



## Truthfulness in contextual information correction

Frédéric Pichon, David Mercier, François Delmotte, Eric Lefèvre

### ► To cite this version:

Frédéric Pichon, David Mercier, François Delmotte, Eric Lefèvre. Truthfulness in contextual information correction. 3rd International Conference on Belief Functions (BELIEF 2014), Sep 2014, Oxford, United Kingdom. pp.11-20. hal-03522197

**HAL Id: hal-03522197**

**<https://hal.science/hal-03522197>**

Submitted on 11 Jan 2022

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

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

# Truthfulness in contextual information correction

Frédéric Pichon, David Mercier, François Delmotte, and Éric Lefèvre

Univ. Lille Nord de France, F-59000 Lille, France

UArtois, LGI2A, F-62400, Béthune, France

`{frederic.pichon,david.mercier,francois.delmotte,eric.lefevre}@univ-artois.fr`

**Abstract.** Recently, a dual reinforcement process to contextual discounting was introduced. However, it lacked a clear interpretation. In this paper, we propose a new perspective on contextual discounting: it can be seen as successive corrections corresponding to simple contextual lies. Most interestingly, a similar interpretation is provided for the reinforcement process. Two new contextual correction mechanisms, which are similar yet complementary to the two existing ones, are also introduced.

**Keywords:** Dempster-Shafer theory, Belief functions, Information correction, Discounting.

## 1 Introduction

Information correction has received quite a lot of attention in recent years in belief function theory (see, *e.g.*, [9, 11]). It is an important question that deals with how an agent should interpret a piece of information received from a source about a parameter  $\mathbf{x}$  defined on a finite domain  $\mathcal{X} = \{x_1, \dots, x_K\}$ . Classically, the agent has some knowledge regarding the reliability of the source and, using the discounting operation [12], he is able to take into account that knowledge and to modify, or *correct*, the initial piece of information accordingly.

Since its inception, the discounting operation has been extended in different ways. Notably, Mercier *et al.* [10, 9] consider the case where one has some knowledge about the reliability of the source, conditionally on different subsets (contexts)  $A$  of  $\mathcal{X}$ , leading to the so-called contextual discounting operation. One may also refine the discounting operation in order to take into account knowledge about the source truthfulness [11]. Of particular interest for the present work is the dual reinforcement operation to contextual discounting introduced in [9]. Mercier *et al.* [9] show that this correction mechanism amounts to the negation [6] of the contextual discounting of the negation of the initial information, but unfortunately they do not go further in providing a clear interpretation for this interesting operation.

In this paper, we study further contextual correction mechanisms. We present (Section 3) a new framework for handling detailed meta-knowledge about source

truthfulness. Using this framework, we then derive the contextual discounting operation (Section 4.1) and its dual (Section 4.2), leading to a new perspective on the former and an interpretation for the latter. We proceed (Section 4.3) with the introduction of two new contextual correction mechanisms, whose interpretations are similar yet complementary to the two existing ones. Background material on belief function theory is first recalled in Section 2.

## 2 Belief function theory: necessary notions

In this section, we first recall basic concepts of belief function theory. Then, we present existing correction mechanisms that are of interest for this paper.

### 2.1 Basic concepts

In this paper, we adopt Smets' Transferable Belief Model (TBM) [14], where the beliefs held by an agent  $Ag$  regarding the actual value taken by  $\mathbf{x}$  are modeled using a belief function [12] and represented using an associated mass function. A mass function (MF) on  $\mathcal{X}$  is defined as a mapping  $m : 2^{\mathcal{X}} \rightarrow [0, 1]$  verifying  $\sum_{A \subseteq \mathcal{X}} m(A) = 1$ . Subsets  $A$  of  $\mathcal{X}$  such that  $m(A) > 0$  are called *focal sets* of  $m$ . A MF having focal sets  $\mathcal{X}$  and  $A \subset \mathcal{X}$ , with respective masses  $w$  and  $1 - w$ ,  $w \in [0, 1]$ , may be denoted by  $A^w$ . A MF having focal sets  $\emptyset$  and  $A \neq \emptyset$ , with respective masses  $v$  and  $1 - v$ ,  $v \in [0, 1]$ , may be denoted by  $A_v$ . The negation  $\bar{m}$  of a MF  $m$  is defined as  $\bar{m}(A) = m(\bar{A})$ ,  $\forall A \subseteq \mathcal{X}$ , where  $\bar{A}$  denotes the complement of  $A$  [6].

Beliefs can be aggregated using so-called combination rules. In particular, the conjunctive rule, which is the unnormalized version of Dempster's rule [5], is defined as follows. Let  $m_1$  and  $m_2$  be two MFs, and let  $m_1 \odot_2$  be the MF resulting from their combination by the conjunctive rule denoted by  $\odot$ . We have:

$$m_1 \odot_2 (A) = \sum_{B \cap C = A} m_1(B) m_2(C), \quad \forall A \subseteq \mathcal{X}. \quad (1)$$

Other combination rules of interest for this paper are the disjunctive rule  $\oslash$  [6], the exclusive disjunctive rule  $\oplus$  and the equivalence rule  $\odot$  [13]. Their definitions are similar to that of the conjunctive rule: one merely needs to replace  $\cap$  in (1) by, respectively,  $\cup$ ,  $\sqcup$  and  $\sqcap$ , where  $\sqcup$  (exclusive OR) and  $\sqcap$  (logical equality) are defined respectively by  $B \sqcup C = (B \cap \bar{C}) \cup (\bar{B} \cap C)$  and  $B \sqcap C = (B \cap C) \cup (\bar{B} \cap \bar{C})$  for all  $B, C \subseteq \mathcal{X}$ . The interpretations of these four rules are discussed in detail in [11].

### 2.2 Correction mechanisms

Knowledge about a source reliability is classically taken into account in the TBM through the *discounting* operation. Suppose a source  $S$  providing a piece

of information represented by a MF  $m_S$ . Let  $\beta$ , with  $\beta \in [0, 1]$ , be  $Ag$ 's degree of belief that the source is reliable.  $Ag$ 's belief  $m$  on  $\mathcal{X}$  is then defined by [12]:

$$m(\mathcal{X}) = \beta m_S(\mathcal{X}) + (1 - \beta), \quad m(A) = \beta m_S(A), \forall A \subseteq \mathcal{X}. \quad (2)$$

Mercier *et al.* [9] consider the case where  $Ag$  has some knowledge about the source reliability, conditionally on different subsets  $A$  of  $\mathcal{X}$ . Precisely, let  $\beta_A$ , with  $\beta_A \in [0, 1]$ , be  $Ag$ 's degree of belief that the source is reliable in context  $A \subseteq \mathcal{X}$  and let  $\mathcal{A}$  be the set of contexts for which  $Ag$  possesses such contextual meta-knowledge.  $Ag$ 's belief  $m$  on  $\mathcal{X}$  is then defined by the following equation known as *contextual discounting* that subsumes discounting (recovered for  $\mathcal{A} = \{\mathcal{X}\}$ ):

$$m = m_S \bigcup_{A \in \mathcal{A}} A_{\beta_A}. \quad (3)$$

In addition, a dual reinforcement process to contextual discounting, called contextual reinforcement hereafter, is introduced in [9]. Let  $m_S$  be a MF provided by a source  $S$ . The contextual reinforcement of  $m_S$  is the MF  $m$  defined by:

$$m = m_S \bigcap_{A \in \mathcal{A}} A^{\beta_A}, \quad (4)$$

with  $\beta_A \in [0, 1]$ ,  $A \in \mathcal{A}$ . Mercier *et al.* [9] show that this correction amounts to the negation of the contextual discounting of the negation of  $m_S$ . However, they do not go further in providing a clear explanation as to what meta-knowledge on the source this correction of  $m_S$  corresponds. One of the main results of this paper is to provide such an interpretation.

### 3 A refined model of source truthfulness

In the correction schemes recalled in Section 2.2, the reliability of a source is assimilated to its relevance as explained in [11]. In [11], Pichon *et al.* assume that the reliability of a source involves in addition another dimension: its truthfulness. Pichon *et al.* [11] note that there exists various forms of lack of truthfulness for a source. For instance, for a sensor, it may take the form of a systematic bias. However, Pichon *et al.* [11] study only the crudest description of the lack of truthfulness, where a non truthful source is a source that declares the contrary of what it knows. According to this definition, from a piece of information of the form  $\mathbf{x} \in B$  for some  $B \subseteq \mathcal{X}$  provided by a relevant source  $S$ , one must conclude that  $\mathbf{x} \in B$  or  $\mathbf{x} \in \overline{B}$ , depending on whether the source  $S$  is assumed to be truthful or not.

In this section, we propose a new and refined model of source truthfulness that allows the integration of more detailed meta-knowledge about the lack of truthfulness of an information source.

#### 3.1 Elementary truthfulness

Assume that a relevant source provides a piece of information on the value taken by  $\mathbf{x}$  of the form  $\mathbf{x} \in B$ , for some  $B \subseteq \mathcal{X}$ . Let us now consider a particular value

$x \in \mathcal{X}$ . Either  $x \in B$  or  $x \notin B$ , that is, the source may tell that  $x$  is possibly the actual value of  $\mathbf{x}$  or it may tell that  $x$  is not a possibility for the actual value of  $\mathbf{x}$ . Furthermore, for each of those two possible declarations by the source about the value  $x$ , one may have some knowledge on whether the source is truthful or not. For instance, one may believe that the source is truthful when it tells that  $x$  is a possibility – in which case one must conclude that  $x$  is possibly the actual value of  $\mathbf{x}$  if the source does tell that  $x$  is a possibility for  $\mathbf{x}$  – and that it lies when it tells that  $x$  is not a possibility – in which case one must conclude that  $x$  is possibly the actual value of  $\mathbf{x}$  if the source does tell that  $x$  is not a possibility for  $\mathbf{x}$ .

To account for such detailed knowledge about the behavior of the source, let us introduce two binary variables  $\mathbf{p}_x$  and  $\mathbf{n}_x$ , with respective frames  $\mathcal{P}_x = \{p_x, \neg p_x\}$  and  $\mathcal{N}_x = \{n_x, \neg n_x\}$ :  $p_x$  (resp.  $\neg p_x$ ) corresponds to the state where the source is truthful (resp. not truthful) when it tells that  $x$  is possibly the actual value for  $\mathbf{x}$ ;  $n_x$  (resp.  $\neg n_x$ ) corresponds to the state where the source is truthful (resp. not truthful) when it tells that  $x$  is not a possibility for the actual value of  $\mathbf{x}$ .

Now, we can define a variable  $\mathbf{t}_x$  with associated frame  $\mathcal{T}_x = \mathcal{P}_x \times \mathcal{N}_x$ , which contains four states  $t_x = (p_x, n_x)$ ,  $\neg t_x^n = (p_x, \neg n_x)$ ,  $\neg t_x^p = (\neg p_x, n_x)$  and  $\neg t_x = (\neg p_x, \neg n_x)$  allowing us to model the global truthfulness of the source with respect to the value  $x$ :  $t_x$  corresponds to the case where the source tells the truth whatever it says about the value  $x$ , in short the source is said to be truthful for  $x$ ;  $\neg t_x^n$  corresponds to the case of a source that lies only when it tells that  $x$  is not a possibility for  $\mathbf{x}$ , which will be called a negative liar for  $x$ ;  $\neg t_x^p$  corresponds to the case of a source that lies only when it says that  $x$  is a possibility for  $\mathbf{x}$ , which will be called a positive liar for  $x$ ;  $\neg t_x$  corresponds to the case where the source lies whatever it says about the value  $x$ , in short the source is said to be non truthful for  $x$ .

There are thus four possible cases:

1. Suppose the source tells  $x$  is possibly the actual value of  $\mathbf{x}$ , *i.e.*, the information  $\mathbf{x} \in B$  provided by the source is such that  $x \in B$ .
  - (a) If the source is assumed to be truthful ( $t_x$ ) or a negative liar ( $\neg t_x^n$ ), then one must conclude that  $x$  is possibly the actual value of  $\mathbf{x}$ ;
  - (b) If the source is assumed to be a positive liar ( $\neg t_x^p$ ) or non truthful ( $\neg t_x$ ), then one must conclude that  $x$  is not a possibility for the actual value of  $\mathbf{x}$ ;
2. Suppose the source tells  $x$  is not a possibility for the actual value of  $\mathbf{x}$ , *i.e.*,  $x \notin B$ .
  - (a) If the source is assumed to be in state  $t_x$  or in state  $\neg t_x^p$ , then one must conclude that  $x$  is not a possibility for the actual value of  $\mathbf{x}$ ;
  - (b) If the source is assumed to be in state  $\neg t_x^n$  or in state  $\neg t_x$ , then one must conclude that  $x$  is possibly the actual value of  $\mathbf{x}$ ;

### 3.2 Contextual truthfulness

Let  $\mathcal{T}$  denote the possible states of  $S$  with respect to its truthfulness for all  $x \in X$ . By definition,  $\mathcal{T} = \times_{x \in \mathcal{X}} \mathcal{T}_x$ .  $\mathcal{T}$  is clearly a big space, however we will be interested in this paper only by a smaller subspace of  $\mathcal{T}$ , which we define below.

Let  $h_A^{t_1, t_2} \in \mathcal{T}$ ,  $A \subseteq \mathcal{X}$ ,  $t_1, t_2 \in \mathcal{T}_x$ , denote the state where the source is in state  $t_1$  for all  $x \in A$ , and in state  $t_2$  for all  $x \notin A$ . For instance, let  $\mathcal{X} = \{x_1, x_2, x_3, x_4\}$ ,  $A = \{x_3, x_4\}$ ,  $t_1 = \neg t_x^p$  and  $t_2 = t_x$ , then  $h_A^{t_1, t_2} = h_{\{x_3, x_4\}}^{\neg t_x^p, t_x} = (t_{x_1}, t_{x_2}, \neg t_{x_3}^p, \neg t_{x_4}^p)$ , *i.e.*, the source is a positive liar for  $x_3$  and  $x_4$ , and is truthful for  $x_1$  and  $x_2$ .

Consider now the following question: what must one conclude about  $\mathbf{x}$  when the source tells  $\mathbf{x} \in B$  and is assumed to be in some state  $h_A^{t_1, t_2}$ ? To answer this question, one merely needs to look in turn at each  $x \in \mathcal{X}$  and to consider 4 cases for each of those  $x \in \mathcal{X}$ : 1)  $x \notin B$  and  $x \notin A$ ; 2)  $x \notin B$  and  $x \in A$ ; 3)  $x \in B$  and  $x \notin A$ ; 4)  $x \in B$  and  $x \in A$ . Table 1 lists, for each of the 4 cases and for all states  $h_A^{t_1, t_2}$ ,  $t_1, t_2 \in \mathcal{T}_x$ , whether one should deduce that a given value  $x \in \mathcal{X}$  is possibly the actual value of  $\mathbf{x}$  or not – the former is indicated by a 1 and the latter by a 0 in columns  $h_A^{t_1, t_2}$ ,  $t_1, t_2 \in \mathcal{T}_x$ .

**Table 1.** Interpretations of the source testimony according to its contextual truthfulness.

| $x \in B$ | $x \in A$ | $\neg t_x^p, \neg t_x^p$ | $t_x, \neg t_x^p$ | $\neg t_x^p, t_x$ | $t_x, t_x$ | $\neg t_x, \neg t_x^p$ | $t_x, \neg t_x^p$ | $\neg t_x, t_x$ | $t_x, t_x$ | $\neg t_x^p, t_x$ | $t_x, \neg t_x^p$ | $\neg t_x, t_x$ | $t_x, t_x$ | $\neg t_x^p, t_x$ | $t_x, \neg t_x^p$ | $\neg t_x, t_x$ | $t_x, t_x$ | $\neg t_x^p, t_x$ | $t_x, \neg t_x^p$ | $\neg t_x, t_x$ | $t_x, t_x$ |
|-----------|-----------|--------------------------|-------------------|-------------------|------------|------------------------|-------------------|-----------------|------------|-------------------|-------------------|-----------------|------------|-------------------|-------------------|-----------------|------------|-------------------|-------------------|-----------------|------------|
| 0         | 0         | 0                        | 0                 | 0                 | 0          | 0                      | 0                 | 0               | 0          | 1                 | 1                 | 1               | 1          | 1                 | 1                 | 1               | 1          | 1                 | 1                 | 1               | 1          |
| 0         | 1         | 0                        | 0                 | 0                 | 0          | 1                      | 1                 | 1               | 1          | 0                 | 0                 | 0               | 0          | 1                 | 1                 | 1               | 1          | 1                 | 1                 | 1               | 1          |
| 1         | 0         | 0                        | 0                 | 1                 | 1          | 0                      | 0                 | 1               | 1          | 0                 | 0                 | 1               | 1          | 0                 | 0                 | 0               | 0          | 1                 | 1                 | 1               | 1          |
| 1         | 1         | 0                        | 1                 | 0                 | 1          | 0                      | 1                 | 0               | 1          | 0                 | 1                 | 0               | 1          | 0                 | 1                 | 0               | 1          | 0                 | 1                 | 1               | 1          |

According to Table 1, when the source is assumed to be in, *e.g.*, state  $h_A^{t_x, \neg t_x^p}$ , *i.e.*, the source is truthful for all  $x \in A$  and a positive liar for all  $x \in \overline{A}$ , then one should deduce that  $x \in \mathcal{X}$  is a possible value for  $\mathbf{x}$  iff  $x \in B$  and  $x \in A$ , and therefore, since this holds for all  $x \in \mathcal{X}$ , one should deduce that  $\mathbf{x} \in B \cap A$ . For instance, consider state  $h_{\{x_3, x_4\}}^{t_x, \neg t_x^p}$  and testimony  $\mathbf{x} \in \{x_1, x_3\}$ , then one should deduce  $\{x_1, x_3\} \cap \{x_3, x_4\} = \{x_3\}$ .

Another interesting state is  $h_A^{\neg t_x^n, t_x}$ , *i.e.*, the source is a negative liar for all  $x \in A$  and truthful for all  $x \in \overline{A}$ , in which case  $x \in \mathcal{X}$  is a possible value for  $\mathbf{x}$  iff  $x \in B$  or  $x \in A$ , and thus one should conclude that  $\mathbf{x} \in B \cup A$ . More generally, as can be seen from Table 1, the couples  $(t_1, t_2) \in \mathcal{T}_x^2$  yields all possible binary Boolean connectives.

Of particular interest in this paper are the states  $h_A^{t_x, \neg t_x^p}$  and  $h_A^{\neg t_x^n, t_x}$ , which have already been discussed, and the states  $h_A^{t_x, \neg t_x}$  (the source is truthful for all  $x \in A$  and non truthful for all  $x \in \overline{A}$ ) and  $h_A^{\neg t_x, t_x}$  (the source is non truthful for all  $x \in A$  and truthful for all  $x \in \overline{A}$ ), which yield respectively  $\mathbf{x} \in B \cap A$  and

$\mathbf{x} \in B \sqcup A$ . Accordingly, we will consider in the sequel only the following subspace  $\mathcal{H} \subseteq \mathcal{T}$ :  $\mathcal{H} = \{h_A^{t_1, t_2} | A \subseteq \mathcal{X}, (t_1, t_2) \in \{(t_x, \neg t_x^p), (\neg t_x^n, t_x), (t_x, \neg t_x), (\neg t_x, t_x)\}\}$ .

Following [11], we can encode the above reasoning by a multivalued mapping  $\Gamma_B : \mathcal{H} \rightarrow \mathcal{X}$  indicating how to interpret the information  $\mathbf{x} \in B$  in each state  $h \in \mathcal{H}$ ; we have for all  $A \subseteq \mathcal{X}$ :

$$\Gamma_B(h_A^{t_x, \neg t_x^p}) = B \cap A, \Gamma_B(h_A^{\neg t_x^n, t_x}) = B \cup A, \Gamma_B(h_A^{t_x, \neg t_x}) = B \sqcap A, \Gamma_B(h_A^{\neg t_x, t_x}) = B \sqcup A.$$

If the knowledge about the source state is imprecise and given by  $H \subseteq \mathcal{H}$ , then one should deduce the image  $\Gamma_B(H) := \bigcup_{h \in H} \Gamma_B(h)$  of  $H$  by  $\Gamma_B$ .

### 3.3 Uncertain testimony and meta-knowledge

More generally, both the testimony provided by the source and the knowledge of  $Ag$  about the source truthfulness may be uncertain. Let  $m_S$  be the uncertain testimony and  $m^{\mathcal{H}}$  the uncertain meta-knowledge. In such case, the *Behavior-Based Correction* (BBC) procedure introduced by Pichon *et al.* [11], can be used to derive  $Ag$  knowledge on  $\mathcal{X}$ . It is represented by the MF  $m$  defined by [11]:

$$m(C) = \sum_{H \subseteq \mathcal{H}} m^{\mathcal{H}}(H) \sum_{B: \Gamma_B(H)=C} m_S(B), \quad \forall C \subseteq \mathcal{X}. \quad (5)$$

For convenience, we may denote by  $f_{m^{\mathcal{H}}}(m_S)$  the BBC of  $m_S$  according to meta-knowledge  $m^{\mathcal{H}}$ , *i.e.*, we have  $m = f_{m^{\mathcal{H}}}(m_S)$  with  $m$  defined by (5).

## 4 Interpretation of contextual corrections

In this section, we propose a new perspective on contextual discounting by recovering it using the framework introduced in Section 3. Then, using a similar reasoning, we provide an interpretation for contextual reinforcement. Finally, we introduce two new contextual correction schemes that are complementary to the two existing ones.

### 4.1 Contextual discounting in terms of BBCs

Let us consider a particular contextual lie among those introduced in Section 3.2: the states  $h_A^{\neg t_x^n, t_x}$ ,  $A \subseteq \mathcal{X}$ , corresponding to the assumptions that the source is a negative liar for all  $x \in A$  and truthful for all  $x \in \overline{A}$ . Among these states,  $h_{\emptyset}^{\neg t_x^n, t_x}$  admits a simpler interpretation: it corresponds to assuming that the source is truthful  $\forall x \in \mathcal{X}$ .

**Theorem 1.** *Let  $m_S$  be a MF. We have,  $\forall A$  and with  $\beta_A \in [0, 1]$ ,  $\forall A \in \mathcal{A}$ :*

$$m_S \bigcirc_{A \in \mathcal{A}} A_{\beta_A} = (\circ_{A \in \mathcal{A}} f_{m_{A, \cup}^{\mathcal{H}}})(m_S), \quad (6)$$

where  $\circ$  denotes function composition (*i.e.*,  $(g \circ f)(x) = g(f(x))$ ) and where  $m_{A, \cup}^{\mathcal{H}}$  is defined by  $m_{A, \cup}^{\mathcal{H}}(\{h_{\emptyset}^{\neg t_x^n, t_x}\}) = \beta_A$ ,  $m_{A, \cup}^{\mathcal{H}}(\{h_A^{\neg t_x^n, t_x}\}) = 1 - \beta_A$ ,  $\forall A \in \mathcal{A}$ .

*Proof.* This theorem can be shown by applying for each  $A \in \mathcal{A}$ ,  $\mathcal{A}$  being finite, the following property:

$$f_{m_{A,\cup}}^{\mathcal{H}}(m_S) = m_S \odot A_{\beta_A}, \quad \forall A \in \mathcal{A}, \quad (7)$$

which is shown as follows.

From (5) and the definition of  $m_{A,\cup}^{\mathcal{H}}$ ,  $\forall C \subseteq \mathcal{X}$ :

$$f_{m_{A,\cup}}^{\mathcal{H}}(m_S)(C) = \beta_A \sum_{B:B=C} m_S(B) + (1 - \beta_A) \sum_{B:B \cup A=C} m_S(B). \quad (8)$$

Which means:

$$f_{m_{A,\cup}}^{\mathcal{H}}(m_S) = \beta_A m_S + (1 - \beta_A)(m_S \odot m_A), \quad (9)$$

with  $m_A$  a MF defined by  $m_A(A) = 1$ .

On the other hand,  $\forall A \in \mathcal{A}$ :

$$m_S \odot A_{\beta_A} = m_S \odot \begin{cases} A \mapsto 1 - \beta_A \\ \emptyset \mapsto \beta_A \end{cases} = \beta_A m_S + (1 - \beta_A)(m_S \odot m_A). \quad (10)$$

□

In other words, contextual discounting, which appears on the left side of (6), corresponds to successive behavior-based corrections – one for each context  $A \in \mathcal{A}$  – where for each context  $A$ , we have the following meta-knowledge: with mass  $\beta_A$  the source is truthful for all  $x \in \mathcal{X}$ , and with mass  $1 - \beta_A$  the source is a negative liar for all  $x \in A$  and truthful for all  $x \in \bar{A}$ .

Successive corrections of an initial piece of information is a process that may be encountered when considering a chain of sources, where the information provided by an initial source may be iteratively corrected by the sources down the chain according to the knowledge each source has on the behavior of the preceding source. The chain of sources problem is an important and complex one, which has received different treatments in logic [4], possibility theory [1] and belief function theory [2, 3]: in particular a solution involving successive corrections, precisely successive discountings, was proposed in [1]. The fact that contextual discounting may be relevant for this problem had not been remarked yet.

## 4.2 Contextual reinforcement in terms of BBCs

Let us consider another kind of contextual lie: the states  $h_A^{t_x, \neg t_x^p}$ ,  $A \subseteq \mathcal{X}$ , corresponding to the assumptions that the source is truthful for all  $x \in A$  and a positive liar for all  $x \in \bar{A}$ . Among these states,  $h_{\mathcal{X}}^{t_x, \neg t_x^p}$  has the same simple interpretation as  $h_{\emptyset}^{\neg t_x^p, t_x}$ .



**Theorem 2.** Let  $m_S$  be a MF. We have,  $\forall A$  and with  $\beta_A \in [0, 1]$ ,  $\forall A \in \mathcal{A}$ :

$$m_S \bigodot_{A \in \mathcal{A}} A^{\beta_A} = (\circ_{A \in \mathcal{A}} f_{m_{A, \cap}^{\mathcal{H}}})(m_S), \quad (11)$$

where  $m_{A, \cap}^{\mathcal{H}}$  is defined by  $m_{A, \cap}^{\mathcal{H}}(\{h_{\mathcal{X}}^{t_x, \neg t_x^p}\}) = \beta_A$ ,  $m_{A, \cap}^{\mathcal{H}}(\{h_A^{t_x, \neg t_x^p}\}) = 1 - \beta_A$ ,  $\forall A \in \mathcal{A}$ .

*Proof.* The proof is similar to that of Theorem 1.  $\square$

Theorem 2 is important as it constitutes the first known interpretation for contextual reinforcement. It shows that, similarly to contextual discounting, contextual reinforcement (left side of (11)) corresponds to successive behavior-based corrections – one for each context. The only difference between the two correction mechanisms is what is assumed with mass  $1 - \beta_A$ : with the former that the source is a negative liar for all  $x \in A$  and truthful for all  $x \in \bar{A}$ , whereas with the latter that the source is truthful for all  $x \in A$  and a positive liar for all  $x \in \bar{A}$ .

*Example 1.* Let us consider a series of three agents: agent 1 reports to agent 2, who reports to agent 3. Let  $m_i$  denote the beliefs of agent  $i$  on  $\mathcal{X} = \{x_1, x_2, x_3\}$  and let  $m_i^{\mathcal{H}}$ ,  $i > 1$ , denote the meta-knowledge of agent  $i$  about agent  $i - 1$ . Furthermore, assume that  $m_2^{\mathcal{H}}(\{h_{\mathcal{X}}^{t_x, \neg t_x^p}\}) = 0.6$  and  $m_2^{\mathcal{H}}(\{h_{\{x_1, x_2\}}^{t_x, \neg t_x^p}\}) = 0.4$ , that is, agent 2 believes with mass 0.6 that agent 1 is truthful for all  $x \in \mathcal{X}$ , and with mass 0.4 that agent 1 is truthful for  $x_1$  and  $x_2$  and a positive liar for  $x_3$ . Suppose further that  $m_3^{\mathcal{H}}(\{h_{\mathcal{X}}^{t_x, \neg t_x}\}) = 0.8$  and  $m_3^{\mathcal{H}}(\{h_{\{x_2, x_3\}}^{t_x, \neg t_x}\}) = 0.2$ . From Theorem 2, we have

$$\begin{aligned} m_2 &= m_1 \bigodot \{x_1, x_2\}^{0.6}, \\ m_3 &= m_2 \bigodot \{x_2, x_3\}^{0.8}, \\ m_3 &= m_1 \bigodot \{x_1, x_2\}^{0.6} \bigodot \{x_2, x_3\}^{0.8}. \end{aligned}$$

### 4.3 Two new contextual correction mechanisms

Contextual discounting and contextual reinforcement are based on corrections induced by simple pieces of meta-knowledge  $m_{A, \cup}^{\mathcal{H}}$  and  $m_{A, \cap}^{\mathcal{H}}$  respectively. In practice, those pieces of meta-knowledge transform a testimony  $\mathbf{x} \in B$  as follows: they both allocate mass  $\beta_A$  to  $B$ , and mass  $1 - \beta_A$  to  $B \cup A$  and to  $B \cap A$ , respectively.

Now, as we have seen in Section 3.2, there exist states  $h_A^{t_1, t_2} \in \mathcal{T}$  that lead to other binary Boolean connectives than the disjunction and the conjunction. This suggests a way to extend contextual discounting and contextual reinforcement. Of particular interest are states  $h_A^{t_x, \neg t_x}$  (the source is truthful for all  $x \in A$  and non truthful for all  $x \in \bar{A}$ ) and  $h_A^{\neg t_x, t_x}$  (the source is non truthful for all  $x \in A$  and truthful for all  $x \in \bar{A}$ ), which yield respectively  $\mathbf{x} \in B \sqcap A$  and  $\mathbf{x} \in B \sqcup A$ . Indeed, the properties satisfied by connectives  $\sqcap$  and  $\sqcup$  allow us to obtain similar relations as those obtained for contextual discounting and contextual reinforcement:

**Theorem 3.** *Let  $m_S$  be a MF. We have,  $\forall \mathcal{A}$  and with  $\beta_A \in [0, 1]$ ,  $\forall A \in \mathcal{A}$ :*

$$m_S \bigcirc_{A \in \mathcal{A}} A^{\beta_A} = (\circ_{A \in \mathcal{A}} f_{m_{A, \sqcap}}^{\mathcal{H}})(m_S), \quad (12)$$

$$m_S \bigcirc_{A \in \mathcal{A}} A_{\beta_A} = (\circ_{A \in \mathcal{A}} f_{m_{A, \sqcup}}^{\mathcal{H}})(m_S), \quad (13)$$

where  $m_{A, \sqcap}^{\mathcal{H}}$  and  $m_{A, \sqcup}^{\mathcal{H}}$  are defined by  $m_{A, \sqcap}^{\mathcal{H}}(\{h_{\mathcal{X}}^{t_x, \neg t_x}\}) = m_{A, \sqcup}^{\mathcal{H}}(\{h_{\emptyset}^{\neg t_x, t_x}\}) = \beta_A$ , and  $m_{A, \sqcap}^{\mathcal{H}}(\{h_A^{t_x, \neg t_x}\}) = m_{A, \sqcup}^{\mathcal{H}}(\{h_A^{\neg t_x, t_x}\}) = 1 - \beta_A$ ,  $\forall A \in \mathcal{A}$ .

*Proof.* The proof is similar to that of Theorem 1.  $\square$

Eqs. (12) and (13) are the  $\sqcap$  and  $\sqcup$  counterparts to Eqs. (6) and (11), which are based on connectives  $\cup$  and  $\cap$ . Hence, if contextual discounting and contextual reinforcement are renamed as  $\sqcup$ -contextual correction and  $\sqcap$ -contextual correction, then Eqs. (12) and (13) may be called  $\sqcap$ -contextual correction and  $\sqcup$ -contextual correction. Let us also stress that although the  $\sqcap$  and  $\sqcup$ -contextual correction mechanisms are based on less classical combination rules than contextual discounting and contextual reinforcement, these two new contextual correction schemes seem to be as reasonable from the point of view of the meta-knowledge that they correspond to. Actually, their interpretations are even simpler since they rely on the classical assumptions of truthfulness and non truthfulness, whereas contextual discounting and contextual reinforcement involve negative and positive lies, which are less conventional. Finally, we note that the computational complexity of the  $\sqcap$  and  $\sqcup$ -contextual correction mechanisms is similar to that of  $\cup$  and  $\cap$ -contextual correction mechanisms: it merely corresponds to the complexity of applying  $|\mathcal{A}|$  combinations by the rules  $\bigcirc$  and  $\bigcirc$ , respectively, where  $|\mathcal{A}|$  denotes the cardinality of  $\mathcal{A}$ .

## 5 Conclusion

Using a new framework for handling detailed meta-knowledge about source truthfulness, a new view on contextual discounting and an interpretation for contextual reinforcement were proposed. In addition, two similar yet complementary contextual correction mechanisms were introduced.

Future work will be dedicated to the application of contextual correction mechanisms. Similarly as contextual discounting [7, 10], their parameters  $\beta_A$ ,  $A \in \mathcal{A}$ , could be obtained from a confusion matrix or learnt from training data, and then they could be used in classification problems. Other potential applications include those involving chain of sources communicating pieces of information between themselves, as is the case in vehicular ad-hoc networks [8].

## References

1. T. Baerecke, T. Delavallade, M.-J. Lesot, F. Pichon, H. Akdag, B. Bouchon-Meunier, P. Capet, and L. Cholvy. Un mode de cotation pour la veille informationnelle en source ouverte. In *6me colloque Veille Stratgique Scientifique et Technologique, Toulouse, France*, 2010.

2. L. Cholvy. Evaluation of information reported: A model in the theory of evidence. In E. Hüllermeier, R. Kruse, and F. Hoffmann, editors, *Proc. of the Int. Conf. on Information Processing and Management of Uncertainty in Knowledge-Based Systems, Dortmund, Germany*, volume 80 of *CCIS*, pages 258–267. Springer-Verlag, 2010.
3. L. Cholvy. Collecting information reported by imperfect information sources. In S. Greco, B. Bouchon-Meunier, G. Coletti, M. Fedrizzi, B. Matarazzo, and R. R. Yager, editors, *Proc. of the Int. Conf. on Information Processing and Management of Uncertainty in Knowledge-Based Systems, Catania, Italy*, volume 299 of *CCIS*, pages 501–510. Springer, 2012.
4. L. Cholvy. Un modèle de cotation pour la veille informationnelle en source ouverte. In *European Workshop on Multi-agent Systems, Toulouse, France*, 2013.
5. A. P. Dempster. Upper and lower probabilities induced by a multivalued mapping. *Annals of Mathematical Statistics*, 38:325–339, 1967.
6. D. Dubois and H. Prade. A set-theoretic view of belief functions: logical operations and approximations by fuzzy sets. *International Journal of General Systems*, 12(3):193–226, 1986.
7. Z. Elouedi, E. Lefèvre, and D. Mercier. Discountings of a belief function using a confusion matrix. In *Int. Conf. on Tools with Artificial Intelligence*, volume 1, pages 287–294, 2010.
8. M. Bou Farah, D. Mercier, E. Lefèvre, and F. Delmotte. A high-level application using belief functions for exchanging and managing uncertain events on the road in vehicular ad hoc networks. *Annals of telecommunications*, 69(3-4):185–199, 2014.
9. D. Mercier, E. Lefèvre, and F. Delmotte. Belief functions contextual discounting and canonical decompositions. *International Journal of Approximate Reasoning*, 53(2):146 – 158, 2012.
10. D. Mercier, B. Quost, and T. Dencœur. Refined modeling of sensor reliability in the belief function framework using contextual discounting. *Information Fusion*, 9(2):246 – 258, 2008.
11. F. Pichon, D. Dubois, and T. Dencœur. Relevance and truthfulness in information correction and fusion. *International Journal of Approximate Reasoning*, 53(2):159–175, 2012.
12. G. Shafer. *A mathematical theory of evidence*. Princeton University Press, Princeton, N.J., 1976.
13. Ph. Smets. The  $\alpha$ -junctions: combination operators applicable to belief functions. In D. M. Gabbay, R. Kruse, A. Nonnengart, and H. J. Ohlbach, editors, *1st Int. J. Conf. on Qual. and Quant. Pract. Reas., Bad Honnef, Germany*, volume 1244 of *Lecture Notes in Computer Science*, pages 131–153. Springer, 1997.
14. Ph. Smets and R. Kennes. The Transferable Belief Model. *Artificial Intelligence*, 66:191–243, 1994.