

网络安全标准实践指南

Practice Guide for Cybersecurity Standards

——个人用户使用人工智能服务安全指南

Personal User Security Guide for Using
Artificial Intelligence Services

v1.0-202608

全国网络安全标准化技术委员会秘书处

Secretariat of the National Technical Committee 260 on
Cybersecurity of Standardization Administration of China

2026年8月

August 2026

本文档可从以下网址获得：

This document can be obtained from the following website:

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

Foreword

《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。

The "Practice Guidelines for Cybersecurity Standards" (hereinafter referred to as the "Practice Guidelines") is a standard related technical document organized and published by the Secretariat of the National Cybersecurity Standardization Technical Committee (hereinafter referred to as the "Cybersecurity Standardization Committee"). It aims to promote cybersecurity related standards and knowledge, and provide standardized practice guidelines around cybersecurity laws, regulations, policies, standards, cybersecurity hotspots, and events.

本文件起草单位：北京中关村实验室、北京快手科技有限公司、中国电子技术标准化研究院、中央网信办数据与技术保障中心、工业和信息化部电子第五研究所、北京航空航天大学、东莞市人工智能与数字经济集团有限公司、阿里云计算有限公司、北京天融信网络安全技术有限公司、蚂蚁科技集团股份有限公司、北京金山办公软件股份有限公司、北京智源人工智能研究院。

This document was drafted by Beijing Zhongguancun Laboratory, Beijing Kwai Technology Co., Ltd., China Electronics Technology Standardization Institute, Central Cyberspace Office Data and Technology Support Center, the Fifth Institute of Electronics of the Ministry of Industry and Information Technology, Beijing University of Aeronautics and



Astronautics, Dongguan Artificial Intelligence and Digital Economy Group Co., Ltd., Alibaba Cloud Computing Co., Ltd., Beijing Tianrongxin Network Security Technology Co., Ltd., Ant Technology Group Co., Ltd., Beijing Jinshan Office Software Co., Ltd., and Beijing Zhiyuan Artificial Intelligence Research Institute.

本文件起草人：刘栋、谷晨、落红卫、马梦娜、余祥臣、刘勇、张妍婷、李芒、张宗洋、罗丽琪、崔栋、申东洋、吴波、王姣、彭骏涛、王龔、林冠辰、欧翹翎、安红云、吕梦、刘颖、杨熙、费凡芮、胡奥迪、张之义、周佳杭、陈海龙。

Drafters of this document: Liu Dong, Gu Chen, Luo Hongwei, Ma Mengna, Yu Xiangchen, Liu Yong, Zhang Yanting, Li Mang, Zhang Zongyang, Luo Liqi, Cui Dong, Shen Dongyang, Wu Bo, Wang Jiao, Peng Juntao, Wang Yan, Lin Guanchen, Ou Qiaoling, An Hongyun, Lv Meng, Liu Ying, Yang Xi, Fei Fanrui, Hu Audi, Zhang Zhiyi, Zhou Jiahang, Chen Hailong.





声 明

Statement

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”。

The copyright of this Practice Guide belongs to the Secretariat of the Cybersecurity Standards Committee. Without written authorization from the Secretariat, no part of the Practice Guide may be copied or translated in any way. Please indicate "Source: Secretariat of the National Technical Committee 260 on Cybersecurity of Standardization Administration of China " when reprinting or quoting the viewpoints and data in this "Practice Guide".





目 录

Catalogue

1 范围	1
scope	1
2 规范性引用文件	1
Normative reference documents	1
3 术语定义	2
Term Definition.....	2
4 概述	2
Overview.....	2
5 安全意识提升	3
Enhancing safety awareness	3
6 安全使用指引	4
Guidelines for Safe Use	4
7 使用行为准则	7
Code of Conduct for Use	7
参考文献	10
References.....	10





1 范围

Scope

本文件提供了个人用户通过互联网使用人工智能服务的安全指引，包括安全意识提升、安全使用指引、使用行为准则等。

This document provides security guidelines for individual users to use AI services through the Internet, including security awareness promotion, safe use guidelines, use codes of conduct, etc.

本文件适用于指导个人用户安全地通过互联网使用人工智能服务，也可为人工智能服务的技术提供者、服务开发者、服务运营者保障个人用户的安全使用提供参考。

This document is applicable to guiding individual users to use AI services safely through the Internet, and can also provide reference for AI service technology providers, service developers, and service operators to ensure the safe use of individual users.

2 规范性引用文件

Normative reference documents

下列文件中的内容通过文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。

The contents of the following documents constitute essential clauses of this document through normative references in the text. Among them, for referenced documents with dates, only the version corresponding to that date is applicable to this document; The latest version (including all modifications) of the referenced document without a date is applicable to this document.

GB/T 25069 信息安全技术术语



GB/T 25069 Information Security Technical Terminology

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

TC260-PG-20266A Network Security Standard Practice Guide - Security Guidelines for Intelligent Agent Deployment and Use

3 术语定义

Term Definition

GB/T 25069 界定的术语和定义适用于本文件。

The terms and definitions defined in GB/T 25069 are applicable to this document.

4 概述

Overview

本文件主要关注个人用户在通过互联网使用人工智能服务时，尤其是利用人工智能技术提供预测、建议、决策，或是生成文本、图片、音频、视频等内容服务，如何提升安全意识、安全使用服务，并给出了建议的使用行为准则，避免因不当使用人工智能服务对自身或他人产生不利影响。

This document primarily focuses on how individual users can enhance their security awareness and safely use artificial intelligence services when accessing such services over the internet—particularly those that leverage AI technologies to provide predictions, recommendations, decision-making, or to generate content such as text, images, audio, and video. It also provides recommended codes of conduct for use, aiming to avoid adverse effects on users themselves or on others resulting from improper use of AI services.



5 安全意识提升

Enhancing safety awareness

个人用户在使用人工智能过程中，需要强化安全意识、建立安全理念，包括：

Individual users need to strengthen their security awareness and establish security concepts when using artificial intelligence, including:

- a) 正确认知人工智能，认识到人工智能的能力具有局限性，可能提供错误信息、出现理解偏差、输出幻觉，或无法按要求完成某些特定任务。

To have a correct understanding of artificial intelligence, it is necessary to recognize that its capabilities have limitations, such as providing incorrect information, causing misunderstandings, producing illusions, or being unable to complete certain specific tasks as required.

- b) 理性对待人工智能，避免盲目接受和信任人工智能生成的结果，增强对虚假、错误、误导性内容的识别能力。

Treat artificial intelligence rationally, avoid blindly accepting and trusting the results generated by artificial intelligence, and enhance the ability to identify false, erroneous, and misleading content.

- c) 适度使用人工智能，将其作为辅助真实生活的工具，避免沉迷依赖，正确认知人工智能情感服务原理，避免通过其过多替代现实交往与真实活动。

Moderate use of artificial intelligence as a tool to assist real life, avoiding addiction and dependence, correctly understanding the principles of artificial intelligence emotional services, and avoiding excessive substitution of real communication and activities through it.



- d) 关注安全风险动态，增强对人工智能安全风险、伦理安全影响的认识，审慎提供隐私信息，意识到人工智能服务可能被恶意利用，例如生成仿冒亲友声音、伪造官方通知等。

Pay attention to security risk dynamics, enhance awareness of artificial intelligence security risks and ethical safety concerns. Be aware of the potential for malicious use of artificial intelligence services, such as generating fake voices of family and friends, forging official notifications, and providing privacy information with caution.

注：人工智能安全风险可参考《人工智能安全治理框架》，人工智能伦理安全影响可参考 TC260-005《人工智能应用伦理安全指引 1.0》

Note: The security risks of artificial intelligence can refer to the "Artificial Intelligence Security Governance Framework", and the ethical and security impacts of artificial intelligence can refer to TC260-005 "Ethical and Security Guidelines for Artificial Intelligence Applications 1.0"

6 安全使用指引

Guidelines for Safe Use

个人用户在使用人工智能服务过程中，应持续强化安全使用措施，其安全使用指引主要包括：

Individual users should continuously strengthen security measures when using artificial intelligence services, and their security usage guidelines mainly include:

- a) 使用从正规渠道获取的人工智能服务，优先选用最新版本；涉及生成式人工智能服务的，应选择已通过生成式人工智能服务备案的大模型，不通过来源不明的中转站、未授权代理等方式使用相关服务。

Use artificial intelligence services obtained from legitimate channels



and prioritize the latest version; For those involving generative artificial intelligence services, large models that have been registered through generative artificial intelligence services should be selected, and relevant services should not be used through unknown sources such as transfer stations or unauthorized agents.

- b) 部署使用智能体时，宜参考 TC260-PG-20266A 《网络安全标准实践指南——智能体部署使用安全指引》，采用适宜的部署方式，并采用相应安全使用的措施。

When deploying and using intelligent agents, it is advisable to refer to TC260-PG-20266A "Network Security Standard Practice Guide - Security Guidelines for Intelligent Agent Deployment and Use", adopt appropriate deployment methods, and adopt corresponding security measures.

- c) 查看用户协议和隐私协议，了解以下信息并进行相应处置：

Review the user agreement and privacy agreement, understand the following information and take appropriate measures:

- 1) 人工智能服务的局限性、适用的人群等；

limitations of artificial intelligence services, applicable populations, etc;

- 2) 会被收集的个人信息以及使用方式等，特别关注个人信息是否会被共享等，并检查是否符合最小化收集原则；

The types of personal information that will be collected and the methods of their use, with particular attention to whether the personal information will be shared, etc., and to verify whether it complies with the principle of minimization in collection;

注 1：个人信息收集、使用的要求可参考 GB/T 35273。

Note 1: The requirements for the collection and use of personal information can refer to GB/T 35273.



- 3) 个人提供的相关数据是否会被用作训练数据，是否具备撤回同意功能；

Whether the relevant data provided by individuals will be used as training data and whether they have the function of withdrawing consent;

- 4) 投诉举报机制和渠道。

Complaint reporting mechanism and channels.

注 2：渠道通常包括但不限于电话、邮件、交互窗口、短信等方式。

Note 2: Channels usually include but are not limited to telephone, email, interactive windows, SMS, and other methods.

- d) 关注人工智能服务申请的访问权限（如麦克风、摄像头、相册、操作用户文件或数据等），关闭非必要的授权。

Pay attention to the access permissions requested for artificial intelligence services (such as microphones, cameras, photo albums, operating user files or data, etc.), and disable non essential authorizations.

- e) 妥善保管敏感信息，在输入信息过程中，审慎提供隐私信息或敏感个人信息，防范隐私信息、商业秘密等意外泄露或被窃取。

Properly protect sensitive information, carefully provide private information or sensitive personal information during the input process, and prevent accidental leakage or theft of private information, trade secrets, etc.

- f) 审慎对待人工智能服务输出的结果：

Prudential treatment of the output of artificial intelligence services:

- 1) 必要时可通过多个人工智能服务或其他渠道对结果进行交叉验证，检查其中是否存在幻觉、偏见、歧视等错误信息，是否存在伪造冒用等虚假信息；

When necessary, cross validation of the results can be



conducted through multiple artificial intelligence services or other channels to check for false information such as illusions, biases, discrimination, and forgery;

- 2) 对于人工智能服务在法律、医疗、就业、教育、理财等方面的输出结果，仅将其作为参考，在涉及财产、人身安全时应通过其他独立渠道进行核实。

For the output results of artificial intelligence services in areas such as law, healthcare, employment, education, and finance, they should only be used as references. When it comes to property and personal safety, they should be verified through other independent channels.

- g) 关注开发者提供的安全更新、补丁等，必要时进行安装和升级。

Pay attention to security updates, patches, etc. provided by developers, and install and upgrade as necessary.

- h) 个人权益受损时及时向主管部门或有关机构进行咨询求助或举报。

When personal rights are damaged, promptly consult and seek help or report to the competent department or relevant institutions.

- i) 如作为未成年人、老年人等用户的监护人，协助其安全使用人工智能服务。

As guardians of users such as minors and the elderly, assist them in safely using artificial intelligence services.

7 使用行为准则

Code of Conduct for Use

个人用户在使用人工智能服务过程中，应主动约束自身使用行为，相关使用行为规范包括：



Individual users should actively regulate their usage behavior when using artificial intelligence services, and relevant usage behavior norms include:

a) 合法使用人工智能，尊重他人尊严与合法权益：

Legitimate use of artificial intelligence, respect for the dignity and legitimate rights and interests of others:

1) 不恶意诱导人工智能产生有害行为、输出有害内容；

Do not maliciously induce artificial intelligence to produce harmful behaviors or output harmful content;

2) 不借助人工智能对他人进行误导、骚扰、操纵；

Not using artificial intelligence to mislead, harass, or manipulate others;

3) 不利用人工智能伪造、冒用他人身份或仿冒权威主体；

Do not use artificial intelligence to forge, impersonate others' identities, or impersonate authoritative entities;

4) 不使用人工智能实施攻击、破坏、盗窃等违法违规行为；

Not using artificial intelligence to carry out illegal and irregular activities such as attacks, destruction, and theft;

5) 不使用人工智能复刻他人画作、小说、歌曲、影视剧本、商标图案等信息，避免侵犯他人知识产权。

Do not use artificial intelligence to replicate other people's paintings, novels, songs, film scripts, trademark patterns, and other information to avoid infringing on others' intellectual property rights.

b) 在使用人工智能服务生成合成内容时：

When using artificial intelligence services to generate synthetic content:

1) 积极履行生成合成内容标识义务，关注相关服务输出内容是否具备生成合成内容标识，不恶意裁剪、删除标识；



Actively fulfill the obligation of generating synthetic content identifiers, pay attention to whether the output content of relevant services has the ability to generate synthetic content identifiers, and do not maliciously cut or delete identifiers;

注：生成合成内容标识的样式见 GB 45438-2025 第 5 章。

Note: The style for generating composite content identifiers can be found in Chapter 5 of GB 45438-2025.

- 2) 如发现生成内容存在准确性、科学性问题，或包含违法不良信息时，不进行二次传播；

If it is found that the generated content has accuracy or scientific issues, or contains illegal and harmful information, it will not be disseminated again;

- 3) 不滥用人工智能输出，如发现模型生成违法不良信息，可通过投诉举报路径反馈至服务提供者，帮助建设良性人工智能生态；

Do not abuse the output of artificial intelligence. If it is found that the model generates illegal and harmful information, it can be reported to the service provider through the complaint reporting path to help build a healthy artificial intelligence ecosystem;

- 4) 在发布人工智能生成合成内容时，遵守发布平台的社区规范或管理要求，主动声明并使用服务提供者提供的标识功能，积极履行内容标识义务。

When publishing artificial intelligence generated composite content, comply with the community norms or management requirements of the publishing platform, actively declare and use the identification function provided by the service provider, and actively fulfill the obligation of content identification.



参 考 文 献

References

- [1] 人工智能安全治理框架
Artificial Intelligence Security Governance Framework
- [2] GB/T 35273 信息安全技术 个人信息安全规范
GB/T 35273 Information Security Technology - Personal Information Security Specification
- [3] GB 45438-2025 网络安全技术 人工智能生成合成内容标识方法
GB 45438-2025 Network Security Technology Artificial Intelligence Generated Synthetic Content Identification Method
- [4] 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 45654-2025 Network Security Technology - Basic Requirements for Security of Generative Artificial Intelligence Services
- [5] TC260-005 人工智能应用伦理安全指引 1.0
TC260-005 Ethical and Safety Guidelines for Artificial Intelligence Applications 1.0