

第5章

内容安全基础

信息内容安全是一个日益受到重视并不断发展的领域。随着信息互联网的普及和技术的飞速发展,我们生活中产生的各种信息及其传播方式日益多样化,给信息内容安全带来了一些新的挑战。信息内容安全涉及多个研究领域,如多媒体信息处理、安全管理、计算机网络和网络应用等,直接和间接地应用各个领域的最新研究成果,为信息内容的安全提供保障。随着信息技术的进步,我们生活中的信息内容已经不再局限于传统的文字、图像和声音,还包括视频、虚拟现实等多媒体形式,这种多样形式使得信息内容面临着更多的风险和威胁,如数据泄露、网络攻击、恶意软件、网络水军、网络暴力和虚假信息等。本章将对信息内容安全基础做一个全面介绍。从信息内容安全的概述、网络舆情与事件检测、网络欺诈与虚假新闻、网络水军与网络暴力以及内容伪造与深度鉴别等角度展开介绍。

5.1

内容安全概述



内容安全概述

随着互联网的迅猛发展,内容安全成为一个备受关注的问题。互联网的出现给人们带来了前所未有的信息自由和便利,互联网与各个行业的融合也日益加深,创造了巨大的经济效益和社会效益,已经成为人们获取信息、互相交流、协同工作的重要途径,并具有应用极为广泛、发展规模最大、贴近人们生活等众多优点。同时,互联网也带来了一些负面影响,如网络欺凌、色情、暴力、恐怖主义等不良内容的泛滥,这些不良内容不仅对个人的心理健康和社会稳定造成威胁,也对企业的声誉和利益构成了危险。与此同时,随着大模型的快速发展,大模型存在的内容安全隐患和风险也越来越明显。大模型可以用于自然语言处理、图像识别、语音生成等众多领域,它们可以生成各种类型的内容,包括文字、图像、音频等,但大模型并不具备判断内容是否合适或符合道德的能力,因此可能导致生成虚假的信息或侵犯他人的知识产权,也可能存在滥用或泄露个人隐私数据的行为。不法分子常使用大模型进行欺诈,严重危害个人及社会的安全。因此,网络信息内容的安全值得广泛关注和深入研究。信息内容安全问题是当今社会面临的严峻挑战之一,这些问题涵盖了国家和社会层面以及组织和个人层面,对社会稳定、个人权益和经济发展产生着重要影响,典型的信息内容安全问题如图 5-1 所示。信息内容安全威胁分为两个层面:国家和社会层面、组织和个人层面。

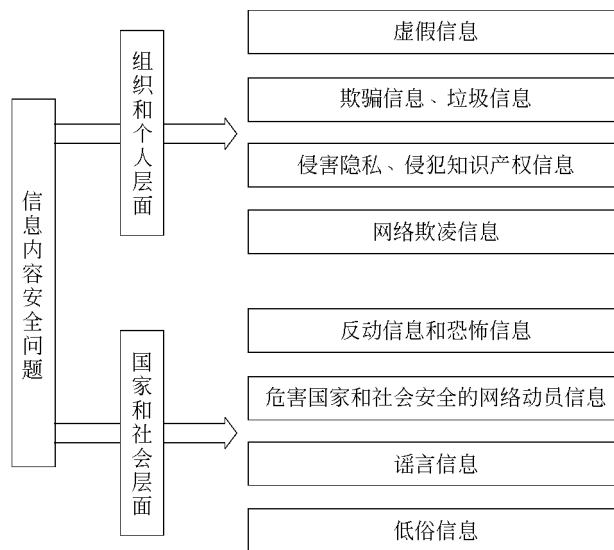


图 5-1 信息内容安全问题

在国家和社会层面上,信息内容安全威胁的种类多样。例如,反动信息可能通过网络传播,用于组织和策划犯罪活动、网络攻击、网络渗透等行为,对国家的政治稳定和社会和谐构成潜在威胁。谣言信息也是一个重要问题,虚假的信息传播可能导致社会恐慌、舆论混乱和社会不稳定。恐怖信息和低俗信息也会对社会产生负面影响,破坏社会道德和价值观。在组织和个人层面上,信息内容安全问题同样严重。虚假信息是一个突出的威胁,包括虚假广告、虚假新闻和虚假宣传等,这些虚假信息可能误导消费者、损害企业信誉,甚至对人们的生命财产安全产生重大影响。欺骗信息也是一个重要问题,如网络诈骗、网络钓鱼等,这些欺骗信息以欺骗手段获取个人敏感信息或财产,给个人和组织带来经济损失和信任危机。垃圾信息的泛滥也给人们的信息获取和处理带来困扰,增加了信息过滤和管理的难度。侵害隐私信息是个人和组织都面临的威胁,如个人隐私泄露、监控和监听等侵犯行为。网络欺凌信息对个人尤为严重,给受害者带来心理创伤和社交困扰。侵犯知识产权信息是对创新和知识产权保护的威胁,可能导致技术窃取、盗版和侵权行为的增加。面对这些威胁,我们需要采取一定的方法来确保信息内容的安全。在国家和社会层面,政府应加强信息监管和法制建设,制定相关法律法规,并加强对反动信息、谣言和恐怖信息的监测和打击,应用先进的技术手段,如人工智能和大数据分析技术,自动检测和过滤有害信息,并建立信息安全监测系统,及时发现、阻止和处理威胁行为。在组织和个人层面,组织应加强信息安全管理,包括加强员工培训、建立安全审查机制和加强网络防护措施。个人应提高信息素养,警惕虚假信息和网络欺凌,采用安全认证机制和加密技术,确保个人信息传输的机密性和完整性。

全球信息化的今天,互联网已经成为人们生活的重要组成部分,对社会、经济和文化产生了深远的影响。互联网的发展呈现出一些明显的趋势和特点,包括开放性、异构性、移动性、动态性和并发性。与此同时,互联网的发展还催生了一系列新型网络形式和服务模式,例如,下一代互联网的研究和探索正在进行中,旨在进一步提升网络的性能和功能。

5G 移动通信网络的商用化已经开启了新的时代,为更快速、可靠的移动互联网连接打下了基础。物联网的兴起将各种物理设备和传感器连接到互联网,实现智能化和自动化的应用。同时,云计算作为一种新的服务模式,提供了基于互联网的存储、计算和应用服务,为个人和企业提供了灵活、可扩展的解决方案。然而,互联网和新兴媒体的迅速发展带来的负面影响需要我们加强重视,例如,大量未经证实或带有偏见的信息在网络上过度泛滥,网络诈骗、网络钓鱼、网络暴力等行为出现,以及利用互联网和新兴媒体传播盗版,对知识产权进行侵犯。因此,在以信息内容为中心的互联网环境下,网络信息内容的安全值得广泛关注和深入研究。

5.1.1 定义

内容安全是信息安全领域中一个至关重要的分支,不同于网络安全、数据安全等,内容安全专注于研究和保护信息的内容,分析并确保内容的合法性、健康性和安全性,保护用户免受有害和不良内容的影响。内容安全的职责包括确保内容符合法律法规、版权规定和其他相关法律要求,涉及上传内容的审核和筛查,防止盗版、侵权、非法内容等问题。内容安全也致力于保护用户的心理和身体健康。内容安全通过内容过滤和审核机制,限制色情、暴力、赌博,以及虚假信息等不良内容的传播。为了保护用户免受恶意内容的侵害,内容安全措施包括识别和阻止恶意代码、网络攻击、网络钓鱼、网络欺诈、网络水军、内容伪造以及采取适当的安全防护措施,确保用户的信息和设备安全。内容安全涵盖了各种类型的信息,包括文本、图像、音频、视频等。通过综合运用技术手段和人工审核机制,内容安全致力于维护一个安全、健康、积极的网络环境。

5.1.2 目的

随着信息技术的发展,我们不可避免地面临着不良信息在网络上迅速扩散的问题,例如,暴力血腥、色情、邪教、赌博等非法信息,这些不良信息不仅严重污染网络环境,而且对社会公共安全、国家安全构成了威胁。内容安全的目的是保护用户免受有害和不良内容的影响,创建一个安全、健康和可信赖的数字媒体环境。信息内容安全可分为两个主要方面:首先是对信息内容的保护,信息内容的保护是确保信息在传输、存储和处理过程中不受未经授权访问、篡改、泄露和破坏的影响,确保内容的保密性、完整性和可用性;其次,信息内容必须符合政治、法律及道德的相关要求。信息内容应符合国家或地区的政治要求,不得传播违法、危害国家安全、煽动暴力或恐怖主义等违反政治准则的信息;信息内容应符合法律法规的要求,不得侵犯他人的知识产权、隐私权,不得传播诽谤、淫秽、暴力等违反法律规定的信息;信息内容应符合社会道德准则,不得传播低俗、歧视、仇恨等违背道德规范的信息,尊重他人的尊严和权益。

5.1.3 重要性

信息内容安全对个人、国家和社会都具有重要意义,事关国家安全、公共安全、文化安全等。对于个人而言,内容安全直接关系到个人的隐私和权益保护。在数字化时代,个人的大量信息被存储、传输和共享,包括个人身份信息、财务数据、健康记录等敏感信息。如

果个人信息受到未经授权的访问、滥用或泄露,个人可能面临身份盗窃、财务损失以及个人声誉受损等风险。因此,内容安全对于个人而言是确保个人隐私和个人安全的重要保障。对于企业而言,内容安全关乎商业机密的保护和商业运作的可持续发展。现代企业依赖于信息技术存储和处理大量的商业数据、客户信息和研发成果,如果企业信息受到未经授权的访问、窃取或篡改,企业可能面临商业机密泄露、品牌声誉受损、经济损失甚至企业倒闭的风险。因此,保护企业的信息内容安全对于企业的持续经营和成功至关重要。对于国家和社会而言,内容安全事关国家安全、社会稳定和公民权益。保护国家重要信息和基础设施的安全是国家的战略需要,涉及国家安全、国防、金融稳定等重大领域。信息泄露或被篡改可能导致国家机密暴露,国家安全受到威胁,甚至对社会秩序和公民权益产生负面影响。除此之外,内容安全还涉及文化安全的保护。文化是一个国家和民族的瑰宝,包括文学、艺术、历史、传统等方方面面。如果不良内容泛滥,可能对文化价值观和社会道德造成冲击,甚至导致文化多样性的丧失。因此,内容安全对于保护和传承文化遗产、维护文化多样性和社会和谐至关重要。总之,内容安全对于维护社会的稳定、保障公民权益和促进可持续发展具有重要意义,需要我们共同努力保障。相关部门通过加强内容安全措施、提高用户意识和建立法律法规,共同创建一个安全、健康和可信赖的数字媒体环境,促进社会的繁荣与进步。

5.1.4 大模型时代的内容安全



大模型时代的内容安全

随着大模型的发展和普及,内容安全问题也凸显出来。大模型的强大能力使其能够生成高质量的文本,但也带来了一系列潜在的风险和挑战。尽管大部分情况下大模型能够产生准确和有用的信息,然而当大模型使用不当时,它会散播一些虚假或令人担忧的内容,从而导致网络欺诈、网络暴力、仇恨言论等问题的加剧,对个人、群体和社会造成伤害。

1. 大模型概述

大模型是一种参数规模庞大且复杂的机器学习模型,具有强大的学习和推理能力,能够解决更复杂的问题。在深度学习领域,大模型通常是指具有数百万到数十亿参数的神经网络模型。神经网络模型需要大量的计算资源和存储空间来进行训练和存储,并且通常需要使用分布式计算和特殊的硬件加速技术实现高效计算。大模型在处理复杂任务时通常有更好的性能和准确度,但需要更多的计算资源和硬件支持。大模型在自然语言处理、图像识别、语音识别等多个领域都有广泛的应用,如图 5-2 所示。例如,在自然语言处理领域,大模型可以用于机器翻译,将一种语言的文本翻译成另一种语言;用于文本生成,如生成新闻文章或消息;用于情感分析,分析文本中的情感倾向,如判断一篇文章是积极的还是消极的。与小模型相比,大模型的优势在于更高的计算效率和更强的处理能力,但同时需要更多的计算资源和存储需求,并且整个过程对外来观察者来说是一个“黑盒”,我们无法准确了解其内部学习的知识和能力体系及其运行的具体规律。因此,我们无法确保模型会按照我们的意图和目标进行工作,并且无法解释其决策的依据和过程,可能会引发一些法律、伦理和责任方面的争议和纠纷,也会带来一定的安全风险。

2. 大模型的使用风险

尽管大模型在应用中具有巨大潜力,但也存在一些问题。其中一个问题是大模型可

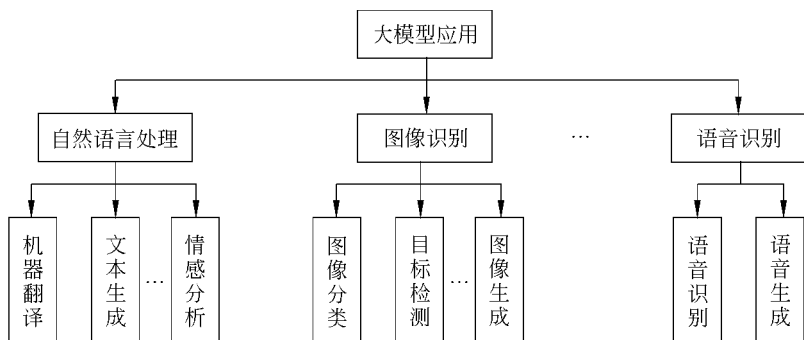


图 5-2 大模型应用示例

能会产生与人类价值观不一致的输出,例如,模型可能会生成不准确、不完整或误导性的信息,可能会误导用户或传播错误的观点,而且大模型也可以被恶意用户滥用来生成虚假信息、恶意内容或进行网络钓鱼等活动,对个人、组织或社会造成危害。大模型的安全风险在文本、图像、语音和视频等多个应用场景中广泛存在,并随着模型的大规模部署而日益严重,导致用户难以信任人工智能系统所做出的决策。更重要的是,大模型在安全风险防范方面相对脆弱,容易受到指令攻击、提示注入和后门攻击等恶意攻击的影响。

此外,在大模型技术的蓬勃发展下,越来越多的企业利用大模型来实现业务赋能。企业通过记录客户的个人资料、购买记录、行为习惯等信息,并运用先进的分析和挖掘技术,为客户提供个性化的商品推送和服务体验。然而,我们必须认识到,大模型技术的不当使用,可能会导致过度收集和违规使用信息数据,导致个人隐私被侵犯。

随着大模型在各个领域的广泛应用,大模型的内容安全风险的范围逐渐扩大,对社会秩序造成的冲击变得更为严重。大模型自身存在着一些风险,下面列举了几种风险类型。

- 虚假信息:大模型可以被不法分子滥用,生成伪造内容、虚假信息、敏感话题或进行网络水军等活动,搅乱社交媒体平台的秩序或破坏公共讨论的健康环境。
- 网络欺诈:大模型的能力可以被黑客或网络犯罪分子利用,进行网络欺诈等行为。例如,犯罪分子可以利用大模型生成逼真的虚假身份、仿冒网站或欺诈性广告,诱导用户泄露个人信息,产生经济损失。
- 数据泄露:大模型通常需要大量的数据来进行训练和优化,其中可能包含用户的个人信息。如果个人信息数据未经适当保护,可能会导致用户隐私的泄露。
- 违法犯罪:大模型可以被恶意用户使用,进行电信诈骗、网络钓鱼等违法犯罪活动,对个人、组织或社会造成危害。
- 不准确或误导性的信息:大模型通过学习大量的文本数据,可能会生成不准确、不完整或误导性的信息,可能会误导用户或传播错误的观点。
- 辱骂仇恨:大模型可以生成一些带有辱骂、脏字脏话、仇恨言论等不当内容来进行网络暴力,从而扰乱社会稳定。
- 偏见和歧视:如果训练数据中存在偏见或歧视,大模型可能会学习到这些偏见并在生成文本时表现出来,可能导致生成的文本具有性别、种族、宗教等方面的偏见或歧视。

3. 大模型的内容安全问题

随着大模型技术的迅速发展和广泛应用,我们进入了一个数据驱动的时代。先进的算法和模型在提高我们生活质量的同时,也引发了一系列前所未有的安全问题。大模型算法本身可能存在缺陷,导致其生成虚假新闻或不正当言论。例如,大模型可以生成没有语法错误和自动翻译痕迹的电子邮件,不法分子利用大模型进行网络钓鱼和电子邮件诈骗,增加诈骗风险。大模型引发的内容安全问题已经成为当前人们关注的焦点,也是人工智能领域面临的挑战之一。应对大模型时代的内容安全问题,可以从以下三个方面展开。

1) 开发管理

在大模型开发阶段,相关人员可以对训练数据进行筛查清洗,同时需要考虑内容安全等因素,发现大模型存在的内容安全问题。例如,相关人员可以通过排除不当、冒犯性或违法的内容,确保数据的质量和安全性。相关人员可以制定合理的规则策略、权限机制来限制训练数据的获取和使用,确保只有经过审核和授权的数据源可以用于模型训练。相关人员也可以引入内容过滤的模块或机制,根据预先设定的规则 and 标准,对生成内容进行即时检测和过滤,减少不当或冒犯性内容的生成,确保模型生成的内容符合政治、法律和道德准则。

2) 攻防技术

相关人员可以使用攻防技术模拟真实攻击行为,发现大模型中的漏洞和弱点,使大模型具备更强的鲁棒性和抵抗能力,应对各种恶意攻击,防止用户数据被泄露、滥用或篡改,维护用户的隐私和个人安全。例如,相关人员可以使用加密技术保护隐私信息,通过数据加密和隐私保护技术,使用户的敏感信息在传输和存储过程中得到保护,防止未经授权的访问和数据泄露,维护用户的隐私权和个人信息安全,增加用户对大模型的信任度,促进用户参与和数据共享。相关人员也可以通过对抗样本生成和对抗性训练,减少对恶意内容的生成和传播,防止模型被攻击者利用进行不良行为,提高模型对恶意输入的识别能力和鲁棒性。

3) 监测报警

相关部门可以提供实时的监测和报警系统,及时发现并处理存在的内容安全问题及漏洞。例如,相关部门可建立实时监控机制,对模型生成的内容进行持续监测和实时分析,通过设定阈值、检测异常模式或使用机器学习算法,及时识别并报警不当或冒犯性内容的出现。同时,相关部门应鼓励用户提供反馈和举报不当或冒犯性内容,建立用户反馈和举报机制,及时响应用户的反馈并采取相应的措施。

5.2

网络舆情与事件检测

随着互联网的普及和社交媒体的兴起,网络舆情事件的规模和影响力不断增大,可能会产生一些负面影响。因此,网络舆情事件检测成为网络信息内容安全管理的重要领域之一。网络舆情事件通常在社交媒体平台上展开,在这些平台上发布的文字、图片和视频等信息内容被关注并广泛传播,使得舆情事件具有了传播速度快和影响范围大的特点。

网络舆情事件往往起源于社交媒体平台上用户发布的内容,一条简短的推文、一张照片或一个视频片段都有可能引发广泛的讨论和关注。当话题或事件引起了足够多的关注和讨论时,这个话题就可以被认定为一个网络舆情事件。网络舆情事件可能对个人、组织或社会产生重大影响,引发公众对某个议题的关注,改变人们的态度和行为。某些网络舆情事件会对相关个体或组织造成名誉损害,甚至会引发社会动荡造成矛盾激化的局面,对社会稳定和秩序构成威胁。为了应对上述负面情况,及时发现和处理网络舆情事件成为一个重要的任务。为此,网络舆情检测技术应运而生,这类技术通过监控和分析社交媒体上用户发表的内容,可以及时发现潜在的网络舆情事件。

5.2.1 网络舆情事件概述

网络舆情是在互联网形成,通过网络传播的群体舆论情绪和言论倾向。网络舆情事件是一种引起广泛关注和讨论,并与特定主题相关的网络舆情现象,它们通常由大量网民在社交媒体、论坛、新闻评论等网络平台上发布的言论、观点和情感所构成,表达了群众的观点、立场和感受。网络舆情事件通常涉及一些重大、敏感或具有争议性的社会事件,可以涉及各种主题,包括社会、政治、商业、科技等。网络舆情事件与传统媒体的舆论事件不同,传统媒体的舆论事件由少数媒体机构操控和引导,而网络舆情事件是由广大网民自发产生和参与。网络舆情事件的主要形式体现在社交媒体的评论、讨论、转发、点赞等交互行为,以及社交媒体的帖子、博客、新闻评论等内容。相比传统舆情事件来说,网络舆情事件更便捷,传播速度更快。

网络舆情事件具有突发性、多样性、实时性等特点。网络舆情事件的突发性是指,网络舆情事件可能源于某一具体事件、言论或行为,能够在短时间内迅速传播扩散,迅速引发公众的关注和参与。网络舆情事件的多样性是指,网络舆情事件涉及的话题和讨论内容多种多样,涵盖社会、政治、经济、文化等领域。网络舆情事件的实时性是指,网络舆情事件受到网络的即时性的影响,使得舆论变化迅速,在短时间内引发公众剧烈的情绪波动。

网络舆情事件对个人、企业、政府和公众等各方面都产生了重大的影响。在个人方面,网络舆情事件可以对个人的声誉和形象产生直接影响,负面的网络舆情可能导致个人声誉受损,社交关系、职业发展受到影响,而正面的网络舆情则可以提升个人的知名度和影响力。在企业方面,网络舆情事件可以对企业的形象和声誉造成直接影响。负面舆论和批评可能导致公众对企业产生怀疑和不信任,进而影响企业品牌价值和市场地位。而正面舆论可以增强企业的声誉和受欢迎程度。网络舆情事件对政府决策也具有一定的影响力。政府需要密切关注舆情事件发展动向,及时了解公众的关切和诉求,以便能够及时做出决策调整。同时,政府也需要通过积极参与和引导舆情方向,及时回应公众的关切和批评,维护社会稳定和公众满意度。网络舆情事件对公众的影响也非常直接和深远,公众可以参与社会热点事件和公共事务,表达自己的意见和观点。然而,舆情事件也可能带来信息过载和虚假信息的传播风险,公众需要具备批判性思维,正确理解和评估舆情事件内容的真实性。



网络舆情
事件概述

网络舆情
事件检测
技术

5.2.2 网络舆情事件检测技术

随着互联网的普及和社交媒体的兴起,网络舆情事件的发生频率和影响力越来越大。网络舆情事件对个人、组织甚至整个社会都具有重要的影响和挑战。因此,网络舆情事件检测技术应运而生。网络舆情事件检测技术旨在通过自动化方法和机器学习等算法,实时监测和分析涉及舆情的信息,发现和追踪潜在的事件。网络舆情事件检测技术可以帮助我们快速捕捉到用户的声音和关注点,及时了解和评估网络舆情事件的规模、趋势和情感倾向,从而及时采取适当的措施应对和管理不良事件,提前预警潜在的危机和风险,采取相应的措施进行危机管理和舆情引导,以最小化负面影响。网络舆情事件检测方法可以按照不同的技术进行划分,几种常见的类型包括基于传统机器学习的事件检测方法、基于深度学习的事件检测方法和基于图神经网络的事件检测方法。

1. 基于传统机器学习的事件检测方法

在深度学习时代到来之前,人们主要依赖传统机器学习方法进行事件检测。基于传统机器学习的事件检测方法主要通过分析文本特征,如关键词、情感倾向、主题等,来识别和分类网络舆情事件。聚类技术是传统机器学习事件检测中常用的方法,可以通过将文本分成不同的分区并判断文本是否属于某个特定的事件。例如,我们将网络舆情文本作为输入数据,将聚类技术应用在事件检测中,聚类算法会根据文本之间的相似性将其分配到不同的聚类簇中,算法通过观察聚类簇的内容,可以快速理解和发现关于特定事件的讨论和观点。

在事件检测中信息内容经常具有不确定性,这是因为舆情信息存在大量噪声和错误单词且用户位置可能不是实际位置。为了确定内容、时间和位置等信息数据,相关研究人员提出了一个位置-时间约束主题模型,如图 5-3 所示^[1]。该模型将每条消息的时间和位置视为附加变量,将位置经度、纬度和时间变量附加到每个令牌上,通过这种方式将不确定的时间、地点和社会内容信息融合在一起。位置-时间约束主题模型使用 KL 散度来度量内容相似性,并使用最长公共子序列度量消息对的链接相似性,将内容相似性与链接相似性集合起来,从而在全局上形式化信息的相似性。

然而,基于传统机器学习的事件检测方法在面对大规模数据和复杂模式时可能存在一些限制,随着深度学习的兴起,人们开始探索使用神经网络来进行事件检测和分类。

2. 基于深度学习的事件检测方法

舆情数据通常具有复杂的语义和结构,传统的手动特征提取方法可能无法捕捉到高度抽象的特征,而深度学习模型可以捕获多层信息,通过多层非线性变换来提取更高级别的语义信息。另外,舆情事件检测通常不仅依赖于文本数据,还可能包含图像、视频、音频等多模态数据,深度学习模型可以灵活地处理和融合多模态数据,提高事件检测的准确性和全面性。常见的深度学习模型方法包括卷积神经网络(Convolutional Neural Networks, CNN)、循环神经网络(Recurrent Neural Network, RNN)、注意力机制、长短期记忆网络(Long Short-Term Memory, LSTM)等,这些深度学习模型方法可以根据任务需求和数据特点进行选择和优化。

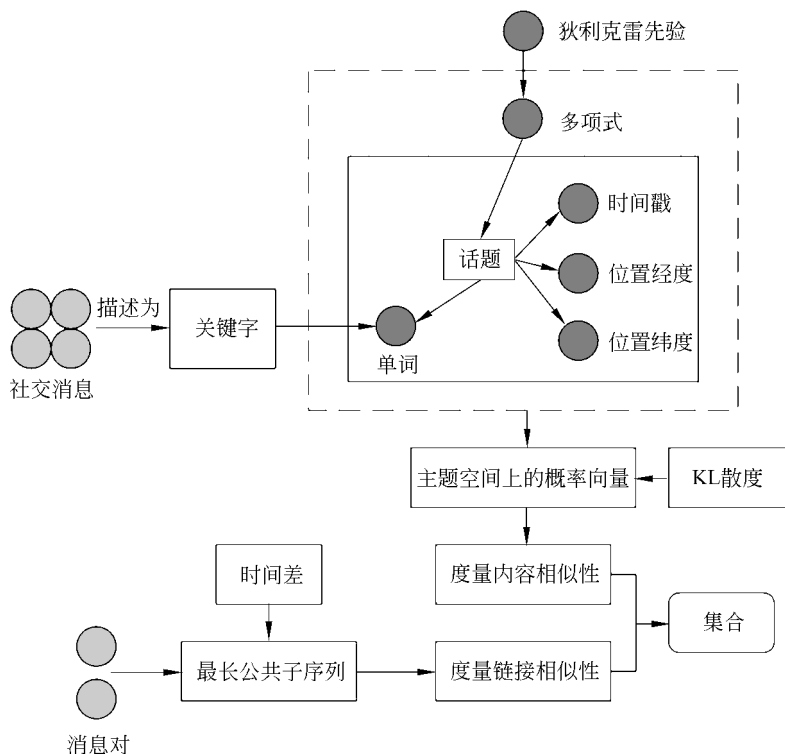


图 5-3 位置-时间约束主题模型架构图

传统的机器学习方法在事件检测方面存在一个挑战,即很难准确地表示事件。然而,基于深度学习的事件检测方法能够有效地解决这个问题。相关研究人员为解决这类问题,针对事件表征问题提出了一种上下文分级的 LSTM 事件表征方法,如图 5-4 所示^[2]。上下文分级的 LSTM 事件表征方法可以捕获多个序列结构,即词级序列和事件级序列,并采用两级层次 LSTM 架构对事件进行增强表示,从而提高舆情事件检测工作的准确

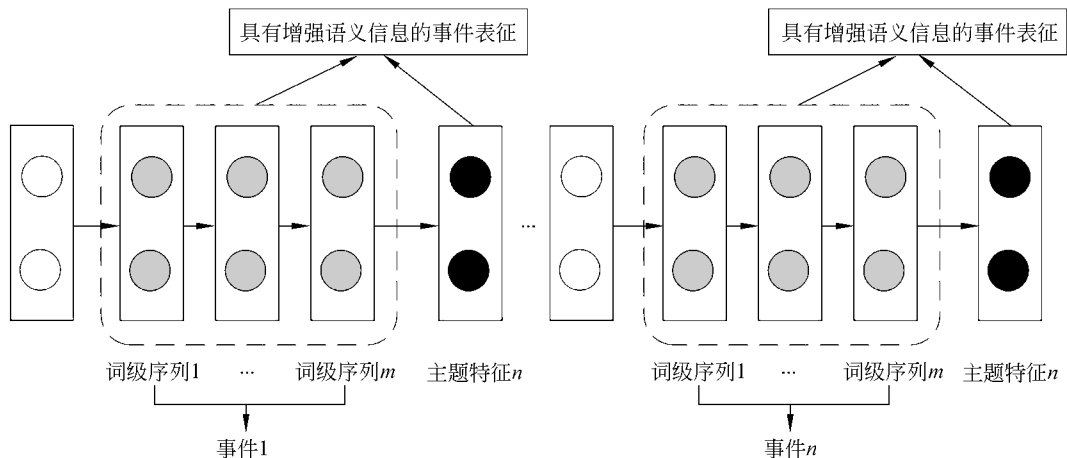


图 5-4 上下文分级的 LSTM 事件表征方法架构图

度。具体而言,第一级层次 LSTM 是对事件进行编码,即词级序列,然后将每个事件映射到嵌入中;第二级层次是对观察到的事件序列进行编码,并与上下文主题特征嵌入向量结合,增强语义信息,从而更好地表征事件的语义关联。

3. 基于图神经网络的事件检测方法

传统的事件检测方法通常面临数据稀疏、特征提取困难和模型泛化能力不足等挑战。而图神经网络技术以图形的形式来表示数据,能够有效地捕捉数据之间的复杂关系,具备强大的表征学习和泛化性能。基于图神经网络的事件检测方法利用图结构的信息学习节点的特征能力,挖掘数据中的隐藏关系和模式,从而实现对事件的准确检测和跟踪。从技术划分上,包括无监督的图神经网络模型、半监督的图神经网络模型以及有监督的图神经网络模型。

在舆情事件检测方面,尽管图神经网络已经取得了一定的成果,但仍然存在一些缺陷,例如,消息建模和长尾分布问题等。针对上述问题,已经有相关的研究工作进行探索 and 解决,其中一种比较先进的图事件检测模型旨在重新定义消息建模方法,并针对长尾分布数据采取相应策略,如图 5-5 所示^[3]。这种图事件检测方法的核心思想是将原始消息建模成具有多元关系的加权异构图,以保存更丰富的连接信息。为了解决消息建模的问题,该图事件检测方法利用多智能体强化学习算法来选择不同关系的邻居节点,并使用邻居聚合方法学习消息的最佳嵌入表示,以更好地表达消息之间的关系和语义信息。为了应对长尾分布数据的问题,该方法采用了平衡采样策略和对比学习机制,对消息表示学习进行增量处理,以克服样本数据量不足导致模型训练不充分的问题。此外,这种方法还引入了深度强化学习引导的 DBSCAN 模型优化事件的聚类效果,包括所需的数量和距离参数,能够更好地发现和识别事件,提高聚类的准确性和效果。

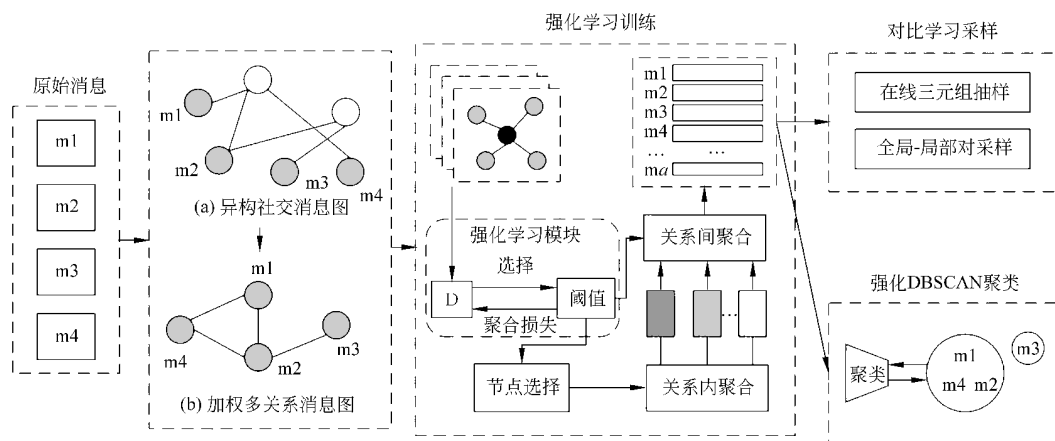


图 5-5 基于图结构的强化聚合事件检测模型架构图

5.2.3 网络舆情事件检测的应用

网络舆情监测与预警系统是网络舆情事件检测技术的典型应用之一。针对互联网媒

