

多 Agent 信念修正研究^{*})

孙召春 高 阳 贾松茂 陈世福

(南京大学计算机软件新技术国家重点实验室 南京210093)

The Research of Multi-agent Belief Revision

SUN Zhao-Chun GAO Yang JIA Song-Mao CHEN Shi-Fu

(State Key Lab for Novel Software Technology, Nanjing University, Nanjing 210093)

Abstract Belief Revision is a theory that studies how to integrate new information into original belief set. Classical BR theory uses AGM frame, but it only resolves problems in single agent BR system. Multi-agent BR faces problems such as the collision of many information sources and how to maximize the logic consistence of multi-agent system. On the basis of game theory model, we form profit matrix under different BR strategies in multi-agent system and try to get the best strategy that satisfies logic consistence of the system through negotiation.

Keywords Belief revision, AGM, Multi-agent belief revision, Game theory

1. 前言

在学习系统或知识库系统等智能应用中,都要解决如何将新知识综合到系统原有的信念集中。最初的研究方法采用真值维护(TMS, Truth Maintenance System),但只能在孤立上下文中起作用^[1]。八十年代中期,Gärdenfors、Alchourron 与 Makinson 等人共同创立了信念修正理论(BR, Belief revision),解决了开放上下文中的知识维护,后继研究者通常将此理论简称为 AGM。

从九十年代中期开始,由于对复杂计算系统的分布化、互连化、开放化、效率和鲁棒性等问题的关心,直接推动多 Agent 系统(MAS, Multi-Agent System)的研究。在 MAS 中为达到共同目标,各 Agent 之间需要相互通信,合作,协调,协商^[2]。当前研究表明信念修正理论是智能 Agent,特别是 BDI 模型的重要逻辑基础之一。但由于多 Agent 系统的复杂性,多 Agent 信念修正(MABR, Multi-agent belief revision)不仅要求每个 Agent 维持自身信念的一致性,而且各 Agent 之间要采取某种策略协作完成全局信念的修正,达到全局一致性。

2. 多 Agent 信念修正

2.1 AGM

八十年代中期,Gärdenfors、Alchourron 和 Makinson 等人共同创立了信念修正理论,用于刻画知识与信念动态演化规律。在该理论中,知识用一阶逻辑语言描述,知识库被表示为一个协调的、逻辑封闭的语句集(通常称之为信念集)^[3]。

在 AGM 框架中定义三个理性原则,即一致性原则(AGM1: Consistency),最小改变原则(AGM2: Minimal Change)和来消息的优先原则(AGM3: Priority to the Incoming Information),并基于三个理性原则定义若干假设以及扩充、约减和修正算子^[7]。

扩充:将语句 a 加入到信念集 K 中,即 $Cn(K \cup \{a\})$,用 $K+a$ 表示。

约减:将语句 a 从信念集 K 中删除,使删除后的语句集仍然是闭包,用 $K-a$ 表示。

修正:若语句 a 与信念集 K 不协调,当 a 加入到 K 中时,需要删除 K 中的一些语句,使得修改后的语句集与 a 协调且为闭包,用 $K \oplus a$ 表示。

三种算子中修正算子处于核心地位。对于修正算子, $\oplus$ AGM 给出了八条理性假设(rationality postulate)^[2]:

($\oplus$ 1) $K \oplus a$ 是一个信念集

($\oplus$ 2) $a \in K \oplus a$

($\oplus$ 3) $K \oplus a \subseteq K+a$

($\oplus$ 4) if $\neg a \in K$, then $K+a \subseteq K \oplus a$

($\oplus$ 5) $K \oplus a = Cn(\perp)$ only if $\vdash \neg a$

($\oplus$ 6) if $a \leftrightarrow b$ then $K \oplus a = K \oplus b$

($\oplus$ 7) $K \oplus (a \wedge b) \subseteq (K \oplus a) + b$

($\oplus$ 8) if $\neg b \in K \oplus a$ then $(K \oplus a) + b \subseteq K \oplus (a \wedge b)$

AGM 框架作为标准信念修正结构被广泛接受,但是它只能用于限定单 Agent 的信念修正行为,而不足以限定令人满意的多 Agent 信念修正中的算子。然而由于经典信念修正理论在描述知识的动态特征和持久性方面具有独特的优势,因此近年来得到广泛研究,并且在形式化和有效、近似计算等方面得到进一步发展。从目前的研究看,对于单 Agent 信念修正理论的进一步研究主要采用两种技术路线:一是符号逻辑的方法;另一种是数值方法^[7]。

2.2 多 Agent 信念修正

随着多 Agent 系统研究的发展,对多 Agent 信念修正研究产生了三种主要模型:MSBR (BR using information from Multiple Sources), DBR (Distributed Belief Revision), MABR (Multi-Agent Belief Revision)^[4]。MSBR 主要解决在多 Agent 系统中,单 Agent 的信息可能来自于多个可能冲突的信息源;而 DBR 是 MSBR 在多 Agent 系统中的一种扩展;MABR 则强调每个 Agent 的信念修正是为了实现多 Agent 系统的共同目标。各种信念修正的层次结构如图1描述。

^{*})本文的研究得到国家自然科学基金资助(编号:60103012)。孙召春 硕士研究生,高 阳 讲师,主要研究领域:分布式人工智能、强化学习。贾松茂 硕士研究生,陈世福 教授,博士生导师,主要研究领域:机器学习、分布式人工智能、数据挖掘。

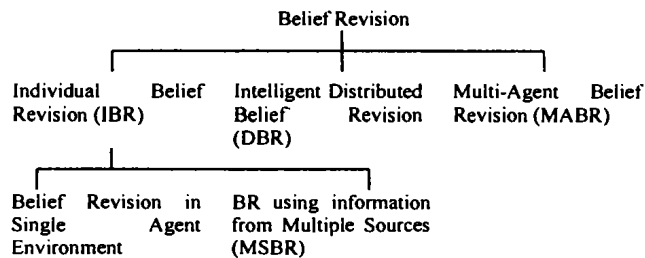


图1 信念修正的层次结构

2.2.1 MSBR MSBR 中仍然执行的是个体信念修正行为^[4]。但单 Agent 信念修正受到三个约束是：修正后信念集的最大一致性，消息的可信度，消息源的可靠性。最大一致性可以由 ATM 系统实现。可信度和可靠性用数值方法实现^[8]。

每个信念可以用 ATMS 节点来表示： $\langle \text{identifier, belief, source, origin set} \rangle$ ；如果在不同节点中存在互为相反的信念，就构成一对矛盾节点。形式如 $\langle \neg, a, -, - \rangle, \langle -, \neg a, -, - \rangle$ 。

在 MSBR 中引入两个数据结构：可靠性集和消息集。令 $S = \{S_1, \dots, S_n\}$ 是消息源集合， $I = \{I_1, \dots, I_m\}$ 是消息集。则定义

可靠性集 = $\{ \langle S_1, R_1 \rangle, \dots, \langle S_n, R_n \rangle \}$ ， R_i 是 0—1 之间的实数，表示 S_i 的可靠性。

消息集 = $\{ \langle I_1, Bel_1 \rangle, \dots, \langle I_m, Bel_m \rangle \}$ ， Bel_i 是 0—1 之间的实数，表示 I_i 的可信度。

另外还有一个集合 $\{ \langle S_1, s_1 \rangle, \dots, \langle S_n, s_n \rangle \}$ ， s_i 是消息源 S_i 给出的消息 I 的子集。可靠性集和该集合共同作为信念函数的输入，消息集是信念函数的输出。

目前解决多消息源冲突的主要方法有 Cantwell 提出的根据消息源的可信度对它们排序，Benferhat 采用的概率逻辑的数据融合，Liberatore 和 Schaefer 采用的仲裁方法等等。从本质上来讲，MSBR 并不是一个分布式方法，因为来自各消息源的数据集成修正是在一个外部模块集中进行的。

2.2.2 DBR 在 DBR 中节点交互以完成任务，因此需要至少两个模块：通信模块，用来选择通信的接收者、讨论者和通讯时间；MSBR 模块。由于节点可能在某种程度上不诚实，因此部分网络中传递的消息可能是不正确的，由此会引起某些节点知识库的矛盾。为了解决这些矛盾，每个节点都装上 MSBR 模块，使得它能够区分消息的可信度和消息源的可靠性。通过比较自己和其他节点的观点，每个节点对自己的系统状态进行改造，其他节点的观点对该节点的影响通过 MSBR 模型采用的规则来实现。

在知识结构上，局部知识和全局知识是分开的。为了抽出全局知识并且保持非集中性，DBR 采用了某种选择形式，即节点组通过交互分布式地选择它们认为应该返回给外部环境的全局输出，而不是由外部管理者来实现。有多种方法进行选择，在极端情况下的两种不同的方法是：

数据驱动：候选者是消息，只有优胜者才能成为全局输出的一部分。

节点驱动：候选者是节点，优胜者是网络中的节点，由优胜者负责组织全局输出。

其他很多方法可以看作是这两种方法的组合。

DBR 和 MSBR 其实采用了同样的机制，但 DBR 的数据集成修正是以分布的方法完成的^[4]。

2.2.3 MABR MABR 强调从系统的角度，每个 Agent 的信念修正为了实现系统的共同目标^[4]，因此 MABR 也可

以叫做智能分布式信念修正 (IDBR)。其与 DBR 的区别是，在 DBR 中存在一个总控模块，将各个节点的信念综合后输出，而 MABR 强调各个 Agent 自身信念的一致性和系统信念的一致性。

Van der Meyden 提出的相互信念修正是一种在同步 MAS 中的语义信念修正理论。相互信念 (mutual belief) 采用理想反省逻辑 (perfect introspection logic) K45n，本质上是一个嵌套过程，在此过程中 Agent 不仅要修正自己关于世界的信念，而且也要修正自己关于其他 Agent 对世界的信念和自己关于其他 Agent 对它的信念，等等。相互信念修正理论的算子满足四个假设^[5]：

- 1) Agent 采用理想反省但可能不一致
- 2) Agent 由于一个常识事件同步修正自己的信念
- 3) 世界是静态的
- 4) 每个 Agent 的修正方法是常识

基于以上四个假设，Agent 可以完全访问其他 Agent 的知识库。当接受了一个新知识，Agent 修正自己的信念，由于修正方法是常识，因此它可以预测其他 Agent 信念的改变。实际上，相互信念修正可以看作是一组单独的信念修正过程一律以并行但分布的风格执行^[4]。相互信念只能处理稳定的同构多 Agent 系统。该模型中的个体信念修正并未满足 AGM 条件，因为它着重于相互性而不是合理性。

Kfir-dahav 等人研究的 MABR 基于异构系统，定义 Agent 的知识库的私有信念域 (PD_i , Private Domain) 和共享信念域 (SD , Shared Domain) 以使每个 Agent 有私有信念同时也有和其他 Agent 共享的信念。一个 Agent 的 PD_i 对其他 Agent 不透明， SD 是每个 Agent 知识库 (Knowledge Base) 的交集。在这种知识结构下，每个 Agent 有自己看世界的角度，同时需要在共享域上协调一致。

因此异构系统的信念修正是两阶段的过程^[6]，第一个阶段是单 Agent 信念修正的过程，如果新的观察包含了关于 $Agent_i$ 的信念集的信息， $Agent_i$ 进行个体信念修正，如果新观察是关于共享域的，那么所有 Agent 都要进行个体信念修正。第二阶段，所有 Agent 协调它们的信念，形成新的 PD_i 和 SD 。

我们认为，MSBR 研究多消息源下的单 Agent 的信念修正行为，而 MABR 研究整个团队的信念修正行为。因此 MSBR 是 MABR 必要的组成成分，但 MABR 与 MSBR 研究的最终目标是不同的。多 Agent 信念修正理论还存在很多亟待解决的问题，如 MAS 中信念的内容和时间没有显式的关系，因此 AGM3 来消息的优先原则在多 Agent 信念修正理论中并不完全适用，并不适用于动态、复杂多 Agent 系统；再如多 Agent 信念修正还没有统一的理论框架，信念修正的可计算问题中构造性成分较大，逻辑方法和数值方法还不能很好地结合在一起，并不适用于动态实时多 Agent 系统中。

3. 基于对策模型的多 Agent 信念修正

我们期望在自利、异构、半竞争的多 Agent 系统中，通过选择单 Agent 的信念修正的不同策略，而实现多 Agent 系统的共同目标。多 Agent 系统采用对策论 (Game theory) 建模，并以此建立多 Agent 信念修正算法。

3.1 对策论建模

在对策模型中每个 Agent 假设都是理性的，试图获得其最大收益；但 Agent 间存在利益冲突。因此我们期望 Agent

间通过合作或协商来消除矛盾、寻求系统的最大收益。首先给出对策论模型定义。

对策论模型由四元组 $\langle A, S, P, N \rangle$ 构造,其中:

A 是 MABR 系统中的 Agent 的集合, $A = \{Agent_1, Agent_2, \dots, Agent_i, \dots, Agent_n\}$, 其中 $i = 1, 2, \dots, n$

S 是每个 Agent 的策略的有限集, $S = \{S_1(i), \dots, S_j(i), \dots, S_m(i)\}$, 其中 $i = 1, 2, \dots, n, j = 1, 2, \dots, m$

P 是 Agent 的赢利函数, $P(S_i(i))$ 表示 $Agent_i$ 采用策略 j 时所获得的赢利值。

N 是协商过程, 采用对策论中求 Nash 平衡的方法来求得最优策略。

Gärdenfors 和 Makinson 认为, 在信念修正的过程中, 决定放弃或保留哪些旧知识的关键因素是 Agent 对旧知识体系的认知信度 (Epistemic Entrenchment)。每个 Agent 都具有对其拥有的知识的认知信度的观念, 信念修正时优先考虑保留认知信度较高的信念, 放弃认知信度较低的信念。对单个 Agent 来说, 考虑认知信度就足够了, 但是在多 Agent 系统中, 各 Agent 通过合作、协调来完成任务, 信念修正时除了考虑自身信念的认知信度, 还要考虑自己的信度对其他 Agent 的影响, 比如在 $Agent_i$ 信度最高的情况下, 可能整个系统的信度反而会降低。但是, 认知信度仍然是衡量系统性能的一个重要因素, 因此我们用认知信度作为赢利值来构造赢得矩阵。并利用 Nash 平衡求解系统的最大赢利值。

3.2 基于对策模型的多 Agent 信念修正

基于上述模型, 考虑只有两个 Agent 的系统中, 当系统加入新知识后多 Agent 信念修正系统的工作过程。下面给出基于对策模型的多 Agent 信念修正算法。

算法:

```
//  $S_i$  为每个 Agent 策略集中的策略
// 求  $Agent_i$  对于策略  $S_i$  的赢利值
For each  $Agent_i$  do
   $\forall j S_i \rightarrow P(S_i(i))$ 
End For
// 以每个 Agent 对于各策略的赢利值构造对策矩阵
 $\forall i, j S_i(1), S_i(2) \rightarrow M(i, j) = (P(S_i(1)), P(S_i(2)))$ 
// 用协商过程求每个策略对下整个系统的赢利值
For each  $M(i, j)$  do
   $N(M(i, j))$ 
End For
// 选择使系统赢利值最大的策略对
 $\exists c, d N(M(c, d))$  is the largest one  $\rightarrow S_c(1); S_d(2)$ 
// 按照该策略对各 Agent 进行信念修正
```

根据上述算法, 假设在一个具体的系统中, $A = \{A_1, A_2\}$, A_1, A_2 相互独立, 即各自的信念对另一 agent 的认知信度没有影响, 信念修正的策略采用标准调整和最大调整。标准调整 (T) 是指加入一条新知识后, 只保留能推导出的知识。最大调整 (M) 是指加入新知识后, 只去掉推出矛盾的信念。求 $Agent_i$ 的认知信度有很多方法, 在各 Agent 互不影响的情况下, 可以用 $Agent_i$ 中每条信念的认知信度的加权平均来求得 $Agent_i$ 的认知信度。

A_1 中的信念集为:

0.9 $a|b$

0.8 $a|d$

0.7 $-c \vdash b \quad c \quad d \quad e$

A_2 中的信念集为:

0.92 $a|c$

0.78 $b \vdash f$

0.6 $-c|e \quad f \quad e$

信念前面的小数表示 Agent 对各信念的认知信度。

假设在状态 1 时, A_1, A_2 均加入一条新知识: $-a$, A_1 用 T 策略得到: $a|b, a|d, d$, 认知信度为: 0.8; 用策略 M 得到: $a|b, a|d, d, e$, 认知信度为 0.78。

同理, A_2 用策略 T 修正得认知信度 0.71, 用 M 策略得认知信度 0.7。

赢得矩阵为:

		A_2	
		T	M
A_1	T	(0.8, 0.71)	(0.8, 0.7)
	M	(0.78, 0.71)	(0.78, 0.7)

用协商机制 N 求出使整个系统的认知信度最优的策略组合 (T, T), 系统按照该策略组合进行信念修正, 到达状态 2。这个例子考虑的是 Agent 互不影响的情况, 每个 Agent 的优化策略的组合就是系统的最优策略。实际情况中, 多 Agent 系统中的 Agent 往往不是相互独立的, 它们之间的信念冲突可能会为认知信度带来影响, 从而影响策略对的选择。

结论 本文首先介绍了经典信念修正理论, 其次深入讨论了多 Agent 信念修正的研究的三个主要模型: DBR、MSBR 和 MABR, 并分析多 Agent 信念修正研究存在的问题; 最后根据这些问题, 提出基于对策论模型的多 Agent 信念修正算法。与其他方法相比, 该算法可以实现异构系统下的信念修正。

但目前只假设多 Agent 系统处于单个离散状态, 而本质上信念修正是一个连续动态的过程。因此前一个状态的系统的认知信度会影响后一个状态系统的认知信度。进一步研究将考虑马尔科夫对策模型, 并采用学习技术, 以实现动态、复杂多 Agent 系统的信念修正。

对多 Agent 信念修正理论深刻的研究将直接推动电子拍卖市场、法庭推理等智能应用的发展。因此多 Agent 信念修正理论是当前 Agent、多 Agent 技术、知识推理和符号逻辑中前沿性的课题, 具有重要的理论价值和应用意义。

参考文献

- 1 de Kleer J. An Assumption-based TMS. *Artificial Intelligence*, 1992, 28(2): 127~162
- 2 Nebel, Base Revision Operations and Schemes: Semantics, Representation, and Complexity. In: Cohn A. G, ed. *Proc. 11th European Conf. on Artificial Intelligence*, 1994
- 3 Jennings N R, Sycara K, Wooldridge M. A Roadmap of Agent Research and Development. *Autonomous Agents and Multi-Agent System*, 1998(1): 275~306
- 4 Liu W, Williams M A. A framework for multi-agent belief revision (part i: The role of ontology). In: Foo N. ed., *12th Australian Joint Conference on Artificial Intelligence, Lecture Notes in Artificial Intelligence*. Sydney, Australia: Springer-Verlag, 1999. 168~179
- 5 van der Meyden R. Mutual Belief Revision. In: *Proc. of the Intl. Conf. on Principles of Knowledge Representation and Reasoning*, Bonn, May 1994. 595~606
- 6 Kfir-dahav N E, Tennenholtz M. Multi-agent belief revision. In: *6th Conf. on Theoretical Aspects to Rationality and Knowledge*, 1996. 175~194
- 7 梁尚敏, 李未. 信念修正的各种方法之比较. *计算机科学*, 1999, 26
- 8 Dragoni A F, Giorgini P, Baffetti M. Distributed Belief Revision vs. Belief Revision in a Multi-Agent environment: first results of a simulation experiment
- 9 Nebel B. How Hard is it to Revise a Belief Base, 1996. 8