

ICS 35.240.01

CCS L77

团 体 标 准

T/ISC XXXX—XXXX

干部全息画像智能系统架构与服务 技术规范

Technical specification for architecture and services
of intelligent cadre holographic portrait system

（征求意见稿）

XXXXXXXXX 发布

XXXXXXXXX 实施

中国互联网协会 发布

目录

前言	1
引言	2
1 范围	3
2 规范性引用文件	3
3 术语和定义	4
3.1 干部全息画像 cadre holographic portrait	4
3.2 画像维度 portrait dimension	4
3.3 关系图谱 relationship graph	4
3.4 智能体路由 agent routing	4
3.5 检索增强生成 retrieval-augmented generation; RAG	4
3.6 评价模型 evaluation model	4
3.7 数据血缘 data lineage	4
3.8 AI 辅助决策 AI-assisted decision-making	4
3.9 可信审计 trusted audit	5
3.10 信创适配 indigenous information technology adaptation	5
4 缩略语	5
5 系统总体架构	6
5.1 架构概述	6
5.2 数据层概述	6
5.3 模型层概述	6
5.4 服务层概述	6
5.5 应用层概述	7
5.6 层间接口要求	7
6 数据采集与融合	7
6.1 多源数据采集	7
6.2 数据分类与标签	7
6.3 数据真实性与有效性核验	8
6.4 数据脱敏与去标识化	8
6.5 数据生命周期管理	8

- 6.6 数据血缘与溯源9
- 7 模型层与智能体路由9
 - 7.1 多规格模型协同9
 - 7.2 意图识别与路由调度9
 - 7.3 模型降级与容错10
 - 7.4 模型配置管理10
- 8 知识图谱与关系推理11
 - 8.1 图谱构建11
 - 8.2 关系建模12
 - 8.3 关系推理12
 - 8.4 置信度管理12
 - 8.5 图谱健康管理12
 - 8.6 图谱质量要求12
- 9 检索增强生成服务13
 - 9.1 混合检索13
 - 9.2 检索质量要求13
 - 9.3 知识增量更新14
 - 9.4 来源引用与溯源14
 - 9.5 图谱增强检索（GraphRAG）14
 - 9.6 自适应检索策略14
 - 9.7 生成质量保障15
 - 9.8 文档分块策略15
- 10 画像生成与评价服务15
 - 10.1 个人画像15
 - 10.2 群体画像16
 - 10.3 评价模型配置16
 - 10.4 评价规则版本管理17
 - 10.5 异常偏差检测与人工复核17
 - 10.6 群体公平性分析18
 - 10.7 AI 辅助边界18

11 智能文书服务18

 11.1 文书模板管理18

 11.2 文书生成19

 11.3 AI生成内容标识19

 11.4 公文格式合规20

 11.5 文书质量要求20

12 服务接口与应用21

 12.1 API 接口规范 21

 12.2 应用集成要求21

 12.3 前端服务要求22

 12.4 并发性能要求22

 12.5 系统可用性要求23

13 安全、可信与合规23

 13.1 访问控制23

 13.2 认证与鉴权23

 13.3 审计与追溯24

 13.4 内容安全24

 13.5 数据安全24

 13.6 等保与分级保护25

 13.7 AI 模型安全26

 13.8 供应链安全26

 13.9 安全事件应急响应26

14 信创适配27

 14.1 国产 AI 处理器27

 14.2 国产操作系统27

 14.3 国产数据库27

 14.4 验证与发布27

15 部署与验收28

 15.1 部署要求28

 15.2 验收测试28

15.3 运维监控28

15.4 测试环境配置要求29

16 智能化演进方向（资料性）35

16.1 智能体互联35

16.2 多模态融合36

16.3 可解释性36

16.4 联邦学习36

16.5 端侧 AI 与模型蒸馏36

16.6 AI Agent 自主规划错误！未定义书签。

16.7 长上下文处理37

附录 A30

附录 B32

附录 C34

C.2 标签体系设计参考35

C.3 评价规则配置参考35

前言

本文件按照 GB/T 1.1—2020《标准化工作导则 第1部分：标准化文件的结构和起草规则》的规定起草。

请注意本文件的某些内容可能涉及专利。本文件的发布机构不承担识别专利的责任。

本文件由中国互联网协会归口。

本文件起草单位：常州市人力资源和社会保障局；江苏省常州技师学院；中国移动通信集团江苏有限公司常州分公司；东南大学计算机科学与工程学院、软件学院、人工智能学院。

本文件主要起草人：周柯敏、左洲鹏、薛晖、王江东、丁国明、周建东、李冉、潘宁、黄平安、周昊、朱士鹏、孙俊、张彩霞。

本文件为首次发布。

引言

干部全息画像智能系统是面向党政机关干部管理工作的专业化人工智能应用系统，旨在通过多源数据融合、知识图谱构建、智能检索与评价分析等技术手段，为干部队伍建设提供数据驱动的辅助决策支持。

本文件旨在规范干部全息画像智能系统的架构设计、技术要求和服务规范，为系统的研发、部署、验收和运维提供统一的技术依据。

本文件与 T/JS AI XXXX《政务大模型本地化部署技术规范》构成“应用+底座”互补关系：T/JS AI XXXX 规定了政务场景下大模型本地化部署的基础设施、模型管理、性能指标等底座层面的技术要求；本文件在上述底座之上，规定干部全息画像系统在数据治理、智能分析、服务接口、安全合规等应用层面的架构要求和技术规范。两标准协同，共同构成政务干部管理领域 AI 应用的完整技术规约体系。

本文件在编制过程中，充分考虑了以下原则：干部管理场景的特殊性——系统服务于公务员队伍建设，涉及高度政治敏感性，应在技术架构层面体现对干部信息安全和规范管理的尊重；AI 辅助边界——系统为组织人事部门提供辅助分析能力，不替代组织决定，任何评价结果仅供参考；轻量化与自主可控——系统宜采用轻量化架构，支持便携式部署，适配国产化软硬件环境，确保数据不出域；评价规则可配置治理——评价维度和权重由组织根据干部管理政策自行配置和版本化管理，系统不预设统一评价标准。

1 范围

本文件规定了干部全息画像智能系统的总体架构、数据采集与融合、模型层与智能体路由、知识图谱与关系推理、检索增强生成服务、画像生成与评价服务、智能文书服务、服务接口与应用、安全可信与合规、信创适配、部署与验收等方面的技术要求。

本文件适用于党政机关干部全息画像智能系统的规划、设计、开发、部署、验收和运维。企事业单位参照使用时，可根据实际情况进行调整。本文件涵盖的合规要求为底线要求。

2 规范性引用文件

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。

GB/T 22239—2019 信息安全技术 网络安全等级保护基本要求

GB/T 25070—2019 信息安全技术 网络安全等级保护安全设计技术要求

GB/T 35273—2020 信息安全技术 个人信息安全规范

GB/T 37668—2019 信息技术 互联网内容无障碍可访问性技术要求与测试方法

GB/T 37988—2019 信息安全技术 数据安全能力成熟度模型

GB/T 39786—2021 信息安全技术 信息系统密码应用基本要求

GB/T 38674—2020 信息安全技术 应用软件安全编程指南

GB/T 41867—2022 信息技术 人工智能 术语

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 9704—2012 党政机关公文格式

GB/T 18336（所有部分） 信息技术 安全技术 IT 安全评估准则

GB/Z 185（所有部分） 人工智能 智能体互联

GB/Z 236—2026 人工智能 检索增强生成系统评价指标与方法

GB/Z 253—2026 人工智能 检索增强生成通用技术要求

T/JSAT XXXX 政务大模型本地化部署技术规范

GM/T 0002—2012 SM4 分组密码算法标准

GM/T 0003—2012 SM2 椭圆曲线公钥密码算法标准

GM/T 0004—2012 SM3 密码杂凑算法标准

GM/T 0054—2018 信息系统密码应用基本要求

3 术语和定义

GB/T 41867—2022 界定的以及下列术语和定义适用于本文件。

3.1 干部全息画像 cadre holographic portrait

基于多源异构数据，通过知识图谱、自然语言处理、评价模型等技术手段，对干部基本信息、履职经历、能力素质、经组织认定的发展潜力情况等进行多维度、立体化、结构化呈现的综合信息模型。

3.2 画像维度 portrait dimension

用于描述和评价干部特定属性的分类框架，如基本信息、学历经历、考核评价、奖惩记录、培训情况等。画像维度由组织根据干部管理需要自行配置，本系统不预设统一的维度标准。

3.3 关系图谱 relationship graph

以图结构表示干部之间、干部与组织之间、干部与事件之间的关联关系的知识表示形式，支持关系的存储、查询和推理。

3.4 模型路由 model routing

根据用户请求意图、复杂度及模型能力标签，将推理请求动态调度至合适模型实例的机制，包括基于规则的静态路由与基于模型预测的动态路由。

3.5 检索增强生成 retrieval-augmented generation; RAG

将外部知识库的检索结果与生成模型相结合，以提高生成内容准确性、时效性和可溯源性的技术方法。

3.6 评价模型 evaluation model

基于组织依规定义的维度、权重和规则，对干部相关数据进行确定性计算的规则化模型。评价规则由组织自行配置和版本化管理；评价模型不包含生成式人工智能的等次评定、优劣排序等评价性判断。人工智能在评价中的作用限于数据提取、特征归纳和辅助计算，其产出经人工确认后方可作为评价计算的输入。

3.7 数据血缘 data lineage

记录数据从采集、处理、存储到应用全生命周期中的来源、变换和流转关系的技术机制，用于支持数据溯源和质量保障。

3.8 AI 辅助决策 AI-assisted decision-making

利用人工智能技术对数据进行分析和推理，为组织人事部门提供参考性分析结果和建议，但最终以人工判断和组织决定为准的决策支持模式。

3.9 可信审计 trusted audit

对系统中数据访问、操作行为、AI 生成过程及结果进行完整记录，确保可追溯、可核查、可问责的审计机制。

3.10 信创适配 indigenous information technology adaptation

系统对国产化处理器、操作系统、数据库、中间件等信息技术应用创新产品的兼容适配能力。

3.11 向量数据库 vector database

用于存储和检索高维向量的专用数据库系统，支持基于相似度的高效检索。

注：本文件中使用的向量数据库特指本地化部署的向量索引系统，不绑定特定厂商产品。

3.12 工作流编排 workflow orchestration

将多个业务步骤按照预定义逻辑组合为可执行流程，并管理其执行、监控和异常处理的能力。

3.13 模型上下文协议 Model Context Protocol; MCP

一种标准化的智能体间通信协议，用于实现模型与外部工具、数据源和服务的安全交互。

3.14 工作秘密 work secret

组织部门在干部选拔任用、考核评价、日常管理等工作产生的，不属于国家秘密但泄露后会妨碍工作正常开展或影响公平公正的信息。

关于工作秘密的管理要求，参见中共中央保密委员会 2019 年印发的《工作秘密管理暂行办法》（中保委发〔2019〕6 号）；《中华人民共和国保守国家秘密法》第六十四条、《中华人民共和国保守国家秘密法实施条例》第七十三条另有规定的，从其规定。

3.15 工具调用 tool calling

以结构化方式描述外部工具的名称、功能、输入参数与输出结构，由大语言模型生成调用请求、运行时执行并回传结果的机制。

3.16 智能体记忆 agent memory

智能体在执行过程中用于存储和检索信息的机制，包括短期记忆（会话上下文与中间状态）和长期记忆（跨会话偏好与历史结论）。

3.17 智能体编排 agent orchestration

对多个智能体的任务分配、协同执行和结果汇聚进行管理的机制，包括主管—专家、分工协作等编排模式。

4 缩略语

下列缩略语适用于本文件。

AI：人工智能（Artificial Intelligence）

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

BM25：最佳匹配 25 检索算法（Best Matching 25）

JWT: JSON 网络令牌 (JSON Web Token)

MCP: 模型上下文协议 (Model Context Protocol)

NLP: 自然语言处理 (Natural Language Processing)

RBAC: 基于角色的访问控制 (Role-Based Access Control)

RAG: 检索增强生成 (Retrieval-Augmented Generation)

RRF: 倒数排序融合 (Reciprocal Rank Fusion)

VLM: 视觉语言模型 (Vision-Language Model)

XAI: 可解释人工智能 (Explainable Artificial Intelligence)

5 系统总体架构

5.1 架构概述

系统应采用分层架构设计，自底向上分为数据层、模型层、服务层和应用层四个逻辑层次。各层次之间应通过标准化接口进行交互，实现松耦合、高内聚的设计目标。

系统架构应满足以下总体要求：

- a) 各层次功能边界清晰，层次间依赖关系明确；
- b) 支持横向扩展，单层次可独立扩容而不影响其他层次；
- c) 支持轻量化部署，可在满足最低硬件要求的单机环境下运行；
- d) 支持数据全部存储和处理在本地环境，不依赖外部云服务。

5.2 数据层概述

数据层负责多源异构数据的采集、存储、清洗和融合管理，核心能力包括：结构化与非结构化数据的统一接入和存储、数据清洗与规范化、数据血缘追踪与版本管理、数据加密与脱敏保护。数据层宜采用扁平化存储结构，通过标签体系区分数据类型。数据层的详细技术要求见第 6 章。

5.3 模型层概述

模型层负责大语言模型、嵌入模型、视觉语言模型等多规格 AI 模型的统一管理和调度，核心能力包括：多规格模型协同工作、模型注册与动态加载、智能路由与意图识别、模型降级与容错、模型配置管理与版本控制、智能体构建、工具接入与编排执行。模型层的详细技术要求见第 7 章。

5.4 服务层概述

服务层应基于模型层能力，封装面向业务场景的功能服务，核心能力包括：知识图谱构建与查询服务、检索增强生成 (RAG) 服务、画像生成与评价服务、智能文书生成服务、 workflow 编排与执行服务。服务层各功能模块之间应通过内部 API 通信，宜采用微服务或模块化架构。服务层的详细技术要求见第 8 章至第 12 章。

5.5 应用层概述

应用层面向终端用户，提供交互界面和业务功能入口，核心能力包括：干部个人画像展示、群体画像分析、智能问答交互、文书自动生成、 workflow 管理与执行监控。应用层的详细技术要求见第 13 章。

5.6 层间接口要求

各层次之间的接口应满足以下要求：

- a) 数据层与模型层之间应采用统一的数据访问接口，支持批量和流式两种数据传递模式；
- b) 模型层与服务层之间应通过标准化 API 进行模型调用，API 应支持同步和异步两种调用模式；
- c) 服务层与应用层之间应采用 RESTful API 或等效的 Web 服务接口；
- d) 各层间接口的数据交换格式宜采用 JSON 或 Protocol Buffers；
- e) 接口应具备版本管理，向后兼容，避免破坏性变更；
- f) 关键接口应具备限流、熔断和降级能力。

6 数据采集与融合

6.1 多源数据采集

系统应支持从多种来源和格式的干部管理数据中进行采集，应满足以下要求：

- a) 应通过 JDBC/ODBC 或等效标准接口接入关系型数据库，并支持 CSV/XLSX 等通用格式的数据导入导出；不绑定特定数据库产品，部署单位可根据实际情况选择符合信创要求的国产或商用数据库。
- b) 应支持半结构化数据源的接入，包括 XML、JSON、PDF 结构化表单等；
- c) 应支持非结构化数据源的接入，包括 Word 文档、扫描件、图片等；
- d) 对于扫描件和图片中的表格数据，宜采用视觉语言模型（VLM）进行智能提取。对于印刷体结构化表格，VLM 字段级提取准确率宜 $\geq 90\%$ ，关键字段（如姓名、身份证号、任职时间等）提取准确率宜 $\geq 95\%$ ；对于复杂版式或手写表格，VLM 字段级提取准确率宜 $\geq 75\%$ ，关键字段提取准确率宜 $\geq 85\%$ 。
- e) 数据采集应支持增量和全量两种模式；
- f) 应提供数据采集任务的配置管理界面，支持授权管理员配置采集规则。
- g) 干部家庭成员、社会关系等信息，应来源于干部人事档案以及《领导干部报告个人有关事项规定》所确定的个人有关事项报告等法定渠道，按规定范围程序采集；不应通过互联网爬取，不应针对干部亲属建立画像。信访举报等未经核实的线索，不得作为画像和评价计算的直接依据。

6.2 数据分类与标签

系统应具备灵活的数据分类能力，应满足以下要求：

- a) 数据类型不预设固定分类，系统应支持 AI 自动识别数据类型并注册新类型；
- b) 新增数据类型时，不应要求修改系统核心代码；

- c) 应通过标签体系实现数据的逻辑分类，标签应支持多级层次结构；
- d) 标签体系应由管理员配置和维护，支持标签的新增、修改、合并和废弃；
- e) 应提供标签使用统计，帮助管理员了解各类数据的使用情况；
- f) 数据分类分级规则应由法定责任人（或其授权人）确认；AI 系统可辅助生成密级建议，但定密、变更和解密的最终决定权归属法定责任人。
- g) 系统处理的数据规模应与业务需求相匹配，避免过度采集。
- h) 系统宜支持数据采集必要性评估，对超出业务需求的数据采集请求提供风险提示。

6.3 数据真实性与有效性核验

系统应对采集的数据进行真实性与有效性核验，应满足以下要求：

- a) 应对关键数据字段（如姓名、身份证号、任职时间等）进行格式合规性校验；
- b) 应支持跨数据源的交叉核验，标记数据不一致的情况；
- c) 宜提供数据异常检测能力，对明显异常的数据（如时间逻辑矛盾、数值超出合理范围等）进行标记和告警；
- d) 数据核验结果应形成报告，供人工审核确认；
- e) 未通过核验的数据应标记为“待确认”状态，不应直接用于画像生成和评价计算。
- f) 数据异常检测的异常识别率宜 $\geq 90\%$ ，基于不低于 500 条标注测试集（含异常样本不少于 100 条）评估，误报率宜 $\leq 10\%$ 。

6.4 数据脱敏与去标识化

系统应对敏感数据进行脱敏和去标识化处理，应满足以下要求：

- a) 应根据数据敏感等级配置脱敏策略，敏感等级划分应由管理员定义；
- b) 脱敏方式应包括但不限于：掩码替换、随机化、泛化等；
- c) 在 AI 模型推理环节，传入模型的上下文数据应经过脱敏处理或访问控制；
- d) 去标识化后的数据不应能通过合理手段还原为原始数据；
- e) 应记录脱敏操作的日志，包括脱敏规则、执行时间和操作人员。

6.5 数据生命周期管理

系统应支持数据全生命周期管理，应满足以下要求：

- a) 应定义数据的保留期限，到期数据应能自动归档或销毁；
- b) 归档数据应保留可恢复能力，恢复操作应经授权审批；
- c) 数据销毁应采用安全删除方式，确保不可恢复；
- d) 应记录数据生命周期各阶段的操作日志。

6.6 数据血缘与溯源

系统应建立完善的数据血缘追踪机制，应满足以下要求：

- a) 应记录每条数据从采集源头到最终应用的完整链路；
- b) 应支持正向溯源（从源头到应用）和反向溯源（从应用到源头）两种查询方向；
- c) 数据血缘信息应包括数据来源、处理规则、变换操作、操作人员和时间戳；
- d) 当数据发生变更时，应自动更新血缘记录；
- e) 宜提供数据血缘的可视化展示功能。

7 模型层与智能体路由

7.1 多规格模型协同

系统应支持多种规格 AI 模型的协同工作，应满足以下要求：

- a) 应支持同时管理多个不同参数规模的模型实例；
- b) 各模型应承担明确的角色分工，如意图识别、文本生成、代码生成、嵌入计算等；
- c) 应为每个模型角色注册能力标签，标签应包括但不限于：模型名称、参数规模、适用任务类型、最大上下文长度、支持的推理格式、硬件加速类型；能力标签体系应支持动态扩展，新模型接入时通过标签描述其能力特征，系统基于标签进行自动匹配和调度；
- d) 应支持模型实例的优先级配置，应建立模型推理优先级队列，高优先级请求优先调度至高规格模型；
- e) 应制定资源调度策略，根据模型负载、加速器资源占用和请求队列深度动态分配计算资源，确保关键任务不被阻塞；
- f) 模型配置信息应持久化存储，系统重启后自动恢复；
- g) 对于混合专家模型（MoE），能力标签应同时标注总参数量和单词元激活参数量；显存/内存资源分配以总参数量为主要依据，单步计算量与时延评估以单词元激活参数量为主要依据。

7.2 意图识别与路由调度

系统应提供智能意图识别和路由调度机制，应满足以下要求：

- a) 应对用户输入进行意图分类，将请求路由至最合适的模型实例；
- b) 意图分类应支持可配置的类别体系，至少包括：知识问答、画像查询、数据分析、文书生成、闲聊应答等类型；
- c) 意图识别模型宜采用轻量级模型，以降低首次响应延迟；
- d) 路由策略应支持基于规则的静态路由和基于模型预测的动态路由两种方式；
- e) 当意图识别置信度低于设定阈值时，应将请求升级至更高规格的模型处理；

f) 路由决策过程应有日志记录，便于问题排查和策略优化。

7.3 模型降级与容错

系统应具备模型降级和容错能力，应满足以下要求：

- a) 当目标模型实例不可用（如资源不足、服务异常）时，应自动降级至预设的备用模型；
- b) 降级策略应可配置，包括降级的优先级顺序和触发条件；
- c) 降级过程中应向用户明确提示当前使用的模型规格变化；
- d) 应具备模型健康检查机制，定期检测各模型实例的可用性；
- e) 模型恢复后应自动切回正常路由策略。

7.4 模型配置管理

系统应提供模型配置管理功能，应满足以下要求：

- a) 应支持模型的新增、删除、参数调整和版本切换；
- b) 模型配置变更应记录操作日志，包括操作人、变更内容和变更时间；
- c) 宜支持模型配置的版本化管理，支持回滚至历史配置；模型配置变更应支持导出备份，恢复操作应经授权审批；
- d) 应提供模型资源使用统计，包括 GPU 显存占用、CPU 利用率、推理请求量等；
- e) 应建立模型版本生命周期管理机制，明确各版本的状态（如活跃、维护、废弃），废弃版本应标记并限制新任务调度；
- f) 模型版本升级或参数重大调整前，应执行回归测试，使用标准测试集验证新版本的功能一致性和性能不低于基线水平，回归测试结果应记录存档。

7.5 智能体基本要求

对声称提供智能体能力的模块，系统应构建以大语言模型为推理核心的智能体，支持在受控范围内自主完成“感知输入与检索—推理规划—工具调用—结果输出”的多步执行闭环；自主执行范围应可配置并符合 10.7 AI 辅助边界；应设最大执行步数与时长上限，具备异常中止与人工接管能力；每次执行应生成结构化执行轨迹（目标、子任务、检索内容、工具调用及参数、中间结果、最终输出、耗时），可追溯可审计。

7.6 工具调用

应以结构化方式（如 JSON Schema）描述工具名称、功能、输入参数、输出结构与权限，由模型生成调用请求、运行时执行并回传结果；工具须经注册授权方可调用，支持白名单、参数校验、越权拦截与全量日志；智能体与工具/数据源连接宜采用 MCP 或等效开放协议，MCP 服务端应支持本地化部署，工具调用安全宜符合 GB/Z 185—2026；工具调用不得绕过既有访问控制与数据脱敏；写操作、不可逆操作、跨系统操作类工具应默认设人工确认。

7.7 智能体记忆

应区分短期记忆（会话上下文与中间状态）与长期记忆（跨会话偏好、历史结论）；记忆内容支持查看、更正、删除，符合留存期限与个人信息保护要求；写入与检索纳入访问控制审计；敏感信息入长期记忆前应脱敏或授权确认；长期记忆召回宜采用向量检索，与第9章 RAG 复用统一检索与溯源能力。

7.8 任务分解与规划

复杂任务应分解为有序子任务并生成执行计划，子任务清单及依赖关系对用户可见；宜采用 ReAct 等“推理—行动”模式依据工具返回动态调整后续步骤；应设最大步数、失败重试与重新规划机制；用户可在执行前/中查看、修改、暂停、终止计划；关键节点与不可逆操作须人工确认；宜在执行后自检反思，但自检不得替代 10.7 的人工确认。

7.9 多智能体编排

宜支持多智能体协同，提供主管—专家/分工协作等编排；智能体间任务传递应明确目标、输入输出契约与权限边界；跨智能体互操作宜遵循 GB/Z 185—2026 及 A2A 等开放协议；完整调用链可追踪可审计；跨智能体工具调用满足 7.6；不得以多智能体协同为由弱化 AI 辅助边界与人工确认。

7.10 智能体评估

应建评测集，指标宜含任务完成率、工具调用成功率与参数正确率、规划合理性、所需步数/轮次、幻觉率、引用一致率；覆盖单工具、多工具串联、多步规划、多智能体协同等难度层级；版本升级按 7.4 f) 回归测试，关键指标不低于基线方可发布。

8 知识图谱与关系推理

8.1 图谱构建

系统应支持干部相关知识图谱的构建和维护，应满足以下要求：

- a) 核心实体包括干部、组织、职位三类，系统应支持；其他实体（如文件、地点、项目、事件等）为可扩展实体，各单位可根据业务需要按需启用；
- b) 应支持多种关系类型的定义，包括结构化关系（如任职关系、隶属关系）和推理关系（如同事关系、同学关系等）；
- c) 实体属性应支持灵活扩展，新增属性字段不应影响已有数据结构；
- d) 应支持从非结构化文本中自动抽取实体和关系，抽取结果应标注置信度；
- e) 图谱数据应支持批量导入和增量更新；
- f) 图谱存储应支持本地化部署，可采用关系型数据库、原生图数据库或其他支持图数据模型的系统；使用图数据库时应支持本地或私有化部署，不强制依赖厂商云服务；不绑定特定产品，数据模型与查询接口与存储实现解耦。

8.2 关系建模

系统应提供灵活的关系建模能力，应满足以下要求：

- a) 关系定义应包含：关系类型、方向性、属性集合、约束条件；
- b) 应支持关系的层次结构，如“上下级”关系可细分为“直接上级”“间接上级”等；
- c) 应支持关系的时序属性，记录关系的生效时间和失效时间；
- d) 推理关系应明确标注推理规则和置信度评分；
- e) 关系建模应限于组织隶属、任职关联、培训同期等业务关系，不应用于推断政治倾向、宗教信仰等敏感属性。其他环节的敏感属性禁止要求见 14.5 j)。

8.3 关系推理

系统应基于知识图谱提供关系推理能力，应满足以下要求：

- a) 应支持路径查询，支持查询两实体之间的关联路径；
- b) 应支持传递性推理，如通过组织隶属关系推导间接关系；
- c) 推理结果应标注推理链路和置信度；
- d) 推理规则应由管理员配置和管理，支持规则的版本化管理；
- e) 应提供推理结果的人工确认机制，置信度低于阈值的推理结果应标记为“待确认”。

8.4 置信度管理

系统应对图谱中的各类信息建立置信度管理机制，应满足以下要求：

- a) 每条实体、关系和推理结果均应关联置信度评分；
- b) 置信度评分应基于数据来源可靠性、交叉验证结果和推理链路长度等因素综合计算；
- c) 应设置置信度阈值，低于阈值的图谱信息在画像展示和评价计算中应进行标记或排除；
- d) 置信度计算规则应由管理员配置，支持动态调整。

8.5 图谱健康管理

系统应提供图谱健康管理功能，应满足以下要求：

- a) 应定期检测图谱数据的一致性，识别孤立节点、悬挂关系等异常情况；
- b) 应提供图谱统计信息，包括实体数量、关系数量、类型分布等；
- c) 应支持图谱数据的备份和恢复；
- d) 宜提供图谱可视化的查询和浏览工具。

8.6 图谱质量要求

知识图谱的构建和推理结果应满足以下质量要求：

- a) 实体识别 F1 值宜 $\geq 75\%$ ；

- b) 关系分类准确率宜 $\geq 70\%$;
- c) 推理路径正确率宜 $\geq 80\%$;
- d) 应使用标注数据集定期评估图谱质量, 评估结果应记录存档;
- e) 当质量指标低于阈值时, 应触发人工复核或模型调优流程。

f) 上述质量指标的评估应基于规模不低于 500 条的标注测试集, 测试集应覆盖主要实体类型和关系类型, 标注应由不少于 2 名独立标注人员完成, 标注一致性 (Cohen's Kappa) 不低于 0.75。

实体抽取评测应采用严格 span 匹配; 报告 micro-F1 与各类 F1、macro-F1; 每类实体测试实例不少于 100 个; 附 bootstrap 95%置信区间。关系抽取评测: Macro-F1 ≥ 0.70 ; 每类关系测试实例不少于 50 个; 关系方向应正确方计入。

9 检索增强生成服务

9.1 混合检索

系统应采用混合检索策略以提升检索质量, 应满足以下要求:

- a) 应同时支持向量语义检索和关键词检索两种方式;
- b) 向量检索应使用嵌入模型将查询和文档映射至向量空间, 支持相似度排序;
- c) 关键词检索宜采用 BM25 或同等效果的算法;
- d) 应采用倒数排序融合 (RRF) 或等效算法将两路检索结果进行融合排序;
- e) 宜在融合排序后增加精排环节, 使用重排序模型 (reranker) 对候选结果进行精细排序;
- f) 向量数据库应支持本地化部署, 选型不应绑定特定厂商产品。
- g) RAG 服务的整体设计和实现宜参考 GB/Z 253—2026《人工智能 检索增强生成通用技术要求》的相关规定。

9.2 检索质量要求

检索服务应满足以下分阶段质量要求:

- a) 基础验收线: 检索结果的准确率 (Precision) 宜不低于 75%, 召回率 (Recall) 宜不低于 70%;
- b) 优化目标线: 经查询改写、重排序等优化后, 检索结果的准确率 (Precision) 宜不低于 85%, 召回率 (Recall) 宜不低于 80%, 达到“良好”等级;
- c) 应支持查询改写 (query rewriting), 对用户原始查询进行语义增强;
- d) 应支持上下文压缩 (context compression), 对检索到的长文本进行信息浓缩, 减少无关信息对生成质量的影响;
- e) 应提供检索质量的持续评估机制, 定期使用标注数据集对检索效果进行评测;
- f) 检索质量的评测应区分基础验收线与优化目标线两个层次, 分别统计和报告。

g) 检索质量评估的测试集规模应不低于 500 条查询，应覆盖知识问答、画像查询、数据分析等主要查询类型，并包含 10%~20%无答案查询；相关性采用二元标注（2=相关，1=边际，0=不相关），以 2 分为裁；评测应在明确截断位置（建议 P@5、R@10）下计算；同时报告答案正确率（宜 $\geq 80\%$ ）与引用一致率（宜 $\geq 90\%$ ）。

9.3 知识增量更新

系统应支持知识库的增量更新，应满足以下要求：

- a) 新增文档应自动触发向量化和索引构建；
- b) 文档更新时应自动删除旧版本索引并建立新版本索引；
- c) 索引更新过程中不应影响在线检索服务；
- d) 应提供索引状态监控，包括文档数量、索引大小、索引健康状态等。

9.4 来源引用与溯源

系统应在生成结果中标注信息来源，应满足以下要求：

- a) RAG 生成的每一条关键事实应标注其来源文档和位置；
- b) 来源引用应包含文档名称、段落位置或页码信息；
- c) 用户应能通过引用链接回溯到原始文档；
- d) 当检索结果置信度不足时，生成内容应明确标注“仅供参考”或类似提示。

9.5 图谱增强检索（GraphRAG）

系统宜利用知识图谱增强检索能力，应满足以下要求：

- a) 宜支持基于知识图谱的关联检索，通过实体和关系路径补充向量检索的语义信息；
- b) 应支持多跳关联查询，能够在图谱上沿关系路径进行多步推理和检索；
- c) 图谱检索结果应与向量检索结果进行融合排序，融合策略应可配置；
- d) 图谱检索应支持查询路径的可视化展示，便于用户理解检索推理过程；
- e) 宜支持子图提取，将查询相关的局部图谱结构作为上下文传递给生成模型。

9.6 自适应检索策略

系统宜具备自适应检索能力，应满足以下要求：

- a) 首次检索后，系统宜对检索结果的置信度进行评估；
- b) 当首次检索结果的置信度低于设定阈值时，系统宜自动触发补充检索，补充检索策略包括但不限于：查询改写、扩展检索范围、切换检索源；
- c) 自适应检索的最大重试次数应可配置，防止无限循环检索；
- d) 补充检索的过程和结果应记录日志，便于检索策略的优化和审计；

e) 宜支持检索策略的自学习，根据历史检索效果持续优化检索参数。

9.7 生成质量保障

RAG 生成环节应建立质量保障机制，应满足以下要求：

a) 生成内容的忠实度（Faithfulness）宜不低于 85%，即生成内容中与检索结果一致的事实性陈述占全部事实性陈述的比例；

b) 生成内容的答案相关性（Answer Relevance）宜不低于 80%，即生成内容中与用户问题相关的信息占全部生成信息的比例；

答案相关性的测量宜采用反推问题—嵌入相似度法（由答案反推若干可能问题，计算与原始问题向量相似度均值），或经标定的等价人工/模型评分。

c) 应提供忠实度和答案相关性的评估方法，支持基于标注数据集的自动化评测；

d) 当生成内容的质量指标低于阈值时，系统应标注风险提示或拒绝生成；

e) 宜支持对生成内容进行事实一致性校验，与知识库中的已知事实进行交叉验证。

f) 忠实度和答案相关性的评估测试集规模应不低于 300 条问答对，应覆盖事实性问答、分析性问答和综合推理等不同难度级别；评估应区分检索正确率、答案正确率和引用一致率三个维度；采用 LLM-as-judge 评测时，应固定 judge 模型版本与评分 prompt，拆分原子陈述，并抽样与人工评分对齐（Kappa $\geq$ 0.8）；报告 bootstrap 95%置信区间。

g) 对涉及具体干部评价、任用倾向、评优处分倾向的提问，系统仅可基于已经核实确认的数据陈述事实，并提示“相关事项由组织按照规定程序研判”，不应生成结论性、推荐性回答，不应输出针对个人的排名、推荐名单或敏感属性推断。

9.8 文档分块策略

系统应支持可配置的文档分块策略，应满足以下要求：

a) 应支持固定长度分块，按设定的字符数或 Token 数进行均匀分块；

b) 宜支持语义边界分块，基于段落、章节等自然语义边界进行分块；

c) 宜支持结构化段落分块，根据文档结构（如标题层级、表格边界）进行智能分块；

d) 分块策略应支持按文档类型分别配置，不同格式的文档可采用不同的分块方案；

e) 分块参数（如块大小、重叠比例）应可配置，并支持通过检索效果评测进行参数调优；

f) 宜支持分块结果的质量评估，检测分块过大、过小或语义断裂等异常情况。

10 画像生成与评价服务

10.1 个人画像

系统应支持干部个人全息画像的生成和展示，应满足以下要求：

- a) 画像应涵盖多维度信息，维度体系由组织根据干部管理需要自行配置；
- b) 画像内容应基于数据层采集的实际数据生成，不应包含无数据支撑的推测性结论；
- c) 画像中的每一条信息应可溯源至原始数据；
- d) 画像应支持时间维度，展示干部信息的动态变化历程；
- e) 画像展示应区分“系统自动生成内容”和“人工标注内容”，两者应有明确的视觉区分；
- f) 应支持画像的导出，导出格式宜支持 WORD、PDF 和 HTML；
- g) 宜支持画像维度间的交叉分析和关联发现，如考核评价与培训经历之间的关联、履职经历与能力素质之间的映射等；
- h) 交叉分析结果应标注数据来源和分析方法，不应生成无数据支撑的关联性结论。

10.2 群体画像

系统应支持群体画像分析功能，应满足以下要求：

- a) 应支持按组织、层级、专业、年龄等维度进行群体统计分析；
- b) 群体统计结果应以图表形式直观展示；
- c) 群体画像应支持对比分析，如不同时期、不同部门的横向和纵向对比；
- d) 群体画像中的个体数据在展示时应遵守数据脱敏要求；
- e) 宜支持梯队分析功能，按组织管理需要分析干部队伍梯队结构和后备力量的群体统计分布；梯队分析应以不针对具体个人的统计形式呈现，不应输出指向特定干部的后备名单、推荐顺序或培养、使用建议。
- f) 宜支持班子结构分析功能，分析领导班子的年龄结构、专业结构、经历互补性等指标；
- g) 宜支持趋势预测功能，基于历史数据对群体结构变化趋势进行分析，为干部队伍建设规划提供参考。
- h) 梯队分析和趋势预测结果应标注为 AI 辅助分析，仅供参考。
- i) 趋势预测输出应说明所采用的统计方法，并附置信区间或误差范围。
- j) 个人画像中的“发展潜力”类内容，仅可展示经组织认定的结论和有数据支撑的事实轨迹，不应由系统生成个人潜力评级、成长预测结论或提拔、培养建议。

10.3 评价模型配置

系统应提供可配置的评价模型，应满足以下要求，参见附录 C：

- a) 评价维度、权重和评分规则应由组织管理员自行配置，系统不预设统一的评价标准；
- b) 评价模型配置应支持版本化管理，每次配置变更应生成新版本；
- c) 配置变更应记录操作人、变更内容、变更时间和变更原因；
- d) 应支持评价模型的回滚，可恢复至任一历史版本；

e) 评价模型配置应支持导入和导出，便于不同组织之间的经验交流；

f) 评价计算过程应透明可审计，每一步计算应可追溯；

g) 对人工智能承担的数据提取、特征归纳环节，应提供特征归因分析，说明各输入特征对归纳结果的贡献方向和程度，宜支持 SHAP (SHapley Additive exPlanations)、LIME (Model-agnostic) 等方法；对大语言模型承担的归纳环节，宜采用基于提示的归因或注意力分析方法。注意力权重仅在经梯度或留一法验证相关后方可作解释性参考，不宜单独作归因依据。人工智能归纳形成的特征，经人工确认前不得进入评价计算。

h) 归因分析结果应向复核人员展示，便于其理解 AI 归纳内容的依据，并与人工确认记录关联留存。

i) 应对基于 AI 的评价结论留存期限进行配置，超过留存期限的评价数据应自动归档或删除；

j) 系统应建立 AI 辅助输出的人工确认机制，明确哪些类型的 AI 输出应经人工确认后方可作为正式决策依据。

k) AI 辅助生成的评价结论应标注其依据数据的覆盖度和置信度，当依据数据不完整时应明确标注。

l) 意图识别 Macro-F1 宜 ≥ 0.85 ；应报告混淆矩阵与级联错误率。

10.4 评价规则版本管理

系统应建立完善的评价规则版本管理机制，应满足以下要求：

a) 每个评价规则版本应包含完整的维度定义、权重配置和计算逻辑；

b) 版本切换时，应明确标注当前生效版本和历史版本；

c) 当评价规则发生变更时，宜提供对存量评价结果的影响分析功能，说明规则变更对已评价对象的影响程度；

d) 对于影响范围较大的规则变更（如维度增减、权重大幅调整），系统应提示管理员进行存量数据的重新评估；

e) 应提供不同版本评价结果的对比功能。

10.5 异常偏差检测与人工复核

系统应提供异常偏差检测机制，支持人工复核，应满足以下要求：

a) 当 AI 评价结果与历史趋势或同类干部平均水平偏离较大时，系统应自动标记并发出告警，偏离检测宜支持与同层级、同类型岗位干部群体的横向比较；

b) 偏离阈值应由管理员配置，支持按维度和指标分别设定；

c) 被标记为偏离较大的评价结果，在进入正式流程前应经人工复核确认；

d) 人工复核应记录复核人、复核时间、复核结论和处理意见；

e) 复核记录应纳入审计日志，长期保存。

10.6 群体公平性分析

系统宜提供群体公平性分析功能，应满足以下要求：

- a) 宜支持按性别、年龄、岗位类别等维度检测评价结果中是否存在系统性偏差；
- b) 公平性检测应基于统计方法，比较不同群体的评价分数分布差异；
- c) 当检测到显著性偏差时，系统应发出告警并提示管理员排查原因；
- d) 公平性分析的结果应记录并纳入评价审计日志；
- e) 公平性检测的维度和阈值应由管理员配置。

f) 公平性检测的样本量应不低于各子群体各 50 条记录，检测应同时报告各组评价分数的均值差异和分布差异。

10.7 AI 辅助边界

系统应明确界定 AI 辅助决策的边界，应满足以下要求：

- a) AI 生成的所有分析结果、评价建议和数据洞察仅供参考，不具有决策效力；
- b) 系统应在所有 AI 生成内容的显著位置标注“AI 辅助分析结果，仅供参考”或等效提示；
- c) 系统不应生成干部任免、岗位调整、职级晋升、评优评先、培训推荐、处分等直接决策建议，只应提供数据支撑和参考分析；
- d) 涉及干部评价的结论性内容，应设置人工确认环节，未经人工确认的 AI 评价结果不应对外输出；
- e) AI 模型不应参与最终的排序、打分或等级划定，评价计算中 AI 的作用限于数据提取、特征归纳和辅助计算，但“辅助计算”不包括等次评定、优劣排序等评价性判断；
- f) 不应仅依据 AI 分析结果做出干部管理相关决定，任何干部管理决策均应有组织人事部门的人工研判环节；
- g) 系统应提供 AI 参与度的透明展示，用户应能了解每个评价环节中 AI 的具体贡献；
- h) 应建立 AI 辅助决策的责任追溯机制，明确每个环节的责任主体。

i) 系统应建立高风险 AI 输出清单并实施强制人工核验，至少包括：

- (1) 干部评价结论性内容；
- (2) 考察材料、监督类文书的评价性段落；
- (3) 异常偏离结果；
- (4) 拟进入正式报告的群体分析、梯队分析和趋势预测结论。

上述内容经人工核验确认前，不得以任何形式（含内部报送、领导参阅）作为正式材料使用。系统应对“内部研判稿”与“正式文稿”分别标注。

11 智能文书服务

11.1 文书模板管理

系统应提供文书模板管理功能，应满足以下要求：

- a) 应支持文书模板的上传、编辑、审核和发布；
- b) 模板应支持变量占位符定义，变量与数据层字段建立映射关系；
- c) 应支持模板版本管理，保留历史版本并支持回滚；
- d) 应支持多种文书格式，包括但不限于 Word（.docx）、PDF 等；
- e) 模板管理应由管理员负责，模板发布前应经审核确认。

11.2 文书生成

系统应支持基于模板和数据自动生成文书，应满足以下要求：

- a) 文书生成应基于模板定义和数据层中的实际数据，不应生成无数据支撑的内容；
- b) AI 辅助生成的文书内容（如考察材料中的评价性段落）应以显著方式标注为 AI 生成初稿；在承办人实质性审核、修改并确认前，不应作为考察结论和干部选拔任用的依据。
- c) 文书生成过程应支持人工预览和修改，修改后重新生成应保留修改记录；
- d) 生成的文书应保留与源数据的关联，便于核验；
- e) 可支持监督类文书的起草辅助，包括干部监督情况报告、廉政意见回复等文书类型的事实归集、条款检索和排版。
- f) 监督类文书应严格基于已经核实确认的数据和事实记录；系统不应生成无数据支撑的内容，不应代替有权机关作出定性意见、廉政评价或者处理建议。监督类文书应经有权机关或者承办人按照规定程序审核签发后方可生效，审核签发记录应纳入审计。
- g) 应支持至少 5 种文档格式的解析，包括 OFD、PDF、Word、Excel 和纯文本格式。

11.3 AI 生成内容标识

系统应对 AI 生成的内容进行标识，并符合《人工智能生成合成内容标识办法》和 GB 45438—2025 的规定：

- a) AI 生成的文书内容应以可识别方式添加显式标识，如文字提示、角标、脚注等；文本显式标识应包含“人工智能（AI）”和“生成（合成）”要素，位于文本起始、末尾或中间适当位置；
- b) AI 生成内容的文件元数据中应添加隐式标识，至少包含生成合成内容属性信息、服务提供者名称或者编码、内容唯一编号；鼓励采用数字水印；
- c) 标识信息应包含生成模型、生成时间和生成依据；文书下载、复制、导出时，文件中应同时含有符合要求的显式标识和隐式标识；
- d) 干部管理文书原则上不得去除显式标识。确因特殊使用场景需提供不含显式标识版本的，应在书面协议中明确使用人的标识义务与使用责任，经审批后提供，相关日志留存不少于六个月，且隐式标识不得去除；
- e) 任何用户不应恶意删除、篡改、伪造、隐匿 AI 生成内容标识；系统应保留标识添加与变更记录。

11.4 公文格式合规

系统生成的公文类文书应符合国家公文格式标准，应满足以下要求：

- a) 公文版面格式应符合 GB/T 9704—2012 的规定；
- b) 公文字体、字号、行间距等排版参数应可配置；
- c) 公文要素（发文字号、签发人、成文日期等）应自动填充并支持人工修正；
- d) 应提供公文格式合规性检查功能，在文书导出前进行格式校验。

11.5 文书质量要求

文书生成服务应满足以下质量要求：

- a) 关键字段填充准确率宜不低于 99%，关键字段包括但不限于姓名、职务、任职时间、文件编号等结构化字段；
- b) 公文格式合规率宜不低于 95%，即生成的公文通过格式校验的比例；
- c) AI 生成内容的事实准确率宜不低于 90%，即 AI 生成的事实性陈述与源数据一致的比例；
- d) 应提供文书质量的自动化检测方法，支持对关键字段、格式合规性和事实准确性的批量检测；
- e) 文书质量检测结果应记录并纳入系统质量管理日志。
- f) 文书质量检测的测试样本应不少于 50 份文书、关键字段实例不少于 1500 个，覆盖主要文书类型；字段准确率判据：点估计 $\geq 99\%$ 且 95% Wilson 置信下界 $\geq 98\%$ ；归一化规则（去空白、统一日期数字格式、全半角）须书面化；源数据为空的可选字段不计入分母；关键字段填充准确率的计算应区分结构化字段和自由文本字段。

12 workflow编排服务

12.1 workflow定义与执行

系统应提供可视化的 workflow编排功能，支持通过拖拽或配置方式定义多步骤业务流程，应满足以下要求：

- a) 应支持流程节点的定义，包括数据输入、处理步骤、条件分支、并行执行和结束节点；
- b) 应支持流程实例的创建、执行、暂停、恢复和终止；
- c) 应对流程执行过程进行状态跟踪，记录每个节点的执行状态、耗时和输出结果；
- d) 应对流程定义进行版本化管理，支持流程定义的发布、回滚和历史版本查看。

12.2 同步与异步模式

workflow编排应支持同步和异步两种执行模式，应满足以下要求：

- a) 同步模式适用于实时交互场景，系统应在单次请求响应周期内完成流程执行并返回结果；
- b) 异步模式适用于长流程或跨部门协作场景，流程执行结果应通过消息通知、回调或轮询方式获取；

- c) 应支持同步模式与异步模式的混合编排，同一流程中可包含同步子步骤和异步子步骤；
- d) 异步流程应提供执行进度查询接口，支持用户了解当前执行状态。

12.3 执行监控与审计

workflow编排服务应满足以下监控与审计要求：

- a) 应提供流程执行的实时监控视图，展示当前执行位置、各节点状态和异常信息；
- b) 应对流程执行记录进行持久化存储，保留时间应满足组织档案管理要求且不少于 3 年。凡构成干部人事档案归档材料的，保管依《干部人事档案工作条例》执行，系统留存、自动销毁规则不得与之抵触。
- c) 应支持流程执行日志的检索和导出，便于事后审计；
- d) 应对流程中的敏感数据操作进行标记和专项审计。

13 服务接口与应用

13.1 REST API 规范

系统应提供符合 RESTful 架构风格的 API 接口，应满足以下要求：

- a) API 设计应采用 RESTful 风格，使用标准 HTTP 方法（GET、POST、PUT、DELETE）对应资源操作；
- b) 应提供清晰的状态码映射，成功响应使用 2xx 系列状态码，客户端错误使用 4xx 系列状态码，服务端错误使用 5xx 系列状态码；
- c) 对于长耗时操作，应支持异步模式：接口返回任务标识（taskId）和处理状态，客户端通过查询接口获取最终结果；
- d) 应支持请求参数校验，对不合法输入返回明确的错误信息和错误码；
- e) 应提供 API 文档（如 OpenAPI/Swagger 格式），包含接口路径、请求方法、参数说明和示例；
- f) API 版本管理应兼容旧版本接口，废弃接口应保留至少一个版本周期。

13.2 应用集成要求

系统应支持与其他系统的集成，应满足以下要求：

- a) 应提供标准化的数据交换接口，支持与其他干部管理信息系统的数据对接；
- b) 集成接口应支持增量数据同步和全量数据同步；
- c) 应支持基于消息队列的异步数据推送；
- d) 集成过程应遵循最小权限原则，只暴露必要的数据和接口；
- e) 宜支持中组部干部信息采集标准格式的数据对接，便于与上级干部管理信息系统的数据交换；
- f) 数据交换格式应提供标准化的映射配置能力，支持不同系统间字段映射的灵活定义。

13.3 前端服务要求

系统前端应满足以下要求：

- a) 应采用前后端分离架构；
- b) 页面首屏渲染时间 P95 宜不超过 3 秒，复杂查询端到端响应时间 P95 宜不超过 10 秒；
- c) 前端应兼容符合国家相关标准的主流浏览器；信创环境宜通过主流国产浏览器（列入信创目录的产品）验证。
- d) 应具备基本的无障碍访问能力；无障碍设计可参照 GB/T 37668—2019 的相关规定执行。
- e) 界面设计应简洁规范，关键操作应有确认提示和操作指引。

13.4 并发与响应性能要求

系统应满足以下并发与响应性能要求。性能验收应在以下基准条件下进行：流式 SSE 输出、预热不少于 50 请求、temperature=0、最大在途=额定并发、软硬件版本/量化精度/网络 RTT/RAG 工况（chunk/top-k/rerank）/测量点与时钟同步逐项列明；向量索引维度不超过 1024（768 为典型值）。基准数值的具体取值作为资料性参考，不构成要求。超出基准条件时，系统应提供性能衰减说明。

- a) 各类性能指标应分别度量：首词元延迟（TTFT）、逐词元延迟（TPOT）、端到端响应时间（E2E）和系统吞吐量（完成率 λ ）；平均在途请求数 L 应满足 $L = \lambda \times W$ （ W 为平均等待+服务时间），验收须验证该式偏差 $\leq 10\%$ ；
- b) 上述指标应分别报告 P50 和 P95 统计值；离线吞吐单列，按 tok/s、tok/s/卡、tok/s/W 单位报告；
- c) 轻量级（主模型 $\leq 7B$ ）：输出 ≤ 1024 tokens 时 E2E P95 ≤ 30 s，首词元延迟 P95 ≤ 5 s；标准负载并发 5 用户、思考时间 10~30 s、闭口泊松到达；
- d) 标准级（主模型 $\leq 14B$ ）：输出 ≤ 512 tokens 时 E2E P95 ≤ 15 s，首词元延迟 P95 ≤ 3 s；标准负载并发 10 用户、思考时间 10~30 s；
- e) 全功能级（含 32B）：输出 ≤ 256 tokens 的短答类 E2E P95 ≤ 10 s，首词元延迟 P95 ≤ 2 s；1024~2048 tokens 长文类统一异步，报告 P50/P95 完成时间；标准负载并发 20 用户；
- f) 供方应提交“输出长度（256/512/1024/2048 tokens） $\times$ 并发”实测 E2E 曲线及 RAG 各阶段耗时分解；
- g) 并发访问时，各用户的请求应得到公平调度，不应出现请求饥饿或死锁现象；
- h) 宜提供并发性能的监控和统计功能，支持在线并发数、响应时间分布、TTFT 和 TPOT 等指标的实时查看。

13.4.1 基准测试工况应固定：流式 SSE 输出；预热不少于 50 请求；采样参数 temperature=0、top_p 固定；最大在途请求数=额定并发；软硬件版本、量化精度（FP16/INT8/INT4）、网络 RTT、RAG 工况（chunk 大小、top-k、是否 rerank）逐项列明；测量点采用客户端时间戳，服务端与客户端时钟同步误差 $\leq 5ms$ ；超时重试应单列不计入基准。

13.4.2 术语定义：

- a) 首词元延迟 TTFT：从请求发出到首个内容词元返回的时延，不含 BOS、不含思考 token（含思考模型的“首思考 token”与“首答案 token”应分别报告）；非流式接口无法直接测量，应采用引擎内部时间戳。
- b) 逐词元延迟 TPOT：生成阶段相邻输出词元的平均间隔，单位 ms/token；按单请求与跨请求聚合分别报告，单位与测量点（引擎侧/客户端）应明示；排队抖动说明应附 P50/P95/P99。

13.5 系统可用性要求

系统应满足以下可用性要求：

- a) 系统月度可用性应不低于 99.5%；
- b) 系统应支持连续运行不低于 72 小时无中断；
- c) 部分降级状态下可用性应不低于 97%；
- d) 系统应具备故障自动恢复能力，关键服务异常后宜能在 60 秒内自动重启；
- e) 宜提供可用性统计报告，记录系统运行时间、中断次数和恢复时间等指标。

14 安全、可信与合规

14.1 访问控制

系统应建立完善的访问控制机制，应满足以下要求：

- a) 应实施基于角色的访问控制（RBAC），支持多角色、多权限的灵活配置；
- b) 应支持基于属性的访问控制（ABAC），根据数据密级、用户属性等条件动态控制访问权限；
- c) 应支持最小权限原则，用户默认不获得任何权限，由管理员按需授权；
- d) 应支持权限的批量管理和继承关系；
- e) 敏感数据和高危操作应支持多人审批机制；
- f) 应支持字段级细粒度权限控制，可配置不同角色对不同数据字段的可见性和可操作性；
- g) 应支持记录级细粒度权限控制，可配置不同角色对特定数据记录的访问范围，如按组织层级限定数据可见范围。

14.2 认证与鉴权

系统应提供安全的认证与鉴权机制，应满足以下要求：

- a) 应支持用户名密码认证，密码策略应包括最小长度、复杂度要求和有效期管理；
- b) 应支持多因素认证，与组织现有的统一身份认证系统对接；
- c) 认证令牌应有时效性，过期自动失效；
- d) 应具备登录失败锁定机制，连续失败次数超过阈值时自动锁定账户；
- e) 宜支持国密算法（SM2/SM3）用于认证过程中的加密和摘要计算。

14.3 审计与追溯

系统应建立全面的审计追溯机制，应满足以下要求：

- a) 应记录用户登录、数据访问、数据修改、文书生成、评价计算、系统配置变更等关键操作；
- b) 审计日志应包含操作人、操作时间、操作类型、操作对象、操作结果和来源 IP 等信息；
- c) 应支持日志分级管理，根据日志内容的数据敏感度设定不同的访问权限和密级。审计日志应不可篡改，支持长期保存；
- d) 应提供审计日志查询和分析功能，支持按时间、用户、操作类型等维度进行筛选；
- e) 应支持审计日志的导出和归档；
- f) AI 生成过程的审计应记录输入提示词、模型选择、输出内容和人工修改情况。

14.4 内容安全

系统应具备内容安全保障能力，应满足以下要求：

- a) 应对 AI 生成的输出内容进行安全过滤，防止生成不当内容；
- b) 应建立敏感词库，敏感词库应支持动态更新；
- c) 用户输入应进行安全检测，防止 SQL 注入和 XSS 等传统 Web 攻击；
- d) 应对 Prompt 注入和越狱攻击进行专项防护，对用户输入进行语义级安全检测，防止通过构造恶意提示词绕过系统安全限制；
- e) 应提供内容安全事件的告警和处置功能；
- f) 应建立政治安全审查机制，对 AI 生成内容中涉及政治敏感性的表述进行识别、过滤和标记。过滤规则应由组织审定，误拦截情形应支持人工复核和申诉；该机制不得替代干部管理工作中正常的组织审查程序。
- g) 应具备事实性核验能力，对 AI 生成的关键事实性陈述应与知识库和源数据进行交叉验证，降低幻觉风险；
- h) 高风险场景（如评价结论生成、文书自动填充等）应设置强制核验环节，未经核验的 AI 生成内容不应直接输出；
- i) 应支持内容安全过滤规则的分级配置，可根据使用场景和安全要求调整过滤严格程度；
- j) 内容安全过滤规则应支持版本化管理，规则变更应记录操作日志。

14.5 数据安全

系统应满足数据安全保护要求：

- a) 数据分类分级应符合《数据安全法》《个人信息保护法》《网络数据安全管理条例》的规定，并参照 GB/T 37988—2019 加强安全能力建设。应建立国家秘密、工作秘密、内部敏感、一般数据四级分类体系，明确范围、规则和访问权限；应识别汇聚、关联后可能属于国家秘密或者重要数据的数据，依法保护；识别为重要数据的，应明确网络数据安全负责人，每年度开展风险评估并报送报告。

b) 个人信息处理应符合《中华人民共和国个人信息保护法》和 GB/T 35273—2020 的要求。党政机关为干部管理目的处理干部个人信息，应以《个人信息保护法》第十三条第一款第（三）项规定的“为履行法定职责或者法定义务所必需”为主要合法性基础，并遵守该法第三十四条规定，不得超出履行法定职责所必需的范围和限度；不得将“取得个人同意”设置为干部考核、任免、监督等管理活动的前置条件，法律法规另有规定的除外。系统应在处理活动开展前完成合法性基础核查并留存记录。

c) 数据传输应采用加密通道，应采用国密算法（SM2/SM3/SM4）进行加密和完整性保护；

d) 数据存储应采用加密存储，密钥管理应符合 GM/T 0054 的规定，建立密钥全生命周期管理机制，涵盖密钥的生成、存储、分发、使用、更新、归档和销毁各环节；

e) 应建立数据备份机制，支持定期备份和灾难恢复；

f) 有下列情形之一的，应事前开展个人信息保护影响评估：处理敏感个人信息；利用个人信息进行自动化分析、决策；委托处理或者向其他处理者提供、公开个人信息；向境外提供个人信息；首次开展或目的方式重大变更；其他对个人权益可能产生重大影响的活动。评估报告和处理情况记录至少保存三年。

g) 画像结果应设定合理的留存期限，超出留存期限的画像数据应及时删除或匿名化处理；

h) AI 推理环节应遵循数据最小化原则，传入模型的上下文数据应限于与当前任务直接相关的必要信息，不应传入与任务无关的个人信息；

i) 敏感个人信息（如身份证件号码、家庭成员信息等）的管理应遵循更严格的保护措施，包括但不限于加密存储、访问审计和使用审批。

j) 系统不应通过数据汇聚、关联分析或者模型推断，生成、输出干部的政治倾向、宗教信仰、民族属性、健康状况（含心理健康）、生物特征识别信息、行踪轨迹等敏感个人信息，不应生成政治忠诚度评分，不应以算法对干部政治表现自动评分或者画像分级。干部政治素质、道德品行等“德”的评价结论，只能由组织按照规定程序作出。法律法规另有规定的，从其规定。

k) 处理干部个人信息应按照《个人信息保护法》第十七条、第三十五条履行告知义务，依法不予告知的除外。

l) 干部对 AI 辅助形成的评价结论，有权要求说明数据来源、评价规则和计算依据，有权提出异议并申请组织复核；对干部权益有重大影响的决定，不得仅依据 AI 自动化分析结果作出。

14.6 等保与分级保护

系统应满足网络安全等级保护相关要求：

a) 系统部署前应完成等级保护定级和备案，安全保护等级应不低于第三级；

b) 安全技术要求应符合 GB/T 22239—2019 第三级及以上的规定；

c) 安全设计应符合 GB/T 25070—2019 第三级及以上的规定；

d) 系统存储、处理、传输涉及国家秘密信息的，应按照国家涉密信息系统分级保护有关规定执行。涉密信息系统分级保护与网络安全等级保护属于不同管理制度，不得相互替代；本系统作为非涉密系统建设运行时，不应存储、处理国家秘密。

e)应定期开展等级保护测评，测评周期不应超过一年。

f)便携式部署不得承载国家秘密；确需承载工作秘密的应经保密评估后方可启用。对于便携式部署场景（如单台工作站临时部署用于演示或试点验证），在确保数据安全和访问控制有效的前提下，可适当降低硬件冗余和集群要求，但固定生产环境部署应满足等保三级的全部要求。

g)系统应按照《密码法》和 GB/T 39786—2021 开展商用密码应用安全性评估；第三级等级保护对象应定期评估，评估结果作为验收和持续运行依据。

14.7 AI 模型安全

系统应建立 AI 模型安全保障机制，应满足以下要求：

a) 应建立模型安全对齐机制，确保模型输出内容与社会主义核心价值观和干部管理政策要求保持一致；

b) 应建立 Prompt 注入和越狱攻击的防御机制，对用户输入进行安全检测和过滤，防止通过构造恶意输入绕过系统安全限制；

c) 宜定期开展红队测试，模拟对抗攻击场景，检测模型的安全防护能力；

d) 应建立模型文件的完整性校验机制，确保模型文件在存储、传输和加载过程中未被篡改；

e) 模型供应链应满足安全要求，模型文件的来源应可追溯，宜建立模型来源验证和可信签名机制；

f) 应建立模型安全事件报告和处置流程，发现模型安全漏洞或异常行为时应及时响应；

g) 高风险场景下宜对模型输入和输出进行双向安全检测，防止通过输入诱导不当输出。

14.8 供应链安全

系统应建立软件供应链安全管理机制，应满足以下要求：

a) 应建立软件物料清单（SBOM），记录系统所依赖的全部第三方组件、库和框架的名称、版本、来源和许可证信息；

b) AI 模型上线前应进行安全评估，评估内容包括模型来源可信度、训练数据安全性、模型行为一致性等；

c) 应建立依赖组件的安全更新机制，定期检测第三方组件的已知安全漏洞，及时更新或修补；

d) SBOM 应保持动态更新，当组件发生变更时自动同步更新清单信息；

e) 宜对关键依赖组件建立可信源验证机制，防止供应链投毒攻击。

14.9 安全事件应急响应

系统应建立安全事件应急响应机制，应满足以下要求：

a) 应制定安全事件应急预案，明确应急组织架构、响应流程和处置措施；

b) 安全事件应进行分级分类管理，根据事件影响范围和严重程度确定响应级别；

c) 发生安全事件后应按规定程序及时报告，报告内容应包括事件描述、影响范围、处置进展和后续措施；

- d) 应定期组织安全事件应急演练，演练频率宜不低于每年一次；
- e) 安全事件处置完成后应形成总结报告，提炼经验教训并更新应急预案。

15 信创适配

15.1 国产 AI 处理器

系统应具备国产化 AI 处理器的适配能力，应满足以下要求：

- a) 系统架构应不绑定特定 GPU 品牌或 CUDA 生态，核心计算应通过抽象层适配不同处理器架构；
- b) 信创适配应分阶段推进：第一阶段应至少完成一种列入国家或行业信创目录且符合相应技术规格的国产 AI 处理器的推理部署适配；第二阶段宜在 12 个月内完成对主流国产 AI 处理器的全面适配；
- c) 模型推理框架应支持或规划支持列入国家或行业信创目录且符合相应技术规格的国产 AI 推理加速库；
- d) 性能指标在国产处理器环境下的验收标准应符合 T/JS AI XXXX 的相关规定。

15.2 国产操作系统

系统应具备国产操作系统的适配能力，应满足以下要求：

- a) 系统应支持或规划支持在列入国家或行业信创目录的主流国产操作系统上运行；
- b) 系统依赖的运行环境（如 Python、数据库等）应有国产操作系统版本；
- c) 系统安装部署脚本应兼容国产操作系统的包管理和路径规范。

15.3 国产数据库

系统应具备国产数据库的适配能力，应满足以下要求：

- a) 数据层应通过抽象层访问数据库，不绑定特定数据库产品；
- b) 应支持或规划支持列入国家或行业信创目录的国产数据库产品；
- c) 数据库操作应使用标准 SQL，避免使用特定数据库的专有语法。

15.4 验证与发布

信创适配应经过规范化的验证流程：

- a) 信创适配验证应在与生产环境一致的国产化环境中进行；
- b) 验证内容应涵盖功能正确性、性能达标性和安全合规性；
- c) 验证结果应形成测试报告，记录测试环境、测试用例和测试结果；
- d) 通过验证的版本方可发布为信创适配版本。

16 部署与验收

16.1 部署要求

系统部署应满足以下要求：

- a) 应提供标准化的部署方案，包括硬件要求、软件环境要求和网络要求；
- b) 应支持轻量化部署，在满足最低硬件配置要求的单机环境下可完成部署；
- c) 宜支持便携式部署，系统可部署于移动存储设备并在不同计算环境间迁移；
- d) 部署过程应尽量简化，宜提供自动化部署脚本或工具；
- e) 部署完成后应进行启动验证，确认各服务正常运行；
- f) 对于便携式部署场景（如单台工作站临时部署用于演示或试点验证），在确保数据安全和访问控制有效的前提下，可适当降低对冗余备份和集群的要求，但应确保单实例的可用性不低于 95%。固定生产环境部署仍应满足本条 a)~e) 的全部要求。

16.2 验收测试

系统验收应进行全面的功能和性能测试，应满足以下要求，应符合附录 A 的规定：

- a) 验收测试应覆盖系统全部功能模块；
- b) 应进行功能正确性测试，验证各功能模块输出结果符合预期；
- c) 应进行性能测试，验证系统响应时间、吞吐量等指标满足要求；
- d) 应进行安全测试，验证访问控制、数据加密、审计日志等安全功能有效；
- e) 应进行兼容性测试，验证系统在目标环境（含信创环境）下的正常运行；
- f) 应进行回归测试，验证功能变更不影响已有功能；
- g) 验收测试结果应形成测试报告，作为系统交付的必要文件。
- h) 应对数据输入规模进行校验，单次请求的数据量超过系统处理能力时应返回明确的错误提示并拒绝处理，防止资源耗尽。

16.3 运维监控

系统应提供运维监控能力，应满足以下要求：

- a) 应支持系统运行状态的实时监控，包括服务健康状态、资源利用率、请求处理量等；
- b) 应支持告警机制，当关键指标超出阈值时自动通知运维人员；
- c) 应支持日志的统一收集和分析；
- d) 应提供运维管理界面，支持服务启停、配置调整、日志查看等常用运维操作；
- e) 应支持系统运行数据的统计分析，为容量规划提供依据。

16.4 部署资源需求分级

系统按资源需求分为三个等级，分级依据为可稳定运行的模型激活参数量、并发能力和上下文支持能力，不绑定特定硬件架构或厂商产品。各等级应满足以下要求：

- a) 轻量级（适用于演示、试点验证或小规模使用）：应能稳定运行单个模型实例（激活参数量不超过 7B），支持不少于 5 个并发推理请求，单模型上下文长度不低于 4096 tokens；
- b) 标准级（适用于部门级日常使用）：应能稳定运行不少于 2 个模型实例（其中主生成模型激活参数量不超过 14B），支持不少于 10 个并发推理请求，上下文长度不低于 8192 tokens；
- c) 全功能级（适用于大规模生产环境）：应能稳定运行不少于 3 个模型实例（含 7B 级、14B 级和 32B 级各一个），支持不少于 20 个并发推理请求，上下文长度不低于 16384 tokens；
- d) MoE 能力等级以单词元激活参数量作为计算规模判定依据，并应同时满足总参数量（含全部专家权重）对应的显存/内存驻留条件；文档应同时标注两者。
- e) 应支持多加速器协同推理，单设备资源不足时支持模型参数自动分配至多设备；
- f) 应优先选用国产 AI 加速卡，满足信创要求。

注：以下硬件配置可作为实现上述各等级的参考，不构成规范性要求：轻量级参考配置为配备不低于 8GB 显存/内存的 GPU 或 NPU 工作站；标准级参考配置为配备不低于 16GB 显存/内存的 GPU 或 NPU 设备；全功能级参考配置为配备 2 张以上高性能加速卡（每张不低于 24GB 显存/内存）或多卡 NPU 设备。具体硬件选型应根据实际模型规格、量化方式和推理引擎的实测结果确定。

附录 A

（规范性）

系统能力验收测试用例表

表 A.1 给出了系统能力验收的主要测试类别和代表性测试用例，每项测试均附量化判定标准。

表 A.1 系统能力验收测试用例表

序号	测试类别	测试内容	对应条款	验收标准
1	架构验证	四层架构完整性	5.1	数据层、模型层、服务层、应用层均能独立运行并通过层间接口通信，4 层完整性通过率 100%
2	数据采集	多源数据接入	6.1	支持结构化、半结构化、非结构化数据至少各一种格式的成功采集，采集成功率 $\geq 95\%$
3	数据核验	数据真实性校验	6.3	异常数据能被标记，核验报告能正确生成，异常识别率 $\geq 90\%$
4	数据血缘	数据溯源	6.6	能从应用端反向追溯至数据来源，溯源路径完整率 $\geq 90\%$
5	模型路由	意图识别与调度	7.2	不同意图类型的请求能被正确路由至对应模型，路由准确率 $\geq 95\%$ ，平均响应时间 $\leq 2s$
6	模型容错	降级与恢复	7.3	主模型不可用时能自动切换至备用模型，切换时间 $\leq 30s$ ，降级期间服务可用率 $\geq 99\%$
7	图谱构建	实体关系抽取	8.1	能从文本中抽取实体和关系并构建图谱，实体识别 F1 $\geq 75\%$ ，关系分类准确率 $\geq 70\%$ ；测试集规模不低于 500 条标注样本
8	图谱推理	关系推理	8.3	能查询两实体之间的关联路径并标注置信度，推理路径正确率 $\geq 80\%$
9	检索服务	混合检索	9.1	向量检索和关键词检索融合后结果排序合理，Precision $\geq 85\%$ ，Recall $\geq 80\%$ ；测试集不低于 200 条查询文档对
10	画像生成	个人画像	10.1	画像内容完整、可溯源，AI 生成内容有标识，字段完整率 $\geq 98\%$
11	评价模型	规则配置	10.3	评价维度和权重可由管理员配置并版本化管理，配置变更生效时间 $\leq 5s$
12	异常检测	偏离告警	10.5	偏离较大的评价结果能被标记并触发告警，告警触发准确率 $\geq 90\%$
13	AI 辅助边界	边界合规	10.7	AI 生成内容有标识，不生成直接决策建议，合规检查通过率 100%
14	文书服务	文书生成	11.2	基于模板和数据能正确生成文书，AI 内容有标识，生成时间 $\leq 10s$
15	公文合规	格式校验	11.4	生成的公文格式符合 GB/T 9704 要求，格式校验通过率 100%

序号	测试类别	测试内容	对应条款	验收标准
16	安全审计	审计日志	14.3	关键操作均有审计记录且不可篡改，审计覆盖率 100%
17	访问控制	权限管理	14.1	不同角色的用户只能访问授权范围内的数据和功能，权限配置准确率 100%
18	数据安全	加密与脱敏	14.5	敏感数据加密存储，展示时自动脱敏，加密覆盖率 100%，脱敏覆盖率 100%
19	信创适配	国产化环境运行	15.1、15.3	在目标信创环境下功能正常、性能达标，功能通过率 100%，性能偏差≤10%
20	部署验收	轻量化部署	16.1	在最低配置环境下能完成部署并正常运行，部署时间≤30min，核心功能验证通过率 100%
21	数据分类	数据分类标签	6.2	AI 自动识别数据类型并注册，分类标签准确率≥80%；高密级数据不降级处理，密级判定一致性 100%
22	数据脱敏	数据脱敏有效性	6.4	敏感数据脱敏后不可还原，脱敏覆盖率 100%，脱敏操作日志完整率 100%
23	知识图谱质量	图谱质量综合评估	8.6	实体识别 F1≥75%，关系分类准确率≥70%，推理路径正确率≥80%，三项指标均达标
24	群体画像	群体统计分析	10.2	支持按组织、层级等维度群体统计，图表展示正确，统计偏差率≤5%
25	评价规则	评价规则版本管理	10.4	支持版本切换和回滚，回滚成功率 100%，版本切换后计算结果一致
26	认证鉴权	认证安全	14.2	支持用户名密码认证和登录失败锁定，锁定阈值可配置，认证令牌过期失效准确率 100%
27	并发性能	并发用户支持	13.1	按部署等级验证：轻量级≥5 并发、标准级≥10 并发、全功能级≥20 并发；首词元延迟和端到端响应时间 P95 满足 §13.4 要求
28	系统可用性	持续运行稳定性	13.5	连续运行 72h 无服务中断，系统可用性≥99.5%
29	AI 幻觉检测	生成内容事实一致性	14.4	系统能标注无数据支撑的生成内容，幻觉检出标注率≥90%，高风险场景强制核验通过率 100%
30	内容安全	内容安全过滤	14.4	敏感词过滤和注入攻击检测有效，敏感内容拦截率≥95%，敏感词库更新响应时间≤24h
31	文书质量	文书生成字段准确率	11.2	自动生成文书字段填充准确率≥99%，模板变量映射正确率 100%
32	系统可用性	72h 连续运行	13.5	系统连续运行 72h，服务可用率≥99.5%，内存泄漏≤50MB/h
33	工作流编排	工作流编排功能	12.1	支持节点定义、流程创建执行暂停恢复终止、版本管理；流程执行状态可追踪；测试样本不少于 3 个不同复杂度的工作流
34	同步异步模式	同步异步执行模式	12.2	同步模式在单次请求周期内返回结果；异步模式支持进度查询；混合编排正常工作；测试样本不少于 5 个流程实例
35	执行监控审计	流程监控和审计功能	12.3	实时监控视图正常展示；执行记录持久化存储不少于 3 年；日志可检索导出；敏感操作标记审计正常

附录 B

（资料性）

参考实现验证说明

本附录给出的实现结构、模块数量与字段规模仅来自单一参考实现，用于说明正文条款的可实现性，不代表推荐规模、不作为符合性判定的依据。符合性判定以正文条款与附录 A 为准。

B.1 原型系统概况

原型系统采用前后端分离架构，后端基于 Flask 框架，前端基于 Vue 框架。系统实现了以下核心能力：

- a) 数据层：支持多源异构数据的统一接入，包括关系型数据库、电子表格、PDF 文档、图片等多种格式；
- b) 模型层：管理多个模型角色，涵盖意图识别、文本生成、嵌入计算、代码生成等任务类型，采用智能路由机制实现多规格模型协同；
- c) 服务层：封装了知识图谱、检索增强生成、画像生成、评价分析、文书生成、工作流编排等多个服务模块；
- d) 应用层：提供多个前端视图和功能组件，覆盖画像展示、智能问答、文书管理、工作流管理等功能。

B.2 关键技术实践

- a) 知识图谱：采用关系型数据库存储，支持多种实体类型、结构化关系和推理关系，通过 VLM 实现多字段的结构化信息提取；
- b) 检索增强生成：采用本地化向量索引系统，实现向量检索与 BM25 关键词检索的 RRF 融合，辅以重排序模型精排、查询改写和上下文压缩；
- c) 安全机制：实现 Session 与 JWT 双轨认证、RBAC 三模式访问控制、五类审计日志、数据脱敏和 AI 生成标识；
- d) 工作流引擎：支持多个预设业务工作流，采用人工确认与自动执行混合节点模式；
- e) 便携部署：支持移动硬盘零配置自适应部署。

B.3 实践对标准编制的支撑

原型系统的开发实践为标准编制提供了以下支撑：

- a) 验证了四层架构的可行性和合理性；
- b) 积累了数据治理、模型管理、安全合规等方面的工程经验；
- c) 识别了 AI 辅助边界、评价规则版本管理、异常偏差检测等关键需求；

d) 为信创适配和轻量化部署提供了技术验证基础。

附录 C

(资料性)

画像维度与标签配置参考

本附录提供了画像维度与标签配置的参考示例，供组织管理员配置评价模型时参考。本附录内容不构成规范性要求，组织可根据自身干部管理需要自行定义和调整。

C.1 画像维度分类参考

表 C.1 给出了画像维度的分类参考。

表 C.1 画像维度分类参考

维度大类	维度子类（示例）	说明
基本信息	自然状况、学历学位、政治面貌	干部基础属性信息
履职经历	任职经历、岗位履历、基层经历	工作经历和岗位变动
考核评价	年度考核、专项考核、民主测评	各类考核结果
能力素质	专业能力、管理能力、学习能力	能力评估与认证
培训情况	培训记录、培训成果、在线学习	培训参与和成效
奖惩记录	表彰奖励、纪律处分	奖惩信息
社会关系	家庭成员、社会关系	按规定采集的关系信息
发展潜力	成长轨迹、专业方向	仅限组织内部研判参考；系统不生成个人发展预测结论，不输出个人推荐顺序。
监督信息	信访举报、组织谈话、重大事项报告、经济责任审计	干部监督相关信息

注：上述维度与干部管理“德能勤绩廉”核心框架的映射关系如下：

德：主要映射至奖惩记录（政治品质、道德品行相关）、社会关系（按规定采集）、监督信息（廉洁自律相关）等维度，考察干部的政治品质和道德品行；

能：主要映射至能力素质、履职经历、培训情况等维度，考察干部的专业能力和履职本领；

勤：主要映射至考核评价（工作态度相关）、培训情况（学习主动性）等维度，考察干部的工作态度和担当精神；

绩：主要映射至考核评价（工作实绩相关）、履职经历（岗位成效）等维度，考察干部的工作实绩和贡献；

廉：主要映射至奖惩记录（廉洁自律相关）、监督信息等维度，考察干部的廉洁自律情况。

组织在配置评价模型时，建议参考上述映射关系，将维度配置与干部管理政策对齐。

C.2 标签体系设计参考

标签体系宜采用层次结构设计：

- a) 一级标签对应画像维度大类；
- b) 二级标签对应画像维度子类；
- c) 三级标签为具体的数据项或评价指标；
- d) 每个标签宜包含标签名称、标签编码、数据类型、取值范围和描述信息。

C.3 评价规则配置参考

评价规则配置宜包含以下要素：

- a) 评价目的说明；
- b) 评价维度和权重分配；
- c) 各维度的评分标准和计算逻辑；
- d) 综合评价方法（如加权求和、等级映射等）；
- e) 规则生效时间和适用范围；
- f) 规则变更历史和变更原因说明。

附录 D

（资料性）

政务智能化发展趋势展望

D.1 跨组织生态互联

注：本章内容为资料性，不构成规范性要求，旨在为系统未来演进提供技术方向参考。

随着政务数字化协作技术的发展，系统未来可考虑以下跨组织演进方向：

- a) 可探索跨组织的标准化数据和服务接口，实现不同组织间系统的安全互联；
- b) 可探索智能体对智能体协议（A2A）在跨组织协作场景中的应用；
- c) 参考 GB/Z 185—2026 智能体互联体系架构相关标准，构建标准化的跨组织通信框架；
- d) 宜探索组织间可信身份互认机制，支持跨系统的身份认证和权限管理；
- e) 宜探索跨组织的服务能力注册与发现机制，实现系统间服务能力的自动发布和匹配；
- f) 宜探索跨组织的工具调用安全机制，确保跨系统交互的输入输出安全可控；
- g) 可探索构建干部管理领域的跨组织协作生态，在符合工作秘密规定并经主管部门批准的前提下实现跨组织数据安全共享。

D.2 多模态融合

系统未来可在以下方面增强多模态处理能力：

- a) 扩展视觉语言模型（VLM）在干部档案处理中的应用，支持手写材料、照片、证书等多模态信息的智能提取；
- b) 宜加强文档智能处理能力，重点突破扫描件 OCR、手写材料识别、印章识别与比对等政务档案场景；
- c) 支持语音交互，实现语音查询和语音播报功能；
- d) 探索视频分析在干部培训考核等场景中的应用；
- e) 建立多模态数据的统一表示和融合分析框架；
- f) 多模态处理能力宜建立量化基准，如手写材料识别准确率、印章匹配准确率等，为技术演进提供可度量的评估标准。

D.3 可解释性

系统未来可加强 AI 决策的可解释性建设：

- a) 引入可解释人工智能（XAI）技术，为模型推理过程提供可理解的解释；
- b) 评价结果宜能展示关键影响因素及其贡献度，可探索 SHAP 值分析、LIME 局部解释、注意力权重可视化等技术路径；
- c) 知识图谱推理链路宜可可视化展示；
- d) 可提供模型行为分析工具，帮助管理人员理解 AI 的决策逻辑和潜在偏差；
- e) 宜建立可解释性的评估框架，定期对模型可解释性水平进行评估和改进。

D.4 联邦学习

在保护数据隐私的前提下，系统未来可探索联邦学习的应用：

- a) 支持跨组织的联邦学习框架，在不共享原始数据的前提下协同训练评价模型；
- b) 联邦学习过程宜满足数据不出域的要求；
- c) 联邦学习模型的聚合和更新过程宜可审计；
- d) 宜评估联邦学习在干部评价模型优化中的适用性和有效性；
- e) 宜探索差分隐私技术的应用，在模型训练过程中引入可控噪声，进一步防止通过模型反推原始数据；
- f) 宜采用安全聚合协议，确保联邦学习过程中各参与方的中间结果在聚合前不可被单独获知；
- g) 宜关注模型投毒防御，建立联邦学习场景下的异常检测和防御机制，防止恶意参与方通过注入有毒数据影响全局模型。

D.5 端侧 AI 与模型蒸馏

系统未来可探索端侧 AI 部署和模型蒸馏技术：

- a) 可通过模型蒸馏技术将大规模模型的知识迁移至轻量化模型，降低对硬件资源的要求；
- b) 宜探索在移动终端或边缘设备上部署轻量化模型，支持离线场景下的基础功能使用；
- c) 端侧模型的精度损失宜控制在可接受范围内，关键功能指标宜不低于云端模型对应指标的 90%；
- d) 模型蒸馏流程宜支持自动化，降低模型适配不同终端设备的人工成本。

D.6 长上下文处理

系统未来可提升长上下文处理能力：

- a) 宜支持更长的上下文窗口，以容纳更多干部档案信息和知识内容；
- b) 可探索长上下文注意力优化技术，如稀疏注意力、分层注意力等，在扩大上下文窗口的同时控制计算成本；
- c) 长上下文场景下宜保持信息提取的准确性，避免“中间丢失”现象（即模型对上下文中间位置的信息关注度不足）；
- d) 宜评估长上下文能力在干部全履历分析、长文档综合研判等场景中的应用价值。