

觉知逻辑的个体信念更新

宋鹏飞 熊卫

摘 要: 本文以包含觉知算子的信念态度逻辑为基础, 研究该逻辑在增加个体信念更新算子后的扩充, 在给出这一扩充的公理系统后, 证明其对于包含觉知的多主体信念库语义模型的可靠性和完全性。进一步, 我们还比较了信念态度逻辑的个体信念更新与命题觉知逻辑的个体信念更新这两种不同的动态过程, 并且证明了这两种动态过程产生的模型具有互模拟关系。

关键词: 觉知; 个体信念更新; 信念库; 互模拟

中图分类号: B81 **文献标识码:** A

1 引言

动态认知逻辑 (Dynamic Epistemic Logic, [8]) 涉及两个主要的信息动态变化: 公开宣告 (Public Announcement) 和个体信念更新 (Private Belief Update)。信息的动态变化不仅会导致知识和信念的改变, 也会导致觉知 (Awareness) 的改变。本文主要讨论个体信念更新¹及其相应的觉知变化, 比如下面的例子:

例 1. 小张住在酒店, 他知道酒店的餐厅有提供咖啡的服务, 但是不确定餐厅今天是否提供咖啡。另外, 小张对于餐厅有提供茶水的服务没有觉知。这时, 小张听到有人说, 今天餐厅不会同时提供咖啡和茶水。得知这个消息之后, 小张虽然仍不知道餐厅今天是否提供咖啡和茶水, 但是觉知到了餐厅有提供茶水的服务这件事。事实上, 今天餐厅既不会提供咖啡, 也不会提供茶水。令 p 为“餐厅今天提供咖啡”, q 为“餐厅今天提供茶水”。小张的知识变化如图 1 所示。

收稿日期: 2021-05-27

作者信息: 宋鹏飞 中山大学逻辑与认知研究所
中山大学哲学系
songpf@mail2.sysu.edu.cn

熊卫 中山大学逻辑与认知研究所
中山大学哲学系
hssxwei@mail.sysu.edu.cn

¹本文中, 如果模型满足定义 7, 则该模型刻画了知识, 此时的信念更新也是知识更新。



图 1

图 1 是两个单主体模型。左边的模型表示小张得到信息之前的认知状态，其中只包含一个原子命题 p ，实心点是 p 为真的可能世界，空心点是 p 为假的可能世界，有下划线的可能世界是当前的可能世界，模型中的箭头表示主体的可通达关系，显然它是等价关系；右边的模型表示小张得到信息之后的认知状态，其中包含两个原子命题 p 和 q ，每个可能世界有两个点，左边的点表示 p 的真值，右边的点表示 q 的真值，实心为真，空心为假，与左侧模型一样，有下划线的可能世界为当前可能世界，箭头表示主体的可通达关系，它也是等价关系。

在模态逻辑中，不确定性可以通过命题在不同可能世界中的不同取值来刻画。对于某一主体在某一可能世界，如果一个命题在所有可通达的可能世界为真，则该主体相信（或知道）该命题为真。例 1 描述了一个很简单的场景，图 1 的可能世界模型刻画了该场景。但是，例 1 与通过模态逻辑来刻画知识和信念的经典方法有所不同。在模态逻辑中，我们假设任一主体都能够觉知到所有相关的原子命题，这一假设被称为封闭世界假设（Closed World Assumption, [3]），而例 1 中的主体得到信息导致模型增加了新的原子命题，显然，其不符合这一假设。因此，研究者尝试修改和完善克里普克模型，以支持主体对命题或公式有不同程度的觉知，以及主体的觉知发生变化。与此同时，在相应的逻辑语言中，也需要在认知逻辑中增加觉知算子以及导致觉知发生变化的动态算子。

形式刻画觉知的方案主要有两种：句法进路（[9–11]）和语义进路（[13, 14, 21, 22]）。这两种方案的差异主要体现在哲学观点的不同。其中，句法进路认为，每一个可能世界都是客观的，所有的原子命题在任一可能世界都有取值；而语义进路则认为，可能世界是主观的，是存在于主体的思想中的，因此，那些没有被主体觉知到的原子命题就会在相应的可能世界上没有取值。具体而言，句法进路在可能世界语义中增加了觉知函数，使每一个主体在每一个可能世界对应一个公式集；而语义进路则是将克里普克语义结构扩展为所有主观可能世界构成的完全格。

由于上述两种方案都是在克里普克语义上加以扩展，因此，在处理个体信念更新的问题时，都会产生同样的问题，即，当某一主体接收到某一条信息，原本的认知模型需要被复制成两个副本，一个副本用来刻画接收到信息的主体的认知状态，另一个副本用来刻画未接收到信息的主体的认知状态。这样做的结果是，在一系列个体信念更新之后，认知模型会发生几何级数的增长。

由罗里尼 (Lorini) 提出的信念态度逻辑 (Logic of Doxastic Attitude, [16, 17]) 能够有效地规避上述问题。信念态度逻辑是以信念库 (Belief Base, [12, 20, 23, 24]) 为基础的逻辑, 而信念库则是信念更新的 AGM 方案 ([1]) 的重要概念。在信念态度逻辑中, 对于个体信念更新, 只需要改变该个体的信念库, 而其他个体的信念库保持不变。由此, 信念态度逻辑给出了个体信念更新的一种“节省的”方案, 避免了复制模型导致的模型复杂度的几何级数增加。至此, 有两种个体信念更新方式, 即, 通过改变信念库而实现的个体信念更新和通过复制模型而实现的个体信念更新, 罗里尼 ([17]) 证明了这两种个体信念更新所产生的模型具有互模拟关系。但是, 罗里尼的理论中并不包含觉知, 因而不能涵盖前文提到的例子。对此, 罗里尼等人在 [18, 19] 提出了包含觉知算子的信念态度逻辑, 并形成了与上述两种方案并列的第三种刻画觉知的方案——信念库进路 (Belief Base Approach)。以此为基础, 本文将在包含觉知算子的情况下证明上述两种信念更新产生的模型具有互模拟关系, 并给出其成立的条件。另外, 虽然 [19] 列出了该逻辑的个体信念更新扩充的公理, 但并未证明这一扩充的可靠性和完全性。对此, 本文将使用句法翻译的方法 ([8]) 证明该扩充相对于多主体信念库语义的可靠性和完全性。

本文的各部分内容安排: 第 2 节到第 4 节是对 [18] 的回顾, 在第 2 节, 我们定义了包含觉知的信念态度逻辑的逻辑语言; 第 3 节给出信念库语义模型的定义, 进而规定不同的模型类; 第 4 节介绍了静态信念态度逻辑的公理系统以及该系统中不同的公理集合对于相应模型类的可靠性和完全性定理; 从第 5 节开始是本文的核心内容, 在介绍了该逻辑在增加个体更新动态算子之后的扩充后 ([19]), 我们用句法翻译的方法证明了该扩充对于信念库语义模型的可靠性和完全性; 第 6 节将信念态度逻辑的个体信念更新与命题觉知逻辑的个体信念更新作比较, 展现了前者在保持模型简洁性方面的优势, 最后我们证明了这两种动态过程产生的模型具有互模拟关系; 第 7 节得出文章的结论并介绍后续工作。

2 逻辑语言

在这一节, 我们回顾了包含觉知的信念态度逻辑 LDAA (Logic of Doxastic Attitudes with Awareness, [18]) 的语言。令 $Atm = \{p, q, \dots\}$ 是可数无限的原子命题集, 令 $Agt = \{1, \dots, n\}$ 是有限的主体集。LDAA 的语言由下面两层定义给出。首先是 $\mathcal{L}_0(Atm, Agt)$:

$$\alpha ::= p \mid \neg\alpha \mid \alpha_1 \wedge \alpha_2 \mid \Delta_i\alpha \mid \bigcirc_i\alpha,$$

其中 $p \in Atm$, $i \in Agt$ 。 $\mathcal{L}_{LDAA}(Atm, Agt)$ 是在 $\mathcal{L}_0(Atm, Agt)$ 的基础上增加隐信念算子得到的语言:

$$\varphi ::= \alpha \mid \neg\varphi \mid \varphi_1 \wedge \varphi_2 \mid \Box_i\varphi \mid \bigcirc_i\varphi,$$

其中, $\alpha \in \mathcal{L}_0(Atm, Agt)$, $i \in Agt$ 。

当前后文明确时, 可以将 $\mathcal{L}_0(Atm, Agt)$ 简写为 $\mathcal{L}_0$, 将 $\mathcal{L}_{LDAA}(Atm, Agt)$ 简写为 $\mathcal{L}_{LDAA}$ 。其他布尔连接词 $\vee, \rightarrow, \leftrightarrow, \top, \perp$ 则是由 $\neg$ 和 $\wedge$ 通过标准的方式定义。公式 $\Delta_i \alpha$ 读作“主体 i 显相信 α 为真”, 公式 $\bigcirc_i \alpha$ 读作“主体 i 觉知到 α ”。算子 Δ_i 可以被叠加, 这说明 $\mathcal{L}_{LDAA}$ 可以表达高阶信念, 例如 $\Delta_i \Delta_j \alpha$, 读作“主体 i 显相信主体 j 显相信 α 为真”。叠加可能是显信念算子和觉知算子的混合, 例如 $\Delta_i \bigcirc_j \alpha$, 读作“主体 i 显相信主体 j 觉知到 α ”。

公式 $\Box_i \varphi$ 读作“主体 i 隐相信 φ 为真”。它的对偶 $\Diamond_i$ 定义为:

$$\Diamond_i \varphi \stackrel{\text{def}}{=} \neg \Box_i \neg \varphi,$$

$\Diamond_i \varphi$ 读作“ φ 与主体 i 的显信念有一致性”。

值得注意的是, 模态词 $\bigcirc_i$ 在逻辑语言的两层定义中都有出现, 而模态词 Δ_i 只出现在第一层。因此, 觉知算子可以出现在显信念算子的辖域中, 而隐信念算子不能出现在显信念算子的辖域中²。另外, 显信念算子和隐信念算子都能出现在觉知算子的辖域中。这是因为, 本文研究的觉知是命题觉知, 即关于原子命题的觉知, 其含义是, 主体觉知到某一公式等价于觉知到构成该公式的所有原子命题。令 $Atm(\varphi)$ 表示出现在公式 φ 中的原子命题的集合, 其递归定义如下:

$$\begin{aligned} Atm(p) &= \{p\}, \\ Atm(\neg \varphi) &= Atm(\varphi), \\ Atm(\varphi_1 \wedge \varphi_2) &= Atm(\varphi_1) \cup Atm(\varphi_2), \\ Atm(\Delta_i \alpha) &= Atm(\alpha), \\ Atm(Y_i \varphi) &= Atm(\varphi), \text{ 其中 } Y \in \{\bigcirc, \Box\}. \end{aligned}$$

令 $\Gamma \subseteq \mathcal{L}_{LDAA}$ 是有限的公式集, 定义 $Atm(\Gamma) = \bigcup_{\varphi \in \Gamma} Atm(\varphi)$ 。

3 形式语义

这一部分回顾了包含觉知的信念库语义。([18]) 不同于克里普克语义的是, 信念库语义中的可通达关系不是模型的初始条件, 而是从信念库中计算得出的。这里首先给出“状态”的定义, 它类似于克里普克语义中的可能世界。

定义 1. 一个状态是一个多元组 $S = (B_1, \dots, B_n, A_1, \dots, A_n, V)$, 其中,

- $B_i \subseteq \mathcal{L}_0$ 是主体 i 的信念库, 其中, $i \in Agt$,

²这里对句法做出限制的原因是, 在下一节的语义定义中, 主体的信念可通达关系是由该主体的显信念计算出的, 也就是说, 隐信念是通过显信念来定义的。因此, 如果隐信念算子出现在显信念算子的辖域中, 将会产生循环定义的问题。规避这一问题方式是用“固定点”的方式来定义隐信念, 而这会显著增加语义的复杂性。

- $A_i = \text{Atm}(B_i)$ 是主体 i 的觉知集合, 其中, $i \in \text{Agt}$,
- $V \subseteq \text{Atm}$ 是该状态中那些为真的原子命题集。

$\mathbf{S}$ 是所有状态的集合。下面的定义给出 $\mathcal{L}_0$ 公式的语义解释。

定义 2. 对任意 $S = (B_1, \dots, B_n, A_1, \dots, A_n, V) \in \mathbf{S}$:

$$\begin{aligned}
 S \models p &\iff p \in V, \\
 S \models \neg \alpha &\iff S \not\models \alpha, \\
 S \models \alpha_1 \wedge \alpha_2 &\iff S \models \alpha_1 \text{ 且 } S \models \alpha_2, \\
 S \models \Delta_i \alpha &\iff \alpha \in B_i, \\
 S \models \bigcirc_i \alpha &\iff \text{Atm}(\alpha) \subseteq A_i.
 \end{aligned}$$

定义 3. 包含觉知的多主体信念库语义模型 MABA (Multi-agent Belief Model with Awareness) 是一个二元组 (S, Cxt) , 其中, $S \in \mathbf{S}$, $\text{Cxt} \subseteq \mathbf{S}$ 。

Cxt 是所有主体的语境或共同背景 (Common Ground, [25]), 指的是所有主体共同分享的信息。从这些信息和各自的显信念中, 主体能够做出推理。下面的定义规定了主体的信念可通达关系是由信念库计算出的。

定义 4. 对任意 $i \in \text{Agt}$, $\mathcal{R}_i$ 是 $\mathbf{S}$ 上的二元关系, 其定义为: 对任意 $S = (B_1, \dots, B_n, A_1, \dots, A_n, V), S' = (B'_1, \dots, B'_n, A'_1, \dots, A'_n, V') \in \mathbf{S}$,

$$(S, S') \in \mathcal{R}_i \text{ 当且仅当 } \forall \alpha \in B_i, S' \models \alpha.$$

在定义了可通达关系的基础上, 可以给出 $\mathcal{L}_{\text{LDA}}$ 公式的语义解释。其中, 布尔公式与标准的定义保持一致, 此处省略。

定义 5. 令 (S, Cxt) 是 MABA, 其中, $S = (B_1, \dots, B_n, A_1, \dots, A_n, V)$ 。有如下等价关系:

$$\begin{aligned}
 (S, \text{Cxt}) \models \alpha &\iff S \models \alpha, \\
 (S, \text{Cxt}) \models \Box_i \varphi &\iff \forall S' \in \text{Cxt}, \text{ 如果 } (S, S') \in \mathcal{R}_i \\
 &\quad \text{那么 } (S', \text{Cxt}) \models \varphi, \\
 (S, \text{Cxt}) \models \bigcirc_i \varphi &\iff \text{Atm}(\varphi) \subseteq A_i.
 \end{aligned}$$

下面两个定义给出了 MABA 的属性, 它们分别对应克里普克模型的自返性和持续性。

定义 6. MABA (S, Cxt) 满足全局可持续性 GC (Global Consistency) 当且仅当, 对于所有 $i \in \text{Agt}$ 和所有 $S' \in (\{S\} \cup \text{Cxt})$, 存在 $S'' \in \text{Cxt}$ 满足 $(S', S'') \in \mathcal{R}_i$ 。

定义 7. $\text{MABA}(S, Cxt)$ 满足信念正确 BC (Belief Correctness) 当且仅当 $S \in Cxt$ 且对于所有 $i \in \text{Agt}$ 和所有 $S' \in Cxt$, $(S', S') \in \mathcal{R}_i$ 。

对于 $X \subseteq \{\text{GC}, \text{BC}\}$, MABA_X 是满足 X 中的所有条件的 MABA 模型类。 $\text{MABA}_\emptyset$ 是所有 MABA 模型构成的模型类, 简记为 **MABA**。很容易证明 $\text{MABA}_{\{\text{BC}\}} = \text{MABA}_{\{\text{GC}, \text{BC}\}}$ 。

令 $\varphi \in \mathcal{L}_{\text{LDAA}}$, φ 对于 MABA_X 模型类有效, 当且仅当, 对所有 $(S, Cxt) \in \text{MABA}_X$, $(S, Cxt) \models \varphi$, 简记为 $\models_{\text{MABA}_X} \varphi$ 。 φ 对于 MABA_X 模型类可满足, 当且仅当, $\neg\varphi$ 对于 MABA_X 模型类不是有效的。

4 公理化

这一部分首先给出 LDAA 逻辑对于不同模型类的公理集, 然后介绍它们关于相应模型类的可靠性和完全性定理及证明思路。

基本的 LDAA 逻辑是在经典命题逻辑的基础上增加以下公理和推理规则的扩充 ([18]):

$$\begin{aligned}
 & (\Box_i \varphi \wedge \Box_i (\varphi \rightarrow \psi)) \rightarrow \Box_i \psi & (\mathbf{K}_{\Box_i}) \\
 & \Delta_i \alpha \rightarrow \Box_i \alpha & (\mathbf{Int}_{\Delta_i, \Box_i}) \\
 & \Delta_i \alpha \rightarrow \bigcirc_i \alpha & (\mathbf{Int}_{\Delta_i, \bigcirc_i}) \\
 & \bigcirc_i \varphi \leftrightarrow \bigwedge_{p \in \text{Atm}(\varphi)} \bigcirc_i p & (\mathbf{AGPP}) \\
 & \frac{\varphi}{\Box_i \varphi} & (\mathbf{Nec}_{\Box_i})
 \end{aligned}$$

对于 $X \subseteq \{\mathbf{D}_{\Box_i}, \mathbf{T}_{\Box_i}\}$, 令 LDAA_X 是 LDAA 逻辑增加下述公理的扩充:

$$\begin{aligned}
 & \neg(\Box_i \varphi \wedge \Box_i \neg\varphi) & (\mathbf{D}_{\Box_i}) \\
 & \Box_i \varphi \rightarrow \varphi & (\mathbf{T}_{\Box_i})
 \end{aligned}$$

对于逻辑 LDAA_X , $X \subseteq \{\mathbf{D}_{\Box_i}, \mathbf{T}_{\Box_i}\}$, 其中任意公式 $\varphi \in \mathcal{L}_{\text{LDAA}}$, 用 $\vdash_{\text{LDAA}_X} \varphi$ 来表示 φ 是 LDAA_X 的定理。对于 $\mathcal{L}_{\text{LDAA}}$ 的公式集 Γ , 如果不存在公式 $\varphi_1, \dots, \varphi_m \in \Gamma$ 满足 $\vdash_{\text{LDAA}_X} (\varphi_1 \wedge \dots \wedge \varphi_m) \rightarrow \perp$, 则 Γ 与 LDAA_X 一致。特别地, φ 与 LDAA_X 一致当且仅当 $\{\varphi\}$ 与 LDAA_X 一致。

很容易得出, $\text{LDAA}_{\{\mathbf{D}_{\Box_i}, \mathbf{T}_{\Box_i}\}}$ 和 $\text{LDAA}_{\{\mathbf{T}_{\Box_i}\}}$ 相等, 因为公理 $\mathbf{D}_{\Box_i}$ 是 $\text{LDAA}_{\{\mathbf{T}_{\Box_i}\}}$ 的定理。

接下来将介绍可靠性和完全性定理, 由于其证明过程较长, 这里只阐述 [18] 的证明思路。可靠性的证明比较简单, 只需要验证每一个公理对于相应的模型类

有效, 每一个推理规则都能保持公式的有效性。对于完全性的证明, 需要定义典范模型, 而典范模型的定义需要克里普克语义结构, 因此, 需要给出 LDAA 的两种类似克里普克结构的语义模型, 它们被称为概念语义和准概念语义, 在这两种语义中, 可通达关系是模型的初始条件, 然后只需证明信念库语义与这两种语义具有等价性。接下来只需要用典范模型的方法证明 LDAA_X 对于具有相应性质的准概念模型具有完全性, 即可得出完全性定理。

定理 1. 令 $X \subseteq \{\mathbf{D}_{\Box_i}, \mathbf{T}_{\Box_i}\}$ 。LDAA_X 对于模型类 $\mathbf{MABA}_{\{cf(x):x \in X\}}$ 是可靠和完全的。([18])

证明. 过程参见 [18] 的第三节和第四节。□

5 LDAA 的个体信念更新

这一部分首先给出 LDAA 在增加个体信念更新算子之后的扩充。([19]) 具体来说, 在 $\mathcal{L}_{\text{LDAA}}(\text{Atm}, \text{Agt})$ 的基础上增加算子 $[+_i\alpha]$, 可以得到语言 $\mathcal{L}_{\text{LDAA-PBE}}(\text{Atm}, \text{Agt})$:

$$\varphi ::= \alpha \mid \neg\varphi \mid \varphi_1 \wedge \varphi_2 \mid \Box_i\varphi \mid \bigcirc_i\varphi \mid [+_i\alpha]\varphi,$$

其中, $i \in \text{Agt}$, $\alpha \in \mathcal{L}_0$, LDAA-PBE 读作“包含觉知与个体信念扩张的信念态度逻辑”, PBE 是 Private Belief Expansion 的简写。

与第 2 节类似, $\mathcal{L}_{\text{LDAA-PBE}}(\text{Atm}, \text{Agt})$ 可被记作 $\mathcal{L}_{\text{LDAA-PBE}}$ 。公式 $[+_i\alpha]\varphi$ 读作“主体 i 的信念库增加 α 后, φ 为真”。在定义 5 的基础上增加如下可满足关系的定义, 动态算子 $[+_i\alpha]$ 可在 MABA 中得到解释。

定义 8. 令 (S, Cxt) 是 MABA, 其中 $S = (B_1, \dots, B_n, A_1, \dots, A_n, V)$ 。有如下等价关系:

$$(S, \text{Cxt}) \models [+_i\alpha]\varphi \iff (S^{+_i\alpha}, \text{Cxt}) \models \varphi,$$

其中, $S^{+_i\alpha} = (B_1^{+_i\alpha}, \dots, B_n^{+_i\alpha}, A_1^{+_i\alpha}, \dots, A_n^{+_i\alpha}, V^{+_i\alpha})$,

$$V^{+_i\alpha} = V,$$

对所有 $j \in \text{Agt}$:

$$\begin{aligned} B_j^{+_i\alpha} &= B_j \cup \{\alpha\} & i &= j, \\ B_j^{+_i\alpha} &= B_j & i &\neq j, \\ A_j^{+_i\alpha} &= \text{Atm}(B_j^{+_i\alpha}). \end{aligned}$$

事件 $+_i\alpha$ 只包含主体 i 得到信息 α 并将 α 增加到其信念库中, 而其他主体的信念库保持不变。因为觉知集合是由信念库计算得出的, 所以, 在这一过程中, 主体 i 的觉知集合间接地受到影响。

另外, 个体信念更新不能保持模型的信念正确 BC 或全局可持续性 GC, 因此, 对个体信念更新的讨论是相对于所有 MABA 构成的模型类 **MABA**。

逻辑 **LDAA-PBE** 在 **LDAA** 的基础上增加如下公理和推理规则的扩充³:

$$\begin{array}{ll}
 [+_i\alpha]p \leftrightarrow p & \text{(AI)} \\
 [+_i\alpha]\neg\varphi \leftrightarrow \neg [+_i\alpha]\varphi & \text{(PE\&N)} \\
 [+_i\alpha](\varphi_1 \wedge \varphi_2) \leftrightarrow ([+_i\alpha]\varphi_1 \wedge [+_i\alpha]\varphi_2) & \text{(PE\&C)} \\
 [+_i\alpha]\Box_j\varphi \leftrightarrow \Box_j\varphi & i \neq j \\
 [+_i\alpha]\Box_i\varphi \leftrightarrow \Box_i(\alpha \rightarrow \varphi) & \text{(PE\&IB)} \\
 [+_i\alpha]\Delta_j\beta \leftrightarrow \Delta_j\beta & i \neq j \text{ 或 } \alpha \neq \beta \\
 [+_i\alpha]\Delta_i\alpha \leftrightarrow \top & \text{(PE\&EB)} \\
 [+_i\alpha]\bigcirc_j\varphi \leftrightarrow \bigcirc_j\varphi & i \neq j \\
 [+_i\alpha]\bigcirc_i\varphi \leftrightarrow \top & \text{Atm}(\varphi) \subseteq \text{Atm}(\alpha) \\
 [+_i\alpha]\bigcirc_i\varphi \leftrightarrow \bigwedge_{p \in \text{Atm}(\varphi) \setminus \text{Atm}(\alpha)} \bigcirc_i p & \text{其他} \quad \text{(PE\&A)} \\
 \frac{\varphi}{[+_i\alpha]\varphi} & \text{(PBEG)}
 \end{array}$$

任意公式 $\varphi \in \mathcal{L}_{\text{LDAA-PBE}}$, 用 $\vdash_{\text{LDAA-PBE}} \varphi$ 来表示 φ 是 **LDAA-PBE** 的定理。根据上述公理, 可以给出归约函数 $\text{red}_{\text{LDAA-PBE}}$ 的递归定义, 该函数将 $\mathcal{L}_{\text{LDAA-PBE}}$ 的公式可以转化为等价的 $\mathcal{L}_{\text{LDAA}}$ 公式。

定义 9. 函数 $\text{red}_{\text{LDAA-PBE}}$ 递归定义如下:

1. $\text{red}_{\text{LDAA-PBE}}(p) = p$
2. $\text{red}_{\text{LDAA-PBE}}(\Delta_j\alpha) = \Delta_j\text{red}_{\text{LDAA-PBE}}(\alpha)$
3. $\text{red}_{\text{LDAA-PBE}}(\neg\varphi) = \neg\text{red}_{\text{LDAA-PBE}}(\varphi)$
4. $\text{red}_{\text{LDAA-PBE}}(\varphi \wedge \psi) = \text{red}_{\text{LDAA-PBE}}(\varphi) \wedge \text{red}_{\text{LDAA-PBE}}(\psi)$
5. $\text{red}_{\text{LDAA-PBE}}(\Box_j\varphi) = \Box_j\text{red}_{\text{LDAA-PBE}}(\varphi)$

³注意到, **LDAA-PBE** 中并不包含个体更新算子叠加的公理, 即对应定义 9 第 13 项的公理。其原因是, 证明 **LDAA-PBE** 的完全性并不是直接通过 **LDAA-PBE** 的公理, 而是利用定义 9 的归约函数将所有 **LDAA-PBE** 公式转化为 **LDAA** 公式。**LDAA-PBE** 新增加的六条公理已经能够确保该归约函数可以将所有的 **LDAA-PBE** 转化为 **LDAA** 公式。

6. $red_{LDAA-PBE}(\bigcirc_j \varphi) = \bigcirc_j red_{LDAA-PBE}(\varphi)$
7. $red_{LDAA-PBE}([+_i \alpha]p) = red_{LDAA-PBE}(p)$
8. $red_{LDAA-PBE}([+_i \alpha]\neg\psi) = red_{LDAA-PBE}(\neg[+_i \alpha]\psi)$
9. $red_{LDAA-PBE}([+_i \alpha](\psi_1 \wedge \psi_2)) = red_{LDAA-PBE}([+_i \alpha]\psi_1 \wedge [+_i \alpha]\psi_2)$
10. $red_{LDAA-PBE}([+_i \alpha]\Box_j \varphi) = red_{LDAA-PBE}(\Box_j \varphi) \quad i \neq j$
 $red_{LDAA-PBE}([+_i \alpha]\Box_i \varphi) = red_{LDAA-PBE}(\Box_i(\alpha \rightarrow \varphi))$
11. $red_{LDAA-PBE}([+_i \alpha]\Delta_j \beta) = red_{LDAA-PBE}(\Delta_j \beta) \quad i \neq j \text{ 或 } \alpha \neq \beta$
 $red_{LDAA-PBE}([+_i \alpha]\Delta_i \alpha) = red_{LDAA-PBE}(\top)$
12. $red_{LDAA-PBE}([+_i \alpha]\bigcirc_j \varphi) = red_{LDAA-PBE}(\bigcirc_j \varphi) \quad i \neq j$
 $red_{LDAA-PBE}([+_i \alpha]\bigcirc_i \varphi) = red_{LDAA-PBE}(\top) \quad Atm(\varphi) \subseteq Atm(\alpha)$
 $red_{LDAA-PBE}([+_i \alpha]\bigcirc_i \varphi) = red_{LDAA-PBE}\left(\bigwedge_{p \in Atm(\varphi) \setminus Atm(\alpha)} \bigcirc_i p\right) \text{ 其他}$
13. $red_{LDAA-PBE}([+_i \alpha][+_j \beta]\varphi) = red_{LDAA-PBE}([+_i \alpha]red_{LDAA-PBE}([+_j \beta]\varphi))$

觉知的归约公式的意思是：通过将 α 增加到信念库中，主体 i 将出现在 α 的原子命题增加到他的觉知集合中；隐信念的归约公式的意思是，主体 i 在信念库中增加 α 后能够推出 φ ，当且仅当，在信念更新之前，主体 i 就能够推出 α 蕴涵 φ ；显信念的归约公式的意思是，主体 i 的信念库增加 α ，其他主体的信念库保持不变。

接下来，我们对 [19] 的内容做补充，证明 LDAA-PBE 相对于 **MABA** 的可靠性和完全性。在证明可靠性和完全性之前，需要首先说明上述归约函数 $red_{LDAA-PBE}$ 给出的等价式是 LDAA-PBE 的定理，即，对于所有的 $\varphi \in \mathcal{L}_{LDAA-PBE}$ ，有 $\vdash \varphi \leftrightarrow red_{LDAA-PBE}(\varphi)$ 。这里存在一个问题，在做归纳证明时，通常需要将归纳假设应用在某一公式的所有子公式上，但是，对于包含个体信念更新算子的公式，这一方法是不够的。例如， $\neg[+_i \alpha]\psi$ 并不是 $[+_i \alpha]\neg\psi$ 的子公式。因此，需要定义如下的公式复杂度函数，然后针对公式复杂度来提出归纳假设。

定义 10. 公式复杂度函数 $c: \mathcal{L}_{LDAA-PBE} \rightarrow \mathbb{N}$ 的定义如下：

$$\begin{aligned}
 c(p) &= 1 \\
 c(\neg\varphi) &= 1 + c(\varphi) \\
 c(\varphi \wedge \psi) &= 1 + \max(c(\varphi), c(\psi))
 \end{aligned}$$

$$\begin{aligned}
c(\Delta_i \alpha) &= 1 + c(\alpha) \\
c(Y_i \varphi) &= 1 + c(\varphi) \text{ 其中 } Y \in \{\bigcirc, \Box\} \\
c([+_i \alpha] \varphi) &= (c(\alpha) + 2) \cdot c(\varphi)
\end{aligned}$$

公式复杂度函数 c 最后一项的参数 2 并不是任意的, 它是保证下面的不等式成立的最小自然数。

引理 1. 对于所有 $\alpha, \beta \in \mathcal{L}_0$, 以及所有 $\varphi, \psi, \psi_1, \psi_2 \in \mathcal{L}_{\text{LDAA-PBE}}$,

1. $c(\psi) \geq c(\varphi)$ φ 是 ψ 的子公式
2. $c([+_i \alpha] \neg \psi) > c(\neg [+_i \alpha] \psi)$
3. $c([+_i \alpha] (\psi_1 \wedge \psi_2)) > c([+_i \alpha] \psi_1 \wedge [+_i \alpha] \psi_2)$
4. $c([+_i \alpha] \Box_i \varphi) > c(\Box_i (\alpha \rightarrow \varphi))$
5. $c([+_i \alpha] \bigcirc_i \varphi) > c \left(\bigwedge_{p \in \text{Atm}(\varphi) \setminus \text{Atm}(\alpha)} \bigcirc_i p \right)$

证明. 只需要逐项验证这些不等式成立即可。 □

上述不等式说明被归约前的公式的复杂度大于归约得到的公式的复杂度⁴。有了这一结果, 我们可以对公式复杂度做数学归纳来证明下面的命题。

命题 1. 令 $\varphi \in \mathcal{L}_{\text{LDAA-PBE}}$ 。可以得到:

$$\vdash_{\text{LDAA-PBE}} \varphi \leftrightarrow \text{red}_{\text{LDAA-PBE}}(\varphi)$$

证明. 对 $c(\varphi)$ 做归纳:

当 φ 是原子命题 p 时, $\vdash_{\text{LDAA-PBE}} p \leftrightarrow p$ 显然成立。

假设对于所有 $c(\varphi) \leq n$ 的 φ , 有 $\vdash_{\text{LDAA-PBE}} \varphi \leftrightarrow \text{red}_{\text{LDAA-PBE}}(\varphi)$, 其中, $n \in \mathbb{N}$ 。

当 φ 的形式如 $\neg \psi$ 、 $\psi_1 \wedge \psi_2$ 、 $\Delta_i \alpha$ 、 $\Box_i \psi$ 、 $\bigcirc_i \psi$ 时, 结论可从归纳假设和引理 1 的第一项得证。

当 φ 是 $[+_i \alpha] p$ 时, 结论可从公理 **AI**、归纳假设、定义 10 得证。

当 φ 是 $[+_i \alpha] \neg \varphi$ 时, 结论可从公理 **PE&N**、归纳假设、引理 1 的第 2 项得证。

当 φ 是 $[+_i \alpha] (\varphi_1 \wedge \varphi_2)$ 时, 结论可从公理 **PE&C**、归纳假设、引理 1 的第 3 项得证。

当 φ 是 $[+_i \alpha] \Box_j \varphi$ 时, 其中 $i \neq j$, 结论可从公理 **PE&IB**、归纳假设、引理 1 的第 4 项得证。

⁴需要注意的是, 引理 1 省略了归约函数 $\text{red}_{\text{LDAA-PBE}}$ 中那些不等关系显然成立的部分。

当 φ 是 $[+_i\alpha]\Box_i\varphi$ 时, 结论可从公理 **PE&IB**、归纳假设、定义 10 得证。

当 φ 是 $[+_i\alpha]\Delta_j\beta$ 时, 结论可从公理 **PE&EB**、归纳假设、定义 10 得证。

当 φ 是 $[+_i\alpha]\bigcirc_j\varphi$ 时, 其中, $i \neq j$ 或 $Atm(\varphi) \subseteq Atm$, 结论可从公理 **PE&A**、归纳假设、定义 10 得证。

当 φ 是 $[+_i\alpha]\bigcirc_j\varphi$ 且并非是前一种情况时, 结论可从公理 **PE&A**、归纳假设、引理 1 的第 5 项得证。

当 φ 是 $[+_i\alpha][+_j\beta]\varphi$ 时, 由 $red_{LDAA-PBE}$ 的定义, 可转化为上述情况。 $\square$

下面的命题说明归纳函数 $red_{LDAA-PBE}$ 确实可以将 $\mathcal{L}_{LDAA-PBE}$ 公式转化为 $\mathcal{L}_{LDAA}$ 公式。

命题 2. 令 $\varphi \in \mathcal{L}_{LDAA-PBE}$ 。则, $red_{LDAA-PBE}(\varphi) \in \mathcal{L}_{LDAA}$ 。

证明. 对 $c(\varphi)$ 做归纳即可给出证明, 此处省略具体步骤。 $\square$

根据上述结果, 我们可以证明 LDAA-PBE 的可靠性和完全性定理。

定理 2. LDAA-PBE 对于模型类 **MABA** 是可靠和完全的。

证明. 对于可靠性, 只需要验证每一个公理都是有效的, 每一个推理规则都能在模型类 **MABA** 保持有效性。因为 LDAA 对于 **MABA** 是可靠的, 只需要验证新增的公理和推理规则, 而其中大多数都是显然的, 这里只证明公理 **PE&A** 是有效的。当 $i \neq j$ 时, 根据语义, $[+_i\alpha]\bigcirc_j\varphi \leftrightarrow \bigcirc_j\varphi$ 显然成立。当 $Atm(\varphi) \subseteq Atm(\alpha)$ 时, 根据语义, $\Delta_i\alpha \rightarrow \bigcirc_i\varphi$ 是重言式。由此可得, $[+_i\alpha](\Delta_i\alpha \rightarrow \bigcirc_i\varphi)$ 是重言式。由前面的等价式可知, 后者等价于 $[+_i\alpha]\Delta_i\alpha \rightarrow [+_i\alpha]\bigcirc_i\varphi$ 。因为 $[+_i\alpha]\Delta_i\alpha \leftrightarrow \top$, 可以得到 $[+_i\alpha]\bigcirc_i\varphi \leftrightarrow \top$ 。对于其他情况, 由 $\bigcirc_i\varphi \leftrightarrow \bigwedge_{p \in Atm(\varphi) \setminus Atm(\alpha)} \bigcirc_i p \wedge \bigwedge_{p \in Atm(\varphi) \cap Atm(\alpha)} \bigcirc_i p$ 可以得到 $[+_i\alpha]\bigcirc_i\varphi \leftrightarrow [+_i\alpha]\bigwedge_{p \in Atm(\varphi) \setminus Atm(\alpha)} \bigcirc_i p \wedge [+_i\alpha]\bigwedge_{p \in Atm(\varphi) \cap Atm(\alpha)} \bigcirc_i p$ 。根据前一种情况, 很容易得出 $[+_i\alpha]\bigwedge_{p \in Atm(\varphi) \cap Atm(\alpha)} \bigcirc_i p \leftrightarrow \top$ 。因此, 有 $[+_i\alpha]\bigcirc_i\varphi \leftrightarrow [+_i\alpha]\bigwedge_{p \in Atm(\varphi) \setminus Atm(\alpha)} \bigcirc_i p$ 。因为 $p \notin Atm(\alpha)$, 对 p 的觉知不受 $[+_i\alpha]$ 事件的影响, 可以得出, $[+_i\alpha]\bigcirc_i\varphi \leftrightarrow \bigwedge_{p \in Atm(\varphi) \setminus Atm(\alpha)} \bigcirc_i p$ 。

对于完全性, 令 $\models_{\mathbf{MABA}} \varphi$ 。根据 LDAA-PBE 的可靠性以及 $\vdash_{LDAA-PBE} \varphi \leftrightarrow red_{LDAA-PBE}(\varphi)$, 可得 $\models_{\mathbf{MABA}} red_{LDAA-PBE}(\varphi)$ 。因为 $red_{LDAA-PBE}(\varphi)$ 不包含任何动态算子, 所以, 由 LDAA 的完全性可得, $\vdash_{LDAA} red_{LDAA-PBE}(\varphi)$ 。由于 LDAA 是 LDAA-PBE 的子系统, 可以得到 $\vdash_{LDAA-PBE} red_{LDAA-PBE}(\varphi)$ 。根据命题 1, 有 $\vdash_{LDAA-PBE} \varphi$ 。 $\square$

6 LDAA-PBE 与命题觉知逻辑的个体信念更新

在这一部分, 我们将针对个体信念更新的问题比较 LDAA-PBE 与命题觉知逻辑 LPA (Logic of Propositional Awareness), 后者首先出现在 [10], 是一般觉知逻辑 LGA (Logic of General Awareness, [9]) 的特殊情况。

LPA 的语言 $\mathcal{L}_{\text{LPA}}(\text{Atm}, \text{Agt})$ 由下面的语法给出:

$$\varphi ::= p \mid \neg\varphi \mid \varphi_1 \wedge \varphi_2 \mid \mathbf{B}_i\varphi \mid \mathbf{A}_i\varphi \mid \mathbf{X}_i\varphi,$$

其中, $p \in \text{Atm}$, $i \in \text{Agt}$ 。公式 $\mathbf{B}_i\varphi$ 读作“主体 i 隐相信 φ 为真”, 公式 $\mathbf{A}_i\varphi$ 读作“主体 i 觉知到 φ ”, 公式 $\mathbf{X}_i\varphi$ 读作“主体 i 显相信 φ 为真”。与前文一致, 其他布尔连接词 $\vee, \rightarrow, \leftrightarrow, \top, \perp$ 由 $\neg$ 和 $\wedge$ 通过标准的方式定义, $\mathcal{L}_{\text{LPA}}(\text{Atm}, \text{Agt})$ 简写为 $\mathcal{L}_{\text{LPA}}$ 。在语义层面, 只需要将 LGA 的语义模型的觉知函数的取值设定为原子命题集的幂集, 就可以得到 LPA 的语义模型。

定义 11. 命题觉知模型 PAM (Propositional Awareness Model) 是一个四元组 $M = (\Omega, \Rightarrow, \rho, \pi)$, 其中,

- Ω 是非空的可能世界集,
- $\Rightarrow: \text{Agt} \times \Omega \longrightarrow 2^\Omega$ 是信念可通达函数,
- $\rho: \text{Agt} \times \Omega \longrightarrow 2^{\text{Atm}}$ 是命题觉知函数,
- $\pi: \text{Atm} \longrightarrow 2^\Omega$ 是赋值函数。

PAM 是所有 PAM 构成的集合。对于 PAM $M = (\Omega, \Rightarrow, \rho, \pi)$ 及 $s \in \Omega$, 二元组 (M, s) 被称为 PAM 的点模型。 $t \in \Rightarrow(i, s)$ 可记作 $s \Rightarrow_i t$ 。下面给出 $\mathcal{L}_{\text{LPA}}$ 的公式相对于 PAM 的点模型的语义解释。

定义 12. 给定 PAM $M = (\Omega, \Rightarrow, \rho, \pi)$ 及可能世界 $s \in \Omega$, 对于 $\mathcal{L}_{\text{LPA}}$ 的公式, 有如下等价关系:

$$\begin{aligned} (M, s) \models p &\iff s \in \pi(p), \\ (M, s) \models \neg\varphi &\iff (M, s) \not\models \varphi, \\ (M, s) \models \varphi \wedge \psi &\iff (M, s) \models \varphi \text{ 且 } (M, s) \models \psi, \\ (M, s) \models \mathbf{B}_i\varphi &\iff \forall t \in \Omega, \text{ 如果 } s \Rightarrow_i t \text{ 那么 } (M, t) \models \varphi, \\ (M, s) \models \mathbf{A}_i\varphi &\iff \text{Atm}(\varphi) \subseteq \rho(i, s), \\ (M, s) \models \mathbf{X}_i\varphi &\iff (M, s) \models \mathbf{B}_i\varphi \text{ 且 } (M, s) \models \mathbf{A}_i\varphi. \end{aligned}$$

对比 LPA 和 LDAA 的定义, 可以看出这两种逻辑是基于不同的哲学出发点。在 LPA 的语义中, 主体的信念可通达关系和觉知集合都是模型的初始条件, 而主体的显信念是由这两个条件计算得出的; 而在 LDAA 的语义中, 模型唯一的初始条件是主体的显信念, 而隐信念和觉知都是由显信念计算得出的。由此可见, LDAA 减少了理论的前提假设, 更符合奥卡姆剃刀的原则。另外, 这两种逻辑刻画了不同意义的显信念概念。在下面给出的翻译函数中也能看出这种差异, LPA 的显信念并没有直接对应 LDAA 的显信念, 而是对应 LDAA 的隐信念和觉知的合取。实

际上, **LDAA** 的显信念表示当前活跃在主体认知中的信念, 即工作记忆中的信念 (主体的信念库即该主体的工作记忆); 而 **LPA** 的显信念表示, 在主体觉知到的原子命题构成的信念中, 该主体相信为真的那些信念, 这样的信念并不一定处在工作记忆中, 但是能被主体所理解。值得重视的是, [18] 建立了由 **LPA** 到 **LDAA** 的多项式嵌入, 这说明后者以更少的前提假设提供了更强的表达力。

接下来, 为了证明互模拟定理, 需要定义 $\mathcal{L}_{\text{LPA}}$ 和 $\mathcal{L}_{\text{LDAA}}$ 之间的翻译函数:

$$tr : \mathcal{L}_{\text{LPA}} \longrightarrow \mathcal{L}_{\text{LDAA}}:$$

$$\begin{aligned} tr(p) &= p \quad \text{其中 } p \in \text{Atm}, \\ tr(\neg\varphi) &= \neg tr(\varphi), \\ tr(\varphi_1 \wedge \varphi_2) &= tr(\varphi_1) \wedge tr(\varphi_2), \\ tr(\mathbf{A}_i\varphi) &= \bigcirc_i tr(\varphi), \\ tr(\mathbf{B}_i\varphi) &= \square_i tr(\varphi), \\ tr(\mathbf{X}_i\varphi) &= \bigcirc_i tr(\varphi) \wedge \square_i tr(\varphi). \end{aligned}$$

$$tr^{\leftarrow} : \mathcal{L}_{\text{LDAA}} \longrightarrow \mathcal{L}_{\text{LPA}}:$$

$$\begin{aligned} tr^{\leftarrow}(p) &= p \quad \text{其中 } p \in \text{Atm}, \\ tr^{\leftarrow}(\neg\varphi) &= \neg tr^{\leftarrow}(\varphi) \quad \text{其中 } \varphi \text{ 不是 } \Delta_i\alpha \text{ 的形式}, \\ tr^{\leftarrow}(\varphi_1 \wedge \varphi_2) &= tr^{\leftarrow}(\varphi_1) \wedge tr^{\leftarrow}(\varphi_2), \\ tr^{\leftarrow}(\bigcirc_i\varphi) &= \mathbf{A}_i tr^{\leftarrow}(\varphi), \\ tr^{\leftarrow}(\square_i\varphi) &= \mathbf{B}_i tr^{\leftarrow}(\varphi), \\ tr^{\leftarrow}(\Delta_i\alpha) &= \mathbf{X}_i tr^{\leftarrow}(\alpha), \\ tr^{\leftarrow}(\neg\Delta_i\alpha) &= \text{null}. \end{aligned}$$

需要注意的是 $tr^{\leftarrow}$ 并不是 tr 的反函数, 因为需要在 $tr^{\leftarrow}$ 中给出 Δ_i 算子的翻译, 而 tr 不会涉及到 Δ_i 算子。实际上, $tr^{\leftarrow}$ 是一个偏函数, 其自变量如果是显信念的否定式, 它的函数值为空。这样做的原因是, 如果将 **LDAA** 的显信念的否定直接翻译成 **LPA** 的显信念否定, 有可能会产生矛盾。根据 [17], **LDAA** 的显信念表示主体的工作记忆, 所以有可能 α 不在主体 i 的工作记忆中, 但 $\square_i\alpha \wedge \bigcirc_i\alpha$ 为真。

以 **PAM** 为基础, 范·迪特玛希 (van Ditmarsch) 等人刻画了信念和命题觉知的动态变化 ([3–5, 7]), 其中, 个体层面的动态在 [5] 得到研究, 其方法是使用行动模型 ([2, 15])。它存在的问题是, 在一系列个体信念更新的操作之后, 通过与行动模型做乘积计算得出的 **PAM** 会呈现几何级数增长。为解决这一问题, 在罗里尼提出的信念态度逻辑中 ([17]), 初始的信念库模型经过一系列个体信念更新

之后,其数量级只会发生相对于信念更新次数的线性增长。然而,罗里尼只给出了以信念库语义模型为基础的个体信念更新过程,且不包含觉知。为应对这一问题,接下来,我们将以 PAM 作为初始模型得出互模拟的结论,其中同时包含觉知的变化。

[18] 的第 3 节给出了 LDAA 的三种语义对于一个有限公式集成立,而将 PAM 转换为 MABA 也需要利用这个语义等价性的结论,因此,这里我们也只考虑 $\mathcal{L}_{LDAA}$ 的有限公式集。

信念库模型的个体信念更新

令 $M = (\Omega, \Rightarrow, \rho, \pi)$ 是 PAM (Ω 有可能是无限集), $\Sigma \subseteq \mathcal{L}_{LPA}$ 是对子公式封闭的公式集, $tr(\Sigma) = \{tr(\varphi) : \varphi \in \Sigma\}$ 。为了证明的方便,令 $\Sigma' = \Sigma \cup \{A_i p : p \in \Sigma, i \in Agt\}$ 。很显然, $tr(\Sigma') \subseteq \mathcal{L}_{LDAA}$ 是一个有限集,且对子公式封闭。

根据 [18] 的定理 3 和第 3 节的语义等价性证明,可以找到一个有限集 Cxt 和一个满射 $\sigma : \Omega \rightarrow Cxt$ 满足,对每一个 $s \in \Omega$,如果 $\sigma(s) = S$,其中 $S = (B_1, \dots, B_n, A_1, \dots, A_n, V)$,那么,

- 对每一个 $p \in \Sigma'$: $tr(p) \in V$ 当且仅当 $s \in \pi(p)$,
- 对每一个 B_i 和每一个 $p \in \Sigma'$: $tr(p) \in Atm(B_i)$ 当且仅当 $p \in \rho(i, s) \cap \Sigma'^5$,
- 对每一个 A_i : $A_i = Atm(B_i)$,
- 对每一个 $t \in \Omega$: 如果 $s \Rightarrow_i t$,那么 $(S, \sigma(t)) \in \mathcal{R}_i$,
- 对每一个 $S' \in Cxt$: 如果 $(S, S') \in \mathcal{R}_i$,那么存在 $t \in \Omega$ 满足 $S' = \sigma(t)$ 。

很容易验证,对于每一个 $S \in Cxt$, (S, Cxt) 是 MABA,对每一个 $\varphi \in \Sigma'$ 和每一个 $s \in \Omega$, $(M, s) \models \varphi$ 当且仅当 $(\sigma(s), Cxt) \models tr(\varphi)$ 。

接下来对 (S, Cxt) 应用个体信念更新。假设主体 i 的信念库增加了 α ,其中, $\alpha \in \mathcal{L}_0(Atm(tr(\Sigma'), Agt), tr^{\leftarrow}(\Delta_i \alpha) \in \Sigma'$ 。然后,根据定义 8,可以得到 $(S^{+i\alpha}, Cxt)$ 。

为了证明互模拟关系,需要将 $(S^{+i\alpha}, Cxt)$ 翻译回 PAM $M' = (\Omega', \Rightarrow', \rho', \pi')$:

- $\Omega' = \{s_{S'} : S' \in S^{+i\alpha} \cup Cxt\}$,
- 对每一个 $i \in Agt$ 和每一个 $s_{S'} \in \Omega'$:
 $\Rightarrow'(i, s_{S'}) = \{s_{S''} : S'' \in Cxt \text{ 且 } (S', S'') \in \mathcal{R}_i\}$,
- 对每一个 $i \in Agt$ 和每一个 $s_{S'} \in \Omega'$: $\rho'(i, s_{S'}) = A'$,
- 对每一个 $p \in Atm(\Sigma')$: $\pi'(p) = \{s_{S'} \in \Omega' : p \in V'\}$ 。

⁵这一项解释了为什么我们令 Σ' 包括 $\{A_i p : p \in \Sigma, i \in Agt\}$ 。如果不这样做,对模型的过滤就不能保证主体 i 在某一信念库的觉知集合与 Σ 和主体 i 在 PAM 的相应可能世界的觉知集合的交集相等。为了证明互模拟关系,需要保证觉知集合在 Σ 的限定内是不变的。

PAM 的个体信念更新

与前一部分相类似, 我们用过滤的方法将 $M = (\Omega, \Rightarrow, \rho, \pi)$ 转换成有限的 PAM $M_{\Sigma'} = (\Omega_{\Sigma'}, \Rightarrow_{\Sigma'}, \rho_{\Sigma'}, \pi_{\Sigma'})$ 。由库伊和雷恩 (Kooi & Renne, [15]) 提出的箭头更新模型 (Arrow Update Model), 这里定义该模型相对于 PAM 的形式。

定义 13. 一个箭头更新模型是一个三元组 $\mathcal{U} = \{\mathbf{O}, \tau, \mathcal{A}\}$, 其中,

- $\mathbf{O}$ 是非空的输出集, 其中每一个元素被称为一个输出,
- $\tau: \text{Agt} \times \mathbf{O} \times \mathbf{O} \rightarrow \mathcal{L}_{\text{LPA}} \times \mathcal{L}_{\text{LPA}}$ 是一个偏函数,
- $\mathcal{A}: \text{Agt} \times \mathbf{O} \rightarrow 2^{\text{Atm}}$ 是觉知变化函数。

为了表达上的方便, 对每一个 $i \in \text{Agt}$ 和每一个 $o', o'' \in \mathbf{O}$, 如果 $\tau(i, o', o'') = (\varphi, \psi)$, 用 $\tau_1(i, o', o'')$ 来表示 φ , 用 $\tau_2(i, o', o'')$ 来表示 ψ 。

箭头更新模型 $(\mathcal{U}, o)$ 作用于 PAM 点模型 (M, s) 之后, 产生一个新的 PAM 点模型, 我们称之为乘积模型。下面是乘积模型的定义。

定义 14. 令 $(\mathcal{U}, o)$ 是一个箭头更新模型, 其中, $\mathcal{U} = \{\mathbf{O}, \tau, \mathcal{A}\}$ 。令 (M, s) 是一个 PAM 点模型, 其中, $M = (\Omega, \Rightarrow, \rho, \pi)$ 。 $(\mathcal{U}, o)$ 作用于 (M, s) 产生的乘积模型是 $(M \otimes \mathcal{U}, (s, o))$, 其中, $M \otimes \mathcal{U} = (\Omega', \Rightarrow', \rho', \pi')$:

- $\Omega' = \Omega \times \mathbf{O}$,
- 对每一个 $i \in \text{Agt}$ 和每一个 $(s', o') \in \Omega'$,
 $\Rightarrow' (i, (s', o')) = \{(s'', o'') \in \Omega' : s' \Rightarrow_i s'', \tau(i, o', o'') \text{ 有定义},$
 $(M, s') \models \tau_1(i, o', o'') \text{ 且 } (M, s'') \models \tau_2(i, o', o'')\}$,
 $\rho'(i, (s', o')) = \rho(i, s') \cup \mathcal{A}(i, o')$,
- 对每一个 $p \in \text{Atm}$, $\pi'(p) = \{(s', o') \in \Omega' : (M, s') \models p\}$ 。

箭头更新模型以 $|\mathbf{O}|$ 为倍数增大了初始模型。具体来说, 对每一个输出 $o' \in \mathbf{O}$ 和 PAM 的每一个可能世界 s' , 会产生 s' 的一个增加了下标 o' 的复制 (s', o') 。而且, 对每一个 $i \in \text{Agt}$ 和每一个输出对 (o', o'') , 箭头更新模型指定了起始条件 $\tau_1(i, o', o'')$ 和目标条件 $\tau_2(i, o', o'')$, 它们分别被 s' 和 s'' 满足, 进而保证在乘积模型中, (s', o') 和 (s'', o'') 具有 i -可通达关系。这样规定使得, 可通达关系在乘积模型中被保留的充分必要条件为相应的两个初始可能世界对应于该输出分别满足起始条件和目标条件。 $\mathcal{A}$ 函数的功能是扩大主体的觉知集合。具体来说, 它将每一个主体和每一个输出对应一个原子命题集, 在乘积模型中, 这一原子命题集将是该主体在被该输出下标的可能世界的觉知集合的子集。直观上讲, 这意味着在乘积模型中, 主体的觉知集合因为增加了一些原子命题作为元素而被扩大。

为了刻画在其他主体的信念不变的情况下某一主体的信念更新, 需要定义一个特殊的箭头更新模型——个体箭头更新模型 (Private Arrow Update Model)。它

只包含两个输出： o_1 和 o_2 。对于输出 o_1 ，他的觉知集合在每一个可能世界都会扩大，而在 φ 为真的可能世界，主体 i 减少他的可通达世界；对于输出 o_2 ，没有任何事发生。

定义 15. 对于主体 i 、公式 φ 、觉知扩张集合 $AE \subseteq Atm$ ，其中 $Atm(\varphi) \subseteq AE$ ，相应的个体更新模型是 $\mathcal{U}^{(\varphi, AE)_i} = (\mathcal{O}, \tau, \mathcal{A})$ ，定义如下：

- $\mathcal{O} = \{o_1, o_2\}$,
- 对于每一个 $k, h \in \{1, 2\}$ 和每一个 $j \in Agt$ ， $\tau(j, o_k, o_h)$ 有定义当且仅当 $(k = 1 \text{ 且 } h = 2) \text{ 或 } k = h = 2$,
- $\tau(i, o_1, o_2) = (\top, \varphi)$ 且 $\tau(i, o_2, o_2) = (\top, \top)$,
- 对所有 $j \neq i$ ， $\tau(j, o_1, o_2) = (\top, \top)$ 且 $\tau(j, o_2, o_2) = (\top, \top)$,
- 如果 $i = j$ 且 $k = 1$ ， $\mathcal{A}(j, o_k) = AE$ ，
否则， $\mathcal{A}(j, o_k) = \emptyset$ 。

令 $\mathcal{U}^{(tr^{\leftarrow}(\alpha), Atm(\alpha))_i} = (\mathcal{O}, \tau, \mathcal{A})$ 是对于主体 i 、公式 $tr^{\leftarrow}(\alpha)$ 、觉知扩张集合 $Atm(\alpha)$ 的个体箭头更新模型。给定 PAM 点模型 (M, s) ，其中 $M = (\Omega, \Rightarrow, \rho, \pi)$ ， $M \otimes \mathcal{U}^{(tr^{\leftarrow}(\alpha), Atm(\alpha))_i}$ 是 (M, s) 和 $(\mathcal{U}^{(tr^{\leftarrow}(\alpha), Atm(\alpha))_i}, o_1)$ 的乘积模型。具体来说，PAM $M \otimes \mathcal{U}^{(tr^{\leftarrow}(\alpha), Atm(\alpha))_i} = (\Omega', \Rightarrow', \rho', \pi')$ 的定义如下：

- $\Omega' = \{(s', o_1) : s' \in \Omega\} \cup \{(s', o_2) : s' \in \Omega\}$,
- 对所有 $s' \in \Omega$ 和所有 $j \in Agt$:
 当 $i = j$ 时， $\Rightarrow'(j, (s', o_1)) = \{(s'', o_2) : s' \Rightarrow_j s'' \text{ 且 } (M, s'') \models \varphi\}$ ，
 当 $i \neq j$ 时， $\Rightarrow'(j, (s', o_1)) = \{(s'', o_2) : s' \Rightarrow_j s''\}$ ，
 $\Rightarrow'(j, (s', o_2)) = \{(s'', o_2) : s' \Rightarrow_j s''\}$ ，
 当 $i = j$ 时， $\rho'(j, (s', o_1)) = \rho(j, s') \cup Atm(\alpha)$ ，
 当 $i \neq j$ 时， $\rho'(j, (s', o_1)) = \rho(j, s')$ 如果 $i \neq j$ ，
 $\rho'(j, (s', o_2)) = \rho(j, s')$ ，
 对所有 $p \in Atm$:
 $\pi'(p) = \{(s', o_1), (s', o_2) : s' \in \pi(p)\}$ 。

在下文中，将 $(M \otimes \mathcal{U}^{(tr^{\leftarrow}(\alpha), Atm(\alpha))_i}, (s, o_1))$ 记作 $(M, s)^{(tr^{\leftarrow}(\alpha), Atm(\alpha))_i}$ 。

互模拟

接下来我们证明两种个体信念更新产生的模型具有互模拟关系。在此之前，有两点需要注意。第一，互模拟关系应该局限在原子命题集的一个有限子集内；第二，与标准的克里普克模型的互模拟不同，这里的互模拟关系应该将模型中的觉知函数考虑在内。范·迪特玛希等人提供了包含觉知的互模拟定义 ([6])，在他们的文章中，这一互模拟被称为标准互模拟。

定义 16. 令 $Q \subseteq Atm$, 给定两个 PAM $M = (\Omega, \Rightarrow, \rho, \pi)$ 和 $M' = (\Omega', \Rightarrow', \rho', \pi')$, 这两个模型的 Q 标准互模拟是关系 $\mathcal{R}(Q) \subseteq \Omega \times \Omega'$ 满足, 对每一个 $(s, s') \in \mathcal{R}(Q)$, 每一个 $i \in Agt$, 和每一个 $p \in Q$:

- 原子命题: $s \in \pi(p)$ 当且仅当 $s' \in \pi'(p)$,
- 觉知: $Q \cap \rho(i, s) = Q \cap \rho'(i, s')$,
- 正向条件: 如果 $s \Rightarrow_i t$, 那么存在 $t' \in \Rightarrow' (i, s')$ 满足 $(t, t') \in \mathcal{R}(Q)$,
- 反向条件: 如果 $s' \Rightarrow_i t'$, 那么存在 $t \in \Rightarrow (i, s)$ 满足 $(t, t') \in \mathcal{R}(Q)$ 。

对于两个 PAM 点模型 (M, s) 和 (M', s') , 其中 $M = (\Omega, \Rightarrow, \rho, \pi)$, $M' = (\Omega', \Rightarrow', \rho', \pi')$, 如果存在互模拟关系 $\mathcal{R}(Q)$, 且 $(s, s') \in \mathcal{R}(Q)$, 则这两个点模型有 Q 互模拟关系, 记做 $(M, s) \simeq_Q (M', s')$ 。

根据 [6] 的命题 21, 可以得到, 如果 $(M, s) \simeq_Q (M', s')$, 那么 (M, s) 和 (M', s') 满足 $\mathcal{L}_{LPA}(Q, Agt)$ 中的相同的公式。

下面的定理是这一部分的核心结论, 这一互模拟关系的左右两边各是一个 PAM 点模型, 其中, 左边的模型是通过信念库语义的个体信念更新得到的, 右边的模型是通过个体箭头更新模型得到的。

定理 3. $(M', s_{S+i\alpha}) \simeq_{Atm(\Sigma)} (M_{\Sigma'}, [s]_{\Sigma'})^{(tr^{\leftarrow}(\alpha), Atm(\alpha))_i}$

证明. 令 $M' = (\Omega', \Rightarrow', \rho', \pi')$, 将 $M \otimes \mathcal{U}^{(tr^{\leftarrow}(\alpha), Atm(\alpha))_i}$ 记做 M'' , 令 $M'' = (\Omega'', \Rightarrow'', \rho'', \pi'')$, 令初始 PAM $M = (\Omega, \Rightarrow, \rho, \pi)$ 。很容易得到:

- $\Omega' = \Omega_{\Sigma'} \cup \{s_{S+i\alpha}\}$,
- $\Omega'' = \{([t]_{\Sigma'}, o_1) : [t]_{\Sigma'} \in \Omega_{\Sigma'}\} \cup \{([t]_{\Sigma'}, o_2) : [t]_{\Sigma'} \in \Omega_{\Sigma'}\}$ 。

定义二元关系 $\mathcal{R}(Atm(\Sigma)) \subseteq \Omega' \times \Omega''$ 如下:

- $(s_{S+i\alpha}, ([s]_{\Sigma'}, o_1)) \in \mathcal{R}(Atm(\Sigma))$,
- 对所有 $[t]_{\Sigma'} \in \Omega_{\Sigma'}$, $([t]_{\Sigma'}, ([t]_{\Sigma'}, o_2)) \in \mathcal{R}(Atm(\Sigma))$ 。

用常规的方式就可以验证 $\mathcal{R}(Atm(\Sigma))$ 定义了 M' 和 M'' 的标准互模拟关系, 因此, $(M', s_{S+i\alpha})$ 具有 $(M'', ([s]_{\Sigma'}, o_1))$ 互模拟关系。□

值得注意的是, 定理 3 的成立需要满足四个条件: 第一, 具有互模拟关系的两个模型都是有限模型, 因此, 需要将初始模型用对子公式封闭的公式集来过滤得到有限模型; 第二, 过滤时, 对子公式封闭的公式集 Σ 需要转化为 Σ' , 从而包括所有主体对 Σ 的原子命题的觉知命题; 第三, 如果 α 是主体 i 的信念库中增加的公式, 那么, $tr^{\leftarrow}(\Delta_i \alpha)$ 是 Σ' 中的公式; 最后, α 是 $\mathcal{L}_0$ 公式。

在证明过程中, 定理 3 与 [17] 的定理 11 使用的方法有一些相同点和不同点。相同点是, 在模型的动态变化过程中, 两者都使用了信念库的更新和个体箭头更

新模型。不同点是,前者以 PAM 作为证明的起始模型,而后者以信念库语义模型作为证明的起始模型,基于 [18] 的由 LPA 到 LDAA 的多项式嵌入,以 PAM 作为证明的起始模型允许我们将 LDAA 的技术直接应用在 LPA 上;由于证明的起始模型的不同,两者在证明过程中需要对模型做不同的处理,前者需要将模型过滤以得到有限模型,进而通过 [18] 的模型转化方法得到 MABA,而后者则是将所有显信念公式作为原子命题来处理,进而从信念库语义模型得到克里普克模型;最重要的是,相比于后者,前者的模型中包含觉知,因此,定义 14 给出的个体箭头更新模型的最后一项是对主体觉知的更新,定义 16 提出了互模拟需要满足觉知相等的条件。

更进一步,定理 3 为 LDAA 和 LPA 的语义在技术方面和哲学方面的差异提供了很好的说明。技术方面,在 LPA 的语义模型 PAM 中,需要通过主体的信念可通达关系来确认该主体的隐信念,而显信念又是通过隐信念和觉知计算出的,因此,主体的个体信念更新意味着需要为每一个这样的主体复制整个模型;反观 LDAA 的语义模型 MABA,每个主体的信念库都是相互独立的,认知模型是从主体的信念库构造出来的,这意味着,当某一主体的显信念发生变化时,认知模型也会自动发生变化。哲学方面,对于 PAM,隐信念和觉知都是模型的初始条件,因此,为实现主体显信念的更新,需要同时改变主体的隐信念和觉知;而 MABA 与之不同,我们只需要更新主体信念库中的公式集即可实现该主体的显信念更新。

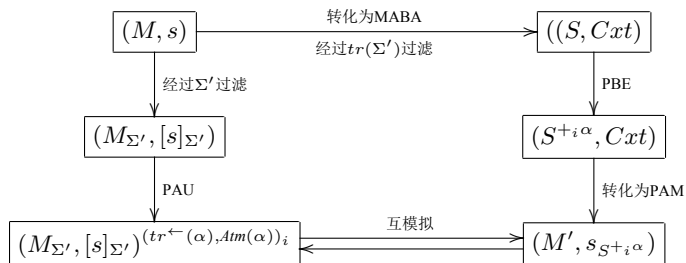


图 2: 图表展示了以 PAM 为初始模型的两种个体信念更新方案的关系。其中, PBE 指信念库的个体信念更新, PAU 指个体箭头更新。

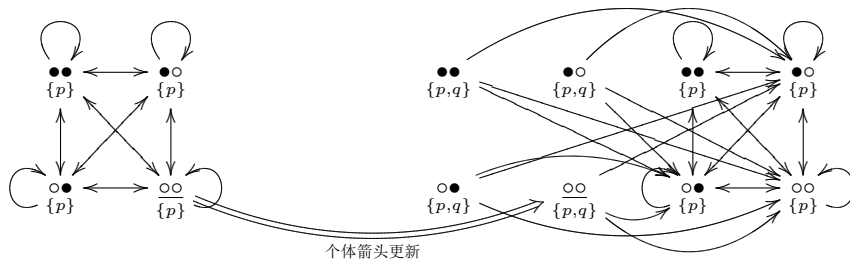


图 3: 这是用个体箭头更新过程来展示例 1 的模型变化过程。图中有两个单主体模型。左边的模型是信念更新之前小张的认知状态,右边的模型信念更新之后的认知状态。大括号中的原子命题表示小张在该可能世界的觉知集合,模型中其他元素的解释与图 1 保持一致。

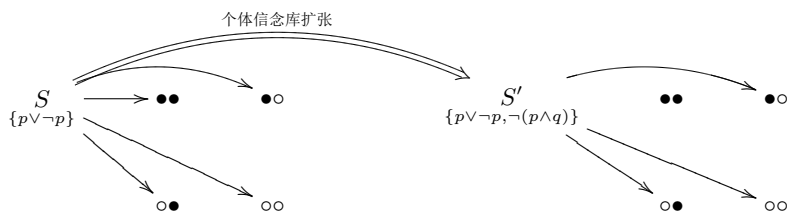


图 4: 这是用信念库语义模型的个体信念更新过程来展示例 1 的模型变化过程。大括号中的公式表示小张信念库中的公式, 图中的其他元素的含义与图 1 保持一致。

图 2 展示了这一节的完整过程。通过本节我们证明了, 尽管 LPA 的显信念算子 X_i 被翻译成 LDAA 的 $\square_i$ 和 $\bigcirc_i$ 算子, 针对 $\triangle_i$ 算子的个体信念更新仍然可以用在基于 PAM 的动态变化中。回到文章开头提出的例子, 图 3 和图 4 分别展示了这两种信念更新的模型变化过程。很容易看出, 基于信念库语义模型的个体信念更新大大简化了模型的转化过程, 避免了模型的复制。

7 结论与后续工作

本文证明了以信念库语义为基础的个体信念更新的动态逻辑 LDAA-PBE 相对于多主体信念库语义的可靠性和完全性。同时, 本文将两种个体信念更新方案做比较, 它们分别是: 基于信念库语义的个体信念更新、基于克里普克语义的个体箭头更新。通过模型转换的方法, 我们证明了这两种方案形成的模型具有互模拟关系。其中, 前一种方案只会使模型的规模发生线性增加, 而后一种方案在每次个体信念更新时都必须复制模型, 从而在一系列更新之后会使模型规模产生几何级数增加。这说明前一种方案在达到相同目的的同时, 能够有效保持模型的简洁性。

在后续工作中, 我们不仅会考虑对原子命题的觉知, 也会研究某一主体对其他主体的觉知, 尝试以信念库语义为基础建立关于身份的觉知逻辑, 而关于这一觉知的动态变化也会纳入该逻辑的讨论范围。

参考文献

- [1] C. E. Alchourrón, P. Gärdenfors and D. Makinson, 1985, "On the logic of theory change: Partial meet contraction and revision functions", *The Journal of Symbolic Logic*, **50**(2): 510–530.
- [2] J. van Benthem, J. van Eijck and B. Kooi, 2006, "Logics of communication and change", *Information and Computation*, **204**(11): 1620–1662.

- [3] H. van Ditmarsch and T. French, 2011, “Becoming aware of propositional variables”, in M. Banerjee and A. Seth (eds.), *Logic and Its Applications*, pp. 204–218, Berlin, Heidelberg: Springer Berlin Heidelberg.
- [4] H. van Ditmarsch and T. French, 2014, “Semantics for knowledge and change of awareness”, *Journal of Logic, Language and Information*, **23**(2): 169–195.
- [5] H. van Ditmarsch, T. French and F. R. Velázquez-Quesada, 2012, “Action models for knowledge and awareness”, *Proceedings of the 11th International Conference on Autonomous Agents and Multiagent Systems*, **Vol. 2**, pp. 1091–1098, Valencia, Spain: International Foundation for Autonomous Agents and Multiagent Systems.
- [6] H. van Ditmarsch, T. French, F. R. Velázquez-Quesada and Y. N. Wáng, 2018, “Implicit, explicit and speculative knowledge”, *Artificial Intelligence*, **256**: 35–67.
- [7] H. van Ditmarsch, T. French, F. R. Velázquez-Quesada and Y. N. Wáng, 2013, “Knowledge, awareness, and bisimulation”, in B. C. Schipper (ed.), *Proceedings of the 14th Conference on Theoretical Aspects of Rationality and Knowledge (TARK 2013)*, Chennai, India, January 7–9, 2013, pp. 61–70.
- [8] H. van Ditmarsch, W. van der Hoak and B. Kooi, 2007, *Dynamic Epistemic Logic*, Dordrecht, The Netherlands: Springer.
- [9] R. Fagin and J. Y. Halpern, 1987, “Belief, awareness, and limited reasoning”, *Artificial Intelligence*, **34**(1): 39–76.
- [10] J. Y. Halpern, 2001, “Alternative semantics for unawareness”, *Games and Economic Behavior*, **37**(2): 321–339.
- [11] J. Y. Halpern and L. C. Rêgo, 2008, “Interactive unawareness revisited”, *Games and Economic Behavior*, **62**(1): 232–262.
- [12] S. O. Hansson, 1993, “Theory contraction and base contraction unified”, *Journal of Symbolic Logic*, **58**(2): 602–625.
- [13] A. Heifetz, M. Meier and B. C. Schipper, 2006, “Interactive unawareness”, *Journal of Economic Theory*, **130**(1): 78–94.
- [14] A. Heifetz, M. Meier and B. C. Schipper, 2008, “A canonical model for interactive unawareness”, *Games and Economic Behavior*, **62**(1): 304–324.
- [15] B. Kooi and B. Renne, 2011, “Generalized arrow update logic”, *Proceedings of the 13th Conference on Theoretical Aspects of Rationality and Knowledge*, pp. 205–211, Groningen, The Netherlands: Association for Computing Machinery.
- [16] E. Lorini, 2018, “In praise of belief bases: Doing epistemic logic without possible worlds”, in S. A. McIlraith and K. Q. Weinberger (eds.), *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence (AAAI-18)*, the 30th Innovative Applications of Artificial Intelligence (IAAI-18), and the 8th AAAI Symposium on Educational Advances in Artificial Intelligence (EAAI-18), New Orleans, Louisiana, USA, February 2–7, 2018, pp. 1915–1922, Palo Alto, CA: AAAI Press.

- [17] E. Lorini, 2020, “Rethinking epistemic logic with belief bases”, *Artificial Intelligence*, **282**: 103233.
- [18] E. Lorini and P. Song, 2020, “Grounding awareness on belief bases”, in M. A. Martins and I. Sedlár (eds.), *Dynamic Logic: New Trends and Applications*, pp. 170–186, Cham: Springer International Publishing.
- [19] E. Lorini and P. Song, “A computationally grounded logic of awareness”, *Journal of Logic and Computation*, forthcoming.
- [20] D. Makinson, 1985, “How to give it up: A survey of some formal aspects of the logic of theory change”, *Synthese*, **62(3)**: 347–363.
- [21] S. Modica and A. Rustichini, 1994, “Awareness and partitioned information structures”, *Theory and Decision*, **37(1)**: 107–124.
- [22] S. Modica and A. Rustichini, 1999, “Unawareness and partitioned information structures”, *Games and Economic Behavior*, **27(2)**: 265–298.
- [23] B. Nebel, 1992, “Syntax based approaches to belief revision”, in P. Gärdenfors (ed.), *Belief Revision*, pp. 52–88, Cambridge: Cambridge University Press.
- [24] H. Rott, 1998, “‘Just because’: Taking belief bases seriously”, in S. R. Buss, P. Hájek and P. Pudlák (eds.), *Logic Colloquium '98: Proceedings of the 1998 ASL European Summer Meeting*, **Vol. 3**, pp. 387–408, Cambridge: Cambridge University Press.
- [25] R. Stalnaker, 2002, “Common ground”, *Linguistics and Philosophy*, **25(5)**: 701–721.

(责任编辑: 映之)

Private Belief Update in Awareness Logic

Pengfei Song Wei Xiong

Abstract

Based on Doxastic Attitude Logic with Awareness, this article studies an extension with private update operator, presents an axiomatization of this extension, and proves soundness and completeness of it with respect to multi-agent belief base models with awareness. Moreover, we compare the private belief base expansion in our logic and the private arrow update in Propositional Awareness Logic, and prove a bisimulation result between two models generated by the two kinds of updates.

Pengfei Song Institute of Logic and Cognition, Sun Yat-sen University
Department of Philosophy, Sun Yat-sen University
songpf@mail2.sysu.edu.cn

Wei Xiong Institute of Logic and Cognition, Sun Yat-sen University
Department of Philosophy, Sun Yat-sen University
hssxwei@mail.sysu.edu.cn