

Consistent update synthesis via privatized beliefs

Thomas Schlögl^{1*}, Roman Kuznets² and Giorgio Cignarale³

¹ Institute of Computer Engineering of TU Wien, Karlsplatz 13 1040 Vienna, Austria

² Institute of Computer Science of the Czech Academy of Sciences, Prague, Czechia

³ Faculty of Informatics, Theory and Logic Group, Vienna, Austria

* Correspondence: tschloegl@ecs.tuwien.ac.at (Schlögl T)

Abstract

Dynamic Epistemic Logic (DEL) formalizes communication in the form of model-transforming updates. Private communication is key in distributed systems, as processes exchanging information about their private local state should not be detectable by others. This focus on privacy clashes with the standard DEL assumption that updates are applied to the whole Kripke model, which is commonly known by all agents, potentially leading to information leakage. The contribution of this paper is twofold: (I) To represent leak-free agent-to-agent communication, we introduce a way to synthesize an action model that stratifies a pointed Kripke model into private agent-clusters, each representing the local knowledge of the processes: Given a goal formula representing private communication, we provide a procedure to construct an action model that: (a) makes the goal formula true; (b) maintains consistency of agents' beliefs, if possible; and (c) ensures minimal change of 'unrelated' beliefs. (II) We introduce a new operation between pointed Kripke models and pointed action models called *pointed updates* that, unlike the product update operation of DEL, maintains only the subset of the world-event pairs that are reachable from the point, without unnecessarily blowing up the model size.

Citation: Schlögl T, Kuznets R, Cignarale G. 2026. Consistent update synthesis via privatized beliefs. *The Knowledge Engineering Review* 41: e006
<https://doi.org/10.48130/ker-0026-0004>

1 Introduction

Epistemic logic (EL)^[1] has been extremely successful in modeling epistemic and doxastic attitudes of agents and groups in multi-agent systems, including distributed systems^[2]. *Dynamic epistemic logic* (DEL)^[3,4] upgrades EL by introducing model-transforming modalities called *updates*. Relational structures such as *action models* in Action Model Logic (AML) and *arrow update models* in the Generalized Arrow Update Logic (GAUL)^[5] are used to represent the evolution of agents' uncertainty under information change in complex communication scenarios. GAUL and AML have proved to be equally *update expressive*^[6]. Thus, without loss of generality, we use the term 'update models' to refer to either action models of AML or arrow update models of GAUL. In the update *synthesis* task, the aim is to (i) find whether there exists an update model that makes a given goal formula φ true and (ii) construct that update model from φ ^[6]. Because the desiderata for this goal are so strong, there is no general guarantee that such synthesis is always possible. For example, in standard public announcement logic^[4] and in arrow update logic^[7] it is not possible in general. However, Hales^[8] proved that it is possible for AML. While it is possible to construct AML or GAUL update models representing completely private communication^[9] (Fig. 2), there is no standardized update synthesis procedure for it.

Our paper is further motivated by two different yet intertwined issues:

- As argued by Artemov^[10], in multi-agent settings, common knowledge of the model (CKM) is required by agents in order to compute higher-order beliefs of other agents. While his argument focuses on uncertainties about facts, it can also be extended to agents' uncertainty about attitudes of other agents. The underlying problem is that, in multi-agent Kripke models, there is an implicit ontological distinction between two kinds of possible worlds: (a) worlds that are actually possible (AP); and (b) worlds that are only virtually possible (VP). While the former worlds constitute, for a given agent, the arena

in which the actual world might lie, the latter kind of worlds are considered only to the extent of computing other agents' beliefs. Artemov argues that without the CKM (comprising both APs and VPs), agents would not be able to compute such higher-order beliefs. Furthermore, avoiding the CKM assumption improves tolerance against local state corruption: upon receiving corrupted information from a faulty agent, a (correct) agent might only deem the local state of the sender as corrupted, instead of being forced to consider a larger part of the accessible Kripke model inconsistent.

- Update models do not naturally represent private agent-to-agent communication, as the product update operation typical of DEL is applied in principle to the full product of all the world-event pairs whose preconditions are matched^[4]. In this sense, these updates are applied globally to the whole model, potentially leading to information leakage. The problem is also identified by Herzig, as 'it is not easy to come up with a meaningful notion of a speaker'^[11], highlighting the fact that local communication might have (undesired) effects on a global model.

Our solution to both issues is to change the structure of the Kripke models under consideration. While Artemov proposes to avoid the CKM by changing the definition of truth for knowledge^[10], we instead suggest a stratified Kripke model structure (more precisely, the kind of structures that will be introduced for privatization are directed acyclic graphs.), without changing any basic definitions: by stratifying a Kripke model into a pointed *privatized Kripke model* we detach each agent's view of the model from any other agent's view of it. In our stratified models, not only do we clearly distinguish between APs and VPs for any agent, but we also avoid the CKM assumption by letting each agent access only a limited and private part of it (a formal definition of privatization is provided in Section 4), without losing the ability to reason about higher-order beliefs of other agents of arbitrary depths. At the same time, while we do not solve the broader issue of agency identified by Herzig, we propose an alternative update operation called *pointed updates*, which does not apply to the full product of

the Kripke model and the action model. In particular, we exploit the novel stratification method (and the distinction between APs and VPs) to discern when to form a world-event tuple and when not. In this sense, our updates are not applied to the whole model, as they are progressively applied only to those worlds that match a certain structure (other than a certain precondition). The proposed stratified structure complies with the crucial notion of a *local view* of an agent, typical of distributed systems, which is only a limited portion of the *global view* of the system, usually inaccessible to agents. The stratification not only represents the idea of each agent having a partial view of a Kripke model, but also distinguishes the actual world from any other worlds, by making it inaccessible. In other words, we endorse a *fallibilistic* assumption, that is, no agent can be sure that the global state of the system is captured in their limited view of it. Naturally, agents might have an accurate view of the system during execution, represented in our model by worlds propositionally equivalent to the actual world that are, however, accessible only in agents' (respective) private clusters.

The aim of this paper is to provide a novel action model synthesis mechanism for AML designed for enforcing private beliefs specified by a quantifier-free goal formula while preserving the consistency of agents' beliefs whenever possible. Our solution to the leak-free, consistent update synthesis task accounts for a limited range of goal formulas, restricted to deterministic belief increase only, i.e., to conjunctions of (positive) modal operators. This limitation is due to the fact that preserving minimality and consistency becomes highly problematic for unrestricted goal formulas: on the one hand, negations of belief operators might introduce ignorance, potentially contradicting previously held beliefs. On the other hand, disjunctions of modal operators have multiple realizations, making the update non-deterministic. Addressing the remaining cases (and their interactions) is left for future work.

Finally, the proposed update mechanism is generally more efficient with respect to AML and GAUL updates as it deletes all states not reachable from the actual world, making the growth in size of the model at worst linear in the size of the goal formula after one update and decreasing starting from the second update based on the same modal syntactic tree. By contrast, it is well known that repeated applications of AML and GAUL updates lead to the exponential growth of the model. It is also known that the model checking problem is PSPACE-complete^[12].

A prominent approach for dealing dynamically with beliefs in Kripke terms is the Dynamic Belief Revision^[13,14], where agents' conditional beliefs are expressed via a plausibility relation ranging over virtual states. This account, however, cannot represent truly private updates, as its models encode both plausibility and epistemic relations that would be revealed by the CKM assumption. In recent years, different answers to the limitations of the CKM assumption in the standard semantics for EL and DEL have been proposed, for example by changing accessibility relations^[10] or by using belief bases as primitives^[15,16]. Our approach, on the other hand, performs private, leakage-free consistent synthesis by structuring standard Kripke models into exclusively accessible parts that are not known to other agents, let alone commonly known.

Existing synthesis methods typically work in a language extended with quantifiers over updates, such as the Arbitrary Action Modal Logic^[8] and Arbitrary Arrow Update Modal Logic^[6]. While quantification over updates allows us to express the synthesis operation within the logical language, a major shortcoming, however, in particular with respect to the applications for distributed systems is that these update operations do not focus on issues of privatization and of minimal change. In addition, most existing synthesis methods do not address belief consistency preservation^[6].

We start by giving a motivating example in Section 2, as well as introducing the basic definitions of DEL and the crucial new definition of pointed update. In Section 3 we propose a solution to the leak-free consistent update synthesis for a limited range of goal formulas and we show its fruitfulness by applying it to our motivating example. In Section 4, we illustrate the properties of the synthesized action models, in particular with respect to privatization properties, crucial for avoiding information leakage while preserving consistency. Finally, conclusions are provided in Section 5.

2 Formal preliminaries and motivation

We use the standard doxastic multimodal language $\mathcal{L}$:

$$\varphi ::= p \mid \neg\varphi \mid (\varphi \wedge \varphi) \mid B_i\varphi$$

where $i \in \mathcal{A} = \{1, \dots, n\}$ (for $n > 1$) and $p \in Prop$, where $Prop$ is the set of propositional atoms. Formula $B_i\varphi$ means *agent i believes φ (to be true)*. The other Boolean connectives are defined as usual and $\hat{B}_i := \neg B_i \neg$ representing the dual modality *considers possible*.

Definition 1 (Kripke frame) A Kripke frame for the set of agents $\mathcal{A}$ is a pair $\langle W, R \rangle$ where $W \neq \emptyset$ is a non-empty set, called domain and $R = (R_1, \dots, R_n)$ consists of binary accessibility relations $R_i \subseteq W \times W$.

Unless specified otherwise, we assume $\mathcal{KD45}$, which is the class of frames that satisfy:

- **transitivity**: if $uR_i v$ and $vR_i w$ then $uR_i w$.
- **euclidean**: if $uR_i v$ and $uR_i w$ then $vR_i w$.
- **seriality**: every state has an outgoing arrow for every $i \in \mathcal{A}$.

These three frame properties correspond to the following axioms^[17]:

- **transitivity** $\Leftrightarrow B_i\varphi \rightarrow B_i B_i\varphi$.
- **euclidean** $\Leftrightarrow \neg B_i\varphi \rightarrow B_i \neg B_i\varphi$.
- **seriality** $\Leftrightarrow B_i\varphi \rightarrow \hat{B}_i\varphi$.

Note that we generally do not assume the factivity of beliefs ($B_i\varphi \rightarrow \varphi$), which in terms of frame properties corresponds to reflexivity ($uR_i u$).

We introduce the standard definitions of DEL^[4]:

Definition 2 (Pointed Kripke model) For a set of agents $\mathcal{A}$, a Kripke model is a triple $\mathcal{M} = \langle S, R, V \rangle$ where $\langle S, R \rangle$ is a Kripke frame with domain S consisting of possible worlds or states, and the valuation function $V: Prop \rightarrow 2^S$ determines the set $V(p) \subseteq S$ of possible worlds where an atom $p \in Prop$ is true. A pointed Kripke model is a pair $(\mathcal{M}, w)$ where $w \in S$ represents the actual world/state.

Truth at world w of model $\mathcal{M}$ is determined by $\mathcal{M}, w \models p$ iff $w \in V(p)$. Boolean connectives are defined as usual. $\mathcal{M}, w \models B_i\varphi$ iff $\mathcal{M}, w' \models \varphi$ for all $w' \in S$ such that $wR_i w'$. A formula φ is false at world w , $\mathcal{M}, w \not\models \varphi$, iff it is not true at w .

Definition 3 (Pointed action model) For a set of agents $\mathcal{A}$, an action model is a triple $\mathcal{U} = \langle E, Q, pre \rangle$ where $\langle E, Q \rangle$ is a Kripke frame with domain E consisting of actions or events, and the precondition function $pre: E \rightarrow \mathcal{L}$ assigns the precondition $pre(\beta) \in \mathcal{L}$ that is necessary for an event $\beta \in E$ to happen. A pointed action model is a pair $(\mathcal{U}, \alpha)$ where $\alpha \in E$ represents the actual event/action.

Definition 4 (Product update) The (restricted modal) product update of a Kripke model $\mathcal{M} = \langle S, R, V \rangle$ with an action model $\mathcal{U} = \langle E, Q, pre \rangle$ is a Kripke model $\mathcal{M} \otimes \mathcal{U} := \langle S', R', V' \rangle$ where

- $S' := \{(v, \beta) \in S \times E \mid \mathcal{M}, v \models pre(\beta)\}$,
- $R'_i := \{((v, \beta), (u, \gamma)) \in S' \times S' \mid (v, u) \in R_i \text{ and } (\beta, \gamma) \in Q_i\}$,
- $V'(p) := \{(v, \beta) \in S' \mid v \in V(p)\}$.

If $S' = \emptyset$, the product update is undefined.

A product update provides the semantics for a communication scenario represented by a pointed action model $(\mathcal{U}, \alpha)$: formula φ is

true after the communication $\mathcal{U}$, i.e., $\mathcal{M}, w \models [\mathcal{U}, \alpha]\varphi$, iff $\mathcal{M} \otimes \mathcal{U}, (w, \alpha) \models \varphi$ or $(w, \alpha) \notin S'$.

Definition 5 (Bisimulation [Kripke frames and Kripke models]) A bisimulation between Kripke frames $\langle W, R \rangle$ and $\langle W', R' \rangle$ is a nonempty binary relation $\mathcal{B} \subseteq W \times W'$, such that for every $v \mathcal{B} v'$ and $i \in \mathcal{A}$:

- **Forth:** if $v R_i u$, then there is $u' \in S'$ such that $v' R'_i u'$ and $u \mathcal{B} u'$;
- **Back:** if $v' R'_i u'$, then there is $u \in S$ such that $v R_i u$ and $u \mathcal{B} u'$.

A bisimulation between Kripke models $\mathcal{M} = \langle S, R, V \rangle$ and $\mathcal{M}' = \langle S', R', V' \rangle$ is a bisimulation $\mathcal{B} \subseteq S \times S'$ between their frames $\langle S, R \rangle$ and $\langle S', R' \rangle$ that additionally satisfies for every $v \mathcal{B} v'$ and $p \in \text{Prop}$:

- **Atoms:** $v \in V(p)$ iff $v' \in V'(p)$.

Similarly a bisimulation between action models $\mathcal{U} = \langle E, Q, \text{pre} \rangle$ and $\mathcal{U}' = \langle E', Q', \text{pre}' \rangle$ is a bisimulation $\mathcal{B} \subseteq E \times E'$ between their frames $\langle E, Q \rangle$ and $\langle E', Q' \rangle$ that additionally satisfies for every $\alpha \mathcal{B} \alpha'$:

- **Pre:** $\text{pre}(\alpha) = \text{pre}'(\alpha')$.

Two pointed Kripke models $(\mathcal{M}, v)$ and $(\mathcal{M}', v')$ (respectively pointed action models $(\mathcal{U}, \alpha)$ and $(\mathcal{U}', \alpha')$) are bisimilar denoted $\mathcal{M}, v \leftrightarrow \mathcal{M}', v'$ (resp. $\mathcal{U}, \alpha \leftrightarrow \mathcal{U}', \alpha'$), iff there exists a bisimulation $\mathcal{B}$ between $\mathcal{M}$ and $\mathcal{M}'$ (resp. $\mathcal{U}$ and $\mathcal{U}'$) s.t. $v \mathcal{B} v'$ (resp. $\alpha \mathcal{B} \alpha'$).

Definition 6 (Modal Equivalence) Pointed Kripke models $(\mathcal{M}, v)$ and $(\mathcal{M}', v')$ are called modally equivalent, notation $\mathcal{M}, v \equiv \mathcal{M}', v'$, whenever $\mathcal{M}, v \models \varphi$ iff $\mathcal{M}', v' \models \varphi$ for all formulas φ .

Theorem 7: (Bisimilarity implies modal equivalence^[18]) If $\mathcal{M}, v \leftrightarrow \mathcal{M}', v'$, then $\mathcal{M}, v \equiv \mathcal{M}', v'$.

2.1 Motivational example

We illustrate our objective by means of the following example:

Example 8 (Balder–Loki–Thor example) Teenage brothers Balder, Loki, and Thor prepare for an exam that, as is commonly known among them, contains one question asking to decide whether p or $\neg p$ is the case. Suppose the correct answer is $\neg p$. The initial model

$\mathcal{M}$ is the left model of Fig. 1. While the three are studying, Loki, unbeknownst to Thor, tries to trick Balder: Loki lies to Balder that he has overheard Thor boasting to know the correct answer to be p . Balder seems to dismiss Loki's claim as a ruse, leaving Loki thinking his trick had no effect. In truth, however, Balder does believe Loki and is now under the impression that Loki agrees with Thor that the answer is p . Balder himself, on the other hand, is not going to take Thor's word on it, given Thor's propensity to boast and act rashly. Thus, using b , l , and t for the brothers, Loki's trick should change Balder's beliefs to achieve the goal formula:

$$\varphi = B_b(B_l p \wedge B_l B_l p \wedge B_l p) \quad (1)$$

without either affecting Loki's or Thor's beliefs or making $B_b p$ true.

Since ignorance of the correct answer is common belief in the initial model $\mathcal{M}$, the public announcement of φ would cause everybody's beliefs to become inconsistent, whereas a private message $B_l p \wedge B_l B_l p \wedge B_l p$ would cause the same inconsistency, but for Balder's beliefs only. This happens whether one uses the world-removing or arrow-removing updates.

Instead, we propose an update method that stratifies the initial model $\mathcal{M}$ into several clusters representing individual beliefs and updating only some of these clusters, depending on the modal structure of the goal formula, resulting in model $\mathcal{M}^U$ in Fig. 1, where updated clusters are represented by light-gray rectangles. Cluster 0 contains only the actual world v , which no agent considers possible. Cluster -1 (the sink) is the copy of the initial model representing (higher-order) beliefs of agents unaffected by the message, including Loki's and Thor's beliefs. In particular, the ignorance of the correct answer is still common belief between Loki and Thor. Balder is also ignorant of the correct answer: his beliefs are represented by cluster 1, so $\mathcal{M}^U, v \not\models B_b p \vee B_b \neg p$. But his ignorance is not anymore common with Loki and Thor. Indeed, according to cluster 2, which represents Balder's beliefs about Thor's beliefs, $\mathcal{M}^U, v \models B_b B_l p$. Cluster 3 plays the same role for Balder's beliefs about Loki's beliefs, making

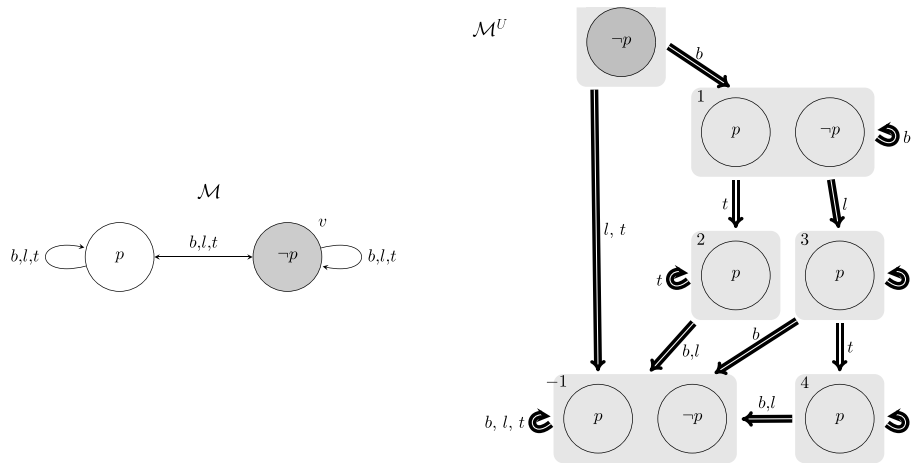


Fig. 1 Left: Initial pointed Kripke model $\mathcal{M}$. Right: Updated model $\mathcal{M}^U$ (the actual world v is dark gray). A formula φ within a circle representing a world means that φ is true in this world. A light-gray rectangle with n in the top left corner is the cluster that is numbered n . A double arrow from such cluster n to cluster m represents a set of accessibility arrows with the same agent's label from every world of cluster n to every world of cluster m .

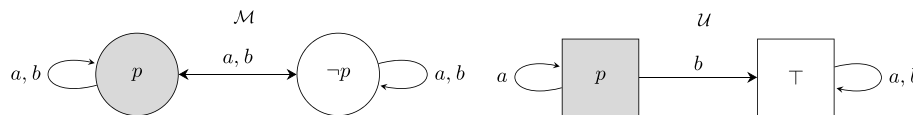


Fig. 2 Left: Initial (pointed) Kripke model $\mathcal{M}$ where agent a and b are uncertain about whether p or $\neg p$ is true. Right: (Pointed) action model $\mathcal{U}$ for a private message p to a : agent a learns p , while b believes that nothing happened. A formula φ within a square representing an event β is the precondition for this event, i.e., $\text{pre}(\beta) = \varphi$.

$M^U, v \models B_b B_l p$. Thus, while himself being ignorant of the answer, Balder thinks that Loki and Thor have chosen for themselves. Finally, cluster 4 is Balder's beliefs about Loki's beliefs about Thor's beliefs so that $M^U, v \models B_b B_l B_l p$. Overall, $M^U, v \models \varphi$. The b -labeled double arrow from cluster 3 to the sink -1 represents two b -labeled arrows from the only world of cluster 3 to each of the two worlds of the sink and ensures that Balder believes that Loki believes that Balder has not changed his beliefs, $M^U, v \models B_b B_l (\neg B_b p \wedge \neg B_b \neg p)$. Similarly, the b, l -labeled double arrows from clusters 2 and 4 to the sink signify that in no scenario where Thor has chosen an answer for himself, does he expect others to follow his example (or even be aware of his choice). It is easy to see that no agent has inconsistent beliefs, $M^U \models \neg B_b \perp \wedge \neg B_l \perp \wedge \neg B_t \perp$, and that Loki's and Thor's beliefs are the same as in the initial model, $M, v \models B_a \psi$ iff $M^U, v \models B_a \psi$ for $a \in \{l, t\}$.

In the following sections we will demonstrate how to synthesize an action model so that after the initial model M is updated with it, we arrive at the desired updated model M^U .

2.2 Pointed updates

It is well known that repeatedly applying product updates for multi-round communication can cause an exponential blow up of the domain of the model, even for simple models and action models, see Figs 2–4. Indeed, every updated model has exactly one world where p is false. Hence, a further update of a model of N worlds with $\mathcal{U}$ creates a model with $2N - 1$ worlds.

The standard (often implicit) solution is to reduce the size of the resulting model by using bisimulation contraction^[19], making the model smaller again. For example, bisimulation contraction deletes all worlds that are unreachable from the resulting actual world, resulting in a bisimilar and hence, modally equivalent^[18], model.

To avoid this two-step procedure, we introduce the notion of *pointed updates*, which incorporates this world-removing operation formally and explicitly:

Definition 9 (Pointed update) Let $(M, w) = (\langle S, R, V \rangle, w)$ be a pointed Kripke model and $(\mathcal{U}, \alpha) = (\langle E, Q, \text{pre} \rangle, \alpha)$ be a pointed action model. If $M, w \models \text{pre}(\alpha)$, then we define the updated pointed Kripke

model $(M \odot \mathcal{U}, (w, \alpha))$ where $M \odot \mathcal{U} := \langle S^{\mathcal{U}}, R^{\mathcal{U}}, V^{\mathcal{U}} \rangle$ as follows. Let $T^{\mathcal{U}} := \{(v, \beta) \in S \times E \mid M, v \models \text{pre}(\beta)\}$. Note that $(w, \alpha) \in T^{\mathcal{U}}$.

• The domain $S^{\mathcal{U}}$ is the smallest subset of $T^{\mathcal{U}}$ containing (w, α) such that:

if $(v, \beta) \in S^{\mathcal{U}}$, $(u, \gamma) \in T^{\mathcal{U}}$, and both $v R_i u$ and $\beta Q_i \gamma$ for some agent i , then $(u, \gamma) \in S^{\mathcal{U}}$;

• $R_i^{\mathcal{U}} := \{(v, \beta), (u, \gamma) \in S^{\mathcal{U}} \times S^{\mathcal{U}} \mid (v, u) \in R_i \text{ and } (\beta, \gamma) \in Q_i\}$;

• $V^{\mathcal{U}}(p) := \{(v, \beta) \in S^{\mathcal{U}} \mid v \in V(p)\}$.

To show that pointed updates produce the same result as product updates but with smaller models, we use the notion of bisimulation.

Definition 10 (G-bisimulation) Pointed Kripke models (M, v) and (M', v') are G -bisimilar for a group $G \subseteq \mathcal{A}$ of agents, notation $M, v \leftrightarrow_G M', v'$, iff for any $a \in G$

• **G-Forth:** if $v R_a u$, then there is $u' \in S'$ such that $v' R'_a u'$ and $M, u \leftrightarrow M', u'$;

• **G-Back:** if $v' R'_a u'$, then there is $u \in S$ such that $v R_a u$ and $M, u \leftrightarrow M', u'$.

The definition of G -bisimilarity for pointed action models is analogous.

Definition 11 (G-indistinguishability) Pointed Kripke models (M, v) and (M', v') are called G -indistinguishable for group $G \subseteq \mathcal{A}$, notation $M, v \equiv_G M', v'$, whenever $M, v \models B_a \varphi$ iff $M', v' \models B_a \varphi$ for all formulas φ and agents $a \in G$.

Theorem 12: (G-bisimilarity implies G-indistinguishability) If $M, v \leftrightarrow_G M', v'$, then $M, v \equiv_G M', v'$.

Proof Because **G-Forth** and **G-Back** (Definition 10) hold, the statement follows from Theorem 7 and Definition 11 as $M, u \equiv M', u'$ in both cases. $\square$

Theorem 13: (Product and pointed updates are bisimilar) Given a pointed Kripke model (M, w) and pointed action model $(\mathcal{U}, \alpha)$ such that $M, w \models \text{pre}(\alpha)$, we have $M \odot \mathcal{U}, (w, \alpha) \leftrightarrow M \otimes \mathcal{U}, (w, \alpha)$.

Proof: Let $M = \langle S, R, V \rangle$, $\mathcal{U} = \langle E, Q, \text{pre} \rangle$, $M \odot \mathcal{U} = \langle S^{\mathcal{U}}, R^{\mathcal{U}}, V^{\mathcal{U}} \rangle$ and $M \otimes \mathcal{U} = \langle S', R', V' \rangle$. By construction, $(w, \alpha) \in S^{\mathcal{U}} \subseteq S'$, and $R^{\mathcal{U}} = R'_{\upharpoonright_{(S^{\mathcal{U}} \times S^{\mathcal{U}})}}$, and $V^{\mathcal{U}} = V'_{\upharpoonright_{S^{\mathcal{U}}}}$. Hence, it is easy to see that $\mathcal{B} := \{(v, \beta), (v, \beta) \mid (v, \beta) \in S^{\mathcal{U}}\}$ is the required bisimulation. $\square$

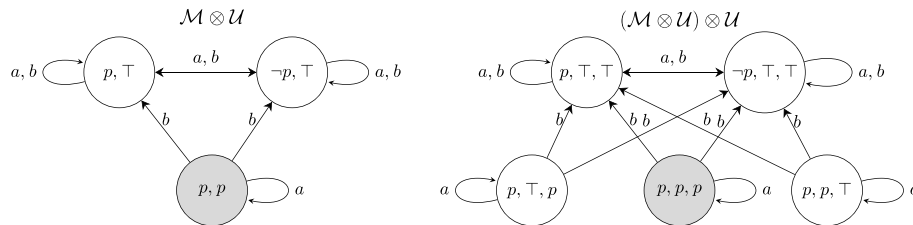


Fig. 3 Left: First product update $M \otimes \mathcal{U}$ of Kripke model M with action model $\mathcal{U}$. Right: Second product update $(M \otimes \mathcal{U}) \otimes \mathcal{U}$ with the same action model $\mathcal{U}$. p is true in all worlds that start from p and false in all worlds that start from $\neg p$.

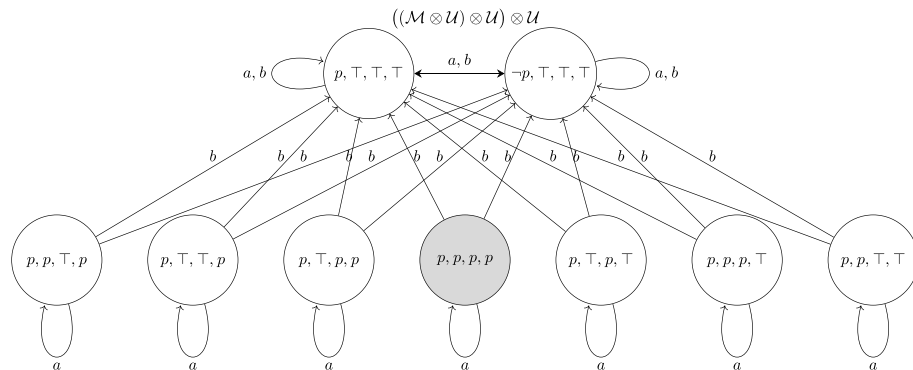


Fig. 4 Third product update $((M \otimes \mathcal{U}) \otimes \mathcal{U}) \otimes \mathcal{U}$ with the same action model $\mathcal{U}$.

Theorem 14: (Updates preserve bisimilarity) Given bisimilar pointed Kripke models $\mathcal{M}, w \leftrightarrow \mathcal{M}', w'$ and bisimilar pointed action models $\mathcal{U}, \alpha \leftrightarrow \mathcal{U}', \alpha'$ such that $\mathcal{M}, w \models \text{pre}(\alpha)$, we have $\mathcal{M} \otimes \mathcal{U}, (w, \alpha) \leftrightarrow \mathcal{M}' \otimes \mathcal{U}', (w', \alpha')$ and $\mathcal{M} \otimes \mathcal{U}, (w, \alpha) \leftrightarrow \mathcal{M}' \otimes \mathcal{U}', (w', \alpha')$. The same holds for G -bisimilarity.

Proof: For product updates, the proof can be found in van Ditmarsch et al.^[4]. For pointed updates, it then follows from Theorem 13. The statements for G -bisimilarity easily follow. $\square$

The advantage of pointed updates can be illustrated by the fact that, while repeated product updates of $\mathcal{M}$ with $\mathcal{U}$ from Fig. 2 lead to the exponential growth of the domain, pointed updates yield the three-world pointed Kripke model depicted left in Fig. 3, no matter how many times the pointed update is performed:

$$((\mathcal{M} \otimes \mathcal{U}) \otimes \mathcal{U}) \odot \dots \odot \mathcal{U} = \mathcal{M} \otimes \mathcal{U}.$$

While bisimulation contraction^[19] can shrink the model obtained by the product update operation, it needs to be invoked on the resulting model in a second step. In other words, first the model needs to be computed (with a potential exponential blowup) and then the bisimulation contraction runs on that model, making the operation computationally costly. On the other hand, the pointed update operation that we introduce runs only once and it produces a model which is no bigger than the one that would be obtained by the product update operation.

In the following section, we use this mechanism to solve the consistent update synthesis problem.

3 Consistent update synthesis: deterministic belief increase

In this section, we propose a solution to the consistent update synthesis task for goal formulas φ representing *deterministic belief increase* (DBI) by generating a pointed action model $\mathcal{U}_\varphi$ such that the result of a pointed update of any given pointed Kripke model with $\mathcal{U}_\varphi$ satisfies φ and inconsistent beliefs (including higher-order beliefs) are not introduced whenever they are possible to avoid. By deterministic belief increase we mean situations when several agents are prescribed additional beliefs (including higher-order beliefs) without creating alternative ways of fulfilling the prescription. Before giving the formal definition, it should be mentioned that there is always a trivial way of creating new beliefs by making the agents' beliefs inconsistent, i.e., making them believe all statements including the desired ones. However, making agents' reasoning inconsistent does not comport with the ideology of minimal change, nor is productive in terms of correcting agents' incidental false beliefs. The main novelty in our method of update synthesis is the aim to preserve consistency whenever possible. Now we give a formal definition of goal formulas representing deterministic belief increase:

Definition 15 The set of target agents of a modal formula φ is defined as follows: $\text{ta}(\varphi) := \emptyset$; $\text{ta}(\neg\varphi) := \text{ta}(\varphi)$; $\text{ta}(\varphi \wedge \psi) := \text{ta}(\varphi) \cup \text{ta}(\psi)$; finally $\text{ta}(B_i\varphi) := \{i\}$.

E.g., $\text{ta}(B_b(B_{i,p} \wedge B_{i,p} \wedge B_{i,p})) = \{b\}$ even though $\text{ta}(B_{i,p} \wedge B_{i,p} \wedge B_{i,p}) = \{i, i\}$.

Definition 16 (DBI goal formulas, DBI normal form) DBI goal formulas φ , or DBI formulas for short, are defined by the following BNF:

$$\varphi ::= B_i\xi \mid B_i(\xi \wedge \varphi) \mid (\varphi \wedge \varphi) \mid B_i\varphi \quad (2)$$

where ξ is any purely propositional formula and $i \in \mathcal{A}$. Thus, each DBI goal formula is a non-empty conjunction of belief operators. Formulas in DBI normal form are DBI formulas φ obtained by restricting the construction as follows: $B_i\xi$ is always DBI normal; $B_i\varphi$ and $B_i(\xi \wedge \varphi)$

are DBI normal iff φ is DBI normal and $i \notin \text{ta}(\varphi)$; $\varphi \wedge \psi$ is DBI normal iff φ and ψ are DBI normal and $\text{ta}(\varphi) \cap \text{ta}(\psi) = \emptyset$.

Lemma 17 For any DBI goal formula φ , there exists a DBI normal formula φ' such that $\mathcal{K}45 \models \varphi \equiv \varphi'$.

Proof We conduct the proof by structural induction over DBI goal formula φ .

Base case: As $\varphi = B_i\xi$ for propositional formula ξ , φ is already DBI normal.

Induction hypothesis: Given goal formula φ , for its proper sub formula ψ there exists a formula ψ' that is both DBI normal and equivalent to ψ .

Induction step:

- If $\varphi = B_i\psi$ or $\varphi = B_i(\xi \wedge \psi)$ for propositional formula ξ , it follows from the induction hypothesis that there exists a formula ψ' that is DBI normal and equivalent to ψ . Suppose by contradiction that $B_i\psi'$ (respectively $B_i(\xi \wedge \psi')$) is not DBI normal, thus $i \in \text{ta}(\psi')$. Hence by Definition 16 $\psi' = (B_i\psi'') \wedge \tau$, where $i \notin \text{ta}(\psi'')$, $i \notin \text{ta}(\tau)$ and both ψ'' and τ are DBI normal. Using the $\mathcal{K}45$ logical equivalences $B_iB_j\theta \equiv B_i\theta$, $B_i(\xi \wedge B_j\theta) \equiv B_i(\xi \wedge \theta)$ and $B_i\theta \wedge B_j\eta \equiv B_i(\theta \wedge \eta)$ we get the following equivalence chains:

$$B_i\psi \equiv B_i\psi' \equiv B_i(B_i\psi'' \wedge \tau) \equiv B_iB_i\psi'' \wedge B_i\tau \equiv B_i\psi'' \wedge B_i\tau \equiv B_i(\psi'' \wedge \tau)$$

$$B_i(\xi \wedge \psi) \equiv B_i(\xi \wedge \psi') \equiv B_i(\xi \wedge B_i\psi'' \wedge \tau) \equiv B_i(\xi \wedge B_i\psi'') \wedge B_i\tau$$

$$\equiv B_i(\xi \wedge \psi'') \wedge B_i\tau \equiv B_i(\xi \wedge \psi'' \wedge \tau).$$

$B_i(\psi'' \wedge \tau)$ respectively $B_i(\xi \wedge \psi'' \wedge \tau)$ are DBI normal.

- Else, if $\varphi = \varphi_1 \wedge \varphi_2$, using the induction hypothesis there exist the two DBI normal formulas φ'_1 and φ'_2 , where φ'_1 is equivalent to φ_1 and φ'_2 is equivalent to φ_2 . Suppose by contradiction that $\text{ta}(\varphi'_1) \cap \text{ta}(\varphi'_2) = G \neq \emptyset$, as otherwise $\varphi'_1 \wedge \varphi'_2$ would already be DBI normal (and equivalent to φ). Hence it must be that $\varphi'_1 = \bigwedge_{i \in G} B_i\psi_1^i \wedge \varphi_1''$ and $\varphi'_2 = \bigwedge_{i \in G} B_i\psi_2^i \wedge \varphi_2''$, where $\text{ta}(\varphi_1'') \cap \text{ta}(\varphi_2'') = \emptyset$ and φ_1'' , φ_2'' as well as all ψ_1^i and ψ_2^i are DBI normal. Thus using the equivalence $B_i\theta \wedge B_j\eta \equiv B_i(\theta \wedge \eta)$ and associativity and commutativity of $\wedge$ we get $\varphi'_1 \wedge \varphi'_2 \equiv \bigwedge_{i \in G} (B_i\psi_1^i \wedge \varphi_1'') \wedge \bigwedge_{i \in G} (B_i\psi_2^i \wedge \varphi_2'')$

$$\equiv \bigwedge_{i \in G} B_i(\psi_1^i \wedge \psi_2^i) \wedge \varphi_1'' \wedge \varphi_2''.$$

In addition, $\bigwedge_{i \in G} B_i(\psi_1^i \wedge \psi_2^i) \wedge \varphi_1'' \wedge \varphi_2''$ is DBI normal. $\square$

Lemma 18 (G -bisimilarity and goal formula preservation) If $\mathcal{M}, v \leftrightarrow_G \mathcal{M}', v'$ and φ is a DBI formula with $\text{ta}(\varphi) \subseteq G$, then $\mathcal{M}, v \models \varphi$ iff $\mathcal{M}', v' \models \varphi$.

Proof It follows from Theorem 12 and the fact that φ is a conjunction of $B_i\psi_i$ for $i \in \text{ta}(\varphi) \subseteq G$. $\square$

Without loss of generality, from now on, we only consider formulas in DBI normal form. For instance, formula (1) is a DBI formula but not DBI normal. Its normal form would be $B_b(B_{i,p} \wedge B_{i,p} \wedge B_{i,p})$. Given such a DBI normal formula φ , we now construct a pointed action model $(\mathcal{U}_\varphi, 0)$ such that for any pointed Kripke model $(\mathcal{M}, w)$, the pointed update $\mathcal{M} \odot \mathcal{U}_\varphi$ is defined, $\mathcal{M} \odot \mathcal{U}_\varphi, (w, 0) \models \varphi$, while other epistemic differences between $(\mathcal{M}, w)$ and $(\mathcal{M} \odot \mathcal{U}_\varphi, (w, 0))$ are minimized. Informally, this means that an update should not influence any beliefs except for those explicitly stated in φ and logically following from φ based on the agents' pre-update beliefs (we adopt the minimality principle in the same spirit of the standard AGM approach^[20] for which an update should lead to the loss of as few previous beliefs as possible^[21]). In particular, if $j \notin \text{ta}(\varphi)$, then j 's beliefs should not be affected even if B_j occurs in φ : updating somebody else's beliefs about j 's beliefs should not affect actual j 's beliefs.

Algorithms 1, 2, 3, and 4 show the construction of an action model for a consistent DBI update.

We illustrate this update synthesis method by applying it to Example 8:

Algorithm 1. SynthesizeActionModel(φ).

```

// initially add the model's point
1:  $E^\varphi := \{0\}$ 
2: for all  $i \in \mathcal{A}$  do
3:    $Q_i^\varphi := \emptyset$ 
4:  $pre(0) := \top$ 
5: CreateActionModel( $\varphi, \langle E^\varphi, R^\varphi, pre^\varphi \rangle, 0$ )
6: AddSink( $\langle E^\varphi, Q^\varphi, pre^\varphi \rangle$ )
7: return ( $\langle E^\varphi, Q^\varphi, pre^\varphi \rangle, 0$ )

```

Algorithm 2. CreateActionModel($\varphi, \langle E^\varphi, Q^\varphi, pre^\varphi \rangle, u$).

```

//  $\langle E^\varphi, Q^\varphi, pre^\varphi \rangle$  are passed by reference
1: if  $\varphi = B_i \xi$  for propositional formula  $\xi$  then
2:    $u' := \text{ComputeAccState}(\langle E^\varphi, Q^\varphi, pre^\varphi \rangle, u, i)$ 
3:    $pre^\varphi(u') := \xi$ 
4: else if  $\varphi = B_i \varphi'$  then
5:    $u' := \text{ComputeAccState}(\langle E^\varphi, Q^\varphi, pre^\varphi \rangle, u, i)$ 
6:   CreateActionModel( $\varphi', \langle E^\varphi, Q^\varphi, pre^\varphi \rangle, u'$ )
7: else if  $\varphi = B_i(\xi \wedge \varphi')$  then
8:    $u' := \text{ComputeAccState}(\langle E^\varphi, Q^\varphi, pre^\varphi \rangle, u, i)$ 
9:   CreateActionModel( $\varphi', \langle E^\varphi, Q^\varphi, pre^\varphi \rangle, u'$ )
10:   $pre^\varphi(u') := \xi$ 
11: else if  $\varphi = \varphi_1 \wedge \varphi_2$  then
12:   CreateActionModel( $\varphi_1, \langle E^\varphi, Q^\varphi, pre^\varphi \rangle, u$ )
13:   CreateActionModel( $\varphi_2, \langle E^\varphi, Q^\varphi, pre^\varphi \rangle, u$ )

```

Algorithm 3. ComputeAccState($\langle E^\varphi, Q^\varphi, pre^\varphi \rangle, u, i$).

```

//  $\langle E^\varphi, Q^\varphi, pre^\varphi \rangle$  are passed by reference
1: create new state  $u'$  (s.t.  $u' \neq 0$  and  $u' \neq -1$ ) with  $pre^\varphi(u') = \top$ 
2:  $E^\varphi := E^\varphi \cup \{u'\}$ 
3:  $Q_i^\varphi := Q_i^\varphi \cup \{(u, u'), (u', u')\}$ 
4: return  $u'$ 

```

Algorithm 4. AddSink($\langle E^\varphi, Q^\varphi, pre^\varphi \rangle$).

```

1: create state  $-1$ , where  $pre^\varphi(-1) = \top$ 
2:  $E^\varphi := E^\varphi \cup \{-1\}$ 
3: for all  $i \in \mathcal{A}$  do
4:    $Q_i^\varphi := Q_i^\varphi \cup \{(-1, -1)\}$ 
5: for all  $s \in E^\varphi$  do
6:   for all  $i \in \mathcal{A}$  do
7:     if  $(\forall s' \in E^\varphi) (s, s') \notin Q_i^\varphi$  then
8:        $Q_i^\varphi := Q_i^\varphi \cup \{(s, -1)\}$ 

```

Example 19 For the DBI normal form $\varphi = B_b(B_i p \wedge B_l(p \wedge B_i p))$ of formula (1) from Example 8, pointed action model $(\mathcal{U}_\varphi, 0)$ constructed according to Algorithm 1 can be found in the left part of Fig. 5. The result of the pointed update of $(\mathcal{M}, v)$ from the left part of Fig. 1 with this $(\mathcal{U}_\varphi, 0)$ can be seen on the right of Fig. 5 (and is isomorphic to the right part of Fig. 1). It is easy to see that $\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0) \models \varphi$. Since all accessibility relations in this updated model $\mathcal{M} \odot \mathcal{U}_\varphi$ are serial, it follows that it is common belief that all agents have consistent beliefs.

In addition, we claim that this update synthesis has been achieved with minimal change to agents' beliefs. Indeed, it is easy to prove using G-bisimilarities for appropriate groups of agents that

$$\begin{array}{ll} \mathcal{M}, v \equiv_{l,t} \mathcal{M} \odot \mathcal{U}_\varphi, (v, 0) & \mathcal{M}, v \equiv_{b,l} \mathcal{M} \odot \mathcal{U}_\varphi, (u, 2) \\ \mathcal{M}, v \equiv_b \mathcal{M} \odot \mathcal{U}_\varphi, (u, 3) & \mathcal{M}, v \equiv_{b,l} \mathcal{M} \odot \mathcal{U}_\varphi, (u, 4) \end{array}$$

where v is the actual world and u is the other world of $\mathcal{M}$ (we omit set braces in the subscript of $\equiv$). The statements in the first column mean that the update is imperceptible for agents l and t and that agent b thinks that agent l does not think b 's beliefs have changed. Similarly, the second column testifies that b thinks that neither t 's beliefs about b and l nor l 's beliefs about what t thinks regarding b or l changed. In other words, the only higher-order beliefs that are affected by the update are those explicitly dictated by φ .

We illustrate the advantage of using pointed updates over the product update operation using Example 8:

Example 20 Consider the same example but using the product update of standard DEL: The result of updating the initial model with the action model obtained from φ is shown in Fig. 6. In particular, the result of the product update has only one additional world with respect to the model obtained via pointed update $\mathcal{M} \odot \mathcal{U}_\varphi$ visible in Fig. 5 (left), which is not reachable from the actual world, namely the p world in the v cluster at the top. However, the difference between the two operations is striking in case of iterated updates: applying to the resulting model the product update operation again with the same action model leads to the Kripke model sketched in Fig. 7, which contains a number of worlds that are not reachable from the actual world but copied in every sub-model nonetheless. Compare this to the pointed update operation, whose result is the same for any number of iterations for the same action model (Theorem 33).

Let us now prove that $\mathcal{U}_\varphi$ performs update synthesis for any DBI normal φ , minimizes belief change of other agents, and preserves consistency whenever possible.

We will use the auxiliary fact that updates do not affect propositional formulas:

Lemma 21 (Propositional invariance) Let (w, α) belong to the domain of $\mathcal{M} \odot \mathcal{U}$ for some Kripke model $\mathcal{M}$ and action model $\mathcal{U}$. Then for any propositional formula ξ ,

$$\mathcal{M}, w \models \xi \iff \mathcal{M} \odot \mathcal{U}, (w, \alpha) \models \xi. \quad (3)$$

Proof For product updates $\otimes$, this is well known (according to van Ditmarsch et al.^[4]). For pointed updates, it follows from Theorem 13. $\square$

Definition 22 (φ -enforcing) We call a pointed action model $\mathcal{U} = \langle E, Q, pre \rangle$ with point $\alpha \in E$ and formula φ in DBI normal form, φ -enforcing iff

- if $\varphi = B_i \xi$, for propositional formula ξ , then for all states $\beta \in E$ s.t. $\alpha Q_i \beta$, $\models pre(\beta) \rightarrow \xi$.
- else if $\varphi = B_i \psi$, then there exists at least one state $\beta' \in E$ s.t. $\alpha Q_i \beta'$ and for all states $\beta \in E$ s.t. $\alpha Q_i \beta$, $(\mathcal{U}, \beta)$ is ψ -enforcing.
- else if $\varphi = B_i(\xi \wedge \psi)$, for propositional formula ξ , then there exists at least one state $\beta' \in E$ s.t. $\alpha Q_i \beta'$ and for all states $\beta \in E$ s.t. $\alpha Q_i \beta$, $(\mathcal{U}, \beta)$ is ψ -enforcing and $\models pre(\beta) \rightarrow \xi$.
- else if $\varphi = \psi_1 \wedge \psi_2$, then $(\mathcal{U}, \alpha)$ is ψ_1 - and ψ_2 -enforcing.

Lemma 23 (Update with φ -enforcing action model) For pointed Kripke model $(\mathcal{M}, v)$, DBI normal formula φ , and φ -enforcing pointed action model $(\mathcal{U}, \alpha)$, if the update $\mathcal{M} \odot \mathcal{U} = \langle S', R', V' \rangle$ is defined (if $\mathcal{M}, v \models pre(\alpha)$), then

$$\mathcal{M} \odot \mathcal{U}, (v, \alpha) \models \varphi \quad (4)$$

Proof We conduct the proof by induction over the recursive structure of φ .

• In the base case, when $\varphi = B_i \xi$, since $(\mathcal{U}, \alpha)$ is φ -enforcing, for all states (w, β) s.t. $(v, \alpha) R'_i(w, \beta)$, $\mathcal{U} \odot \mathcal{M}, (w, \beta) \models \xi$. Hence $\mathcal{U} \odot \mathcal{M}, (v, \alpha) \models B_i \xi$ by semantics of B_i .

• If $\varphi = B_i \psi$, since for all i -accessible states β from α , $(\mathcal{U}, \beta)$ is ψ -enforcing using the IH we get that for all i -accessible states (w, β)

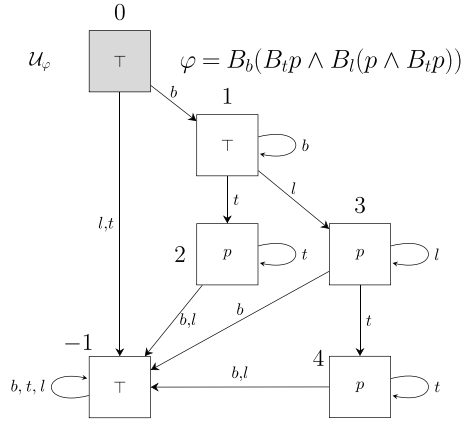


Fig. 5 Successful update synthesis $\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0) \models \varphi$ by applying pointed update with pointed action model $(\mathcal{U}_\varphi, 0)$ for DBI formula (1) from Example 8 to $(\mathcal{M}, v)$ from Fig. 1 (left).

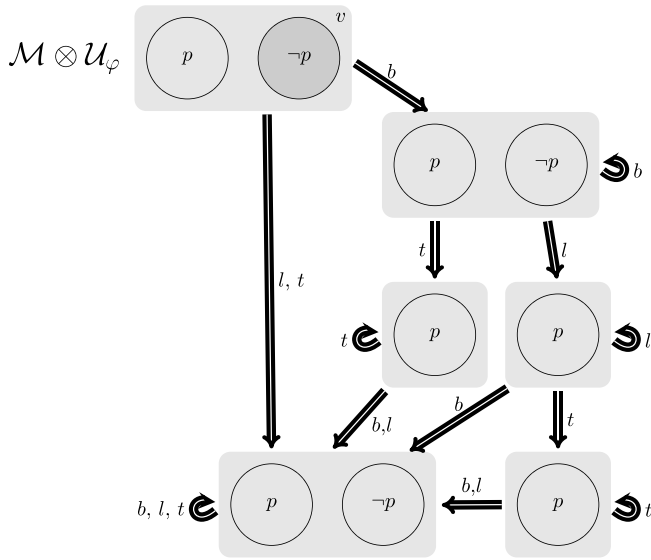


Fig. 6 Balder-Loki-Thor updated model using product updates $\mathcal{M} \otimes \mathcal{U}_\varphi$.

from (v, α) , $\mathcal{M} \odot \mathcal{U}, (w, \beta) \models \psi$. Hence $\mathcal{U} \odot \mathcal{M}, (v, \alpha) \models B_i \psi$ by semantics of B_i .

- If $\varphi = B_i(\xi \wedge \varphi')$ the reasoning is simply a combination of the previous two cases.

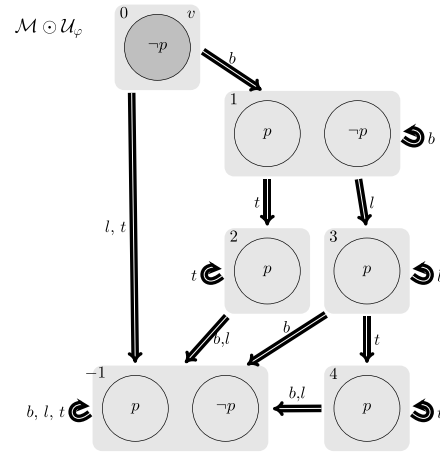
- If $\varphi = \psi_1 \wedge \psi_2$, as $(\mathcal{U}, \alpha)$ is ψ_1 - and ψ_2 -enforcing using the IH we get that $\mathcal{U} \odot \mathcal{M}, (v, \alpha) \models \psi_1$ and $\mathcal{U} \odot \mathcal{M}, (v, \alpha) \models \psi_2$ from which $\mathcal{U} \odot \mathcal{M}, (v, \alpha) \models \psi_1 \wedge \psi_2$ follows. $\square$

Lemma 24 ($\mathcal{U}_\varphi$ is φ -enforcing) For DBI normal formula φ , the pointed action model $(\mathcal{U}_\varphi, 0)$ synthesized by Algorithm 1 is φ -enforcing and $\text{pre}^\varphi(0) = \top$.

Proof We start the proof by induction over the structure of φ and the action model $\mathcal{U}_\varphi$ recursively constructed by Algorithm 2. We prove the statement for the resulting model $\mathcal{U}_\varphi$ after the function call $\text{CreateActionModel}(\varphi, \mathcal{U}, u)$, where $\mathcal{U} = \{0, \emptyset, \text{pre}\}$, where $\text{pre}(0) = \top$ following from code lines 1.–4. of Algorithm 1.

- In the base case where $\varphi = B_i \xi$, it trivially follows from Algorithm 2 code lines 1.–3. that $\mathcal{U}_\varphi$ is φ -enforcing, as a single new i -accessible state u' from u is constructed with $\text{pre}^\varphi(u') = \xi$.

- If $\varphi = B_i \varphi'$ from Algorithm 2 code lines 4.–6. in particular using the IH for line 6. it follows that $(\mathcal{U}_\varphi, u')$ is φ' -enforcing. Since φ is DBI normal and therefore $i \notin \text{ta}(\varphi')$ and the fact that recursive calls in



Algorithm 2 only ever occur for newly created states, u' is the only state i -accessible from u . Hence $\mathcal{U}_\varphi$ is φ -enforcing.

- If $\varphi = B_i(\xi \wedge \varphi')$ by Algorithm 2 code lines 7.–10. the reasoning is a combination of the previous two cases.

- If $\varphi = \varphi_1 \wedge \varphi_2$ by Algorithm 2 code lines 11.–13. in particular using the IH, for the recursive call on line 12. we get that the intermediate model $\mathcal{U}'_\varphi$ after the execution of code line 12. is φ_1 -enforcing. Since $\text{ta}(\varphi_1) \cap \text{ta}(\varphi_2) = \emptyset$ by DBI normality of φ and the fact that recursive function calls are only made for newly created states, we can be certain that the recursive call in line 13. operates on a separate branch in the model than the call in line 12 did. Thus we conclude that $\mathcal{U}'_\varphi$ is not only φ_2 -, but continues to be φ_1 -enforcing as well.

By the simple observation that no recursive calls of Algorithm 2 are made for existing states, therefore also not for state 0, it remains by Algorithm 1 code line 4 that $\text{pre}^\varphi(0) = \top$, as Algorithm 4 does not modify the precondition of prior existing states. $\square$

Theorem 25: (Update synthesis success) For any DBI normal formula φ and any pointed Kripke model $(\mathcal{M}, v)$, the pointed update of $(\mathcal{M}, v)$ with $(\mathcal{U}_\varphi, 0)$ is defined and

$$\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0) \models \varphi. \quad (5)$$

Proof Follows from Lemma 23 and Lemma 24. $\square$

The sequence of modal operators that one encounters, when traversing the syntax tree of formula φ (including the empty list ε) starting from its root, we call *modal operator sequence* or simply *modality sequence*. For example, given $\varphi = B_i(B_j p \wedge B_r q)$, the set of modal operator sequences of φ , would be $\{\varepsilon, B_i, B_i B_j, B_i B_r\}$.

To formulate statements about minimal change, we introduce the notion of independent formulas.

Definition 26 (Independent formulas) Let φ be a DBI normal formula and its corresponding action model $\mathcal{U}_\varphi = \langle E^\varphi, Q^\varphi, \text{pre}^\varphi \rangle$ be constructed according to Algorithm 1. A modal formula θ is in $\top$ -shape with respect to event $\alpha_0 \in E^\varphi$ iff for each sequence of modal operators $B_{i_1} \dots B_{i_k}$ used in the construction of θ , including the empty sequence ε with $k=0$, and for the unique corresponding sequence $\alpha_0 Q_{i_1}^\varphi \alpha_1 Q_{i_2}^\varphi \alpha_2 \dots \alpha_{k-1} Q_{i_k}^\varphi \alpha_k$ of events $\alpha_j \in E^\varphi$ all $\text{pre}^\varphi(\alpha_j) = \top$ for $0 \leq j \leq k$ (note that the uniqueness of such a sequence formally follows from Lemma 48 and Theorem 49). Formula θ is called independent of φ iff it is in $\top$ -shape w.r.t. $0 \in E^\varphi$. If θ is not independent of φ , it is dependent on φ .

Example 27 For formula $\varphi = B_b(B_i p \wedge B_l(p \wedge B_t p))$ from Example 8, we have that $\theta = \neg B_b B_b p \vee B_t B_t p$ is independent from φ because the modal operator sequences $\varepsilon, B_b, B_b B_b, B_t$, and $B_t B_t$ correspond

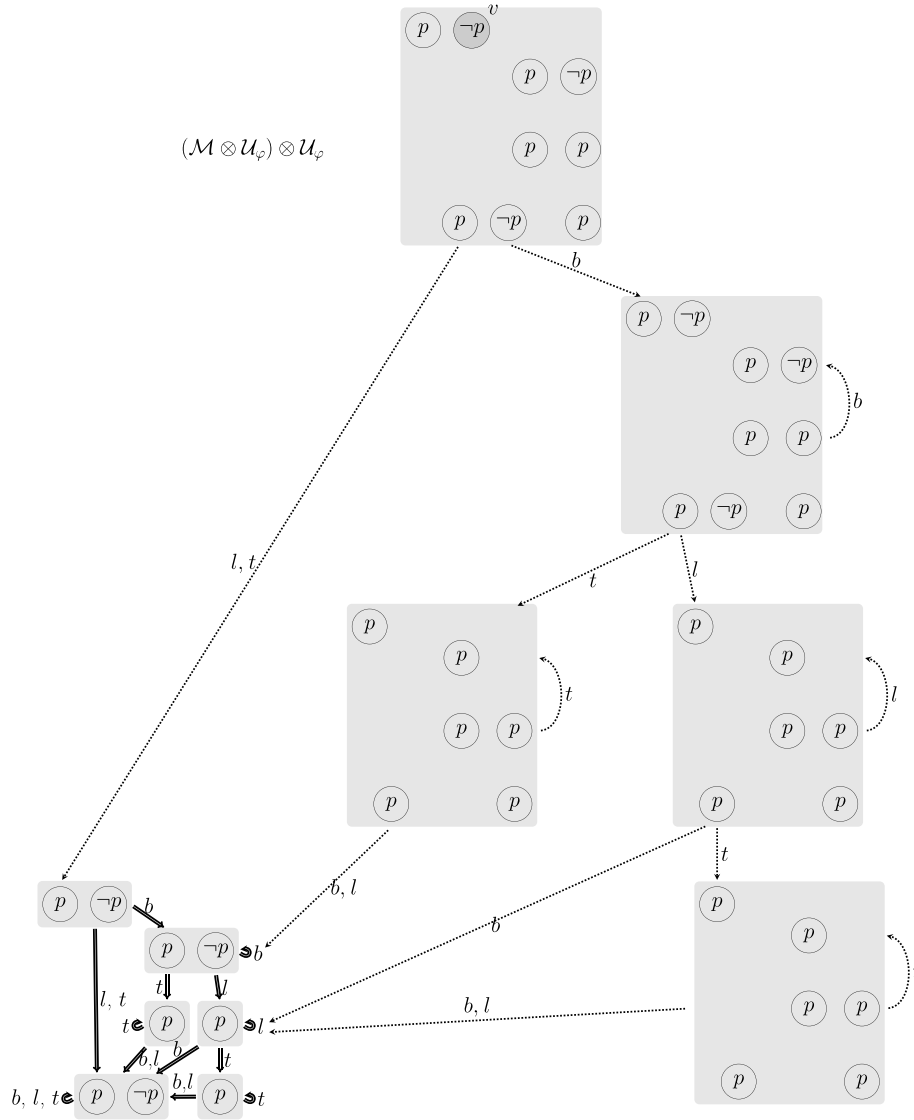


Fig. 7 Sketch of the result of the iterated product update operation $(\mathcal{M} \otimes \mathcal{U}_\varphi) \otimes \mathcal{U}_\varphi$. Grey rectangles are a compact way to represent sub-models originating from worlds in the previous model. Dotted arrows are meant to capture the direction of the arrows and the agents involved. Unlike double arrows, they do not represent arrows from and to all worlds in the rectangle, but only from and to some of these worlds. We leave to the reader their identification. Finally, note that the sink is a copy of the previous model.

to event sequences $0, 0Q_b^\varphi 4, 0Q_b^\varphi 4Q_b^\varphi 4, 0Q_t^\varphi -1$, and $0Q_t^\varphi -1Q_t^\varphi -1$ and preconditions for events $0, 4$, and -1 are all $\top$. On the other hand, $\eta = \neg B_b B_t p$ is dependent on φ because modality sequence $B_b B_t$ corresponds to event sequence $0Q_b^\varphi 4Q_t^\varphi 3$, and $\text{pre}^\varphi(3) = p \neq \top$. Both θ and η were true in $\mathcal{M}, v$ from Fig. 1, left. While the pointed update with $(\mathcal{U}_\varphi, 0)$ keeps θ true, formula η becomes false in $\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0)$ (see Fig. 5, right).

In other words, given a formula φ in DBI normal form, a formula θ is independent of φ if the sequences of modal operators in the formula θ have a corresponding sequence of events in $\mathcal{U}_\varphi$ and the precondition of those events is always $\top$. A formula is then dependent in case at least one of the preconditions for those events is different from $\top$. The main idea is that the pointed update $\mathcal{U}_\varphi$ will keep the formulas that are independent of φ true after the update.

Theorem 28 (Update synthesis minimality) *If formula θ is independent of a DBI normal formula φ , then for any pointed Kripke model $(\mathcal{M}, v)$,*

$$\mathcal{M}, v \models \theta \iff \mathcal{M} \odot \mathcal{U}_\varphi, (v, 0) \models \theta. \quad (6)$$

Proof Let $\mathcal{M} = \langle S, R, V \rangle$. We prove by induction on the construction of θ that, for any $\alpha \in E^\varphi$ and any $u \in S$, if θ is in $\top$ -shape with respect to α , then $\mathcal{M}, u \models \theta$ iff $\mathcal{M} \odot \mathcal{U}_\varphi, (u, \alpha) \models \theta$. (6) is an instance of this induction statement for $\alpha = 0$ and $u = v$. Note that if θ is in $\top$ -shape with respect to α , then $\text{pre}^\varphi(\alpha) = \top$ because of the empty sequence ε , hence, the pointed update of $(\mathcal{M}, u)$ with $(\mathcal{U}_\varphi, \alpha)$ is defined for all $u \in S$. For propositional atoms, the statement follows from the definition of pointed updates. The cases for Boolean connectives are straightforward. It remains to show the induction statement for $\theta = B_i \eta$. Let $\beta \in E^\varphi$ be the unique event such that $\alpha Q_i^\varphi \beta$. Then $\text{pre}^\varphi(\beta) = \top$ because of sequence B_i of θ and, additionally, η is in $\top$ -shape with respect to β . By IH, for any $u \in S$, we have $\mathcal{M}, u \models \eta$ iff $\mathcal{M} \odot \mathcal{U}_\varphi, (u, \beta) \models \eta$. We have $\mathcal{M}, u \models B_i \eta$ iff $\mathcal{M}, w \models \eta$ for some $u R_i w$. By IH, this is equivalent to $\mathcal{M} \odot \mathcal{U}_\varphi, (w, \beta) \models \eta$, which is equivalent to $\mathcal{M} \odot \mathcal{U}_\varphi, (u, \alpha) \models B_i \eta$ because $(u, \alpha) R_i^\varphi (w, \gamma)$ iff $u R_i w$ and $\gamma = \beta$. $\square$

Corollary 29 *Let $\varphi = \bigwedge_{j \in G} B_j \omega_j$ be a DBI normal formula, where $\emptyset \neq G \subseteq \mathcal{A}$.*

1. If $i \notin \text{ta}(\varphi)$, i.e., if $i \notin G$, then i 's beliefs are unaffected by pointed update $\mathcal{U}_\varphi$, i.e., $\mathcal{M}, v \models B_i \sigma$ iff $\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0) \models B_i \sigma$ for any formula σ and any pointed Kripke model $(\mathcal{M}, v)$.

2. If $i \in \text{ta}(\varphi)$, but $\omega_i = \bigwedge_{j \in H} B_j \pi_j$ has no propositional component, then i 's propositional beliefs are unaffected by $\mathcal{U}_\varphi$, i.e., $\mathcal{M}, v \models B_i \chi$ iff $\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0) \models B_i \chi$ for any propositional formula χ and any pointed Kripke model $(\mathcal{M}, v)$.

Theorem 30 (Update synthesis consistency preservation) Let φ be a DBI normal formula. Pointed update with $\mathcal{U}_\varphi$ can only cause (higher-order) inconsistent beliefs iff the respective (higher-order) beliefs originally excluded the respective propositional preconditions: for any modality sequence $B_{i_1} \dots B_{i_k}$ with $i_j \neq i_{j+1}$ for any j and any pointed Kripke model $(\mathcal{M}, v)$,

$$\mathcal{M}, v \models \hat{B}_{i_1}(\text{pre}^\varphi(\alpha_1) \wedge \hat{B}_{i_2}(\text{pre}^\varphi(\alpha_2) \wedge \dots \hat{B}_{i_k} \text{pre}^\varphi(\alpha_k))) \iff \mathcal{M} \odot \mathcal{U}_\varphi, (v, 0) \not\models B_{i_1} \dots B_{i_k} \perp \quad (7)$$

for the unique sequence $0Q_{i_1}^\varphi \alpha_1 Q_{i_2}^\varphi \alpha_2 \dots Q_{i_k}^\varphi \alpha_k$ of events $\alpha_j \in E^\varphi$.

Proof Let $\mathcal{M} = \langle S, R, V \rangle$ and $\mathcal{M} \odot \mathcal{U}_\varphi = \langle S^\varphi, R^\varphi, V^\varphi \rangle$ for the pointed update of $(\mathcal{M}, v)$ with $(\mathcal{U}_\varphi, 0)$. The left statement holds iff there is a sequence $vR_{i_1} s_1 R_{i_2} s_2 \dots R_{i_k} s_k$ of states from S such that $\mathcal{M}, s_j \models \text{pre}^\varphi(\alpha_j)$ for $j = 1, \dots, k$. It is easy to observe (by induction on k) that this is equivalent to $(s_j, \alpha_j) \in S^\varphi$ for $j = 1, \dots, k$ and $(v, 0)R_{i_1}^\varphi(s_1, \alpha_1)R_{i_2}^\varphi(s_2, \alpha_2) \dots R_{i_k}^\varphi(s_k, \alpha_k)$. The equivalence to the right statement now follows from the fact that $\mathcal{M} \odot \mathcal{U}_\varphi, (s_k, \alpha_k) \not\models \perp$ and the uniqueness of the sequence of α_j 's such that $0Q_{i_1}^\varphi \alpha_1 Q_{i_2}^\varphi \alpha_2 \dots Q_{i_k}^\varphi \alpha_k$. $\square$

Corollary 31 Let φ be a DBI normal formula.

1. If φ fits either of the clauses of Cor. 29 for agent i , then $\mathcal{U}_\varphi$ preserves i 's consistency, i.e., $\mathcal{M}, v \not\models B_i \perp$ iff $\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0) \not\models B_i \perp$ for any $(\mathcal{M}, v)$.

2. If $\varphi = B_i(\xi \wedge \bigwedge_{k \in H} B_k \pi_k) \wedge \bigwedge_{j \in G} B_j \omega_j$ where $i \notin G$ and ξ is propositional, then $\mathcal{U}_\varphi$ makes i 's beliefs inconsistent if and only if i originally does not consider ξ to be possible, i.e., $\mathcal{M}, v \models \hat{B}_i \xi$ iff $\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0) \not\models B_i \perp$ for any $(\mathcal{M}, v)$.

Definition 32 (Idempotence) We call an action model $(\mathcal{U}, \alpha)$ idempotent iff pointed updates $(\mathcal{M} \odot \mathcal{U}, (w, \alpha))$ and $(\mathcal{M} \odot \mathcal{U}) \odot \mathcal{U}, ((w, \alpha), \alpha)$ are defined and isomorphic for any pointed Kripke model $(\mathcal{M}, w)$.

Theorem 33 For any DBI normal formula φ and any pointed Kripke model $(\mathcal{M}, v)$, $\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0)$ is idempotent.

Proof The updates $\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0)$ and $(\mathcal{M} \odot \mathcal{U}_\varphi) \odot \mathcal{U}_\varphi, ((v, 0), 0)$ are always defined by Theorem 25.

Since $\mathcal{U}_\varphi$ has an out-tree structure, meaning every state is reachable by a directed walk from the point, we conduct the proof by forward induction over this structure beginning from the root 0.

Ind. statement.: Given state $((w, \alpha), \beta)$ in the model $(\mathcal{M} \odot \mathcal{U}_\varphi) \odot \mathcal{U}_\varphi, \beta = \alpha$.

Ind. hyp.: Given state $((w, \alpha), \beta)$ for all parent states $((w', \alpha'), \beta')$ the inductive statement holds.

Base case: Point $(v, 0)$ is the only state that can be combined with 0, since by Algorithm 1 0 has no incoming edges in $\mathcal{U}_\varphi$.

Ind. step: Suppose by IH the inductive statement holds for state $((w, \alpha), \alpha)$, but not for its i -accessible state $((x, \beta), \gamma)$, meaning $\gamma \neq \beta$. Note that this i -accessible state must exist, since $(\mathcal{M} \odot \mathcal{U}_\varphi) \odot \mathcal{U}_\varphi, ((v, 0), 0)$ could not be smaller than $\mathcal{M} \odot \mathcal{U}_\varphi, (v, 0)$, as all preconditions in $\mathcal{U}_\varphi$ are propositional by Algorithm 2. Since $((w, \alpha), \alpha)R_i^{\mathcal{U}_\varphi}((x, \beta), \gamma)$ it must be the case that also $(w, \alpha)R_i^{\mathcal{U}_\varphi}(x, \beta)$ by Definition 9. Since φ is DBI normal, every state in $\mathcal{U}_\varphi$ has no more than one j -accessible state, hence β is the only i -accessible state from α . Therefore if $((w, \alpha), \alpha)R_i^{\mathcal{U}_\varphi}((x, \beta), \gamma)$ it must be that $\gamma = \beta$. $\square$

4 The role of privatization

The minimality of the update synthesis method we have described relies on a particular structure of the action model used for the update. We call this structure *privatized*. In this section, we give an explicit definition and show why it serves the purpose of preserving beliefs whenever possible. We begin with the notion of modal syntactic tree, which represents the nesting of modalities within a formula as a tree and provides the basis of the formal definition of privatization. Note that the structure of the action model returned by Algorithm 1 (almost) exactly follows this tree structure already, so there was no need to introduce it explicitly.

Definition 34 (Modal syntactic tree) The modal syntactic tree $\mathcal{T}_\varphi$ of a goal formula φ is an out-tree with a single unlabeled root, while all non-root nodes are labeled with modal operators representing the nesting of modalities in φ . Hence it is essentially φ 's syntax tree, where however all non-modal operators are left out.

Some examples of modal syntactic trees are provided in Fig. 8: In $\mathcal{T}_{B_i \xi}$, the root has one child-leaf labeled B_i . Both $\mathcal{T}_{B_i \varphi}$ and $\mathcal{T}_{B_i(\xi \wedge \varphi)}$ are obtained by labeling the root of $\mathcal{T}_\varphi$ with B_i and making it the only child of the new root. Finally, $\mathcal{T}_{\varphi \wedge \psi}$ is obtained by taking the disjoint union of $\mathcal{T}_\varphi$ and $\mathcal{T}_\psi$ and identifying their roots.

Definition 35 (Modal-syntactic-tree root paths) For formula φ we define $\text{RootP}(\varphi)$ as the set of paths $((\text{root}, \alpha_1, i_1), \dots, (\alpha_{l-1}, \alpha_l, i_l))$ of length $l \geq 0$ in $\mathcal{T}_\varphi$ starting from the root, where i_k is the label of α_k for $k = 1, \dots, l$.

Definition 36 (Kripke-frame root walks) For a pointed Kripke frame $(\langle S, R \rangle, w)$ we similarly define the set of root walks $\text{RootW}(\langle S, R \rangle, w)$ as the set of all walks starting from the root (point) w . We further define the restriction $\text{RootW}_{\text{nsr}}(\langle S, R \rangle, w)$ as the subset of $\text{RootW}(\langle S, R \rangle, w)$, where we exclude walks that contain at least two successive edges for the same agent.

Definition 37 (Agent sequence) For a walk $\sigma = ((\alpha_1, \alpha_2, i_1), \dots, (\alpha_l, \alpha_{l+1}, i_l))$ from Defs. 35 or 36 we define $\text{AgSeq}(\sigma) := (i_1, \dots, i_l)$. In particular, $\text{AgSeq}(\varepsilon) := \varepsilon$. We extend this definition to sets, where for a set of walks Σ , $\text{AgSeq}(\Sigma)$ is the set of corresponding agent sequences.

The set of all agent sequences of length l we denote by $\mathcal{A}^l$. The set of all agent sequences of length l without any successively repeating agents we denote by $\mathcal{A}_{\text{nsr}}^l$.

We now define clusters and walk accessible parts of the model.

Definition 38 (Clusters and reachable states) For a Kripke frame $\mathcal{F} = \langle W, R \rangle$, world $w \in W$, and sequence $(i_1, \dots, i_l) \in \mathcal{A}^l$ of agents, we introduce the cluster of path-accessible worlds

$$\mathcal{C}_{\mathcal{F}, w}^{i_1, i_2, \dots, i_l} := \{u \in W \mid (\exists u_2, \dots, u_l \in W) wR_{i_1} u_2 R_{i_2} \dots u_l R_{i_l} u\}. \quad (8)$$

In particular, for the empty sequence, $\mathcal{C}_{\mathcal{F}, w}^\varepsilon = \{w\}$. We also define the reachability operator

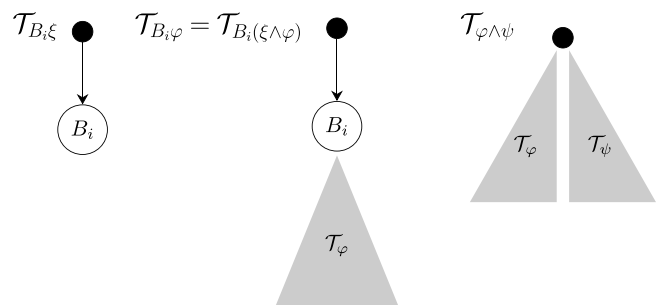


Fig. 8 Modal syntactic trees.

$$R(\mathcal{F}, w) := \bigcup_{l=0}^{\infty} \bigcup_{\rho \in \mathcal{A}^l} C_{\mathcal{F}, w}^{\rho} \quad (9)$$

For a set $W' \subseteq W$, we define $R(\mathcal{F}, W') := \bigcup_{w' \in W'} R(\mathcal{F}, w')$.

Definition 39 For action model $\mathcal{U} = \langle E, Q, pre \rangle$ and state $u \in E$, we define the u -sub model $\mathcal{U}^u = \langle E^u, Q^u, pre^u \rangle$ as the sub model of $\langle E, Q, pre \rangle$ that is reachable from u , where $E^u = R(\mathcal{U}, u)$, $Q^u = Q|_{(E^u \times E^u)}$ and $V^u = V|_{E^u}$.

Lemma 40 For DBI goal formula φ , action model $\langle E, Q, pre \rangle$, for which $u \in E$ is a leaf node (potentially with a reflexive loop), let $\langle E^{\varphi}, Q^{\varphi}, pre^{\varphi} \rangle$ be the action model synthesized by a call of [Algorithm 2](#) `CreateActionModel`($\varphi, \langle E, Q, pre \rangle, u$), then the size of its u -sub model $\langle E^{\varphi^u}, Q^{\varphi^u}, pre^{\varphi^u} \rangle$ is at worst linear in the size of the modal syntactic tree of φ .

Proof We conduct the proof by structural induction over the DBI goal formula φ .

Base case: $\varphi = B_i \xi$ for purely propositional formula ξ . The resulting u -sub action model consists only of the root u and a newly created node u' , with an at least an i -arrow going from u to u' and a reflexive loop for u' :

$$E^{\varphi} = \{u, u'\}, Q_i^{\varphi} \supseteq \{(u, u'), (u', u')\}.$$

Induction hypothesis: Given goal formula φ , for its proper sub formula φ' and state $u' \in E$, let $\langle E^{\varphi'}, Q^{\varphi'}, pre^{\varphi'} \rangle$ be the action model synthesized by a call of [Algorithm 2](#) `CreateActionModel`($\varphi', \langle E, Q, pre \rangle, u'$), then the size of the u' -sub model $\langle E^{\varphi'^{u'}}, Q^{\varphi'^{u'}}, pre^{\varphi'^{u'}} \rangle$ is at worst linear in the size of the modal syntactic tree of φ' .

Induction step:

- If $\varphi = B_i \varphi'$ or $\varphi = B_i(\xi \wedge \varphi')$ for purely propositional formula ξ . By [Algorithm 2](#) it follows that only a single i -accessible new state u' is added, after which the algorithm is called recursively (on code line 6. respectively 9.) for u' and φ' . Since after the recursive call, the u' -sub model $\langle E^{u'}, Q^{u'}, pre^{u'} \rangle$ is linear in the size of the modal syntactic tree of φ' , by assumption of the IH $\langle E^u, Q^u, pre^u \rangle$ is linear in the size of the modal syntactic tree of φ as well, as the difference between $\langle E^u, Q^u, pre^u \rangle$ and $\langle E^{u'}, Q^{u'}, pre^{u'} \rangle$ is only a single state.

- Else if $\varphi = \varphi_1 \wedge \varphi_2$. Using the IH for both recursive calls (on code lines 12. and 13.), it directly follows that $\langle E^u, Q^u, pre^u \rangle$ is linear in the size of both φ_1 's and φ_2 's, thus also φ 's modal syntactic tree. $\square$

Theorem 41 The action model $\langle E^{\varphi}, Q^{\varphi}, pre^{\varphi} \rangle$ synthesized by a call `SynthesizeActionModel`(φ) of [Algorithm 1](#) for DBI goal formula φ is at worst linear in the size of the modal syntactic tree of φ .

Proof Follows from [Algorithm 1](#), [Algorithm 4](#), and Lemma 40. $\square$

Definition 42 We define $\mathcal{Q}_{\text{nsr}}$ as the largest subset of $\mathcal{Q}$ s.t. in any $\varphi \in \mathcal{Q}_{\text{nsr}}$ no two subsequent nestings of modalities with the same agent occur, meaning formulas like $B_i B_i(p \wedge B_i q)$ are disallowed.

Note that this restriction of the language does not affect any of the results in Section 3 because of the following obvious proposition.

Proposition 43 Every DBI formula is in $\mathcal{Q}_{\text{nsr}}$.

Before we provide a formal definition of privatization we want to first give an informal definition and also clarify why we introduce two notions.

Given a formula $\varphi \in \mathcal{Q}_{\text{nsr}}$ a pointed Kripke frame is weakly privatized with respect to φ iff for every agent sequence of any root path of φ 's modal syntactic tree, all states are accessible via such an agent sequence from the point are not accessible via another agent sequence (without successively repeating agents). A Kripke frame is strongly privatized iff it is weakly privatized and for every agent sequence for any root path of φ 's modal syntactic tree there exists at least one state accessible via this agent sequence.

The reason why we make a distinction between weak and strong privatization is that when applying a strongly privatized action model

the resulting epistemic model need not necessarily be strongly but only weakly privatized, since whole clusters could get removed as a result of the update.

Definition 44 (Privatized) We call a pointed Kripke frame $(\mathcal{F}, w) = (\langle S, R \rangle, w)$ (strongly) privatized with respect to formula $\varphi \in \mathcal{Q}_{\text{nsr}}$ iff for any root path $\sigma \in \text{RootP}(\varphi)$

$$C_{\mathcal{F}, w}^{\text{AgSeq}(\sigma)} \neq \emptyset \text{ and } \left(\forall \rho \in \bigcup_{l=0}^{\infty} \mathcal{A}_{\text{nsr}}^l \setminus \{\text{AgSeq}(\sigma)\} \right) (C_{\mathcal{F}, w}^{\text{AgSeq}(\sigma)} \cap C_{\mathcal{F}, w}^{\rho}) = \emptyset. \quad (10)$$

Definition 45 (Weakly Privatized) Similarly we call a pointed Kripke frame $(\mathcal{F}, w) = (\langle S, R \rangle, w)$ weakly privatized w.r.t. formula $\varphi \in \mathcal{Q}_{\text{nsr}}$ iff for any root path $\sigma \in \text{RootP}(\varphi)$

$$\left(\forall \rho \in \bigcup_{l=0}^{\infty} \mathcal{A}_{\text{nsr}}^l \setminus \{\text{AgSeq}(\sigma)\} \right) (C_{\mathcal{F}, w}^{\text{AgSeq}(\sigma)} \cap C_{\mathcal{F}, w}^{\rho}) = \emptyset. \quad (11)$$

Definition 46 (Walk accessibility) For a pointed Kripke frame $(\mathcal{F}, w) = (\langle W, R \rangle, w)$ we define walk accessibility operator

$$\text{WAcc}_{\text{nsr}}(u, \mathcal{F}, w) := \{\sigma \in \text{RootW}_{\text{nsr}}(\mathcal{F}, w) \mid \pi_2 \pi_{|\sigma|} \sigma = u\} \quad (12)$$

to be the set of agent-alternating walks starting from w and ending in u . By π_k we denote the k th projection function.

Walk accessibility yields an alternative equivalent definition of privatized frames:

Lemma 47 (Privatized alt.) A pointed Kripke frame $(\mathcal{F}, w)$ is (strongly) privatized w.r.t. $\varphi \in \mathcal{Q}_{\text{nsr}}$ iff for any root path $\sigma \in \text{RootP}(\varphi)$

$$C_{\mathcal{F}, w}^{\text{AgSeq}(\sigma)} \neq \emptyset \text{ and } \left(\forall s \in C_{\mathcal{F}, w}^{\text{AgSeq}(\sigma)} \mid \text{AgSeq}(\text{WAcc}_{\text{nsr}}(s, \mathcal{F}, w)) \right) = 1. \quad (13)$$

Proof: Follows from Defs. 37, 44, and 46. $\square$

Lemma 48 (Weakly privatized alt.) A pointed Kripke frame $(\mathcal{F}, w)$ is weakly privatized w.r.t. $\varphi \in \mathcal{Q}_{\text{nsr}}$ iff for any root path $\sigma \in \text{RootP}(\varphi)$

$$\left(\forall s \in C_{\mathcal{F}, w}^{\text{AgSeq}(\sigma)} \mid \text{AgSeq}(\text{WAcc}_{\text{nsr}}(s, \mathcal{F}, w)) \right) = 1. \quad (14)$$

Proof: Follows from Defs. 37, 45 and 46. $\square$

Theorem 49 For DBI formula φ , the pointed action model $(\mathcal{U}_{\varphi}, 0)$ as constructed by [Algorithm 1](#) is privatized w.r.t. φ .

Proof by induction on the recursive construction of $\tilde{\mathcal{U}}_{\varphi}$ by [Algorithm 2](#).

Ind. hyp.: Given DBI goal formula φ , its proper subformula φ' , and action model $\mathcal{U}$ that does not contain any out-neighbors for agents in $\text{ta}(\varphi')$, the u' -sub model $(\tilde{\mathcal{U}}_{\varphi'}, u')$ after a recursive call `CreateActionModel`($\varphi', \mathcal{U}, u'$) is privatized w.r.t. φ' , when not considering any loops at its point u' .

Base case: For formula $B_i \xi$, where ξ is propositional, since $\text{AgSeq}(\text{RootP}(B_i \xi)) = \{\varepsilon, (i)\}$, we only need to check all states walk-accessible via these two sequences. Since in $\tilde{\mathcal{U}}_{B_i \xi}$ the state $m \notin \{0, -1\}$ is the only state accessible via an i -edge from the root 0, as the 0 itself has no incoming arrows, $\text{AgSeq}(\text{WAcc}_{\text{nsr}}(m, \tilde{\mathcal{U}}_{B_i \xi}, 0)) = \{i\}$. Furthermore the root 0 is only walk-accessible via the empty sequence ε , therefore $\text{AgSeq}(\text{WAcc}_{\text{nsr}}(0, \tilde{\mathcal{U}}_{B_i \xi}, 0)) = \{\varepsilon\}$. Lastly as the sink -1 is not walk-accessible via any of the two sequences ε and (i) in $\text{AgSeq}(\text{RootP}(B_i \xi))$ we conclude that $(\tilde{\mathcal{U}}_{B_i \xi}, 0)$ is indeed privatized w.r.t. $B_i \xi$.

Ind. step: If $\varphi = B_i \psi$ for DBI normal formula ψ , then $(\tilde{\mathcal{U}}_{\psi}, u')$ the u' -sub model returned by the recursive call of [Algorithm 2](#) on line 6 is privatized w.r.t. ψ , when considered without the reflexive i loop of u' by the IH. Since this recursive call on line 6. was the last line to execute, it follows that $\tilde{\mathcal{U}}_{\varphi} = \tilde{\mathcal{U}}_{\psi}$. As $i \notin \text{ta}(\psi)$ and the i -edge from u to u' is the only outgoing edge connecting u with the u' -sub model $(\tilde{\mathcal{U}}_{\psi}, u')$, we conclude that $(\tilde{\mathcal{U}}_{\varphi}, u)$ is privatized with respect to φ .

If $\varphi = B_i(\xi \wedge \psi)$ for propositional formula ξ and DBI normal formula ψ , the reasoning is analogous to the previous case, as privatization is a frame property and thus independent of the additional precondition added.

If $\varphi = \varphi_1 \wedge \varphi_2$, for DBI normal formulas φ_1 and φ_2 , by the IH for the recursive function call on line 12. $(\tilde{\mathcal{U}}_{\varphi_1}, u)$ is privatized w.r.t. φ_1 . Since $\text{ta}(\varphi_1) \cap \text{ta}(\varphi_2) = \emptyset$ it follows that $\text{RootP}(\varphi) = \text{RootP}(\varphi_1) \sqcup \text{RootP}(\varphi_2)$, hence by the IH for the recursive call on line 13. the resulting model $(\tilde{\mathcal{U}}_{\varphi}, u)$ is not only privatized w.r.t. φ_2 , but remains privatized w.r.t. φ_1 as well, thus is privatized w.r.t. φ .

Since by the initial call of Algorithm 2 on line 5 in Algorithm 1 $\mathcal{U}$ consists only of the state 0 without any in- or outgoing edges and since the point of the resulting model $\tilde{\mathcal{U}}_{\varphi}$ doesn't contain any incoming edges (in particular no loops), $\tilde{\mathcal{U}}_{\varphi}$ is privatized w.r.t. φ . It remains to investigate the execution of the call of Algorithm 4 on line 6, which however merely adds j -arrows from states to the new -1 state if they don't already have outgoing j -edges for every agent j , therefore $\mathcal{U}_{\varphi}$ is indeed privatized w.r.t. φ . $\square$

Lemma 50 Let the pointed update $(\mathcal{M} \odot \mathcal{U}, (w, \alpha))$ of a pointed Kripke model $(\mathcal{M}, w)$ with a pointed action model $(\mathcal{U}, \alpha)$ be defined.

For any $\rho \in \bigcup_{l=0}^{\infty} \mathcal{A}^l$,

$$(x, \beta) \in C_{\mathcal{M} \odot \mathcal{U}, (w, \alpha)}^{\rho} \implies \beta \in C_{\mathcal{U}, \alpha}^{\rho} \text{ and } x \in C_{\mathcal{M}, w}^{\rho}. \quad (15)$$

Proof Let $\mathcal{M} = \langle S, R, V \rangle$, $\mathcal{U} = \langle E, Q, \text{pre} \rangle$, and $\mathcal{M} \odot \mathcal{U} = \langle S', R', V' \rangle$. We use induction on $l = |\rho|$.

Base case $l = 0$: Since $\mathcal{A}^0 = \{\varepsilon\}$ the statement follows trivially from Defs 38 and 9.

Ind. step: Consider any agent sequence $\bar{\rho} = \rho \circ i$ of length $l+1$ and state $(x, \beta) \in C_{\mathcal{M} \odot \mathcal{U}, (w, \alpha)}^{\bar{\rho}}$. By Definition 38, there must exist a state $(v, \gamma) \in C_{\mathcal{M} \odot \mathcal{U}, (w, \alpha)}^{\rho}$ such that $(v, \gamma) R'_i(x, \beta)$, in particular, $v R_i x$ and $\gamma Q_i \beta$ by Definition 9. By IH for ρ of length l , we have $v \in C_{\mathcal{M}, w}^{\rho}$ and $\gamma \in C_{\mathcal{U}, \alpha}^{\rho}$. Thus, by Definition 38 both $x \in C_{\mathcal{M}, w}^{\bar{\rho}}$ and $\beta \in C_{\mathcal{U}, \alpha}^{\bar{\rho}}$. $\square$

Theorem 51 For a pointed action model $(\mathcal{U}, \alpha)$ privatized w.r.t. DBI formula φ and a pointed epistemic Kripke model $(\mathcal{M}, w)$, if their pointed update exists, then $(\mathcal{M} \odot \mathcal{U}, (w, \alpha))$ is weakly privatized with respect to φ .

Proof Suppose by contradiction that the pointed action model $(\mathcal{U}, \alpha)$ is privatized w.r.t. φ , but for some pointed Kripke model $(\mathcal{M}, v)$, the update $(\mathcal{M} \odot \mathcal{U}, (w, \alpha))$ exists, however is not weakly privatized w.r.t. φ . This means that there exists a root path $\sigma \in \text{RootP}(\varphi)$, agent sequence $\rho \in \bigcup_{l=0}^{\infty} \mathcal{A}_{n, sr}^l \setminus \text{AgSeq}(\sigma)$ and state $(x, \beta) \in C_{\mathcal{M} \odot \mathcal{U}, (w, \alpha)}^{\rho}$ and $(x, \beta) \in C_{\mathcal{M} \odot \mathcal{U}, (w, \alpha)}^{\text{AgSeq}(\sigma)}$. By Lemma 50 this implies that $\beta \in C_{\mathcal{U}, \alpha}^{\text{AgSeq}(\sigma)}$ and $\beta \in C_{\mathcal{U}, \alpha}^{\rho}$. However this contradicts the assumption that $(\mathcal{U}, \alpha)$ is privatized w.r.t. φ , as $C_{\mathcal{U}, \alpha}^{\text{AgSeq}(\sigma)} \cap C_{\mathcal{U}, \alpha}^{\rho} = \emptyset$ by Definition 44. $\square$

Theorem 52 For DBI formula φ and any pointed Kripke model $(\mathcal{M}, v)$, the pointed update between $(\mathcal{U}_{\varphi}, 0)$ (as synthesized by Algorithm 1) and $(\mathcal{M}, v)$ is defined and $(\mathcal{M} \odot \mathcal{U}, (w, \alpha))$ is weakly privatized.

Proof Follows immediately from Theorems 25, 49, and 51. $\square$

5. Conclusions

We presented a way to synthesize a pointed action model for a limited range of goal formulas that: (i) makes the goal formula true in the resulting model; (ii) privatizes any Kripke model to which it is

applied to, thus breaking the common knowledge of the model assumption and simulating totally private communication; (iii) preserves consistency whenever possible, marking a difference from, e.g., van Ditmarsch et al.^[6]; and (iv) has minimal side effects, in the sense that it changes as few beliefs unspecified in the goal formula as possible. In addition, the synthesized pointed action models are combined with pointed Kripke models via the new pointed update operation, which does not apply to the whole model globally, but rather it respects the structure of the privatization introduced in the synthesis. Since the pointed update is a subset of the full product update operation, the proposed update mechanism is quite efficient: only the initial update increases the size of the model at worst linearly in the size of the goal formula, while all subsequent updates via goal formulas with the same belief structure (or substructure thereof) only decrease the model size. Compare this with the exponential blow up after iterated updates using the standard product updates.

Compared with existing synthesis methods our work presents the following advantages: the synthesis methods proposed in the literatures^[6,8] do not focus on issues of privatization and of minimal change, which is a strong shortcoming when considering applications to distributed systems. Recall that the goal of the synthesis tasks is, given a goal formula φ , to output an update model that, applied to any initial model, outputs a model where φ holds. For some models however, φ could hold vacuously, in case e.g., there are no relations in the resulting model for a given agent involved in the formula. In our case, this is not possible: because of the proposed stratification of the model, an agent does not result in believing inconsistencies unless precisely instructed so by a sequence of updates. In other words, our synthesis method addresses belief consistency preservation, while other existing methods do not. While it is possible to construct AML or GAUL update models representing completely private communication^[9] and thus it is possible to synthesize them, there is no standardized update synthesis procedure for it. Finally, AAUL was proved to be undecidable^[22]. On the other hand, our language is a restriction of the standard doxastic language, and is hence decidable. Consequently, while not as powerful, it is computationally more appealing for the synthesis task.

We aim at extending the update mechanism so as to allow a wider range of goal formulas, introducing negations of modal operators, which might introduce ignorance, and disjunctions of modal operators, which have multiple possible realizations, while preserving the properties of privatization, minimality and consistency. Furthermore, we plan to introduce group updates, such as public announcements, providing full granularity in the update design.

Author contributions

The authors confirm their contributions to this study as follows: development of synthesis algorithms: Schlögl T; related research: Cignarale G; formal proofs: Schlögl T, Kuznets R; draft manuscript preparation: Cignarale G, Kuznets R. All authors reviewed the results and approved the final version of the manuscript.

Data availability

Data sharing is not applicable to this article as no datasets were generated or analyzed during this work.

Acknowledgments

This research was supported by the Austrian Science Fund (FWF) through projects ByZDEL (Grant No. P 33600-N) and A Logical

Framework for Graded Deontic Reasoning (10.55776/PAT2141924). We are grateful to Hans van Ditmarsch, Stephan Felber, Krisztina Fruzsá, Rojo Randrianomentsoa, Hugo Rincón Galeana, and Ulrich Schmid for multiple illuminating and inspiring discussions.

Conflict of interest

The authors declare that they have no conflict of interest.

Dates

Received 13 December 2024; Revised 8 November 2025; Accepted 9 March 2026; Published online 30 June 2026

References

- [1] Hintikka J. 1962. *Knowledge and belief: an introduction to the logic of the two notions*. Ithaca: Cornell University Press. <https://hdl.handle.net/2324/1000461361>
- [2] Fagin R, Halpern J, Moses Y, and Vardi M. 1995. *Reasoning About Knowledge*. US: MIT Press. [10.7551/mitpress/5803.001.0001](https://doi.org/10.7551/mitpress/5803.001.0001)
- [3] Plaza J. 1989. Logics of public communications. In *Proceedings of the Fourth International Symposium on Methodologies for Intelligent Systems: Poster Session Program*, eds. Emrich ML, Pfeifer MS, Hadzikadic M, Ras Z. Oak Ridge National Laboratory. pp. 201–216
- [4] van Ditmarsch H, van der Hoek W, Kooi B. 2007. *Dynamic epistemic logic*. Vol. 337. Dordrecht: Springer. doi: [10.1007/978-1-4020-5839-4](https://doi.org/10.1007/978-1-4020-5839-4)
- [5] Kooi B, Renne B. 2011. Generalized arrow update logic. In *TARK XIII, Theoretical Aspects of Rationality and Knowledge: Proceedings of the Thirteenth Conference (TARK 2011), Groningen, The Netherlands, July 12–14, 2011*. US: Association for Computing Machinery. pp. 205–211 [10.1145/2000378.2000403](https://doi.org/10.1145/2000378.2000403)
- [6] van Ditmarsch H, van der Hoek W, Kooi B, Kuijter L. 2020. Arrow update synthesis. *Information and Computation* 275:104544
- [7] Kooi B and Renne B. 2011. Arrow update logic. *The Review of Symbolic Logic* 4:536–559
- [8] Hales J. 2013. Arbitrary action model logic and action model synthesis. *2013 28th Annual ACM/IEEE Symposium on Logic and Computer Science (LICS), 25–28 June 2013, New Orleans, USA*. USA: IEEE. pp. 253–262 doi: [10.1109/LICS.2013.31](https://doi.org/10.1109/LICS.2013.31)
- [9] Baltag A, Renne B. 2016. Dynamic epistemic logic. In *The Stanford Encyclopedia of Philosophy*, ed. Zalta EN. Winter 2016 Edition. Metaphysics Research Lab, Stanford University. <https://plato.stanford.edu/archives/win2016/entries/dynamic-epistemic>
- [10] Artemov S. 2020. Observable Models. In *Logical Foundations of Computer Science. LFCS 2020. Lecture Notes in Computer Science*, eds. Artemov S, Nerode A. Vol. 11972. Cham: Springer. pp. 12–26. doi: [10.1007/978-3-030-36755-8_2](https://doi.org/10.1007/978-3-030-36755-8_2)
- [11] Herzig A. 2017. Dynamic epistemic logics: promises, problems, shortcomings, and perspectives. *Journal of Applied Non-Classical Logics* 27:328–341
- [12] Aucher G, Schwarzentruher F. 2013. On the complexity of dynamic epistemic logic. *TARK 2013: Theoretical Aspects of Rationality and Knowledge, Proceedings of the 14th Conference — Chennai, India*, ed. Schipper BC. Davis, US: University of California. pp. 19–28 doi: [10.48550/arXiv.1310.6406](https://doi.org/10.48550/arXiv.1310.6406)
- [13] Baltag A, Smets S. 2008. A qualitative theory of dynamic interactive belief revision. In *Logic and the Foundations of Game and Decision Theory (LOFT 7), Texts in Logic and Games*, eds. Bonanno G, van der Hoek W, Wooldridge M. Vol. 3. Netherlands: Amsterdam University Press. pp. 11–58 www.jstor.org/stable/j.ctt46mz4h.4
- [14] van Ditmarsch H. 2013. Revocable Belief Revision. *Studia Logica* 101:1185–1214
- [15] Lorini E. 2019. Exploiting belief bases for building rich epistemic structures. *Proceedings Seventeenth Conference on Theoretical Aspects of Rationality and Knowledge, Toulouse, France, 17–19 July 2019*, ed. Moss LS. Vol. 297. Electronic Proceedings in Theoretical Computer Science. Open Publishing Association. pp. 332–353 doi: [10.4204/EPTCS.297.21](https://doi.org/10.4204/EPTCS.297.21)
- [16] Lorini E. 2020. Rethinking epistemic logic with belief bases. *Artificial Intelligence* 282:103233
- [17] van Ditmarsch H, Halpern J, van der Hoek W, Kooi B. 2015. *Handbook of epistemic logic*. London, UK: College Publications.
- [18] Blackburn P, de Rijke M, Venema Y. 2001. *Modal Logic*. Vol. 53. UK: Cambridge University Press. doi: [10.1017/CBO9781107050884](https://doi.org/10.1017/CBO9781107050884)
- [19] Wen X, Yu Q, Liu Y. 2013. Multi-agent epistemic explanatory diagnosis via reasoning about actions. *IJCAI '13: Proceedings of the Twenty-Third international joint conference on Artificial Intelligence, Beijing, China, August 3–9, 2013*. USA: AAAI Press. pp. 1183–1190 <https://dl.acm.org/doi/10.5555/2540128.2540298>
- [20] van Benthem J, Smets S. 2015. Dynamic Logics of belief change. In *Handbook of Epistemic Logic*, ed. van Ditmarsch H, Halpern JY, van der Hoek W, Kooi B. College Publications. pp. 313–393 doi: [10.3166/jancl.17.129-155](https://doi.org/10.3166/jancl.17.129-155)
- [21] Hansson S. 2022. Logic of belief revision. In *The Stanford Encyclopedia of Philosophy*, ed. Zalta EN. US: Metaphysics Research Lab, Stanford University. <https://plato.stanford.edu/archives/spr2022/entries/logic-belief-revision>
- [22] van Ditmarsch H, van der Hoek W, Kuijter L. 2017. The undecidability of arbitrary arrow update logic. *Theoretical Computer Science* 693:1–12



Copyright: © 2026 by the author(s). Published by Maximum Academic Press, Fayetteville, GA. This article is an open access article distributed under Creative Commons Attribution License (CC BY 4.0), visit <https://creativecommons.org/licenses/by/4.0/>.