

基于 Multi-Agent 的网络谣言传播事件中 信息主体信任识别研究*

滕 婕¹ 夏志杰¹ 罗梦莹² 王筱莉¹

(1. 上海工程技术大学 管理学院 上海 201620; 2. 华东师范大学 工商管理学院 上海 200241)

摘要: [目的/意义] 全媒体背景下,网络舆情的爆发使谣言借助通信渠道快速扩散,群体辟谣效果下产生各种质量层次不齐的信息。为了对信息主体进行信任识别,构建信任识别模型。[方法/过程] 运用 Agent 关系网络,结合 Agent 间的历史交互信息,提出基于 Multi-Agent 的信任识别模型,考虑到个体异质性,引入动态权重因子,将直接信任和间接信任进行组合,计算出目标主体的综合信任值,从中寻找出最为可靠的信息主体。[结果/结论] 网络谣言事件下构建的信任识别模型对恶意信息主体能够有效识别;使用动态权重因子计算的综合信任值敏感性增强;信任判别网络中 Agent 的数量并不影响信任模型的有效性。

关键词: 网络舆情;网络谣言;谣言传播;Multi-Agent;信息主体;信任识别

中图分类号:G206

文献标识码:A

文章编号:1002-1965(2020)03-0105-10

引用格式:滕 婕,夏志杰,罗梦莹,等. 基于 Multi-Agent 的网络谣言传播事件中信息主体信任识别研究[J]. 情报杂志, 2020, 39(3): 105-114.

DOI:10. 3969/j. issn. 1002-1965. 2020. 03. 015

Research on Trust Recognition of Information Subject in Network Rumor Propagation Event Based on Multi-Agent

Teng Jie¹ Xia Zhijie¹ Luo Mengying² Wang Xiaoli¹

(1. School of Management, Shanghai University of Engineering Science, Shanghai 201620;

2. School of Business Administration, East China Normal University, Shanghai 200241)

Abstract: [Purpose/Significance] Under the background of all-media, the outbreak of network public opinion causes rumors to spread rapidly, and all kinds of rumor-refuting information with different quality levels are produced accordingly. In order to judge the trust of the Agent, the trust recognition model is constructed. [Method/Process] A trust recognition model is proposed based on Multi-Agent, which combines the historical information and the relational network of Multi-Agent. This paper considers the influence of individual heterogeneity, introducing dynamic weight factor, combining direct trust with indirect trust. And the comprehensive trust values of target subject are calculated. This model can find out the most reliable information subject. [Result/Conclusion] It finds the following conclusion: the trust recognition model can effectively identify the malicious information Agents; the sensitivity of comprehensive trust value is enhanced using dynamic weight factor; the number of Agent in interpersonal networks does not affect the effectiveness of the trust model.

Key words: network public opinion; network rumors; rumors spread; Multi-Agent; information subject; trust recognition

收稿日期:2019-07-19

修回日期:2019-11-17

基金项目:国家社会科学基金项目“非常规突发事件中社交媒体不实信息的群体干预模式研究”(编号:14BTQ026);国家自然科学基金青年项目“新媒体中考考虑群体差异的谣言传播机理及干预策略研究”(编号:71503163);国家哲学社会科学基金一般项目“大数据背景下辟谣信息生成传播机理与辟谣效果研究”(编号:19BGL234);教育部人文社会科学青年基金项目“突发事件中社交媒体可信信息的特征识别研究”(编号:17YJCZH199)资助的研究成果之一。

作者简介:滕 婕(ORCID:0000-0002-8064-6551),女,1993年生,硕士研究生,研究方向:信息管理与信息系统;夏志杰(ORCID:0000-0003-3354-0989),男,1978年生,博士,教授,硕士生导师,研究方向:信息管理与信息系统;罗梦莹(ORCID:0000-0002-6497-4213),女,1990年生,博士研究生,研究方向:信息管理;王筱莉(ORCID:0000-0003-2774-300X),女,1981年生,博士,讲师,研究方向:信息传播、应急管理理论。

0 引言

随着信息通信技术的迅猛发展,全媒体呈现出融合型媒体形态,平民化、个性化、交互性强等特征尤为突出,政府等权威机构发布的辟谣信息往往出现滞后问题^[1],传统的可信度评估标准受到 Web2.0 背景的强烈影响,在这种标准中,网民参与辟谣的热情空前高涨,越来越多的网民通过自身的知识和经验,以及用户间的协作共享或竞争等机制,形成群体智慧下的辟谣活动^[2]。然而,在群体参与的辟谣过程中,每一个体由于精力、时间等限制,所接触的信息量有限但却是纷繁复杂的,个体面对着各种质量层次不齐的辟谣信息,难以辨别各种信息信任状况,谣言传播势头难遏^[3]。因此,网络信息主体的有效信任识别不仅有利于防范谣言误导公众,还有利于净化信息传播环境。

为了解决网络中的信息安全等问题,以往研究对于谣言的控制问题,往往通过各种模型参数的设定实现群体间信息传播的交互模拟^[4-5]进行演化仿真研究,通过观察仿真结果来探索大数据环境下信息的传播规律^[6]。然而,从宏观层面研究群体交互下的信息传播规律,很难准确反映个体的观点决策行为,缺少描述网络信息传播中每一个体决策的整个过程规律^[7]。实际上,从个体微观层面了解参与人际关系网络形成的局部交互的复杂信息行为过程,从个体角度掌握对信息可信度判断决策的过程对有效解决危机事件下恶意信息误导等问题具有决定性作用。因此,网络谣言传播事件下,从微观角度了解辟谣信息在个体人际关系网络中的信任识别和传播原则,构建信息主体信任识别模型,对个体辨别网络上辟谣信息的真实性具有重要意义。

本研究基于 Multi-Agent 模型来研究危机事件下网络信息主体的信任识别问题。首先,借鉴 Personalized 传统的信任模型,对网络谣言事件中应急响应的主体进行不同状态角色的划分。其次,考虑到个体异质性的问题,通过 Agent 间历史交互信息,引入动态权重因子,将直接信任和间接信任进行组合呈现更有效的信任评价,构建考虑异质性的信任识别模型,以期通过网络信息个体的有效信任识别和辟谣信息高质量传播提供理论与方法支持。

1 理论基础与模型假设

1.1 理论基础

1.1.1 信任识别模型中信息主体状态划分 危机事件的应急响应过程中,采用 Multi-Agent 建模思想将危机事件中的信息主体抽象为智能体 Agent,首先信任识别模型的分析要确定网络舆情事件应急响应

中的主体,即 Agent,主体之间的交互行为是 Agent 行为,在此基础上构建 Agent 交互的规则和算法。信任识别模型的信息主体包括源主体、中间主体、推荐主体和目标主体四类^[8]。根据 Agent 在网络舆情事件爆发时充当的角色,本研究中的源主体 B 通常为网民个体,目标主体 S 通常为发布辟谣信息的辟谣信息发布者,推荐主体 A 指的是与目标主体 S 有过交互经验,且具有一定影响力的意见领袖。网络舆情事件下,三类主体在应急响应辟谣的过程中扮演不同的角色,发挥着不同的作用。

网民 B 是网络舆情事件中舆情信息的生产者、接收者和传播者,网络舆情事件下接收到网络不同版本辟谣信息的网民并不一定能够辨别出信息的真实性,之后的很长一段时间公众都处于信息的真空期^[9],此时,网络上就会出现意见领袖的影子,网民易受意见领袖观点的影响。网民因其自身从众性特点,极易发生群体极化现象,是网络舆情事件应急响应下急需进行引导和调节的群体;

意见领袖 A 是微博上经常为他人提供信息、意见和评论的一些用户,包括名人、媒体和网络舆情事件下出现的草根意见领袖等^[10]。危机事件的信息传播主要由意见领袖推动^[11],网络意见领袖数量众多,媒介素养良莠不齐,他们往往会对信息发表观点,对危机事件下的辟谣信息进行信任判别,既而影响到网民对辟谣信息的信任度。但部分微博意见领袖的失范行为,会在舆论引导中产生负面效应。

辟谣信息发布者 S 包括个人和组织机构的官方类微博两种,随着新媒体时代的到来,网民通过自媒体可以在任意时间、地点发表言论和感想,每个人都是信息的生产者,因此辟谣信息发布者是不仅局限于政府等相关机构,而是指网络上发布辟谣信息的任何个人或组织机构。辟谣信息在舆情危机中具有稳定公众情绪,引领舆论发展方向的重要作用,公众对辟谣信息的信任状况会对群体行为产生重要的影响。

1.1.2 信任关系网的建立 公众对信息的态度和观点是综合素质下对事件的决策行为,由于个体异质性的影响,不同的个体会产生差异化的决策,个体的观点表现在对信息的信任判别^[12]。网络舆情事件下,网络谣言传播迅速,不同的信息主体也会发布各种辟谣信息,在开放的 Agent 系统中,网民需要从潜在的 Agent 中选择最为可信的 Agent 进行交互^[13]。危机情景下,信任值较高的 Agent 可以迅速维护网络环境的稳定,控制负面情感的传播,依据 Agent 表现出来的特征信息,判断 Agent 的可信任状况。

Agent 的表现特征分为直接信任和间接信任,如图 1 所示,A 与 B、B 与 S₁、B 与 S₂两主体之间进行直

接交互,通过直接交互经验建立的信任即为直接信任^[14]。对于网民来说,直接信任是最相关也是最可靠的信任信息。尽管依据 Agent 间的直接信任关系计算信任值是一个有效的方法,但单纯根据直接交互记录得到目标 Agent 的信任值是带有较强主观意愿色彩的,另外当直接交互次数不足时,利用直接信任无法准确计算信任值。因此,Agent 可以通过关系网络寻求其他 Agent 的意见进行信任判断。A 与 S₁、A 与 S₂ 是通过其他 Agent 的推荐信息进行间接的交互,这种由推荐信息计算得到的就是间接信任。本研究考虑到个体差异性,引入动态权重因子对直接信任与间接信任进行综合,得到信息主体的整体信任状况,推荐信息由 Agent 关系网络获得。此外,由于意见领袖或者辟谣信息发布者的信息行为是非静态的,网民根据以往的交互记录对未来行为的判别缺乏精确性,因此,本文将信息主体之间的信任判别记录进行不同时间窗口的分割,记为 T_i,其中 T_i 表示目前的时间窗口。

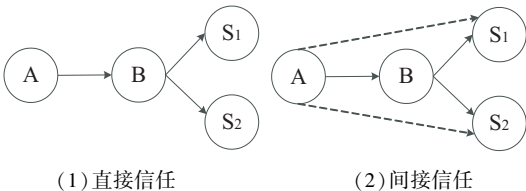


图1 信任关系网

1.2 模型的基本假设 在传统的 Personalized 信任模型中,信任程度表示类型的二值法具有实际的使用意义,其通过与信息主体的历史交互经验对信息主体进行信任判别,即,“0”代表信任,“1”代表不信任^[15]。另外,通过确认不同因素对个体的影响方式和影响程度是本研究的重点和难点,为了便于研究,作出以下假设:

假设 1:个体在进行信任识别判断过程中,由于时间、精力等限制,所接触的信息主体数量是有限的,即,个体周围的信息主体数量是有限的。

假设 2:不同信息主体之间通过交互的历史经验进行信任识别。信任识别机制的构建是建立在个体行为评估的基础上,收集信息主体以往的信任交互历史记录,综合考察对可信性产生影响的多种因素,依据信任关系的评估属性集进行建模,从而实现信任度的计算和授权决策。

假设 3:个体如果缺乏对目标主体的交互经验,将会寻求其他信息主体的意见和判别结果。意见领袖是具有代表性的网民集合,假设网民 B_i 的观点对网民 B_j 产生了某种影响,那么对于网民 B_j 所处的这一局部人际关系环境下, B_i 就相当于有某种影响的意见领袖。

假设 4:个体异质性体现在对信任判别结果的信心程度和可接受的最大误差程度。由于个体在信息获

取和识别能力等方面存在较大的差异,通过对信任判别结果的两个参数(可接受的最大误差程度 η 和信心程度 γ)反映个体异质性。

2 基于 Multi-Agent 的网络谣言中信息主体信任识别模型

2.1 基于 Multi-Agent 的网络谣言中信任机制构建 本研究中信任识别模型是基于网络舆情事件发生时,个体在面对网络各种多样复杂的辟谣信息时所做出的信任决策过程。由于个体在信息获取和识别能力等方面存在较大的差异,因此,基于 Personalized 信任模型,考虑到个体的异质性,引入动态权重因子,从微观层面对个体的信任识别过程进行建模,提出一个网络网络谣言事件信息主体信任识别机制。

信任模型的基本思路如下:首先,使用目标主体中的信息查询存在共同评价推荐者,其次,基于推荐者与个体的公共评价计算直接可信度,并融入个体异质性特征对其进行调整,然后,利用直接可信度和推荐者的评价计算间接信任值,最后,通过直接信任与间接信任的有效聚合得到综合信任值。模型的流程如下图 2 所示:

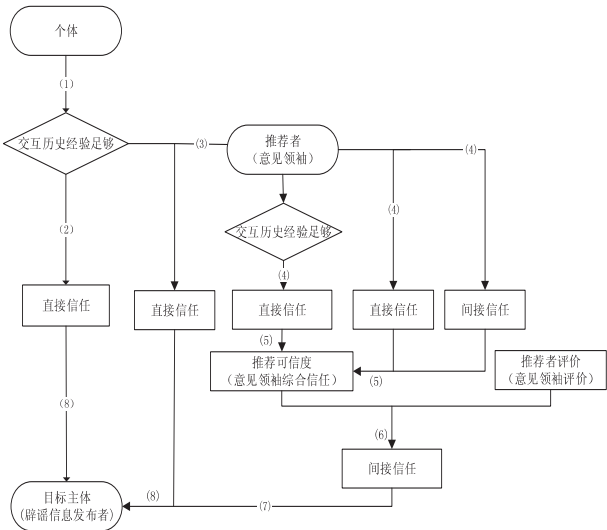


图2 信任机制流程图

a. 个体查询是否与目标主体拥有足够多的交互经验。如果是,则执行步骤 2,否则执行步骤 3。

b. 个体基于自身与目标的交互经验,利用动态的遗忘因子计算直接信任度,然后执行步骤 8。

c. 个体向推荐者(意见领袖)寻求帮助,来获取目标主体的间接信任。

d. 用户基于对意见领袖的交互经验计算其推荐行为的直接和间接信任度。

e. 由步骤 4 结果获取推荐者的推荐可信度。

f. 通过融合不同推荐者的推荐可信度和推荐评价获取间接信任。

g. 利用动态加权因子将直接信任和间接信任结合起来。

h. 完成信任建模。

2.2 基于 Multi-Agent 的网络谣言信任识别模型

2.2.1 意见领袖的直接信任值

意见领袖的直接信任是网民和意见领袖通过以往与辟谣信息发布者进行直接交互的历史记录和经验,网民依据公共评价的判别对的特征对意见领袖进行信任判别。其中,网民和意见领袖对辟谣信息发布者的信任结果分别用向量 R_B, S_i 和 R_A, S_i 表示, (r_A, S_i, r_B, S_i) 表示信任判别对, r_B, S_i 和 r_A, S_i 分别为 R_B, S_i 和 R_A, S_i 在同一时间窗口 T_i 内的二分类数值,分别表示网民和意见领袖对辟谣信息发布者进行信任判断的结果, $T = \{T_1, T_2, \dots, T_n\}$, $T_1 = 1$ 是当前最近一次的时间窗口。个体观点的选择是内外环境因素综合下的决策行为,观点的选择结果表现在对信息的信任判别^[16]。所以,当网民的观点决策为相信辟谣信息发布者发布的辟谣信息,则,令 $r_B, S_i = 1$; 当网民观点决策是不信任辟谣信息时,令 $r_B, S_i = 0$, 同理,意见领袖观点是相信辟谣信息时,令 $r_A, S_i = 1$, 当 $r_A, S_i = 0$, 表示意见领袖认为辟谣信息不可信。意见领袖 A' 的直接信任度反映的是意见领袖 A' 向网民 B' 提供关于辟谣信息发布者 S' 的真实评论的概率,因为意见领袖的信息不完整,估计该概率的最佳方法是使用概率的期望值。 N_{all}^A 表示网民和意见领袖对辟谣信息发布者 S' 进行信任评价的信任判别对总数目, N_{Si}^A 表示网民和意见领袖对辟谣信息发布者 S' 进行共同信任判断的总数目。

$$N_{all}^A = \sum_{i=1}^{n1} N_{Si}^A \quad (1)$$

N_{consis}^A 代表网民与意见领袖 A' 对辟谣信息信任判断结果相一致的总数目,则判断结果不一致的总数目为 $N_{all}^A - N_{consis}^A$ 。因此,意见领袖的直接信任值 $Tr_{die}(A)$ 为:

$$\begin{cases} \alpha = N_{consis}^A + 1 \\ \beta = N_{all}^A - N_{consis}^A + 1 \\ Tr_{die}(A) = E(P(A)) = \frac{\alpha}{\alpha + \beta} \end{cases} \quad (2)$$

$P(A)$ 指的是网络舆情事件下,意见领袖 A' 针对谣言进行公平的信任判别的概率, $E(P(A))$ 为 β 分布的期望值,也为直接可信度计算数值。

2.2.2 意见领袖的间接信任值 网络舆情事件下,意见领袖会根据自身以往的经验对辟谣信息进行最初的可信判断,网民和意见领袖并不总是对同一谣言进行信任判别,即信任判别对 (r_A, S_i, r_B, S_i) 并不总是存在,当网民以往对辟谣信息发布者缺少充足的交互经验时,就会依赖于意见领袖的信任判别^[17]。网络

舆情事件发生时,在不同的时间窗口中,意见领袖 A' 对于辟谣信息发布者进行信任判别的总数目用 N_{all}^A 表示, N_{fair}^A 是意见领袖 A' 对于辟谣信息发布者进行公平的信任判别的总数目,则恶意攻击下的信任判别总数目为 $N_{all}^A - N_{fair}^A$, 与意见领袖的直接可信度相似,意见领袖的间接可信度也是由意见领袖 A' 针对谣言进行公平的信任判别的概率决定,意见领袖提供的公平的信任判别比例越高,间接可信度就越大。意见领袖的间接信任值 $Tr_{die}(A)$ 为:

$$\begin{cases} \alpha' = N_{fair}^A + 1 \\ \beta' = N_{all}^A - N_{fair}^A + 1 \\ Tr_{indie}(A) = \frac{\alpha'}{\alpha' + \beta'} \end{cases} \quad (3)$$

2.2.3 意见领袖的综合信任值

信息判别来源主要包括直接信任和间接信任,本研究中考虑到个体的异质性,通过设置权重因子将两者结合,得到在危机事件下意见领袖发布信息的综合信任值 $Tr(A)$ 。

$$Tr(A) = \varepsilon * Tr_{die}(A) + (1 - \varepsilon) * Tr_{indie}(A) \quad (4)$$

如果一个网民与辟谣信息发布者有充足的交互经验,则会通过以往交互经验,赋予直接信任 ε 一个较大的权重,直接信任的权重随着直接交互次数的增多而增大,当它超过阈值时,用户完全相信自己的经验,即,当 $\varepsilon = 1$ 时,该网民将会只依赖于对意见领袖的直接信任 $Tr_{indie}(A)$; 反之,当直接交互次数不足时,在进行信任评估时,应当考虑间接信任。 ε 的大小由公式(5)计算可得:

$$\varepsilon = \begin{cases} \frac{N_{all}^A}{N_{lim}}, N_{all}^A < N_{lim} \\ 1, N_{all}^A \geq N_{lim} \end{cases} \quad (5)$$

其中, N_{all}^A 表示网民和意见领袖对辟谣信息发布者进行信任评价的信任判别对总数目, N_{lim} 是网民认为与辟谣信息发布者以往的直接交互经验足够时,需要的最低交互次数。对于临界值 N_{lim} 的确定,本研究利用基于 Chernoff Bound 的理论方法,它受到网民可接受的误差和信心程度的影响:

$$N_{lim} = -\frac{1}{2\eta^2} \ln \frac{1-\gamma}{2} \quad (6)$$

其中, η 是网民可接受的最大误差程度, γ 代表信心程度。可接受误差程度越大,所需最低交互次数越少;信心程度越高,要求的最低交互次数越大。

2.2.4 辟谣信息发布者的直接信任值

辟谣信息发布者的直接信任是网民根据以往与辟谣信息发布者进行直接交互的历史记录和经验,从而对其进行的信任判别,在直接信任关系中,网民与辟谣信息发布者交互的次数增加,两主体间的信任关系会更加确定。网络舆情事件发生时,网民对辟谣信息发布者的所有

信任判别结果用向量 $R_{B,S}$ 表示, R_{B,S_i} 是在某一时间窗口 T_i 内的二分类数值, $T = \{T_1, T_2, \dots, T_n\}$, $T_1 = 1$ 是当前最近一次的时间窗口。网民与辟谣信息发布者以往的交互经验并不是总与当前的行为相关,辟谣信息发布者可能会随时间而改变行为。因此,对于以往的交互记录的信任状况,需要通过一定的衰减来反映数据的时效性,近期信任判别的重要性优于较早时期记录^[18],本研究引用遗忘因子 λ 表示这种衰减现象。基于网民 B 与辟谣信息发布者 S 的直接交互经验,可以得到 B 对 S 的直接信任:

$$Tr_{die}(S') = \frac{\sum_{i=1}^n N_{true,i}^B \lambda^{i-1} + 1}{\sum_{i=1}^n (N_{true,i}^B + N_{false,i}^B) \lambda^{i-1} + 2} \quad (7)$$

其中,在某一时间窗口 T_i 内, $N_{true,i}^B$ 表示网民相信辟谣信息的总数目; $N_{false,i}^B$ 表示网民不信任辟谣信息的总数目。

2.2.5 辟谣信息发布者的间接信任值 辟谣信息发布者的间接信任是网民对信息的判断不依赖于以往与辟谣信息发布者进行直接交互的历史记录和经验,而是通过参考其他信息主体的意见,从而对辟谣信息发布者进行的信任判别,即,当网民 B 对于辟谣信息发布者 S 没有充足的交互经验时, B 往往求助于意见领袖 A 从而对 S 完成评估。意见领袖 $A' = \{A'_1, A'_2, \dots, A'_k\}$ 对辟谣信息发布者 $S' = \{S'_1, S'_2, \dots, S'_l\}$ 进行信任判别,不同信息的信任判别结果被划分到不同的时间窗口下。网络舆情事件下,意见领袖 A_j 在时间窗口 T_i 内的所有信任判别记录中,认为辟谣信息是可信任的总数目为 $N_{true,i}^{A_j}$, 认为辟谣信息不可信任的总数目为 $N_{false,i}^{A_j}$ 。考虑到意见领袖可能存在对辟谣信息发布者进行恶意攻击信息,因此,将有过恶意判别的意见领袖做出的信任判别结果赋予较小的权重,降低综合信任值,通过意见领袖以往的信任判别记录计算信任判别的贴现函数,贴现函数对恶意判别行为具有有效的识别和处理^[14]。

$$D_{true,i}^{A_j} = \frac{2Tr(A_j) N_{true,i}^{A_j}}{(1 - Tr(A_j))(N_{true,i}^{A_j} + N_{false,i}^{A_j}) + 2} \quad (8)$$

$$D_{false,i}^{A_j} = \frac{2Tr(A_j) N_{false,i}^{A_j}}{(1 - Tr(A_j))(N_{true,i}^{A_j} + N_{false,i}^{A_j}) + 2} \quad (9)$$

其中, $Tr(A_j)$ 为意见领袖 A_j 的综合可信度。通过贴现函数可知辟谣信息发布者的间接可信度为:

$$Tr_{indie}(S) = \frac{[\sum_{j=1}^k \sum_{i=1}^n D_{true,i}^{A_j} \lambda^{i-1}] + 1}{[\sum_{j=1}^k \sum_{i=1}^n (D_{true,i}^{A_j} + D_{false,i}^{A_j}) \lambda^{i-1}] + 2} \quad (10)$$

2.2.6 辟谣信息发布者的综合信任值 网络舆情事态下,辟谣信息发布者的综合信任和意见领袖的

综合信任相似,由网民与辟谣信息发布者的直接交互经验和意见领袖的间接推荐信息可以得到辟谣信息发布信息的综合信任值 $Tr(S)$ 。

$$Tr(S) = \varepsilon' * Tr_{die}(S) + (1 - \varepsilon') * Tr_{indie}(S) \quad (11)$$

如果一个网民与辟谣信息发布者有充足的交互经验,权重 ε' 的值偏大,直接信任的权重随着直接交互次数的增多而增大,当它超过阈值时,用户完全相信自己的经验。即,当 $\varepsilon' = 1$ 时,该网民将会只依赖于对辟谣信息发布者的直接信任 $Tr_{die}(S)$;反之,当直接交互次数不足时,在进行信任评估时,应当考虑间接信任。 ε' 的大小由公式(12)计算可得:

$$\varepsilon' = \begin{cases} \frac{N_{all}^B}{N_{lim}}, N_{all}^B < N_{lim} \\ 1, N_{all}^B \geq N_{lim} \end{cases} \quad (12)$$

其中, N_{all}^B 表示网民对辟谣信息发布者进行信任评价的总数目, N_{lim} 是网民认为与辟谣信息发布者以往的直接交互经验足够时,需要的最低交互次数,由公式(6)求得临界值 N_{lim} 。对于临界值 N_{lim} 的确定,仍然采用 Zhang 提出的基于 Chernoff Bound 理论的方法。

3 算例分析

为了证明模型的有效性,本研究提供例子来详细描述模型中各个部分的计算过程,并对比了不同信任判别记录的意見領袖和辟谣信息发布者的信任建模。

假设信任网络中存在三名意见领袖 $A' = \{A'_1, A'_2, A'_3\}$, 五名辟谣信息发布者 $S' = \{S'_1, S'_2, S'_3, S'_4, S'_5\}$, 和—网民 B' , 系统的关系网络图如图 3 所示。实线表示主体之间产生直接信任关系,虚线指主体之间是产生间接信任的关系。网络舆情下,主体之间的信任判别按时间划分到不同的时间窗口,时间窗口记为 T_i , 其中 T_1 表示目前的时间窗口。

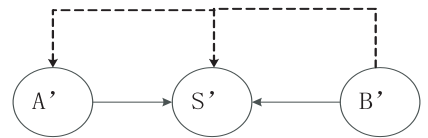


图3 信任关系结构图

3.1 意见领袖信任值 网络舆情情境下,意见领袖传播信息的过程在某种意义上就是帮助网民筛选整理信息的过程,意见领袖传播的信息是否会对网民产生心理和行为影响取决于网民 B' 对意见领袖 A' 的信任度。当网民 B' 与辟谣信息发布者 S' 进行交互时,会对其进行信任判别,如表 1 所示;当网民 B' 与辟谣信息发布者 S' 没有足够的交互经验或对辟谣信息有质疑时,往往会寻求意见领袖 A' 的观点,意见领袖 A' 根据以往与辟谣信息发布者进行直接交互的历史记录和经验,从而对其进行信任判别,如表 2 所示。信任判别过

程按时间划分到不同的时间窗口。

表 1 网民 B 对辟谣信息发布者的信任判别

	T _i	T ₁	T ₂	T ₃	T ₄	T ₅
S ₁		1	1	1	1	0
S ₂		-	1	1	1	1
S ₃		1	0	1	0	1
S ₄		-	0	-	1	-
S ₅		1	-	-	-	1

1)“-”是指危机事件下,网民在时间窗口 T_i内未对辟谣信息进行信任判别;2)“1”是指危机事件下,网民在时间窗口 T_i下认为辟谣信息是可信的;3)“0”表示危机事件下,网民在时间窗口 T_i下认为辟谣信息不可信

表 2 意见领袖 A 对辟谣信息发布者的信任判别

A _j	A ₁					A ₂					A ₃				
T _i	T ₁	T ₂	T ₃	T ₄	T ₅	T ₁	T ₂	T ₃	T ₄	T ₅	T ₁	T ₂	T ₃	T ₄	T ₅
S ₁	0	0	0	0	0	1	1	1	0	0	1	1	1	1	1
S ₂	0	0	0	0	0	0	1	1	1	1	1	1	1	1	1
S ₃	-	-	-	-	-	1	-	1	0	0	1	1	1	1	1
S ₄	0	0	0	0	0	1	0	0	1	-	1	1	1	1	1
S ₅	0	0	0	0	0	1	1	0	0	1	-	-	-	-	-

通过将网民和意见领袖对辟谣信息的信任判别结果进行比较,网民 B 和意见领袖 A 对辟谣信息发布者 S 的共同判别对 $N_{all}^{A1} = 13$, $N_{all}^{A2} = 17$, $N_{all}^{A3} = 16$ 。然而网民 B 与意见领袖 A 对不同时间窗口下的辟谣信息具有不一样的信任判别结果,如表 3 所示,意见领袖 A₁、A₂、A₃的直接可信度是不同的,A₃的直接可信度 0.83 为最高,A₁的直接可信度 0.07 是最低的,这说明意见领袖 A₃在危机事件下的辟谣信息信任判别相比另外两位意见领袖更为公平,A₁很可能为辟谣信息发布者的恶意攻击者。

表 3 意见领袖的直接和间接可信度

A _j	A ₁	A ₂	A ₃
$N_{true}^{A_j}$	0	12	14
α	1	13	15
β	14	6	3
$Tr_{die}(A)$	0.07	0.68	0.83
$N_{fair}^{A_j}$	0	14	20
α'	1	15	21
β'	21	10	1
$Tr_{indie}(A)$	0.05	0.60	0.95

此外,对于意见领袖 A 的间接可信度值的计算,由表 2 可知,意见领袖 A 对于辟谣信息发布者进行信任判别的总数目 $N_{all}^{A1} = 20$, $N_{all}^{A2} = 23$, $N_{all}^{A3} = 20$,“1”表示是意见领袖 A 对于辟谣信息发布者进行公平的信任判别,“0”表示是意见领袖 A 对于辟谣信息发布者进行恶意攻击下的信任判别,A₁的间接可信度最低, $Tr_{indie}(A1) = 0.05$,A₃的间接可信度 $Tr_{indie}(A3) = 0.95$,这说明在危机情境下,A₃最有可能提供公平公正的信任判别,A₁最有可能提供错误信息。

综合可信度是由意见领袖的直接可信度和间接可

信度得到,因此,需要确定直接和间接可信度的权重 ε ,而 ε 的值是 η 、 γ 和 $N_{all}^{A_j}$ 决定, $N_{all}^{A1} = 13$, $N_{all}^{A2} = 17$, $N_{all}^{A3} = 16$,设信心程度 γ 为固定值 0.8,网民可接受的不同误差程度下的意见领袖综合信任度如表 4 所示。结果表明,网络舆情事件下,意见领袖 A₁、A₂和 A₃同时推荐辟谣信息相关信息的信任态度时,A₃对辟谣信息的评价最为公平和可靠,相反,A₁对辟谣信息的信任评价最不可信,原因可能是其对辟谣信息发布者存有较大偏见。

表 4 不同误差程度下的意见领袖综合信任值

η	0.05	0.10	0.15
N_{min}	461	115	51
$Tr(A_x)$	0.051	0.052	0.055
$Tr(A_y)$	0.603	0.612	0.627
$Tr(A_z)$	0.946	0.933	0.912

不同的网民与辟谣信息发布者交互的历史记录和经验不一样,即使对于相同的意见领袖,也会具有不同的信任值,为了反映网民的异质性对意见领袖信任值的差异性,本研究假设在不同的时间窗口下,网民 B 对辟谣信息发布者 S 也存在有交互记录,信任判别结果如表 5 所示。

表 5 网民 B 对辟谣信息发布者的信任判别

T _i	T ₁	T ₂	T ₃	T ₄	T ₅
S ₁	1	1	-	-	1
S ₂	1	-	-	1	-
S ₃	0	1	-	-	-
S ₄	1	1	-	-	-
S ₅	0	0	-	-	-

通过相同的计算过程,可以知道网民 B'对意见领袖的直接可信度、间接可信度和综合信任值,令 $\eta = 0.10, \gamma = 0.8$,网民 B'对意见领袖的信任值如表 6 所示。结果显示,网民 B'与 B 的直接信任值是不一样的,虽然两者对于意见领袖的间接信任值没有变化,其最终的综合信任值仍具有差异性。

表 6 网民 B'对意见领袖的信任值

A_j	A_1	A_2	A_3
$Tr_{die}(A_j)$	0.07	0.26	0.50
$Tr_{indie}(A_j)$	0.05	0.60	0.95
$Tr(A_j)$	0.052	0.550	0.887

最后,为了检验模型的鲁棒性,在表 2 意见领袖进行的信任判别结果中,假定“1”表示意见领袖进行的恶意的信任判别,“0”表示意见领袖进行的公平的信任评判,则网民 B'对意见领袖 A_1 、 A_2 和 A_3 的间接信任值发生了改变,如表 7 所示。网民在可接受的不同大小的信任误差 η 内,对意见领袖的直接可信度也在随之变化,如表 8 所示,研究模型通过调节直接与间接可信度的权重 ε 的大小,能够有效反映意见领袖的信任值。

表 7 意见领袖恶意评判下的间接信任值

A_j	A_1	A_2	A_3
$N_{fair}^{A_j}$	20	9	0
α'	21	10	1
β'	1	15	21
$Tr_{indie}(A_j)$	0.95	0.40	0.05

表 8 鲁棒性检验的意见领袖综合信任值

η	0.1	0.15	0.2
N_{min}	115	51	29
$Tr(A_1)$	0.85	0.73	0.55
$Tr(A_2)$	0.44	0.49	0.56
$Tr(A_3)$	0.16	0.29	0.48

3.2 辟谣信息发布者信任值 网络舆情事件发生时,谣言在短时间内快速扩散,辟谣信息发布者 S'传播的辟谣信息是否会消减危机事件的负面影响取决于网民 B'对 S'的信任度。为了便于运算,假设网民 B'并未与辟谣信息发布者 S_3 的交互信任记录,且危机事件下,网民 B'浏览了辟谣信息发布者 S_3 发布的辟谣信息,因此,根据公式(7)所得网民 B'对辟谣信息发布者的直接信任值为:

$$Tr_{die}(S_3) = \frac{0 + 1}{0 + 0 + 2} = 0.5$$

(13)

由于网民对于辟谣信息发布者缺乏充足的交互经验,往往寻求意见领袖的建议,意见领袖在危机事件应急响应中对网民的信息行为起到了重要的调节作用^[19]。在意见领袖的信任测试例子中,误差值 $\eta = 0.15$ 时,网民对 A_1 、 A_2 和 A_3 的信任值分别为 0.055、

0.627 和 0.912, A_1 的信任值最低,且与另两位意见领袖信任值差异过大,因此,网民只会受到意见领袖 A_2 和 A_3 的信任判别结果的影响,根据两者对 S_3 发布的辟谣信息的信任判别结果,可以求得不同时间窗口下评判辟谣信息为可信和不可信的总数目,如表 9 所示,信任判别的贴现函数如表 10 所示。令遗忘因子 $\lambda = 0.9$,根据公式(8)、(9)和(10),可计算出辟谣信息发布者 S_3 的间接信任值 $Tr_{indie}(S_3)$ 为:

$$Tr_{indie}(S_3) = ((0.528 * 0.9^{1-1} + 0.528 * 0.9^{3-1}) + (\sum_{i=1}^5 0.874 * 0.9^{i-1}) + 1 / (0.528 * 0.9^{1-1} + \sum_{i=3}^5 0.528 * 0.9^{i-1}) + (\sum_{i=1}^5 0.874 * 0.9^{i-1}) + 2) = 0.762$$

(14)

表 9 意见领袖 A_1 和 A_2 对 S_3 的信任分析

T_i	T_1	T_2	T_3	T_4	T_5
$N_{true,i}^{A_2}$	1	0	1	0	0
$N_{false,i}^{A_2}$	0	0	0	1	1
$N_{true,i}^{A_3}$	1	1	1	1	1
$N_{false,i}^{A_3}$	0	0	0	0	0

表 10 信任判别的贴现函数

T_i	T_1	T_2	T_3	T_4	T_5
$D_{true,i}^{A_2}$	0.528	0	0.528	0	0
$D_{false,i}^{A_2}$	0	0	0	0.528	0.528
$D_{true,i}^{A_3}$	0.874	0.874	0.874	0.874	0.874
$D_{false,i}^{A_3}$	0	0	0	0	0

网民 B'并未与辟谣信息发布者 S_3 有过交互,因此,直接可信度的权重 $\varepsilon' = 0$,辟谣信息发布者 S_3 的综合信任值 $Tr(S_3)$ 为:

$$Tr(S_3) = 0 * 0.5 + (1 - 0) * 0.762 = 0.762$$

(15)

网络舆情事件下,网民对辟谣信息的信任度是意见领袖影响的结果,从 S_3 的综合信任值可以看出,意见领袖 A_3 对辟谣信息的信任判别更值得信赖,而 A_2 对网民的影响较低。当不考虑意见领袖的影响作用时, S_3 的间接可信值为:

$$Tr'_{indie}(S_3) = \frac{(1 * 0.9^{1-1} + 1 * 0.9^{3-1}) + (\sum_{i=1}^5 1 * 0.9^{i-1}) + 1}{(1 * 0.9^{1-1} + \sum_{i=3}^5 1 * 0.9^{i-1}) + (\sum_{i=1}^5 1 * 0.9^{i-1}) + 2} = 0.743$$

(16)

此时, S_3 的综合信任值 $Tr'(S_3)$ 为:

$$Tr'(S_3) = 0 * 0.5 + (1 - 0) * 0.743 = 0.743$$

(17)

比较 S_3 的综合信任值 $Tr(S_3)$ 和 $Tr'(S_3)$ 可以发现,通过本研究模型运算下的综合信任值和实际的信

信任值非常接近,这说明模型可以有效的预测信任值的大小。

4 仿真结果与分析

本实验通过分析网民、意见领袖和辟谣信息发布者的一些特性进行仿真,并模拟 Agent 信息主体之间的交互行为,从而验证信任建模的有效性。实验所使用的工具为 Matlab R2014b,它具有简单易用、可操作性强和数据处理能力好等特点。

4.1 意见领袖信任测试 网络舆情情境下,信任模型是处于开放式多参与者的环境下,意见领袖信任是信任模型中最重要的部分,由于网民很难与每一个辟谣信息发布者都发生交互信息,往往会求助意见领袖的意见和其信任判定结果,但是由于危机事件下,网民易受到一些恶意意见领袖的信任判别的影响,从而影响网民对辟谣信息的识别。因此,在意见领袖信任测试实验中,考虑了三种典型的恶意意见领袖类型^[20]。

a. 意见领袖 A_0 对辟谣信息进行公正的信任评判。

b. 不管辟谣信息发布辟谣信息真假情况如何,意见领袖 A_1 对辟谣信息判别结果总是为真。

c. 意见领袖 A_2 往往会提供与实际情况相反的判别结果。

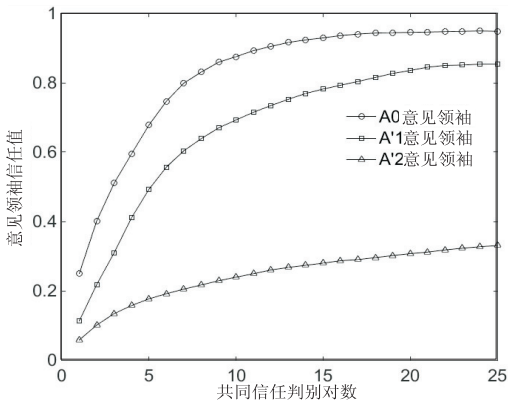


图4 意见领袖信任测试

基于三种意见领袖测试进行 25 个信任判别周期内信任值的变化。实验的结果如图 4 所示。仿真结果显示,三个意见领袖的信任值在最初都非常低,随着网民与其共同信任判别对的增加,不同类型的意见领袖信任值产生了较大差异,意见领袖的信任判别与现实情况越不相符,其信任值越低。其次,尽管意见领袖 A_1 提供的信任评判结果与 A_0 对辟谣信息进行公正的信任评判很接近,但也无法达到和 A_0 同样高度的信任值,然而对于总是提供公正判别结果的意见领袖,信任值可以随着共同信任判别对数的增多而提高。因此,本研究模型通过 Agent 关系网络寻找意见领袖,如果意见领袖出现恶意信息引领行为,都会存在其信用记录中,影响信用值,对恶意攻击他人或提供虚假信息的

信息主体可以进行有效识别。所以,意见领袖恶意信任评判的成本是很大的,信任模型可在一定程度上避免恶意意见领袖的出现。

4.2 辟谣信息发布者信任测试 由于在群体参与的辟谣过程中,个体因精力、时间等限制,所接触的信息主体有限,且往往会受到意见领袖观点的影响。其次,因为网民和意见领袖并不总是对每一个辟谣信息发布者具有交互历史,网民和意见领袖对辟谣信息发布者的信任判别会根据以往的交互经验对任意 k' 个辟谣信息发布者进行信任判别。因此,为了模拟现实中个体在进行观点决策和信任判别时的场景,令 $l = 10, k = 80, k' = 8$ 。在以下两种场景中对信任模型中辟谣信息发布者展开信任测试模拟。场景一:其中,意见领袖中的 30% 是恶意的信任判别者,即,他们总是作出与实际情况相反的信任判别;场景二:其中,意见领袖中的 60% 是恶意的信任判别者。通过比较辟谣信息发布者间接信任值和综合信任值,结果如图 5 所示。

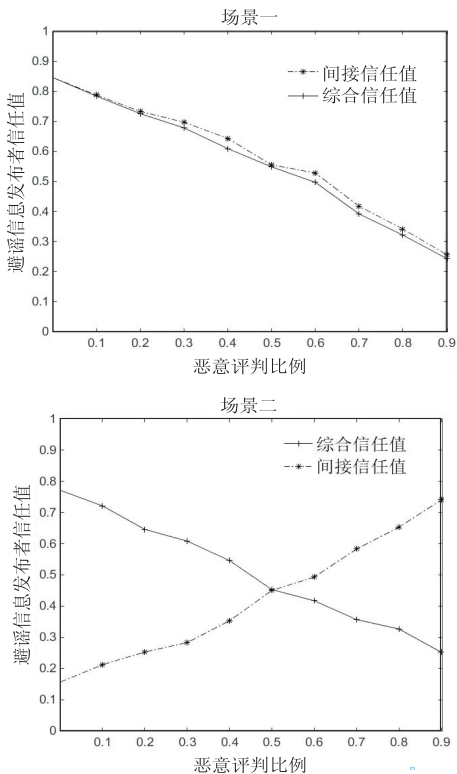


图5 辟谣信息发布者信任测试

在恶意意见领袖的比例较小时(场景一),虽然意见领袖恶意评判的概率一直在增加,但间接信任值和综合信任值的趋势一直是递减趋势,且两者的递减幅度并未有明显的差别;但是当意见领袖中的 60% 是恶意的信任判别者(场景二),即,他们总是作出的信任判别总是与实际情况相反,当使用综合信任值衡量辟谣信息发布者的可信信任问题时,随着意见领袖恶意评判的概率的增加,信任值递减,当使用间接信任值衡量辟谣信息发布者的信任问题时,随着意见领袖恶意评

判的概率的增加,信任值逐渐增加。显然,综合信用值比间接信用更贴近现实,容错性更低。这是因为间接信任只是网民根据意见领袖意见,得到的辟谣信息发布者信任判别,判别结果只取决于意见领袖评判。而本模型下的综合信任是在间接信任基础上加入了网民本身与辟谣信息发布者交互历史的信任判别,并且两种来源的判别比重取决于网民的信心程度和可接受最大误差程度,充分考虑到了个体的异质性,更加符合现实实际情况。

4.3 延展性测试 本文选取“塑料紫菜”这一典型食品安全网络谣言事件为例,根据微博“塑料紫菜”为关键字,通过“集搜客”爬虫软件获取相关博文量,博文的数量作为信息发布主体的数量,统计整理 2017 年 2 月 27 日 01:00-21:00 的信息发布主体数据,随着网络舆情的持续发展,信息主体范围和数量会有大幅增加,即,参与人际关系网络 Agent 的数量会受到事件重要性的影响发生改变。为了验证这种情况下本研究的信任模型仍然有效,进行了模型延展性测试,统计整理得到,信息发布者数量从 01:00 的 16 达到 21:00 时刻的 814,网民根据意见领袖对辟谣信息发布者的观点进行信任判别,意见领袖作出的信任判别有 23% 是恶意评判,假设每次增加的间隔设置为 80,仿真结果如图 6 所示。

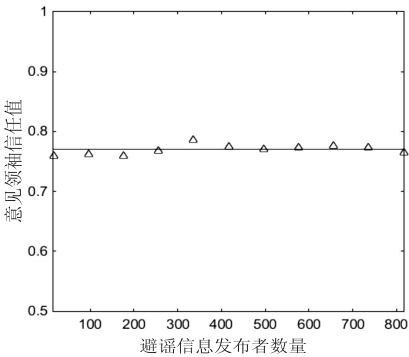


图 6 模型延展性测试

从图 6 结果得到,在信息通信技术快速发展的时代,随着网络舆情的发展,参与辟谣的网民数量的增多,辟谣信息发布者的数量虽然会影响到对辟谣信息发布者的信任判别总数目,但结果显示,本研究模型具有较好的包容性,随着信息发布者数量的变化,意见领袖的信任值并未发生变化,表明参与 Agent 人际关系网络的信息主体的数量不会影响到信任模型的有效性。

5 结 论

本研究基于 Multi-Agent 模型把网络舆情事件中的主体进行划分,以 Agent 之间的直接交互历史,结合 Agent 关系网络中意见领袖对辟谣信息发布者的信任

判别信息,以动态权重因子将直接信任和间接信任进行组合,最后,考虑了恶意意见领袖进行的信息判别问题计算信任值时。根据本文研究得出结论如下:

第一,网络谣言事件下,信任识别模型的构建可对恶意攻击他人或提供虚假信息的信息主体进行有效识别。在网络谣言事件发生时,政府等权威机构应重视和充分理解采用该信任识别类系统所带来的实际效益。相关网络舆情应急部门只有真正重视网络信任识别系统的重要功能和作用,才能有效深度挖掘 Agent 人工智能信任识别的特征,真正的使得基于 Agent 的信任识别模型能在舆情管理中得到实际开发和应用,使网民在危机事件发生时,不会受到网络谣言的影响,形成社会恐慌局面。另一方面,关系网络中信息流动性较大,对已被证实网络上出现的虚假且危害性较大的信息,政府等权威应急机构可通过不同渠道的宣传普及方式,通过关系网的信息传播,揭示虚假信息的危害性及其不良意图,从而抑制谣言的传播。

第二,动态权重因子的引入对 Agent 主体进行信任识别的效果更贴近实际。在有限群体中,为构建危机情境下明确的信任体系,政府等相关机构应该建立开放性的信息共享、沟通平台,使群体成员能够掌握更多其他个体的信息历史交互情况,进而更好实现信任关系的筛选和建立。另外,考虑到网络意见领袖对网民信任判别的影响,政府等应急部门要培养代表官方权威信息传播的意见领袖,意见领袖所在的关系网络是一个发生信息交互扩散的有限群体,在公共危机事件下,政府部门通过发挥意见领袖作用,建立相应的激励机制,例如对于进行公正信任评判的意见领袖,可以通过信任值排名给予其相应的物质奖励,同时,对意见领袖的失范行为也要进行一定程度的惩罚。

第三,参与人际关系网络 Agent 的数量并不影响信任模型的有效性。在公共危机事件爆发时,应对社会化媒体谣言传播的方法主要依靠政府和媒体发布的官方权威信息辟谣者参与网络谣言控制,在实践中往往可操作性和及时性不强,一味地追求信息的权威性,只依靠政府机构发布权威信息进行谣言环境下的辟谣工作,将信息局限在官方媒体中,将不利于信息的快速扩散和信任关系网络的有效形成。鉴于此,政府应充分调动起社会各个机构平台的积极性,使更多的团体或个人作为官方权威信息的拥护者,当公共危机事件发生时,网络上均匀分布权威信息,辟谣的相关信息会被公众较快接收和获取。

本研究根据人际交互经验研究了考虑个体异质性特征的信息主体识别模型,但在网络谣言环境中不同信息主体重要性和影响力存在一定的差异,这些问题都是以后需要进一步研究的重要内容。

参考文献

- [1] 唐雪梅,赖胜强.突发事件中政府对网络谣言的辟谣策略研究——以太伏中事件为例[J].情报杂志,2018,37(9):95-99.
- [2] 占欣,夏志杰,罗梦莹,等.影响群体智慧抑制社会化媒体谣言传播的因素研究[J].图书馆,2018(8):85-90.
- [3] Metzger M J, Flanagan A J, Medders R B. Social and heuristic approaches to credibility evaluation online[J]. Journal of Communication, 2010, 60(3):413-439.
- [4] 张亮,杨闪,任立肖.基于元胞自动机的互联网迷因模仿意愿传播仿真[J].情报杂志,2018,37(4):114-119.
- [5] 潘新.基于复杂网络的舆情传播模型研究[D].大连:大连理工大学,2010.
- [6] 张鹏,赵动员,谢毛迪,等.基于强关系的移动社交网络信息传播机理建模与仿真研究[J].情报科学,2019,37(3):105-111.
- [7] Flicker L. The influence of opinion leaders[J]. An independent review, 2012, 35(3):74-79.
- [8] 甘早斌,丁倩,李开,等.基于声誉的多维度信任计算算法[J].软件学报,2011,22(10):2401-2411.
- [9] 王虎,洪巍.食品安全网络谣言应对研究——以“塑料紫菜”事件为例[J].电子政务,2017(12):22-30.
- [10] Lazarsfeld P F, Berelson B, Gaudeth H. The People's Choice[M]. New York: Columbia University Press, 1948:434-445.
- [11] 曹学艳,段飞飞,方宽,等.网络论坛视角下突发事件舆情的关键节点识别及分类研究[J].图书情报工作,2014,58(4):65-70.
- [12] Starcke K, Brand M. Decision making under stress: A selective review[J]. Neuroscience & Biobehavioral Reviews, 2012, 36(4):1228-1248.
- [13] 伍京华,张富娟,许陈颖.基于Agent的情感劝说的信任识别模型研究[J].管理工程学报,2019,33(2):219-226.
- [14] 张高旭.基于多属性评价的信任模型研究[D].天津:天津大学,2016.
- [15] Zhang J, Cohen R. Design of a mechanism for promoting honesty in e-marketplaces[C]. Proceedings of the National Conference on Artificial Intelligence. AAAI Press, 2007, 22(2):1495.
- [16] 黄传超,胡斌,闫钰炜,等.网络暴力下突发事件中观点决策与舆情反转[J].管理工程学报,2019,33(1):252-258.
- [17] 彭泽洲.基于社会网络与声誉信任机制的移动多Agent系统信任模型[J].计算机应用与软件,2012,29(8):190-192.
- [18] He D, Chen C, Chan S, et al. Re Trust: Attack-resistant and lightweight trust management for medical sensor networks[J]. IEEE Transactions on Information Technology in Biomedicine, 2012, 16(4):623-632.
- [19] 王佳敏,吴鹏,沈思.突发事件中意见领袖对网民的情感影响建模研究[J].情报杂志,2018,37(9):120-126.
- [20] Dellarocas C. Immunizing online reputation reporting systems against unfair ratings and discriminatory behavior[C]. Proceedings of the 2nd ACM conference on Electronic commerce. ACM, 2000:150-157.

(责编:王平军;校对:王菊)