

2025 年国家网络安全宣传周 “网络安全技术高峰论坛主论坛” 在昆明举行

9月15日,2025年国家网络安全宣传周“网络安全技术高峰论坛主论坛”在云南省昆明市滇池国际会展中心举行。中央网信办副主任、国家网信办副主任杨建文,云南省省委常委、昆明市委书记刘洪建出席活动并致辞,云南省政协党组成员、副主席王以志主持活动。

杨建文指出,当前,新一轮科技革命和产业变革加速演进,以人工智能、卫星互联网、量子科技等为代表的网络信息技术快速迭代,在驱动引领经济社会发展的同时,也带来了网络系统安全、内容安全、数据安全风险的交织叠加。面对新形势、新变化、新要求,迫切需要构建适应大网络安全工作格局的网络安全保障能力。一是构建纵深布防、高效协同的国家网络安全防御能力,实现对

各类重大网络安全威胁的“一点发现、全网防御”。二是构建包容审慎、双向赋能的网络安全技术治理能力,为新兴未来产业健康发展开辟良好赛道。三是构建自主创新、开放合作的网络安全产业基础能力,让全球搭乘我国互联网发展“快车”。四是构建体系完备、产教融合的网络安全人才培育能力,培养带动更多的高水平人才。

刘洪建代表昆明市委、市政府对莅临主论坛的各位嘉宾表示欢迎,对大家长期以来给予昆明的支持和帮助表示感谢。他说,昆明作为我国面向南亚东南亚开放的重要门户、“数字云南”建设的核心枢纽,既是数字经济发展的热土,也是网络安全防护的前沿。近年来,昆明市认真贯彻落实习近平总书记关于网

络强国的重要思想,在网络空间安全领域积极探索、主动作为,基础设施升级完善,防护能力整体跃升,数字引擎动能澎湃,通信枢纽链接全球,网络综合治理效能持续提升。昆明市将以此次论坛为契机,充分吸收好、运用好论坛形成的各项成果,进一步完善网络安全防护体系,加快打造面向南亚东南亚的网络安全产业高地,为守护好国家网络空间的西南通道、谱写网络强国云南篇章作出昆明贡献。

论坛上,中国工程院院士吴世忠、信息安全管理国际工作组创始人爱德华·汉弗莱斯(Edward·Humphreys),公安部十一局负责同志、香港数字政策专员、部分高校和企业负责人等分别围绕“人工智能安全治理国

际合作”、“信息安全标准国际合作实践与思考”等主题作主旨演讲,分享经验成果、碰撞智慧火花。

论坛上集中发布一批网络安全重要成果,全国网络安全标准化技术委员会发布《人工智能安全治理框架》2.0版、《政务大模型应用安全规范》,国家计算机网络应急技术处理协调中心、国家工业信息安全发展研究中心分别作了解读,国家计算机网络应急技术处理协调中心发布2025年人工智能技术赋能网络安全应用测试结果。

本次高峰论坛由中央宣传部、中央网信办等十部门联合主办,相关政府部门负责同志、知名专家学者、企业机构代表和媒体记者等近千人参加了活动。

据昆明发布

国内首次 AI 大模型安全“体检”结果出炉 15 款大模型及应用产品接受测试

都市时报全媒体记者 郑荣行

9月16日,在第22届中国网络安全年会(暨国家网络安全宣传周网络安全协同防御分论坛)上,国内首次针对人工智能大模型的实网众测检验结果正式发布,大模型特有漏洞177个,占比超过60%。

本次活动累计发现 各类安全漏洞 281 个

此次活动由中央网信办网络安全协调局指导,国家计算机网络应急技术处理协调中心主办,相关平台和基地提供技术支持。活动发布了国内首次针对 AI 大模型的实网众测检验结果,共动员559名白帽子,对国内15款大模型及应用产品进行了全面漏洞测试。测试产品涵盖基础大模型、垂域大模型、智能体、模型开发平台等多种类型,既包含单模态大模型,也涵盖多模态大模型,具有广泛的代表性。

从结果来看,腾讯混元大模型、百度文心一言、阿里巴巴通义 App、智谱华章智谱清言等产品发现的漏洞风险较少,展现出较高安全防护水平。活动还涌现出梁宏宇、马宝新、刘铭等一批优秀白帽子,他们涵盖网安企业人才、科研院校师生及社会人士,从外部攻击者视角完成多维度实战化测试,即便在厂商防护下仍挖掘出高危及严重等级漏洞,彰显专业能力与奉献精神。

值得关注的是,本次活动累计发现



众测结果发布现场

供图

各类安全漏洞281个,其中,大模型特有漏洞177个,占比超60%,凸显AI大模型面临传统安全领域之外的新兴风险。典型漏洞风险包括部分产品不当输出类漏洞危害严重、信息泄露类漏洞多发、提示注入类漏洞普遍、无限制消耗类攻击防护不足及传统安全漏洞仍存等。鉴于AI大模型用户量大、使用率高,此类漏洞

若被恶意利用,将严重影响国内AI大模型产品生态。

建议加快制定 AI 漏洞分类分级标准

对此,活动明确AI大模型安全治理需各方发力,持续加强安全防护措施应对各

类漏洞风险,加快制定AI漏洞分类分级及场景风险等级标准,通过众测模式汇聚社会白帽力量构筑共治生态,加强内生安全治理,并将安全融入AI系统全生命周期。

业内人士指出,随着人工智能应用场景不断拓展,新的风险和攻击方式仍将持续出现,AI大模型安全治理工作任务重道远,亟需各方共同努力,从源头筑牢安全基石。