

一种新的类 BAN 逻辑模态语义模型^{*})

——兼论类 BAN 逻辑的语法缺陷

谢鸿波 周明天

(电子科技大学-卫士信息安全联合实验室 成都 610054)

摘要 由于类 BAN 逻辑缺乏明确而清晰的语义,其语法规则和推理的正确性就受到了质疑。本文定义了安全协议的计算模型,在此基础上定义了符合模态逻辑的类 BAN 逻辑“可能世界”语义模型,并从语义的角度证明了在该模型下的类 BAN 逻辑语法存在的缺陷,同时,指出了建立或改进类 BAN 逻辑的方向。

关键词 类 BAN 逻辑,模态逻辑,形式语义

On Semantics Model of BAN-like Logic and Flaws in its Syntax Rules

XIE Hong-Bo ZHOU Ming-Tian

(Information Security United Lab. of UESTC—WESTONE, Chengdu 610054)

Abstract For lack of explicit and definite semantics in BAN like logics, their correctness of syntax rules and reasoning is under suspicion. In this paper, a computing model for security protocols is defined. Based on it, semantics model about possible word of BAN like logic is defined conforming with traditional modal logic. And from the viewpoint of semantics, some syntax flaws in BAN-like logic are proved. Also, how to construct a new logic or improve BAN-like logics is pointed out.

Keywords BAN-like logic, Modal logic, Formal semantics

1 引言

自从 20 世纪 80 年代末 Burrow, Abadi, Needham 提出用 BAN 逻辑^[1]来分析安全协议的安全性之后,基于逻辑和信念的形式化分析就成为安全协议安全性分析的焦点。由于 BAN 逻辑的不完善,许多学者在 BAN 逻辑的基础上提出了自己的验证逻辑,如 AT 逻辑^[2], SVO 逻辑^[3], GNY 逻辑^[4], MB 逻辑^[5]等。这些逻辑被统称为类 BAN 逻辑。

类 BAN 逻辑对安全协议的分析是针对协议消息展开的,它根据协议参与者在完成协议动作(发送/接收消息)后产生的信任关系,来判断安全协议是否能达到既定目的。类 BAN 逻辑确实发现了一些协议以前没有发现的漏洞^[1~5],然而它仍然是不完善的逻辑。Mao Wenbo^[7]等从分析步骤、逻辑语法等方面系统地分析了类 BAN 逻辑四个主要缺陷。Nesset^[6]指出 BAN 逻辑有致命缺陷:“它(BAN 逻辑)解决了谁得到和确认一个密钥的问题(即:认证密钥),但没能解决谁该得到密钥(即:密钥机密性)”。类 BAN 逻辑的形成正是 BAN 逻辑不断完善的过程,一方面是消除逻辑语法上不合理或不正确的公理或逻辑推理规则,另一方面试图给出更严格的逻辑语义以证明类 BAN 逻辑语法自身的正确性和可靠性。

要完善类 BAN 逻辑并使之更好地分析安全协议的安全性,就必须分析它的缺陷以及产生这些缺陷的原因。目前,这方面的研究主要是在分析类 BAN 逻辑的语法结构上,包括公理和初始化步骤等。由于类 BAN 逻辑本质上是一种模态逻辑,如果不从模态逻辑语义模型上分析类 BAN 逻辑,就很

难完善类 BAN 逻辑或者指出类 BAN 逻辑完善的方向。

本文的主要工作是,在 SVO 等人工作的基础上,给出类 BAN 逻辑完整而明确的模态逻辑语义模型,定义了类 BAN 逻辑的“可能世界”及“可达关系”,并在这样的模态框架下,给出了类 BAN 逻辑的赋值函数,通过该模型框架,本文从类 BAN 逻辑的本质——模态逻辑的角度分析了类 BAN 逻辑的语法缺陷。本文没有给出新的类 BAN 逻辑语法,但指明了语法结构完善的方向。

本文第 2 节定义类 BAN 逻辑的模态逻辑语义模型,第 3 节在本文定义的模态逻辑语义模型下分析现有类 BAN 逻辑的语法缺陷,最后小结本文的工作,并指出类 BAN 逻辑完善的方向。

2 类 BAN 逻辑的语义模型

BAN 逻辑提出后,因其缺乏语义证明而受到许多学者的批评^[6,7]。类 BAN 逻辑被相继提出^[2~5],其中 AT 逻辑是最早定义类 BAN 逻辑语义模型的,而 SVO 逻辑被认为是这方面最有代表性的工作。

在 AT 逻辑中,用 (r, k) 表示可能世界(其中 r 是协议的一轮运行, k 是这轮运行中的第 k 个时刻),并在 (r, k) 下定义了基本结构的真值条件。AT 逻辑用 hide 操作符对“不可读”的加密消息做了处理,通过迭代算法计算出“好的”运行集合 G ,并以此定义了可能世界之间的关系 $R: (r, k) \sim_i (r', k')$ 。在可能世界 (r, k) 及其关系 R 下,AT 逻辑定义了“信任”的真值条件,从而给出了类 BAN 逻辑的模态语义模型。

SVO 逻辑所定义的可能世界和基本结构的真值条件与

^{*}) 本文受电子科学基金 514500101DZ02 课题资助。谢鸿波 博士研究生,感兴趣的领域是网络与信息系统安全、协议分析、电子商务、电子政务等。周明天 博士生导师,主要研究领域为计算机网络、分布对象技术、并行分布处理和网络与信息系统安全。

AT 逻辑一样。但它强调了消息的可理解性,并通过 Comprehension 操作符来计算可理解的消息。与 AT 逻辑不同的是,它是通过本地状态 $r_i(t)$ 和可理解的消息来定义可能世界的关系的。SVO 逻辑还简单地证明了该语义模型的可靠性,使得 SVO 逻辑语法的正确性有了保证,SVO 逻辑更符合模态逻辑。

然而,无论是 AT 逻辑还是 SVO 逻辑都没有对类 BAN 逻辑定义清晰、完整的模态逻辑语义模型,并在这个模型下建立有效的逻辑语言。由于其可能世界及其关系的定义依赖某些严格的限制^[2,3](比如“好的”运行、恒为真的初始信任等),且没有确定意义的真值函数定义,因此,他们定义的模态逻辑语义模型自身存在着某些重要缺陷,其逻辑语法不能很好地分析安全协议的安全属性,也没有用该模型去建立有效的逻辑语言。

本文定义一个清晰、完整,并符合模态逻辑的类 BAN 逻辑语义模型。

2.1 计算模型

安全协议是由一组有限的实体($P_1, P_2, \dots, P_n$)完成的。这些实体在协议规范下产生相应的动作,在协议结束时得到所需要的安全属性。设实体 P_i 可以完成的动作集合为 $\Phi = \{send(m), receive(), newdate()\}$,其中 $send(m)$ 表示 P_i 按协议规范发送了一个消息, $receive()$ 表示 P_i 接收到一个符合协议规范的消息,并将该消息加入到本地消息集合 $\chi = \{m_i | i = 1, 2, \dots, k\}$ 中,其中 i 表示本轮运行的第 i 个时刻, $newdate()$ 表示 P_i 按协议规范产生新的原始数据(一个完整的消息或消息的一部分)。在协议的一轮运行中, P_i 完成的动作组成了它的本地动作序列集合 $\zeta = \{a_i | i = 1, 2, \dots, k, a \in \Phi\}$,其中 i 表示本轮运行的第 i 个时刻。

协议的一轮运行实际上是协议的一个运行实例。它由一系列状态组成,每个状态的发生时间都以一个整数时刻 i 表示,其中运行的初始状态被赋予一个时刻 $k, k \leq 0$,第 i 个时刻被赋予值 $k + (i - 1)$ 。文中用全局状态的有限集合 $S = \{S(i) | i = 1, 2, \dots, k\}$ 表示协议的一轮运行,其中 $S(i)$ 表示该轮运行的第 i 个时刻的全局状态,它由该时刻的实体的本地状态多元组 $(s_1, s_2, \dots, s_n)$ 构成。在本文中,第 i 个时刻的本地状态由本地消息集合和本地动作序列集合的二元组表示: $s_{i-k} = (\zeta, \chi)$ 。在安全协议中,实体 P_i 的安全属性由其本地状态集合 $S_i(r) = \{s_{i-k}(r) | k = 1, 2, \dots, N\}$ 组成的全局状态域计算得出。

由于 P_i 是在本地消息集合与本地动作序列集合上计算安全属性的,而本地动作序列集合根据协议规范可以很明确地确定,因此, P_i 对本地消息集合中的消息的理解就成了计算其安全属性最关键的问题了。本文用一个 6 元组来定义消息的消息属性, $M_p = (P_n, Sid, Mid, S, R, Mc)$,其中 P_n 表示协议名字, Rid 表示协议的某轮运行, Mid 表示消息号, S 表示消息的发送者, R 表示消息预期接收者, Mc 表示消息内容。从消息属性的定义可以看出,消息属性可以唯一确定某个协议某轮运行的某个消息。

P_i 对本地消息的理解其实就是对该消息的消息属性的理解。一般情况下, P_i 可以从协议规范中得到 P_n, Mid, S, R 的值,而 Rid 一般不能从协议规范中得到。实际上对消息属性中这些元素的理解一般都来自协议消息,即 Mc 。因为通常的安全协议中不会明确地指出协议运行轮次号、消息号、发送者、接收者等属性的值,这些属性都是隐含地包含在协议消息

中,所以对协议消息内容的理解是至关重要的。下面给出本文“可理解”消息集合的计算方法。

用 $Comp_i(r, i)$ 表示实体 P_i 在第 r 轮协议运行到第 i 个时刻的可理解消息集合。集合 $Comp_i(r, i)$ 的计算是一个迭代过程, $Comp_i(r, 0)$ 包含 P_i 在第 r 轮运行开始时可理解的所有明文消息。 $Comp(M)$ 表示由消息 M 得到的“可理解”消息集合,如果 P_i 在第 r 轮运行的第 i 个时刻收到消息 M ,那么 $Comp_i(r, i) = Comp_i(r, i - 1) \cup Comp(M)$ 。因此, $Comp_i(r, i)$ 的计算变成计算第 i 个时刻的 $Comp(M)$ 。

一个消息 M 可以看成是若干个子消息集合的集合。消息可以分为三类:

1. C : 表示明文子消息集合,如:常量、用户名等。
2. E : 经过密码算法处理过的消息集合,如:加密计算 $\{.\}_k$, 哈希计算 $H(.)$ 等。
3. S : 与共享秘密结合的消息集合,如 $\langle . \rangle_s$: 只有知道了 s 才能确定括号中的内容。

在定义了三类子消息后,计算 $Comp(M)$ 的步骤如下:

1. 令 $M' = M$
 2. 将消息 M' 根据定义分为 3 类,即 $M' = \{C, E_1, \dots, E_k, S_1, \dots, S_n\}$
 3. 计算对 C 的理解
for all $c \in C$
if P_i 不能理解 c
then 用 $*$ 表示 c
end for
 $M' = M' - C$
 4. 计算对 E 的理解
for $i = 1$ to k
if P_i 对 E_i 不能进行反密码计算
then 用 $*$ 表示 E_i and $M' = M' - E_i$
else 将 P_i 对 E_i 进行反密码计算的结果放入 M'
end for
 5. 计算对 S 的理解
for $i = 1$ to n
if P_i 不知道 S_i 中的共享秘密
then 用 $*$ 表示 S_i and $M' = M' - S_i$
else 将 P_i 对 S_i 分离共享秘密后的结果放入 M'
end for
 6. 如果 $M' \neq \phi$, 则返回 2, 重复执行 2~6; 否则结束。
- 完成以上步骤之后,消息 M 变为 $M(*)$, 即: $Comp(M) = M(*)$ 。

以上定义了类 BAN 逻辑的计算模型,下面将在这个计算模型的基础上,定义类 BAN 逻辑的语义模型。

2.2 语义模型构造

类 BAN 逻辑系统自身的正确性是需要通过形式化的语义分析的。Syverson 指出^[3],逻辑语义的重要作用在于为评估形式推理系统提供了一个可靠的方法。逻辑推论是命题真值的推理过程,而形式推理则是符号的演变过程。由于符号是没有真值意义的,因此必须将逻辑推论与形式推理对应起来,要做到这一点必须关注两个问题:

- 可靠性,即: $\Gamma \vdash A \Rightarrow \Gamma \models A$, 如果符号集合可以推导出公式 A , 那么在论域 Γ 中命题是有效的(即为真)。
- 完备性,即: $\Gamma \models A \Rightarrow \Gamma \vdash A$, 如果在论域 Γ 中命题 A 是有效的(即为真),那么可以从符号集合 Γ 推导出公式 A 。

对于类 BAN 逻辑来说,“可靠性”是最重要的,因为类 BAN 逻辑其实就是符号的形式推理逻辑,而它的推理结论要成为有意义的命题,就必须满足“可靠性”。而逻辑语义可用于证明形式逻辑的可靠性和完备性。

对于类 BAN 逻辑,本文定义了如下的语义模型。

定义 1 二元组 $\langle W, R \rangle$ 是一个框架, 当且仅当 W 是可能世界的非空集合, R 是可能世界上的任意二元关系 $R \in W \times W$ 。

定义 2 三元组 $\langle W, R, V \rangle$ 是一个模态语义模型, 当且仅当二元组 $\langle W, R \rangle$ 是一个框架, V 是在这个框架上的真值指派。

由定义 1、2 可以看出, 一个给定的框架, 其固定的是可能世界集合和这些可能世界之间的可达关系, 而赋值函数 (即: 真值指派) 决定了在每个可能世界中能得到怎样的事实。

基于 2.1 节的计算模型, 类 BAN 逻辑的语义模型可以按照如下的方法来构建。

在类 BAN 逻辑中, 协议参与者 P_i 的可能世界就是该协议的某一轮运行, 因此, P_i 在该轮运行的全局状态 $S_i(r)$ 就是它的一个可能世界, 用集合 $W_p = \{S_i(r) | r \in \mathcal{R}\}$ 来表示 P_i 的可能世界集合, 显然它是个非空集合。对于这个可能世界集合, 本文定义了这样的可达关系。

定义 3 设 $s_{i-k}(r) \in S_i(r), s'_{i-k}(r') \in S_i(r')$, 那么有关系 $R_p: s_{i-k}(r) \sim s'_{i-k}(r')$, 当且仅当 $s'_{i-k}(r'), s_{i-k} \in \text{CompS}(S_i)$, 其中 $\text{CompS}(S_i)$ 是 P_i 可理解状态集合。

定义 4 $\text{CompS}(S_i)$ 是 P_i 可理解状态集合, $\forall (s_{i-k}(r) = (\zeta, \chi)) \in \text{CompS}(S_i)$, 当且仅当:

1. ζ 符合协议规范; 2. $\chi \in \text{Comp}_i(r, i)$ 。

由此, 本文定义了类 BAN 逻辑的语义框架 $\langle W, R \rangle$ 。对于实体 P_i 而言, 语义框架为 $\langle W_p, R_p \rangle$, 可能世界为 W_p 是由 P_i 的协议运行轮次的全局状态构成的, 因此, 由 P_i 参与的该协议的所有运行情况都包含在这个可能世界集合中, 而可达关系 R_p 是建立在能被 P_i 所理解的状态之上的, 因为不被 P_i 理解的状态将不被 P_i 接受, 也就不能对 P_i 的信任关系产生影响, 而类 BAN 逻辑就是对信任的推导。可见, 这样的框架是符合类 BAN 逻辑的。由关系 R_p 的定义可以看出这是一个等价关系, 图 1 表示了其可能世界的这种关系。

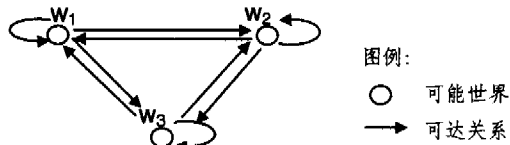


图 1 可能世界及其可达关系

这是个典型的 S_5 模态逻辑系统^[8]。在这样的模态系统中, 命题“必然 A , 那么 A ”必为真, 即: $\models \Box A \rightarrow A$ 。这就涉及到语义模型中赋值函数的定义, 它决定了命题在框架中的真值条件。下面, 本文给出类 BAN 逻辑的赋值函数定义。

定义 5 设 $V_{S_i, P}(\phi)$ 表示对于实体 P_i 来说, 命题 ϕ 在可能世界 $S_i(r)$ 中的真值。其中 $V_{w, P}(\phi) = 1$, 当且仅当 $\forall S'_i(r) \in W_p, S'_i(r) R_p S_i(r), \models_{S'_i} \phi$, 否则 $V_{w, P}(\phi) = 0$ 。

根据定义 5, 命题 A 在可能世界 $S_i(r)$ 为真的条件是: 在 P_i 的可能世界集合 W_p 中, 对于 $S_i(r)$, 在所有满足可达关系 R_p 的可能世界 $S'_i(r)$ 里, 命题 A 都为真, 否则命题 A 为假。

这样的真值定义是符合安全协议的实际运用环境的。安全协议实际上是运行在一个复杂而不安全的环境中, 除了严格地按照协议规范运行协议这样的可能世界外, 还有重放攻击、假冒攻击等情况下产生的可能世界。“实体 P_i 信任 A ”这样的合式命题在没有严格真值定义的情况下, 在所有可能世界中的真值并不总是有效的 (即为真)。如果这是安全协议运行的结果, 那么这就是协议被攻击的地方。如果它是类 BAN

逻辑的初始信任, 那么由不可靠的初始信任经过形式推理得到的结论也是不可靠的, 这就是为什么有些安全协议被类 BAN 逻辑验证之后仍然被发现有协议漏洞的原因^[7]。根据定义 5 给出的真值定义, 可以保证对于满足该真值条件的合式命题在实体 P_i 的模态语义框架 $\langle W_p, R_p \rangle$ 中都为真, 从而保证了逻辑推理的正确性。

至此, 本文定义了一个完整、清晰, 且符合模态逻辑的语义模型 $\langle W, R, V \rangle$ 。下一节, 将在本文的语义模型框架下, 分析类 BAN 逻辑的语法缺陷。

3 类 BAN 逻辑的语法缺陷

类 BAN 逻辑一般包括三部分: 基本结构、初始命题、推理规则。其基本结构定义了逻辑的原子公式, 包括诸如: “实体 P 相信 X ”、“实体 P 看到 X ”、“实体 P 说过 X ”、“ K 是实体 P 和实体 Q 共享密钥”等。初始命题是类 BAN 逻辑作者自定义的为真的公式, 其形式一般为: $\frac{A}{B}$, 即: 如果 A , 那么 B 。推理规则有两个, 一个是分离规则: 如果 $\vdash A \rightarrow B, \vdash A$, 那么 $\vdash B$; 一个是必然规则: 如果 $\vdash A$, 那么 $\vdash P$ 相信 A 。使用类 BAN 逻辑分析协议的过程为:

1. 确定初始假设集合 (即: 初始信任);
2. 用基本结构形式化协议消息 (即: 理想化协议);
3. 根据初始命题确定每个协议步骤之后产生的新的协议运行命题集合;
4. 对“初始假设集合”, “初始命题”, “协议运行命题集合”使用推理规则得到实体的信任集合, 并比较该信任集合与协议目标是否一致。

其推理模型如图 2 所示。

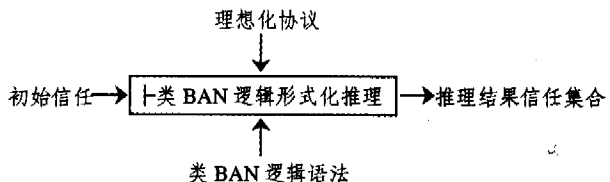


图 2 类 BAN 逻辑基本模型

类 BAN 逻辑的形式化推理结论一般是形如这样的公式: “实体 P 信任 A ”。在类 BAN 逻辑中, 它的含义为: 如果 P 信任 A , 那么 A 为真。类 BAN 逻辑得到这样的结论是在一定条件下的, 即: 在所有实体都正确按照协议规范执行协议步骤, 并且没有攻击者参与的可能世界里。这样的条件限制显然不符合安全协议实际的运行环境, 从而导致了“ P 信任 A , 而 A 实际并不为真”的情况出现, 使得类 BAN 逻辑的推论结论并不总是可靠。

本文把这样的情况称为“信任真实性”问题。即: 类 BAN 逻辑的形式化推理得到的信任, 并不是真实的信任。这种语法上的缺陷是因为类 BAN 逻辑缺乏完整的语义模型, 没有明确的赋值函数定义, 使得命题有效成立的真值条件在类 BAN 逻辑中是比较模糊的。推理结论的信任不真实性, 其实只是一个表象, 真正的原因是逻辑语法中的“初始命题”并不是构建在语义模型基础之上的, 因此很多“初始命题”本身就在“信任真实性”问题。

在 BAN 逻辑中, 消息意义命题用于判断消息的发送者是谁, 并获得该消息所表达的信任。其命题表示为:

$$\frac{P \models P \stackrel{K}{\leftrightarrow} Q, P \triangleleft \{X\}_K}{P \models Q \mid \sim X}, \text{ 即: } P \text{ 相信 } K \text{ 是 } P \text{ 与 } Q \text{ 共享的密钥,}$$

且 P 看到了用 K 加密的消息 X , 那么 P 相信 Q 说过 X 。根据模态语义中真值条件的函数定义, 当且仅当 $\forall S'_i(r) \in W_p, S'_i(r) R_p S_i(r), \vdash_{S'_i}^w \frac{P \models P \xleftrightarrow{K} Q, P \triangleleft \{X\}_K}{P \models Q \mid \sim X}$, 时, 这个命题为真。现在考察这样一种可能世界 w_1 , 在 w_1 中协议攻击者重放了以前 P 产生的消息 $\{X\}_K$, 显然世界 w_1 满足定义 3 定义的关系 R_p , 所以消息意义命题也应该在世界 w_1 中为真。然而, 在这个可能世界中该命题的前件为真, 而后件为假 (因为 Q 实际上没有说过 $\{X\}_K$, 也就没说过 X), 即:

$V_{w_1,p}(P \models P \xleftrightarrow{K} Q, P \triangleleft \{X\}_K) = 1, V_{w_1,p}(P \models Q \mid \sim X) = 0$, 从而 $V_{w_1,p}(\frac{P \models P \xleftrightarrow{K} Q, P \triangleleft \{X\}_K}{P \models Q \mid \sim X}) = 0$ 。可见, 该命题不满足模态语义逻辑的真值条件定义, 从而产生了不真实的信任。

管辖权命题则用于表达协议中的信任关系。其命题表示为: $\frac{P \models Q \Rightarrow X, P \models Q \mid \sim X}{P \models X}$, 即: P 相信 Q 对于 X 的产生具有可信权威, 且 P 相信 Q 说过 X , 那么 P 相信 X 。根据模态语义中真值条件的函数定义, 当且仅当 $\forall S'_i(r) \in W_Q, S'_i(r) R_Q S_i(r), \vdash_{S'_i}^w \frac{Q \Rightarrow X, Q \mid \sim X}{X}$, 然后由实体 P 对它运用必然规则, 那么这个命题为真。这个推理过程存在定义域的不一致, 第一步的推理是在 Q 的可能世界 W_Q 及其可达关系 R_Q 中进行的, 最终得到的却是 P 的信任关系。也就是说, 在建立这个信任的过程中, 模态语义框架已经由 Q 的框架 $\langle W_Q, R_Q \rangle$ 变成了 P 的框架 $\langle W_p, R_p \rangle$, 因此对 P 而言这并不是个真实的信任。

GNV 逻辑为判定消息是否具有可识别性增加了若干“初始命题”。这些新增的命题同样存在着“信任真实性”的问题。在此, 仅举一例来分析其语法缺陷。在 GNV 逻辑中有这样一个推理规则: R6 规则: $\frac{P \Rightarrow XH(X)}{P \models \phi(X)}$, 即: P 拥有消息 X 的哈希, 那么 P 相信 X 是可识别的。这将导致一个很可笑的逻辑错误。应用 GNV 逻辑的 P1 规则: $\frac{P \triangleleft X}{P \Rightarrow X}$, 即: P 看到 X , 那么 P 拥有 X ; P4 规则: $\frac{P \Rightarrow X}{P \Rightarrow H(X)}$, 即: P 拥有 X , 那么 P 拥有 X 的哈希; 和 R6 规则可以得到公式: $\frac{P \triangleleft X}{P \models \phi(X)}$, 即: P 看到 X , 那么 P 相信 X 是可识别的。造成这个错误的原因是规则 R6 本身不符合 2.1 节中对可理解消息的计算。根据计算, 如果 $X \notin \text{Comp}_i(r, i)$, 那么即使 P 拥有了 X 的哈希, 它也不能理解消息 X , 所以消息 X 对它来说仍然是不可识别的。

SVO 逻辑中消息关联命题: $\frac{P \xleftrightarrow{K} Q, R \triangleright \{X^Q\}_K}{Q \mid \sim X, Q \triangleleft K}$, 即: K 是 P 和 Q 共享密钥, 收到来自 Q 的加密信息, 那么 Q 说过 X , 且 Q 拥有 K 。这个命题与 BAN 逻辑中的消息意义命题相似, 但不同的是这个命题中没有信任关系。消息 $\{X^Q\}_K$ 很特别, 它表明是来自 Q 的加密信息, 这个消息结构中添加了消息来源符号, 然而这并不是协议规定的消息格式, 它来自类 BAN 逻辑的诚实性假设, 即: 拥有密钥 K 的实体可以被认为能够诚实地指出该消息来自谁。这个实体诚实性假设恰恰限制了命题的可能世界集合, 使得该命题在满足关系的所有可能世界里并非都为真, 如同 BAN 逻辑的消息意义命题一样。

其他类 BAN 逻辑都有同样类型的问题, 由于篇幅限制, 本文不一一列举。通过上面简单的分析, 可以清楚地看到由

于缺乏完整、明确的语义模型, 使得类 BAN 逻辑的语法规则不是在具有确定真值条件的赋值函数下定义其法则的真值, 从而导致了“信任真实性”的问题。

缺乏语义模型而导致的语法缺陷, 还存在于基本结构集合和语法推论断言集合中。目前类 BAN 逻辑的基本结构定义得还不完善, 因此, 还无法描述所有的协议消息, 从而限制了它的使用范围, 而有限的断言集合也限制了逻辑的分析能力。类 BAN 逻辑在这方面的改进显得有些艰难, 因为所有基本结构以及断言中的谓词都必须给出真值条件, 否则, 用这些存在“信任真实性”问题的基本结构和断言作为推理的前提和结论, 其语法自身的正确性就会被质疑。可见由于没有具有确定真值条件的赋值函数, 因此缺乏相应的机制来添加新的基本结构和断言谓词, 限制了类 BAN 逻辑的使用范围。

通过上面的分析可以看出, 类 BAN 逻辑的语法缺陷主要是因为目前的类 BAN 逻辑不是在清晰、完整的模态逻辑模型基础上建立的, 它没有明确的真值条件定义。因此, 要解决“信任真实性”问题, 应该以所有可能世界为论域, 在有明确定义的赋值函数下重新计算有效的最小“初始命题”集合, 或者在模态语义模型 $\langle W, R, V \rangle$ 的基础上, 严格地按照其中赋值函数 V 的定义来构建一套新的类 BAN 逻辑。

结论和进一步的工作 类 BAN 逻辑以其简洁、易懂的证明风格引起广泛关注。然而, 由于缺乏语义描述, 其语法的正确性被广大学者质疑。以 AT、SVO 逻辑为代表, 一些学者给类 BAN 逻辑作出模态逻辑语义的解释, 但都没有给出类 BAN 逻辑完整而清晰的模态语义模型, 并在此基础上建立类 BAN 逻辑语法。因此, 类 BAN 逻辑语法及其形式化推理的正确性, 仍然有待于证明。从这个角度上来说, 类 BAN 逻辑还是个很不完善的逻辑。

本文提出了一个合理的协议计算模型, 在此基础上定义了一个完整、清晰, 具有典型模态语义特点, 并符合安全协议实际运行环境的模态逻辑语义。此外, 还用所定义的模态逻辑语义模型分析了目前类 BAN 逻辑的语法缺陷。本文没有在该模态逻辑语义模型的基础上构建一套完整的逻辑语法和形式化推理规则, 但指出了在这个模态逻辑模型下如何改进或建立类 BAN 逻辑语法规则。这是我们下一步要做的工作。

参考文献

- 1 Burrows M, Abadi M, Needham R. A logic of authentication: [Technical Report 39]. Digital Systems Research Center, 1989
- 2 Abadi M, Tuttle M. A semantics for a logic of authentication. In: Proc. of the Tenth ACM Symposium on Principles of Distributed Computing, ACM Press, Aug. 1991. 201~216
- 3 Syverson P F, van Oorschot P C. On unifying some cryptographic protocol logics. In: Proc. of the IEEE Computer Society Symposium on Research in Security and Privacy, 1994. 14~28
- 4 Li Gong, Needham R, Yahalom R. Reasoning about belief in cryptographic protocols. In: Proc. of the 1990 IEEE Computer Society Symposium on Research in Security and Privacy, 1990. 234~248
- 5 Mao W, Boyd C. Towards formal analysis of security protocols. In: Proc. of the Computer Security Foundation Workshop VI, 1993. 147~158
- 6 Nessett D M. A critique of the burrows, abadi, and needham logic. Operating Systems Review, 1990, 24(2): 35~38
- 7 Mao W. An augmentation of of BAN-like logics. In: 8th IEEE Computer Security Foundations Workshop. IEEE Computer Society Press, 1995. 44~55
- 8 陆汝钊.《人工智能》下册. 科学出版社, 1996