

Truth-Tracking by Belief Merge

Nina Gierasimczuk¹, Thomas Skafte²

¹*Department of Applied Mathematics and Computer Science, Technical University of Denmark*

²*Department of Computer Science, University of Luxembourg*

Abstract

In this paper we investigate truth-tracking properties of iterated belief revision methods. We focus on the multi-agent context when several agents revise their plausibility spaces simultaneously and arrive at joint beliefs through merge. We investigate the convergence properties of such a process: given a distributed sequence of information that truthfully and completely describes the actual world, will the process converge to true beliefs about the actual world? We focus on two methods that were previously proven to work for single-agents truth-tracking based on belief revision: conditioning and lexicographic revision. We show that some of the positive results carry on to our multi-agent setting. We introduce a novel kind of belief merge operation that is memory-based. We also study how to deal with occasional erroneous information by applying priority merge techniques.

Keywords

Truth-tracking, belief change, reasoning in multi-agent systems, belief merge, non-monotonic logic

1. Introduction

Belief revision theory offers a variety of ways to model the process of adapting knowledge to new information. These methods are often captured by a set of rationality postulates expressing their desired properties. For instance, one such postulate states that, after revision with a new piece of information, the resulting belief should be consistent with it. A similar approach is common in social choice theory, where the preferences of multiple agents are aggregated to produce a common, agreed-upon decision. Here, rationality postulates can, for instance, require that no agent should be able to dominate the outcome, i.e., that there should be no dictator. A different view of the rationality of belief revision takes learnability into account (see, e.g., [1]). The idea is that belief revision methods can be evaluated with respect to how well they track the truth: assuming that among the agent's uncertainty range there is a possible world s that accurately describes the true state of affairs, and the agent is given optimal conditions for learning about it, will she converge to s by using that belief revision method? In [2], it was shown that two popular belief revision methods are, under certain idealised conditions, effective for learning: conditioning (first explicitly discussed in [3]) and lexicographic revision (introduced in [4] and further developed in [5]). In this paper, we explore these positive results in the context of multi-agent learning, where several agents revise their beliefs in parallel and jointly converge to the truth. The collective aspect of this type of learning is captured using operators that merge the individual beliefs of agents at a given stage of belief revision. Ideally, the collective belief should converge to the designated true state.

How do different belief revision methods perform in such a collective learning scenario? There are several ways to address this question. In this paper, we assume that our agents start with the same prior plausibility ordering and are trying to learn the identity of the same actual world. They are also homogeneous, i.e., they each use the same belief revision method, but they differ in the information about the actual world they receive. It is easy to imagine situations in which their individual beliefs and conjectures about the real world are incorrect, while collectively they accumulate enough information to arrive at the truth.

24th International Workshop on Nonmonotonic Reasoning, July 17–19, 2026, Lisbon, Portugal

✉ nigi@dtu.dk (N. Gierasimczuk); thomas.skafte@uni.lu (T. Skafte)

id 0000-0001-5081-4676 (N. Gierasimczuk)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

Our objective is to introduce a framework for analysing the truth-tracking capacity of merging methods for groups of learning agents. Our aim is to identify merging operators that enable groups of learning agents to jointly behave as a single learning agent, either by mimicking such behaviour at each stage of learning or, at least, in terms of the ability to converge to the truth. We analyse two classical belief revision methods that have been shown to be suitable for learning in the single-agent case: conditioning and lexicographic revision [6, 7, 2]. We demonstrate the usefulness and robustness of our ‘truth-tracking by belief merge’ framework by applying it to the case of conditioning. The second part of the paper is more exploratory: we discuss the more complex case of lexicographic revision by considering a variety of methods and examples. We then propose a special type of memory-based belief merge operator, which we conjecture is suitable for collective truth-tracking with lexicographic agents. In the final part, we present a preliminary investigation of which collective learning methods can handle limited cases of incorrect information.

The paper is structured as follows. We begin Section 2 by recalling the setup of truth-tracking by belief revision from [2], which we extend to the multi-agent case in Section 3. In Section 4, we focus on groups of conditioning agents and show that conditioning remains a good learning method in the collective context by specifying an appropriate form of belief merge, which we call knowledge-intersection merge. In Section 5, we turn to lexicographic revision. We show that, although it is a good single-agent learning method, it does not perform well when combined with some seemingly reasonable merge methods. We then discuss a new merge method that could allow the single-agent universality of lexicographic revision to carry over to the multi-agent context. In Section 6, we shift our attention to streams that are closed under negation, i.e., for any proposition that can be observed, its negation is also observable. We discuss how one can partially relax the soundness of streams for learning by using priority merge. We outline some of the related work in Section 7. We conclude with a discussion of possible future directions in Section 8.

2. Notation and basic concepts

Let us start with the basic notion of *epistemic space* $\mathbb{S} = (S, \mathcal{O})$, with S a set of possible worlds, and $\mathcal{O} \subseteq \mathcal{P}(S)$ the set of observations (we assume both these sets are at most countable). A *plausibility space* is a $\mathbb{B} = (S, \mathcal{O}, \preceq)$ where $\preceq \subseteq S \times S$ is a plausibility preorder¹ on S . Given an epistemic space $\mathbb{S} = (S, \mathcal{O})$ we define the set of all plausibility spaces over $\mathbb{S}$ as $\mathbb{B}_{\mathbb{S}} = \{(S, \mathcal{O}, \preceq) \mid \preceq \text{ is a preorder on } S\}$. In this paper we will be concerned with one epistemic space being (initially) shared between multiple agents, hence the subscript $\mathbb{S}$ will be often suppressed from our notation. Given a plausibility space $\mathbb{B} = (S, \mathcal{O}, \preceq)$, $\min(\mathbb{B}) = \{s \in S \mid \text{there is no } t \in S, \text{ such that } (t, s) \in \prec\}$. As usual in belief revision literature, such a set of minimal possible worlds according to $\preceq$ determines the agent’s beliefs.

Observations For any $s \in S$, $\mathcal{O}_s = \{O \in \mathcal{O} \mid s \in O\}$ is the set of observations true in s , in other words information that is consistent with the world s . We assume that for any $s, t \in S$ if $s \neq t$, then $\mathcal{O}_s \neq \mathcal{O}_t$, i.e., there are no two worlds that make exactly the same observations true. An epistemic space is *negation closed* if for any observation its complement can also be observed, in other words, if $O \in \mathcal{O}$, then $\overline{O} \in \mathcal{O}$. A *stream* of observations ε is an infinite sequence of elements of $\mathcal{O}$. Let $n \in \mathbb{N}$. We will denote with $\varepsilon[n]$ an initial segment of ε of length $n + 1$, ε^n stands for the n -th element of ε , and $\text{set}(\varepsilon)$ for the set of elements enumerated in ε . Given $s \in S$, a stream is *sound* for s if $\text{set}(\varepsilon) \subseteq \mathcal{O}_s$ and it is *complete* for s if $\mathcal{O}_s \subseteq \text{set}(\varepsilon)$. The above notation will analogously apply to finite sequences of observations, $\sigma \in \mathcal{O}^*$. Given data sequence σ_1 and a data sequence or stream σ_2 , $\sigma_1 * \sigma_2$ stands for a concatenation of σ_1 with σ_2 .

Revision methods A *one-step revision method* is a function R_1 , which on the input of a plausibility space and an observation outputs a new, revised, plausibility space. In this paper we will be interested

¹A preorder on S is a reflexive and transitive binary relation. We will write $s \prec t$ iff $s \preceq t$ and $t \not\preceq s$. Moreover, the expression $(s, t) \in \preceq$ (or $s \preceq t$) means that s at least as plausible as t , which in the figures is represented as an arrow from t to s .

in two revision methods: conditioning and lexicographic revision.

Definition 1.

One-step conditioning is defined as:

$$\text{Cond}_1((S, \mathcal{O}, \preceq), O) = (S', \mathcal{O}^{S'}, \preceq^{S'}),$$

where $S' = S \cap O$, $\mathcal{O}^{S'} = \{O \cap S' \mid O \in \mathcal{O}\}$ and $\preceq^{S'} = \preceq \cap (S' \times S')$;²

One-step lexicographic revision is defined as:

$$\text{Lex}_1((S, \mathcal{O}, \preceq), O) = (S, \mathcal{O}, \preceq'),$$

where $(s, t) \in \preceq'$ if one of the following holds: $(s, t) \in \preceq^O$, $(s, t) \in \preceq^{\overline{O}}$ or $(s \in O \text{ and } t \notin O)$.

The above revisers are one-step functions. We will iterate them, i.e., apply to sequences of observations, by applying them recursively in the following way:

$$R(\mathbb{B}, \sigma * O) := R_1(R(\mathbb{B}, \sigma), O),$$

for $R \in \{\text{Cond}, \text{Lex}\}$.

Learners A learner is a function L that on the input of an epistemic space and a sequence of observations outputs a subset of the set of possible worlds, i.e., a conjecture. In other words, $L((S, \mathcal{O}), \sigma) \subseteq S$. We will say that an a learner L learns $s \in S$ if for any stream ε sound and complete for s , there is a $k \in \mathbb{N}$ such that for all $n \geq k$, $L(\mathbb{S}, \varepsilon[n]) = \{s\}$. We will say L learns $\mathbb{S}$ if it learns all $s \in S$. Finally, we will say that $\mathbb{S}$ is *learnable* if there is an L that learns it. A learner L is *universal* if it can learn any epistemic space that is learnable (by any learning function).³

Given a revision method R and a plausibility $\preceq$, a revision-based learner is defined to, at each stage, output the set of minimal worlds in the plausibility space (after it was adjusted by the revision method on the basis of incoming information):

$$L_R^{\preceq}((S, \mathcal{O}), \sigma) = \min(R((S, \mathcal{O}, \preceq), \sigma)).$$

We will say that $\mathbb{S} = (S, \mathcal{O})$ is learnable by revision method R if there is plausibility order $\preceq$ on S such that $L_R^{\preceq}$ learns $\mathbb{S}$.

In [2] it was shown that conditioning and lexicographic revisers are universal on sound and complete data streams. The prior plausibility assignments that guarantee the universality order the epistemic space in a way that allows the belief revision method to be *non-overgeneralizing*, meaning that initially the plausibility relation will prefer worlds that ‘tightly’ fit the incoming data. This ‘Occam’s Razor’-like approach produces a plausibility space where, as long as the incoming information is sound and complete with respect to the true world, at some point the agent begins to prioritize the real world over others. Crucially, for infinite spaces, such a prior plausibility order might have to be non-well-founded (i.e., for a time the plausibility give the ‘best’ possible world). Specifically, the appropriate plausibility which guarantees convergence is the so-called specialisation preorder $\preceq_{\text{occ}}$, given by $s \preceq_{\text{occ}} t$ iff $\mathcal{O}_s \subseteq \mathcal{O}_t$.

Example 1. Let $\mathbb{S} = (S, \mathcal{O})$ be an epistemic space, defined as follows: $S = \{x, y\}$ and $\mathcal{O} = \{p = \{x, y\}, q = \{y\}\}$, see Figure 1a. A plausibility preorder on this space can be given in at least two ways, as in Figures 1b and 1c. 1b prefers the state where both observations are true. This is not optimal for learning: if the actual world is x , the agent will never have a reason to start preferring x over y , because there is no observation true in y that is false at x . 1c shows a better preorder: if x is the actual world, then the agent will remain correct no matter what new sound information she receives. If y is the actual world then at some point she will receive the observation $\{y\}$ and she will have a reason to revise the preorder.

²We will use this notation throughout the paper: for a set $X \subseteq S$, we will write $\mathcal{O}^X$ and $\preceq^X$ to denote the restriction of $\mathcal{O}$ and $\preceq$ to X , respectively.

³Sound and complete streams create perfect conditions for learning. We will later relax this assumption by dropping soundness. For now, when ‘learning’ is used without qualification, we mean learnability from sound and complete data streams.

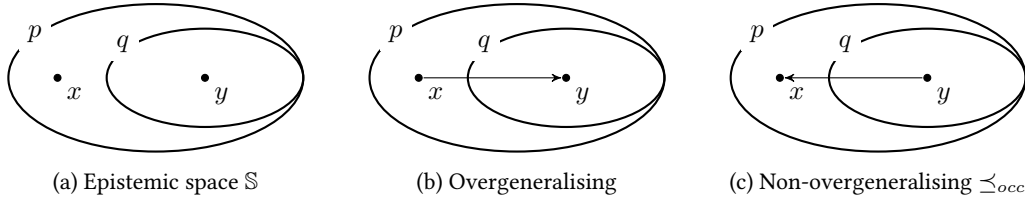


Figure 1: $\mathbb{S}$ from Example 1 with different plausibility preorders. Arrows point to more plausible worlds (lesser according to $\preceq$). For clarity of exposition, we omit reflexive arrows.

3. Collective learning by revision and merge

We will now extend our setup to multiple agents. Throughout we will work under the following assumptions. There are finitely many agents, $1, \dots, n$ and they all start with an identical plausibility space. One possible world, s , is in fact the actual world, and the agents do not know which one it is. They then preform their individual inquiry, by receiving information, each agent from her own data stream. These individual streams, while sound, are not necessarily complete with respect to s . The whole team however is assumed to collectively receive the complete picture, i.e., the sum of the contents of all of the streams is the set of all and only true observations at s . At each stage of their individual inquiry, the plausibility spaces are *merged*, by some plausibility merge function. Let us discuss some examples of such operations.

Plausibility merge A plausibility *merge* is a function that on the input of finitely many plausibility spaces outputs a new plausibility space. We will take $E = (\mathbb{B}_1, \dots, \mathbb{B}_n)$ to denote a *plausibility profile* that represents plausibility spaces of our n agents, all based on the same plausibility space $\mathbb{B} = (S, \mathcal{O}, \preceq)$. We consider the following intuitive candidates for merging methods:

- *knowledge-intersection merge* given by $f^{K\cap}(E) = (\bigcap_{i=1}^n S_i, \mathcal{O}^{S'}, \preceq^{S'})$.
- *plausibility-intersection merge* given by $f^{P\cap}(E) = (S, \mathcal{O}, \bigcap_{i=1}^n \preceq_i)$.
- *majority merge* given by $f^{maj}(E) = (S, \mathcal{O}, maj)$, where $(s, t) \in maj$ if $(s, t) \in \preceq_i$ for majority of plausibility spaces in E .

Each of the agents will revise their plausibility space by separate data stream. We then define a *data profile* $\Sigma = (\varepsilon_1, \dots, \varepsilon_n)$ of data streams giving inputs to each of the n agents individually. We extend the notation we specified for data streams to data profiles, such that $\Sigma[k] = (\varepsilon_1[k], \dots, \varepsilon_n[k])$, $\Sigma^k = (\varepsilon_1^k, \dots, \varepsilon_n^k)$, $set(\Sigma) = \bigcup_{i=1}^n set(\varepsilon_i)$. A collective version $\mathbf{R}$ of the single-agent revision method R (named with a bold-face $\mathbf{R}$, to distinguish from the single-agent revision method) on E and Σ at step k is defined as follows:

$$\mathbf{R}(E, \Sigma[k]) = \mathbf{R}((\mathbb{B}_1, \dots, \mathbb{B}_n), (\varepsilon_1[k], \dots, \varepsilon_n[k])) = (R(\mathbb{B}_1, \varepsilon_1[k]), \dots, R(\mathbb{B}_n, \varepsilon_n[k])).$$

When all plausibility spaces in E are the same and equal to $\mathbb{B}$ (this will be always the case in the initial step), instead of $\mathbf{R}(E, \Sigma[k])$ we will write $\mathbf{R}(\mathbb{B}, \Sigma[k])$.

Note that this definition effectively makes our learning homogenous (all agents apply the same reviser) and synchronous (all agent apply their updates at the same discrete pace). When referring to an initial segment of a profile of finite data sequences, the notation $\Sigma[k] = (\sigma_1[k], \dots, \sigma_n[k])$ will only make sense for k smaller or equal than the length of the shortest among the n sequences. There are several ways to overcome this limitation. One is to allow, as an element of data sequences and streams, an empty symbol λ that signifies an absence of input. Another, non-equivalent way is to redefine the sequences and streams so that at steps where the original does not give any input, the new version repeats the last-seen element. Finally, another and probably the least invasive way is to assume that there is a special ‘global’ observation $\mathbf{O} = S$, that corresponds to ‘no learning data’, as this observation is true in all possible worlds, so it does not carry any learning power.

A collective learner $\mathbf{L}$ (we again use a bold-face symbol to distinguish it from the single-agent kind) is applied to an epistemic space and a finite initial segment of a data profile and outputs a conjecture, i.e., $\mathbf{L}(\mathbb{S}, \Sigma[k]) \subseteq S$. Let us select one possible world $s \in S$ to be the true state. We will say that the data profile Σ is sound and complete with respect to s iff $\bigcup_{i=1}^n \text{set}(\varepsilon_i) = \mathcal{O}_s$. We will say that the collective learning method $\mathbf{L}$ learns $s \in S$ if for every Σ sound and complete for s , there is a $k \in \mathbb{N}$, such that for all $n \geq k$, $\mathbf{L}(\mathbb{S}, \Sigma[n]) = \{s\}$.

Let us now turn to how to base a learning method on concrete revisers and mergers. Note, that the result from collective revision $\mathbf{R}$ is a plausibility profile and so a merge f can be applied to it, yielding a joint plausibility space. That (joint) plausibility space can then be inspected for minimal elements that give a (collective) conjecture. We then obtain a collective learner based on merge f , revision method R and the prior plausibility preorder $\preceq$:

$$\mathbf{L}_{R,f}^{\preceq}((S, \mathcal{O}), \Sigma[k]) = \min(f(\mathbf{R}((S, \mathcal{O}), \preceq), \Sigma[k])).$$

We will say that a collective revision $\mathbf{R}$ and merge f learn $\mathbb{S}$, if there exists a prior plausibility assignment $\preceq$, such that $\mathbf{L}_{R,f}^{\preceq}$ learns $\mathbb{S}$.

4. Learning with conditioning and knowledge-intersection merge

We will first look at the multi-agent version of the conditioning revision method. Several observations concerning conditioning will be useful for us here. Firstly, it is straight-forward to see that conditioning with a sequence of observations can be replaced by conditioning with a single observation: a conjunction (intersection) of the elements of σ .

Lemma 1. *Let $(S, \mathcal{O}, \preceq)$ be a plausibility space and σ be a data sequence, then:*

$$\text{Cond}((S, \mathcal{O}, \preceq), \sigma) = \text{Cond}((S, \mathcal{O}, \preceq), \bigcap \text{set}(\sigma)).$$

The above observation directly implies that conditioning is *set-driven*, i.e., the effect of conditioning does not depend on the order in which the information is presented. Set-drivenness is an important property of learning functions in computational learning theory, introduced in [8] and further studied in [9].

Lemma 2. *For any two data sequences σ and σ' such that $\text{set}(\sigma) = \text{set}(\sigma')$, we have that*

$$\text{Cond}((S, \mathcal{O}, \preceq), \sigma) = \text{Cond}((S, \mathcal{O}, \preceq), \sigma')$$

Lemma 3. *Let $\Sigma = (\sigma_1, \dots, \sigma_n)$ be a data profile consisting of n data sequences, and let σ be a data sequence such that $\bigcup_{i=1}^n \text{set}(\sigma_i) = \text{set}(\sigma)$ (i.e., the profile and the sequence enumerate the same set of data). Then we have that $f^{K \cap}(\text{Cond}((S, \mathcal{O}, \preceq), \Sigma)) = \text{Cond}((S, \mathcal{O}, \preceq), \sigma)$.*

Proof. Note that $\text{Cond}((S, \mathcal{O}, \preceq), \sigma) = (S', \mathcal{O}^{S'}, \preceq^{S'})$, where $S' = \bigcap \text{set}(\sigma)$. Then, we get the following sequence of transformations:

$$f^{K \cap}(\text{Cond}(S, \mathcal{O}, \preceq), \Sigma) = f^{K \cap}(\text{Cond}((S, \mathcal{O}, \preceq), \sigma_1), \dots, \text{Cond}((S, \mathcal{O}, \preceq), \sigma_n)) = (S'', \mathcal{O}^{S''}, \preceq^{S''})$$

where $S'' = \bigcap_{i=1}^n (\bigcap \text{set}(\sigma_i))$ (by Lemma 1), which can be rewritten as $\bigcap (\bigcup_{i=1}^n \text{set}(\sigma_i))$, and then (by the assumption) is the same as $\bigcap \text{set}(\sigma)$. So, $S' = S''$. $\square$

Proposition 1. *The collective learner $\mathbf{L}$ based on Cond and $f^{K \cap}$ collectively learns $(S, \mathcal{O}, \preceq)$ if and only if the learner L based on Cond learns $(S, \mathcal{O}, \preceq)$.*

Proof. ($\Leftarrow$) Assume that L learns $(S, \mathcal{O}, \preceq)$ and, for contradiction, that $\mathbf{L}$ does not. That means that there is an $s \in S$ and sound and complete $\Sigma = (\varepsilon_1, \dots, \varepsilon_n)$ for s such that $\mathbf{L}$ does not converge to $\{s\}$. Let us consider such s and Σ . We construct a special data stream for s from Σ :

$$\varepsilon = \varepsilon_1^0, \varepsilon_2^0, \dots, \varepsilon_n^0, \varepsilon_1^1, \dots, \varepsilon_n^1, \varepsilon_1^2, \dots$$

First note that ε is sound and complete for s . Then, by Lemma 3, for all $k \in \mathbb{N}$ $f^{K \cap}(\mathbf{Cond}((S, \mathcal{O}, \preceq), \Sigma[k])) = \mathbf{Cond}((S, \mathcal{O}, \preceq), \varepsilon[nk])$. But that means that $\mathbf{Cond}$ never converges to $\{s\}$ on a ε for s . Contradiction.

($\Rightarrow$) Assume that $\mathbf{L}$ learns $(S, \mathcal{O}, \preceq)$ and, for contradiction, that L does not. That means that there is an $s \in S$ and ε for s such that L does not converge to $\{s\}$. Let us consider such an s and sound and complete ε . Consider the following special data profile $\Sigma = (\varepsilon_1, \dots, \varepsilon_n)$ for s constructed from ε in the following way: for any $k \in \{1, \dots, n\}$, $\varepsilon_k = \varepsilon^{k-1}, \varepsilon^{n+k-1}, \varepsilon^{2n+k-1}, \varepsilon^{3n+k-1}, \dots$. First note that Σ is sound and complete for s . Then, by Lemma 3, for all $k \in \mathbb{N}$, $f^{K \cap}(\mathbf{Cond}((S, \mathcal{O}, \preceq), \Sigma[k])) = \mathbf{Cond}((S, \mathcal{O}, \preceq), \varepsilon[nk])$. But that means that $\mathbf{L}$ never converges to $\{s\}$ on a Σ for s . Contradiction. $\square$

The construction that brings us back-and-forth between the data streams and data profiles in the above proof is very useful. Depending on the direction of its application we will call it *cyclic decomposition* (when we transition from a data stream to a data profile) or *cyclic unification* (when transitioning from data profile to data stream).

5. Learning with lexicographic revision and plausibility merge

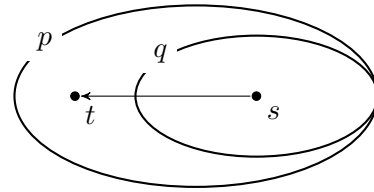
Conditioning removes the worlds that are inconsistent with the incoming information, and so it is a quite radical revision method. Lexicographic revision, on the other hand, is of a softer kind. It revises the preorder, rather than the set of states, which stays the same regardless of what information is received. The intersection merge on agents that iteratively perform lexicographic revision is then vacuous—it will always output the initial set of possible worlds. We then need a more sensitive kind of mergers that operate on the individual preorders rather than on the set of possible worlds. Applying the familiar ‘intersection methodology’ on the individual preorders will be our starting point. Let us recall the notion of plausibility-intersection merge, that we have described in Section 3: $f^{P \cap}(E) = (S, \mathcal{O}, \bigcap_{i=1}^n \preceq_i)$.

Note that both lexicographic revision method and the plausibility intersection merge work only on the plausibility preorder and leave the remaining components (S and $\mathcal{O}$ unchanged). We will hence simplify our notation, whenever possible, and refer to the plausibility space with $\preceq$, assuming S and $\mathcal{O}$ implicitly. So, we will write, for instance: $\text{Lex}(\preceq, \varepsilon[k]) = \preceq'$, instead of $\text{Lex}((S, \mathcal{O}, \preceq), \varepsilon[k]) = (S, \mathcal{O}, \preceq')$.

The intuition behind the plausibility intersection merger is that as the individual agents become more and more accurate with their plausibility representation (i.e., as the actual word becomes more and more plausible, given the sound data streams presented to the agents), the joint, shared part of their plausibility preorders would somehow reflect that progress. The example below shows that in the present setting that intuition is wrong.⁴

Example 2. Consider two agents working on a plausibility space $(S, \mathcal{O}, \preceq)$ containing worlds s and t along with two observations $\{p, q\}$. Let s be the actual world (denoted by $\underline{s}$).

- $S = \{s, t\}$
- $\mathcal{O} = \{p, q\}$, with $p = \{s, t\}$, $q = \{s\}$
- $\preceq = \{(t, s), (s, s), (t, t)\}$
- $\text{set}(\varepsilon_1) = \{p, q\}$
- $\text{set}(\varepsilon_2) = \{p\}$



The plausibility assignment initially prioritises t . The agents 1, 2 start with this initial plausibility space and they each receive their respective data streams ε_1 (consisting of observations p and q) and ε_2 (consisting only of the, repeated, observation p). At some finite stage k , 1 and 2 will have seen instances of all information available in their respective streams. Applying lexicographic revision individually will result in their individual plausibility diverging. For agent 1 $\text{Lex}(\preceq, \varepsilon_1[k]) = \{(s, t), (s, s), (t, t)\}$, and for agent 2 $\text{Lex}(\preceq, \varepsilon_2[k]) = \{(t, s), (s, s), (t, t)\}$. The intersection the two individual plausibility preorders will yield

⁴In the Figures below, we omit the reflexive part of $\preceq$ as it is not subject to change, and so it does not play any role in our reasoning.

a reflexive relation, $\{(s, s), (t, t)\}$, which does not single out one most plausible world. Hence, even though the agents jointly received a sound and complete data stream for the actual world s , they ended up with a plausibility order which does not favour s .

One might think that for larger groups of agents, such effects can be overcome by a solution familiar from voting scenarios—a merge that sanctions only those pairs on which the majority of the agents agree: $(s, t) \in maj$ iff $(s, t) \in \preceq_i$ for majority of plausibility spaces in E . It is easy to demonstrate that in our setting there is no mechanism to ensure that the majority would indeed receive the more informative data, i.e., no guarantee that the majority would be right, as the following example shows:

Example 3. Consider again the plausibility space from Example 2, but this time assume that we have three agents 1, 2, and 3. Agents 1 and 2 will be as in Example 2, and the agent 3 will be a copy of agent 2, i.e., receiving only the information p . After all of the agents have seen everything there is to see in their respective streams the resulting preorder will be as the initial $\preceq$, as the majority of agents (2 and 3) agree that t is more plausible than s . Since none of the agents will change their mind further, the minimal state will continue to be t , even though s is the actual world, and collectively the agents were given a sound and complete data stream for s .

What makes the plausibility intersection and majority merges insufficient, is that some agent's acquired evidence about the real plausibility 'drowns' in the abundance of the less-informed majority. The agents who lag behind in their updates (agents 2 and 3 in Examples 2 and 3) are taken into account to the same extent as the ones who actually did 'more work' on their plausibility spaces (like agent 1 in our Examples). As a result, whatever outstanding progress was made on discovering the actual world through updating plausibility, can be erased by the more idle agents. To overcome this problem, we need a different type of merge. Before we introduce them let us make two observations about set-drivenness of the method Lex. This allows comparing Cond and Lex with respect to how they depend on the order of received information.

Proposition 2. Lex is not set-driven.

The above fact can be demonstrated with the following example.

Example 4. Let the plausibility space $\mathbb{B} = (S, \mathcal{O}, \preceq)$, with $S = \{s, t, r\}$, $\mathcal{O} = \{p, q\}$, $\preceq = \{(r, t), (s, r), (s, t), (s, s), (t, t), (r, r)\}$ be the initial plausibility space of two agents 1 and 2, where $p = \{s, t\}$ and the true world is t . The top part of Figure 2 depicts that space. Below it, two vertical sequences of updates can be seen, the column on the left-hand side depicts two steps of updates on the sequence (p, q) , the right-hand side on (q, p) . The final outcomes differ in the relationship between s and r (red arrows).

We will now define a new type of merge, which prioritizes the parts of the preorders that were newly acquired during the learning process, relative to some initial, jointly agreed upon plausibility. We hypothesize, it is a good candidate for a merge that would match lexicographic revision.

Definition 2. A descendant merger is a function that, unlike previously discussed merging methods, apart from the plausibility profile of the group of agents, takes as an input the initial plausibility space:

$$f^d(\preceq, (\preceq_i)_{i \in \{1, \dots, n\}}) = \bigcup_{i \in \{1, \dots, n\}} \text{diff}(\preceq, \preceq_i) \cup (\text{cons}((\preceq_i)_{i \in \{1, \dots, n\}}) \setminus (\preceq \cup \preceq^{-1})) \cup \{(s, s) \mid s \in S\},$$

where:

- $\text{diff}(\preceq, \preceq_i) = \{(s, t) \mid (s, t) \in \preceq_i \text{ and } (t, s) \in \preceq\}$;
- $\text{cons}((\preceq_i)_{i \in \{1, \dots, n\}}) = \{(s, t) \mid \text{there is } i \text{ s.t. } (s, t) \in \preceq_i \text{ and there is no } j \text{ s.t. } (t, s) \in \preceq_j\}$.

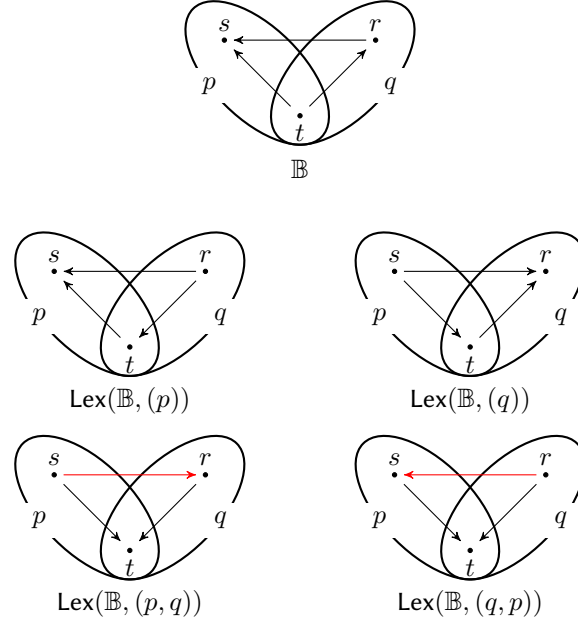


Figure 2: Lex is not set-driven

Intuitively, the descendant merger takes as the first argument some jointly agreed on (initial) plausibility order and, as the second argument, a plausibility profile. The operator `diff` checks how each agent in the profile revised the pairs initially comparable in $\preceq$. If an agent's plausibility for the pair was flipped with respect to the original $\preceq$, that new pair is collected; then we take the union over all agents of such flipped pairs. In other words, the first part of the equation collects the sum of all *changes* applied by the agents individually to the *initially comparable* pairs. As for the second part of the equation, note that the agents might have acquired comparability of states that were initially, according to $\preceq$, incomparable. The middle element of the equation above takes pairs (s, t) that are incomparable according to $\preceq$ (note the exclusion of the set $\preceq \cup \preceq^{-1}$) for which there is a group *consensus* (*cons*) on how they should be ordered (i.e., there is at least one agent whose plausibility supports (s, t) , and no agent has the opposite ordering (t, s)). The last part of the equations ensures that f^d outputs a reflexive relation.

Our conjecture regarding the characterization of learning with lexicographic revision and merge is as follows. Assume that a group of lexicographic revisers agrees to jointly start with a learnable epistemic space $(S, \mathcal{O})$ equipped with an Occamist plausibility preorder $\preceq_{occ}$, which is exactly the so-called specialization preorder: $(s, t) \in \preceq_{occ}$ iff $\mathcal{O}_s \subseteq \mathcal{O}_t$ (see [2, 10, 6]), and they receive a sound and complete data profile. We consider a collective learner $\mathbf{L}$ which outputs the initial plausibility preorder as long as `diff` is empty for all agents, and otherwise outputs the result of descendant merge f^d applied to the individual outputs. We conjecture that such an $\mathbf{L}$ collectively learns $(S, \mathcal{O}, \preceq_{occ})$. Moreover, we conjecture that under such learning conditions the descendant merge always produces a preorder. We leave the verification of these claims to the follow-up work. For the time being, let us consider some instructive examples.

Example 5. First let us consider the learning scenario from Example 2. For the step $k \in \mathbb{N}$, such that $set(\varepsilon_1[k]) = \{p, q\}$ and $set(\varepsilon_2[k]) = \{p\}$, we get that $\text{diff}(\preceq, \text{Lex}(\preceq, \varepsilon_1[k])) = \{(s, t)\}$ and $\text{diff}(\preceq, \text{Lex}(\preceq, \varepsilon_2[k])) = \emptyset$. So, $\bigcup_{i \in \{1, 2\}} \text{diff}(\preceq, \preceq_i) = \{(s, t)\}$. Also $\text{cons}((\preceq_i)_{i \in \{1, \dots, n\}}) \setminus (\preceq \cup \preceq^{-1}) = \emptyset$. Finally, after adding all reflexive pairs, we obtain a preorder, in which the true world s is the most plausible one. Since the streams can only repeat the previously seen information, that will never change.

To trace the steps of descendant merger in more detail let us consider a more complex scenario.

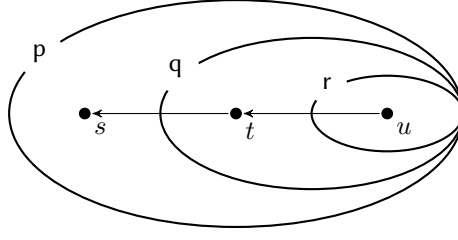


Figure 3: Plausibility space from Example 6, the transitive arrow from u to s is implicitly assumed.

Example 6. Consider the plausibility space $(S, \mathcal{O}, \preceq)$ in Figure 3 composed of the following elements: $S = \{s, t, u\}$, $\mathcal{O} = \{p, q, r\}$, with $p = \{s, t, u\}$, $q = \{t, u\}$, $r = \{u\}$, and $\preceq = \{(s, t), (t, u), (s, u), (s, s), (t, t), (u, u)\}$. The actual world is u .

We will consider two agents: 1 and 2, with their corresponding data streams, such that $\text{set}(\varepsilon_1) = \{p, q\}$ and $\text{set}(\varepsilon_2) = \{q, r\}$. Note that collectively the agents receive a sound and complete data stream for u , but individually neither of them gets a complete stream. For simplicity, let us assume that the streams enumerate the information efficiently, i.e., $\varepsilon_1[1] = (p, q)$ and $\varepsilon_2[1] = (q, r)$. Then, $\text{Lex}(\preceq, p) = \preceq$, and so $\text{diff}(\preceq, \text{Lex}(\preceq, p)) = \emptyset$ (agent 1 did not learn anything new). On the other hand, $\text{Lex}(\preceq, q) = \preceq' = \{(t, u), (u, s), (t, s), (s, s), (t, t), (u, u)\}$, and so $\text{diff}(\preceq, \preceq') = \{(u, s), (t, s)\}$. Together the agents learned: $\emptyset \cup \{(u, s), (t, s)\} = \{(u, s), (t, s)\}$. Further, they did not learn anything new about the relationship between u and t , nor about any other pairs (since all states were initially comparable). Hence, $f^d(\preceq, \text{Lex}(\preceq, \Sigma[0])) = \{(u, s), (t, s), (s, s), (t, t), (u, u)\}$. In the next step each agent continues their own inquiry: for agent 1, $\text{Lex}(\preceq, (p, q)) = \{(u, s), (t, s), (t, u), (s, s), (t, t), (u, u)\}$, and so $\text{diff}(\preceq, \text{Lex}(\preceq, (p, q))) = \{(u, s), (t, s)\}$ (agent 1 learned the same information that agent 2 learned in the previous step). For agent 2 we get: $\text{Lex}(\preceq, (q, r)) = \preceq'' = \{(u, t), (u, s), (t, s), (s, s), (t, t), (u, u)\}$, and so $\text{diff}(\preceq, \preceq'') = \{(u, t), (u, s), (t, s)\}$. Together the agents learned $\{(u, t), (u, s), (t, s)\}$. Hence, $f^d(\preceq, \text{Lex}(\preceq, \Sigma[1])) = \{(u, t), (u, s), (t, s), (s, s), (t, t), (u, u)\}$. The actual world u is the unique minimal element. Since the streams can only repeat the previously given information, neither of the sets of minimal elements of the individual preorders will change at any later stage.

Note that our pre-orders do not have to be total. The following example illustrates that fact. It also includes a non-trivial calculation of the set cons.

Example 7. Let us consider a case similar to the one in Example 4, and interpret Figure 4 as two agents: agent 1 learning on sequence (p, q) on the left-hand side and agent 2 learning on (p, p) sequence on the right-hand side (the actual world is t). First consider the starting plausibility preorder, which now needs to be augmented—the specialization pre-order $\preceq_{\text{occ}}$ that should be collectively imposed in the epistemic space to guarantee learning is depicted in top of Figure 4 below. Note that the states s and r are not comparable under $\preceq_{\text{occ}}$.

After the first step of inquiry, both agents learned that (t, r) with respect to the original $\preceq_{\text{occ}}$ and they also learned about a previously incomparable pair: they both have that (s, r) . Hence, the f^d at this stage will output $\{(t, r), (s, r), (s, s), (t, t), (r, r)\}$. In the next step, agent 1 learns q and so, with respect to the initial $\preceq_{\text{occ}}$ learns (t, s) and (t, r) , and (with respect to the initially incomparable worlds) (r, s) . Agent 2 is still at the stage in which with respect to $\preceq_{\text{occ}}$ it only learned (t, r) and with respect to the initially incomparable worlds (s, r) . At this stage the union of diff-sets for both agents is $\{(t, s), (t, r)\}$. The cons-set is now empty: the learned orderings of worlds incomparable according to the original $\preceq$ are opposite for both agents: (r, s) for agent 1 and (s, r) for agent 2 (red arrows). So, that this stage f^d outputs $\{(t, s), (t, r), (s, s), (t, t), (r, r)\}$. The most plausible world is the actual world t , and the resulting plausibility relation is reflexive and transitive, but not total.

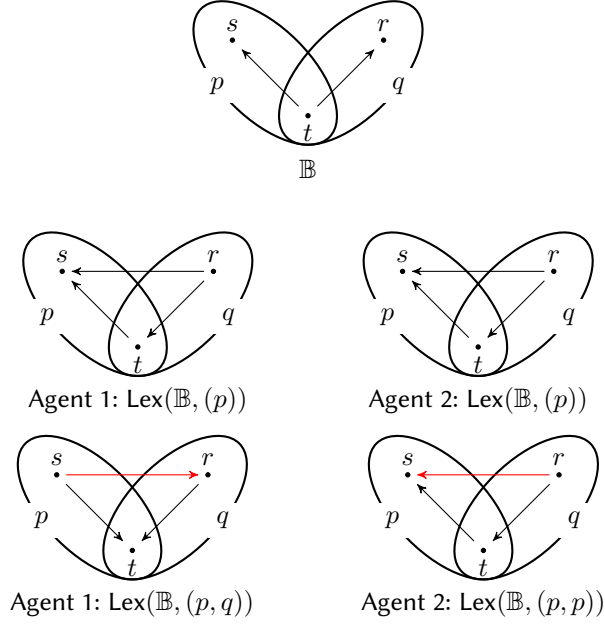


Figure 4: Descendant merge with non-trivial consensus set cons . The top is the starting plausibility pre-order $\preceq_{occ}$

6. Self-correcting streams

The requirement that each stream in the data profile ε_i is sound with respect to the true world s is very strong. For example if we take the perspective of observations coming from physical sensors then noise and errors are expected. To ease this requirement we apply the concept of self-correcting⁵ streams.

Definition 3.

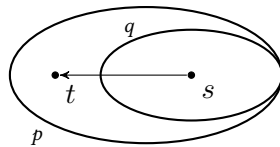
Let $\mathbb{S} = (S, \mathcal{O})$ be a negation-closed epistemic space, then a data stream (or sequence) ε is self-correcting with respect to world s if every incorrect observation is corrected by the later observations. Stated formally,

1. there is $k \in \mathbb{N}$ such that for all $n \geq k$, $s \in O_n$, and
2. for every $i \in \mathbb{N}$ such that $s \notin O_i$, there is some $k > i$, $O_k = \overline{O_i}$.

A data profile Σ is self-correcting when each data stream $\varepsilon_i \in \Sigma$ is self-correcting, and collectively-correcting when any incorrect observation is rectified at a later step in at least one data stream (or sequence).

In [2] it is shown that a single-agent learner using lexicographic revision still retains universality on negation-closed epistemic spaces with self-correcting and complete data streams. Let us now consider collectively-correcting data profiles.

Example 8. Consider the plausibility space of Example 2 but extend the set of observations to be negation-closed and instead of s being the true world, we set it to be t .



⁵In [2] self-correcting streams are called *fair*. We change the name in order to avoid confusion with various interpretations of ‘fairness’ in belief merge and social choice theory.

Let the two agents 1 and 2 receive the data streams $\varepsilon_1 = (q, p, \dots)$ and $\varepsilon_2 = (p, \bar{q}, \dots)$ respectively, which forms a collectively-correcting and complete data profile $\Sigma = (\varepsilon_1, \varepsilon_2)$ for t . Using the lexicographic revision the two agents will at $\Sigma[1]$ have the conflicting plausibility orders of $\text{Lex}(\preceq, \varepsilon_1[1]) = \{(s, t), (s, s), (t, t)\}$ and $\text{Lex}(\preceq, \varepsilon_2[1]) = \{(t, s), (s, s), (t, t)\}$. Applying the descendant merger on the profile of agents gives $\{s, t\}$, and so it selects the incorrect world to be minimal.

As we see a negation closed epistemic space that is learnable by Lex on self-correcting and complete streams is not necessarily learnable by Lex and f^d on collectively-correcting and complete data profiles. The issue is that when a conflict occurs between two plausibility orders, as in the example, the descendent merger simply selects the one that is the opposite of what the prior plausibility order without considering which is received later. When switching to collectively-correcting data profiles such temporal distinctions become important.

6.1. Handling errors with priority merge

As demonstrated above the location of errors and their corrections in the data profile is very important for the possibility of learning. It would be convenient to be able to identify the agents more responsible for erroneous observations. The concept of priority merge (see, e.g., [11]) could be useful here: the ‘more erroneous’ agents could be de-prioritized in the merge. A natural order of this kind can be induced from self-correcting streams, in which any occurrence of an error is truthfully accounted for later on. Sequential decomposition of such a stream over a set of revising agents is what we are looking for. Let us restrict for the time being to the finite case and assume that we have a self-correcting data sequence σ of length $m = n \cdot j$, we can then construct a finite data profile $\Sigma[j]$ of n data sequences, each of length j :

$$\text{if } \sigma = (\underbrace{p_1, \dots, p_j}_{\sigma_1}, \dots, \underbrace{p_k, \dots, p_{(i-1)j+1}}_{\sigma_i}, \dots, \underbrace{p_{m-j+1}, \dots, p_m}_{\sigma_n}), \text{ then } \Sigma[j] = (\sigma_1, \dots, \sigma_i, \dots, \sigma_n).$$

Because the sequence σ is self-correcting, any error is either corrected at a later stage in the sequence it appears in, σ_i , or in a posterior sequence σ_{i+k} , $i + k \leq n$, which gives a new type of a collectively-correcting data profile. This allows us to impose a priority on the agents such that for any two sequences σ_i, σ_j when $i < j$ the agent j will be prioritized. We will call such a plausibility profile *relevance increasing*, as the agents appearing later in the profile are more relevant for learning.

Merging on relevance increasing plausibility profiles still takes into account newly learned preferences, but since we allow corrections and errors to occur in two different agents we can run into disagreements. These conflicts are then handled by prioritizing agents as the relevance increasing property suggests.

Definition 4. We define a strict priority merge with defaulting operation $\rtimes$ in the following way:

$$\rtimes(\preceq_1, \preceq_2) = \preceq_1 \cup \{(s, t) \in \preceq_2 \mid (t, s) \notin \preceq_1\}.$$

We can then define descendant priority merge as:

$$f^\rtimes(\preceq, (\preceq_i)_{i \in \{1, \dots, n\}}) = \rtimes(\preceq_n \setminus \preceq, \rtimes(\preceq_{n-1} \setminus \preceq, \rtimes(\dots \rtimes (\preceq_1 \setminus \preceq, \preceq))))).$$

Intuitively, the strict priority merge with defaulting prioritizes everything that the later agents learned more than what the ones listed earlier in the profile learned.

The descendent priority merger selects relations that differ from the prior plausibility order in the same way that the descendent merger does, but when a conflict occurs it will prioritize the latter agent in the profile.

Example 9. Consider again the case in the Example 8. Applying the priority merge with defaulting on the two preorders $\preceq_1 = \{(s, t), (s, s), (t, t)\}$ and $\preceq_2 = \{(t, s), (s, s), (t, t)\}$ that the two agents arrived at will give: $f^\rtimes(\preceq, (\preceq_1, \preceq_2)) = \{(t, s), (s, s), (t, t)\}$. However, switching the two agents provides the opposite result $f^\rtimes(\preceq, (\preceq_2, \preceq_1)) = \{(s, t), (s, s), (t, t)\}$

When the data profile is ordered according to relevance we can handle conflicts between agents with the priority merge operation.

A construction we want to highlight is by the *cyclic decomposition* of a single infinite stream, the method we introduced and used in Section 4. It results in a specific distribution of observations that is not initially suited for descendant priority merge. We notice that when performing cyclic decomposition of ε , observations ε^0 and ε^n will be the first two observations for the first agent, see Table 1. We cannot rely on the relevance increasing property in the same manner, since observations $\varepsilon^1, \dots, \varepsilon^{n-1}$ will all be given to agents listed later in the profile than the agent that receives observation ε^n . However, the n observations that are distributed over the agents within one time step of the data profile Σ will be listed in the desired (self-correcting) order. We could then apply the strongly data retentive merge at each step, and iteratively apply the merge to next layers in a relevance increasing way.

$$\begin{array}{c|c}
 \begin{array}{c} \vdots \\ f(\preceq, \preceq_2 * \mathbf{R}(\mathbb{B}, (\varepsilon^{2n} \ \varepsilon^{2n+1} \ \dots \ \varepsilon^{3n-1}))) \\ f(\preceq, \preceq_1 * \mathbf{R}(\mathbb{B}, (\varepsilon^n \ \varepsilon^{n+1} \ \dots \ \varepsilon^{2n-1}))) \\ f(\preceq, \mathbf{R}(\mathbb{B}, (\varepsilon^0 \ \varepsilon^1 \ \dots \ \varepsilon^{n-1}))) \end{array} & \begin{array}{c} \vdots \\ = \preceq_3 \\ = \preceq_2 \\ = \preceq_1 \end{array} \\
 \hline
 \begin{array}{c} \varepsilon_1 \quad \varepsilon_2 \quad \dots \quad \varepsilon_n \end{array} &
 \end{array}$$

Table 1

Iterative merging on data profile constructed by cyclic decomposition.

7. Related Work

In this paper we presented a multi-agent extension of the symbolic single-agent learning by belief-revision framework from [6, 7, 2]. In the latter, the authors have shown the universality of their learners using concepts from formal learning theory, e.g., finite tell-tales. We build upon these results and investigate the power of the collective learners that operate on distributed streams and plausibility orders.

Multi-agent belief revision on plausibility spaces has been previously studied in the context of Dynamic Epistemic Logic in [12, 13], where notions of action and belief were used to give an account of reasoning about other agents. For now, in our approach we set aside the important issues of complex epistemic interactions between multiple agents identified, e.g., in social epistemology, but we are planning to address them in future work.

Belief merge strategies and *probabilistic* truth-tracking have been studied in preference aggregation, where [14, 15] develop a logical framework ensuring that merged belief profiles respect integrity constraints while representing agents' information, in contrast to our focus on convergence to a postulated actual world generating observations. In [16] the authors distinguish a synthetic view of merging, aimed at coherence, and an epistemic view assessing truth recovery; under assumptions of independent, homogeneous, and probabilistically reliable agents, they use Condorcet's Jury Theorem to analyse majority- and distance-based operators and their convergence to the true world under integrity constraints. Extending this, [17] replaces exact truth with truthlikeness, shifting attention from truth identification to closeness to truth and enabling evaluation when the true world is excluded by constraints, while [18] replace probabilistic convergence with worst-case robustness guarantees, providing a non-Bayesian justification for the improved reliability of belief fusion as the number of sources increases.

In a symbolic setting, truth-tracking with *non-expert information sources* has been recently investigated in [19, 20]. The authors step away from the assumption of soundness of the streams, by introducing conditional observations of the form 'source i reports that φ is possible in case c '. This approach allows for modelling non-expert streams and investigating truth-tracking capacities of different merging methods and various conditioning methods.

8. Conclusions and future work

In this paper, we investigate the truth-tracking properties of iterated belief revision methods in a multi-agent context. Learning takes place in a distributed fashion, and a common conjecture of the team of agents is reached by resorting to various kinds of plausibility merge operations. We focus on two methods that have previously been shown to be reliable for truth-tracking in the single-agent context: conditioning and lexicographic revision. We show that the results for conditioning carry over to the multi-agent case. We also outline a method that appears suitable for lexicographic revision. Our descendant merge operation is a novel approach to learning-oriented combination of beliefs. We also study how to handle occasional erroneous information by applying priority merge techniques.

The avenues for future work are numerous. Apart from completing the task of proving the universality of lexicographic revision with descendant merge on both sound and erroneous streams, we are interested in the following directions.

Our approach in this paper has been centred on specific merging operators rather than on generalising *families of merging operators* that satisfy certain useful properties. The first steps in this direction would be to identify the truth-tracking-friendly properties of plausibility mergers. Since collective learners are pairs consisting of a learning method and a plausibility merger, the characterisations expected of the mergers should most likely be formulated in terms of preserving the properties of single-agent learning methods. A similar question arises for minimal revision [21]. While it was shown in [2] that it does not, in general, generate universal learning methods, [22] focuses on characterising the epistemic spaces for which minimal revision is well-behaved. This direction can be further developed toward the multi-agent case along the lines of our framework.

So far, we have assumed that all streams in a data profile have the same domain of observations $\mathcal{O}$. However, this need not be the case. Allowing them to differ would provide a means of combining information from multiple sources into a set of observations Ω . A merging operator would then need to consider worlds that neither agent could represent individually, thereby increasing expressive power. This approach naturally aligns with combining information from different sensor types; for a human, such an example would be sight and hearing.

Finally, in the field of *robotics*, one of the central problems is that of *state estimation*—a robot cannot determine how to reach a desired goal without first knowing its current state. Since its introduction in 1960, the Kalman filter has been the primary workhorse for such problems. Its success partly stems from the robustness of combining multiple sources of observations to reduce noise in individual sensors [23]. In a sense, the concept of truth-tracking can be seen as an abstraction of state estimation, moving from numerical state spaces to possible worlds, with the added benefit of handling non-numerical data such as more complex epistemic notions.

Declaration on Generative AI

The authors have not employed any Generative AI tools.

References

- [1] K. T. Kelly, The learning power of belief revision, in: TARK'98: Proceedings of the 7th Conference on Theoretical Aspects of Rationality and Knowledge, Morgan Kaufmann Publishers Inc., San Francisco, CA, USA, 1998, pp. 111–124.
- [2] A. Baltag, N. Gierasimczuk, S. Smets, Truth-tracking by belief revision, *Studia Logica: An International Journal for Symbolic Logic* 107 (2019) 917–947. doi:10.1007/s11225-018-9812-x.
- [3] C. E. Alchourrón, P. Gärdenfors, D. Makinson, On the logic of theory change: Partial meet contraction and revision functions, *The Journal of Symbolic Logic* 50 (1985) 510–530. doi:10.2307/2274239.

- [4] A. C. Nayak, Iterated belief change based on epistemic entrenchment, *Erkenntnis* 41 (1994) 353–390. doi:10.1007/BF01130759.
- [5] A. Darwiche, J. Pearl, On the logic of iterated belief revision, *Artificial Intelligence* 89 (1997) 1–29. doi:10.1016/S0004-3702(96)00038-0.
- [6] N. Gierasimczuk, *Knowing One’s Limits. Logical Analysis of Inductive Inference*, Ph.D. thesis, Universiteit van Amsterdam, The Netherlands, 2010.
- [7] A. Baltag, N. Gierasimczuk, S. Smets, Belief revision as a truth-tracking process, in: *Proceedings of the 13th Conference on Theoretical Aspects of Rationality and Knowledge, TARK XIII*, Association for Computing Machinery, New York, NY, USA, 2011, p. 187–190. doi:10.1145/2000378.2000400.
- [8] D. Angluin, Inductive inference of formal languages from positive data, *Information and Control* 45(2) (1980) 117–135. doi:10.1016/S0019-9958(80)90285-5.
- [9] S. Lange, T. Zeugmann, Set-driven and rearrangement-independent learning of recursive languages, *Mathematical Systems Theory* 29 (1996) 599–634. doi:10.1007/BF01301967.
- [10] A. Baltag, N. Gierasimczuk, S. Smets, On the solvability of inductive problems: A study in epistemic topology, in: R. Ramanujam (Ed.), *Proceedings Fifteenth Conference on Theoretical Aspects of Rationality and Knowledge, TARK 2015*, Carnegie Mellon University, Pittsburgh, USA, June 4–6, 2015, volume 215 of *EPTCS*, 2015, pp. 81–98. doi:10.4204/EPTCS.215.7.
- [11] Z. Chistoff, N. Gratzl, O. Roy, Priority merge and intersection modalities, *The Review of Symbolic Logic* 15 (2022) 165–196. doi:10.1017/S1755020321000058.
- [12] A. Baltag, S. Smets, Dynamic belief revision over multi-agent plausibility models, in: *Proceedings of LOFT*, volume 6, University of Liverpool, 2006, pp. 11–24.
- [13] A. Baltag, S. Smets, A qualitative theory of dynamic interactive belief revision, in: H. Arló-Costa, V. F. Hendricks, J. van Benthem (Eds.), *Readings in Formal Epistemology: Sourcebook*, Springer International Publishing, Cham, 2016, pp. 813–858. doi:10.1007/978-3-319-20451-2_39.
- [14] S. Konieczny, R. P. Pérez, Merging with integrity constraints, in: A. Hunter, S. Parsons (Eds.), *Symbolic and Quantitative Approaches to Reasoning and Uncertainty*, Springer Berlin Heidelberg, Berlin, Heidelberg, 1999, pp. 233–244. doi:10.1007/3-540-48747-6_22.
- [15] S. Konieczny, R. P. Pérez, Merging information under constraints: A logical framework, *Journal of Logic and Computation* 12 (2002) 773–808. doi:10.1093/logcom/12.5.773.
- [16] P. Everaere, S. Konieczny, P. Marquis, The epistemic view of belief merging: Can we track the truth?, in: H. Coelho, R. Studer, M. J. Wooldridge (Eds.), *ECAI 2010 - 19th European Conference on Artificial Intelligence*, Lisbon, Portugal, August 16–20, 2010, *Proceedings*, volume 215 of *Frontiers in Artificial Intelligence and Applications*, IOS Press, 2010, pp. 621–626. doi:10.3233/978-1-60750-606-5-621.
- [17] S. D’Alfonso, Belief merging with the aim of truthlikeness, *Synthese* 193 (2016) 2013–2034. doi:10.1007/s11229-015-0825-y.
- [18] J. Karge, S. Rudolph, The More the Worst-Case-Merrier: A Generalized Condorcet Jury Theorem for Belief Fusion, in: *Proceedings of the 19th International Conference on Principles of Knowledge Representation and Reasoning*, 2022, pp. 205–214. doi:10.24963/kr.2022/21.
- [19] J. Singleton, R. Booth, Who’s the Expert? On Multi-source Belief Change, in: *Proceedings of the 19th International Conference on Principles of Knowledge Representation and Reasoning*, 2022, pp. 331–340. doi:10.24963/kr.2022/33.
- [20] J. Singleton, R. Booth, Truth-tracking with non-expert information sources, *J. Artif. Intell. Res.* 81 (2024) 619–641. doi:10.1613/jair.1.15273.
- [21] C. Boutilier, Iterated revision and minimal change of conditional beliefs, *Journal of Philosophical Logic* 25 (1996) 263–305. doi:10.1007/BF00248151.
- [22] E. Baccini, Z. Christoff, N. Gierasimczuk, R. Verbrugge, Who is afraid of minimal revision?, in: *Proceedings Twentieth Conference on Theoretical Aspects of Rationality and Knowledge (TARK 2025)*, volume 437 of *Electronic Proceedings in Theoretical Computer Science*, 2025, pp. 301–320. doi:10.4204/EPTCS.437.25.
- [23] T. D. Barfoot, *State Estimation for Robotics*, 2 ed., Cambridge University Press, 2024.