

AI Control Standard

Prüfbarer AI-Kontrollstandard für Finanzinstitute

Dokumenttyp: Prüfbarer AI-Kontrollstandard für Finanzinstitute

Version: V2.0 · Status: Published Version · Datum: 12.07.2026

Owner: Martin Bischoff · 4T2 Solutions GmbH

Public · Frei verfügbares Rahmenwerk



Dokumentkontrolle

Feld	Wert
Dokumentname	AI Control Standard
Version	V2.0
Status	Published Version
Vertraulichkeit	Public · Frei verfügbares Rahmenwerk
Owner	Martin Bischoff
Review-Zweck	Publizierte Version

Inhaltsverzeichnis

1. Zweck, Zielgruppe und Verwendung.....	4
Zweck.....	4
Zielgruppe.....	4
Beantwortete Managementfrage.....	4
Beziehung zur Dokumentenfamilie.....	4
Erwartete Ergebnisse und Abgrenzung.....	4
Methodische Grundlage.....	4
2. Kontrolllogik, Profilklassen und Trigger.....	4
3. AI Control Applicability Statement.....	5
4. Rollen im Kontrollmodell.....	5
5. Safeguard-Feldmodell.....	6
6. Governance und Accountability Safeguards.....	6
7. Classification und Risk Profiling Safeguards.....	7
8. Data, Privacy und Confidentiality Safeguards.....	9
9. Betriebsfreigabe und Nutzungsfreigabe.....	10
10. Model Risk und Evaluation Safeguards.....	12
11. Supplier und Third-Party Safeguards.....	13
12. Secure AI Safeguards.....	14
13. Incident, Monitoring und Reporting Safeguards.....	14
14. AI Literacy, Human Oversight und Communication Safeguards.....	15
15. Assurance, Testing und Evidence Safeguards.....	17
16. Deviation, Compensating Controls und Risk Acceptance.....	17
17. Mapping zu Regulierung und Frameworks.....	17
18. Evidence und Nachweisführung.....	18
19. Glossar.....	18
20. Quellen.....	19
21. Haftungs- und Anwendungshinweis.....	21
Copyright und offene Lizenz.....	21
Transparenz zur Erstellung.....	21

1. Zweck, Zielgruppe und Verwendung

Zweck

Der Control Standard für Artificial Intelligence (AI) übersetzt die Entscheidungen des AI Governance Framework for Financial Institutions in prüfbare Safeguards. Er legt fest, welche Evidence vorliegen muss, wie Second Line und Revision diese prüfen und wann eine Anforderung nicht erfüllt ist.

Zielgruppe

Er richtet sich an First Line, Second Line, Interne Revision und externe Prüfer.

Beantwortete Managementfrage

Woran lässt sich nachvollziehbar erkennen, dass eine Freigabeentscheidung angemessen umgesetzt und dokumentiert wurde?

Beziehung zur Dokumentenfamilie

Für operative Workflows ist das AI Operations Handbook massgeblich. Technische Mindestanforderungen und technische Evidence definiert die AI Security Baseline. Das Framework bleibt die verbindliche Quelle für Rollen, Lifecycle, Risikoappetit sowie Betriebsfreigabe und Nutzungsfreigabe. Dieser Standard ersetzt weder diese Dokumente noch eine rechtliche Einzelfallbeurteilung.

Erwartete Ergebnisse und Abgrenzung

Der Standard gilt für AI Use Cases, AI-Systeme,¹ AI-Plattformen und AI-Funktionen im Geltungsbereich des Frameworks. Profilklassen und Trigger bestimmen die Kontrolltiefe. Das Ergebnis ist ein belastbares Applicability Statement mit zugeordneten Safeguards, Evidence, Testergebnissen und dokumentierten Abweichungen.

Methodische Grundlage

Erklärende Passagen folgen Digital.gov und der 4T2 Professional Prose. Die strukturierten Control Statements verwenden die Special Publication 800-53 des U.S. National Institute of Standards and Technology (NIST SP 800-53) als offenes Muster für atomare Kontrollergebnisse. Evidenz- und Prüffelder orientieren sich methodisch an NIST SP 800-53A; stabile IDs und getrennte Felder an der Open Security Controls Assessment Language (OSCAL). Diese Nutzung ist methodisch und begründet weder die Übernahme des NIST-Katalogs noch eine Zertifizierung.

2. Kontrolllogik, Profilklassen und Trigger

Profilklassen steuern die Kontrolltiefe. Trigger aktivieren zusätzliche Safeguards. Das Applicability

¹ AI-System im Sinne von Art. 3 Nr. 1 AI Act. Die interne Scope-Prüfung darf breiter ansetzen, bis die rechtliche Einordnung abgeschlossen ist.

Statement dokumentiert, welche Safeguards gelten, warum einzelne Safeguards nicht gelten und welche Abweichungen akzeptiert wurden.

In der Anwendung beginnt ein Team nicht beim ersten Safeguard und arbeitet die Liste bis zum Ende ab. Es bestimmt zuerst Profilkategorie und Trigger, leitet daraus die relevanten Safeguard-Familien ab und dokumentiert die Auswahl im Applicability Statement. So bleibt der Standard streng genug für Prüfung und Assurance, ohne jeden AI Use Case mit derselben Kontrolltiefe zu behandeln.

Profilklasse	Bedeutung	Kontrolllogik
B - Basic	geringe interne Wirkung	Basisnachweise und Owner reichen meist aus.
E - Enhanced	vertrauliche Daten, relevante Prozesse oder Supplier	Zusätzliche Datenschutz-, Supplier-, Security- und Monitoring-Safeguards.
C - Critical	kritische Funktion, hohe fachliche Wirkung oder starke Abhängigkeit	Verstärkte kritische Prüfung, Managementsicht und technische Evidence.
H - High Risk	Hochrisiko-AI oder vergleichbare Wirkung	Vollständiger erhöhter Kontrollsatz mit Managemententscheidung.
X - Excluded / prohibited	verbotene Praktik oder nicht vertretbare Nutzung	Keine Freigabe. Stop und Eskalation.

Trigger sind keine Freitextnotizen. Sie müssen im Applicability Statement nachvollziehbar zu Safeguards führen.

3. AI Control Applicability Statement

Das AI Control Applicability Statement² ist der zentrale Kontrollentscheid. Es zeigt:

- welche Profilkategorie gilt,
- welche Trigger aktiv sind,
- welche Safeguards gelten,
- welche Safeguards nicht anwendbar sind,
- welche Abweichungen kompensiert werden,
- welche Restrisiken befristet akzeptiert wurden.

Nichtanwendbarkeit braucht eine Begründung. Risk Acceptance braucht Owner, Laufzeit, Genehmigung, Kompensationskontrolle und Review-Datum.

4. Rollen im Kontrollmodell

Wirksame Kontrollen brauchen eine eindeutige Zuordnung zwischen fachlicher Verantwortung, technischer Nachweisführung und kritischer Prüfung durch die Kontrollfunktionen. Die Rollen im Kontrollmodell zeigen deshalb, wer Evidence verantwortet und wer ihre Belastbarkeit prüft.

Rolle	Kontrollfunktion
AI Owner	verantwortet Use-Case-Evidence und fachliche Kontrollumsetzung

² Das AI Control Applicability Statement ist das interne Steuerungsinstrument für Profilkategorie, Trigger, anwendbare Safeguards, begründete Nichtanwendbarkeit und Abweichungen. Es ist nicht mit einer gesetzlich definierten Konformitätserklärung gleichzusetzen.

Rolle	Kontrollfunktion
System Owner / IT Operations	liefert Betriebs- und technische Systemnachweise
Data Owner	bestätigt Datenquelle, Klassifikation und zulässige Nutzung
CISO	prüft technische Risiken und Security Evidence kritisch
Compliance / Legal	prüft regulatorische Einordnung und Rechtsrisiken kritisch
Datenschutz	prüft personenbezogene Daten, DSFA und Betroffenenrechte kritisch
Risk Management	prüft Risikoappetit, Risikoprofil und Modellrisiko kritisch
Interne Revision	prüft unabhängig Design und Wirksamkeit

Workflow-Handoffs gehören in das AI Operations Handbook. Technische Evidence Owner werden in der AI Security Baseline konkretisiert.

5. Safeguard-Feldmodell

Jeder Safeguard verwendet dasselbe Feldmodell.

Das Feldmodell trennt Erwartung und Prüfung. Die Mindestanforderung beschreibt, was erreicht sein muss. Evidence zeigt, woran dies erkennbar ist. Test Type und Fail Criteria machen daraus eine nachvollziehbare Kontrollaussage statt einer allgemeinen Empfehlung.

Feld	Bedeutung
Zweck	Warum der Safeguard existiert
Applicability	Profilklasse oder Trigger
Mindestanforderung	prüfbare Erwartung
Evidence	Nachweis, der vorliegen muss
Test Type	Review-, Stichproben-, Walkthrough- oder technische Prüfung
Fail Criteria	klare Bedingung, wann der Safeguard nicht erfüllt ist

6. Governance und Accountability Safeguards

Governance Safeguards verankern jeden AI Use Case in einer nachvollziehbaren Verantwortungs- und Entscheidungsstruktur. Fehlt diese Grundlage, können nachgelagerte fachliche oder technische Kontrollen ihre Wirkung nicht verlässlich entfalten.

Geprüft wird, ob Verantwortung, kritische Prüfung und Eskalationsrechte vor einer Freigabe eindeutig feststehen. Erst damit wird die Entscheidung für Second Line und Revision belastbar.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI GOV-01	Das Institut benennt für jeden AI Use Case einen verantwortlichen AI Owner.	B, E, C, H	Owner Record	Evidence Review	kein verantwortlicher Owner
AI GOV-02	Das Institut benennt für jedes AI-System einen verantwortlichen System Owner.	alle Systeme	System Record	Evidence Review	keine Betriebsverantwortung

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI GOV-03	Der AI Owner wendet den genehmigten AI Risk Appetite auf den AI Use Case an.	E, C, H	Risk Appetite Check	Challenge Review	Nutzung ausserhalb Appetite ohne Entscheidung
AI GOV-04	Das Institut trennt First Line, Second Line und unabhängige Prüfung nachvollziehbar.	E, C, H	Rollen- und Prüfungsnachweis	Design Review	First und Second Line nicht getrennt
AI GOV-05	Das Institut weist einer befugten Rolle ein wirksames Stop- und Eskalationsrecht zu.	C, H	Eskalationsnachweis	Walkthrough	keine Stop- oder Eskalationsrolle
AI GOV-06	Das Institut blockiert, isoliert oder beendet eine AI-Nutzung mit Ausschlussstatus X technisch wirksam.	X-Entscheidung	X Decision Record, X Closure Review	Technical Test	Nutzung trotz Ausschlussstatus möglich
AI PAC-01	Das Institut prüft vor Entwicklung, Beschaffung, eigenem Marktauftritt und wesentlichen Änderungen, ob es Provider wird. Ein positiver Entscheid blockiert eine reine Deployer-Freigabe.	Eigen-/Auftragsentwicklung, eigener Name oder Marke, Rebranding, Zweck- oder wesentliche Änderung	Provider Activation Record	Legal / Compliance Review	Providerstatus ungeklärt oder Deployer-Freigabe trotz positivem Gate
AI PAC-02	Das Institut genehmigt bei bestätigter Hochrisiko-Providerrolle vor Release einen Provider Compliance Plan mit den erforderlichen Deliverables aus Art. 9–21, 43, 47–49 und 72–73 AI Act sowie den einschlägigen Anhängen.	bestätigte Hochrisiko-Providerrolle	Provider Compliance Plan, Deliverable Checklist, Release Decision	Evidence Review	Release ohne genehmigten Plan oder erforderliches Deliverable
AI PAC-03	Das Institut stoppt bei einem Importeur-, Distributor-, Annex-I-Produkthersteller-, eigenen GPAI-Modellprovider- oder formellen Real-World-Testing-Signal den Release bis zum dokumentierten Spezialentscheid.	atypisches Rollensignal	Specialist Handoff Closure	Legal / Compliance Review	Marktbereitstellung, Inbetriebnahme oder Test ohne Spezialfreigabe

7. Classification und Risk Profiling Safeguards

Klassifikation und Risikoprofil bestimmen, welche Kontrolltiefe ein AI Use Case benötigt. Die Safeguards sichern ab, dass regulatorische Rolle, Risikotreiber und Revalidierungsbedarf begründet statt lediglich formal erfasst werden.

Die Prüfung zeigt, ob die gewählte Profilkategorie die tatsächliche Wirkung abbildet und alle relevanten

Trigger berücksichtigt. Davon hängen Umfang, Tiefe und Zeitpunkt der weiteren Kontrollen ab.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI CLS- 01	Das Institut prüft jeden Kandidaten vor dem Profilentcheid auf AI-Relevanz.	alle Kandidaten vor Profilentcheid	Intake und Scope Record	Evidence Review	AI-Relevanz nicht geprüft
AI CLS- 02	Das Institut bestimmt je Rechtseinheit und Einsatzland die Operatorrolle und bewertet mögliche Rollenwechsel.	relevante Rechtseinheiten und Einsatzländer	Operator Role Assessment	Legal Review	Provider-, Deployer-, Importeur-, Distributor- oder Produktherstellerrolle unklar; Art.-25-Rollenwechsel nicht bewertet
AI CLS- 03	Das Institut prüft Kandidaten und Revalidierungen auf verbotene Praktiken und stoppt ungeklärte Verdachtsfälle.	alle Kandidaten und Revalidierungen	Prohibited-Practice Check	Challenge Review	Verdacht ohne Stop-Entscheid
AI CLS- 04	Das Institut begründet und dokumentiert die Profilkategorie B, E, C oder H oder den Ausschlussstatus X.	alle Profil- und X-Entscheide	Profilklassen- oder X-Entscheid	Applicability Review	Entscheidung nicht begründet
AI CLS- 05	Das Institut prüft und dokumentiert den Trigger für eine Grundrechte-Folgenabschätzung (Fundamental Rights Impact Assessment, FRIA).	H und relevante Annex-III-Trigger	FRIA Trigger Record ³	Legal / Compliance Review	FRIA fehlt trotz Trigger
AI CLS- 06	Das Institut bewertet General-Purpose AI (GPAI) und die daraus entstehenden Downstream-Risiken.	GPAI, Modellintegration	Provider Information Record	Evidence Review	Anbieterinformationen fehlen
AI CLS- 07	Das Institut erfasst und überwacht die anwendbaren Revalidierungstrigger.	B, E, C, H	Trigger Matrix	Evidence Review	Trigger ohne Review-Entscheid
AI CLS- 08	Das Institut beurteilt und dokumentiert den gesetzlichen Registrierungstrigger.	Annex-III-System, Providerstatus, Art.-6-Abs.-3-Ausnahme oder öffentlicher Deployerbezug	Registration Decision Record	Legal / Compliance Review	Art.-49-Trigger ungeklärt
AI CLS- 09	Das Institut schliesst eine erforderliche EU-Datenbankregistrierung vor dem massgeblichen Einsatzzeitpunkt ab.	bestätigter Art.-49-Trigger	EU Database Registration Evidence	Evidence Review	Inverkehrbringen, Inbetriebnahme oder Nutzung trotz fehlender Registrierung

3 Die Fundamental Rights Impact Assessment nach Art. 27 AI Act ist nicht bei jeder Hochrisiko-AI automatisch erforderlich. Der Trigger hängt insbesondere von der Art des Betreibers oder einer Nutzung nach Anhang III Nr. 5(b) oder 5(c) ab.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI-CLS-10	Das Institut dokumentiert je Rechtseinheit, Einsatzland und Systemversion den AI-Act-Scope, Ausnahmen und jede Operatorrolle.	alle Kandidaten	Scope and Role Assessment	Legal / Compliance Review	territorialer oder sachlicher Scope, Ausnahme oder Rolle ungeklärt
AI-CLS-11	Das Institut prüft versioniert die Relevanz von Anhang I, jede einschlägige Nummer von Anhang III und einen möglichen Art.-6-Abs.-3-Entscheid. Die Begründung bewertet signifikante Risiken, materiellen Entscheidungseinfluss, den gesetzlichen Ausnahmetatbestand und die Registrierungsfolge.	möglicher Hochrisikobezug	High-Risk Classification Record, Article 6(3) Reasoning Sheet	Legal / Compliance Review	Anhangspfad, Nummer, substanzielle Ausnahmebegründung oder Folgepflicht fehlt
AI-REG-01	Das Institut entscheidet die Registrierungspflicht getrennt für Provider, Art.-6-Abs.-3-Fälle und die besonderen öffentlichen Deployerfälle.	möglicher Art.-49-Trigger	Registration Decision Record	Legal / Compliance Review	Registrierung pauschal angenommen oder relevanter Trigger nicht geprüft
AI-REG-02	Das Institut prüft vor Einreichung Vollständigkeit und Qualität der einschlägigen Felder aus Anhang VIII.	bestätigter Registrierungstrigger	Annex VIII Field Quality Review	Evidence Review	Pflichtfeld unvollständig, widersprüchlich oder ungeprüft
AI-REG-03	Das Institut reicht die Registrierung fristgerecht ein, referenziert sie intern und aktualisiert sie bei relevanten Änderungen.	bestätigter Registrierungstrigger	EU Database Registration Evidence, Register Link	Evidence Review	fehlende, verspätete oder veraltete Registrierung

8. Data, Privacy und Confidentiality Safeguards

Datenbezogene Safeguards verbinden die fachliche Datennutzung mit Datenschutz und Vertraulichkeit. Sie schaffen belastbare Entscheidungen darüber, welche Daten unter welchen Schutzbedingungen für den AI Use Case verwendet werden dürfen.

Second Line und Revision prüfen, ob Datenherkunft, zulässige Nutzung und Schutzbedarf belegt sind. Die Freigabe bleibt ohne diese Evidence nicht nachvollziehbar.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI-DAT-01	Das Institut bestimmt und dokumentiert die Datenklassen der AI-Nutzung.	B, E, C, H	Data Classification Record	Evidence Review	Datenklasse fehlt
AI-DAT-02	Der AI Owner bindet den zuständigen Data Owner in die Datenentscheidung ein.	E, C, H, Retrieval-Augmented Generation (RAG)	Data Owner Sign-off	Challenge Review	Datenverantwortung unklar

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI DAT-03	Das Institut prüft bei personenbezogenen Daten Rechtsgrundlage und Anforderungen der Datenschutz-Grundverordnung (DSGVO).	personenbezogene Daten	Privacy Assessment	Datenschut zprüfung	Rechtsgrundlage fehlt
AI DAT-04	Das Institut prüft die Trigger für eine Datenschutz-Folgenabschätzung (DSFA) und für Art. 22 DSGVO.	relevante Personenwirkung	DSFA Screening	Datenschut zprüfung	Trigger nicht bewertet
AI DAT-05	Das Institut bestimmt den Vertraulichkeitsbedarf und die zulässige Verarbeitung.	vertrauliche Daten	Confidentiality Decision	Evidence Review	Schutzbedarf ohne Entscheidung
AI DAT-06	Das Institut definiert und belegt Löschung und Aufbewahrung der relevanten Daten und Artefakte.	E, C, H	Retention Evidence	Evidence Review	kein Lösch- oder Aufbewahrungsnachweis
AI DAT-07	Das Institut prüft GPAL-Eingaben vor externer Übertragung auf unzulässige Daten und Secrets.	externe GPAL, Coding-Assistenten	Ingress Policy, Blocking Test, Exception Record	Technical Test	unzulässige Daten oder Secrets passieren den Kontrollpunkt
AI FRIA-01	Das Institut entscheidet den Art.-27-Trigger insbesondere für Anhang III N° 5(b) und 5(c) sowie weitere gesetzlich erfasste Deployerfälle.	H und einschlägige Annex-III-Nutzung	FRIA Trigger Record	Legal / Compliance Review	Trigger unvollständig oder pauschal aus Hochrisiko abgeleitet
AI FRIA-02	Das Institut dokumentiert vor der ersten Nutzung Prozess und Zweck, Dauer und Häufigkeit, betroffene Gruppen und spezifische Grundrechtsschäden.	bestätigter Art.-27-Trigger	FRIA Record	Evidence Review	Pflichtinhalt nach Art. 27 Abs. 1(a)–(d) fehlt
AI FRIA-03	Das Institut dokumentiert Human Oversight und Schutzmassnahmen einschliesslich Governance- und Beschwerdewegen und genehmigt die FRIA vor der ersten Nutzung.	bestätigter Art.-27-Trigger	Approved FRIA Record	Challenge Review	Oversight, Massnahme, Beschwerdeweg oder Freigabe fehlt
AI FRIA-04	Das Institut meldet das Ergebnis mit dem vorgesehenen Template an die Marktüberwachungsbehörde, soweit keine gesetzliche Ausnahme greift.	meldepflichtige FRIA	FRIA Notification Evidence oder begründete Ausnahme	Legal / Compliance Review	erforderliche Mitteilung oder Ausnahmebegründung fehlt
AI FRIA-05	Das Institut verwendet eine bestehende FRIA nur für hinreichend ähnliche Fälle und aktualisiert sie bei Änderungen oder Überalterung; eine DSFA bleibt getrennt verknüpft.	bestehende oder geänderte FRIA	FRIA Reuse/Change Decision, DSFA Link	Revalidation Review	ungeprüfte Wiederverwendung, veraltete FRIA oder Gleichsetzung mit DSFA

9. Betriebsfreigabe und Nutzungsfreigabe

Eine technisch freigegebene Plattform macht einen konkreten AI Use Case noch nicht nutzbar. Diese Safeguards halten Betriebsfreigabe, Nutzungsfreigabe und begrenztes Testen auseinander und

sichern die spätere Revalidierung ab.

Die Prüfung bestätigt, dass beide Entscheide eigenständig vorliegen und sich auf den tatsächlichen Betrieb beziehungsweise den konkreten AI Use Case beziehen. Eine Plattformfreigabe darf die fachliche Nutzungsentscheidung nicht ersetzen.

ID	Control Statement	Applicability
AI OPS-01	Das Institut betreibt produktive AI-Systeme nur mit dokumentierter Betriebsfreigabe.	produktive Systeme
AI OPS-02	Das Institut erlaubt produktive AI-Nutzungen nur mit dokumentierter Nutzungsfreigabe.	produktive AI Use Cases
AI OPS-03	Das Institut entscheidet getrennt über Betriebsfreigabe und Nutzungsfreigabe.	alle produktiven AI Use Cases
AI OPS-04	Der AI Owner dokumentiert Dauer, Daten, Nutzerkreis und Abbruchkriterien eines Piloten.	Pilot / Test
AI OPS-05	Das Institut legt für relevante AI Use Cases ein verbindliches Revalidierungsdatum fest.	E, C, H

Die Hochrisiko-Deployer-Safeguards werden als eigener Tabellenblock geführt, damit die langen gesetzlichen Aussagen und ihre Anwendbarkeit lesbar bleiben.

ID	Control Statement	Applicability
AI DEP-01	Das Institut nutzt und überwacht ein Hochrisiko-System gemäss der freigegebenen Betriebsanleitung.	Hochrisiko-Deployer
AI DEP-02	Das Institut weist die menschliche Aufsicht kompetenten, Hochrisiko-Deployer geschulten, befugten und unterstützten Personen zu.	
AI DEP-03	Das Institut stellt bei von ihm kontrollierten Inputdaten deren Relevanz und hinreichende Repräsentativität für den Zweck sicher.	Hochrisiko-Deployer mit Inputkontrolle
AI DEP-04	Hat das Institut Grund zur Annahme eines Risikos nach Art. 79 Abs. 1 bei anleitungsgemässer Nutzung, setzt es die Nutzung unverzüglich aus. Es informiert ohne unangemessene Verzögerung den Provider oder Distributor und die zuständige Marktüberwachungsbehörde.	Hochrisiko-Deployer mit Risikosignal nach Art. 26 Abs. 5
AI DEP-05	Das Institut bewahrt automatisch erzeugte Logs unter seiner Kontrolle mindestens sechs Monate oder nach einer längeren anwendbaren Regel auf.	Hochrisiko-Deployer mit kontrollierten Logs
AI DEP-06	Das Institut informiert vor Arbeitsplatznutzung die Arbeitnehmervertretung und betroffene Beschäftigte, soweit Art. 26 dies verlangt.	Hochrisiko-AI am Arbeitsplatz
AI DEP-07	Das Institut unterstützt die DSFA und informiert betroffene natürliche Personen über eine einschlägige Annex-III-Hochrisikoentscheidung.	DSFA- oder Art.-26-Abs.-11-Trigger

ID	Control Statement	Applicability
AI DEP-08	Das Institut bewahrt die vollständige Deployer-Evidence auf und arbeitet bei begründeten Behördenanfragen im Rahmen der gesetzlichen Pflichten mit.	Hochrisiko-Deployer
AI DEP-09	Bei einem schwerwiegenden Vorfall informiert das Institut sofort zuerst den Provider, danach den Importeur oder Distributor und die zuständige Marktüberwachungsbehörde. Ist der Provider nicht erreichbar, dokumentiert es den Versuch und aktiviert den gesetzlichen Ersatzpfad.	Hochrisiko-Deployer mit schwerwiegendem Vorfall

Die zugehörigen Nachweise und Prüfentscheidungen bleiben über die Control-ID eindeutig verbunden.

ID	Evidence	Test Type	Fail Criteria
AI OPS-01	Betriebsfreigabe-Record	Evidence Review	Betrieb ohne Freigabe
AI OPS-02	Nutzungsfreigabe-Record	Evidence Review	Nutzung ohne Freigabe
AI OPS-03	Register-Link	Applicability Review	Plattformfreigabe ersetzt Nutzungsfreigabe
AI OPS-04	Pilot Record	Walkthrough	Pilot ohne Dauer, Daten- oder Nutzergrenzen
AI OPS-05	Revalidierungsdatum	Evidence Review	kein Review-Datum
AI DEP-01	Instruction Conformance Record	Walkthrough	Nutzung weicht ungeprüft von der Betriebsanleitung ab
AI DEP-02	Oversight Assignment and Competence Record	Evidence Review	Aufsicht ohne Kompetenz, Befugnis oder Unterstützung
AI DEP-03	Input Control Record	Data Quality Review	kontrollierbare Inputs ungeprüft oder ungeeignet
AI DEP-04	Risk Suspension and Notification Record	Incident Walkthrough	Nutzung nicht unverzüglich ausgesetzt oder erforderlicher Adressat nicht informiert
AI DEP-05	Retention Rule, Log Export Test	Technical Test	Aufbewahrung zu kurz, ungeklärt oder Export nicht möglich
AI DEP-06	Workforce Information Record	Evidence Review	Nutzung vor erforderlicher Information
AI DEP-07	DSFA Handoff, Affected-Person Notice	Legal / Compliance Review	erforderliche Information oder Datenschutz-Handoff fehlt
AI DEP-08	Authority Cooperation Record	Evidence Review	erforderliche Information oder Zusammenarbeit nicht nachweisbar
AI DEP-09	Serious Incident Notification Record, Contact Attempt Evidence	Incident Walkthrough	Reihenfolge, sofortige Meldung oder Ersatzpfad fehlt

10. Model Risk und Evaluation Safeguards

Modellrisiken entstehen aus Leistungsgrenzen, Datenabhängigkeiten und Veränderungen im Betrieb. Die Safeguards verlangen deshalb eine risikogerechte Validierung und laufende Nachweise, mit denen unerwünschte Wirkung rechtzeitig erkannt werden kann.

Prüfbar wird damit, ob Leistungsgrenzen bekannt sind und Änderungen rechtzeitig einen erneuten Entscheid auslösen. Kritische Modelle benötigen dafür belastbare Validation Evidence und definierte Schwellen.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI MRM-01	Das Institut verknüpft bankfachlich oder aufsichtlich relevante Modelle mit dem Modellinventar.	bankfachlich oder aufsichtlich relevante Modelle	Model Inventory Link	Evidence Review	Modellreferenz fehlt
AI MRM-02	Das Institut validiert kritische oder modellrisikorelevante Modelle risikogerecht.	C, H, Modellrisiko	Validation Report	Model Review	Validierung fehlt
AI MRM-03	Das Institut definiert messbare Schwellen und Reaktionen für Modell- und Datendrift.	lernende oder datenabhängige Modelle	Monitoring Plan	Evidence Review	keine Drift-Schwellen
AI MRM-04	Das Institut prüft bei Personenwirkung Bias- und Fairness-Risiken.	Personenwirkung	Fairness Evidence	Challenge Review	Diskriminierungsrisiko ungeprüft
AI MRM-05	Das Institut protokolliert Overrides so, dass Muster und Wirksamkeit ausgewertet werden können.	entscheidungsnahe Nutzung	Override Log	Sample Test	Overrides nicht nachvollziehbar
AI MRM-06	Das Institut regelt und dokumentiert Modellwechsel und Modellstilllegung.	Modellwechsel / Stilllegung	Stilllegungs-Evidence	Evidence Review	alte Version nicht nachvollziehbar

11. Supplier und Third-Party Safeguards

Externe AI-Leistungen verlagern die Verantwortung des Instituts nicht auf den Anbieter. Die Safeguards machen Abhängigkeiten, Datenverwendung und Exit-Fähigkeit prüfbar, damit Drittparteirisiken über den gesamten Lebenszyklus steuerbar bleiben.

Die Prüfung klärt, ob das Institut auch bei Veränderungen des Suppliers entscheidungs- und handlungsfähig bleibt. Freigabefähig ist eine externe Leistung nur, wenn wesentliche Abhängigkeiten, Vertragsrechte und Exit-Voraussetzungen belegt sind.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI SUP-01	Das Institut schliesst vor Nutzung einer externen Leistung ein Supplier Assessment ab.	externe Leistung	Supplier Assessment	Evidence Review	Assessment fehlt
AI SUP-02	Das Institut prüft den Bezug zum Digital Operational Resilience Act (DORA) und einen erforderlichen Eintrag im Informationsregister.	Informations- und Kommunikationstechnologie-Drittpartei (ICT-Drittpartei)	DORA / Outsourcing Check, DORA Register Link, Arrangement-ID ⁴	Compliance Review	DORA-Bezug oder erforderlicher Registereintrag ungeklärt

⁴ DORA ist für ICT-Drittdienstleistungen und ICT-Risiken im jeweiligen Anwendungsbereich massgeblich. Eine DORA-Prüfung ersetzt keine gesonderte bankaufsichtliche Auslagerungsprüfung, wenn diese zusätzlich ausgelöst wird.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI SUP-03	Das Institut erfasst die für die Leistung eingesetzten Subprozessoren.	externe Datenverarbeitung	Subprocessor Record	Evidence Review	Subprozessoren unbekannt
AI SUP-04	Das Institut entscheidet, ob und unter welchen Bedingungen Anbieter Institutsdaten für Training verwenden dürfen.	externe AI / GPAI	Training Exclusion Evidence	Contract Review	Trainingsnutzung unklar
AI SUP-05	Das Institut bewertet Exit-Fähigkeit und Portabilität vor der Freigabe.	E, C, H	Exit Evidence	Challenge Review	kein belastbarer Exit
AI SUP-06	Das Institut bewertet Konzentrationsrisiken kritischer Supplier.	kritische Supplier	Concentration Record	Risk Review	kritische Abhängigkeit ungeprüft

12. Secure AI Safeguards

AI-spezifische Angriffspfade müssen in die technische Kontrollumgebung des Instituts übersetzt werden. Diese Safeguards verbinden den AI Use Case mit der AI Security Baseline und verlangen Evidence für besonders exponierte Architekturen wie RAG und Agenten.

Second Line und Revision beurteilen hier nicht einzelne technische Konfigurationen, sondern die Vollständigkeit und Belastbarkeit der technischen Evidence. Die detaillierten Anforderungen bleiben in der AI Security Baseline.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI SEC-01	Das Institut ordnet jedem technischen AI-System die anwendbaren Requirements der AI Security Baseline zu.	alle technischen Systeme	Security Evidence Package	Evidence Review	Baseline nicht bewertet
AI SEC-02	Das Institut erstellt für exponierte AI-Systeme ein aktuelles Threat Model.	C, H, Agent, RAG, extern	Threat Model	Security Review	Angriffspfade ungeprüft
AI SEC-03	Das Institut prüft Prompt- und Output-Risiken generativer AI-Systeme.	GenAI	Test Evidence	Security Test Review	Prompt-/Output-Risiko ungeprüft
AI SEC-04	Das Institut begrenzt und überwacht Rechte von Agenten und Tool-Integrationen.	Agent / Tool-calling / Model Context Protocol (MCP)	Tool Inventory	Technical Evidence Review	kritische Tools ohne Limits
AI SEC-05	Das Institut kontrolliert Berechtigungen, Herkunft und Qualität der RAG-Quellen.	RAG	Knowledge Base Evidence	Technical Evidence Review	Berechtigungen oder Quellenqualität ungeprüft
AI SEC-06	Das Institut führt für kritische generative AI-Systeme risikobasiertes Red Teaming oder Abuse Testing durch.	C, H, kritische GenAI	Red Team Report	Security Review	kritischer AI Use Case ohne Abuse Test

13. Incident, Monitoring und Reporting Safeguards

Kontrollen bleiben nur wirksam, wenn Abweichungen im Betrieb erkannt und in steuerbare Entscheidungen überführt werden. Die Safeguards verbinden Monitoring, Incident-Triage, Meldeprüfung und Management Reporting zu einer nachvollziehbaren Evidence-Kette.

Die Prüfung zeigt, ob Signale zu rechtzeitigen Entscheidungen führen und Findings bis zum Abschluss verfolgt werden. Fehlende Eskalation oder nicht bearbeitete Massnahmen stellen die Wirksamkeit der Freigabe infrage.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI INC-01	Das Institut definiert messbare Monitoring-Signale und Eskalationsschwellen.	E, C, H	Monitoring Plan	Evidence Review	keine Kontrollsignale
AI INC-02	Das Institut unterhält für produktive AI-Systeme einen dokumentierten AI-Incident-Triage-Pfad.	alle produktiven Systeme	Incident Triage Record	Walkthrough	kein AI-Triage-Pfad
AI INC-03	Das Institut prüft regulatorische und vertragliche Meldepflichten parallel zur Incident-Behandlung.	relevante Incidents	Reporting Decision Record	Walkthrough	Meldepflichten nicht parallel geprüft
AI INC-04	Das Institut verfolgt Lessons Learned bis zur wirksamen Anpassung von Verantwortlichkeiten, Abläufen oder Kontrollen.	Incidents / Findings	Action Tracker	Evidence Review	keine Massnahmenverfolgung
AI INC-05	Das Institut aktiviert bei bestätigter Providerrolle einen Post-Market-Monitoring-Plan und den gesetzlichen Provider-Serious-Incident-Prozess; für reine Deployer gelten AI DEP-04 und AI DEP-09.	Provider-Aktivierungsentscheid positiv	Post-Market Plan, Provider Incident Record	Legal / Compliance Review	Providerfall ohne Plan, Meldeuhr, Zuständigkeit oder Evidence
AI REP-01	Das Institut stellt prüfbare Evidence für Key Performance, Risk und Control Indicators bereit.	E, C, H	Reporting Evidence	Evidence Review	Reportingdaten fehlen
AI REP-02	Das Institut berichtet mindestens quartalsweise an die Geschäftsleitung und mindestens jährlich an den Verwaltungsrat; rote Ereignistrigger werden unverzüglich eskaliert.	AI-Portfolio und aktive AI Use Cases	Management Reporting Calendar, Escalation Record	Evidence Review	Mindestfrequenz oder ereignisbasierte Eskalation fehlt

14. AI Literacy, Human Oversight und Communication Safeguards

Menschliche Aufsicht setzt neben Wissen auch klare Befugnisse und überprüfbare Eingriffsmöglichkeiten voraus. Diese Safeguards prüfen daher nicht nur Kommunikation und Training, sondern auch deren praktische Wirksamkeit im jeweiligen Nutzungskontext.

Entscheidend ist, ob verantwortliche Personen Ergebnisse beurteilen, eingreifen und eine Nutzung stoppen können. Teilnahmebestätigungen allein reichen als Evidence nicht aus.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI LIT-01	Das Institut belegt die rollenbezogene AI-Kompetenz der relevanten Personen.	alle relevanten Rollen	Training Evidence	Evidence Review	Rollentraining fehlt
AI LIT-02	Das Institut kommuniziert zulässige Nutzung, Grenzen und Meldewege an betroffene Nutzer.	produktive Nutzung	Communication Record	Evidence Review	Nutzungsrahmen unbekannt

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI LIT-03	Das Institut weist entscheidungsnahen AI-Nutzungen eine befugte menschliche Aufsicht zu.	entscheidungsnahe Nutzung	Oversight Assignment	Walkthrough	Aufsicht ohne Befugnis
AI LIT-04	Das Institut testet, ob die menschliche Aufsicht Ergebnisse verwerfen, eingreifen und stoppen kann.	C, H	Oversight Test	Challenge Review	kein Eingriff möglich
AI LIT-05	Das Institut misst die Wirksamkeit der AI Literacy über Teilnahmequoten hinaus.	E, C, H	Effectiveness Metrics	Evidence Review	nur Teilnahmequote ohne Wirksamkeit
AI TRA-01	Das Institut entscheidet je System und Nutzung, welcher Tatbestand aus Art. 50 gilt, welche Rolle verpflichtet ist und welche Ausnahme verwendet wird.	mögliche direkte Interaktion, Emotion/Biometrie oder generierte Inhalte	Transparency Decision Record	Legal / Compliance Review	Tatbestand, Rolle, Zeitpunkt oder Ausnahme ungeklärt
AI TRA-02	Das Institut informiert natürliche Personen klar, zugänglich und spätestens bei der ersten Interaktion, wenn es Provider eines direkt interagierenden AI-Systems ist und die AI-Natur nicht offensichtlich ist.	bestätigter Art.-50-Abs.-1-Trigger	Interaction Notice and Test	User Journey Test	Hinweis fehlt, ist verspätet oder unverständlich
AI TRA-03	Das Institut informiert exponierte Personen klar und zugänglich über eingesetzte Emotions- oder biometrische Kategorisierungssysteme.	bestätigter Art.-50-Abs.-3-Trigger	Exposure Notice Record	Evidence Review	erforderlicher Hinweis fehlt
AI TRA-04	Das Institut kennzeichnet Deepfakes und AI-generierte oder manipulierte Texte von öffentlichem Interesse; eine Ausnahme wegen menschlicher Prüfung und redaktioneller Verantwortung wird begründet.	bestätigter Art.-50-Abs.-4-Trigger	Publication and Exception Record	Publication Review	Disclosure fehlt oder Ausnahme unbegründet
AI TRA-05	Das Institut versieht synthetische Inhalte bei eigener Providerrolle technisch machbar mit wirksamer, interoperabler, robuster und zuverlässiger maschinenlesbarer Markierung und prüft deren Erkennbarkeit.	Provider und Art.-50-Abs.-2-Trigger	Marking Design and Detection Test	Technical Test	Markierung oder dokumentierte Machbarkeitsbegründung fehlt
AI EXPL-01	Das Institut entscheidet, ob Art. 86 für eine individuell nachteilige, rechtlich oder ähnlich erheblich wirkende Entscheidung auf Basis eines erfassten Annex-III-Systems gilt und Gesundheit, Sicherheit oder Grundrechte nachteilig berührt.	möglicher Art.-86-Trigger	Explanation Applicability Decision	Legal / Compliance Review	Trigger, nachteilige Wirkung oder Subsidiarität ungeklärt

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI EXPL- 02	Das Institut benennt einen zugänglichen Anfrageweg und eine verantwortliche Stelle; eine interne SLA wird nicht als gesetzliche AI-Act-Frist dargestellt.	bestätigter oder möglicher Art.-86-Trigger	Request Route and SLA Record	Walkthrough	kein Anfrageweg, Owner oder irreführende Frist
AI EXPL- 03	Das Institut gibt eine klare und aussagekräftige Erklärung zur Rolle des Hochrisiko-Systems und zu den Hauptelementen der Entscheidung frei und bewahrt sie als Evidence auf.	bestätigter Art.-86-Anspruch	Approved Explanation Response	Quality Review	Antwort unverständlich, nicht freigegeben oder nicht nachweisbar

15. Assurance, Testing und Evidence Safeguards

Assurance prüft, ob Safeguards passend ausgewählt wurden und wirksam arbeiten. Sie prüft nicht blind jedes Kontrollfeld für jeden AI Use Case.

Testart und Prüftiefe richten sich nach Profilkategorie, Trigger und Kontrollziel. Das Ergebnis muss eine reproduzierbare Aussage über Design und Wirksamkeit ermöglichen.

ID	Control Statement	Applicability	Evidence	Test Type	Fail Criteria
AI REGW- 01	Das Institut führt einen datierten Regulatory Watch Record mit Quelle, formellem Rechtsstatus, Anwendungsdatum, betroffenen Controls, Owner und Umsetzungsfrist.	relevante Rechts- oder Guidance-Änderung	Regulatory Watch Record	Legal / Compliance Review	Status, Auswirkung, Owner oder Frist fehlt
Test Type		Zweck			
Evidence Review		prüft, ob erforderliche Nachweise existieren			
Applicability Review		prüft, ob Auswahl und Nichtanwendbarkeit logisch sind			
Challenge Review		prüft kritische Entscheidungen durch Second Line			
Walkthrough		prüft Rollen, Handoffs und Ereignisreaktion			
Sample Test		prüft Wirksamkeit anhand aktiver AI Use Cases			
Technical Evidence Review		prüft Nachweise aus der Security Baseline			

16. Deviation, Compensating Controls und Risk Acceptance

Abweichungen sind nur zulässig, wenn sie dokumentiert, befristet und genehmigt sind.

Kompensationskontrollen müssen die gleiche Risikowirkung adressieren. Für verbotene Praktiken gibt es keine Risk Acceptance.

Die prüfende Stelle beurteilt, ob Genehmigung, Kompensationskontrolle, Laufzeit und Review-Datum vollständig sind. Eine offene oder unbefristete Abweichung erfüllt diese Anforderung nicht.

17. Mapping zu Regulierung und Frameworks

Mappings erklären, aus welchem Regelungs- oder Framework-Kontext eine Kontrollerwartung stammt. Sie unterstützen die Nachvollziehbarkeit, ersetzen aber weder die Prüfung der konkreten Anwendbarkeit noch die Auslegung der massgeblichen Quelle.

Mapping area	Control Standard treatment
AI Act	Safeguards für Scope, Rollen, Hochrisiko, Deployerpflichten, FRIA, Registrierung, Transparenz, Erläuterung, Provider-Aktivierung und Incidents
DORA	Safeguards für ICT-Risiko, Supplier, Resilienz, Incident-Evidence und das Informationsregister nach Art. 28 Abs. 3
DSGVO	Safeguards für Rechtsgrundlage, DSFA, Art. 22, TOMs und Löschung
KWG / MaRisk	nur bei bankfachlicher Steuerungswirkung, IKS, Modellrisiko oder Auslagerung
BSI / NIST / OWASP	Offene Qualitäts- und Kontrollreferenzen für angemessene Safeguards

18. Evidence und Nachweisführung

Evidence wird am jeweiligen AI Use Case oder AI-System referenziert und bleibt über den Lifecycle auffindbar. Ein Evidence Package bündelt die für einen Entscheid erforderlichen Records und Links, ohne führende Register oder Quellsysteme zu duplizieren. Evidence Review bezeichnet die dokumentierte Prüfung dieser Nachweise.

Vollständigkeit allein belegt keine Wirksamkeit. Evidence muss aktuell, dem richtigen Owner zugeordnet und für den angegebenen Safeguard aussagekräftig sein. Findings, Ausnahmen und Folgeentscheide bleiben mit dem ursprünglichen Evidence Package verknüpft.

19. Glossar

Begriff	Verbindliche Bedeutung
Applicability Statement	dokumentierter Entscheid über Profilkategorie, Trigger, anwendbare Safeguards, Nichtanwendbarkeit und Abweichungen
Provider-Aktivierungsentscheid	release-blockierender Kontrollentscheid, der bei bestätigter Providerrolle einen Provider Compliance Plan verlangt; englische Referenz: Provider Activation Gate
Requirements and Applicability Register	kontrollierte Zuordnung atomarer Pflichten zu Rolle, Trigger, Control, Workflow, Evidence und Quellenstand
AI	Artificial Intelligence; in der Dokumentfamilie einheitlich verwendete Bezeichnung für künstliche Intelligenz.
Evidence	prüfbarer Nachweis für einen Safeguard oder Entscheid
Evidence Package	referenzierte Gesamtheit der für einen Entscheid erforderlichen Evidence
Evidence Review	dokumentierte Prüfung von Aktualität, Aussagekraft und Zuordnung der Evidence
Safeguard	prüfbare Controllerwartung dieses Standards
Trigger	definiertes Merkmal, das zusätzliche Prüfung oder Safeguards auslöst
Risk Acceptance	befristete, begründete und genehmigte Annahme eines Restrisikos durch die zuständige Entscheidungsrolle; mit Owner und Review-Datum dokumentiert
Betriebsfreigabe	Entscheid über den sicheren Betrieb einer Plattform oder eines AI-Systems unter definierten Bedingungen
Nutzungsfreigabe	Entscheid über einen konkreten AI Use Case, seine Daten, Nutzer und Wirkung

Begriff	Verbindliche Bedeutung
Revalidierung	erneute Bewertung nach definiertem Zeitpunkt oder wesentlicher Änderung
Stilllegung	kontrollierte Beendigung einer AI-Nutzung oder eines AI-Systems
FRIA	Fundamental Rights Impact Assessment; Grundrechte-Folgenabschätzung nach Art. 27 AI Act bei bestätigtem Trigger.
DSFA	Datenschutz-Folgenabschätzung nach Art. 35 DSGVO bei voraussichtlich hohem Datenschutzrisiko.
DSGVO	Datenschutz-Grundverordnung; Rechtsrahmen für die Verarbeitung personenbezogener Daten.
DORA	Digital Operational Resilience Act; europäischer Rechtsrahmen für digitale operationale Resilienz im Finanzsektor.
GPAI	General-Purpose AI; vielseitig einsetzbares AI-Modell mit möglichen Downstream-Risiken.
ICT	Information and Communication Technology; Informations- und Kommunikationstechnologie im DORA-Kontext.
KPI / KRI / KCI	Leistungs-, Risiko- und Kontrollindikatoren für Management Reporting und Kontrollwirksamkeit.
RAG	Retrieval-Augmented Generation; Verbindung generativer AI mit kontrollierten Wissensquellen.
MCP	Model Context Protocol; Protokoll zur Anbindung von Modellen an Tools und Kontextquellen.
CISO	Chief Information Security Officer; Führungsrolle für Informationssicherheit und zugehörige Sicherheitsrisiken.
Digital.gov Plain Language	Offene Methode der U.S. General Services Administration für verständliche und testbare Erklärungstexte.
4T2 Professional Prose	4T2 Schreibmethode für klare fachliche Erklärungen ausserhalb strukturierter Control Statements.
NIST SP 800-53	Offener NIST-Kontrollkatalog, der als Muster für atomare Kontrollergebnisse dient.
NIST SP 800-53A	Offene NIST-Prüfmethode, die als Muster für Evidence und Verifikation dient.
OSCAL	Offenes NIST-Modell für stabile, strukturierte und maschinenlesbare Kontrollinformationen.

Rollen, Lifecycle und weitere Führungsbegriffe sind im AI Governance Framework for Financial Institutions definiert. Technische Begriffe stehen in der AI Security Baseline.

20. Quellen

Primary regulatory sources validated on 2026-07-11:

- Regulation (EU) 2024/1689 (AI Act), especially Articles 4-6, 9-15, 25-27, 49, 50 and 71 as well as Annex VIII: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- Council of the EU, final adoption of the Digital Omnibus on AI on 29 June 2026; publication, entry into force and resulting applicability dates remain under regulatory watch:

<https://www.consilium.europa.eu/en/press/press-releases/2026/06/29/artificial-intelligence-council-gives-final-green-light-to-simplify-and-streamline-rules/>

- Regulation (EU) 2022/2554 (DORA), especially Articles 5-14 and 28-30, and Commission Implementing Regulation (EU) 2024/2956: <https://eur-lex.europa.eu/eli/reg/2022/2554/oj/eng>
- Regulation (EU) 2016/679 (GDPR), especially Articles 5, 6, 22, 25, 28, 32 and 35: <https://eur-lex.europa.eu/eli/reg/2016/679/oj/eng>
- BaFin MaRisk, Circular 06/2024 (BA), only for conditional banking-governance mappings: https://www.bafin.de/SharedDocs/Veroeffentlichungen/DE/Rundschreiben/2024/rs_06_2024_MaRisk_BA.html

Supporting control and testing references:

- Digital.gov Plain Language Guide Series: <https://digital.gov/guides/plain-language>
- NIST SP 800-53 Rev. 5 and NIST SP 800-53A Rev. 5:
<https://csrc.nist.gov/pubs/sp/800/53/r5/upd1/final> and
<https://csrc.nist.gov/pubs/sp/800/53/a/r5/final>
- NIST Open Security Controls Assessment Language (OSCAL):
<https://pages.nist.gov/OSCAL/>
- NIST AI RMF 1.0 and current NIST AI RMF resources: <https://www.nist.gov/itl/ai-risk-management-framework>
- BSI AI Finance Test Criteria Catalogue:
https://www.bsi.bund.de/SharedDocs/Downloads/EN/BSI/KI/AI-Finance_Test-Criteria.html
- OWASP Top 10 for LLM Applications 2025: <https://genai.owasp.org/resource/owasp-top-10-for-llm-applications-2025/>

NIST marks AI RMF 1.0 as under revision. It remains a supporting framework, not a binding source, and must stay under source watch.

21. Haftungs- und Anwendungshinweis

Dieser Standard wurde mit fachlicher Sorgfalt und gegen den angegebenen Rechtsstand erstellt. Eine Gewähr für Vollständigkeit, Fehlerfreiheit oder dauerhafte Aktualität wird nicht übernommen. Er ist weder Rechtsberatung noch regulatorische Freigabe, Zertifizierung oder Konformitätsbescheinigung.

Welche Safeguards gelten und ob ihre Evidence wirksam ist, muss für jede Rechtseinheit, jedes Einsatzland, jede Operatorrolle, Systemversion und konkrete AI-Nutzung gesondert entschieden und geprüft werden. Die Verantwortung bleibt beim anwendenden Institut.

Copyright und offene Lizenz

© 2026 4T2 Solutions GmbH. Text und von 4T2 Solutions GmbH erstellte Diagramme dieses Dokuments stehen unter der Creative Commons Namensnennung 4.0 International Lizenz (CC BY 4.0): <https://creativecommons.org/licenses/by/4.0/>. Kopieren, Bearbeiten und weitere Nutzung sind erlaubt, sofern 4T2 Solutions GmbH als Quelle genannt, die Lizenz verlinkt und vorgenommene Änderungen kenntlich gemacht werden. Rechte Dritter sowie Namen, Logos und Kennzeichen von 4T2 Solutions GmbH sind davon ausgenommen.

Transparenz zur Erstellung

Martin Bischoff hat dieses Dokument mit Unterstützung generativer AI erstellt. AI wurde für Formulierungsentwürfe, redaktionelle Überarbeitung, grafische Gestaltung sowie die strukturierte Prüfung auf inhaltliche Vollständigkeit und Konsistenz eingesetzt. Auswahl, fachliche Bewertung und Freigabe erfolgten durch Martin Bischoff. Die AI-Unterstützung ersetzt keine unabhängige Rechts-, Wirksamkeits- oder Konformitätsprüfung.