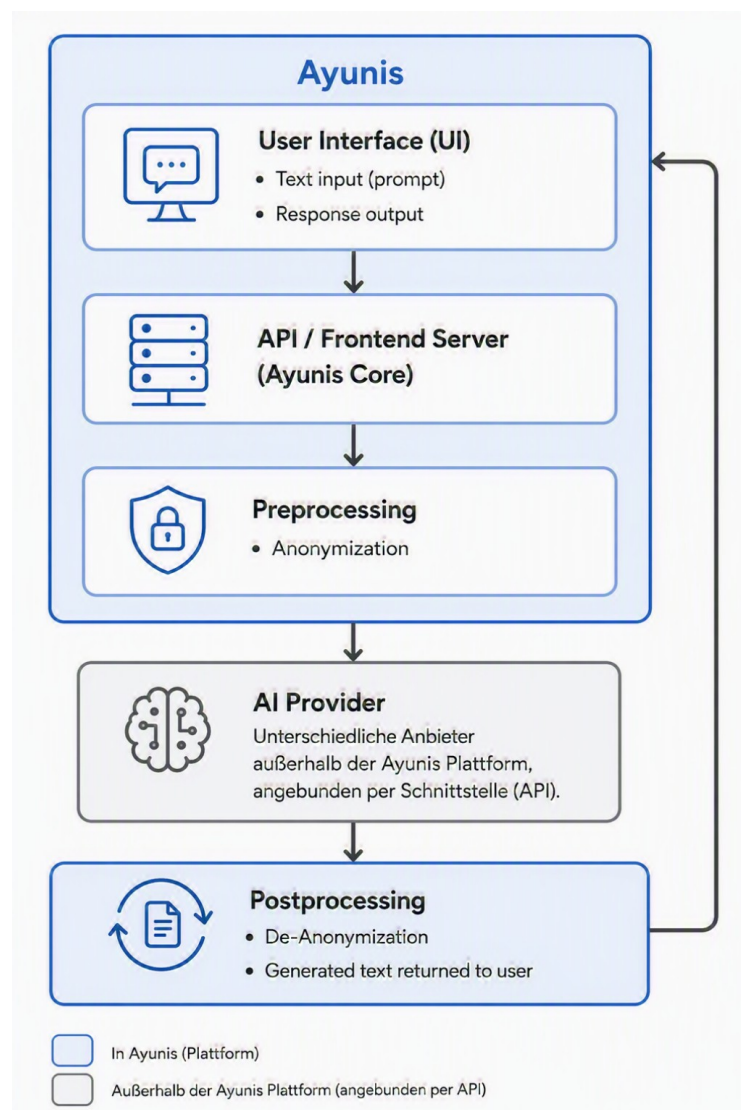


Datenschutz- und KI-Konzept Ayunis

Ayunis ist ein KI-gestützter Dienst zur Vereinfachung und Unterstützung von Geschäftsabläufen in öffentlichen Einrichtungen. Das vorliegende Dokument beschreibt die datenschutzrechtlichen, technischen, organisatorischen und AI-Act-bezogenen Maßnahmen, die für den Betrieb von Ayunis etabliert sind. Ziel ist es, Kunden eine nachvollziehbare Übersicht über die Grundsätze, Schutzmaßnahmen und Governance-Mechanismen bereitzustellen, mit denen ein angemessenes Sicherheits- und Compliance-Niveau nach dem Stand der Technik sichergestellt wird.

Nutzerdaten werden ausschließlich in Datenbanken in Deutschland gespeichert. Daten, die zur Verarbeitung an Sprachmodelle unterschiedlicher Anbieter weitergeleitet werden, werden nicht für Trainingszwecke verwendet und auch nicht in Drittländer übertragen. Alle Daten werden ausschließlich innerhalb der EU verarbeitet. Hierzu bestehen entsprechende vertragliche Regelungen mit den eingesetzten Anbietern. Die nachfolgende Grafik beschreibt generativ den Datenfluss.



Dieses Dokument dient als kundenfähige Beschreibung des Datenschutz- und AI-Act-Konzepts für Ayunis. Es stellt die wesentlichen Regelungen zur Governance, Risikosteuerung, Transparenz, Datensicherheit und regulatorischen Einordnung dar. Die beschriebenen Maßnahmen werden fortlaufend überprüft und im Rahmen der Weiterentwicklung des Systems sowie bei geänderten regulatorischen oder technischen Anforderungen angepasst.

Ayunis Core

Nutzer interagieren über eine Chat-Oberfläche mit Large Language Models (LLMs) unterschiedlicher Anbieter. Dabei ist zu beachten, dass die LLMs bei Anbietern in isolierten Umgebungen betrieben werden, so dass ein Drittlandtransfer der Daten ausgeschlossen werden kann. Die (Kunden) Administratoren eines Kundenkontos können festlegen, welche Anbieter innerhalb der bereitgestellten Konfiguration genutzt werden. Neben Textnachrichten können auch Dokumente, Websites und Internetrecherchen in Konversationen einbezogen werden. Ayunis ermöglicht es Nutzern, Internetrecherchen KI-gestützt durchzuführen und die dabei gewonnenen Informationen im Rahmen der Konversation weiterzuverarbeiten. Der Dienst ermöglicht die Extraktion und KI-gestützte Bearbeitung von Inhalten aus Dokumenten. Darüber hinaus können über definierte Schnittstellen weitere Tools aus Drittsystemen eingebunden und Prompt-Bibliotheken verwaltet werden. Die Nutzung erfolgt innerhalb einer geregelten Zweckbestimmung und unter Berücksichtigung definierter datenschutzrechtlicher, technischer und organisatorischer Leitplanken. Vor Weitergabe der Daten an das LLM erfolgt eine Anonymisierung durch Ayunis selbst (Privacy by Design).

1. Regulatorische Einordnung nach AI Act

1.1 Ausschluss verbotener Praktiken

Ayunis ist nicht für verbotene Praktiken im Sinne des Art. 5 AI Act ausgelegt. Der Dienst wird weder für manipulative, ausnutzende oder sozial bewertende Zwecke noch für unzulässige biometrische Kategorisierung bereitgestellt. Im Rahmen des Produkt- und Freigabeprozesses wird sichergestellt, dass neue Funktionalitäten, Zielgruppen oder Einsatzszenarien auf ihre regulatorische Zulässigkeit hin überprüft werden.

1.2 Regulatorische Einordnung des Dienstes

Ayunis Core ist nach der dokumentierten Zweckbestimmung so ausgestaltet, dass insbesondere Transparenzpflichten, Governance-Anforderungen sowie Anforderungen an Risikosteuerung, Datensicherheit und menschliche Kontrolle berücksichtigt werden. Die regulatorische Einordnung des Dienstes erfolgt entlang der jeweiligen Zweckbestimmung, der vorgesehenen Einsatzszenarien sowie der technischen Ausgestaltung, die vom Kunden / Nutzer der Software festgelegt werden. Änderungen an Funktionen, Integrationen oder Anwendungsfällen werden im Rahmen etablierter Freigabe- und Reviewprozesse erneut bewertet.

Für sensible oder regulierte Einsatzkontexte gelten zusätzliche Prüf- und Freigabeanforderungen, die der Kunde bewerten muss. Hierdurch wird sichergestellt, dass Anwendungsfälle mit erhöhten regulatorischen Anforderungen gesondert bewertet und nur im Rahmen definierter organisatorischer und technischer Leitplanken umgesetzt werden.

1.3 Etablierte AI-Act-Grundsätze

- Sicherstellung eines angemessenen AI-Literacy-Niveaus der beteiligten Beschäftigten und Verantwortlichen,
- transparente Gestaltung der Interaktion mit dem KI-System,
- klare Governance und dokumentierte Verantwortlichkeiten,
- strukturierte Steuerung von Risiken, Änderungen und Einsatzgrenzen,
- angemessene Maßnahmen zu Genauigkeit, Robustheit und Cybersicherheit.

2. AI-Governance und Verantwortlichkeiten

AI-Act-Compliance ist als querschnittliche Aufgabe organisiert, und betrifft Datenschutz, IT-Sicherheit, Produktverantwortung, Entwicklung, Betrieb, rechtliche Bewertung und Support. Für Ayunis besteht ein Governance-Modell, das Verantwortlichkeiten, Freigabeprozesse, Eskalationswege und Überprüfungsmechanismen für Änderungen an System, Modellen, Schnittstellen und Einsatzszenarien festlegt:

- Festlegung verantwortlicher Funktionen für Datenschutz, Informationssicherheit und AI-Act-Compliance,
- dokumentierte Freigabeprozesse für neue Modelle, Tools, Schnittstellen und wesentliche Funktionsänderungen,
- definierte Eskalationswege bei Auffälligkeiten, Sicherheitsvorfällen oder Compliance-Risiken,
- regelmäßige Überprüfung von Zweckbestimmung und zulässigen Einsatzszenarien,
- Nachweisführung über getroffene Maßnahmen, Bewertungen und Freigaben.

Diese Governance-Logik entspricht dem lebenszyklusbezogenen und kontrollorientierten Ansatz einer belastbaren AI-Act-Compliance-Struktur.

3. AI Literacy

Ayunis stellt sicher, dass die mit Betrieb, Nutzung, Betreuung oder Freigabe des Systems befassten Personen über ein angemessenes Maß an KI-Kompetenz verfügen. Dies betrifft insbesondere Administratoren, Produktverantwortliche, Entwickler, Support- und Betriebsteams sowie weitere Personen, die maßgebliche Entscheidungen über Einsatz, Konfiguration oder Freigabe treffen.

Die AI-Literacy-Maßnahmen umfassen insbesondere die Vermittlung von Kenntnissen zur Funktionsweise und zu den Grenzen generativer KI, zu typischen Fehlerbildern wie Halluzinationen, zu Risiken aus Prompt Injection und unsicheren Tool-Anbindungen, zu Anforderungen an Transparenz und menschliche Kontrolle sowie zu datenschutzrechtlichen und regulatorischen Rahmenbedingungen.

Die Durchführung von Schulungs- und Sensibilisierungsmaßnahmen wird dokumentiert und regelmäßig überprüft.

4. Transparenz gegenüber Nutzern

Da Nutzer über eine Chat-Schnittstelle mit dem System interagieren, ist die Interaktion mit einem KI-System transparent ausgestaltet. Nutzer werden in geeigneter Weise darüber informiert, dass Antworten ganz oder teilweise KI-generiert sein können und inhaltliche Fehler nicht ausgeschlossen werden können. Soweit KI-generierte Inhalte gegenüber Dritten weiterverwendet werden, gelten ergänzende Vorgaben zur Kennzeichnung und zur fachlichen Prüfung:

- klarer Hinweis innerhalb der Benutzeroberfläche auf die Nutzung eines KI-Systems,
- Hinweise auf mögliche Fehler, Unvollständigkeit oder Halluzinationen,
- Vorgaben zur menschlichen Prüfung bei relevanten Entscheidungen,
- Regelungen zur Kennzeichnung KI-generierter Inhalte in geeigneten Anwendungsfällen.

5. Human Oversight und zulässige Einsatzgrenzen

Für den Einsatz von Ayunis muss der Kunde organisatorische Leitplanken sicherstellen, damit KI-Ausgaben nicht ungeprüft zu rechtlich, wirtschaftlich oder tatsächlich erheblichen Entscheidungen führen. KI-generierte Inhalte dienen der Unterstützung und unterliegen in relevanten Anwendungsfällen einer angemessenen fachlichen Kontrolle:

- KI-Ausgaben dienen der Unterstützung und nicht der unkontrollierten Endentscheidung,
- fachliche Prüfung ist erforderlich, wenn Inhalte gegenüber Dritten verwendet oder in operative Prozesse übernommen werden,
- besonders sensible oder regulierte Anwendungsfälle unterliegen gesonderten Freigabeanforderungen.

6. Risikomanagement für Ayunis Core

Für Ayunis besteht ein etabliertes Risikomanagement, das technische, organisatorische, datenschutzrechtliche und AI-Act-bezogene Risiken erfasst, bewertet und behandelt. Dabei werden insbesondere Risiken aus der Nutzung generativer KI, aus der Anbindung externer Modelle und Systeme sowie aus der konkreten Zweckbestimmung des Dienstes berücksichtigt. Die Risikosteuerung erfolgt fortlaufend und wird bei Änderungen an Funktionen, Integrationen oder Einsatzszenarien aktualisiert.

Etablierte Maßnahmen zur Risikosteuerung

- Hinweise an Nutzer zur möglichen Fehleranfälligkeit KI-generierter Inhalte,
- Sicherheitsfunktionen im Webdienst zur Begrenzung schädlicher Aktionen,
- kontrollierter Zugriff von Modellen auf erforderliche und definierte Tools,
- strukturierte Ausgabeformate und Validierungsmechanismen zur Reduzierung von Folgefehlern,
- organisatorisch definierte Einsatzgrenzen und Freigabeanforderungen für sensible Anwendungsfälle,
- Reviewprozesse für Änderungen an Modellen, Tools, Prompt-Bibliotheken und Schnittstellen,
- regelmäßige Neubewertung relevanter Risiken und Vorfälle.

7. Präzisierte Bewertung einzelner technischer Risiken

Prompt Injection

Interne Tools sind mit definierten Ein- und Ausgabeformaten ausgestaltet. Risiken aus nutzerseitig konfigurierten Dritt-Systemen werden durch vertragliche, organisatorische und technische Regelungen begrenzt. Ergänzend bestehen Kontrollen, um missbräuchliche Tool-Responses und unsichere Konfigurationen zu reduzieren.

Sensitive Information Disclosure

Das System verarbeitet die vom Nutzer eingebrachten Daten auf Grundlage definierter Nutzungs- und Berechtigungsregeln. Ergänzende Maßnahmen zu Datenminimierung, Rollen- und Berechtigungskonzepten sowie Nutzungsbeschränkungen tragen dazu bei, die Eingabe unzulässiger oder besonders sensibler Inhalte zu begrenzen.

Supply Chain

Für eingesetzte Modellanbieter und relevante Drittdienste bestehen dokumentierte Auswahl-, Freigabe- und Bewertungsprozesse. Dabei werden insbesondere Vertragslage, Datenverarbeitung, Sicherheitsniveau, Verfügbarkeit, Änderungsmanagement und Transparenz der eingesetzten Anbieter berücksichtigt.

Data & Model Poisoning

Risiken aus manipulierten Eingabedaten, kompromittierten Dokumenten, bösartigen Websites, Embeddings und unsicheren Tool-Inputs werden im Rahmen des Risikomanagements berücksichtigt. Soweit keine eigenen Modelle trainiert oder feinjustiert werden, reduziert sich das Risiko klassischer Modellvergiftung zusätzlich.

Improper Output Handling

Strukturierte Ausgaben, Validierungsmechanismen und organisatorische Prüfpflichten tragen dazu bei, dass KI-Ausgaben nicht ungeprüft in operative oder rechtlich relevante Prozesse übernommen werden.

Excessive Agency

Modelle erhalten ausschließlich Zugriff auf definierte und erforderliche Funktionen. Direkte Zugriffe auf besonders kritische Systemkomponenten werden durch technische und organisatorische Begrenzungen ausgeschlossen.

System Prompt Leakage

Systemprompts, Konfigurationslogiken und interne Steuerinformationen werden als schützenswerte Informationen behandelt und unterliegen angemessenen technischen und organisatorischen Schutzmaßnahmen.

Vector and Embedding Weaknesses

Bei der Nutzung externer Embedding- oder Vektordienste werden Auswahl, Sicherheitsniveau und vertragliche Einbindung der Anbieter nachvollziehbar dokumentiert und im Rahmen des Freigabeprozesses bewertet.

Misinformation

Da Halluzinationen systembedingt nicht vollständig ausgeschlossen werden können, bestehen Nutzerhinweise, Einsatzgrenzen sowie fachliche Prüfpflichten für relevante Weiterverwendungen KI-generierter Inhalte.

Unbounded Consumption

Technische Begrenzungen wie Rate Limits und weitere Steuerungsmechanismen reduzieren Missbrauchs- und Kostenrisiken und werden im Rahmen des Betriebs fortlaufend überprüft.

8. Zusätzliche technische und organisatorische Maßnahmen zum Datenschutz und zur Informationssicherheit

1. Anonymisierung der Prompts:

- a) **Automatische Maskierung (Anonymisierung von Inhalten):** Zur Reduzierung von Datenschutzrisiken werden Inhalte vor der Weitergabe an externe KI-Modelle automatisiert anonymisiert. Die Maskierung erfolgt immer dann, wenn eine Konversation im sogenannten Anonymous-Modus geführt wird. Dieser Modus ist aktiv, wenn:
 - der Nutzer die Konversation ausdrücklich anonym anlegt oder
 - ein KI-Modell verwendet wird, das ausschließlich anonymisierte Daten verarbeitet (anonymousOnly).

In diesem Fall werden sowohl:

- die vom Nutzer eingegebenen Inhalte (Prompts) als auch
- Ergebnisse aus angebundenen Tools sowie Inhalte der Internetrecherchen

vor der Verarbeitung durch das KI-Modell automatisiert maskiert, sodass keine direkt identifizierbaren personenbezogenen Daten übermittelt werden.

- b) **Umgang sensibler Eingaben:** Ayunis blockiert grundsätzlich keine Nutzereingaben. Stattdessen wird der Schutz sensibler Daten durch technische und organisatorische Mechanismen sichergestellt:
 - i. **Nutzung anonymisierter Modelle:** KI-Modelle können so konfiguriert werden, dass sie ausschließlich mit anonymisierten Daten arbeiten (anonymousOnly). Bei Verwendung solcher Modelle wird die Maskierung automatisch für alle zugehörigen Konversationen erzwungen. Sensible Informationen werden in diesem Fall ausschließlich in anonymisierter Form verarbeitet.
 - ii. **Organisationsspezifische Ausnahmen (Whitelist):** Administratoren können definieren, dass bestimmte Datenkategorien oder Inhalte nicht anonymisiert werden (z. B. der Name der eigenen Organisation). Diese Ausnahmen werden über konfigurierbare Regeln gesteuert, die auf bestimmten Datenkategorien (PII – personenbezogene Daten) sowie optionalen Mustern (z. B. Regex) basieren.
- c) **Klassifizierung von Prompts:** Ayunis führt ausschließlich einer Klassifizierung im Hinblick auf personenbezogene Daten (PII) durch. Es erfolgt keine inhaltliche Bewertung der Prompts hinsichtlich Themen, Stimmung oder Risiken. Die Klassifizierung dient somit ausschließlich dazu, personenbezogene Daten zu erkennen und – sofern erforderlich – geeignete Schutzmaßnahmen wie Maskierung anzuwenden.

2. **Verschlüsselung der Datenübertragung:** Der Zugriff auf die Anwendung sowie die Kommunikation mit den Frontend erfolgen ausschließlich über HTTPS. Die TLS-Terminierung wird über einen vorgelagerten Reverse-Proxy umgesetzt, der selbst betrieben wird. Dadurch wird die externe Kommunikation zentral abgesichert und kontrolliert. Authentifizierungs-Cookies werden mit den Sicherheitsattributen „HttpOnly“ und Secure sowie einer restriktiven SameSite-Konfiguration gesetzt, um die Risiken aus unbefugtem Zugriff, Cross-Site-Scripting und Cross-Site-Request-Forgery zu reduzieren. Ergänzend ist die CORS-Konfiguration auf zugelassene Ursprünge beschränkt; sicherheitsrelevante HTTP-Header werden serverseitig gesetzt. Verbindungen zu externen KI-Diensten, insbesondere zu den angebundenen LLMs, erfolgen ausschließlich verschlüsselt über deren HTTPS-Endpunkte. Die Kommunikation zwischen internen Systemkomponenten findet innerhalb des nicht öffentlich erreichbaren Container-Netzwerks der Instanz statt.
3. **Schutz ruhender Daten:** Besonders schützenswerte Geheimnisse werden anwendungsseitig verschlüsselt gespeichert. Dies betrifft insbesondere Zugangsdaten externer Integrationen, die mit AES-256-GCM verschlüsselt abgelegt und nur bei technisch erforderlicher Nutzung entschlüsselt werden. Passwörter und API-Schlüssel werden nicht im Klartext gespeichert, sondern ausschließlich als Hash verarbeitet beziehungsweise gespeichert; ergänzende Regelungen hierzu ergeben sich aus den Maßnahmen zur Verwaltung von API-Schlüsseln und Geheimnissen sowie zum Passwortschutz. Backups werden verschlüsselt gespeichert. Für die zusätzliche Absicherung der Datenbestände im laufenden Betrieb ist die Einführung einer dateisystemseitigen Verschlüsselung mittels LUKS für Datenträger und Datenbankbestände vorgesehen.
4. **Verwaltung von API-Schlüsseln und Geheimnissen:** API-Schlüssel der Anwendung werden nicht im Klartext gespeichert, sondern ausschließlich als Hash verarbeitet und abgelegt. Der Klartextwert wird nur einmalig bei der Erstellung angezeigt und in späteren Übersichten maskiert dargestellt, sodass eine nachträgliche Einsichtnahme ausgeschlossen ist. Zugangsdaten externer Integrationen werden verschlüsselt gespeichert und sind nach der Erstellung nicht mehr einsehbar, sondern nur noch änderbar. Betriebsgeheimnisse wie Signaturschlüssel, Datenbankzugangsdaten und Mailzugangsdaten werden über geschützte Konfigurationsmechanismen verwaltet. Der Zugriff auf diese Konfigurationen ist auf berechtigte Administratoren beschränkt und dient der kontrollierten Verwaltung sensibler technischer Geheimnisse.
5. **Rollen- und Berechtigungskonzept:** Der Zugriff auf Funktionen und Daten erfolgt auf Grundlage definierter Benutzerrollen und Berechtigungen. Dabei bestehen zwei Berechtigungsebenen:
 - a) organisationsbezogene Rollen, insbesondere Administrator und Benutzer, sowie
 - b) eine plattformbezogene Administrationsrolle für Super-Administratoren.
 Die Durchsetzung der Berechtigungen erfolgt serverseitig über eine zentrale Berechtigungsprüfung, sodass Zugriffsentscheidungen nicht allein clientseitig erfolgen. Administrative Endpunkte sind durchgängig rollengeschützt. Daten werden mandantenbezogen beziehungsweise organisationsbezogen getrennt verarbeitet, um eine unzulässige organisationsübergreifende Einsichtnahme oder Nutzung von Daten zu verhindern.
6. **Schutz privilegierter Zugänge:** Administrative und privilegierte Endpunkte der Anwendung sind durchgängig mit serverseitigen Rollenprüfungen abgesichert. Die Authentifizierung erfolgt über signierte Tokens mit definierter Gültigkeitsdauer, der Zugang ist zudem an eine bestätigte E-Mail-Adresse gebunden. Optional kann der Zugriff zusätzlich über eine organisationsbezogene IP-Allowlist eingeschränkt werden. Im Produktivbetrieb ist eine Begrenzung der Anmeldeversuche aktiv, um automatisierte oder missbräuchliche Login-Versuche zu erschweren. Eigene privilegierte Zugänge zur Betriebsumgebung, insbesondere Server-, Datenbank- und Infrastrukturzugriffe, sind durch Mehr-Faktor-Authentifizierung geschützt und auf berechtigte Personen beschränkt.
7. **Datenminimierung und Zweckbindung:** Die Verarbeitung personenbezogener Daten erfolgt zweckgebunden und auf das für Bereitstellung, Betrieb und Nutzung der Anwendung erforderliche Maß begrenzt. Über die erforderlichen personenbezogenen Stammdaten, insbesondere Name, E-

Mail-Adresse und Organisationszugehörigkeit, hinaus stellt Ayunis einen PII-Anonymisierungsmechanismus bereit. In aktivierten Anonymitätsmodi werden personenbezogene Inhalte vor der Übermittlung an externe KI-Modelle automatisch durch Platzhalter ersetzt, sodass diese den Modellen nicht im Klartext zugehen. Soweit Nutzungsereignisse technisch erfasst werden, beschränkt sich dies auf inhaltsfreie Metadaten zur Nutzung der Anwendung, etwa die Anzahl bestimmter Nutzeraktionen wie das Absenden eines Chats sowie den zugehörigen Zeitstempel. Inhalte der Kommunikation werden hierbei nicht protokolliert.

8. **Passwortschutz:** Passwörter werden nach anerkanntem Verfahren mit verschlüsselt und niemals im Klartext gespeichert. Eine Ausgabe von Passwörtern in API-Antworten ist ausgeschlossen. Zusätzlich gilt eine Passwort-Komplexitätsrichtlinie, die insbesondere eine Mindestlänge sowie die Verwendung von Groß- und Kleinbuchstaben und Ziffern vorsieht. Dadurch wird das Risiko einer Kompromittierung durch unzureichend geschützte oder schwache Passwörter reduziert.
9. **Lösch- und Aufbewahrungskonzept:** Für die Verarbeitung bestehen definierte Aufbewahrungs- und Löschfristen. Konversationen und Chats werden standardmäßig nicht automatisiert gelöscht, die Aufbewahrung kann jedoch organisationsbezogen auf Anfrage hin konfiguriert werden. Hierfür stehen Aufbewahrungszeiträume von 30, 90, 180, 365 oder 730 Tagen zur Verfügung. Dokumente und Wissensdatenbanken sind an die jeweilige Nutzung beziehungsweise Ressource gekoppelt und werden mit der zugehörigen Ressource gelöscht.

Benutzerkonten werden nach Account- oder Vertragsende innerhalb von 30 Tagen gelöscht. Verschlüsselte Backups werden für 30 Tage rollierend aufbewahrt. Anwendungs-Logs werden für 30 Tage gespeichert. Fehler- und Monitoring-Daten werden für 90 Tage vorgehalten. Für einzelne Ressourcen wie Konversationen, Wissensdatenbanken und Dokumente bestehen explizite Löschfunktionen, die auch zugehörige Dateien im Objektspeicher entfernen.

Löschanträge betroffener Personen nach Art. 17 DSGVO können per E-Mail an den Support gestellt werden und werden innerhalb von 30 Tagen umgesetzt. Dabei ist zu berücksichtigen, dass Daten erst nach Ablauf der Backup-Rotation von 30 Tagen vollständig aus den Sicherungen entfernt sind.

10. **Protokollierung und Nachvollziehbarkeit:** Betriebs- und sicherheitsrelevante Ereignisse werden strukturiert protokolliert. Die Protokollierung im Produktivbetrieb beinhaltet Informationen im Benutzer- und Organisationskontext zur Nachvollziehbarkeit von Ereignissen. Relevante Fehler- und Monitoring-Informationen werden an das eingesetzte Fehler- und Monitoring-System weitergegeben. Das Monitoring wird in der eigenen Ayunis Infrastruktur in Deutschland betrieben; eine Übermittlung in ein Drittland findet damit nicht statt. Inhalte von Nachrichten sowie Klartext-Geheimnisse werden bei der Protokollierung nicht erfasst. Vor der Übermittlung an das Monitoringsystem werden sensible Datenfelder entfernt. Anwendungs-Logs werden für 30 Tage aufbewahrt, Monitoring-Daten werden für 90 Tage gespeichert.

9. Zusammenfassende Einordnung

Ayunis wird auf Grundlage eines etablierten Datenschutz-, Informationssicherheits- und AI-Act-Konzepts betrieben. Dieses umfasst insbesondere Governance-Regelungen, Transparenz-mechanismen, Human-Oversight-Leitplanken, Risikomanagement, technische und organisatorische Schutzmaßnahmen sowie definierte Prozesse für die weitere Entwicklung und fortlaufende Überprüfung des Systems. Hierdurch wird sichergestellt, dass regulatorische, technische und organisatorische Anforderungen im Rahmen der Zweckbestimmung des Dienstes angemessen berücksichtigt werden.

Das Dokument wird fortlaufend überprüft und insbesondere bei neuen Modellen oder Modellanbietern, neuen Tool-Anbindungen oder Agentenfunktionen, Erweiterungen der Zweckbestimmung, Einsätzen in sensiblen oder regulierten Fachkontexten sowie bei Änderungen regulatorischer Leitlinien oder Standards aktualisiert.

