



Syntactic Possibilistic Goal Generation

Célia da Costa Pereira, Andrea G. B. Tettamanzi

► To cite this version:

Célia da Costa Pereira, Andrea G. B. Tettamanzi. Syntactic Possibilistic Goal Generation. ECAI 2014 - 21st European Conference on Artificial Intelligence, Aug 2014, Prague, Czech Republic. pp.711-716, 10.3233/978-1-61499-419-0-711 . hal-01078151

HAL Id: hal-01078151

<https://hal.science/hal-01078151>

Submitted on 28 Oct 2014

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

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

Copyright

Syntactic Possibilistic Goal Generation

Célia da Costa Pereira and Andrea G. B. Tettamanzi¹

Abstract. We propose syntactic deliberation and goal election algorithms for possibilistic agents which are able to deal with incomplete and imprecise information in a dynamic world. We show that the proposed algorithms are equivalent to their semantic counterparts already presented in the literature. We show that they lead to an efficient implementation of a possibilistic BDI model of agency which integrates goal generation.

1 Introduction

Researchers in the field of AI are growing more and more aware of the importance of integrating the goal generation process when representing the agent's capabilities [12]. Indeed, the idea of fixing *a priori* the goals that an agent must achieve is acceptable only for few domains where all possible goals can be determined in advance. On the other hand, a suitable framework for representing the agent's capabilities should consider the fact that an agent must also choose which goals to pursue from the collection of goals which have been generated previously [14, 15, 5].

In most real world situations, information available for an agent is incomplete. This means that the set of the agent's beliefs represents an imprecise description of the real world. Possibility theory [17] is well-suited for modeling uncertain or vague information by means of a possibility distribution. In [5], the authors used such a distribution to propose a semantic representation of beliefs allowing thus to express that some worlds (interpretations) are more plausible for a BDI agent than others. As pointed out by Dubois *et al.* [9], this semantic view of beliefs is also the one developed in the theory of belief change by Gärdenfors [10]. In that theory, the logical consequences of the belief base [13] represents the full set of the agent's beliefs — its belief set. In this semantic setting, adding new information to a belief set comes down to some worlds which were previously possible becoming impossible. This means that the more information is available, the smaller the set of worlds considered possible by the agent and the more precise the information held by the agent. It is important to observe that, while this kind of semantic representation is well-suited to a theoretical treatment of the matter, it is not at all adapted to the implementation of an agent framework. This is why, to use a syntactic possibilistic representation of beliefs, in [6] the authors propose an equivalent syntactic possibilistic belief-change operator.

An algorithm for generating the goals of an agent based on the semantic representation of beliefs and desires was proposed in [5]. However, to the best of our knowledge, no one has yet proposed an algorithm for goal generation based on the syntactic counterparts of the possibilistic representations of beliefs and desires. This paper aims at bridging this gap and at making an efficient implementation

of the agent's goal generation and adoption possible. We provide the syntactic version of the possibilistic deliberation algorithm and of the goal-election algorithm proposed in [5].

The rest of the paper is organized as follows. Section 2 presents some basic definitions in possibility theory. Section 3 presents the representation of graded beliefs both in the syntactic and in the semantic settings. Section 4 presents the components of the possibilistic model. Section 5 presents a syntactic representation of desires while in Section 6 an algorithm for generating such kind of desires is proposed. Section 7 presents the two algorithms used for selecting as goals the maximally justified desires among the maximally possible desires. Finally, Section 8 concludes the paper.

2 Background

In this section, we present some basic definitions in possibility theory that will be used throughout the paper.

2.1 Language and Interpretations

We adopt a classical propositional language to develop the theoretical framework used to represent information manipulated by an agent.

Definition 1 (Language) Let $\mathcal{A}$ be a finite set of atomic propositions and let $\mathcal{L}$ be the propositional language such that $\mathcal{A} \cup \{\top, \perp\} \subseteq \mathcal{L}$, and, $\forall \phi, \psi \in \mathcal{L}$, $\neg \phi \in \mathcal{L}$, $\phi \wedge \psi \in \mathcal{L}$, $\phi \vee \psi \in \mathcal{L}$.

Additional connectives can be defined as useful shorthands for combination of connectives of $\mathcal{L}$, e.g., $\phi \supset \psi \equiv \neg \phi \vee \psi$.

We will denote by $\Omega = \{0, 1\}^{\mathcal{A}}$ the set of all interpretations on $\mathcal{A}$. An interpretation $\omega \in \Omega$ is a function $\omega : \mathcal{A} \rightarrow \{0, 1\}$ assigning a truth value p^ω to every atomic proposition $p \in \mathcal{A}$ and, by extension, a truth value ϕ^ω to all formulas $\phi \in \mathcal{L}$; $\omega \models \phi$ means that $\phi^\omega = 1$ (ω is a model of ϕ); if $S \subseteq \mathcal{L}$ is a set of formulas, $\omega \models S$ means $\omega \models \phi$ for all $\phi \in S$; $S \models \phi$ means that $\forall \omega \models S$, $\omega \models \phi$. The notation $[\phi]$ denotes the set of all models of formula $\phi \in \mathcal{L}$: $[\phi] = \{\omega \in \Omega : \omega \models \phi\}$. Likewise, if $S \subseteq \mathcal{L}$ is a set of formulas, $[S] = \{\omega \in \Omega : \forall \phi \in S, \omega \models \phi\} = \bigcap_{\phi \in S} [\phi]$.

2.2 Possibility Theory

Fuzzy sets [16] are sets whose elements have degrees of membership in $[0, 1]$. Possibility theory is a mathematical theory of uncertainty that relies upon fuzzy set theory, in that the (fuzzy) set of possible values for a variable of interest is used to describe the uncertainty as to its precise value. At the semantic level, the membership function of such set, π , is called a *possibility distribution* and its range is $[0, 1]$. By convention, $\pi(\omega) = 1$ means that it is totally possible for ω to be the real world, $1 > \pi(\omega) > 0$ means that ω is only somehow possible, while $\pi(\omega) = 0$ means that ω is certainly not the real world.

¹ Univ. Nice Sophia Antipolis, CNRS, I3S, UMR 7271, 06900 Sophia Antipolis, France, email: {celia.pereira, andrea.tettamanzi}@unice.fr

A possibility distribution π is said to be normalized if there exists at least one interpretation ω_0 s.t. $\pi(\omega_0) = 1$, i.e., there exists at least one possible situation which is consistent with the available knowledge.

Definition 2 (Measures) A possibility distribution π induces a possibility measure Π , its dual necessity measure N , and a guaranteed possibility measure Δ . They all apply to a classical set $A \subseteq \Omega$ and are defined as follows:

$$\Pi(A) = \max_{\omega \in A} \pi(\omega); \quad (1)$$

$$N(A) = 1 - \Pi(\bar{A}) = \min_{\omega \in \bar{A}} \{1 - \pi(\omega)\}; \quad (2)$$

$$\Delta(A) = \min_{\omega \in A} \pi(\omega). \quad (3)$$

In words, $\Pi(A)$ expresses to what extent A is consistent with the available knowledge. Conversely, $N(A)$ expresses to what extent A is entailed by the available knowledge. The guaranteed possibility measure [8] estimates to what extent *all* the values in A are actually possible according to what is known, i.e., any value in A is at least possible at degree $\Delta(A)$.

3 Syntactic and Semantic Representations

A possibilistic belief base is a finite set of weighted formulas $B = \{(\phi_i, \alpha_i), i = 1, \dots, n\}$, where α_i is understood as a lower bound of the degree of necessity $N([\phi_i])$ (i.e., $N([\phi_i]) \geq \alpha_i$). Here, $B(\phi_i) = \alpha_i$ means that the degree to which formula ϕ_i belongs to the set B is α_i . Formulas with $\alpha_i = 0$ are not explicitly represented in the belief base, i.e., only a belief which is somehow believed/accepted by the agent is explicitly represented. The higher the weight, the more certain the formula.

From belief base B the degree of belief $\mathbf{B}(\phi)$ of any arbitrary formula $\phi \in \mathcal{L}$ may be computed as follows [2]:

$$\mathbf{B}(\phi) = \max\{\alpha : B_\alpha \models \phi\}, \quad (4)$$

where $B_\alpha = \{\phi : B(\phi) \geq \alpha\}$, with $\alpha \in [0, 1]$, is called an α -cut of B . The meaning of Equation 4 is that the degree to which an agent believes ϕ is given by the maximal degree α such that ϕ is entailed only by the formulas whose degree of membership in the base is at least α . This is the *syntactic* representation of graded beliefs [2] that we will use in this paper.

Alternatively, one may regard a *belief* as a necessity degree induced by a normalized² possibility distribution π on the possible worlds $\omega \in \Omega$ [2]: $\pi : \Omega \rightarrow [0, 1]$, where $\pi(\omega)$ is the possibility degree of interpretation ω . In this case, the degree to which a given formula $\phi \in \mathcal{L}$ is believed can be calculated as $\mathbf{B}(\phi) = N([\phi]) = 1 - \max_{\omega \not\models \phi} \pi(\omega)$, where N is the necessity measure induced by π . This is the *semantic* representation of graded beliefs proposed in [2] which has been used in [5].

The syntactic and the semantic representations of graded beliefs are equivalent [7]. Therefore, they may be used interchangeably as convenience demands. This means that, given a belief base B such that, for all α , B_α is consistent, one can construct a possibility distribution π such that, for all $\phi \in \mathcal{L}$, $N([\phi]) = \max\{\alpha : B_\alpha \models \phi\}$, where B_α is the α -cut of base B . In particular, π may be defined as follows: for all $\omega \in \Omega$,

$$\pi(\omega) = 1 - \max\{\alpha : B_\alpha \models \neg\phi_\omega\}, \quad (5)$$

² Normalization of a possibility distribution corresponds to consistency of the beliefs.

where ϕ_ω denotes the *minterm* of ω , i.e., the formula satisfied by ω only. Notice that π is normalized. Indeed, since, by hypothesis, for all α , B_α is consistent, there exists an interpretation $\omega^* \in \Omega$, such that, for all $\alpha \in (0, 1]$, $\omega^* \models B_\alpha$; therefore, $\pi(\omega^*) = 1$, because no formula ϕ exists such that $\omega^* \models \neg\phi$ and $B(\phi) > 0$.

4 A Possibilistic BDI Model

The possibilistic BDI model of agency we adopt is an adaptation of the one used in [5]. The main difference is that we replace the semantic representation of beliefs and desires with syntactic representations in the form of a belief base B and a desire base D . Figure 1 provides a schematic illustration of the model.

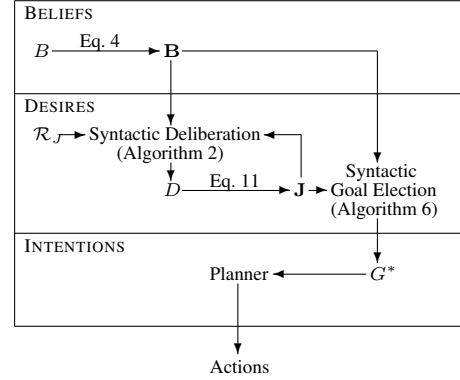


Figure 1. A schematic illustration of the proposed BDI model. The meaning of the symbols is explained in the text.

The agent receives information from the outside world and thus forms its beliefs. The “program” of the agent consists of a number of desire-generation rules, contained in a rule-base $\mathcal{R}_J$. These, together with the beliefs, determine a set of goals, which are then fed into a planner module to compute the actions to be performed by the agent.

The internal mental state of the agent is completely described by a possibilistic belief base B and by a set of desire-generation rules $\mathcal{R}_J$. The set $\mathbf{B}$ of the formulas believed by the agent is computed from B according to Equation 4.

The set $\mathbf{J}$ of the agent’s justified desires is generated dynamically through a deliberation process which applies the rules in $\mathcal{R}_J$ to the current beliefs and justified desires to produce a desire base D , from which the fuzzy set of justified desires $\mathbf{J}$ is computed according to Equation 11.

Finally, the agent rationally *elects* its goals G^* from the justified desires $\mathbf{J}$ as the most desirable of the possible sets of justified desires, according to its beliefs. The agent then plans its actions to achieve the elected goals G^* by means of a planner module, whose discussion lies outside of the scope of this paper.

In the following sections, we give technical details about all the components of the model and about syntactic deliberation and goal election algorithms. To help the reader, we will use an adaptation of the running example of [5], which follows.

Example Dr. A. Gent has submitted a paper to ECAI 2014 he has written with his co-author I. M. Flaky, who has promised to go to Prague to present it if it is accepted. Dr. Gent knows that, if the paper is accepted, publishing it (which is his great desire), means to pay the

conference registration (for his co-author or for himself) and then be ready to go to Prague to present it, in case I. M. is unavailable.

If the paper is accepted (a), Dr. Gent is willing to pay the registration (r); furthermore, if the paper is accepted and Dr. Flaky turns out to be unavailable (q), he is willing to go to Prague to present it (p). Finally, if he knows the paper is accepted and wishes to present it, he will desire to have a hotel room (h) and a plane ticket reserved (t).

Then, one fine day, Dr. Gent receives the notification of acceptance of his paper: the source is fully trustworthy—the program chair of ECAI 2014. However, soon after learning that the paper has been accepted, Dr. Flaky rushes into Dr. Gent's office to inform him that he is no more available to go to Prague; as always, Dr. Gent does not completely trust what Dr. Flaky tells him, as he is well accustomed to his changing mind. A few weeks later, rumors get to Dr. Gent's ear that all the hotels in Prague are full (f); although Dr. Gent considers this news as yet unverified, he takes notice of it. Let's assume that, combined with his *a priori* beliefs, like that if the planes are all booked out (b), he might not succeed in reserving a flight, this yields Dr. Gent's beliefs, represented by the following base:

$$B = \{(b \supset \neg t, 0.9), (f, 0.2), (q, 0.75), (a, 1), (f \supset \neg h, 1), (p \supset (a \wedge r), 1)\}. \quad (6)$$

5 Representing Desires

We may regard desires as expression of preference for some states of affairs over some others. Therefore, from a semantic point of view, such preference may be encoded as an assignment of a qualitative utility $u(\omega) \in [0, 1]$ to every world $\omega \in \Omega$: $u(\omega) = 0$ means that ω is among the least preferred worlds, $u(\omega) = 1$ means that ω is among the most preferred worlds, and $u(\omega) > u(\omega')$ means that ω is preferred to ω' . Such qualitative utility u is thus, formally, a possibility distribution.

Guaranteed possibility measure, Δ , can be used to associate a preference degree to arbitrary formulas [3]. The set of the agent's justified desires, $\mathbf{J}$, a fuzzy set of formulas in the language of choice, is thus defined based on possibility distribution u , which, unlike π , needs not be normalized, since desires may very well be inconsistent, as follows: for all formulas $\phi \in \mathcal{L}$,

$$\mathbf{J}(\phi) = \Delta([\phi]) = \min_{\omega \models \phi} u(\omega). \quad (7)$$

$\mathbf{J}$ may be extended to sets of formulas in the following way. Let $S \subseteq \mathcal{L}$ be a set of formulas, then $\mathbf{J}(S) = \min_{\phi \in S} \mathbf{J}(\phi)$.

The basic mechanism which determines how desires arise, i.e., which desires are justified and to which extent, is rule-based and may be described in terms of desire-generation rules.

Definition 3 (Desire-Generation Rule) *A desire-generation rule r is an expression of the form $\beta_r, \psi_r \Rightarrow_D^+ \phi$,³ where $\beta_r, \psi_r, \phi \in \mathcal{L}$. The unconditional counterpart of this rule is $\alpha \Rightarrow_D^+ \phi$, with $\alpha \in (0, 1]$.*

The intended meaning of a conditional desire-generation rule is: “an agent desires every world in which ϕ is true at least as much as it believes β_r and desires ψ_r ”. The intended meaning of an unconditional rule is straightforward: the degree to which the agent desires ϕ is α .

Given a desire-generation rule r , we shall denote $\text{rhs}(r)$ the formula on the right-hand side of r .

³ Note that the implication used to define a desire-generation rule is not the material implication.

Example (continued) Dr. Gent's $\mathcal{R}_J$ may be described by the following desire-generation rules:

$$\begin{aligned} R_1 : & \quad a, \quad p \Rightarrow_D^+ t \wedge h, \\ R_2 : & \quad a \wedge q, \quad \perp \Rightarrow_D^+ p, \\ R_3 : & \quad a, \quad \perp \Rightarrow_D^+ r. \end{aligned}$$

The degree of activation of a desire-generation rule depends on the degree to which its left-hand side is satisfied, i.e., the degree to which β_r is believed and ψ_r is desired.

Definition 4 (Rule Activation) *Let $r = \beta_r, \psi_r \Rightarrow_D^+ \phi$ be a desire-generation rule. The degree of activation of r , $\text{Deg}(r)$, is given by*

$$\text{Deg}(r) = \min\{\mathbf{B}(\beta_r), \mathbf{J}(\psi_r)\}.$$

For an unconditional rule $r = \alpha_r \Rightarrow_D^+ \phi$, $\text{Deg}(r) = \alpha_r$.

A semantic deliberation algorithm was proposed in [5] which, given a belief base $\mathbf{B}$ induced by a possibility distribution π and a set of desire-generation rules $\mathcal{R}_J$, computes the corresponding possibility distribution u .

In order to replace such algorithm with a syntactic deliberation algorithm, the first step is to replace the qualitative utility u with a desire base D , which, by analogy with the belief base B , will be represented as a fuzzy set of formulas. However, whereas the membership degrees of B are interpreted and treated as necessity degrees, the membership degrees of formulas of D are to be interpreted as minimum guaranteed possibilities. In D , each piece of information $[\phi_i, \alpha_i]$ expresses that any world satisfying ϕ_i is considered satisfactory for the agent to at least a degree α_i .

Let $\text{supp}(D) = \{\phi : D(\phi) > 0\}$. One important property that a desire base should obey, which is a consequence of the membership degrees representing guaranteed possibilities is that, for all formulas $\phi, \psi \in \text{supp}(D)$,

$$\text{if } \phi \models \psi \text{ then } D(\phi) \geq D(\psi). \quad (8)$$

Exactly as a belief base B induces a corresponding possibility distribution π_B , a desire base D induces a corresponding qualitative utility u_D as follows (cf. Definition 12 in [4]): for all $\omega \in \Omega$,

$$u_D(\omega) = \max_{\phi: \omega \models \phi} D(\phi). \quad (9)$$

If ω does not satisfy any formula ϕ in D then ω is not satisfactory at all for the agent. Formally, if there is no $[\phi_i, \alpha_i] \in D$ such that $\omega \models \phi_i$ then $u_D(\omega) = 0$.

The above definition may be understood as follows: every formula ϕ occurring in the desire base may be regarded as the representative of $[\phi]$, the set of its models. It has been proven in [9] that u_D , defined as per Equation 9, is the most specific possibility distribution satisfying $D(\phi) = \Delta([\phi])$ for all formulas $\phi \in \text{supp}(D)$.

Now, let $\mathbf{J}_D$ be the (fuzzy) set of justified desires in D . By definition, the degree of justification of all formulas occurring in the base must be identical to their degree of membership in the base, i.e., for all formulas ϕ ,

$$\text{if } D(\phi) > 0 \text{ then } \mathbf{J}_D(\phi) = D(\phi).$$

If $D(\phi) = \alpha$, it means that $\min_{\omega \in [\phi]} u_D(\omega) = \alpha$, or, in other terms, that, $\forall \omega \in [\phi], u_D(\omega) \geq \alpha$. Therefore, we get Equation 9.

The next step is to show how, given a desire base D , the degree of justification of any arbitrary desire formula ψ may be calculated. Based on the definition of $\mathbf{J}$, we may write, for all formulas ψ ,

$$\mathbf{J}_D(\psi) = \min_{\omega \models \psi} u_D(\omega) = \min_{\omega \models \psi} \max_{\phi: \omega \models \phi} D(\phi). \quad (10)$$

We now have to eliminate all references to the models ω from the above formula in order to make it “syntactic”. We do not want references to the models because we do not want to be obliged to enumerate all the interpretations explicitly in order to compute $\mathbf{J}_D(\psi)$. This is indeed the reason why the syntactic view for goal generation proposed in this paper should lead to a more efficient implementation than the semantic view already present in the literature. We will show this assertion later in the paper.

One way to obtain the elimination of all references to the models is to construct a set of formulas $\mathcal{P}(D)$, the “partition” of Ω according to D , containing all the $2^{\|\text{supp}(D)\|}$ conjunctions of positive or negated formulas occurring in D .

Proposition 1 *Let $\text{supp}(D) = \{\phi_1, \phi_2, \dots, \phi_n\}$ and let $\mathcal{P}(D) = \{\xi_0, \xi_1, \xi_2, \dots, \xi_{2^n-1}\}$ with*

$$\begin{aligned}\xi_0 &= \neg\phi_1 \wedge \neg\phi_2 \wedge \dots \wedge \neg\phi_n, \\ \xi_1 &= \phi_1 \wedge \neg\phi_2 \wedge \dots \wedge \neg\phi_n, \\ &\vdots \\ \xi_{2^n-1} &= \phi_1 \wedge \phi_2 \wedge \dots \wedge \phi_n.\end{aligned}$$

Then, with the convention that $\max \emptyset = 0$,

$$\begin{aligned}\mathbf{J}_D(\psi) &= \min_{\substack{\xi \in \mathcal{P}(D) \\ \psi \wedge \xi \neq \perp}} \max_{\substack{\phi \in \text{supp}(D) \\ \xi \models \phi}} D(\phi) = \\ &= \min_{\substack{i=0, \dots, 2^n-1 \\ \psi \wedge \xi_i \neq \perp}} \max_{\substack{j=1, \dots, n \\ \xi_i \models \phi_j}} D(\phi_j).\end{aligned}\quad (11)$$

Proof: The models of the formulas in $\mathcal{P}(D)$ form a partition of Ω :

$$\bigcup_{i=0}^{2^n-1} [\xi_i] = \Omega, \quad \forall i, j, i \neq j, \xi_i \wedge \xi_j = \perp.$$

The qualitative utility u_D is constant over each $[\xi_i]$: for all $\omega \models \xi_i$,

$$u_D(\omega) = \max_{\substack{j=1, \dots, n \\ \xi_i \models \phi_j}} D(\phi_j). \quad (12)$$

Therefore, instead of minimizing over all $\omega \in \Omega$, it is sufficient to minimize over all $\xi_i \in \mathcal{P}(D)$. Moreover, since we are minimizing $u_D(\omega)$ over all models of ψ , we should only consider those ξ_i such that $[\xi_i] \cap [\psi] \neq \emptyset$, i.e., such that $\xi_i \wedge \psi \neq \perp$. $\square$

Example (continued) Let’s assume Dr. Gent’s desire base is

$$D = \{(t \wedge h, 0.75), (p, 0.75), (r, 1)\}.$$

We will see later how such base may be derived from the desire-generation rules. To be able to compute $\mathbf{J}(\psi)$ for any arbitrary formula ψ , we may pre-compute $\mathcal{P}(D)$ and, for all $\xi_i \in \mathcal{P}(D)$, the corresponding term $\alpha_i = \max_{\substack{\phi \in \text{supp}(D) \\ \xi_i \models \phi}} D(\phi)$ in Equation 11:

$$\begin{aligned}\xi_0 &= \neg(t \wedge h) \wedge \neg p \wedge \neg r, & \alpha_0 &= \max \emptyset = 0, \\ \xi_1 &= (t \wedge h) \wedge \neg p \wedge \neg r, & \alpha_1 &= \max\{0.75\} = 0.75, \\ \xi_2 &= \neg(t \wedge h) \wedge p \wedge \neg r, & \alpha_2 &= \max\{0.75\} = 0.75, \\ \xi_3 &= (t \wedge h) \wedge p \wedge \neg r, & \alpha_3 &= \max\{0.75, 0.75\} = 0.75, \\ \xi_4 &= \neg(t \wedge h) \wedge \neg p \wedge r, & \alpha_4 &= \max\{1\} = 1, \\ \xi_5 &= (t \wedge h) \wedge \neg p \wedge r, & \alpha_5 &= \max\{0.75, 1\} = 1, \\ \xi_6 &= \neg(t \wedge h) \wedge p \wedge r, & \alpha_6 &= \max\{0.75, 1\} = 1, \\ \xi_7 &= (t \wedge h) \wedge p \wedge r, & \alpha_7 &= \max\{0.75, 0.75, 1\} = 1.\end{aligned}$$

Now if, for instance, we want to compute $\mathbf{J}(\psi)$, we will have to compute the minimum, for i such that $\psi \wedge \xi_i \neq \perp$, of the α_i . Therefore, we have, for example, $\mathbf{J}(t) = 0$ and $\mathbf{J}(r) = 1$.

Based on Equation 11, computing the degree of justification of formula ψ given the desire base D requires $O(n2^n)$ entailment checks, where n is the size of the desire base. Checking whether a formula entails another formula is a logical reasoning problem which may be reduced to the *satisfiability problem*, whose computational complexity varies depending on the specific logic considered, but does not depend on n . For instance, satisfiability in propositional logic (also known as Boolean satisfiability) is NP-complete [11]; concept satisfiability in description logics goes from polynomial to NEXPTIME-complete [1].

6 Generating Desires

We are now ready to present a syntactic deliberation algorithm, which calculates the desire base given a set of desire-generation rules $\mathcal{R}_J$ and a belief set $\mathbf{B}$.

We first recall the semantic deliberation algorithm presented in [5]. Let $\mathcal{R}_J^\omega = \{r \in \mathcal{R}_J : \omega \models \text{rhs}(r)\}$ denote the subset of $\mathcal{R}_J$ containing just the rules whose right-hand side would be true in world ω and $\text{Deg}_\mu(r)$ the degree of activation of rule r calculated using μ as the qualitative utility assignment. Given a mental state $\mathcal{S} = \langle \pi, \mathcal{R}_J \rangle$, the following algorithm computes the corresponding qualitative utility assignment, u .

Algorithm 1 (Semantic Deliberation)

INPUT: $\pi, \mathcal{R}_J$. OUTPUT: u .

1. $i \leftarrow 0$; for all $\omega \in \Omega$, $u_0(\omega) \leftarrow 0$;
2. $i \leftarrow i + 1$;
3. For all $\omega \in \Omega$,

$$u_i(\omega) \leftarrow \begin{cases} \max_{r \in \mathcal{R}_J^\omega} \text{Deg}_{u_{i-1}}(r), & \text{if } \mathcal{R}_J^\omega \neq \emptyset, \\ 0, & \text{otherwise;} \end{cases}$$

4. if $\max_\omega |u_i(\omega) - u_{i-1}(\omega)| > 0$, i.e., if a fixpoint has not been reached yet, go back to Step 2;
5. For all $\omega \in \Omega$, $u(\omega) \leftarrow u_i(\omega)$; u is the qualitative utility assignment corresponding to mental state $\mathcal{S}$.

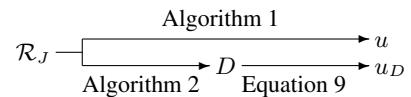
Let $\text{Deg}_X(r)$ be the degree of activation of rule r based on the desire base X . When $X = \emptyset$, the degree of justification of all desire formulas is zero. The belief base B does not change during the deliberation process.

Algorithm 2 (Syntactic Deliberation)

INPUT: $B, \mathcal{R}_J$. OUTPUT: D .

1. $i \leftarrow 0$; $D_0 \leftarrow \emptyset$;
2. $i \leftarrow i + 1$;
3. $D_i \leftarrow \{(\text{rhs}(r), \text{Deg}_{D_{i-1}}(r)) : r \in \mathcal{R}_J\}$;
4. if $D_i \neq D_{i-1}$, i.e., if a fixpoint has not been reached yet, go back to Step 2;
5. $D \leftarrow D_i$ and the deliberation has finished.

We now have to prove that this syntactic deliberation algorithm is equivalent to the semantic deliberation algorithm presented in [5]. The hypothesis is summarized in the following diagram:



Proposition 2 *We can prove that $u_D = u$.*

Proof: We may proceed by induction: we will consider the sequence $\{u_i\}_{i=0,1,\dots}$ of the possibility distributions constructed by Algorithm 1 and the sequence $\{u_{D_i}\}_{i=0,1,\dots}$, whose elements u_{D_i} are the possibility distributions induced by the desire bases D_i constructed by Algorithm 2, and we will prove that, if $u_{i-1} = u_{D_{i-1}}$, then $u_i = u_{D_i}$.

Now, $u_{i-1} = u_{D_{i-1}}$ means that $\text{Deg}_{u_{i-1}}(r) = \text{Deg}_{D_{i-1}}(r)$ for all rules r . By Equation 9, for all $\omega \in \Omega$, we may write

$$\begin{aligned} u_{D_i}(\omega) &= \max_{\phi: \omega \models \phi} D_i(\phi) \\ &= \max_{r \in \mathcal{R}_J: \omega \models \text{rhs}(r)} \text{Deg}_{D_{i-1}}(r) \\ &= \max_{r \in \mathcal{R}_J} \text{Deg}_{u_{i-1}}(r) = u_i(\omega). \end{aligned}$$

This proves the induction step. Finally, it is straightforward to verify that $u_0 = u_{D_0}$, therefore the two sequences will be identical and so their limits, u and u_D , and this concludes the proof. $\square$

Example (continued) Let us apply Algorithm 2 to Dr. Gent's mental state: we obtain

$$\begin{aligned} D_0 &= \emptyset, \\ D_1 &= \{(p, 0.75), (r, 1)\}, \\ D_2 &= \{(t \wedge h, 0.75), (p, 0.75), (r, 1)\}, \\ D_3 &= \{(t \wedge h, 0.75), (p, 0.75), (r, 1)\} = D_2. \end{aligned}$$

Therefore,

$$D = \{(t \wedge h, 0.75), (p, 0.75), (r, 1)\}. \quad (13)$$

7 Generating Goals

In [5], the assumption was made that a rational agent would select as goals the maximally justified desires among the maximally possible desires. In other words, a rational agent should first restrict attention only to those desires that it would be most *normal* (i.e., unsurprising, likely, ...) to expect they might come true and then decide to actively pursue those, among them, that have the highest qualitative utility.

In order to write a goal election algorithm according to such assumption, we need to define the set of desires possible to a given degree.

Definition 5 *Given $\gamma \in (0, 1]$, $J_\gamma = \{\phi \in \text{supp}(\mathbf{J}) : \Pi([\phi]) \geq \gamma\}$ is the (classical) subset of $\text{supp}(\mathbf{J})$ containing only those desires whose overall possibility is at least γ . We recall that $\Pi([\phi]) = 1 - \mathbf{B}(\neg\phi)$.*

We now define a goal set for a given level of possibility γ , as the set of the maximally justified γ -possible desires.

Definition 6 (Goal set) *The γ -possible goal set is*

$$G_\gamma = \begin{cases} \arg \max_{S \subseteq J_\gamma} \mathbf{J}(S) & \text{if } J_\gamma \neq \emptyset, \\ \emptyset & \text{otherwise.} \end{cases}$$

We denote by γ^* the maximum possibility level such that $G_\gamma \neq \emptyset$. Then, the goal set elected by a rational agent will be

$$G^* = G_{\gamma^*}, \quad \gamma^* = \max_{G_\gamma \neq \emptyset} \gamma. \quad (14)$$

Let $\text{Img}(\pi)$ be the level set⁴ of possibility distribution π and $\text{Img}(u)$ be the level set of qualitative distribution u . Notice that $\text{Img}(u)$ and $\text{Img}(\pi)$ are both finite, independently of Ω being finite, as proven in [5].

The following two algorithms, adapted from [5], allow an agent to compute G_γ for a given possibility lower bound γ , and the optimal goal set G^* , based on a semantic representation of beliefs and desires as two possibility distributions, π and u . We will call them *semantic*, to distinguish them from the two algorithms that we are going to propose to replace them, which will assume a syntactic representation of both beliefs and desires.

Algorithm 3 (Semantic Computation of G_γ)

INPUT: π, u . OUTPUT: G_γ .

1. $\delta \leftarrow \max \text{Img}(u)$;
2. *determine the least specific formula ϕ such that $\mathbf{J}(\phi) \geq \delta$ as follows:*

$$\phi \leftarrow \bigvee_{u(\omega) \geq \delta} \phi_\omega,$$

where ϕ_ω denotes the minterm of ω , i.e., the formula satisfied by ω only;

3. *if $\Pi([\phi]) \geq \gamma$, terminate with $G_\gamma = \{\phi\}$; otherwise,*
4. $\delta \leftarrow \max\{\alpha \in \text{Img}(u) : \alpha < \delta\}$, *0 if no such α exists;*
5. *if $\delta > 0$, go back to Step 2;*
6. *terminate with $G_\gamma = \emptyset$.*

Algorithm 4 (Semantic Goal Election)

INPUT: π, u . OUTPUT: G^* .

1. $\gamma \leftarrow \max \text{Img}(\pi) = 1$, *since π is normalized;*
2. *compute G_γ by Algorithm 3;*
3. *if $G_\gamma \neq \emptyset$, terminate with $\gamma^* = \gamma$, $G^* = G_\gamma$; otherwise,*
4. $\gamma \leftarrow \max\{\alpha \in \text{Img}(\pi) : \alpha < \gamma\}$, *0 if no such α exists;*
5. *if $\gamma > 0$, go back to Step 2;*
6. *terminate with $G^* = \emptyset$: no goal may be elected.*

Proposition 3 *The syntactic versions of Algorithms 3 and 4 are Algorithms 5 and 6 given below.*

Algorithm 5 (Syntactic Computation of $G_{\bar{\gamma}}$)

INPUT: B, D . OUTPUT: $G_{\bar{\gamma}}$.

1. $\delta \leftarrow \max \text{Img}(D)$;
2. *if $\min_{\psi \in D_\delta} \mathbf{B}(\neg\psi) \leq \bar{\gamma}$, terminate with $G_{\bar{\gamma}} = D_\delta$; otherwise,*
3. $\delta \leftarrow \max\{\alpha \in \text{Img}(D) : \alpha < \delta\}$, *0 if no such α exists;*
4. *if $\delta > 0$, go back to Step 2;*
5. *terminate with $G_{\bar{\gamma}} = \emptyset$.*

Algorithm 6 (Syntactic Goal Election)

INPUT: B, D . OUTPUT: G^* .

1. $\bar{\gamma} \leftarrow 0$;
2. *compute $G_{\bar{\gamma}}$ by Algorithm 5;*
3. *if $G_{\bar{\gamma}} \neq \emptyset$, terminate with $\gamma^* = 1 - \bar{\gamma}$, $G^* = G_{\bar{\gamma}}$; otherwise,*
4. $\bar{\gamma} \leftarrow \min\{\alpha \in \text{Img}(B) : \alpha > \bar{\gamma}\}$, *1 if no such α exists;*
5. *if $\bar{\gamma} < 1$, go back to Step 2;*
6. *terminate with $G^* = \emptyset$: no goal may be elected.*

Proof: We begin by observing that $\text{Img}(u) = \text{Img}(D)$, i.e., the level set of the desire base is the same as the level set of the corresponding qualitative utility. Furthermore, since we now have a desire base, we

⁴ The level set of a possibility distribution π is the set of $\alpha \in [0, 1] : \exists \omega$ such that $\pi(\omega) = \alpha$.

may replace the construction of ϕ based on the minterms in Step 2 of Algorithm 3 with a more straightforward

$$\phi \leftarrow \bigvee_{D(\psi) \geq \delta} \psi = \bigvee_{\psi \in D_\delta} \psi,$$

that is, the disjunction of all the formulas in the δ -cut of the desire base. This also suggests that, instead of returning $G_\gamma = \{\phi\}$, it is equivalent, but more intuitive, to return $G_\gamma = D_\delta$. Finally, instead of testing the condition $\Pi([\phi]) \geq \gamma$ in Step 3, it is equivalent to test the condition $\mathbf{B}(\neg\phi) \leq 1 - \gamma$. Now, by the DeMorgan laws and by the properties of necessity,⁵

$$\mathbf{B}(\neg\phi) = \mathbf{B}\left(\bigwedge_{\psi \in D_\delta} \neg\psi\right) = \min_{\psi \in D_\delta} \mathbf{B}(\neg\psi),$$

which allows us to avoid constructing ϕ explicitly and to test directly the condition $\min_{\psi \in D_\delta} \mathbf{B}(\neg\psi) \leq 1 - \gamma$.

This also suggests that we define $\bar{\gamma} = 1 - \gamma$ as the “impossibility of the goals”, and reformulate the goal election algorithm as a search for the least impossible set of maximally justified goals. $\square$

Since (i) Algorithm 6 iterates over the level set of the belief base and (ii) Algorithm 5, which is called as a subroutine at each iteration of Algorithm 6, loops over the level set of the desire base, and (iii) the most complex task performed at each iteration of Algorithm 5 is computing the degree of belief of each negated desire formula ψ in a δ -cut of the desire base, we may conclude that the computational cost (in number of entailment checks) of syntactic goal election is

$$O(\|\text{Img}(B)\| \cdot \|\text{Img}(D)\|) \cdot C,$$

where C is the number of entailment checks needed for computing $\mathbf{B}(\phi)$ for an arbitrary formula ϕ , given the belief base B , which is done using Equation 4, thus giving $C = O(\|\text{Img}(B)\|)$. Furthermore, we may observe that $\|\text{Img}(B)\| \leq \|B\| = m$ and $\|\text{Img}(D)\| \leq \|D\| = n$. Therefore, we may conclude that carrying out the syntactic goal election requires $O(m^2n)$ entailment checks.

The termination of Algorithms 1 and 4 is proved in [5]; the termination of Algorithms 2 and 6 is a direct consequence of their termination.

Example (continued) We may now apply Algorithm 6 to elect the goals of Dr. Gent, given that his belief base is the one given in Equation 6 and his desire base is the one given in Equation 13; therefore, $\text{Img}(B) = \{0.2, 0.75, 0.9, 1\}$ and $\text{Img}(D) = \{0.75, 1\}$.

We begin by calling Algorithm 5 with $\bar{\gamma} = 0$: δ is set to $\max \text{Img}(D) = 1$, and the corresponding δ -cut of D is in fact the core of D , $D_1 = \{(r, 1)\}$. Now, $\mathbf{B}(\neg r) = 0 \leq \bar{\gamma}$; therefore $G_0 = \{r\}$ and Algorithm 6 terminates immediately with $\gamma^* = 1 - \bar{\gamma} = 1$, $G^* = G_0 = \{r\}$, i.e., Dr. Gent will elect as his goal just to register to ECAI 2014.

8 Conclusions

We have proposed a syntactic representation for desires within a possibilistic BDI model of agency; we have shown its equivalence to the semantic representation based on a qualitative utility; we have provided the syntactic equivalent of the deliberation algorithm, which generates the set of justified desires given a set of desire-generation

rules and a belief base. We have then provided the syntactic equivalent of the goal election algorithm, which generates the goals as the maximally justified desires among the maximally possible desires.

The cost of computing the degree of justification of a formula and of electing the goals has been given in terms of the basic operation of checking whether a formula entails another formula. Even though the cost of the former task grows exponentially with to the size of the desire base, in practice it is expected to be feasible, given that the size of the desire base depends on the number of desire-generation rules and it is hard to think of applications that would call for a large number of such rules. On the other hand, goal generation is polynomial in the size of the belief and desire bases. Compare this to what happens when using the semantic representation, where $2^{\|\mathcal{C}\|}$ interpretations have to be explicitly represented and iterated over to compute the degree of justification of a formula. In our opinion, a syntactic computation of beliefs, desires, and goals is the only viable alternative to implement possibilistic agents based on languages whose semantics involve an infinite number of possible worlds. This is why we believe the results here presented are a first and important step towards the practical implementation of a possibilistic BDI framework.

REFERENCES

- [1] *The Description Logic Handbook: Theory, implementation and applications*, eds., Franz Baader, Diego Calvanese, Deborah McGuinness, Daniele Nardi, and Peter Patel-Schneider, Cambridge, 2003.
- [2] S. Benferhat, D. Dubois, H. Prade, and M-A. Williams, ‘A practical approach to revising prioritized knowledge bases’, *Studia Logica*, **70**(1), 105–130, (2002).
- [3] S. Benferhat and S. Kaci, ‘Logical representation and fusion of prioritized information based on guaranteed possibility measures: application to the distance-based merging of classical bases’, *Artif. Intell.*, **148**(1–2), 291–333, (2003).
- [4] S. Benferhat and S. Kaci, ‘Logical representation and fusion of prioritized information based on guaranteed possibility measures: Application to the distance-based merging of classical bases’, *Artificial Intelligence*, **148**(1–2), 291–333, (2003).
- [5] C. da Costa Pereira and A. Tettamanzi, ‘Belief-goal relationships in possibilistic goal generation’, in *ECAI 2010*, pp. 641–646, (2010).
- [6] C. da Costa Pereira and A. Tettamanzi, ‘A syntactic possibilistic belief change operator for cognitive agents’, in *IAT 2011*, pp. 38–45, (2011).
- [7] D. Dubois, J. Lang, and H. Prade, ‘Possibilistic logic’, in *Handbook of Logic in Artificial Intelligence and Logic Programming*, 439–513, Oxford University Press, New York, (1994).
- [8] D. Dubois and H. Prade, ‘An overview of the asymmetric bipolar representation of positive and negative information in possibility theory’, *Fuzzy Sets Syst.*, **160**(10), 1355–1366, (2009).
- [9] H. Dubois, P. Hajek, and H. Prade, ‘Knowledge-driven versus data-driven logics’, *Journal of Logic, Language and Information*, **9**(1), 65–89, (2000).
- [10] P. Gärdenfors, *Knowledge in Flux: Modeling the Dynamics of Epistemic States*, MIT Press, 1988.
- [11] M. R. Garey and D. S. Johnson, *Computers and Intractability: A guide to the Theory of NP-Completeness*, Freeman, New York, 1979.
- [12] M. Hanheide, N. Hawes, J. L. Wyatt, M. Göbelbecker, M. Brenner, K. Sjöö, A. Aydemir, P. Jensfelt, H. Zender, and G.-J. M. Kruijff, ‘A framework for goal generation and management’, in *AAAI Workshop on Goal-Directed Autonomy*, (2010).
- [13] B. Nebel, ‘A knowledge level analysis of belief revision’, in *Proceedings of the KR’89*, pp. 301–311, San Mateo, (1989).
- [14] A. Pokahr, L. Braubach, and W. Lamersdorf, ‘A goal deliberation strategy for bdi agent systems.’, in *MATES*, volume 3550, pp. 82–93, (2005).
- [15] J. Thangarajah, J. Harland, and N. Yorke-Smith, ‘A soft COP model for goal deliberation in a BDI agent’, in *CP’07*, (2007).
- [16] L. A. Zadeh, ‘Fuzzy sets’, *Information and Control*, **8**, 338–353, (1965).
- [17] L. A. Zadeh, ‘Fuzzy sets as a basis for a theory of possibility’, *Fuzzy Sets and Systems*, **1**, 3–28, (1978).

⁵ $N(A \cap B) = \min\{N(A), N(B)\}$.