

Belief revision: from theory to practice

ALDO FRANCO DRAGONI

Istituto di Informatica, University of Ancona, Italy

Abstract

Belief revision is the process of rearranging a knowledge base to preserve global consistency whilst accommodating incoming information. Early approaches to belief revision used symbolic model theoretic methods, considering the problem as one of changing a logical theory. More recent approaches have adopted qualitative syntactic methods, taking them into the area of truth maintenance systems, and numerical mathematical methods, thus moving into the mainstream literature of uncertainty management. Multi-agent systems, in which information may come from a variety of human and artificial sources with different degrees of reliability, seem to be a natural domain for belief revision. The aim of this paper is to give a synoptic perspective of this composite subject from the clear air of the high theoretical peaks down to the muddy plain of practical algorithms.

1. Introduction

The body of beliefs (facts and rules) accumulated in the course of time by a knowledge-based system interacting with a complex and dynamic world is destined to evolve. Information may be perceived directly from the environment, or may come from a variety of human or artificial sources with different degrees of reliability. Some of these pieces of information integrate and corroborate the previously held corpus of sentences about the world, but others might cause serious conflicts with the established knowledge. In this case, the eventual acquisition of the new evidence should be accompanied by a partial or total reduction of the credibility of the conflicting pieces of knowledge. If the system's collection of beliefs is not a flat set of facts but contains rules, finding such conflicts and determining all the sentences involved in the contradictions can be hard because knowledge is only partially explicit.

Since the seminal, philosophical and influential works of Gärdenfors et al. (Alchourrón et al., 1985; Gärdenfors, 1988), the ideas on “belief revision” have been progressively refined (Gärdenfors, 1992) and ameliorated toward normative, effective and quasi-computable paradigms (Benferhat et al., 1993; Nebel, 1994; Williams, 1995).

We will try to give a unitary perspective of this composite subject, starting from the foundational papers on belief revision and updating (section 2). We will see how this ideal line joined the pragmatic framework of the “reason maintenance” approaches to deal with finite sets of sentences (section 3), and how different they are from the classical numerical methods which conceive belief revision as *uncertainty* revision (section 4).

We conclude by introducing a model for belief revision in a multi-source environment capable of being employed in practical systems (section 5). Its main characteristics are that it:

1. links assumption-based (symbolic) reasoning and uncertainty-management (numerical) techniques;
2. substitutes the “Priority of the Incoming Information” principle with the “Recoverability” principle.

2. Belief revision and updating

2.1 Preliminaries

If a contradictory knowledge space has always to be revised, then the contradiction must always be detectable. This means that we have to adopt a decidable language to represent beliefs. This constraint could be relaxed, for instance by considering “weakly-consistent” a set of beliefs from which a contradiction has never been derived (Dragoni, 1993). However, following the majority of the authors and for didactic reasons, along the paper we will use a propositional language L to represent knowledge. Sometimes we will abandon the original terminology in force of the following notational conventions:

$\wedge \vee \neg$: standard connectives

Ξ : set of the propositional letters of L

K : a subset of L

Ω : set of the possible interpretations for L

$\omega \models K$ ($\omega \not\models K$): the interpretation $\omega \in \Omega$ is a model (not a model) for K

$[p]$: set of the models of p

If $K = \{p_1, p_2, \dots, p_n\}$ then $[K] = [p_1] \cap [p_2] \cap \dots \cap [p_n]$ ($= [p_1 \wedge p_2 \wedge \dots \wedge p_n]$)

$K \vdash p$ ($K \not\vdash p$): p is (is not) a logical consequence of K , i.e., $[K] \subseteq [p]$ ($[K] \not\subseteq [p]$)

$\text{Th}(K)$: deductive closure of K

If $K = \{p_1, p_2, \dots, p_n\}$ then $K \downarrow p$ is the set of the maximal subsets $K' \subseteq K$ such that $[K'] \not\subseteq [p]$.

If the new information reports of some modification in the current state of a dynamic world, then the consequent change in the representation of the world is called “updating”. If the new information reports of new evidence regarding a static world whose representation was approximate, incomplete or erroneous, then the corresponding change is called “revision”.

2.2 Sentence-based revision

Alchourrón et al. (1985) introduced rational principles and postulates to whom changes of logical theories should obey. The *knowledge space* K is a deductively closed set of sentences of L . A finite subset B of K such that $K = \text{Th}(B)$ is a *knowledge base* (or *base*) for K . The authors distinguish between three kinds of change of K caused by the arrival of the sentence p :

- *Expansion*, K^+p , adds p and its logical consequences to K .
- *Contraction*, K^-p , deletes p from K ; this may involve the elimination of other sentences from K .
- *Revision*, K^*p , adds p and its logical consequences to K restoring the consistency of the new knowledge space; when p is inconsistent with K some sentences must be retracted from K .

Expansion is defined simply by $K^+p = \text{Th}(K \cup \{p\})$. When $(K \cup \{p\})$ is inconsistent, $\text{Th}(K \cup \{p\}) = L$ so that it is necessary to make a revision. K^*p and K^-p cannot be directly related with K because there can be many sets of sentences susceptible to be eliminated from K . Alchourrón et al. gave eight “rationality postulates” for each kind of change, inspired by the following three principles:

- *Consistency*: every change must yield a consistent knowledge space.
- *Minimal Change*: every change should alter as little as possible the knowledge space.
- *Priority to the Incoming Information*: incoming information always belongs to the revised knowledge space.

Here are the postulates for revision.

K*1. For each p and K , K^*p is still a knowledge space

K*2. $p \in K^*p$

K*3. $K^*p \subseteq K^+p$

K*4. If $\neg p \notin K$ then $K^+p \subseteq K^*p$

- K*5. $K*p$ is inconsistent iff p is inconsistent
 K*6. If p and q are logically equivalent then $K*p = K*q$
 K*7. $K*(p \wedge q) \subseteq (K*p)^+ q$
 K*8. If $\neg q \notin K*p$ then $(K*p)^+ q \subseteq K*(p \wedge q)$

K*2 reflects the principle of priority to the incoming information. K*3 and K*4 say that revision is just expansion when p is consistent with K . K*5 reflects the principle of consistency and allows the knowledge space K to be inconsistent. K*6 says that revision does not depend on syntax. K*7 and K*8 regard revision after composite information; they are not much intuitive (Nebel, 1989) and can be replaced by the following (Darwiche and Pearl, 1994):

$$C1. (K*p)*p \wedge q = K*(p \wedge q)$$

As Lehmann (1995) showed, “C1 is an elegant, slightly more powerful, version of the postulates K*7 and K*8”; the idea is that if one learns, first, some partial information (p) and then the full information ($p \wedge q$), the partial information should not make a difference.

These axioms describe the rational properties to which revision should obey, but they do not suggest how to perform it. A first operative definition (Gärdenfors, 1988) came from Lévy’s identity:

$$K*p = (K - \neg p)^+ p$$

that defines revision in terms of contraction.

Following the principle of minimal change, contraction can be characterized starting from the maximal subsets of K that fail to imply p . A subset K' of K is “maximal w.r.t. to the property of not implying p ” iff $p \notin \text{Th}(K')$ and for each K'' such that $K' \subset K'' \subseteq K$, $p \in \text{Th}(K'')$.

A first tentative approach to defining a contraction function is that of choosing only one element of $K \downarrow p$ by means of an appropriate selection function S . Formally, such “maxi-choice” contraction is defined:

$$K^- p = \begin{matrix} n \\ K \end{matrix} S(K \downarrow p) \quad \begin{matrix} \text{if } p \text{ is not a tautology} \\ \text{otherwise} \end{matrix}$$

Maxi-choice revision (revision based on the maxi-choice contraction) satisfies the first six postulates, but the last two only conditionally (see Gärdenfors, 1992). Furthermore, such a revision leads to a knowledge space that contains too many sentences, since, for every q , either $q \in K*p$ or $\neg q \in K*p$ (Lévy, 1994).

The “full-meet contraction” selects all the elements of $K \downarrow p$:

$$K^- p = \begin{matrix} n \\ K \end{matrix} (K \downarrow p) \quad \begin{matrix} \text{if } K \downarrow p \neq \emptyset \\ \text{otherwise} \end{matrix}$$

Full-meet revision yields a knowledge space that contains only the logical consequences of p !

A compromise between maxi-choice and full-meet contraction is the “partial-meet contraction”:

$$K^- p = \begin{matrix} n \\ K \end{matrix} S(K \downarrow p) \quad \begin{matrix} \text{if } K \downarrow p \neq \emptyset \\ \text{otherwise} \end{matrix}$$

Even the partial-meet revision unconditionally satisfies only the first six postulates, but the last two can be regarded as constraints on S . If the elements of $K \downarrow p$ are ordered by a preferential relation κ and S selects those maximal w.r.t. κ , then the partial-meet revision is defined “relational” and satisfies K*7. If κ is transitive, then the “relational-transitive” partial-meet revision satisfies K*8. Every revision satisfying all the eight postulates is a relational-transitive revision (Nebel, 1989).

2.3 The Epistemic Entrenchment

Every revision satisfying all the eight postulates relies on a preferential order over the formulae in Gärdenfors (1988) conceived an ordering κ_{EE} , called “Epistemic Entrenchment”, that envisages the

logical dependencies of the formulae in K ; it depends on K but it applies to all the formulae of L . $p \kappa_{EE} q$ means that p is less entrenched (more exposed to eventual changes) than q .

κ_{EE} satisfies the following postulates:

- EE1. κ_{EE} is transitive
- EE2. For all $p, q \in L$, if $p \vdash q$ then $p \kappa_{EE} q$
- EE3. For all $p, q \in L$, either $p \kappa_{EE} p \wedge q$ or $q \kappa_{EE} p \wedge q$
- EE4. If K is consistent, then $p \notin K$ iff $\exists q, p \kappa_{EE} q$
- EE5. If $\exists q \in L, q \kappa_{EE} p$ then p is a tautology

EE2 accounts for minimal change: if either q or p must be retracted from K , then it will be a smaller change to give up p rather than q since, to retract q , it will be necessary to abolish p too. Equivalent formulae are equally entrenched hence κ_{EE} is reflexive. Since, to retract $p \wedge q$, we need to give up at least one of the two, EE3 says that either p or q has the same degree of epistemic entrenchment than $p \wedge q$ (it must hold also EE2). As a consequence, every couple of formulae p and q are comparable (through $p \wedge q$). EE4 and EE5 give the lowest degree of entrenchment to formulae that do not belong to K and the highest one to the tautologies.

An epistemic entrenchment specifies a partial-meet contraction (and revision):

$$C^-: \quad q \in K^- p \text{ iff } q \in K \text{ and, either } p <_{EE} q \vee p, \text{ or } p \text{ is a tautology}$$

$K^- p$ contains only the formulae of K that have a greater degree of epistemic entrenchment than p .

Example. Suppose you see a man enjoying an ice cream. There are only two ice cream shops, A and B, so at least one of the two must be open. Let a stand for “A is open” and b for “B is open”. Suppose you see the lights on in A, so that you can presume a . It holds that:

$$\star <_{EE} a <_{EE} a \vee b \text{ since, } a \vdash a \vee b \neg$$

The knowledge space is $K = \text{Th}(a)$. Suppose now that you are informed that A is closed ($\neg a$). You need to retract a from K . Since $a \vee b$ is more entrenched than a it will not be retracted, so $K^- a = \text{Th}(a \vee b)$ and $K^* \neg a = \text{Th}(\neg a \wedge (a \vee b)) = \text{Th}(\neg a \wedge b)$.

From an implementational point of view, there are three problems with such a revision:

1. It does not directly deal with finite bases, which are the only practical representations of knowledge spaces.
2. κ_{EE} depends on K , so it is difficult to iterate the revision because the ordering defined on $K^* p$ could be different from the one defined on K .
3. The choice of a particular ordering κ_{EE} satisfying the postulates EE1 ÷ EE5 is arbitrary; as Gärdenfors wrote (1990b): “[the postulates] leave the main problem unsolved: what is a reasonable metric for comparing different epistemic states?”.

Section 3 presents and discusses some relevant approaches and problems with the definition of a revision operator for finite bases. Mary-Anne Williams (1995) showed how the first two problems can be solved (section 3.1). Regarding the last problem, the opinion implicitly expressed in sections 4 and 5 is that, practically, such an implementable and reasonable metric can be provided only by numerical approaches.

2.4 Updating

If the world that the knowledge space is aimed to represent is dynamic and the new information reports some changes in the state of the world, then the AGM postulates for revision are inappropriate.

Example. In the “ice-cream” example, suppose you had not seen the lights on in A, so that all you know is $K = \text{Th}(a \vee b)$. Suppose you see a man closing ice-cream shop A; the incoming

information p is again $\neg a$. From K^*3 and K^*4 , since p does not contradict K , $K^*p = K^+p = \text{Th}((a \vee b) \wedge \neg a)$. Since $b \in K^*p$, we can conclude that ice-cream shop B is open!. This result would be justified if the last information ($\neg a$) had completed the description of a static situation; but now the case is that someone changed the situation by closing A, so that the conclusion “B is open” is not acceptable.

From this simple example, it follows that the AGM postulates defining expansion as a particular case of revision do not apply to updating.

Katsuno and Mendelzon (1991a) proposed the following Postulates for Updating ($^\circ$):

- $K^\circ 1.$ $K^\circ p \vdash p$
- $K^\circ 2.$ If $K \vdash p$ then $K^\circ p$ is equivalent to K
- $K^\circ 3.$ If K and p are satisfiable then $K^\circ p$ is satisfiable
- $K^\circ 4.$ If K is equivalent to K' and p is equivalent to p' then $K^\circ p$ is equivalent to $K'^\circ p'$
- $K^\circ 5.$ $(K^\circ p) \wedge q \vdash K^\circ(p \wedge q)$
- $K^\circ 6.$ If $K^\circ p \vdash q$ and $K^\circ q \vdash p$ then $K^\circ p$ is equivalent to $K^\circ q$
- $K^\circ 7.$ If K is complete then $(K^\circ p) \wedge (K^\circ q) \vdash K^\circ(p \vee q)$
- $K^\circ 8.$ $(K \vee K')^\circ p$ is equivalent to $(K^\circ p) \vee (K'^\circ p)$

$K^\circ 2$ is weaker than K^*3 and K^*4 , since it imposes that $K^\circ p$ is equivalent to K only when p is *derivable* from K (which therefore remains unchanged), not simply when p is *consistent* with K . $K^\circ 3$ states only that $K^\circ p$ is consistent if K is consistent, while its AGM counterpart states that the revised knowledge space is consistent even when the original one is inconsistent, i.e., revision must restore consistency, while updating is requested only to preserve it.

These differences are uninfluential when K is a *complete* description of the world, i.e., when for any sentence p of L , either $K \vdash p$ or $K \vdash \neg p$.

2.5 Model-based revision

Revision can also be described at the semantic level: “from $[K]$ and $[p]$, find $[K^*p]$.” The Consistency Principle imposes that $[K^*p] \neq \emptyset$. The Priority to the Incoming Information imposes that $[K^*p] \subseteq [p]$. Minimal Change here means that some kind of distance between the elements of $[p]$ and the elements of $[K]$ should be minimized, i.e., $[K^*p]$ must contain only the elements of $[p]$ which are “closest” to those of $[K]$ according to some reasonable metric.

Maxi-choice revision selects only one model of p , i.e. $[K^*p] = \omega_p$, where $\omega_p \in [p]$. Full-meet revision selects all the models of p , i.e., $[K^*p] = [p]$. Partial-meet revision is a compromise between the two: in general, it selects more than one model of p , but not all of them.

Example. In the previous example, Ξ contains only two propositional letters, a and b . K contains $a \vee b$ and a , i.e., $K = \text{Th}(a)$. The incoming information p is $\neg a$.

	ω_1	ω_2	ω_3	ω_4
Ω	a, b	$a, \neg b$	$\neg a, b$	$\neg a, \neg b$
$[K]$				
$[p]$				
maxi-choice 1 $[K^*p]$				
maxi-choice 2 $[K^*p]$				
full-meet $[K^*p]$				

Partial-meet revision can be one among maxi-choice 1, maxi-choice 2 and full-meet, yielding, respectively, $K^*p = \text{Th}(\neg a \wedge b)$, $K^*p = \text{Th}(\neg a \neg b)$ and $K^*p = \text{Th}(\neg a)$.

Katsuno and Mendelzon (1991a, b) showed that every revision satisfying all eight AGM postulates implies the existence of a total preordering κ_K on Ω such that all the models of K have the same highest priority (lowest position in κ_K) w.r.t those interpretations that do not belong to $[K]$, and:

$$[K^*p] = \{\omega \in [p] \mid \omega \text{ is minimal in } \kappa_K \upharpoonright [p]\} \stackrel{\text{def}}{=} \min([p], \kappa_K)$$

i.e., revision selects the models of p which are closest to those of K w.r.t. κ_K . κ_K sorts interpretations while κ_{EE} orders sentences, but they are very closely related:

- both depend on K
- κ_K can be defined from κ_{EE} :
 $\omega \kappa_K \omega' \prec \neg p \kappa_{EE} \neg q$, where ω is the only model of p and ω' is the only model of q
- κ_{EE} can be defined from κ_K :
 $p \kappa_{EE} q \prec$ for all ω and ω' such that $\omega \in \min([p], \kappa_K)$ and $\omega' \in \min([q], \kappa_K)$, $\omega \kappa_K \omega'$

Katsuno and Mendelzon also showed that revision satisfies $K^{\circ}1$ – $K^{\circ}8$ (i.e., is an updating) iff, for each interpretation ω in $[K]$, there exists a preordering κ_{ω} over Ω such that:

$$[K^{\delta}p] = \min_{\omega \in [K]}^b([p], \kappa_{\omega})$$

where $\min([p], \kappa_{\omega})$ represents the models of p which are minimal w.r.t. κ_{ω} .

While κ_K depends on K , κ_{ω} depends on the particular interpretation ω . Chou and Winslett (1994, p. 160) define this operation as *local* in the sense that each model ω of K is updated independently of the others. Winslett's (1988) Possible Model Approach gives an example of such preorderings. Given a model ω of K , the interpretations ω' in Ω are sorted according to the set $\text{Diff}(\omega, \omega')$ of propositional letters on whose truth values ω' disagrees with ω : $\omega' \kappa_{\omega} \omega''$ iff $\text{Diff}(\omega, \omega') \rho \text{Diff}(\omega, \omega'')$.

Example. In the previous example, Ω has four interpretations, $\omega_1, \omega_2, \omega_3$ and ω_4 , with the following difference sets:

$$\begin{aligned} \text{Diff}(\omega_1, \omega_2) &= \text{Diff}(\omega_3, \omega_4) = \{b\} \\ \text{Diff}(\omega_1, \omega_3) &= \text{Diff}(\omega_2, \omega_4) = \{a\} \\ \text{Diff}(\omega_1, \omega_4) &= \text{Diff}(\omega_2, \omega_3) = \{a, b\} \end{aligned}$$

K has three models, ω_1, ω_2 and ω_3 , so there are three possible preorderings: $\kappa_{\omega_1}, \kappa_{\omega_2}$ and κ_{ω_3} :

$$\begin{aligned} \kappa_{\omega_1} &= \omega_1, \{\omega_2, \omega_3\}, \omega_4 \\ \kappa_{\omega_2} &= \omega_2, \{\omega_1, \omega_4\}, \omega_3 \\ \kappa_{\omega_3} &= \omega_3, \{\omega_1, \omega_4\}, \omega_2 \end{aligned}$$

where $\{\}$ means that the interpretations enclosed are not comparable. Since $[p] = \{\omega_3, \omega_4\}$:

	ω_1	ω_2	ω_3	ω_4
Ω	a, b	$a, \neg b$	$\neg a, b$	$\neg a, \neg b$
$\min(\{\omega_3, \omega_4\}, \kappa_{\omega_1})$				
$\min(\{\omega_3, \omega_4\}, \kappa_{\omega_2})$				
$\min(\{\omega_3, \omega_4\}, \kappa_{\omega_3})$				
$[K^{\circ}p]$				

i.e., $K^{\circ}p = \text{Th}(\neg a)$; the fact that B is open (b) cannot be derived.

Another quite reasonable local criterion to order interpretations is the *cardinality* of the difference. Given a model ω of K , the interpretations ω' in Ω are sorted according to the cardinality of $\text{Diff}(\omega, \omega')$: $\omega' \kappa_{\omega} \omega''$ iff $|\text{Diff}(\omega, \omega')| < |\text{Diff}(\omega, \omega'')|$.

Example. In the previous example:

$$\begin{aligned} |Diff(\omega_1, \omega_1)| &= |Diff(\omega_2, \omega_2)| = |Diff(\omega_3, \omega_3)| = |Diff(\omega_4, \omega_4)| = 0 \\ |Diff(\omega_1, \omega_2)| &= |Diff(\omega_3, \omega_4)| = |Diff(\omega_1, \omega_3)| = |Diff(\omega_2, \omega_4)| = 1 \\ |Diff(\omega_1, \omega_4)| &= |Diff(\omega_2, \omega_3)| = 2 \end{aligned}$$

There are three possible non-strict orderings:

$$\begin{aligned} \kappa_{\omega_1} &= \omega_1, (\omega_2, \omega_3), \omega_4 \\ \kappa_{\omega_2} &= \omega_2, (\omega_1, \omega_4), \omega_3 \\ \kappa_{\omega_3} &= \omega_3, (\omega_1, \omega_4), \omega_2 \end{aligned}$$

where $()$ means that the interpretations enclosed hold the same position. Apart from the difference between $\{\}$ and $()$, these are the same as before; even under the cardinality criterion, $K^\circ p = Th(\neg a)$.

If the incoming information had been $[p] = \{\omega_2, \omega_3\}$ or $[p] = \{\omega_1, \omega_4\}$, then we would have to select more than one interpretation for each κ_ω . Prioritizing the propositional letters in Ξ according to some degree of “importance”, can be useful to embed these preorderings into a total and strict order.

Updating is monotonic: if $K\omega \ K'$ then $K^\circ p\omega \ K'^\circ p$.

2.6 Foundations vs coherence theories of belief revision

The distinction between coherence and foundations approaches to belief revision (Doyle, 1992) relies on the different meaning and importance attributed to justifications (Dixon and Foo, 1993). According to Harman (1986), foundational belief revision presupposes that some beliefs depend on others for their justification; these can in turn depend on others and so on, until the foundational assumptions which depend only on themselves. The revision consists in eliminating all the beliefs that are no longer sufficiently justified, and in the addition of new beliefs that are either introduced as basic assumptions or are justified by other justified beliefs.

Coherence-based belief revision does not rely on such a lattice of justifications. A piece of knowledge is believed simply if there are no reasons to doubt it. Justifications are necessary only for beliefs that are challenged by new conflicting evidence or by competing more coherent, complete or well established theories. The AGM one is an example of such a revision.

The archetypes of foundational belief revision systems are the “Truth (or Reason) Maintenance Systems”, as the (J)TMS of Doyle (1979) and the ATMS by de Kleer (1986). The main difference between the two is that the former works with a single set of assumptions at a time, while the latter manages all the alternative sets of assumptions at the same time.

Harman and Gärdenfors support the coherence approaches to belief revision; according to them, the foundational systems do not reflect the way we revise our opinions and are too computationally expensive. Doyle (1992) notices that the two approaches are not so different as it would seem at a first glance. The following are the main criticisms advanced by Harman and Gärdenfors, with the corresponding refutations by Doyle:

- Humans seldom trace the reasons for their beliefs and often maintain them even if the bases are destroyed. Doyle controverts that even if the TMS’s way of reasoning is not anthropomorphic, it is a mechanism computationally feasible that produces rather reasonable, useful and practical results.
- Harman complains that the foundation approaches do not respect his “principle of conservatism”: “one is justified in continuing fully to accept something in the absence of a special reason not to”. Doyle seems to not accept this principle of rationality and notices that Gärdenfors adopts the term “conservatism” to mean “minimal change”. In this sense, even the (A)TMSs are conservative since they minimize the set of beliefs to give up.
- Gärdenfors defends the coherence approaches from the accusation of not being able to give explanations. He advocates the epistemic entrenchment as a way to extract the reasons for the

beliefs. Unfortunately, Gärdenfors' ordering is not so able in that, in particular, Doyle points out that it cannot manage adequately multiple reasons for a same belief: if p and q are independent but equally entrenched beliefs, and both of them explain r , then $K^* \neg p$ gives up also q , barring the evaluation of the alternative explanations.

- Gärdenfors says that keeping a trace of all the derivations and using them to revise beliefs is too costly for an automatic reasoner. Furthermore, changing only the basic assumptions is too restrictive, since nothing ensures that these are more important than the derived sentences. Doyle confutes that coherence approaches work with deductive closures which are much more intractable than a finite number of derivations. Even the epistemic entrenchment should be defined over the entire language L (this point has been refuted by Williams when she proved that ordering a finite set of assumptions implies defining a *finite* epistemic entrenchment over the entire language—see below). It is reasonable taking into account only a finite number of basic assumptions, since it is these that influence the entire corpus of beliefs.

3. Revision for finite bases

AGM revision respects Dalal's (1988) "principle of irrelevance of the syntax", according to whom syntactically different but logically equivalent formulae represent the same knowledge space. Since the calculus of deductive closures is infeasible, the partisans of "syntax-dependent" belief revision consider knowledge spaces made up of a limited number of sentences. They claim that asserting facts is more important than deriving others from them.

3.1 Partial entrenchment rankings

Bernard Nebel devoted a substantial part of his studies to the definition of a revision applicable to finite sets of sentences (1989, 1991, 1992, 1994). The knowledge space is defined by one of its bases B , which is a consistent, but not deductively closed, set of sentences. This line (which seems influenced by the ATMS) shows both coherence and the foundation aspects, coherence-based revision is applied to the original pieces of knowledge which appear in the base; while foundation-based revision is (implicitly) applied to the pieces of knowledge which will eventually be derived as logical consequences of the selected base(s).

Nebel (1989) redefines for such a base the operations of contraction $^{\xi}$ and revision $^{\wedge}$:

$$B^{\xi} p = \begin{array}{c} \begin{array}{c} \text{\tiny 8 } \psi \\ \text{\tiny } \omega \\ \text{\tiny } W \\ \text{\tiny } C \\ \text{\tiny } C \in (B \vartheta p) \end{array} \\ \text{\tiny } \approx \\ \text{\tiny } \approx \end{array} \begin{array}{c} \wedge (B \vee \neg p) \\ B \end{array} \quad \begin{array}{l} \text{if } p \text{ is not a tautology} \\ \text{otherwise} \end{array}$$

$$B^{\wedge} p = (B^{\xi} \neg p) \wedge p$$

He showed that the contraction $^{\xi}$ over B is equivalent to a partial-meet contraction on $K = \text{Th}(B)$ where the selection function $S_B(K \downarrow p)$ takes only the elements of $(K \downarrow p)$ that contain as many elements of B as possible. In this case, $K^- p = \text{Th}(B) \neg p = \text{Th}(B^{\xi} p)$. $^{\xi}$ is relational, for the relation $X \angle Y \stackrel{\text{def}}{=} (X \vartheta B) \theta (Y \vartheta B)$, so $^{\xi}$ satisfies the postulate K^*7 . However, this relation is not complete. Let $<$ be a total ordering on the propositions of B . We can embed $\angle$ in $<$ defining a new total ordering $\angle'$: when X and Y are incomparable by set-inclusion, we assume Y to be larger than X by $\angle'$ iff there is an element in Y which is larger by $<$ than any element in X which is not in Y . A partial-meet contraction defined by using the selection function

$$S_{B,\kappa}(K \downarrow p) \stackrel{\text{def}}{=} \{C \in (K \downarrow p) \mid \text{\tiny 8 } C' \in (K \downarrow p): (C' \vartheta B) \angle' (C \vartheta B)\}$$

satisfies all the eight rationality postulates. If the ordering $<$ is strict, then the intersection of the elements of $(K \downarrow p)$ selected by $S_{B,\kappa}$ coincides with a sole element of $(B \downarrow p)$, so the partial meet on K

corresponds to a maxi-choice on B . The analogous of maxi-choice applied to bases does not yield the undesirable result that it does for knowledge spaces.

Two problems are left open.

- Revision is sensitive to syntactic changes of the formulae.
- The revision's outcome $B' = (B^\sharp \neg p) \wedge p$ is a single formula, not a set of formulae as the initial base, so it is difficult to iterate the revision. If, subsequently, it would be necessary to revise B' on the light of q , and $B' \vdash \neg q$, then $B' \vartheta \neg q$ would be the empty set and $B'' = (B' \vee \neg q) \wedge q = (B' \vee q) \wedge (\neg q \vee q) = B' \vee q$. Since $B' \vdash \neg q$, we would be compelled to accept q by rejecting B' .

Mary-Anne Williams (1995) showed that both these problems “can be solved in theoretically satisfying ways wholly within the AGM paradigm”. She pointed out that, while the AGM postulates do not uniquely determine a revision, an epistemic entrenchment relation does. Belief revision means “epistemic entrenchment revision”: the newcoming information *transmutes* the old epistemic entrenchment into a new one which, in turn, yields a different contraction and, hence, a new revised knowledge space. She defined *partial entrenchment ranking* (p.e.r.) on a base B in L , a function $B: B\omega \{0, 1, \dots, \omega\}$ such that:

- PER1. $\{q \in B \mid B(p) < B(q)\} \not\vdash p$
 PER2. $B(p) = 0$ if p is a contradiction
 PER3. $B(p) = \omega$ iff p is a tautology

Let B_+ denote the set of the consistent sentences in B , and let $K_+ = \text{Th}(B_+)$. $\text{cut}(p)$ is defined as the subset of B_+ made of all the sentences q such that $\{r \in B_+ \mid B(r) > B(q)\} \not\vdash p$. Williams proved that a p.e.r B (on B) can be regarded as a specification of a finite epistemic entrenchment on K_+ which is an ordering κ_{EE} such that $p \kappa_{EE} q$ iff either $p \notin K_+$ or $\text{cut}(q) \theta \text{cut}(p)$.

The *degree of acceptance* of a sentence p w.r.t. B , called $\text{degree}(p)$, is defined to be the largest j such that $\{q \in B_+ \mid B(q) \leq j\} \vdash p$ if $p \in K_+$, $\text{degree}(p) = 0$ otherwise. It can be computed by the following algorithm:

$\text{degree}(p) := \omega + 1$
 $\text{do } \text{degree}(p) := \text{degree}(p) - 1$
 $\text{until } \text{degree}(p) = 0 \text{ or } \{q \in B \mid \text{degree}(p) \leq B(q)\} \vdash p$

AGM revision can be iterated by transmuting partial entrenchment rankings. The transmutation is provoked by the income of a consistent and not-tautological sentence p with a degree of firmness i . A transmuted p.e.r. $B_{(p,i)}^A$ is defined to be one such that:

1. $B_{(p,i)}^A(p) = i$
2. $K_+^A = \text{Th}(\{q \in B_+ \mid \text{degree}(\neg p) < B(q)\} \vdash \{p\})$

She introduced the following particular transmutation, called *adjustment*:

$$B_{(p,i)}^- = \begin{cases} \rho B_{(p,i)}^- & \text{if } i < \text{degree}(p) \\ (B_{(\neg p, i)}^-)^+_{(p,i)} & \text{otherwise} \end{cases}$$

where

$$B_{(p,i)}^-(q) = \begin{cases} \rho i & \text{if } B(q) > i \text{ and } \text{degree}(p) = \text{degree}(p \vee q) \\ B(q) & \text{otherwise} \end{cases}$$

for all $q \in B$, and

$$B_{(p,i)}^+(q) = \begin{cases} \neg B(q) & \text{if } B(q) > i \\ \neg \text{degree}(\neg p \vee q) & \text{if } B(q) \leq i < \text{degree}(\neg p \vee q) \text{ or } p \eta q \\ \neg \text{degree}(\neg p \vee q) & \text{otherwise} \end{cases}$$

for all $q \in (B \setminus \{p\})$. Intuitively, an adjustment involves minimal changes to B such that p is accepted with degree i .

Example.

B	$degree$	$B_{(p,3)}^\Lambda$	$B_{(p,r)}^\Lambda$	$B_{(p,4)}^\Lambda$
$\neg p \vee q$ 3	$\neg p \vee p$ 3	p 3	$\neg p \vee q$ 3	$\neg p$ 4
r 2	r 2	$\neg p \vee q$ 3	r 2	$\neg p \vee q$ 4
p 1	p 1	q 3	$\neg p \vee r$ 2	$\neg p \vee r$ 4
(a)	q 1	r 2	p 0	r 2
	$\neg p \vee r$ 2	$\neg p \vee r$ 2	q 0	q 0
	$\neg p$ 0	$\neg p$ 0	$\neg p$ 0	p 0
	$\neg q$ 0	$\neg q$ 0	$\neg q$ 0	$\neg q$ 0
	$\neg r$ 0	$\neg r$ 0	$\neg r$ 0	$\neg r$ 0
	(b)	(c)	(d)	(e)

- (a) a finite partial entrenchment ranking on a base in L
- (b) the corresponding degree of acceptance of some formulae of L
- (c) a transmutation: more evidence for p is acquired so that we decide to increase the degree of acceptance of p from 1 to 3
- (d) a transmutation: the contraction of p
- (e) a transmutation: the acceptance of $\neg p$ with degree 4.

Neither partial entrenchment rankings nor transmutations are syntax-dependent. The problem with them is a matter of complexity; Williams (1995) wrote “a full complexity analysis [of adjustment] is yet to be conducted”, but actually, we would be very surprised if it would be less than exponential in the size of the base.

3.2 Ordering sentences and sets of sentences

Nebel (1994) considers revision schemes with preference information that has a size polynomial in the size of the base. An “epistemic relevance ordering” is an ordering κ_{ER} that stratifies a base B into n priority classes $B_1, \dots, B_n$ as follows:

$$q \in B_1 \text{ iff } \exists p \in B : q \kappa_{ER} p,$$

$$q \in B_{i+1} \text{ iff } \exists p \in B_i : (p <_{ER} q \wedge \neg(\exists s \in B : p <_{ER} s <_{ER} q)).$$

The intuitive meaning of $p \kappa_{ER} q$ is that q is at least relevant, important or credible as p . κ_{ER} does not respect the logical contents of the sentences as κ_{EE} and partial entrenchments do. The justification seems to rely on the logical paradoxes of the material implication: a rule $\neg q \vee p$ should not necessarily be considered more important than p just because $p \vdash \neg q \vee p$.

Let $\overline{B}_j$ be the union of the highest strata down to the j th one.

Nebel defines $B \rightarrow p$ as the set of the maximal subsets of B that fail to imply p and *contain as many sentences of the highest priority as possible*:

$$B \rightarrow p = fB' \theta B_j B' \not\vdash p, \exists C, j \neg(B' \sqsubset \overline{B}_j \rho C \sqsubset \overline{B}_j) \leftarrow C \sqsubset \overline{B}_j \vdash pg$$

Example. Consider the following stratified base B .

B_3	$(e \wedge a) \omega$	b
B_2	e	a
B_1	c	

The following are the elements of $B\vartheta b$.

	B_3	B_2	B_1
B^1	$(e \wedge a)\omega \ b$	e	c
B^2	$(e \wedge a)\omega \ b$	a	c
B^3	$(e \wedge a)\omega \ b$	e	
B^4	$(e \wedge a)\omega \ b$	a	
B^5	$(e \wedge a)\omega \ b$		c
B^6	$(e \wedge a)\omega \ b$		
B^7		$e \ a$	c
B^8		$e \ a$	
B^9		e	c
B^{10}		a	c
B^{11}		e	
B^{12}		a	
B^{13}			c

The set of the maximal elements of $B\vartheta b$ is $B \dashv b = \{B^1, B^2\}$.

Let $B = \{q_1, \dots, q_m\}$, with $m \leq n$, the base B where its m elements are sorted by κ_{ER} and the sentences in a same stratum occupy adjacent positions. The elements of $B \dashv p$ can be computed by repeatedly forcing chronological backtraking in the following algorithm:

```

 $B \dashv p := \emptyset$ 
for  $i = m$  to  $1$  do
   $B \dashv p := B \dashv p \cup \{q_i\}$  if  $B \dashv p \cup \{q_i\} \not\models p$ 
   $B \dashv p := B \dashv p$  otherwise

```

which is, basically, the same algorithm presented in Dragoni (1992).

The corresponding “prioritized meet base revision” $B \Phi p$, is defined as:

$$B \Phi p = \bigcup_{B' \in (B \dashv p)} Th(B')^{A^+ p}$$

There are two problems with this revision:

1. It does not satisfy all the AGM postulates.
2. It is still computationally hard.

Alternatively, one may conceive each stratum as made of a single formula, the conjunction of all its elements; in this case $B \dashv p$ contains only one element.

Example. For the base of the previous example, $B \dashv p = \{B^5\}$.

Given a base B , let an *argument* for p be a subset A of B such that $A \vdash p$ and all the sentences in A are necessary to derive p . The previous contraction removes the least relevant sentence in each argument for p , but sacrifices also the sentences that belong to the same stratum.

Nebel defines an other kind of revision, $B \Omega q$, called “cut base revision”:

$$B \Omega q = Th(\{p \in B \mid p \in B_j, \overline{B_j} \not\models \neg q\}^+ q)$$

This revision satisfies the first six postulates, but is more drastic than prioritized meet. It cuts away (hence the name?) all classes that have a priority equal or lower to the class that is “responsible” for an inconsistency. For this reason, one may argue that cut base revision violates the principle of minimal change.

In general, we could adopt various criteria to sort and select the elements of $B\vartheta p$. Let $B' = B'_1 \cup \dots \cup B'_n$ and $B'' = B''_1 \cup \dots \cup B''_n$ be two consistent subsets of B , where $B'_i = B'_i \cup B_i$ and

$B''_i = B'' \upharpoonright B_i$. Benferhat et al. (1993) consider three ways to translate κ_{ER} into a preference relation τ on 2^B :

- *best-out ordering*. $B'' \tau_S^{bo} B'$ iff the most credible of the sentences in $B \setminus B''$ is more credible than the most credible of the sentences in $B \setminus B'$. This ordering is complete and its maximal elements are all the consistent subsets of B that contain $\overline{B}_j$, where j is the lowest index such that $\overline{B}_j$ is consistent.
- *inclusion-based ordering*. $B'' \tau_S^{inc} B'$ iff $\exists i$ such that $B'_i \sigma B''_i$ and for any $j < i$, $B'_j = B''_j$. It is equivalent to that proposed by Brewka (1989). This preordering is strict but partial; its maximal consistent elements are of the form $A = A_1 \upharpoonright \dots \upharpoonright A_n$, where for every $i = 1, \dots, n$, $A_i \upharpoonright \dots \upharpoonright A_n$ is a maximally consistent subset of $B_i \upharpoonright \dots \upharpoonright B_n$. These elements are also maximal for τ_S^{bo} .
- *lexicographic ordering*. $B'' \tau_S^{lex} B'$ iff $\exists i$ such that $|B'_i| > |B''_i|$ and for any $j < i$, $|B'_j| = |B''_j|$, and $B'' = \tau_S^{lex} B'$ iff for any j , $|B'_j| = |B''_j|$. This preordering is complete; its maximal consistent elements are of the form $A = A_1 \upharpoonright \dots \upharpoonright A_n$, where for every $i = 1, \dots, n$, $A_i \upharpoonright \dots \upharpoonright A_n$ is a cardinality-maximally consistent subset of $B_i \upharpoonright \dots \upharpoonright B_n$.

$B \downarrow p$ contains the elements of $B \vartheta p$ maximal w.r.t. inclusion-based ordering.

Example. For the base of the previous example, $B \downarrow p$ is sorted as follows:

<i>best-out</i>	<i>inclusion-based</i>	<i>lexicographic</i>
$B^1, B^2, B^3, B^4, B^5, B^6$	$B^1 \quad B^1 \quad B^2 \quad B^2$	B^1, B^2
$B^7, B^8, B^9, B^{10}, B^{11}, B^{12}, B^{13}$	$B^3 \quad B^3 \quad B^4 \quad B^4$	B^3, B^4
	$B^5 \quad B^5 \quad B^5 \quad B^5$	B^5
	$B^6 \quad B^6 \quad B^6 \quad B^6$	B^6
	$B^7 \quad B^7 \quad B^7 \quad B^7$	B^7
	$B^8 \quad B^8 \quad B^8 \quad B^8$	B^8
	$B^{10} \quad B^9 \quad B^{10} \quad B^9$	B^9, B^{10}
	$B^{12} \quad B^{11} \quad B^{12} \quad B^{11}$	B^{11}, B^{12}
	$B^{13} \quad B^{13} \quad B^{13} \quad B^{13}$	B^{13}

These ways to sort sets of sentences do not respect their logical contents. As Benferhat et al. (1995) say: “we do not assume any (in)dependence relations between beliefs. . . [they] are put in the knowledge base as they are and as they come from their sources of information”. These authors gave also another characterization of the dichotomy between coherence and foundation theories. While coherence theories insist on restoring consistency, foundations theories accept inconsistency and cope with it. Coherence theories propose to give up some formulas of B in order to get one or several consistent sub-bases, and to apply classical entailment on these conclusions. Foundation theories retain all available information but each plausible conclusion inferred from B is justified by a strong reason for believing in it. In this sense, foundation theories are not revisions at all! They gave two example of such approaches, the *argumentation inference* and the *safely supported inference*. The argumentation inference suggests that q can be inferred from an inconsistent B if an argument A for q exists such that the least credible sentence in A is more credible than the least credible sentence in any other argument A' for $\neg q$ contained in B . The safely supported inference suggests that q can be inferred from an inconsistent B if a stratum j exists such that $\overline{B}_j$ contains an argument A for q such that its sentences are not involved in any of the contradictions that $\overline{B}_j$ (eventually) contains (although A may contain sentences contradicted by other sentences belonging to strata lower than j). The authors also propose an interesting insight into the concept of syntax dependency. Provided that these approaches are syntax-dependent, it could be interesting trying to establish how much they are! The authors propose four different degrees of syntax sensitivity. Unfortunately, both the proposed kinds of inference do not occupy a good position in that scale.

4. Revision as numerical treatments of uncertainty

The knowledge space of a complex knowledge-based system does not suffer only from inconsistency; it can also be affected by imprecision/vagueness (for instance; “The President has not been *prolific*”) and uncertainty (for instance: “*I believe that* the President has only a daughter”). Fuzzy logic (Zadeh, 1965; see also Dubois and Prade, 1990) is a well known way to deal with imprecise properties and relations. In this paper, we do not deal with imprecise/vague knowledge, but we survey some of the most important frameworks for reasoning under uncertainty.

Our goal is that of finding a reasonable way to sort sentences or interpretations according to their degree of certainty, in order to direct the change imposed by the inconsistencies. Numerical distributions of credibility over sentences or interpretations play the same role that κ_{EE} , κ_{κ} , p.e.r. and κ_{ER} play in the symbolic frameworks. Generally, numerical approaches do not respect logical dependencies among the sentences.

Logics of uncertainty often represent the knowledge space K and the incoming information p in terms of their models, called “possible worlds”. It is supposed that $[K]$ contains the world that corresponds to the real one, but we don’t know which world it is. However, not all the worlds in $[K]$ are equal candidates to represent the real one, hence each of them is associated with a “weight” $d(\omega)$ that expresses its degree of realism. Hence, a knowledge space is represented not simply by $[K]$, but by an assignment function $d(\omega): \Omega \rightarrow [0,1]$ such that $d(\omega')=0$ for each $\omega' \notin [K]$. The numerical assignment is an easy way to prioritize the interpretations.

The arrival of p normally means that the real world belongs to $A = [p]$. This arrival changes d into a new assignment (new prioritization) d' . Imposing the priority to the incoming information means assigning $d'(\omega)=0$ to each $\omega \notin A$. Minimizing this change means minimizing some kind of distance between d and d' .

This semantic way to represent knowledge spaces does not distinguish between foundational assumptions and derived sentences. Let $\Delta \rho \Omega$ be the set of worlds such that $d(\omega) > 0$. Any sentence $q \in L$, such that $[q] \theta \Delta$, is justified to belong to the knowledge space K , where $[K] = \Delta$. The believability of q depends on the assignment $d(\omega)$ for each $\omega \in [q]$, in a way that is specific of the particular formalism adopted. So, we could not regard these approaches as foundational, unless we see the models as the basic elements, the sentences as the derived ones and the logical entailment $\models$ as the (only) justification. In our opinion, this conception of foundations is all but unreasonable.

The coherence view of belief revision states that a belief should not be abandoned if there is no reasons to doubt of it, i.e., if it is compatible with the incoming information. Here the problem is that of establishing what a belief is intended to be: a *sentence* or a *model*. A sentence q is compatible with A if $[q] \cap A \neq \emptyset$. A model ω is compatible with A if $\omega \in A$. Consider the following simple example:

Example. $K = \{\neg a \vee b, a\}$ and the incoming information p is $\neg b$.

	ω_1	ω_2	ω_3	ω_4
Ω	a, b	$a, \neg b$	$\neg a, b$	$\neg a, \neg b$
$[a]$				
$[\neg a \vee b]$				
$[K] = [(\neg a \vee b) \wedge a]$				
$[p]$				
$[a \wedge p]$				
$[(\neg a \vee b) \wedge p]$				
$[K \wedge p]$				

K is incompatible with p , so we have to abandon it. However, if we regard K as a set of two beliefs, $\neg a \vee b$ and a , then coherence imposes to maintain one of the two since both are (separately) compatible with p . If K is regarded as a set of possible worlds, then the only belief is $a \wedge b$, which is not compatible with p and must be dropped.

None of the following formalisms can be regarded as a coherence sentence-based approach to

belief revision. All of them can be regarded as coherence model-based approaches to belief revision, since none of them put at zero the weight of a model of p .

4.1 The probabilistic framework

The Bayesian probabilistic approach has been the starting point, and still is a widely accepted reference, for the treatment of uncertainty (Pearl, 1988; Shafer and Pearl, 1990). The knowledge space is characterized by a *probability measure* P on 2^Ω , whose fundamental property is *additivity*: $8A, B \theta \Omega, A \uparrow B \neq \emptyset \leftarrow P(A \uparrow B) = P(A) + P(B)$. $P(\Omega) = 1$, so if $\bar{A} = \Omega - A$ then $P(A) + P(\bar{A}) = 1$. We might also consider the probability distribution $p(\omega)$ that assigns a probability degree to each world in Ω , where $P(A) = \sum_{\omega \in A} p(\omega)$. $p(\omega) = 0$ signifies that, certainly, ω is not a possible world. $p(\omega) = 1$ means that ω is surely the real world. An incoming information $A \rho \Omega$ changes the probability measure of any subsets B of Ω through the very famous Bayes' Conditioning Rule:

$$P(B|A) = \frac{P(B \uparrow A)}{P(A)} = \frac{P(A \uparrow B)P(B)}{P(A)}$$

which can also be expressed in terms of probability distribution:

$$p(\omega|A) = \frac{p(\omega)}{P(A)} \quad \text{if } \omega \in A \\ = 0 \quad \text{otherwise}$$

This modification is defined only for $P(A) > 0$, hence it is not applicable when A is judged impossible by the previously determined probability measure P . Bayesian conditioning obeys the principle of priority to incoming information; it increases the probability of the not impossible worlds belonging to A to the prejudice of those external to A which become impossible.

Example. Suppose it is known that in a basket there are only apples (a), only bananas (b) or only pears (c). $K = \{(a \wedge \neg b \wedge \neg c) \vee (\neg a \wedge b \wedge \neg c) \vee (\neg a \wedge \neg b \wedge c)\}$, $\Xi = \{a, b, c\}$ and Ω contains eight possible worlds. Suppose that the current knowledge space is characterized by the probability distribution $p(\omega)$ reported in the table. Let $\neg a$ be the incoming information, i.e. $A = \{\omega_4, \omega_6, \omega_7, \omega_8\}$. $P(\neg a) = P(A) = p(\omega_4) + p(\omega_6) + p(\omega_7) + p(\omega_8) = 0.5$, so the probability distribution is modified as reported in the row $p(\omega|A)$.

	ω_1	ω_2	ω_3	ω_4	ω_5	ω_6	ω_7	ω_8
ω	a,b,c	a,b, $\neg$ c	a, $\neg$ b,c	$\neg$ a,b,c	a, $\neg$ b, $\neg$ c	$\neg$ a,b, $\neg$ c	$\neg$ a, $\neg$ b,c	$\neg$ a, $\neg$ b, $\neg$ c
[K]								
$p(\omega)$	0	0	0	0	0.5	0.3	0.2	0
A								
$p(\omega A)$	0	0	0	0	0	0.6	0.4	0

As Katsuno and Mendelzon showed (section 2.3), a revision satisfying all the eight AGM postulates relies on the existence of a total preordering κ_κ on Ω such that all the models of K have the same highest priority and revision selects the elements of A which have the best position in κ_κ (hence are *closest* to those in [K]). $p(\omega)$ induces an ordering κ_κ on Ω , but it distinguishes the models of K ($\kappa_\kappa = \omega_5 < \omega_6 < \omega_7$ in the example) and gives the same lowest priority ($p(\omega) = 0$) to the interpretations outside [K]. Probabilistic revision selects *all* the elements of A which are in [K], not the highest in κ_κ (ω_6 in the example). It does not satisfy the AGM postulates, especially because it fails to address the notion of minimal change.

Lewis developed a probabilistic rule that deals with the principle of minimal change (see Dubois and Prade, 1994). It moves the “masses” associated to the worlds outside the newcoming information to those inside which are closest to them according to some metric defined over Ω . This means that if $\omega \notin A$, then its weight $p(\omega)$ is added to the world $\omega_A \in A$ closest to it. If $\omega \in A$, then

$\omega_A = \omega$. This rule, called “imaging” is defined as: $8\omega', p_A(\omega') = \sum_{\omega'=\omega_A}^P p(\omega)$. It can change an impossible world into a possible one, i.e., one may have $p_A(\omega) > 0$ while $p(\omega) = 0$, and it can turn a sure fact into an uncertain one, i.e., one may have $P_A(B) < 1$ while $P(B) = 1$.

Example. In the previous example, the only world outside A that belongs to $[K]$ (i.e., has a probability different from 0) is ω_5 . Adopting as metric the number of literals with different sign, i.e., the $|Diff(\omega_5, \omega)|$, among the worlds belonging to A the closest to ω_5 is ω_8 .

	ω_1	ω_2	ω_3	ω_4	ω_5	ω_6	ω_7	ω_8
$p(\omega)$	0	0	0	0	0.5	0.3	0.2	0
$p_A(\omega)$	0	0	0	0	0	0.3	0.2	0.5

Note that $p(\omega_8) = 0$ but $p_A(\omega_8) = 0.5$.

As Katsuno and Mendelzon showed (section 2.5), a revision is an updating (i.e., satisfies K^*1-K^*8) iff, for each interpretation ω in $[K]$, there exists a preordering κ_ω over Ω such that the revision selects only their respective maximal elements. Now, when a world ω in $[K]$ adds its weight to its “closest” world ω_A belonging to A , ω selects ω_A as the maximal element according to κ_ω , so imaging in an updating. In the previous example, the most probable world was that in which there were apples in the basket. After the incoming information reporting that there are no apples in the basket, the most probable world becomes that in which the basket is empty, not that in which there are bananas, as in the Bayesian conditioning case. This result is clearly in accord with the updating operation.

Jeffrey extended the Bayesian approach to the case that the newcoming information is uncertain, hence relaxing the principle of priority to the incoming information. $A\rho\Omega$ is attached with a probability α . The probability measures on 2^Ω are modified as follows:

$$P(Bj(A, \alpha)) = \alpha P(BjA) + (1 - \alpha)P(Bj\bar{A})$$

The worlds belonging to A become more probable but less ($\alpha < 1$) than in the certain case, while the worlds outside A that were not considered impossible become less probable, but still not impossible.

Example. Suppose that A is affected by a degree of uncertainty $\alpha = 0.7$. $\bar{A} = \{\omega_1, \omega_2, \omega_3, \omega_5\}$ and $P(\bar{A}) = 0.5$. The new probability distribution is $p_A(\omega|(A, \alpha))$.

	ω_1	ω_2	ω_3	ω_4	ω_5	ω_6	ω_7	ω_8
$p(\omega)$	0	0	0	0	0.5	0.3	0.2	0
$p_A(\omega (A, \alpha))$	0	0	0	0	0.3	0.42	0.28	0

The worlds that were considered impossible are still absolutely improbable, but the worlds outside the newcoming information still have a degree of probability.

Let us turn from the semantic connotation of knowledge spaces to the syntactic one. In the probabilistic framework the probability of a sentence p is simply the probability measure $P([p])$. So, for instance, in the last example the probability of $b \vee c$ (in the basket there are bananas or pears) after A is $P([b] \vee [c]) = 0.7$. $P([p]) = 0$ if p is a contradiction and $P([p]) = 1$ if p is a tautology.

It is interesting to see if the probabilistic approach generates an epistemic entrenchment. In effect, the probability measure is, obviously, transitive: if $P([p]) \kappa P([q])$ and $P([q]) \kappa P([r])$ then $P([p]) \kappa P([r])$. EE2 is verified since $p \vdash q$ means $[p] \theta [q]$ so that $P([p]) \kappa P([q])$ (it is always easier to retract p than q). EE4 is verified; in fact, $p \notin K$ means $P([p]) = 0$, and if K is consistent then there exist some impossible worlds so that $P([p]) = 0$ iff $8q(P([p]) \kappa P([q]))$. EE5 is verified since if $8q \in L$

$P([q]) \kappa P([p])$ then $[p] = \Omega$. Unfortunately, EE3 is generally unsatisfied since $[p \wedge q] \theta [p]$ and $[p \wedge q] \theta [q]$ so that $P([p \wedge q]) \kappa P([p])$ and $P([p \wedge q]) \kappa P([q])$; normally it is easier to retract a conjunction than any of its conjuncts. As an example, consider the probability distribution $p(\omega)$ of the first example in section 4.1. $[a] = \{\omega_1, \omega_2, \omega_3, \omega_5\}$, $[b] = \{\omega_1, \omega_2, \omega_4, \omega_6\}$, $[a \wedge b] = \{\omega_1, \omega_2\}$, $P([a]) = 0,5 > P([b]) = 0,3 > P([a \wedge b]) = 0$.

4.2 The possibilistic framework

In the possibility theory (Dubois and Prade 1990, 1991, 1992, 1994), the knowledge space is represented by a *possibility distribution* $\pi(\omega): \Omega \rightarrow [0,1]$. $\pi(\omega) > \pi(\omega')$ means that ω is more plausible than ω' . $\pi(\omega) = 0$ signifies that ω is an impossible world, while $\pi(\omega) = 1$ does not mean that ω is the real world, as in the probabilistic framework, but only that nothing hampers it to be so. In Ω there can be many worlds with $\pi(\omega) = 1$ and the property of additivity does not hold. A knowledge space is:

- *consistent* if $\exists \omega_0 \in \Omega$, $\pi(\omega_0) = 1$
- *complete* if $\exists \omega_0 \in \Omega$, $\pi(\omega_0) = 1$ and $\forall \omega \in \Omega$, $\omega \neq \omega_0 \rightarrow \pi(\omega) = 0$; ω_0 is the real world
- *vacuous* if $\forall \omega \in \Omega$, $\pi(\omega) = 1$; state of total ignorance
- *absurd* if $\forall \omega \in \Omega$, $\pi(\omega) = 0$; all the worlds are considered impossible.

A possibility distribution π is more specific than another π' if $\forall \omega \in \Omega$, $\pi(\omega) \leq \pi'(\omega)$ and $\exists \omega' \in \Omega$, $\pi(\omega') < \pi'(\omega')$.

Given a possibility distribution $\pi(\omega)$ we can define a *possibility measure* $\Pi(A) = \max_{\omega \in A} \pi(\omega)$ that satisfies the axiom $\Pi(A \cup B) = \max(\Pi(A), \Pi(B))$. It measures the degree of consistency of A with the knowledge space represented by $\pi(\omega)$. $\Pi(A) = 0$ means that A is impossible, $\Pi(A) = 1$ means that A is totally consistent with π , $\Pi(A) = \Pi(\bar{A}) = 1$ expresses the total ignorance about A and $\bar{A}$. The certainty degree of A is measured by the *necessity function* $N(A) = 1 - \Pi(\bar{A})$ whose characteristic axiom is $N(A \cap B) = \min(N(A), N(B))$. A is certain when $N(A) = 1$. The possibilistic approach captures the idea of total ignorance under the form of the vacuous knowledge space for which $\Pi(A) = 1 \forall A \subseteq \Omega$; this is not possible in the probabilistic framework, since even if we assign the same probability $p(\omega) = 1/|\Omega|$ to each world in Ω , the probability associated to each set A (i.e., to each proposition of the language) will be proportional to $|A|$.

When the incoming information $A \subseteq \Omega$ is consistent with $\pi(\omega)$, i.e. $\Pi(A) = 1$, the *expansion* yields the new possibility distribution $\pi_A^+(\omega) = \pi(\omega)$ if $\omega \in A$, $\pi_A^+(\omega) = 0$ otherwise. The worlds inside A maintain their values of possibility while the others are put to zero.

Example. Consider the possibility distribution $\pi(\omega)$ listed in the table. Let $\neg b$, i.e., $A = \{\omega_3, \omega_5, \omega_7, \omega_8\}$, be the incoming information. $\Pi(A) = 1$, so we need to perform a simple expansion, yielding the distribution $\pi_A^+(\omega)$ reported in the table:

	ω_1	ω_2	ω_3	ω_4	ω_5	ω_6	ω_7	ω_8
$\pi(\omega)$	0	0	0	0	1	0.8	0.7	0
$\pi_A^+(\omega)$	0	0	0	0	1	0	0.7	0

When the incoming information $A \subseteq \Omega$ is not consistent with $\pi(\omega)$ neither impossible, i.e., $0 < \Pi(A) < 1$, we need to perform the *revision*; this is defined as:

$$\forall B \subseteq \Omega, \Pi(A \cap B) = \Pi(B|A) \wedge \Pi(A)$$

where $*$ stands for the *min* or the product $\odot$.

If $*$ = *min*, then there are many solutions for $\Pi(B|A)$ and we choose the least specific one:

$$\begin{aligned} \Pi(B|A) &= 1 && \text{if } \Pi(A \cap B) = \Pi(A) > 0 \\ &= \Pi(A \cap B) && \text{otherwise} \end{aligned}$$

In particular, $\Pi(B|A) = 0$ if $A \nabla B = \emptyset$, i.e., as in the probabilistic framework, all the worlds outside A are considered impossible. The conditional necessity function is defined dually as $N(B|A) = 1 - \Pi(\overline{B}|A)$. In terms of possibility distribution the revision is defined as:

$$\begin{aligned} \pi(\omega|A) &= 1 && \text{if } \pi(\omega) = \Pi(A), \omega \in A \\ &= \pi(\omega) && \text{if } \pi(\omega) < \Pi(A), \omega \in A \\ &= ' && \text{if } \omega \in A \end{aligned}$$

If $* = \ominus$, then the revised possibility measures obey the equation:

$$\pi(B|\ominus), \quad \Pi(B|A) = \frac{\Pi(A \nabla B)}{\Pi(A)}$$

In terms of possibility distribution:

$$\begin{aligned} \pi(\omega|A) &= \frac{\pi(\omega)}{\Pi(A)} && \text{if } \omega \in A \\ &= ' && \text{otherwise} \end{aligned}$$

Both the rules require that $\Pi(A) \neq 0$.

Example. Consider the possibility distribution $\pi(\omega)$ listed in the table. Let $\neg a$, i.e., $A = \{\omega_4, \omega_6, \omega_7, \omega_8\}$, be the incoming information. $\Pi(A) = 0.8$, so we need to perform a revision, yielding the distribution $\pi_{min}(\omega|A)$ if $* = min$ and the distribution $\pi_{\ominus}(\omega|A)$ if $* = \ominus$.

	ω_1	ω_2	ω_3	ω_4	ω_5	ω_6	ω_7	ω_8
$\pi(\omega)$	0	0	0	0	0.9	0.8	0.7	0
$\pi_{min}(\omega A)$	0	0	0	0	0	1	0.7	0
$\pi_{\ominus}(\omega A)$	0	0	0	0	0	1	0.875	0

The possible worlds are the same, but their possibility values are greater in the latter case.

We can define the counterpart of imaging in the possibilistic framework; let $Cl(\omega') \ominus \overline{A}$ be the set of the worlds external to A closest to ω' :

$$\begin{aligned} \pi_A^{\delta}(\omega') &= \max_{\omega \in Cl(\omega') \ominus \overline{A}} \pi(\omega) && \text{if } \omega' \in A \\ &= ' && \text{otherwise} \end{aligned}$$

i.e., we assign to each $\omega' \in A$ the maximum value among its own one and the values of the worlds outside A which are closest to it.

Example. The example used for the imaging in the probabilistic framework works exactly the same here, since $\pi_A^{\delta}(\omega_8) = \max_{\omega \in \{\omega_5, \omega_8\}} \pi(\omega) = 0.5$.

The possibilistic approach can be extended to the case that the incoming information A is uncertain, i.e. $N(A) < 1$. Such uncertainty can be interpreted in two ways:

1. $N(A) = \alpha$ is regarded as a constraint for the revised distribution $\pi'(\omega)$, that must satisfy $N'(A) = \alpha$;
2. $N(A) = \alpha$ is regarded as an extra piece of information that may be useful or not to refine the current knowledge space; α is considered a degree of priority of A .

Interpretation 1 is in the spirit of Jeffrey's rule and its correspondent rule is defined as:

$$\begin{aligned} \pi(\omega|A, \alpha) &= \pi(\omega|A) && \text{if } \omega \in A \\ &= (1 - \alpha)^A \pi(\omega|\overline{A}) && \text{otherwise} \end{aligned}$$

where $*$ stands for min or $\ominus$ depending on the fact that $\pi(\omega|A)$ is ordinal or Bayesian. The constraints

$N(A) = \alpha$ imposes that after the revision $\Pi(A) = 1$ and $\Pi(\bar{A}) = 1 - \alpha$, which means leaving a possibility to the worlds outside A . The worlds inside A are revised as A would be certain, while those outside A are revised as the incoming information would be $\bar{A}$, but the possibility values are reduced of $(1 - \alpha)$.

Interpretation 2. leads to consider A as a fuzzy-set F whose membership function is defined as

$$\begin{aligned}\mu_F(\omega) &= 1 && \text{if } \omega \in A \\ &= 1 - \alpha && \text{otherwise}\end{aligned}$$

The possibility distribution changes as follows:

$$\begin{aligned}\pi(\omega|F) &= \pi(\omega|A) && \text{if } \omega \in A \\ &= \pi(\omega) \wedge (1 - \alpha) && \text{otherwise}\end{aligned}$$

where $*$ stands for $\min$ or Θ . As before, the worlds inside A are revised as it would be sure and the others have a possibility degree other than zero.

Example. Let us consider the example for possibility revision but with a certainty degree $\alpha = 0.6$ for A . They hold $\Pi(A) = 0.8$, $\Pi(\bar{A}) = 0.9$ and $N(A) = 0.1$. If we adopt the interpretation 1. of α , then we obtain the revised distribution $\pi_{\min}^1(\omega|A)$ if $*$ = $\min$ and the distribution $\pi_{\Theta}^1(\omega|A)$ if $*$ = Θ . In both cases, after the revision we have $\Pi(A) = 1$, $\Pi(\bar{A}) = 1 - \alpha$ and $N(A) = \alpha$. If we adopt the interpretation 2. of α , then we obtain the revised distribution $\pi_{\min}^2(\omega|A)$ if $*$ = $\min$ and the distribution $\pi_{\Theta}^2(\omega|A)$ if $*$ = Θ . After the revision is not necessarily the case that $N(A) = \alpha$. The sure effect of uncertainty is that the worlds outside A are not necessarily put to zero (see ω_5).

However, it is never the case that impossible worlds come to possibility.

It is time to move from the semantic connotation of knowledge space to the syntactic one. If $\Delta\rho\Omega$ represents a knowledge space, then a set of formulae K , such that $[K] = \Delta$, is a syntactic

	ω_1	ω_2	ω_3	ω_4	ω_5	ω_6	ω_7	ω_8
$\pi(\omega)$	0	0	0	0	0.9	0.8	0.7	0
$\pi_{\min}^1(\omega A)$	0	0	0	0	0.4	1	0.7	0
$\pi_{\Theta}^1(\omega A)$	0	0	0	0	0.4	1	0.875	0
$\pi_{\min}^2(\omega A)$	0	0	0	0	0.4	1	0.7	0
$\pi_{\Theta}^2(\omega A)$	0	0	0	0	0.36	1	0.875	0

representation of the same knowledge space. In *possibilistic logic*, any sentence p is attached with the number $\alpha \in [0, 1]$ such that $N([p]) \lambda \alpha$, i.e., α is the lower bound for the degree of necessity of the set of worlds in which p is true (Dubois et al., 1994). $(p, 1)$ means that p is certain. p unknown is expressed as $(p, 0)$ and $(\neg p, 0)$. The inference rule is the following extension of the resolution principle to weighted clauses:

$$(c, \alpha), (c', \alpha') \vdash (Res(c, c'), \min(\alpha, \alpha'))$$

where c and c' are propositional clauses and $Res(c, c')$ their resolvent.

Let $B = \{(p_i, \alpha_i), i = 1, \dots, n\}$ be a *possibilistic base*. Every completely ordered base can be transformed into a possibilistic one, since the unitarity of the interval is arbitrary. $[(p, \alpha)]$ is a fuzzy set on Ω defined by:

$$\begin{aligned}\mu_{[p, \alpha]}(\omega) &= 1 && \text{if } \omega \in [p] \\ &= 1 - \alpha && \text{otherwise}\end{aligned}$$

This membership function is the least specific possibility distribution such that $N([p]) \lambda \alpha$. The set of possible worlds, in which every $p_i \in B$ is true, has the following possibility distribution:

$$\forall \omega \in \Omega, \quad \pi(\omega) = \min_{i=1, \dots, n} \max(\mu_{[p_i, \alpha_i]}(\omega), 1 - \alpha_i)$$

The consistency of B is equivalent to the consistency of the classic base obtainable from B removing the weights. It can be checked by means of resolution, trying to generate an empty clause with a weight $\alpha > 0$, i.e., $B \vdash (\star, \alpha)$. By definition $B \vdash (p, \alpha)$ means that $B \vdash (\neg p, 1) \vdash (\star, \alpha)$. If a proposition can be derived from different sets of clauses, then its degree of necessity is the highest one obtainable, consequently the *inconsistency degree* of B is defined as $Inc(B) = \max\{\alpha \mid B \vdash (\star, \alpha)\}$. If B is inconsistent then $Inc(B) > 0$ and nonmonotonic inferences can be drawn from B , precisely all the propositions (p, α) such that $\alpha > Inc(B)$ (cf. Nonfjall and Larsen (1992) for efficient methods). Only the clauses in B with a degree of necessity less than α will be involved in the inference of (p, α) . A subset of B is consistent if it contains only clauses with a degree of necessity higher than $Inc(B)$.

The *expansion* of B by p is $B \vdash \{(p, 1)\}$, provided that $Inc(B \vdash \{(p, 1)\}) = 0$. In fact, the possibility distribution on the possible worlds that verify $B \vdash \{(p, 1)\}$ is $\pi'(\omega) = \min(\pi(\omega), \mu_{[p]}(\omega)) = \pi_p^+(\omega)$.

Let B be consistent, while $B' = B \vdash \{(p, 1)\}$ be inconsistent with $Inc(B') = \alpha > 0$. It holds:

$$B \vdash f(p, 1)g \vdash (q, \beta) \quad \text{with } \beta > \alpha \quad \text{iff} \quad N([q]/[p]) > \alpha$$

where $N([q]/[p])$ is the necessity measure of $[q]$ after revising $[B]$ by $[p]$. The possibility distribution $\bar{\pi}'(\omega)$ induced from the consistent part of B' made of sentences whose weight is higher than α is defined as:

$$\begin{aligned} \bar{\pi}'(\omega) &= \pi(\omega) & \text{if } \omega \in [p] \text{ and } \pi'(\omega) < 1 - \alpha \\ &= 1 & \text{if } \omega \in [p] \text{ and } \pi'(\omega) = 1 - \alpha \\ &= \pi'(\omega) & \text{otherwise} \end{aligned}$$

Hence $\bar{\pi}'(\omega)$ is the result of revising $\pi(\omega)$ by $[p]$, i.e., $\pi(\omega|[p])$.

This revision is drastic, since it rejects all the sentences (p_i, α_i) with $\alpha_i < \alpha$, even if they were not involved in the derivation of any inconsistencies, and replaces them with $(p, 1)$. The result is the same produced by a revision based on epistemic entrenchment or on Best-Out ordering.

A less drastic revision scheme is based on the choice of preferred subsets of B that fail to imply $(\neg p, \alpha)$ for any $\alpha > 0$. We may take advantage of the ordering to make the selection of the preferred subbases that will be in turn ordered by a inclusion-based or lexicographical scheme, as described in section 3.2 (Benferhat et al., 1993; cf. also Brewka, 1989). However, such revision schemes cannot be expressed at the semantic level.

4.3 The belief-function framework

This approach (Shafer, 1976) assigns a basic probability directly to the subsets of Ω , called the *frame of discernment*, $m: 2^\Omega \rightarrow [0, 1]$, with the constraints $m(\emptyset) = 0$ and $\sum_{A \in 2^\Omega} m(A) = 1$. m is called *basic probability assignment (bpa)*. If $m(A) > 0$ then A is a *focal element*. The *belief* and *plausibility* functions on the subsets of Ω are defined as

$$\begin{aligned} Bel(A) &= \sum_{B \subseteq A} m(B) \\ Pl(A) &= 1 - Bel(\bar{A}) = \sum_{B \cap A \neq \emptyset} m(B) = \sum_{B \cap (\Omega - A) = \emptyset} m(B) \end{aligned}$$

$Bel(A)$ measures the persuasion that the real world is inside A ; it may be that there is no evidence that directly supports A (see the case $\{\omega_1, \omega_3\}$ in the next example), but it cannot be excluded ($Bel(A) \neq 0$ and $Pl(A) \neq 0$) because there is evidence in favour of some of its subsets. These functions are not additive: $Bel(A) + Bel(\bar{A}) \leq 1$ and $Pl(A) + Pl(\bar{A}) \geq 1$. The knowledge is:

- *certain* and *precise* if $m(T) = 1$, $\sum_{A \in 2^\Omega} m(A) = 0$ and T is singleton
- *certain* and *imprecise* as before but T is not singleton
- *consistent* if all the focal elements are nested
- *inconsistent* if all the focal elements are disjoint
- *void* if $m(\Omega) = 1$ and $\sum_{A \in 2^\Omega} m(A) = 0$.

Example. Suppose that $\Omega = \{\omega_1, \omega_2, \omega_3\}$.

2^Ω	$\{\omega_1\}$	$\{\omega_2\}$	$\{\omega_3\}$	$\{\omega_1, \omega_2\}$	$\{\omega_1, \omega_3\}$	$\{\omega_2, \omega_3\}$	$\{\omega_1, \omega_2, \omega_3\}$
m	.2	.1	.1	.2	0	.3	.1
Bel	.2	.1	.1	.5	.3	.5	1
Pl	.5	.7	.5	.9	.9	.8	1

The incoming information $A\rho\Omega$ changes the values of credibility and plausibility of the subsets of Ω . In terms of plausibility such modifications are defined by the “Dempster Rule of Conditioning”:

$$Pl(B|A) = \frac{Pl(A \sqcap B)}{Pl(A)}$$

$$Bel(B|A) = 1 - Pl(B|A)$$

The subsets disjoint from A will become implausible. In terms of credibility, revision could be done by means of the “Geometric Rule of Conditioning”:

$$Bel_g(B|A) = \frac{Bel(A \sqcap B)}{Bel(A)}$$

$$Pl_g(B|A) = 1 - Bel_g(B|A)$$

The subsets disjoint from A will become incredible. The Dempster Rule of Conditioning is not applicable when $Pl(A) = 0$, i.e., when A is impossible, while the geometric rule is not applicable when $Bel(A) = 0$, i.e., when A is incredible.

2^Ω	$\{\omega_1\}$	$\{\omega_2\}$	$\{\omega_3\}$	$\{\omega_1, \omega_2\}$	$\{\omega_1, \omega_3\}$	$\{\omega_2, \omega_3\}$	$\{\omega_1, \omega_2, \omega_3\}$
$Bel(\Delta A)$	.22	.44	0	1	.22	.44	1
$Pl(\Delta A)$	.56	.78	0	1	.56	.78	1
$Bel_g(\Delta A)$	.4	.2	0	1	.4	.2	1
$Pl_g(\Delta A)$	.8	.6	0	1	.8	.6	1

Example. In the previous example, suppose that $A = \{\omega_1, \omega_2\}$. $Pl(A) = 0.9$, $Bel(A) = 0.5$, so we can apply both the Dempster’s and Geometric Rule of Conditioning.

This framework also deals with uncertain inputs. They are treated as new *bpas* on the elements of 2^Ω . The change will consist of merging the two evidences (the prior and the new) about the same real situation. Given a frame of discernment Ω and the two *bpas* m_1 and m_2 , the new combined *bpa* is defined by the following Dempster Rule of Combination:

$$m(B) = m_1(B) \Phi m_2(B) = \frac{m_1(A_1) \Delta m_2(A_2)}{m_1(A_1) \Delta m_2(A_2)} \quad \begin{matrix} X \\ A_1 \sqcap A_2 = B \\ A_1 \sqcap A_2 \neq \emptyset \end{matrix}$$

Example. Referring to the previous example, suppose that someone tells us that the real world is ω_1 or ω_2 , but he cannot discriminate between the two, i.e., we have the new evidence represented by $m_A(\{\omega_1\}) = m_A(\{\omega_2\}) = 0.5$. The new combined *bpa* $m \Phi m_A$ and the corresponding new belief function $Bel(\Delta A)$ are reported in the following table:

2^Ω	$\{\omega_1\}$	$\{\omega_2\}$	$\{\omega_3\}$	$\{\omega_1, \omega_2\}$	$\{\omega_1, \omega_3\}$	$\{\omega_2, \omega_3\}$	$\{\omega_1, \omega_2, \omega_3\}$
m_A	.5	.5	0	0	0	0	0
$m \Phi m_A$	.42	.58	0	0	0	0	0
$Bel(\Delta A)$	.42	.58	0	1	.42	.58	1

This rule reinforces concurring evidence and weakens conflicting evidence. It can be applied only if the evidence is independent and refers to the same frame of discernment.

Because of the commutativity of the product, the rule is independent from the sequence of the pieces of information, so it violates the principle of priority to the incoming information! Those who consider this a transgression introduce the “extended Jeffrey rule”:

$$Pl(Bj(F_{\epsilon}, m_{\epsilon})) = \prod_{A \in \Omega} m_{\epsilon}(A)^{\frac{Pl_1(A \uparrow B)}{Pl_1(A)}}$$

This rule is clearly asymmetric and depends on the chronological order of the evidence.

As a final remark, we say that the main problem with the belief function formalism is the computational complexity of the Dempster’s Rule of Combination; its straightforward application is exponential in the frame of discernment and the amount of evidence. However, much effort has been spent in reducing the complexity of that rule. Such methods range from “efficient implementations” (Kennes 1992) to “qualitative approaches” (Parsons, 1994), through “approximate techniques” with statistical methods as the Montecarlo sampling algorithm (Wilson, 1991; Moral and Wilson, 1996).

5. Recoverability vs. priority to the incoming information

5.1 The principle of recoverability

Almost all the models for belief revision presented so far recognize the rationality of the “priority to the incoming information” principle. However, as exceptions, we have seen that the Dempster–Shafer approach puts the new information on the same grounds as the previous knowledge space, and both the probabilistic and possibilistic frameworks can be extended to deal with uncertain inputs.

We highlight two points about this principle.

1. While giving priority to the incoming information is acceptable when updating the representation of an evolving world, it is not generally justified when revising the representation of a static situation. In this case, the chronological sequence of the informative acts has almost nothing to do with their credibility or importance; we think in particular of multi-agent domains, in which information comes from various sources (Shafer and Srivastava, 1990; Dragoni and Di Manzo, 1995; Dragoni, 1996). In these cases, it seems reasonable to treat all the available pieces of information as if they had been collected at the same time.
2. Accepting the incoming information (hoping that it is not inconsistent) and remaining consistent implies throwing away part of the previously held knowledge, but this change should not be irrevocable. Let $K*p$ be the cognitive state revised after the new information p . For each cognitive state K , and sentences p and q such that $K \vdash p$ and $K*q \not\vdash p$, there can always be another piece of information r such that $(K*q)*r \vdash p$ even if $r \not\vdash p$. An obvious case should be $r = \neg q$.

Rejecting the incoming information does not necessarily mean leaving the knowledge space unchanged since, in general, it alters the distribution of the masses. Obviously, a rejectable information should have little power to change the “status quo”. As a consequence of this alteration, a different selection of preferred sentences might result. Among the newly preferred beliefs, one may again find some previously discarded pieces of knowledge.

Summing up, to make practical and useful belief revision in a multi-agent environment, we substitute the priority to the incoming information principle with the following one:

- *principle of recoverability*: any previously held piece of knowledge should belong to the current knowledge space whenever it is consistent with it.

This principle does not hold for updating. The meaning of the word “revision” changes from “dealing with new information” to “dealing with a new broader set of pieces of information”. Our sentence-based model for belief revision considers two knowledge repositories:

1. The *knowledge background* (KB), which is the set of all pieces of knowledge available to the reasoning agent; since it can be inconsistent, it cannot be used as a whole to support reasoning and decision processes.
2. The *knowledge base* $B \theta KB$, which is the maximally consistent, currently preferred piece of knowledge that should be used for reasoning and decision supporting; since it is maximally consistent, it may contain pieces of knowledge with low degrees of credibility.

The incoming information p , with its weight of evidence, is confronted not just with B , but with the overall KB. A first advantage in doing so is that the degrees of credibility of the sentences in $KB \bowtie \{p\}$ are reviewed on a broader and less prejudicial basis. This is not a transmutation (section 3.1), because the weight of p may also change. The main advantage in doing so is that we can rescue sentences from KB by virtue of the maximal consistency of B . On the basis of the new ordering induced by p , the system chooses another base $B' \theta KB \bowtie \{p\}$ that could be syntactically equal to B (meaning that p has been rejected) but, in general, it will have a different credibility distribution than B . B' may not contain p even in the case that B' and B are syntactically different, since the new ordering may induce the choice of a base which is different from the preceding one, but still contains sentences incompatible with p : the maximal consistency of B' may induce the rescue of sentences from KB, even if p has been rejected.

Suppose that KB is consistent, so that $B \eta KB$. Let $KB' = KB \bowtie \{p\}$. If KB' is consistent too, then B is simply expanded by p : $B' = B \bowtie \{p\} = KB \bowtie \{p\} = KB'$. If KB' is inconsistent, then it will be necessary to start a revision process that will determine the new base B' .

If $p \in B'$, then $B' = C \bowtie \{p\}$, where $C \in KB \vartheta \neg p$. If $p \notin B'$, then $B' = B$.

Suppose now that KB is inconsistent. If we revise only B by p , we could not recover information from KB. For instance, partial-meet base revision (section 3.1) would select some $B' \in B \vartheta \neg p$, but it will always be possible to find out some $B'' \in KB \vartheta \neg p$ such that $B' \theta B''$.

Example. Let $KB = \{a, \neg a \vee b, \neg b\}$, $B = \{a, \neg a \vee b\}$ and $p = \neg a$.

$B \vartheta \neg p = \{\{\neg a \vee b\}\}$, so $B' = \{\neg a \vee b\}$. Base revision would yield $B^*p = \{\neg a \vee b, \neg a\}$. $KB \vartheta \neg p = \{\{\neg a \vee b, \neg b\}\}$, so $B'' = \{\neg a \vee b, \neg b\}$. Now it holds that $B' \rho B''$. The revision yields $KB^*p = \{\neg a \vee b, \neg b, \neg a\}$.

The arrival of p caused the recover of $\neg b$, that was not in the previous base B but in KB.

Computationally, our route to belief revision consists of five steps:

1. Detection of the minimally inconsistent subsets of $KB \bowtie \{p\}$, here called *nogoods*.
2. Generation of the maximally consistent subsets of $KB \bowtie \{p\}$, here called *goods*.
3. Revision of the credibility weights of the sentences in $KB \bowtie \{p\}$.
4. Choice of a preferred maximally consistent subset of $KB \bowtie \{p\}$ as the new revised base B' .
5. Selection of the derived sentences which are derivable from B' .

As already pointed out, when p is consistent with B , not necessarily $B' = B \bowtie \{p\}$ (expansion), since the new credibility distribution produced at step 3 may yield a totally different choice at step 4. So, the rejection of the priority to the incoming information principle implies that K*4 and K*5 no longer hold (if p is inconsistent it will be part of none of the goods produced at step 2, so it will never be part of a base).

Steps 1, 2 and 5 deal with consistency and derivation, and act on the symbolic part of the information. Operations are in ATMS style; to find out nogoods and goods, we adopt (and adapt) a well-known set-covering algorithm that will be presented in the next section. Even in the propositional case, determining all the minimal inconsistencies can be very hard. However, such a

condition can be relaxed. A reasonable way is to leave out (better “leaving in”) contradictions involving only sentences with low degrees of credibility (see Nonfjall and Larsen (1992) for efficient methods). The consequence is that some of the goods are not really consistent, but if you don’t mind some negligible sentences... As a final remark on this argument, it must be said that in practical applications dealing with commonsense knowledge (for instance, Dragoni and Di Manzo, 1995), such minimal inconsistencies are not automatically derivable, but they can be provided interactively by the user.

Steps 3 and 4 deal with uncertainty and work with the numerical weight of the information. Both contribute to the choice of the revised knowledge space, so their reasonableness should be evaluated as a couple. The numerical formalisms presented are able to perform both of them. In fact, they attach weights to sets of possible worlds, so the credibility of a single sentence p is determined in the same way as the credibility of a set of sentences B by the weights attached to $[p]$ and $[B]$, respectively. However, flexibility is an advantage in separating the two steps; for instance, depending on the characteristics of the knowledge domain under consideration and the kind of task and/or decision that should be taken on the basis of the revision outcome, the selection function could also consider one (or a combination) of the methods described in section 3.2.

To perform step 3, probabilistic methods with uncertain inputs (section 4.1) seem inadequate for the strong dependence that they impose on the credibility of a sentence and that of its negation. The “drowning problem” is a limit of the possibilistic approach (Benferhat et al., 1993). We see that the belief-function formalism could work well.

Step 4 translates such ordering on the *sentences* in KB into an ordering on the *goods* of KB. The best classified good is selected as the preferred revised knowledge base. If the ordering on KB is not strict, then there can be multiple preferred goods. In this case, we could take their intersection as the revised knowledge base (Benferhat et al., 1993; Roos, 1992); however, the intersection is not maximally consistent, and this means that all the conflicting pieces of knowledge with the same credibility will be rejected.

In this framework, even new information compatible with B can rescue previously rejected pieces of knowledge $R \rho KB$, simply by determining some upsetting between the credibility of a set $S \rho B$ and the credibility of R . It is unreasonable to decrease the credibility of S upon the arrival of information p that is not in conflict with S , but what may happen is that p supports R , increasing their credibility w.r.t. S .

Another question is: step 4 should consider only the *implicit* ordering of the sentences in KB (relative classification without the numerical weights), or could take advantage of the *explicit* ordering (numerical weights). The first approach seems closer to human cognitive behaviour (which normally refrains from numerical calculus). The second one seems more informative (takes into account not only relative positions but also detachings) and, hence, more rational. As an example of a numerical way to perform step 4, ordering the goods according to their average credibility seems reasonable and easy to calculate. The main difference with the inclusion-based method is that the preferred good may not contain the most credible sentence. However, it has its own drawbacks. In the short run (with goods made of few sentences), there is a strange and unintuitive dependence of the effect of the incoming information on the cardinality of the goods in which it enters. In the long run, the degrees of credibility of the goods become very close to each other, so that the differences turn out to be rather insignificant.

5.2 Steps 1 and 2

We adopt the algorithm presented in Reiter (1987) to calculate model-based diagnoses (which is similar to the one presented by Chou and Winslett, 1994).

We recall here the basic definitions to keep the paper self-contained. Let F be a collection of sets. A *hitting-set* for F is a set $H \rho \bigcup_{S \in F} S$ such that $H \cap S \neq \emptyset$ for each $S \in F$. A hitting-set is *minimal* if none of its proper subsets is a hitting-set for F .

Example. Let $F = \{\{1,5,7\}, \{2,5,6\}, \{1,2\}\}$. The minimal hitting sets for F are: $\{1,2\}$, $\{1,5\}$, $\{1,6\}$, $\{5,2\}$, $\{7,6,1\}$ and $\{7,2\}$.

An *HS-Tree* for F is a smallest labelled tree T such that:

1. its root is labelled with P if $F = \emptyset$; otherwise its root is labelled with an element of F ;
2. if n is a node of T , let $H(n)$ be the set of the labels of the arcs on the path from the node to the root. If n is labelled with P , then it has no successors in T . If n is labelled with an element Σ of F , then for each element $\sigma \in \Sigma$, n has a successors node n_σ joined with n by an arc labelled with σ . The label of n_σ is a set $S \in F$ such that $S \supset H(n) = \emptyset$ if it exists, otherwise n is labelled with P .

If n is labelled with P , then $H(n)$ is a hitting-set for F . The set of all the $H(n)$ of the nodes labelled with P includes all the minimal hitting-sets for F . An HS-Tree, generated and pruned according to the following five rules, minimizes the accesses to F while preserving all the paths, from the leafs to the root, corresponding to the minimal hitting sets.

1. The HS-Tree must be generated breadth first.
2. If a node n is labelled with S and n' is another node such that $H(n') \supset S$, then label n' with S .
3. If a node n is labelled with P and n' is another node such that $H(n) \theta H(n')$, then n' must be pruned.
4. If a node n has already been generated and n' is another node such that $H(n) = H(n')$, then n' must be pruned.
5. If a node n is labelled with S and n' is labelled with S' with $S' \rho S$, then for each $\sigma \in S - S'$, the arc starting from n and labelled with σ must be pruned.

For an HS-Tree for F generated and pruned with these rules, the set $\{H(n) \mid n \text{ is a node of } T \text{ labelled with } P\}$ is the collection of all the minimal hitting sets for F .

Given KB , the collection of the goods and the collection of the nogoods are dual. If we remove from KB exactly one element for each nogood, what remains is a good (Roos, 1992). Hence, the collection of the goods can be found by calculating all the minimal hitting-sets for the collection N of the nogoods, and keeping the complement of each of them w.r.t. KB .

If F is a collection of sets and $F' \theta F$ is the collection of all the minimal elements (w.r.t. set inclusion) of F , then F and F' have the same minimal hitting-sets. This simplifies our task, since we do not need to calculate the collection N of the nogoods (i.e., *minimally* inconsistent subsets of KB), but just a collection $M \subseteq N$ of inconsistent subsets of KB , which is much easier.

It can be proved that, given a collection F of sets, in any HS-Tree T relative to F , generated and pruned with the rules 1–5, every minimal element of F appears as the label of at least one node of T . So, if M contains all the nogoods N , then all of them will be used as labels of the nodes of T . During the generation of a node n , the rule 5 imposes to check if there is a label S which is a superset of the just calculated label S' for n , so we simply have to eliminate such an S if it exists. Finding the nogoods is useful for step 3 of the revision process, since the method that we adopt (the belief function formalism) revises the weights, especially on the basis of *minimal* contradictions.

The collection $M \subseteq N$ can be generated while generating the HS-Tree. After reducing in clausal normal form all the sentences in KB , we start a refutation process on it. If we find the empty clause, we label the root of T with the set of clauses in KB that have been involved in the refutation. Each clause σ_m in this node labels an arc toward a successor node. To calculate the label of this successor node we start a tentative refutation on $KB \setminus \sigma_m$. If $KB \setminus \sigma_m$ is consistent, then the node is labelled with P , otherwise it is labelled with the proper subset of $KB \setminus \sigma_m$ from which the empty clause has been derived. The process iterates until there are no more successor nodes.

Summarizing, the following algorithm generates all the goods and the nogoods in KB . $Inc(C)$ is a function on set of clauses that returns P if C is consistent, returns a set $D \theta C$ if D is inconsistent.

1. $NG := \emptyset$, $G := \emptyset$
2. when a node n needs a new label, give it the label $Inc(KB \setminus H(n))$ and add it to NG
3. when rule 5 applies, eliminate S from NG

4. for each $H(n)$ such that n is labelled with P , put $KB \setminus H(n)$ in G
5. return NG and G

NG is the collection of the nogoods and G is the collection of the goods in KB .

5.3 Steps 3 and 4

We adopt the Dempster–Shafer approach because, as it treats all the pieces of information as if they had been provided at the same time, it does not obey the priority to the incoming information principle. As the probabilistic and possibilistic approaches, even Dempster–Shafer’s approach is able to assign a degree of credibility both to single sentences and to goods (i.e., sets of sentences). However, as we shall see, the way it evaluates the credibility of a good is not quite satisfactory.

Being especially interested in multi-agent domains, we adopt the belief-function formalism, in the special guise in which Shafer and Srivastava (1990) apply it to auditing. In that work, it is regarded as a process with, basically, the following I/O:

INPUT: list of couples $\langle \text{source, piece of information} \rangle$
 list of couples $\langle \text{source, reliability} \rangle$
 OUTPUT: list of couples $\langle \text{piece of information, credibility} \rangle$

Both the reliability of any source and the credibility of any piece of information are in the range $[0,1]$, but while reliability is regarded as the *probability* that the source is faithful, credibility is a *belief-function* on the piece of information.

The reader should refer to Shafer and Srivastava (1990) for a less synthetic description and a “philosophical” discussion of this subject. We present the method in a way more related to the classic description introduced in section 4.3.

Let $S = \{s_1, \dots, s_n\}$ be the set of the sources, and let kb_i be the subset of KB received from s_i . Each source s_i is associated with a *reliability* R_i , that is regarded as the *a priori probability* that the source is faithful. The main idea with this multi-source version of the belief function framework is that *a reliable source cannot give false information, while an unreliable source can give correct information*; the hypothesis that s_i is reliable is compatible only with the models of kb_i , while the hypothesis that s_i is unreliable is compatible with the overall Ω . Each source s_i is an evidence for KB and generates the following *bpa* m_i on 2^Ω :

$$m_i(X) = \begin{cases} R_i & \text{if } X = [kb_i] \\ 1 - R_i & \text{if } X = \Omega \\ 0 & \text{otherwise} \end{cases}$$

All these *bpas* will then be combined through the Dempster Rule of Combination. From the combined *bpa* m , the credibility of a sentence p of L is given, as usual, by:

$$Bel(p) = \sum_{B \models p} m(B)$$

Following Shafer and Srivastava, we defined the *a priori* reliability of a source as the *probability* that the source is reliable. These degrees of probability are “translated” by the Theory of Evidence into belief-function values on the given pieces of information. After the combination of the evidences, we may also want to estimate the *a posteriori* reliability of the sources. To be congruent with the *a priori* reliability, also the *a posteriori* reliability must be a *probability* value, not a belief-function one. This is the reason why we adopt Bayesian Conditioning instead of the Theory of Evidence to calculate it. Let us see in detail how it works here.

Let us consider the hypothesis that only the sources belonging to $\Phi \subseteq S$ are reliable. If the sources are independent, then the probability of this hypothesis is $R(\Phi) = \prod_{s_i \in \Phi} R_i \prod_{s_j \in \Phi^c} (1 - R_j)$. We could calculate this “combined reliability” for any subset of S . It holds that $\sum_{\Phi \subseteq S} R(\Phi) = 1$. Possibly, the sources belonging to a certain Φ cannot all be considered reliable, because they gave contradictory information, i.e., a set of information items s such that $[s] = \emptyset$. In this case, the combined

reliabilities of the remaining subsets of S are subjected to the Bayesian Conditioning so that they sum up again to “1”, i.e., we divide each of them by “ $1 - R(\Phi)$ ”. In the case where there are more subsets of S , say $\Phi_1, \dots, \Phi_l$, containing sources which cannot all be considered reliable, then $R(\Phi) = R(\Phi_1) + \dots + R(\Phi_l)$. We define the revised reliability NR_i of a source S_i as the sum of the conditioned combined reliabilities of the “surviving” subsets of S containing S_i .

An important feature of this method to recalculate the sources’ reliability is that if S_i is involved in contradictions, then $NR_i \neq R_i$, otherwise $NR_i = R_i$.

The main problem with Dempster’s Rule of Combination is its computational complexity. However, we realized that two properties hold:

1. Sentences not involved in (minimal) contradictions and received from a single source do not contribute to the mechanism, in the sense that the degrees of credibility of the other sentences do not depend on their presence. Their degree of credibility is equal to the new degree of reliability of its source.
2. Multiple (minimal) contradictions involving sentences received exclusively from exactly the same sources are redundant; all the sentences from the same source receive the same degree of credibility, independently of the number and cardinality of the contradictions.

Property (1) implies that we can leave out of the process those sentences received from a single source that are not involved in contradictions. Property (2) says that what is important is that a set of sources was contradictory, not how many times nor about what or about how many sentences. This also allows us to leave out of the process some sentences involved in contradictions; this is significant in situations like that of two sources systematically in contradiction only with each other.

The belief-function formalism is able to attach directly a degree of credibility to any good g , regarding it as a single conjunctive sentence, hence bypassing step 4 in our framework. Unfortunately, it holds the following result. Let kb be the subset of KB received *exclusively* from a source s , and let G be a good. If there exist two partitions of kb , say kb' and kb'' , such that only the sentences of kb' belong to G , then $Bel(G) = 0$. This event is all but infrequent, so that the belief-function formalism is a bad estimator of the goods’ credibility, since it does not discriminate sufficiently among them. It attaches a non-null credibility only to goods supported by sources that have never been contradictory regarding any information provided. Besides, it is unreasonable to put at zero the credibility of a good only because it contains pieces of information received from a source that gave other pieces of information that we do not believe (are false in the good).

5.4 Step 5

This step is not particularly significant since, theoretically, it simply consists in applying classical entailment on the preferred good to deduce a plausible conclusion from it. We adopted an ATMS and stored each sentence derived by the Theorem Prover with an *origin set* (Martins and Shapiro, 1988), i.e., a set of basic assumptions which are all *necessary* to derive it. Practically, the last step consists in selecting from the derived sentences all those whose origin set is a subset of the preferred good. We could relax the definition of origin set to that of a set of basic assumptions used to derive the sentence. This is easier to compute, and does not have pernicious consequences; the worst that can happen is that, this relaxed origin set being a superset of the real one, it is not sure that it will be a subset of the preferred good as the real one is, and so some derived logical consequences of the preferred good may be not recognized (at first).

6. Example

The following example is taken from the investigative domain. It is an extreme simplification of a case of the Inquiry Support System presented in Dragoni (1997), which embodies the model for belief revision presented in the previous section. The text in *courier* has been just translated into English directly from the Input/Output text files (which were written in Italian). In this example,

goods are printed out ordered according to the “best-out” method; however, we have also provided their “average credibility” and their belief-function value. The case starts with three sources, B, M and A, that refer three facts regarding A (which is the defendant in the trial). The sources are assigned the same degree of reliability (0.5). The other two special sources are OBS (a fictitious source that gives only verified facts; reliability 0.9) and U (which represents the user and introduces only hypothetical facts and/or rules; reliability 0.3).

INITIAL INPUT:

B asserts: 'A is a partner of R'

M asserts: 'A was driving the car of S'

A asserts: it is false that 'A knows S'

"a priori" reliability of A, B and M: 0.5

There are no contradictions; there is one good and the degrees of reliability don't change.

					A is a partner of R 0.500	
					A was driving the car of S 0.500	
					it is false that A knows S 0.500	
GOOD 1					Bel 0.125	Average 0.500
=====						
"a posteriori" rel.		FALL		OLD		NEW
=====						
	B	0.000		0.500		0.500
	A	0.000		0.500		0.500
	M	0.000		0.500		0.500

NEW INFORMATION:

U asserts: 'A was using the car of S' **implies** 'A knows S'

"a priori" reliability of U: 0.3

					A is a partner of R 0.500	
					A was driving the car of S 0.459	
					it is false that A knows S 0.459	
GOOD 1					Bel 0.189	Average 0.473
=====						
					A is a partner of R 0.500	
					A was driving the car of S 0.459	
					A knows S or	
					it is false that A was driving the car of S 0.243	
					A knows S <i>deriv</i>	
GOOD 2					Bel 0.081	Average 0.401
=====						
					A is a partner of R 0.500	
					it is false that A knows S 0.459	
					A knows S or	
					it is false that A was driving the car of S 0.243	
					<i>it is false that A was driving the car of S</i> <i>deriv</i>	
GOOD 3					Bel 0.081	Average 0.401
=====						
"a posteriori" rel.		FALL		OLD		NEW
=====						
	B	0.000		0.500		0.500
	A	0.041		0.500		0.459
	M	0.041		0.500		0.459
	U	0.057		0.300		0.243

The hypothetical rule introduced by U causes a contradiction with the testimonies of A and MA. Their reliabilities decrease, but the preferred good does not change (in this case). Note that the second and third goods have a non-empty context (one derived sentence)

NEW INFORMATION:

A asserts: 'L is a friend of A'
OBS asserts: 'L is a cousin of S'
"a priori" reliability of OBS: 0.9

			L is a cousin of S		0.900
			A is a partner of R		0.500
			L is a friend of A		0.459
		A was driving the car of S		0.459	
		it is false that A knows S		0.459	
GOOD 1	Bel	0.189	Average		0.556
=====					
			L is a cousin of S		0.900
			A is a partner of R		0.500
			L is a friend of A		0.459
		A was driving the car of S		0.459	
			A knows S or		
		it is false that A was driving the car of S		0.243	
			A knows S		deriv
GOOD 2	Bel	0.081	Average		0.512
=====					
			L is a cousin of S		0.900
			A is a partner of R		0.500
			L is a friend of A		0.459
		it is false that A knows S		0.459	
			A knows S or		
		it is false that A was driving the car of S		0.243	
		it is false that A was driving the car of S		deriv	
GOOD 3	Bel	0.081	Average		0.512
=====					

"a posteriori" rel.		FALL		OLD		NEW	
OBS		0.000		0.900		0.900	
B		0.000		0.500		0.500	
A		0.041		0.500		0.459	
M		0.041		0.500		0.459	
U		0.057		0.300		0.243	

Incoming information that does not contradict nor confirm anything (like 'A is associated with R' and 'L is friend of A' at this stage) belongs to all the goods; its belief-function value is exactly the *a posteriori* reliability of its source (respectively 0.5 for B and 0.4594594595 for A) as results from the Bayesian Conditioning.

Sources that give information not contradicted nor confirmed by others (like OBS and B at this stage) maintain their *a priori* reliability. Of course, these general rules simplify enormously the computation.

NEW INFORMATION:

U asserts: 'L is friend of A' **and** 'L is cousin of S' **implies** 'A knows S'

			L is a cousin of S		0.892
			A is a partner of R		0.500
		A was driving the car of S		0.496	
			L is a friend of A		0.417
		it is false that A knows S		0.417	
GOOD 1	Bel	0.184	Average		0.544
=====					
			L is a cousin of S		0.892
			A is a partner of R		0.500
		A was driving the car of S		0.496	
			L is a friend of A		0.417
			A knows S or		
		it is false that L is a friend of A or			
		it is false that L is a cousin of S		0.184	
			A knows S or		
		it is false that A was driving the car of S		0.184	
			A knows S		deriv

GOOD 2	Bel	0.000	Average	0.445			
=====							
		L is a cousin of S		0.892			
		A is a partner of R		0.500			
		A was driving the car of S		0.496			
		it is false that A knows S		0.417			
		A knows S	or				
		it is false that L is a friend of A	or				
		it is false that L is a cousin of S		0.184			
		it is false that L is a friend of A		deriv			
GOOD 3	Bel	0.000	Average	0.498			
=====							
		L is a cousin of S		0.892			
		A is a partner of R		0.500			
		L is a friend of A		0.417			
		it is false that A knows S		0.417			
		A knows S	or				
		it is false that A was driving the car of S		0.184			
		it is false that A was driving the car of S		deriv			
GOOD 4	Bel	0.000	Average	0.482			
=====							
		L is a cousin of S		0.892			
		A is a partner of R		0.500			
		it is false that A knows S		0.417			
		A knows S	or				
		it is false that L is a friend of A	or				
		it is false that L is a cousin of S		0.184			
		A knows S	or				
		it is false that A was driving the car of S		0.184			
		it is false that A was driving the car of S		deriv			
		it is false that L is a friend of A		deriv			
GOOD 5	Bel	0.000	Average	0.435			
=====							
		A is a partner of R		0.500			
		A was driving the car of S		0.496			
		L is a friend of A		0.417			
		it is false that A knows S		0.417			
		A knows S	or				
		it is false that L is a friend of A	or				
		it is false that L is a cousin of S		0.184			
		it is false that L is a cousin of S		deriv			
GOOD 6	Bel	0.000	Average	0.403			
=====							
		A is a partner of R		0.500			
		L is a friend of A		0.417			
		it is false that A knows S		0.417			
		A knows S	or				
		it is false that L is a friend of A	or				
		it is false that L is a cousin of S		0.184			
		A knows S	or				
		it is false that A was driving the car of S		0.184			
		it is false that A was driving the car of S		deriv			
		it is false that L is a cousin of S		deriv			
GOOD 7	Bel	0.009	Average	0.340			
=====							
"a posteriori" rel.		FALL		OLD		NEW	
	B		0.000		0.500		0.500
	M		0.004		0.500		0.496
	OBS		0.008		0.900		0.892
	A		0.083		0.500		0.417
	U		0.116		0.300		0.184

The new hypothetical rule introduced by U (the user of the system) yields another contradiction (with 'L is friend of A' and 'L is cousin of S'). The reliability of the three sources

involved (U, A and OBS) slightly decrease. Also the credibility of *all* the information items they provided decreases, not only the credibility of the propositions directly involved in the contradiction (see the information ‘A was using the car of S’ **implies** ‘A knows S’ that falls from 0.2432432432 to 0.1836734694).

Until now, the hypothetical rules introduced by the user have been discarded from the preferred good. The hypothesis ‘A knows S’ belongs to the context (derived sentences) of the third good. Suppose now that a letter of recommendation is written by A to S. This can be taken as a proof that ‘A knows S’.

NEW INFORMATION:

OBS asserts: ‘A knows S’

			A knows S	or	
	it is false that L is a friend of A	or			
	it is false that L is a cousin of S		0.871		
		A knows S	or		
it is false that A was driving the car of S		0.871			
	A knows S		0.843		
	L is a cousin of S		0.829		
	A is a partner of R		0.500		
A was driving the car of S		0.493			
	L is a friend of A		0.078		

GOOD 1	Bel	0.000	Average	0.641
=====				

			A knows S	or	
	it is false that L is a friend of A	or			
	it is false that L is a cousin of S		0.871		
		A knows S	or		
it is false that A was driving the car of S		0.871			
	L is a cousin of S		0.829		
	A is a partner of R		0.500		
	it is false that A knows S		0.078		
it is false that A was driving the car of S		deriv			
it is false that L is a friend of A		deriv			

GOOD 2	Bel	0.000	Average	0.630
=====				

			A knows S	or	
	it is false that L is a friend of A	or			
	it is false that L is a cousin of S		0.871		
		A knows S	or		
it is false that A was driving the car of S		0.871			
	A is a partner of R		0.500		
	L is a friend of A		0.078		
	it is false that A knows S		0.078		
it is false that A was driving the car of S		deriv			
it is false that L is a cousin of S		deriv			

GOOD 3	Bel	0.014	Average	0.480
=====				

			L is a cousin of S		0.829
			A is a partner of R		0.500
	A was driving the car of S		0.493		
	L is a friend of A		0.078		
	it is false that A knows S		0.078		

GOOD 4	Bel	0.000	Average	0.396
=====				

			A knows S	or	
	it is false that L is a friend of A	or			
	it is false that L is a cousin of S		0.871		
		L is a cousin of S		0.829	
		A is a partner of R		0.500	
	A was driving the car of S		0.493		
	it is false that A knows S		0.078		
it is false that L is a friend of A		deriv			

GOOD 5	Bel	0.000	Average	0.554
=====				
			A knows S or	
		it is false that L is a friend of A or		
		it is false that L is a cousin of S	0.871	
		A is a partner of R	0.500	
		A was driving the car of S	0.493	
		L is a friend of A	0.078	
		it is false that A knows S	0.078	
		<i>it is false that L is a cousin of S</i>	<i>deriv</i>	
GOOD 6	Bel	0.000	Average	0.404
=====				
			A knows S or	
		it is false that A was driving the car of S	0.871	
		L is a cousin of S	0.829	
		A is a partner of R	0.500	
		L is a friend of A	0.078	
		it is false that A knows S	0.078	
		<i>it is false that A was driving the car of S</i>	<i>deriv</i>	
GOOD 7	Bel	0.000	Average	0.471
=====				
"a posteriori" rel.		FALL	OLD	NEW
	B	0.000	0.500	0.500
	M	0.007	0.500	0.493
	U	0.010	0.300	0.290
	OBS	0.071	0.900	0.829
	A	0.422	0.500	0.078

The previously preferred good now is the least preferred! The hypothetical rules introduced by the user are now far more credible than how reliable is the user himself. A is now the least reliable information provider.

7. Conclusions

After having examined some of the main symbolic and numerical formalisms for belief revision, we have presented a method that tries to integrate both approaches towards more practical and useful results. This method comes from researches in multi-agent domains, in which agents need to assess the credibility of information coming from different sources. It disconceives the principle of "priority to the incoming information" and follows the principle of "recoverability": any previously held piece of information p must belong to the current knowledge base B whenever p is consistent with B . Pieces of knowledge may not only be abandoned (non-monotonicity), but also rescued (recoverability) after new incoming information. To permit this, it is necessary to maintain two knowledge repositories: the knowledge background KB , which is the (inconsistent) collection of the pieces of knowledge available; and the knowledge base B , which is a maximally consistent and preferred subset of KB . The selection of B among the many possible maximally consistent subsets of KB will be performed on the basis of a numerical credibility ordering on the sentences in KB . Besides recoverability, by doing so we overcome many limitations of other classic ways to belief revision, in particular:

- the revision can be iterated
- inconsistent incoming information does not yield inconsistent revised knowledge spaces
- the numerical revision is performed on a broader base (the overall KB)
- the revision is more flexible; for instance, the incoming information could be rejected even if it is consistent with the current knowledge base
- the complete numerical ordering renders the revision as least drastic as possible

Furthermore, the splitting between the symbolic treatment of the inconsistencies and the numerical revision of the credibility weights provides a clear understanding of what is going on and lucid explanations for the choices. Numerical revision is performed in the belief-function framework. We also gave suggestions to soothe the impact with the computational complexity of this formalism.

References

- Alchourrón, CE, Gärdenfors, P, and Makinson, D, 1985. "On the logic of theory change: partial meet contraction and revision functions" *Journal of Symbolic Logic*, **50**, 510–530.
- Benferhat, S, Cayrol, C, Dubois, D, Lang, J and Prade, H, 1993. "Inconsistency management and prioritized syntax-based entailment" *Proc. 13th Inter. Joint Conf. on Artificial Intelligence* 640–645.
- Benferhat, S, Dubois, D and Prade, H, 1995. "How to infer from inconsistent beliefs without revising?" *Proc. 14th Inter. Joint Conf. on Artificial Intelligence* 1449–1455.
- Brewka, G, 1989. "Preferred subtheories: an extended logical framework for default reasoning" *Proc. 11th Inter. Joint Conf. on Artificial Intelligence* 1043–1048.
- Chou Tymothy, S-C, and Winslett, M, 1994., "A model-based belief revision system" *Journal of Automated Reasoning* **12** 157–208.
- Dalal, M, 1988. "Investigations into a theory of knowledge base revision" *Proc. 7th National Conf on Artificial Intelligence* 475–479.
- Darwiche, A and Pearl, J, 1994. "On the logic of iterated base revision" Fagin, R, (ed.), *Proc. 5th Conf on Theoretical Aspects of Reasoning about Knowledge* 519–525. Morgan Kaufmann.
- de Kleer, J, 1986. "An assumption based truth maintenance system" *Artificial Intelligence* **28** 127–162.
- Dixon, S, and Foo, N, 1993. "Connections between the ATMS and AGM belief revision" *Proc. 13th Inter. Joint Conf. on Artificial Intelligence* 534–539.
- Doyle, J, 1992. "Reason maintenance and belief revision: foundation versus coherence theories" in Gärdenfors, P (ed.), *Belief Revision* Cambridge University Press.
- Doyle, J, 1979. "A truth maintenance system" *Artificial Intelligence* **12** (3) 231–272.
- Dragoni, AF, 1992. "A model for belief revision in a multi-agent environment" in Werner, E and Demazeau, Y. (eds.), *Decentralized A. I.* 3 North Holland.
- Dragoni, AF, Giorgini, P and Puliti, P 1994. "Distributed belief revision versus distributed truth maintenance" *Proc. 6th IEEE Conf. on Tools with AI* IEEE Press.
- Dragoni, AF and Di Manzo, M, 1995. "Supporting complex inquiries" *International Journal of Intelligent Systems* **10** (11) 959–986.
- Dragoni AF, Mascaretti, F and Puliti, P, 1995. "A Generalized approach to consistency-based belief revision in Gori, M and Soda, G. (eds.), *Topics in Artificial Intelligence: Proc. Conference of the Italian Association for Artificial Intelligence. LNAI 992* Springer-Verlag.
- Dragoni, AF, 1996. "Distributed decision support systems under limited degrees of competence: a simulation study" in *Decision Support Systems* special issue on "Intelligent Agents as a Basis for Decision Support Systems" North Holland.
- Dragoni, AF, 1997. "Maximal consistency and theory of evidence in the police inquiry domain" *Artificial Intelligence and Law Journal* special issue on "Time and Evidence" (submitted).
- Dubois, D, Lang, J and Prade, H, 1994. "Possibilistic logic" in DM Gabbay et al. (eds.), *Handbook of Logic in Artificial Intelligence and Logic Programming*, vol. 3, 439–513, Oxford University Press.
- Dubois, D and Prade, H, 1994. "A survey of belief revision and update rules in various uncertainty models" *International Journal of Intelligent Systems* **9** 61–100.
- Dubois, D and Prade, H, 1990. "An introduction to possibilistic and fuzzy logics" in G Shafer and J Pearl (eds.), *Readings in Uncertain Reasoning* Morgan Kaufmann.
- Dubois, D and Prade, H, 1991. "Epistemic entrenchment and possibilistic logic" *Artificial Intelligence* **50** 223–239.
- Dubois, D and Prade, H, 1992. "Belief change and possibility theory" in Gärdenfors, P (ed.), *Belief Revision* Cambridge University Press.
- Gärdenfors, P, 1988. "Knowledge in Flux: Modeling the Dynamics of Epistemic States" MIT Press.
- Gärdenfors, P, 1990a. "Belief revision and non monotonic logic: two sides of the same coin?" *Proc. 9th European Conference on Artificial Intelligence* John Wiley & Sons.
- Gärdenfors, P, 1990b. "The dynamics of belief systems: foundations vs. coherence theories" *Revue Internationale de Philosophie* (44) 24–46.
- Gärdenfors, P, 1992. "Belief revision: an introduction" in Gärdenfors, P (ed.), *Belief Revision* Cambridge University Press.
- Harman, G, 1986. "Change in View: Principles of Reasoning" MIT Press.

- Katsuno, H and Mendelzon, AO, 1991a. "On the difference between updating a knowledge base and revising it" in Allen, J, Fikes, R and Sandewall, E (eds.), *Proc. 2nd Inter. Conf. on Principles of Knowledge Representation and Reasoning* Morgan Kaufmann, 387–394.
- Katsuno, H, and Mendelzon, AO, 1991b. "Propositional knowledge base revision and minimal change" *Artificial Intelligence* **52** 263–294.
- Kennes, R, 1992. "Computational aspects of the möbius transform of a graph" *IEEE Transactions in Systems, Man and Cybernetics* **22** 201–223.
- Laskey, KB and Lehner, PE, 1990. "Belief maintenance: an integrated approach to uncertainty management" in G Shafer and J Pearl (eds.), *Readings in Uncertain Reasoning* Morgan Kaufmann.
- Léa Sombé 1994. "A glance at revision and updating in knowledge bases" *International Journal of Intelligent Systems* **9** 1–27.
- Lehmann, D, 1995. "Belief revision, revised" *Proc. 14th Inter. Joint Conf. on Artificial Intelligence* 1534–1540.
- Lévy, F, 1994. "A survey of belief revision and updating in classical logic" *International Journal of Intelligent Systems* **9** 29–59.
- Martins, JP and Shapiro, SC, 1988. "A Model for belief revision" *Artificial Intelligence* **35** 25–79.
- Moral, S and Wilson, N, 1996. "Importance sampling Monte-Carlo algorithms for the calculation of Dempster–Shafer belief" *Proceeding of IPMU'96* Granada.
- Nebel, B, 1989. "A knoweledge level analysis of belief revision" in RJ Brachman, HJ Levesque and R. Reiter, (eds.), *Proc. Inter. Conf. on Principles of Knowledge Representation and Reasoning*, Morgan Kaufmann, 301–311.
- Nebel, B, 1991. "Belief revision and default reasoning: syntax-based approaches" in J. Allen, R. Fikes and E. Sandewall, (eds.), *Proc. 2nd Inter. Conf. on Principles of Knowledge Representation and Reasoning* Morgan Kaufmann.
- Nebel, B, 1992. "Syntax based approaches to belief revision" in P. Gärdenfors (ed.), *Belief Revision* Cambridge University Press.
- Nebel, B, 1994. "Base revision operations and schemes: semantics, representation, and complexity" in AG Cohn (ed.), *Proc. 11th European Conference on Artificial Intelligence* John Wiley & Sons.
- Nonfjall, H and Larsen, HL, 1992. "Detection of potential inconsistencies In knowledge bases" *International Journal of Intelligent Systems* **7** 81–96.
- Parsons, S, 1994. "Some qualitative approaches to applying the Demster–Shafer theory" *Information and Decision Technologies* **19** 321–337.
- Pearl, J, 1988. "Probabilistic Reasoning for Intelligent Systems Morgan Kaufmann.
- Reiter, R, 1987. "A theory of diagnosis from first principles" *Artificial Intelligence* **53**.
- Roos, N, 1992. "A logic for reasoning with inconsistent knowledge" *Artificial Intelligence* **57** 69–103.
- Shafer, G. 1976. *A Mathematical Theory of Evidence* Princeton University Press.
- Shafer, G, 1990. "Belief functions" in G Shafer and J Pearl (eds.), *Readings in Uncertain Reasoning* Morgan Kaufmann.
- Shafer, G and Pearl, J, 1990. *Readings in Uncertain Reasoning*, Morgan Kaufmann.
- Shafer, G and Srivastava R, 1990. "The Bayesian and belief-function formalisms a general perspective for auditing" in G Shafer and J Pearl (eds.), *Readings in Uncertain Reasoning* Morgan Kaufmann.
- Williams, MA, 1995. "Iterated theory base change: a computational model" *Proc. 14th Inter. Joint Conf. on Artificial Intelligence* 1541–1547.
- Wilson, N, 1991. "A Monte-Carlo algorithm for Dempster–Shafer belief" in BD D'Ambrosio, P. Smets and P. Bonissone (eds.), *Uncertainty in Artificial Intelligence, Proc. Seventh Conference*. Morgan Kaufmann, 414–417.
- Winslett, M, 1988. "Reasoning about action using a possible model approach" *Proc. 7th National Conf on Artificial Intelligence* 89–93.
- Zadeh, LA, 1965. "Fuzzy sets" *Information and Control* **8** 338–353.
- Zhu, Q and Lee, ES, 1993. "Dempster–Shafer approach in propositional logic" *International Journal of Intelligent Systems* **8**.