logic*. Technical Report MIT/LCS/TM-117, Laboratory for Computer Science, MIT, Cambridge, Massachusetts, 1978.
- [38] J. S. Rosenschein. *Rational Interaction: Cooperation among Intelligent Agents*. Ph. D. thesis, Department of Computer Science, Stanford University, Stanford, California, 1986.
- [39] J. S. Rosenschein and M. R. Genesereth. *Communication and Cooperation*. Technical Report 84-5, Heuristic Programming Project, Department of Computer Science, Stanford University, Stanford, California, 1984.
- [40] S. J. Rosenschein. Plan synthesis: a logical perspective. In *Proceedings of the Seventh International Joint Conference on Artificial Intelligence*, pages 331-337, Vancouver, Canada, August 1981.
- [41] J. R. Searle. *Intentionality: An Essay in the Philosophy of Mind*. Cambridge University Press, New York City, New York, 1983.
- [42] J. R. Searle. Collective Intentionality, ms., 1987.

