



全国网络安全标准化技术委员会
National Technical Committee 260 on Cybersecurity of SAC

TC260-PG-20268A

网络安全标准实践指南

——人工智能应用安全指引 总则

(v1.0-202609)

全国网络安全标准化技术委员会秘书处

2026 年 09 月

本文档可从以下网址获得：

www.tc260.org.cn/



全国网络安全标准化技术委员会
National Technical Committee 260 on Cybersecurity of SAC



前 言

《网络安全标准实践指南》（以下简称《实践指南》）是全国网络安全标准化技术委员会（以下简称“网安标委”）秘书处组织制定和发布的标准相关技术文件，旨在围绕网络安全法律法规政策、标准、网络安全热点和事件等主题，宣传网络安全相关标准及知识，提供标准化实践指引。

本文件起草单位：中国电子技术标准化研究院、国家广播电视总局广播电视科学研究院、中国交通通信信息中心、国家卫生健康委统计信息中心、教育部教育管理信息中心、应急管理部大数据中心、中国工商银行、中国农业银行、北京中关村实验室、浦江国家实验室、国家工业信息安全发展研究中心、国家计算机网络应急技术处理协调中心、中央网信办数据与技术保障中心、中国网络空间研究院、工业和信息化部电子第五研究所、中国信息安全测评中心、国家信息技术安全研究中心、北京火山引擎科技有限公司、北京快手科技有限公司、阿里云计算有限公司、华为技术有限公司、小米科技有限责任公司、蚂蚁科技集团股份有限公司等。

本文件主要起草人：姚相振、郝春亮、张妍婷、赵韡、王磊、杨伟平、贺敏、王志伟、张震、宋首友、陈朴、李岳峰、张卫伟、吕飞霄、李子威、宋昊、郭建领、谷晨、彭骏涛、邵萌、赵冉、王博、吴巍、申东洋、吴波、李琦、马梦娜、刘栋、孟令宇、刘北水、王盈、杜渐、戴明、落红卫、初翎祯、林冠辰、卢春景、郭晓霞、程佩哲、许啸、张荣、柳嘉琪、李寒雨、李申、谭龙、吴子坚、田小龙、曹岳、吕莹楠、王立夫、于新颖、黄冬秋。



声 明

本《实践指南》版权属于网安标委秘书处，未经秘书处书面授权，不得以任何方式抄袭、翻译《实践指南》的任何部分。凡转载或引用本《实践指南》的观点、数据，请注明“来源：全国网络安全标准化技术委员会秘书处”。



摘 要

以高水平人工智能安全标准保障高质量发展，全面提升各行业领域人工智能应用安全水平，依据《中华人民共和国网络安全法》《中华人民共和国数据安全法》《中华人民共和国个人信息保护法》等法律，以及《生成式人工智能服务管理暂行办法》《人工智能生成合成内容标识办法》《互联网信息服务算法推荐管理规定》《互联网信息服务深度合成管理规定》等规定，制定人工智能应用安全指引系列实践指南。

人工智能应用安全指引系列文件包含通用文件及行业领域文件两类。通用文件提供通用性安全指导，适用于各行业领域开展人工智能应用活动。行业领域文件给出适用于特定行业领域的实践指引。

本《实践指南》属于通用文件。行业领域开展人工智能应用活动，在符合本《实践指南》的基础上，还应进一步满足行业领域文件相关内容。



目 录

1 范围	1
2 规范性引用文件	1
3 术语和定义	2
4 缩略语	3
5 概述	3
6 基本原则	4
7 人工智能应用通用安全指引	5
7.1 前期规划阶段	5
7.2 设计开发阶段	6
7.3 验证确认阶段	7
7.4 部署阶段	9
7.5 运行和监控阶段	10
7.6 持续验证评估阶段	12
7.7 退役下线阶段	12
7.8 其他	13
附 录 A（规范性） 术语表	15
附 录 B（规范性） 缩略语表	18
参 考 文 献	19





1 范围

本文件规定了人工智能应用安全总则，包括基本原则以及人工智能应用的前期规划、设计开发、验证确认、部署、运行和监控、持续验证评估、退役下线等各阶段的通用安全指引。

本文件适用于各行业领域组织开展人工智能应用的安全风险防范与管理，可为有关监管部门和第三方测试评估机构提供参考。

2 规范性引用文件

下列文件中的内容通过本文件中的规范性引用而构成本文件必不可少的条款。其中，注日期的引用文件，仅该日期对应的版本适用于本文件；不注日期的引用文件，其最新版本（包括所有的修改单）适用于本文件。

GB/T 25069 信息安全技术 术语

GB/T 41867 信息技术 人工智能 术语

GB/T 42888-2023 信息安全技术 机器学习算法安全评估规范

GB 45438-2025 网络安全技术 人工智能生成合成内容标识方法

GB/T 45652-2025 网络安全技术 生成式人工智能预训练和优化训练数据安全规范

GB/T 45654-2025 网络安全技术 生成式人工智能服务安全基本要求

GB/T 45674-2025 网络安全技术 生成式人工智能数据标注安全



规范

GB/T 45958-2025 网络安全技术 人工智能计算平台安全框架

TC260-005 人工智能应用伦理安全指引 1.0

TC260-PG-202515A 网络安全标准实践指南——生成式人工智能服务安全应急响应指南

TC260-PG-20261A 网络安全标准实践指南——人工智能加速芯片安全技术规范

TC260-PG-20267A 网络安全标准实践指南——个人用户使用人工智能服务安全指南

TC260-PG-20266A 网络安全标准实践指南——智能体部署使用安全指引

3 术语和定义

GB/T 25069 和 GB/T 41867 以及本文件附录 A 中界定的术语和定义适用于本文件。

3.1 人工智能 artificial intelligence; AI

〈学科〉人工智能系统(B.2)相关机制和应用的研究和开发。

[来源: GB/T 41867—2022, 3.1.2]

3.2 人工智能系统 artificial intelligence system

针对人类定义的给定目标,产生诸如内容、预测、推荐或决策等输出的一类工程系统。

注 1: 该工程系统使用人工智能(B.1)相关的多种技术和方法,开发表征数据、知识、过程等的模型,用于执行任务。



注2：人工智能系统具备不同的自动化级别。

[来源：GB/T 41867—2022, 3.1.8]

3.3 人工智能应用 artificial intelligence application

具有功能特性的人工智能的使用，该人工智能在利益相关方场景中运行以实现预期结果。

[来源：ISO/IEC 5339:2024(en), 3.1]

4 缩略语

本文件附录 B 中的缩略语适用于本文件。

5 概述

本文件作为人工智能应用安全指引的总则，聚焦人工智能应用前期规划、设计开发、验证确认、部署、运行和监控、持续验证评估、退役下线全生命周期中的共性安全问题，提出人工智能应用安全基本原则，提供通用安全指引，为相关方开展人工智能应用安全工作提供参考。开展人工智能应用时，坚持造福人类、促进发展，以服务社会和可持续发展为导向，在开展应用前明确其社会价值，避免人工智能技术的盲目应用。

人工智能应用全过程通常涉及数据提供方、建设部署方、运行管理方、访问使用方等角色。数据提供方负责提供模型训练以及产品或服务运行过程中使用的数据；建设部署方负责创建或调整人工智能模型、算法、系统，开展相关设计开发、验证确认和部署；运行管理方



负责将人工智能产品或服务投入运行，并承担运维、监控等工作；访问使用方直接使用人工智能产品或服务，获得输出结果并基于输出结果实现特定目标。同一组织承担多个角色时，可结合实际活动统筹履行相应职责。

6 基本原则

人工智能应用安全基本原则包括：

a) **以人为本，可管可控。**坚持以人为本，确保人工智能最终控制权归属于人类，避免人工智能应用脱离人类控制。

b) **识别风险，分级保护。**面向人工智能应用全过程进行安全风险识别，对人工智能应用进行分级保护，实现资源的合理配置与精准的安全治理。

c) **提高能力，覆盖全程。**对人工智能安全能力成熟度进行评估，提高人工智能应用全过程的安全能力，并结合各阶段特点采取相应措施。

d) **加强测评，客观验证。**加强人工智能应用安全测评，客观验证人工智能应用安全状态及防护措施的有效性，为人工智能应用的监管提供依据。

e) **面向变化，动态调整。**建立动态适应人工智能技术迭代与应用风险变化的调整机制，更新技术与管理相关安全防护措施，提升人工智能应用的安全性、可靠性。



7 人工智能应用通用安全指引

7.1 前期规划阶段

前期规划安全指引包括：

- a) 综合分析人工智能应用对国家安全和经济社会发展的影响，充分考虑伦理安全要求，研判开展相应人工智能应用的必要性和合理性。

注：对伦理安全的影响参考 TC260-005《人工智能应用伦理安全指引 1.0》。

- b) 围绕人工智能技术内生安全风险、技术应用安全风险、以及应用衍生安全风险等方面，开展安全风险识别以及风险分类。

注：安全风险识别参考《人工智能安全治理框架》。

- c) 按照安全风险识别结果，逐条目分析风险发生概率以及风险影响严重程度，按照就高从严原则综合研判确定人工智能应用安全分级。

注：安全分级参考我国人工智能应用安全分类分级有关国家标准，分为低安全风险、一般安全风险、较大安全风险、重大安全风险、特别重大安全风险五级。

- d) 梳理人工智能应用前期规划、设计开发、验证确认、部署、运行和监控、持续验证评估、退役下线等阶段所涉及的建设部署方、运行管理方，对每个相关方分别按照其所涉及的阶段开展人工智能安全能力成熟度评估，判定成熟度级别。

注：能力成熟度评估参考我国人工智能安全能力成熟度评估有关国家标准。

- e) 根据人工智能应用安全分级结果，以及应用各阶段所涉及的建设部署方、运行管理方人工智能安全能力成熟度等级，针对每



个相关方开展人工智能安全能力提升或更换，使各单位人工智能安全能力成熟度等级均不低于应用安全分级。

- f) 结合人工智能应用场景、识别安全风险情况，以及建设部署方、运行管理方人工智能安全能力成熟度情况，确定人工智能应用安全措施，制定监测预警以及应急响应方案。

注：监测预警及应急响应方案制定参考 TC260-PG-202515A《网络安全标准实践指南——生成式人工智能服务安全应急响应指南》。

- g) 考虑人工智能应用对生命财产安全的影响，体现人类最终控制原则，规划一键接管、版本回退、紧急关停等安全措施，保障极端情况下干预止损的敏捷性。
- h) 明确人工智能应用安全的第一责任人，建立覆盖前期规划、设计开发、验证确认、部署、运行和监控、持续验证评估、退役下线等各阶段的安全责任机制，以及覆盖模型选型、数据源选取、部署方式选择等关键决策及关键操作记录，明确各相关方的安全责任边界和安全协同机制，支撑重大安全事件的责任追溯。

7.2 设计开发阶段

设计开发安全指引包括：

- a) 在算法规则、模型框架、系统架构等角度，提升人工智能技术可解释性、公平性、鲁棒性、可靠性、透明性、隐私保护、价值对齐等内生安全能力设计。



- b) 基于开源资源进行二次开发的，遵循相应开源协议规范，对所使用的开发框架、代码等进行安全审计，关注开源框架安全及漏洞相关问题，识别并修复潜在的安全漏洞。
- c) 针对模型研发和二次开发所使用的训练数据来源，在来源选择、不同来源搭配、来源管理和追溯等方面提高安全水平。

注：参考 GB/T 45654-2025《网络安全技术 生成式人工智能服务安全基本要求》中相关部分内容。

- d) 规范训练数据标注流程，采用交叉标注、结果审计等质量控制方法，提升标注准确性、安全性和可靠性，降低个体差异和个人偏见对标注质量的影响。

注：参考 GB/T 45674-2025《网络安全技术 生成式人工智能数据标注安全规范》中相关部分内容。

- e) 具备数据安全管理机制，包括对各模态训练数据的清洗过滤，重点去除违法不良信息和个人隐私信息；对数据源进行有效的安全检查，防止被投毒等。

注：参考 GB/T 45652-2025《网络安全技术 生成式人工智能预训练和优化训练数据安全规范》中相关部分内容。

- f) 使用安全透明的训练推理计算框架，提升模型算法训练和推理环境的安全性。

注：参考 GB/T 45958-2025《网络安全技术 人工智能计算平台安全框架》中 5.1 资源层安全功能要求。

7.3 验证确认阶段

验证确认安全指引包括：



- a) 结合设计开发阶段对人工智能技术的内生安全能力设计，根据确定的可靠性、公平性、透明性、可解释性等指标开展验证确认。

注：参考 GB/T 42888-2023《信息安全技术 机器学习算法安全评估规范》中 5.1.3 a) 构建测试数据集。

- b) 采用人工验证、自动化评估验证、交叉验证等方式，根据人工智能应用过程的类型和特点，开展安全测试：

- 1) 涉及生成内容的，测试生成内容安全防护措施是否有效，能否防止生成违法有害的内容、混淆事实误导用户的内容、不真实不准确的内容，以及防止生成内容被干扰或篡改；

注：参考 GB/T 45654-2025《网络安全技术 生成式人工智能服务安全基本要求》中 5.2 的要求。

- 2) 涉及控制物理装置的，测试能否避免与其他物体发生任务计划外的碰撞，且在任务执行过程中不出现剧烈抖动或者频繁卡顿的现象，避免对环境安全、人身安全、设备功能安全等造成损害；

- 3) 涉及重大决策的，对可解释性、透明性以及输入信息遭受干扰时决策结果的鲁棒性进行测试，挖掘后门攻击等安全风险；

- 4) 涉及未成年人等特殊群体的，测试是否结合特殊群体特点和权益保护需求，在数据处理、内容交互、功能设计及风险处置等方面采取差异化保护措施，并验证相关措施的有效性。



- c) 安全测试结果未达到规划的安全能力的，当通过模型调优等手段能够满足要求时，采取相应措施提升安全能力并重新开展安全测试，直至符合安全能力需求；经评估需要调整设计方案时，重新进行设计开发。

7.4 部署阶段

部署安全指引包括：

- a) 采用安全护栏等安全防护技术手段，识别拦截违法不良内容、防范提示词注入攻击、防止执行链路失控等。安全防护技术手段使用第三方服务的，对其安全技术能力和持续服务能力进行测试评估。
- b) 对建设部署所需的软硬件设备、第三方工具等进行安全检测，对未修复的已知漏洞进行修复，无法及时修复的，采取有效的缓解措施。建立漏洞追溯机制，跟踪相关软硬件安全漏洞、缺陷信息，防范供应链植入后门。
- c) 根据部署环境，采取针对性安全保护措施：
 - 1) 部署于公有云时，对传输和存储的数据进行加密保护，通过计算隔离的安全容器运行环境提升安全防护能力，支持操作审计与日志服务等功能以满足安全审计需求。在高敏感场景下，宜使用机密计算技术；



- 2) 部署于私有云时，做好数据分类分级保护。加强授权访问控制，实施差异化权限管理配置，对于算力、存储、应用程序等分别配置控制权限；
 - 3) 部署于本地服务器时，确保本地服务器环境具备安全防护机制和风险控制措施，具备防火墙、入侵检测、防御系统、恶意代码防护等安全组件。具备严格的访问控制和权限管理机制，确保各类用户、服务和进程在授权范围内对资源进行访问和操作；
 - 4) 部署于终端设备时，所处理的用户数据原则上不宣传出终端设备，确需传输的宜提前征得用户同意，并进行个人信息保护影响评估，确保对用户权益没有显著负面影响。
- d) 基于可控、持续、可扩展、兼容等方面考虑进行算力配置及资源选择，提升算力安全性。

注：人工智能加速芯片安全性参考 TC260-PG-20261A《网络安全标准实践指南——人工智能加速芯片安全技术规范》。

7.5 运行和监控阶段

运行和监控安全指引包括：

- a) 对用户输入信息进行安全检测，在不满足安全规则时进行相应处置。根据人工智能产品或服务的类型和特点，采取技术措施，确保输出内容的合格率、准确性、可靠性满足应用要求。在具备舆论传播属性的服务场景设置与服务场景规模相匹配数量的监看人员，确保内容安全满足要求。



注：参考 GB/T 45654-2025《网络安全技术 生成式人工智能服务安全基本要求》中 6.5 有关要求。

- b) 在具备舆论传播属性的服务场景，按有关要求及强制性国家标准做好生成合成内容标识工作，通过国家标准公共服务平台等途径验证标识有效性。

注 1：满足《人工智能生成合成内容标识办法》以及 GB 45438-2025《网络安全技术 人工智能生成合成内容标识方法》的要求。

注 2：参考人工智能生成合成内容标识服务平台
<https://www.gcmark.com>。

- c) 对人工智能产品或服务的输入内容持续监测，防范恶意输入攻击，并建立常态化监测测评手段以及应急管理措施，发现安全风险时及时管控。
- d) 制定应急预案，发生安全事件时及时管控，特别针对重大及以上安全事件具备紧急关停或紧急切换备用产品或服务能力，具备版本回退等方面能力。涉及物理空间与人交互的场景，设置便于人工操作的关停方式、防止失控时无法停止。
- e) 面向当前人工智能不可解释性、幻觉等技术特点，围绕人工智能安全问题定责难的现状，鼓励在涉及人身安全的应用场景，事前针对重大人身安全事故建立无条件救济保障措施，在发生重大事故时进行第一时间、无理由的人道主义救济，全面筑牢人身安全保障。
- f) 制定信息内容交互行为规范、安全运营机制、投诉反馈机制、技术防护能力等，防范人工智能产品或服务被不当或恶意利用生成、发布、传播虚假有害信息风险。

注：参考 TC260-PG-20267A《网络安全标准实践指南——个人用户使用人工智能服务安全指南》。



7.6 持续验证评估阶段

持续验证评估安全指引包括：

- a) 每年以及在常态监测预警中发现人工智能应用安全风险存在变化时，开展 6.3 中的安全测试及安全调整，出现新风险项时对安全要求进行更新。
- b) 发生重大风险事件或应用场景、智能化水平和服务规模发生显著变化时，首先按照 6.3 开展安全测试，检查安全设计方案的有效性，失效则应退回前期规划阶段，重新开展安全设计进行版本更新。

7.7 退役下线阶段

退役下线安全指引包括：

- a) 围绕退役下线的必要性、可行性以及影响范围进行安全评估，形成安全退役下线工作方案，妥善处置基础设施、数据、系统，并向访问使用方、相关合作主体等告知退役下线安排。法律法规规定需向有关主管部门报告、备案或者履行其他手续的，向有关主管部门等同步该方案。
- b) 保障产品或服务安全平稳停止，在停止产品或服务前保证原有的访问使用方完成迁移工作，避免因退役下线骤停带来人身、财产、社会安全问题。
- c) 安全处置人工智能算法、模型等相关数据，包括但不限于模型文件、训练数据、权重文件、超参数配置、训练脚本、推理服



务配置等，确保无法被恢复或加载。安全处置产品或服务相关的数据，包括但不限于用户数据、业务数据等，确保不可恢复。对于确需留存的数据，按照合法、正当、必要原则保留最小范围实施归档，明确归档留存期限，并确保归档存储的安全隔离和权限控制。

7.8 其他

其他相关安全要求如下。

- a) 应按照等级保护有关要求开展系统安全建设以及测评工作。

注：参考标准包括但不限于 GB/T 22239 《信息安全技术 网络安全等级保护基本要求》、GB/T 22240 《信息安全技术 网络安全等级保护定级指南》等。

- b) 涉及个人信息处理的，应按照个人信息保护有关要求取得用户个人同意、开展个人信息安全影响评估、个人信息安全合规审计等。

注：参考标准包括但不限于 GB/T 35273 《信息安全技术 个人信息安全规范》、GB/T 39335 《信息安全技术 个人信息安全影响评估指南》等。

- c) 应按照数据安全保护有关要求开展数据分类分级、数据安全风险评估、数据安全能力成熟度评估、数据出境安全风险评估等。

注：参考标准包括但不限于 GB/T 37988 《信息安全技术 数据安全能力成熟度模型》、GB/T 43697 《数据安全技术 数据分类分级规则》、GB/T 45577 《数据安全技术 数据安全风险评估方法》、GB/T 46068 《数据安全技术 个人信息跨境处理活动安全认证要求》等。

- d) 涉及密码应用的，应按照密码应用有关要求开展密码应用设计、密码应用测评等。

注：参考标准包括但不限于 GB/T 39786 《信息安全技术 信息系统密码应用基本要求》、GB/T 43206 《信息安全技术 信息系统密码应用测评



要求》、GB/T 43207 《信息安全技术 信息系统密码应用设计指南》等。

- e) 涉及关键信息基础设施的，应按照关键基础设施有关要求开展关基安全风险评估等。

注：参考标准包括但不限于 GB/T 39204 《信息安全技术 关键信息基础设施安全保护要求》、关键信息基础设施安全保护能力指标体系有关国家标准等。





附录 A

(规范性)

术语表

A.1 机器学习 machine learning

通过计算技术优化模型参数的过程,使模型的行为反映数据或经验。

[来源: GB/T 41867—2022, 3.2.10]

A.2 深度学习 deep learning

通过训练具有许多隐层的神经网络来创建丰富层次表示的方法。

注:深度学习是机器学习的一个子集。

[来源: GB/T 41867—2022, 3.2.27]

A.3 大模型 large-scale model

基于大量数据训练得到,具有复杂计算架构,能处理复杂任务,且具备一定泛化性的深度学习模型。

注:大模型的参数量由其功能和模态决定,一般不低于1亿。大模型训练使用的数据总量受参数量的影响,达到收敛的大模型的参数量的对数与其训练数据总量的对数成正比。

[来源: GB/T 45288.1—2025, 3.1]

A.4 大模型平台 large-scale model platform

为开发或使用大模型提供各类资源的软硬件平台。

注:大模型平台不包含大模型。

[来源: GB/T 45288.3—2025, 3.1]

A.5 生成式人工智能服务 generative artificial intelligence service



利用生成式人工智能技术向公众提供生成文本、图片、音频、视频等服务的内容。

[来源: GB/T 45654—2025, 3.1]

A.6 基础模型 foundation model

可用于或可适配于一个或多个领域中广泛任务的人工智能模型。

注1: 构建基础模型的典型方法是对大量数据应用监督式机器学习或自监督式机器学习。

注2: 基础模型可作为各种应用、任务和使用场景的一部分, 其使用不必涉及生成式人工智能。

[来源: ISO/IEC 22989:2022/DAmD 1(en), 3.3.19]

A.7 智能体 agent

能够自主感知环境、制定决策、采取行动实现特定目标的智能系统, 一般具有记忆、规划、使用工具等基本能力。

[来源: 《人工智能安全治理框架》2.0, 术语 14]

A.8 语料库 corpus

自然语言数据的集合。

注: 语料库可用于多种活动, 例如文本分析或术语工作。

[来源: ISO 1087:2019, 3.6.4]

A.9 对抗样本 adversarial example

在输入数据中通过故意添加细微干扰获得的、可导致机器学习算法模型以高置信度给出错误输出的样本。

[来源: ISO/IEC TR 29119-11:2020(en), 3.1.7]

A.10 对抗攻击 adversarial attack

通过构造微扰数据等输入样本, 使人工智能模型产生错误输出或行



为的攻击方式。

[来源：《人工智能安全治理框架》2.0，术语 13]

A.11 数据投毒 data poisoning

攻击者篡改、注入错误、误导数据，“污染”模型的概率分布，进而造成准确性、可信度下降的行为。

[来源：《人工智能安全治理框架》2.0，术语 12]





附 录 B

(规范性)

缩略语表

AI: 人工智能 (Artificial Intelligence)

API: 应用程序编程接口 (Application Programming Interface)

BIOS: 基本输入输出系统 (Basic Input Output System)

DDoS: 分布式拒绝服务攻击 (Distributed Denial of Service)

MFA: 多因素认证 (Multi-factor Authentication)

PC: 个人计算机 (Personal Computer)

SBOM: 软件物料清单 (Software Bill of Materials)

SDLC: 软件开发生命周期 (Software Development Life Cycle)

SLA: 服务等级协议 (Service Level Agreement)

TEE: 可信执行环境 (Trusted Execution Environment)

TPM: 可信平台模块 (Trusted Platform Module)



参 考 文 献

- [1] 中华人民共和国网络安全法
- [1.1] GB/T 22239 信息安全技术 网络安全等级保护基本要求
- [1.2] GB/T 22240 信息安全技术 网络安全等级保护定级指南
- [1.3] GB/T 39204 信息安全技术 关键信息基础设施安全保护要求
- [1.4] GB/T 43698 网络安全技术 软件供应链安全要求
- [2] 中华人民共和国密码法
- [2.1] GB/T 39786 信息安全技术 信息系统密码应用基本要求
- [2.2] GB/T 43207 信息安全技术 信息系统密码应用设计指南
- [3] 中华人民共和国个人信息保护法
- [3.1] GB/T 35273 信息安全技术 个人信息安全规范
- [4] 中华人民共和国数据安全法
- [4.1] GB/T 37988 信息安全技术 数据安全能力成熟度模型
- [4.2] GB/T 41479 信息安全技术 网络数据处理安全要求
- [4.3] GB/T 43697 数据安全技术 数据分类分级规则
- [5] ISO/IEC 22989 Information technology — Artificial intelligence — Artificial intelligence concepts and terminology
- [6] ISO/IEC 42001 Information technology — Artificial intelligence — Management system
- [7] ISO/IEC 5338 Information technology — Artificial intelligence — AI system life cycle processes

[8] ISO/IEC 5339 Information technology — Artificial intelligence —
Guidelines for AI applications

