

MRIS v2.0 — Mythos-resistente Informationssicherheit

MRIS

Wirksamkeitsbewertung bestehender ISO 27001 Annex A Controls unter Gen-AI-beschleunigter Offensive
Framework-übergreifend · Evidenzbasiert

Version 2.0 | August 2026

Autor: Richard Peddi (mit AI-Unterstützung)

ZIELGRUPPE

CISOs, ISMS-Verantwortliche und Sicherheitsarchitekten, die ihre bestehenden Umsetzungen von Controls gegen die neue Realität agentischer und Gen-AI-beschleunigter Angriffe prüfen müssen

Lizenz und Haftungsausschluss

© 2026 Richard Peddi

Dieses Werk ist lizenziert unter der Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0).

Sie dürfen dieses Werk:

- teilen – das Material in jedwedem Format oder Medium vervielfältigen und weiterverbreiten
- bearbeiten – das Material remixen, verändern und darauf aufbauen
- für beliebige Zwecke nutzen, auch kommerziell

Unter folgenden Bedingungen:

Namensnennung – Sie müssen angemessene Urheber- und Rechteangaben machen, einen Link zur Lizenz beifügen und angeben, ob Änderungen vorgenommen wurden.

Weitergabe unter gleichen Bedingungen – Wenn Sie das Material remixen, verändern oder anderweitig direkt darauf aufbauen, dürfen Sie Ihre Beiträge nur unter derselben Lizenz wie das Original verbreiten.

Lizenztext: [\[https://creativecommons.org/licenses/by-sa/4.0/\]](https://creativecommons.org/licenses/by-sa/4.0/)(<https://creativecommons.org/licenses/by-sa/4.0/>)

Haftungsausschluss

Dieses Schnellheft stellt eine technische und organisatorische Referenz dar. Es ersetzt keine individuelle Risikoanalyse, keine rechtliche Beratung und keine auditspezifische Prüfung. Die hier vorgenommene Bewertung der Wirksamkeit einzelner Controls in ihrer typischen Umsetzung bezieht sich auf den zum Erstellungszeitpunkt öffentlich dokumentierten Stand agentischer AI-Bedrohungen. Die Bewertung kann sich bei neuer Evidenz verschieben.

Hinweis zur Erstellung

Dieses Dokument wurde unter Einsatz von AI-Werkzeugen erstellt. Auswahl und Bewertung der Quellen, die fachlichen Einstufungen sowie die inhaltliche Verantwortung liegen beim Autor.

Zitierempfehlung

Peddi, Richard (2026): MRIS – Mythos-resistente Informationssicherheit, Version 2.0. [mris.info](https://www.mris.info)

Über den Autor

Richard Peddi ist Fachinformatiker für Systemintegration und geprüfter IT-Spezialist, Certified IT-Business Manager (Operative Professional) und hält einen M.Sc. in Information Security Management. Er ist unter anderem Lead Auditor für ISO/IEC 27001. Beruflich war er als Security Consultant im KRITIS-Umfeld und als CISO tätig, bevor er seine heutige Position als Group CISO übernahm. MRIS entsteht unabhängig von seiner beruflichen Tätigkeit.

Inhaltsverzeichnis

Lizenz und Haftungsausschluss	2
Haftungsausschluss	2
Hinweis zur Erstellung	2
Zitierempfehlung	2
Über den Autor	2
Inhaltsverzeichnis	3
Vorwort	6
Management Summary	7
Zentraler Befund	7
Die Bedrohungslage hat sich verschoben	7
Das erfordert eine Anpassung des Sicherheitsmanagements	7
1 Einleitung	7
1.1 Warum dieses Schnellheft existiert	7
1.2 Die zentrale These: Methodik intakt, Eingaben destabilisiert	8
1.3 Der zentrale Claim	8
1.4 Position im ISMS-Stack und Grenzen des Schnellhefts	8
2 Feststellungen: Die vier Verschiebungen der Angriffsrealität	9
2.1 Kollaps des Patch-Gap-Fensters	9
2.2 Kompression der Zeitachse	9
2.3 Fragmentierung des Bedrohungsereignisses	10
2.4 Entkopplung von Fähigkeit und Akteur	10
2.5 Entwicklungen seit der Erstveröffentlichung: Vom Ein-Anbieter-Ereignis zum Multi-Vendor-Markt	10
2.6 Konsequenz für ISO/IEC 27005	11
3 Methodik der Bewertung	11
3.1 Das Vier-Kategorien-Raster	11
3.1.1 Standfest	11
3.1.2 Teilweise degradiert	11
3.1.3 Reine Reibung	11
3.1.4 Nicht betroffen	11
3.1.5 Hinweis zur Kontextabhängigkeit der Klassifikation	12
3.2 Bewertungskriterien	12
3.3 Framework-Stack und Priorisierung	12
3.4 Quellenbasis	13
3.5 Grenzen	14
Teil II – Hauptanalyse: Die 93 Controls der ISO/IEC 27002:2022	14
4 Standfeste Controls – das belastbare Fundament	15
4.1 Auswahllogik	15
4.2 Controls im Detail	15
4.3 Zusammenfassung	21

5 Teilweise degradierte Controls – Wirksamkeit mit Vorbehalt	22
5.1 Übersicht und Degradationsmuster	22
5.2 Controls im Detail	22
5.3 Zusammenfassung	35
6 Reine Reibung – Controls, deren klassische Implementierungsform ersetzt oder fundamental gehärtet werden muss.....	36
6.1 Übersicht und Risikoprofil.....	36
6.2 Controls im Detail	36
6.3 Zusammenfassung	38
7 Nicht betroffene Controls – neutrale Basis.....	38
7.1 Controls ohne Mythos-Relevanz.....	38
7.2 Übersicht der 23 nicht betroffenen Controls.....	39
Teil III – Konvergenzanalyse: Lücken und Synthese.....	40
8 Lückenanalyse – was andere Frameworks explizit fordern, das die ISO 27002 nur implizit abdeckt	40
8.1 Methodik der Lückenanalyse	40
8.2 Cluster 1: Post-Quantum-Strategie und Crypto-Agility.....	40
8.3 Cluster 2: Supply-Chain-Transparenz und SBOM-Pflicht	41
8.4 Cluster 3: Container, Confidential Computing und Multi-Tenancy.....	41
8.5 Cluster 4: Verbindliche Meldefristen und Root-Cause-Analyse.....	42
8.6 Cluster 5: Kontinuierliche Prüfung zusätzlich zu periodischen Audits	42
8.7 Cluster 6: Phishing-resistente Identität und Zero-Trust-Architektur	42
8.8 Cluster 7: Automatisierungspflicht und Resilienz-Testing	43
8.9 Zusammenfassung: Das Lücken-Delta der ISO 27002.....	43
9 Synthese – Katalog der Mythos-Härtungs-Controls.....	43
9.1 Zweck und Anwendungsweise des Katalogs.....	43
9.2 Dreizehn Mythos-Härtungs-Controls	44
9.3 Ableitung aus den Lückenclustern	49
9.4 Kompakt-Übersicht des MHC-Katalogs.....	50
9.5 Reifegrad-Stufen pro MHC.....	51
Teil IV – Handlungsempfehlungen und Ausblick.....	52
10 Handlungsempfehlungen für die ISMS-Anpassung.....	52
10.1 Geltungsbereich und Grenzen der Empfehlungen.....	52
10.2 Sofort-Neubewertung bestehender Controls	53
10.3 Strukturelle Härtung: Übernahme der MHC ins SoA	53
10.4 Reifung des ISMS-Prozesses.....	54
10.5 Vorstandskommunikation und Risikodialog	54
10.6 Mythos-relevante Kennzahlen.....	55
10.7 Zusammenfassung der Empfehlungen	56
11 Reflexion und Ausblick.....	56
11.1 Grenzen dieser Arbeit	56
11.2 Was dieses Schnellheft nicht leistet	56

11.3 Erwartete Weiterentwicklungen.....	56
11.4 Bedrohungshorizont	57
11.5 Schlussbemerkung	57
Teil V – Ergänzungs-Layer: Weiterführende Frameworks	57
Anhänge – Arbeitsmaterialien	58
Anhang A – Bewertungsmatrix aller 93 Controls.....	59
Anhang B – Framework-Mapping für Kernthemen	61
Anhang C – MHC-Katalog als Standalone-Arbeitsblatt	62
Anhang D – Indikatoren und Monitoring-Quellen	63
Anhang E – RACI-Modell für die MHC-Umsetzung	64
Hinweise zur Anwendung	65
Anhang F – Risikobewertungs-Brücke nach ISO/IEC 27005	65
F.1 Eingabe: MRIS-Bewertungsergebnis je Control	65
F.2 Verarbeitungsschritte im Risikoprozess.....	65
Schritt 1: Risiko-Reassessment (ISO/IEC 27005:2022 Kap. 7.3)	65
Schritt 2: Behandlungsoptionen (ISO/IEC 27005:2022 Kap. 8.1)	66
Schritt 3: Restrisiko-Bewertung und Akzeptanz (ISO/IEC 27005:2022 Kap. 8.6)	66
F.3 Beispielhafte Anwendung	66
F.4 Ausgabe in das ISMS	67
Anhang G – KPI-Definitionen und Mess-Standardisierung	67
Glossar	70
Quellenverzeichnis.....	71
Normen und regulatorische Quellen	71
NIST-Publikationen	71
Threat Intelligence und Industrie-Publikationen	72
Versionshinweise	73

Vorwort

Anthropic dokumentierte im November 2025 mit GTG-1002 erstmals einen Cyberangriff, in dem Claude Code 80 bis 90 Prozent der taktischen Angriffsarbeit über rund 30 Ziele autonom ausführte. Im April 2026 bestätigten Claude Mythos Preview und das Glasswing-Defensivprogramm: AI-Modelle finden Schwachstellen und verketteten sie zu funktionsfähigen Exploits in einem Tempo, das klassische Reaktionszyklen strukturell überfordert.

Damit stellt sich die Kernfrage dieses Schnellhefts: Welche der heute in einem ISMS verankerten Controls halten dieser Realität noch stand – und welche nicht?

MRIS setzt methodisch auf ISO/IEC 27001:2022 und der Risikomethodik der ISO/IEC 27005:2022 auf, bewertet alle 93 Controls aus ISO/IEC 27002:2022 gegen die Mythos-Bedrohungslage, gleicht sie primär mit NIST, BSI CS:2026, DORA, CRA, NIS2-Richtlinie und NIS2-Durchführungsverordnung ab und liefert einen Katalog von dreizehn Mythos-Härtungs-Controls (MHC) mit Reifegrad-Stufen und audit-fähigen Schwellwerten. Weitere herangezogene Normen und Regelwerke – darunter BSI IT-Grundschutz, BSI TR-02102 und TR-03183, ISO/IEC 42001, MITRE ATT&CK und D3FEND, OWASP, CIS Controls, SLSA, FIDO2/WebAuthn und die DSGVO – sind in Kapitel 3.3 und 3.4 vollständig ausgewiesen.

MRIS ist kein vollwertiges ISMS-Framework und ersetzt weder ISO/IEC 27001 noch ISO/IEC 27005. Es ergänzt das bestehende ISMS um einen Wirksamkeits- und Realitäts-Layer (siehe Kap. 1.4 zur Positionierung im ISMS-Stack).

Im April 2026 hat eine Arbeitsgruppe aus Cloud Security Alliance, SANS Institute, [un]prompted und OWASP Gen AI Security Project das Strategiepapier „The AI Vulnerability Storm: Building a Mythos-ready Security Program“ vorgelegt. MRIS und Mythos-Ready sind komplementär: Mythos-Ready definiert die Risikolandschaft auf Industrie-Konsensebene, MRIS übersetzt sie in eine Control-by-Control-Bewertung und einen audit-fähigen Härtungskatalog.

Seit Mai 2026 ist diese Fähigkeit keine Ein-Anbieter-Beobachtung mehr: Mit OpenAI Daybreak und Microsoft MDASH ist die agentische Schwachstellen-Entdeckung Multi-Vendor-Realität, und mit Claude Fable 5 ist die Mythos-Klasse allgemein verfügbar. Kapitel 2.5 ordnet die Entwicklungen seit der Erstveröffentlichung ein.

Dieses Schnellheft entsteht als unabhängiges Werk außerhalb meiner beruflichen Tätigkeit. Es wird unter mrinfo.info veröffentlicht; Anregungen und Kritik sind ausdrücklich erwünscht.

Mein Dank gilt Daniel Heldt für die kritische fachliche Durchsicht der Version 1.8 und die daraus hervorgegangenen Verbesserungen.

Richard Peddi

Juli 2026

Management Summary

Zentraler Befund

Die etablierten Methoden der Informationssicherheit sind nicht falsch geworden – aber eine Annahme, auf der sie beruhen, stimmt nicht mehr: dass ein Angriff für den Angreifer langsam, aufwendig und teuer ist. Genau diese Bremswirkung entfällt, sobald Angreifer künstliche Intelligenz einsetzen. Ein erheblicher Teil des heutigen Schutzniveaus ruht damit auf einer Grundlage, die nicht mehr trägt.

Die Bedrohungslage hat sich verschoben

Vier Veränderungen sind auf Leitungsebene relevant:

- **Tempo:** Die Zeit zwischen dem Bekanntwerden einer Schwachstelle und ihrer Ausnutzung schrumpft von Wochen auf Stunden. Bisher ausreichende Reaktionsfenster sind zu kurz geworden.
- **Automatisierung:** Angreifer lassen den Großteil der Arbeit von Software autonom erledigen – schneller, als ein Team reagieren kann.
- **Tarnung:** Angriffe werden in viele kleine, für sich unauffällige Schritte zerlegt, die einzeln keinen Alarm auslösen.
- **Verfügbarkeit:** Angriffsfähigkeiten, die früher nur staatlichen Akteuren offenstanden, sind heute breit zugänglich. Der Kreis ernstzunehmender Angreifer wächst.

Das erfordert eine Anpassung des Sicherheitsmanagements

Die Antwort ist keine neue Sicherheitsstrategie, sondern eine gezielte Nachjustierung des bestehenden Managementsystems für Informationssicherheit (ISMS). Wo Schutz bisher darauf gesetzt hat, dass Angriffe für den Gegner zu aufwendig sind, muss er auf belastbarere Grundlagen umgestellt werden. Dieses Dokument zeigt strukturiert und priorisiert, welche vorhandenen Maßnahmen weiterhin tragen, welche nachgeschärft werden müssen und wo gezielt ergänzt werden sollte. Der Aufwand baut auf dem Bestehenden auf und ist überschaubar – er erfordert aber die bewusste Entscheidung, ihn jetzt anzugehen.

1 Einleitung

1.1 Warum dieses Schnellheft existiert

Der Anlass dieses Schnellhefts ist eine konkrete und öffentlich bestätigte Veränderung der Angriffsrealität. Anthropic hat zwischen August 2025 und April 2026 in mehreren Threat-Intelligence-Berichten dokumentiert, wie Claude-Modelle in realen Kampagnen eingesetzt wurden: als autonome Orchestratoren einer Cyberspionagekampagne gegen rund 30 globale Ziele, zur Exploit-Generierung sowie zur automatisierten Erstellung von Phishing-Inhalten und Schadsoftware. Mit der Veröffentlichung des Glasswing-Programms und der ersten Bewertung von Claude Mythos Preview im April 2026 hat der Hersteller öffentlich erklärt, dass Modelle mittlerweile Schwachstellen in Code in einem Tempo finden und ausnutzen können, das klassische Reaktionsprozesse strukturell überfordert.

Für die Informationssicherheit bedeutet das einen Bruch, der bereits eingetreten ist. Die Angreiferseite operiert in Zeitfenstern, die vor drei Jahren als unrealistisch galten: Patch-Reversing zu funktionierendem Exploit in Stunden statt Tagen, Intrusion-Operationen mit 80 bis 90 Prozent autonomer Ausführung, Zerlegung komplexer Angriffe in einzeln legitim erscheinende Mikroschritte.

Damit stellt sich die Frage, die jede ISMS-Verantwortung beantworten muss: Welche der heute in einem ISMS verankerten Controls halten dieser Realität noch stand? Und wo täuscht die Existenz eines Controls eine Schutzwirkung vor, die der Mythos-Angreifer längst nicht mehr respektiert?

Dieses Schnellheft liefert die Antwort als systematische Bewertung aller 93 Controls aus ISO/IEC 27002:2022, ergänzt um die Perspektiven der fünf primären Vergleichsframeworks: NIST, BSI-C5:2026-Prüfkatalog, DORA, CRA sowie NIS2-Richtlinie mit Durchführungsverordnung (EU) 2024/2690; der vollständige Framework-Stack einschließlich der sekundären Referenzen ist in Kapitel 3.3 ausgewiesen. Der Mehrwert liegt in der Konvergenzanalyse: Wo ein Framework eine Anforderung stellt oder expliziter formuliert, die die ISO so nicht

kennt oder nur bedingt impliziert und die unter Mythos-Bedingungen standhält, liegt ein Härtungsvorschlag für ISO-zentrierte ISMS vor.

Was im April 2026 als Anthropic-zentriertes Ereignis begann, ist seit Mai 2026 ein Multi-Vendor-Markt (Kap. 2.5). Die Kernthesen dieses Schnellhefts bleiben davon unberührt; ihre empirische Basis ist breiter geworden.

Begriffsklärung: „MRIS“ (Mythos-resistente Informationssicherheit) bezeichnet das hier vorgestellte, unabhängige Rahmenwerk; das vorliegende Schnellheft ist dessen Referenzfassung. Wo im Folgenden von „MRIS“ die Rede ist, ist stets das Rahmenwerk gemeint, nicht dieses einzelne Dokument.

1.2 Die zentrale These: Methodik intakt, Eingaben destabilisiert

ISO/IEC 27005:2022 definiert Risiko als „Auswirkung von Unsicherheit auf Ziele“ (3.1.3); das *Risikoniveau* (3.1.15) ergibt sich aus der Kombination von Schadensausmaß (consequence) und Eintrittswahrscheinlichkeit (likelihood). Für die Abschätzung der Eintrittswahrscheinlichkeit absichtlicher Risikoquellen nennt die Norm (Kap. 7, vgl. Anhang A.2.3) insbesondere die Motivation (Kosten-Nutzen-Abwägung des Angriffs), die Fähigkeiten und Ressourcen des Akteurs sowie die Ausnutzbarkeit der Schwachstelle. Diese Methodik ist nicht falsch – was sich verändert hat, sind ihre Eingaben.

Hinweis zur Autoren-Ergänzung: Zu den Risikogrößen, die ISO/IEC 27005:2022 für absichtliche Quellen heranzieht – Eintrittswahrscheinlichkeit, Schadensausmaß und die Angreiferfähigkeit (Motivation, Fähigkeiten, Ressourcen) – führt dieses Schnellheft als weitere Dimension die *Zeit zwischen Bedrohungsereignis und organisatorischer Reaktion* ein. Die Norm nennt diesen Zeitfaktor nicht explizit als Risikoparameter; er dient im Folgenden als zusätzliches Bewertungskriterium, weil Mythos-Angriffe die Reaktionszeit strukturell unter die menschlich handhabbare Schwelle drücken.

Mythos-Klasse-AI senkt für viele Angriffsklassen die Kapazitätsbarriere zwischen Gelegenheitstäter und fortgeschrittenem Akteur erheblich, komprimiert die Zeit zwischen Schwachstellen-Veröffentlichung und Exploit auf Stunden, fragmentiert Angriffe in einzeln nicht alarmierende Mikroereignisse und entwertet systematisch Controls, deren Schutzwirkung nur auf begrenzter Angreifergeduld beruht. Physische und hardware-nahe Angriffe – etwa Fault-Injection an Chips oder der Zugriff auf isolierte Systeme – bleiben davon weitgehend unberührt. Die zentrale These lautet daher: Die ISO/IEC 27005-Methodik bleibt anwendbar, aber die bisher eingehenden Wahrscheinlichkeits- und Wirksamkeitsparameter müssen neu bewertet werden – jedes Control daraufhin, ob es seine Schutzwirkung aus einer harten Barriere oder nur aus Reibung bezieht und ob diese Reibung unter Mythos-Bedingungen noch genügt.

1.3 Der zentrale Claim

Wenn NIST, C5, DORA, CRA, NIS2-Richtlinie oder die NIS2-Durchführungsverordnung eine Anforderung stellt oder expliziter formuliert, die die ISO 27002 so nicht kennt oder nur bedingt impliziert, und diese Anforderung unter Mythos-Bedingungen standhält, liegt darin ein konkreter Härtungsvorschlag für jedes ISO-zentrierte ISMS. Die CISO-Aufgabe verschiebt sich von Compliance-Tiefe zur wirksamkeitsorientierten Control-Suche.

1.4 Position im ISMS-Stack und Grenzen des Schnellhefts

MRIS ist kein Management-System-Standard und beansprucht nicht, ein vollständiges ISMS abzubilden. Diese Selbstverortung ist explizit Teil der Methodik.

MRIS als Wirksamkeits- und Realitäts-Layer auf einem ISMS-Fundament. MRIS setzt ein etabliertes ISMS voraus – typischerweise nach ISO/IEC 27001:2022, dem BSI IT-Grundschutz, dem BSI C5 2026 oder einem gleichwertigen Rahmen – und beantwortet auf dieser Basis eine einzige Frage: Wirken die heute implementierten Controls noch gegen Gen-AI-beschleunigte Offensive? Ohne ISMS-Fundament ist MRIS nicht anwendbar.

Was MRIS leistet:

- Wirksamkeits-Bewertung aller 93 Controls aus ISO/IEC 27002:2022 gegen die Mythos-Bedrohungslage (Teil II).
- Lückenanalyse zu primären Frameworks (BSI C5:2026, NIST, DORA, CRA, NIS2 mit Durchführungsverordnung) (Kap. 8).
- Audit-fähigen Katalog von dreizehn Mythos-Härtungs-Controls mit Reifegrad-Stufen, Tools und Schwellwerten (Kap. 9).

- Operationalisierung in vier Anhängen: Bewertungsmatrix (A), Standalone-MHC-Arbeitsblatt (C), RACI-Modell (E), KPI-Definitionen (G). Bewertungsmatrix und MHC-Katalog stehen zusätzlich als maschinenlesbare Fassung (YAML, CC BY-SA 4.0) bereit.

Was MRIS bewusst nicht leistet:

- Kein vollständiges ISMS nach ISO/IEC 27001 Clause 4–10 (Kontext, Führung, Planung, Unterstützung, Betrieb, Bewertung, Verbesserung). Diese Anforderungen werden vorausgesetzt, nicht ersetzt.
- Kein eigener Risikomanagement-Prozess. MRIS liefert über die Bewertungs-Kategorien (Kap. 3) und die Risikobewertungs-Brücke (Anhang F) Eingaben für den Risikoprozess nach ISO/IEC 27005, ersetzt diesen aber nicht.
- Kein vollständiges PDCA-Management-System. MRIS liefert über die Reifegrad-Stufen (Kap. 9.5), die Versionierung und die Continuous-Audit-Empfehlungen (MHC-10) Bausteine für Plan-Do-Check-Act, ersetzt aber kein eigenständiges Management-System.
- Kein Compliance-Mapping-Werkzeug. Die Framework-Vergleiche der Einzelcontrols sind Wirksamkeits-Querverweise, kein vollständiges Anforderungs-zu-Norm-Mapping mit Audit-Trail. Für Zertifizierungs- und regulatorische Nachweise sind dafür spezialisierte GRC-Werkzeuge zu nutzen.
- Keine organisationsspezifische Policy-Hierarchie. Policies, Verfahrensanweisungen und Arbeitsanweisungen müssen in der Organisation entwickelt werden; MRIS liefert dafür inhaltliche Anker, aber keine Templates.

Korrekte Nutzung. MRIS ist als Layer auf einem bestehenden ISMS gedacht: Das ISMS liefert Governance, Risikomanagement-Prozess, Policy-Hierarchie und PDCA-Zyklus. MRIS liefert die Wirksamkeits-Bewertung der einzelnen Controls gegen Mythos und konkrete, audit-fähige Härtungsempfehlungen. Die Trennung beider Layer ist methodisch entscheidend – MRIS soll ISMS-Frameworks nicht duplizieren, sondern punktgenau ergänzen, wo diese Frameworks aufgrund ihrer technologie-agnostischen Konzeption bei Gen-AI-beschleunigten Bedrohungen an ihre Grenzen stoßen.

2 Feststellungen: Die vier Verschiebungen der Angriffsrealität

Die folgenden vier Feststellungen sind keine Hypothesen. Sie beruhen auf den öffentlich dokumentierten Threat-Intelligence-Berichten von Anthropic aus dem Zeitraum August 2025 bis April 2026 sowie auf den Defensiv-Empfehlungen des Anthropic-Security-Engineering-Teams im Rahmen des Glasswing-Programms.

2.1 Kollaps des Patch-Gap-Fensters

Klassische Vulnerability-Management-Programme rechnen mit einem Zeitfenster zwischen Patch-Veröffentlichung und Verfügbarkeit eines funktionierenden Exploits. Dieses Fenster wurde in der Praxis meist in Tagen bis Wochen gemessen. Es erlaubte gestaffelte Patch-Zyklen, manuelle Freigabeprozesse und selektives Priorisieren nach Kritikalität.

Anthropic beschreibt Patch-Reversing zu funktionierendem Exploit als genau die Art mechanischer Analyse, in der aktuelle Modelle herausragen. Die veröffentlichte Defensivempfehlung fordert konsequent, internet-exponierte Systeme binnen 24 Stunden (Richtwert) nach Verfügbarkeit eines Exploits zu patchen. Für jedes Control, das auf einen geordneten Patch-Zyklus mit manueller Freigabe setzt, ist die Wirksamkeitsannahme gebrochen.

2.2 Kompression der Zeitachse

ISO 27005 wie auch viele Incident-Response-Playbooks setzen voraus, dass zwischen Erkennung, Triage und Entscheidung menschlich handhabbare Zeit liegt. Diese Zeit ist nicht mehr gegeben.

Im dokumentierten GTG-1002-Fall führte Claude Code 80 bis 90 Prozent der taktischen Angriffsschritte autonom aus, bei Anfrageraten, die für menschliche Operatoren nicht erreichbar sind. Kontrollen mit Human-in-the-Loop in der unmittelbaren Reaktionskette werden nicht wegen mangelnder Qualität unwirksam, sondern weil ihre Reaktionszeit nicht mehr in derselben Größenordnung wie die Angriffsgeschwindigkeit liegt.

2.3 Fragmentierung des Bedrohungsereignisses

Das Risikomodell der ISO 27005 setzt ein diskretes, identifizierbares Bedrohungsereignis voraus, gegen das Controls greifen. Diese Voraussetzung hält unter agentischen Angriffsmustern nicht mehr. Anthropic beschreibt für GTG-1002 explizit, dass der Angreifer den Gesamtangriff in Teilaufgaben zerlegte – Vulnerability-Scanning, Credential-Validierung, Lateral Movement, Datenextraktion –, die jeweils als legitime technische Anfragen erschienen.

Das Risiko materialisiert sich erst in der Aggregation. Controls, die auf der Evaluation diskreter Ereignisse basieren – signaturbasierte Erkennung ohne Verhaltenskontext, Rate-Limits mit naiven Schwellwerten, Policy-basiertes Alerting auf Einzelaktionen – verlieren Wirksamkeit, obwohl sie technisch funktionieren. Sie sehen, was sie sehen sollen. Sie sehen nicht, dass die Summe einen Angriff ergibt.

2.4 Entkopplung von Fähigkeit und Akteur

Die klassische Likelihood-Bewertung stützt sich wesentlich auf die Einschätzung der Angreiferfähigkeit: Ein Nation-State-Akteur galt gegen einen mittelständischen deutschen Softwarehersteller als unwahrscheinlich, nicht weil das Interesse fehlte, sondern weil die technische Kapazität für eine individualisierte Kampagne knapp war. Diese Kopplung zwischen Akteurstyp und verfügbarer Kapazität hat Mythos-Klasse-AI weitgehend aufgelöst.

Die Threat-Intelligence-Berichte zeigen, dass Akteure mit geringer eigener technischer Kompetenz mittlerweile in der Lage sind, Operationen auszuführen, die bislang nur professionellen APT-Gruppen zugeschrieben worden wären. Die Likelihood-Achse in jeder Risikomatrix verschiebt sich nach oben – nicht wegen neuer Bedrohungsakteure, sondern weil die Kapazitätsbarriere bestehender Akteure gefallen ist.

2.5 Entwicklungen seit der Erstveröffentlichung: Vom Ein-Anbieter-Ereignis zum Multi-Vendor-Markt

Zwischen Mai und Juli 2026 wurde aus der Anthropic-zentrierten Beobachtungslage ein Markt. Vier Entwicklungen sind für die Bewertungen dieses Schnellhefts relevant:

OpenAI Daybreak (11. Mai 2026). OpenAI kombiniert die GPT-5.5-Modellfamilie mit „Codex Security“ als agentischem Harness zu einem System für Bedrohungsmodellierung, Schwachstellen-Entdeckung, Exploit-Validierung in isolierten Umgebungen und Patch-Erzeugung; der Zugang ist gestuft (u. a. „GPT-5.5-Cyber“ für autorisiertes Red-Teaming).

Microsoft MDASH (12. Mai 2026). Der „Multi-Model Agentic Scanning Harness“ orchestriert über 100 spezialisierte Agenten über ein Ensemble aus Frontier- und destillierten Modellen. Erste Volumen-Evidenz: 16 zuvor unbekannte Schwachstellen im Windows-Netzwerk- und Authentisierungs-Stack, darunter vier kritische RCE (u. a. CVE-2026-33824, CVSS 9.8), gepatcht im Mai-Patchday; im öffentlichen CyberGym-Benchmark 88,45 % gegenüber 83,1 % der Mythos Preview.

General Availability der Mythos-Klasse (9. Juni / 1. Juli 2026). Mit Claude Fable 5 ist ein Modell der Mythos-Klasse allgemein verfügbar — mit Cyber-Schutzklassifikatoren; die Variante ohne diese Schutzstufen (Claude Mythos 5) bleibt Glasswing-Partnern vorbehalten. Für die Bewertungslage entscheidend: Die Fähigkeitsklasse diffundiert in die Breite, der Verteidigungsvorsprung einer exklusiven Preview entfällt schrittweise. Das verschärft die in 2.1 und 2.2 hergeleitete Exploit-Ökonomie.

Offenlegungswellen (ab Juli 2026). Project Glasswing meldete für die erste Programmphase über 10.000 High-/Critical-Findings; bis 22. Mai 2026 waren 1.596 Funde an Maintainer von 281 Open-Source-Projekten offengelegt, davon 97 gepatcht. Der angekündigte öffentliche 90-Tage-Bericht (Juli 2026) und auslaufende Embargofristen führen in H2 2026 zu CVE-Wellen. Parallel meldet die NVD strukturelle Überlast (Einreichungen +263 % gegenüber 2020; Enrichment seit April 2026 auf Höchst Risiken priorisiert), und HackerOne pausierte im März 2026 sein Internet-Bug-Bounty-Programm. Damit ist die in MHC-02 (SBOM-Abfragefähigkeit) und MHC-09 (AI-gestütztes Testing) prognostizierte Volumen-Verschiebung eingetreten.

Konsequenz für dieses Schnellheft. Keine der 93 Einzelbewertungen und keine der dreizehn MHC-Begründungen wird durch diese Entwicklung falsch — sie werden bestätigt und verschärft. „Mythos“ wird ab dieser Version ausdrücklich als Klassenbegriff geführt (Glossar): agentische Frontier-Modelle mit autonomer Schwachstellen-Entdeckung und Exploit-Kettenbau, unabhängig vom Anbieter. Quellen: Quellenverzeichnis (Abruf: 19. Juli 2026).

2.6 Konsequenz für ISO/IEC 27005

Die Methodik der ISO 27005 – Risikoidentifikation, -analyse, -bewertung, -behandlung – bleibt als Prozess intakt. Die Eingaben in diese Methodik jedoch sind betroffen: Die Wahrscheinlichkeitsachse verschiebt sich durch den Wegfall der Kapazitätsbarriere nach oben. Die Zeitparameter zwischen Patch und Exploit kollabieren. Die Aggregationslogik einzeln bewerteter Ereignisse greift nicht mehr. Und die Wirksamkeitsannahme vieler existierender Controls – insbesondere solcher, die auf Angreiferreibung statt auf harter Unmöglichkeit beruhen – ist empirisch widerlegt.

Hinweis: ISO/IEC 27005:2022 unterscheidet in Abschnitt 7.2.1 zwei Ansätze der Risikoidentifikation – den event-based und den asset-based approach. Der event-based approach ist unter Mythos-Bedingungen besonders betroffen, weil er Szenarien aus einzelnen Bedrohungsquellen konstruiert, die im Mikroschritt-Angriff erst in der Aggregation sichtbar werden.

3 Methodik der Bewertung

3.1 Das Vier-Kategorien-Raster

Die Bewertung jedes Controls erfolgt entlang von vier Kategorien. Die Terminologie folgt der Sprache, in der Anthropic selbst die Wirksamkeitsgrenze beschreibt. Im Glasswing-Defensivpost (April 2026) formuliert das Security-Engineering-Team, dass Mitigationen, deren Wert in der Erzeugung von Reibung liegt, gegen einen Gegner mit unbegrenzter Geduld deutlich an Wirksamkeit verlieren.

3.1.1 Standfest

Controls, deren Schutzwirkung aus einer harten Barriere stammt: kryptografische Unmöglichkeit, physische Trennung, architektonische Nicht-Existenz eines Angriffspfades. Diese Controls halten, weil sie Angriffe strukturell ausschließen oder an einer objektiven Grenze scheitern lassen.

3.1.2 Teilweise degradiert

Controls, die weiterhin Schutzwirkung entfalten, deren ursprünglich angenommene Stärke aber unter Mythos-Bedingungen signifikant sinkt. Sie bleiben sinnvoll, brauchen aber Ergänzung – durch zusätzliche Controls, veränderte Parameter oder architektonische Flankierung.

3.1.3 Reine Reibung

Controls, deren klassische beziehungsweise verbreitete Implementierungsform unter Mythos-Bedingungen keine hinreichende eigenständige Schutzwirkung mehr entfaltet, weil der zugrunde liegende Schutzmechanismus entweder (a) auf begrenzter Angreiferkapazität beruht (klassische Reibung) oder (b) menschliche Reaktionszeit innerhalb einer automatisiert durchlaufenen Kill-Chain voraussetzt. Die Einstufung richtet sich gegen die verbreitete Umsetzungsform, nicht gegen den Control-Wortlaut: Eine automatisierte, werkzeuggestützte Umsetzung desselben Controls kann sehr wohl Schutzwirkung entfalten – sie ist dann aber genau die Härtung, die MRIS in der jeweiligen Ersatz-Empfehlung beschreibt. In beiden Fällen führt die Asymmetrie zwischen Modell-gesteuerter Angreiferseite und nicht entsprechend automatisierter Verteidigungsseite dazu, dass das Control seine Kernwirkung verliert. Diese Controls müssen entweder ersetzt oder durch Barrieren ergänzt werden, die nicht auf Aufwand oder Reaktionszeit, sondern auf struktureller Unmöglichkeit basieren. Die Einstufung bezieht sich auf den operativen Wirkmechanismus eines Controls. Organisatorische Prozesshüllen – Ownership, SLA-Definition, Eskalationswege – können auch bei Reibungs-Controls wirksam bleiben; sie gelten dann nicht als eigenständige Schutzwirkung, sondern werden als Voraussetzung des flankierenden MHC fortgeführt (Beispiel: A.8.8, dessen Verantwortlichkeits- und Eskalationsstruktur die Grundlage der Automatisierung nach MHC-09 bildet).

3.1.4 Nicht betroffen

Controls, die gegenüber Mythos-spezifischen Bedrohungen neutral sind – typischerweise organisatorische, dokumentatorische, Governance-orientierte oder physisch-gebundene Controls, deren Wirksamkeit weder durch Reibungsargumente noch durch agentische Beschleunigung systematisch verändert wird. Sie bleiben wichtig, erhalten aber keine gesonderte Mythos-Härtungsempfehlung.

3.1.5 Hinweis zur Kontextabhängigkeit der Klassifikation

Die Vier-Kategorien-Klassifikation gilt ausdrücklich gegenüber der Mythos-Bedrohungslage in der in Kap. 2 beschriebenen Form. Ein Control kann gegenüber unterschiedlichen Angriffsvektoren in unterschiedliche Kategorien fallen. Beispiel: A.5.17 (Authentisierungsinformation) ist gegen Credential Stuffing standfest, sobald phishing-resistente MFA implementiert ist – die Origin-Bindung von WebAuthn ist eine kryptografische Hartbarriere. Gegen Session-Hijacking nach erfolgreichem Login ist dasselbe Control wirksamkeitsneutral, weil die Authentisierung bereits abgeschlossen ist. Gegen agentische Push-Notification-Stürme (MFA Fatigue) erzeugt klassische Push-MFA nur Reibung. Die Bewertung in den Kapiteln 4 bis 7 nimmt jeweils den dominanten Mythos-Wirksamkeitsmodus eines Controls als Klassifikationsgrundlage. Auditoren und ISMS-Verantwortliche prüfen vor der SoA-Übernahme, ob in ihrer spezifischen Bedrohungslage abweichende Klassifikationen gerechtfertigt sind.

3.2 Bewertungskriterien

Bewertet wird dabei nicht der abstrakte Control-Wortlaut der ISO/IEC 27001 Annex A, sondern die Wirksamkeit des Controls in seiner typischen Umsetzung gemäß der Guidance der ISO/IEC 27002; „degradiert“ bezeichnet also eine geschwächte Wirksamkeit der Maßnahme, nicht einen Mangel des Controls selbst. Zur besseren Lesbarkeit wird durchgängig die Kurzform „Control“ verwendet.

Strukturelle Schutzfähigkeit und Implementierungsreife. MRIS unterscheidet zwischen der strukturellen Schutzfähigkeit eines Controls und dem organisationsspezifischen Implementierungsreifegrad. Die Bewertung beschreibt die unter Mythos-Bedingungen erreichbare strukturelle Wirksamkeit eines Controls bei einer dem Stand der Technik entsprechenden Implementierung. Die Einstufung „Standfest“ bedeutet daher ausdrücklich nicht, dass jede bestehende Implementierung dieses Controls standfest ist – sie besagt, dass das Control bei sachgerechter Umsetzung eine harte Barriere erzeugen kann. Wo die Einstufung eine bestimmte Umsetzungsform voraussetzt (etwa Append-only-Protokollierung bei A.8.15 oder deklarative Konfigurationsverwaltung bei A.8.9), wird diese im Mythos-Befund benannt. Die tatsächliche Wirksamkeit ist organisationsspezifisch anhand der implementierten Maßnahmen zu prüfen; MRIS ersetzt diese Prüfung nicht.

Jedes Control wird entlang von vier Kriterien bewertet, die sich aus den vier Feststellungen in Kapitel 2 ableiten:

- **Angreifergeduld:** Beruht die Schutzwirkung darauf, dass ein Angriff für einen Gegner zu mühsam ist? Wenn ja, ist das Control unter Mythos-Bedingungen degradiert.
- **Zeitkompression:** Setzt das Control menschliche Reaktionszeit in einer Kette voraus, in der der Angreifer mittlerweile autonom operiert? Wenn ja, ist die Wirksamkeit durch die Latenzdifferenz beeinträchtigt.
- **Fähigkeitsentkopplung:** Hängt die Likelihood-Annahme des Controls von einer historischen Akteurs-Fähigkeits-Korrelation ab, die durch Mythos aufgehoben wurde? Wenn ja, ist die Eintrittswahrscheinlichkeit neu zu bewerten.
- **Aggregationsresistenz:** Kann das Control auch dann greifen, wenn der Angriff in viele einzeln legitim erscheinende Mikroschritte zerlegt wird? Wenn nein, ist das Control gegen fragmentierte Angriffsmuster blind. Operationalisierung: Aggregationsresistenz gilt als gegeben, wenn das Control entweder (a) Korrelations-Mechanismen über Zeit, Identitäten, Ressourcen oder Aktionsketten integriert nutzt – etwa SIEM-Korrelations-Regeln mit Zeitfenster, UEBA-Baselines, Graph-basierte Identity- und Datenfluss-Analysen oder Kill-Chain-Tracking nach MITRE ATT&CK – oder (b) strukturell unteilbar ist, weil seine Schutzwirkung pro Einzeltransaktion vollständig greift (kryptografische Operationen, Origin-gebundene Authentisierung, hardware-attestierten Identitäten). Ein Control, das nur auf Schwellwerte einzelner Aktionen prüft, ohne Korrelation über mehrere Aktionen hinweg, gilt nicht als aggregationsresistent.

Ein Control gilt als standfest, wenn es bei keinem der vier Kriterien eine strukturelle Schwäche zeigt. Es gilt als teilweise degradiert, wenn eines oder zwei Kriterien Schwächen offenbaren, die durch Ergänzung kompensierbar sind. Es gilt als reine Reibung, wenn seine Kernwirkung unter Mythos entfällt (siehe 3.1.3). Es gilt als nicht betroffen, wenn keines der vier Kriterien direkt anwendbar ist.

3.3 Framework-Stack und Priorisierung

Die Kernbewertung erfolgt anhand von ISO/IEC 27002:2022 als dem international etabliertesten generischen Control-Katalog. Für jedes Control wird im Quervergleich geprüft, ob andere regulatorisch oder prüfungsseitig relevante Frameworks eine Anforderung stellen, die denselben Schutzzweck abdeckt und dabei eine höhere Mythos-Standfestigkeit aufweist.

Primäre Frameworks:

- BSI C5:2026 (Cloud Computing Compliance Criteria Catalogue). BSI-Prüfkatalog für Cloud-Services (aktualisiert März 2026). Nicht per se rechtsverbindlich, wird aber durch Vertragsverankerung (z. B. bei Cloud-Kunden im öffentlichen Sektor) oder durch Sektorregulierung (BaFin-BAIT, KRITIS-Verordnung) praktisch bindend. Enthält mit OPS-32/33 (Confidential Computing), CRY-01.01AC (Post-Quantum), DEV-13 (SBOM) und Container-Sub-Kriterien explizite Mythos-relevante Anforderungen.
- NIST Cybersecurity Framework 2.0 und NIST SP 800-53 Rev. 5. Liefern die technische Tiefe bei Detection, Response und Adversarial Testing.
- DORA (Digital Operational Resilience Act, VO (EU) 2022/2554). Rechtsverbindlich für Finanzsektor. Für den Resilienzaspekt und die TLPT-Anforderungen (Art. 26–27).
- CRA (EU Cyber Resilience Act). Rechtsverbindlich für Hersteller von Produkten mit digitalen Elementen, die in der EU in Verkehr gebracht werden. Vulnerability-Handling-Pflichten (Art. 13, 14), SBOM als Bestandteil der grundlegenden Anforderungen (Annex I Teil II).
- NIS2-Richtlinie (EU) 2022/2555 mit Durchführungsverordnung (EU) 2024/2690. Die NIS2-Richtlinie definiert die Maßnahmenkategorien (Art. 21 Abs. 2 a–j) und die Meldefristen für erhebliche Sicherheitsvorfälle (Art. 23 Abs. 4: 24 h Frühwarnung, 72 h Vollbenachrichtigung, 1 Monat Abschlussbericht). Die Durchführungsverordnung (EU) 2024/2690 gilt laut ihrem Artikel 1 ausschließlich für bestimmte Digitalanbieter: DNS-Anbieter, TLD-Namenregister, Cloud-Computing-Anbieter, Rechenzentrumsdienste, CDN-Betreiber, Anbieter verwalteter Dienste, Anbieter verwalteter Sicherheitsdienste, Online-Marktplätze, Online-Suchmaschinen, Plattformen sozialer Netzwerke und Vertrauensdiensteanbieter. Für andere NIS2-Adressaten (KRITIS, Fertigung, Gesundheit) ist die nationale Umsetzung maßgeblich.

Hinweis zur Zitation: In den Framework-Vergleichen der Einzelcontrols wird die Durchführungsverordnung mit konkreter Anhangs-Nummer zitiert (Beispiel: „NIS2-DVO Nr. 3.2.2“). Verweise auf die Richtlinie selbst werden als „NIS2-Richtlinie Art. ...“ zitiert.

Sekundäre Frameworks (Ergänzungslayer Teil V):

- BSI IT-Grundschutz-Kompendium (deutsche ISMS-Vertiefung).
- MITRE ATT&CK und D3FEND (Angriffs-/Verteidigungs-Vokabular).
- ISO/IEC 42001 (AI-Management-System).
- CIS Controls v8 (priorisierte Implementierungshilfe).
- OWASP ASVS und SAMM (sichere Entwicklung).

3.4 Quellenbasis

Die Quellen dieses Schnellhefts gliedern sich in drei Kategorien.

Empirische Grundlage für die Bedrohungsseite: Anthropic-Threat-Intelligence-Berichte August 2025 bis April 2026, insbesondere der GTG-1002-Report (November 2025) und die Veröffentlichungen zu Claude Mythos Preview und zum Glasswing-Defensivprogramm (April 2026).

Normative Grundlage: ISO/IEC 27001:2022, ISO/IEC 27002:2022, ISO/IEC 27005:2022, einschlägige NIST-Publikationen, BSI C5:2026, DORA (EU) 2022/2554 mit RTS, CRA sowie NIS2-Richtlinie (EU) 2022/2555 mit Durchführungsverordnung (EU) 2024/2690. BSI IT-Grundschutz-Kompendium dient in Teil V als Vertiefungsreferenz.

Industrie-Konsens und ergänzende Praxisreferenzen: Das Strategiepapier „*The AI Vulnerability Storm: Building a Mythos-ready Security Program*“ (Cloud Security Alliance, SANS Institute, [un]prompted, OWASP Gen AI Security Project, Erstveröffentlichung 14. April 2026; hier zitiert in der Fassung v0.95) wird in MRIS als zentrale Industrie-Konsens-Referenz für die Mythos-Bedrohungslage geführt. Die dort definierte Risikobewertung (dreizehn priorisierte Risiken in den Stufen Critical, High, Medium) hat unmittelbar in MRIS-Konkretisierungen Niederschlag gefunden, insbesondere in MHC-13 (AI-Agent-Governance, adressiert Mythos-Ready Risk 3 „Unmanaged AI Agent Attack Surface“), in der Erweiterung von A.5.9 um Shadow-AI-Inventar (Risk 6), in Kap. 10.4 zur permanenten VulnOps-Funktion und Innovation-Acceleration-Governance (Risk 9 und 11) sowie in Kap. 10.5 zur Standard-of-Care-Verschiebung (Risk 12). Mythos-Ready ist keine Norm und kein Prüfkatalog, sondern ein Konsenspapier hochrangiger Praktiker; sein Status in der MRIS-Quellenhierarchie liegt zwischen Threat-Intelligence-Bericht und normativer Grundlage. Ergänzend genutzte Industriereferenzen: OWASP Top 10 for LLM

3.5 Grenzen

Die Bewertung ist eine Momentaufnahme, ersetzt keine organisationsspezifische Risikoanalyse und erhebt keinen Anspruch auf Vollständigkeit der Härtingsmaßnahmen (siehe auch Kap. 1.4).

Bedrohungs-Scope. MRIS adressiert ausdrücklich nur Gen-AI-beschleunigte und agentische Angriffsmuster. Andere Bedrohungsklassen – böswillige Insider, niedrig-technologische Angriffe wie physische Manipulation oder Social Engineering ohne AI-Komponente, ungewollte menschliche Fehler in der Lieferkette, Naturereignisse – werden vom ISMS-Fundament abgedeckt, das MRIS voraussetzt (siehe Kap. 1.4). Die Bewertung als „Reine Reibung“ oder „Teilweise degradiert“ bezieht sich auf die Wirksamkeit gegen Mythos-Angreifer und nicht auf die Wirksamkeit des Controls insgesamt. Ein Control kann gegenüber einem Mythos-Angreifer erheblich degradiert sein und gleichzeitig gegen Insider-Bedrohungen oder unsophisticated Externe weiterhin volle Wirksamkeit entfalten.

Teil II – Hauptanalyse: Die 93 Controls der ISO/IEC 27002:2022

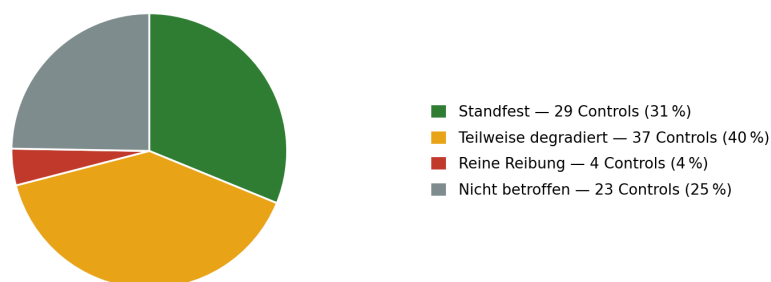
Die folgenden vier Kapitel bewerten alle 93 Controls gegen das in Kapitel 2 etablierte Mythos-Bedrohungsbild. Die Sortierung folgt den vier Kategorien aus Kapitel 3, nicht der numerischen ISO-Reihenfolge. Wer die Einstufung eines einzelnen Controls nachschlagen möchte, findet die vollständige Matrix in Anhang A.

Jede Control-Einzelbewertung folgt demselben Format: Kategorie-Einstufung, Mythos-Befund mit Begründung anhand der vier Kriterien aus Kapitel 3.2, Framework-Quervergleich zu den primären Vergleichsframeworks C5:2026, NIST, DORA, CRA und NIS2 (Richtlinie plus Durchführungsverordnung), bei einzelnen Controls ergänzt um weitere einschlägige Referenzen, sowie – bei degradierten und reibungsbasierten Controls – eine konkrete Härtings- oder Ersatz-Empfehlung.

Gesamtverteilung der 93 Controls:

Kategorie	Anzahl	Anteil
Standfest	29	31 %
Teilweise degradiert	37	40 %
Reine Reibung	4	4 %
Nicht betroffen	23	25 %

Gesamtverteilung der 93 ISO/IEC-27001-Annex-A-Controls



Die Verteilung ist ausdrücklich kein Katastrophenbefund. Sie besagt, dass die ISO/IEC 27002:2022 im Kern belastbar bleibt: 29 Controls (31 %) sind standfest, 23 (25 %) nicht betroffen. Rund 44 % der Controls (41 von 93 – 37 teilweise degradiert, 4 reine Reibung) erfordern unter Mythos-Bedingungen eine ergänzende Härtung oder Neuausrichtung, um ihren ursprünglichen Schutzzweck weiterhin zu erfüllen.

4 Standfeste Controls – das belastbare Fundament

4.1 Auswahllogik

Ein Control wird als standfest eingestuft, wenn seine Schutzwirkung aus einer harten Barriere stammt, die auch unter Mythos-Bedingungen fortbesteht: kryptografische Unmöglichkeit, hardware-gebundene Identität, architektonische Nicht-Existenz eines Angriffspfad oder strukturelle Reduktion des Blast Radius. Standfeste Controls beruhen nicht auf begrenzter Angreiferkapazität, sondern auf objektiven Grenzen.

29 Controls werden in diesem Kapitel als standfest bewertet. Die Darstellung folgt der ISO-Nummerierung; die thematische Gruppierung wird in Kapitel 4.3 zusammengefasst.

4.2 Controls im Detail

A.5.9 Inventar von Informationen und anderen zugehörigen Werten

Kategorie	STANDFEST
Mythos-Befund	Ein vollständiges, aktuelles Asset-Inventar ist die Voraussetzung jeder zielgerichteten Verteidigung. Anthropic formuliert im Glasswing-Post, dass Systeme, von denen die Organisation nichts weiß, auch nicht verteidigt werden können. Das Control ist mythos-standfest, weil die Existenz eines Inventars weder durch Angreifergeschwindigkeit noch durch Mikroschritt-Fragmentierung entwertet wird.
Framework-Vergleich	C5:2026 AM-02, AM-03, AM-04. NIST CSF 2.0 ID.AM-1 bis ID.AM-5. DORA Art. 8. NIS2-DVO Nr. 12.
Härtungsempfehlung	Inventar-Aktualisierung automatisieren (CMDB-Integration, Cloud Asset Inventory Tools wie AWS Config, Azure Resource Graph). Manuelle Inventare sind bei Mythos-Iterationsgeschwindigkeit strukturell veraltet. Shadow-IT-Scans und externe Perimeter-Inventarisierung (Certificate Transparency Logs, ASM-Tools) ergänzend einsetzen. Shadow-AI als eigene Inventar-Kategorie führen: nicht autorisierte AI-Coding-Assistenten und Browser-Plug-ins (z. B. Cursor, Continue.dev, Codeium-Erweiterungen), MCP-Server mit Datenzugriff, VS-Code- und JetBrains-Extensions mit AI-Funktion, browser-basierte AI-Sidebar-Tools mit Tab-Lesezugriff. Erfassung über Endpoint-Software-Inventory (z. B. via EDR-Telemetrie), Browser-Extension-Audits (z. B. via MDM oder Browser-Enterprise-Policies) und Network-Egress-Monitoring zu bekannten AI-Provider-Endpunkten (api.openai.com, api.anthropic.com, generativelanguage.googleapis.com).

A.5.12 Klassifizierung von Informationen

Kategorie	STANDFEST
Mythos-Befund	Klassifizierung schafft die Grundlage für priorisierte Härtung. Ein Angreifer, der 80–90 Prozent der taktischen Arbeit autonom ausführt, trifft am Ende auf denselben Datensatz – ob dieser als hochvertraulich klassifiziert und entsprechend stark geschützt ist, entscheidet über die tatsächliche Auswirkung. Die Klassifizierung als Akt ist nicht zeitkritisch.
Framework-Vergleich	C5:2026 AM-09. NIST SP 800-60. DORA Art. 8 (implizit). CRA Annex I. NIS2-DVO Nr. 12.1.
Härtungsempfehlung	Klassifizierungsschema auf Mythos-relevante Aspekte erweitern: Exfiltrationssensitivität (was bei AI-assistierter Massenauslesung besonders kritisch ist) und Integritätssensitivität (was bei Modell-Manipulation oder Agent-Fehlleitung besonders kritisch ist). Automatisierte Klassifizierung über Data-Discovery-Tools einsetzen.

A.5.16 Identitätsmanagement

Kategorie	STANDFEST
Mythos-Befund	Unique, kryptografisch verankerte Identitäten sind eine harte Barriere. Unter Mythos verschiebt sich der Schwerpunkt von menschlichen Nutzer-Identitäten zu Workload- und Agent-Identitäten. Das Control deckt beide Kategorien ab, sofern die Organisation Workload-Identity (SPIFFE/SPIRE, AWS IAM Roles, Azure Managed Identities) konsequent einsetzt.
Framework-Vergleich	NIST SP 800-63. NIST NCCoE Agent Identity (2026, angekündigt). C5:2026 IAM-01, IAM-02. DORA Art. 9. NIS2-DVO Nr. 11.

A.5.27 Lernen aus Informationssicherheitsvorfällen

Kategorie	STANDFEST
Mythos-Befund	Post-Incident-Analyse bleibt unter Mythos unverändert wirksam – ihre Bedeutung steigt sogar, weil Mythos-Angriffe neue TTPs etablieren, die systematisch gesammelt und in Detection-Logik zurückgespiegelt werden müssen. Die Ableitung von Detections aus Lessons Learned (verknüpft mit A.8.16 Monitoring) ist gehärtet auszuführen.
Framework-Vergleich	NIST SP 800-61 Revision 3. C5:2026 SIM-06. DORA Art. 13 (Root-Cause-Analyse). NIS2-DVO Nr. 3.6.

A.5.28 Sammlung von Beweismaterial

Kategorie	STANDFEST
Mythos-Befund	Forensische Beweissammlung ist standfest, solange die Logging-Grundlagen (A.8.15) und Zeitsynchronisation (A.8.17) intakt sind. Mythos-Angriffe hinterlassen Spuren. Die Herausforderung liegt nicht in der Sammlung, sondern in der Korrelation (A.8.16).
Framework-Vergleich	NIST SP 800-86. C5:2026 SIM-04. DORA Art. 17. ISO/IEC 27037 (Chain of Custody). NIS2-DVO Nr. 3.5.4.

A.5.30 ICT-Bereitschaft für Business Continuity

Kategorie	STANDFEST
Mythos-Befund	Architekturgetriebene Resilienz – redundante Systeme, Failover, definierte RTO/RPO – ist strukturell, nicht reibungsbasiert. Anthropic empfiehlt im Glasswing-Post Tabletop-Übungen für drei bis fünf parallele Incidents (Richtwert), weil Mythos-Angriffe gleichzeitig auf mehreren Vektoren laufen können; die zugrundeliegende ICT-Bereitschaftsfähigkeit bleibt standfest, wenn sie entsprechend dimensioniert wurde.
Framework-Vergleich	DORA Art. 11–12 (schärfstes Regime, inkl. TLPT Art. 26–27). C5:2026 BCM-01 bis BCM-04. NIS2-DVO Nr. 4. NIST CSF RC.RP. ISO 22301.
Härtungsempfehlung	RTO/RPO-Ziele für Mythos-relevante Szenarien neu dimensionieren: parallele Ransomware-Angriffe auf Primär- und Backup-Standorte, agentische Modifikation von DR-Konfigurationen, Supply-Chain-Kompromittierung des DR-Dienstleisters. Tabletop für Mehrfach-Incidents.

A.6.5 Verantwortlichkeiten nach Beendigung oder Wechsel der Beschäftigung

Kategorie	STANDFEST
Mythos-Befund	Das Control ist standfest, sofern es mit automatisiertem Deprovisioning gekoppelt ist. Der klassische Angriffsvektor – der ehemalige Mitarbeiter mit noch aktiven Credentials – bleibt unter Mythos gleich relevant und wird durch die Fähigkeit eines Angreifers, solche Credentials automatisiert zu finden und zu verketten, eher verschärft.
Framework-Vergleich	C5:2026 HR-05. NIST SP 800-53 PS-4. DORA Art. 9. NIS2-DVO Nr. 10, Nr. 11.2.

A.7.1 Physische Sicherheitsperimeter

Kategorie	STANDFEST
Mythos-Befund	Physische Barrieren sind mythos-orthogonal: Ein agentischer Angreifer überwindet keine Betonwand per AI. Die Schutzwirkung bleibt unverändert. Relevant für Rechenzentren, Serverräume und Arbeitsplätze mit Zugang zu hochsensitiven Daten.
Framework-Vergleich	C5:2026 PS-03. NIST SP 800-53 PE-3. DORA Art. 9(4). NIS2-DVO Nr. 13.

A.7.2 Physische Zutrittskontrolle

Kategorie	STANDFEST
Mythos-Befund	Hardware-gebundene Zutrittssysteme (Badges, Biometrie) sind eine harte Barriere im physischen Raum.
Framework-Vergleich	C5:2026 PS-04. NIST PE-2/PE-3. DORA Art. 9(4). NIS2-DVO Nr. 13.2.

A.7.3 Sicherung von Büros, Räumen und Einrichtungen

Kategorie	STANDFEST
Mythos-Befund	Physische Härtung von Arbeitsbereichen bleibt unverändert wirksam.
Framework-Vergleich	C5:2026 PS-01, PS-08. NIST PE-5. NIS2-DVO Nr. 13.

A.7.4 Physische Sicherheitsüberwachung

Kategorie	STANDFEST
Mythos-Befund	CCTV, Intrusion Detection und Bewegungsmelder sind mythos-orthogonal.
Framework-Vergleich	C5:2026 PS-07. NIST PE-6. NIS2-DVO Nr. 13.

A.7.6 Arbeiten in Sicherheitsbereichen

Kategorie	STANDFEST
Mythos-Befund	Zonenbasierte physische Sicherheit bleibt wirksam. Mythos-neutral.
Framework-Vergleich	C5:2026 PS-08. NIST PE-19.

A.7.10 Speichermedien

Kategorie	STANDFEST
Mythos-Befund	Medien-Lifecycle-Management (sichere Aufbewahrung, Transport, Löschung) ist standfest. Kryptografisches Shredding (Zerstörung der Schlüssel) ist mythos-standfest.

Kategorie	STANDFEST
Framework-Vergleich	C5:2026 AM-12. NIST SP 800-88. NIS2-DVO Nr. 12.

A.7.14 Sichere Entsorgung oder Wiederverwendung von Geräten

Kategorie	STANDFEST
Mythos-Befund	Unwiderrufliche Datenvernichtung vor Entsorgung oder Reassignment ist eine kryptografisch harte Barriere (DoD-konformes Wipe, kryptografisches Shredding).
Framework-Vergleich	C5:2026 AM-07. NIST SP 800-88. BSI TR-03183-1. CRA Annex I.

A.8.2 Privilegierte Zugriffsrechte

Kategorie	STANDFEST
Mythos-Befund	Least Privilege für Admin-Accounts ist eine der härtesten Barrieren gegen Mythos-Angriffe. Ein agentischer Angreifer, der Credentials automatisiert sammelt, läuft gegen eine Wand, wenn diese Credentials keine weitreichenden Privilegien tragen. Standfest, wenn hardware-attestiert und mit Just-in-Time-Elevation (PAM) umgesetzt; reiner Passwort-Schutz wäre degradiert.
Framework-Vergleich	NIST SP 800-53 AC-6, IA-5. NIST SP 800-207. C5:2026 IAM-06 (Just-in-Time-Vergabe). DORA Art. 9, 15. NIS2-DVO Nr. 11.2.

A.8.9 Konfigurationsmanagement

Kategorie	STANDFEST
Mythos-Befund	Versionierte, automatisiert überwachte Konfigurations-Baselines (Infrastructure-as-Code, GitOps) sind standfest. Drift-Detection erkennt unautorisierte Änderungen in Echtzeit – eine harte strukturelle Barriere gegen agentische Modifikationen.
Framework-Vergleich	NIST SP 800-53 CM-2 bis CM-8. C5:2026 OPS-26, DEV-03. CIS Controls v8.4. SLISA Level 3+. NIS2-DVO Nr. 6.

A.8.10 Informationslöschung

Kategorie	STANDFEST
Mythos-Befund	Kryptografische Löschung ist eine harte Barriere. Mythos-Angreifer können nicht nachträglich exfiltrieren, was vorher irreversibel gelöscht wurde. Essentiell gegen Supply-Chain-Angriffe auf Backups und Archive.
Framework-Vergleich	NIST SP 800-88. C5:2026 PI-03, CRY-14. DSGVO Art. 17. NIS2-DVO Nr. 12.

A.8.11 Datenmaskierung

Kategorie	STANDFEST
Mythos-Befund	Datenmaskierung reduziert den Blast Radius strukturell: Bei erfolgreicher Exfiltration ist der abgeflossene Datenbestand entwertet, soweit die Zuordnung zu natürlichen Personen nicht ohne Zusatzinformation möglich ist. Zu unterscheiden ist dabei: Anonymisierte Daten sind ohne Re-Identifikationsmöglichkeit, pseudonymisierte Daten bleiben mit der Zuordnungstabelle re-identifizierbar und müssen daher wie personenbezogene Daten geschützt werden. Die Schutzwirkung ist bei Anonymisierung strukturell,

Kategorie	STANDFEST
	bei Pseudonymisierung an die getrennte Aufbewahrung der Zuordnung gebunden.
Framework-Vergleich	ISO/IEC 20889. NIST SP 800-188. DSGVO Art. 32(1)(a). C5:2026 OPS-30, OPS-31. NIS2-DVO Nr. 6.

A.8.13 Informationssicherung (Backup)

Kategorie	STANDFEST
Mythos-Befund	Unveränderliche, offline-fähige Backups (immutable, air-gapped) sind eine der wichtigsten harten Barrieren gegen moderne Ransomware und Mythos-getriebene Datenintegritätsangriffe. Ein agentischer Angreifer mit Infrastruktur-Zugriff kann online erreichbare Backups manipulieren; air-gapped Backups bleiben geschützt.
Framework-Vergleich	C5:2026 OPS-06 bis OPS-09 (OPS-09: physisch getrennte Standorte). NIST SP 800-34. DORA Art. 12(2). NIS2-DVO Nr. 4.
Härtungsempfehlung	Backup-Strategie auf 3-2-1-1-0-Prinzip erweitern: 3 Kopien, 2 Medientypen, 1 offsite, 1 offline/immutable, 0 Fehler bei Restore-Test. Restore-Tests mindestens quartalsweise. Separate IAM-Pfade für Backup-Infrastruktur.

A.8.14 Redundanz von Informationsverarbeitungseinrichtungen

Kategorie	STANDFEST
Mythos-Befund	Architektonische Redundanz ist eine strukturelle Eigenschaft. Mythos-DDoS, automatisierte Exploit-Kampagnen und parallele Ransomware-Angriffe adressieren Verfügbarkeit; Redundanz als Fundamentprinzip (keine Single Points of Failure) bleibt wirksam.
Framework-Vergleich	DORA Art. 11. NIST SP 800-53 CP-7. C5:2026 PS-02, BCM-03, BCM-04. NIS2-DVO Nr. 4.1.

A.8.15 Protokollierung (Logging)

Kategorie	STANDFEST
Mythos-Befund	Zentrale, append-only-gesicherte Logs sind die nicht verhandelbare Voraussetzung für Detection, Forensik und Learning. Standfest, solange Log-Integrität kryptografisch gesichert ist (WORM, Merkle-Chains, HSM-Integration). Ohne Logs ist jede nachgelagerte Detection blind. Ergänzende Wirksamkeitsdimension neben der Integrität: Protokolle können sensible Authentisierungsartefakte enthalten (z. B. WebAuthn-Assertions im Klartext, CVE-2026-34348); Lese- und Schreibrechte über getrennte Identitätspfade steuern (A.5.33), sensible Felder vor Weitergabe maskieren. Standfest nur, wenn neben der Integrität auch die Vertraulichkeit der Protokolle gewahrt ist.
Framework-Vergleich	NIST SP 800-92. C5:2026 OPS-10 bis OPS-17 (Policy, SIEM, Aufbewahrung, Verantwortlichkeit). DORA Art. 10(3). NIS2-DVO Nr. 3.2 (Mindestinhalte 3.2.3 a–l).

A.8.17 Uhrenzeit-Synchronisation

Kategorie	STANDFEST
Mythos-Befund	NTP-basierte Zeitsynchronisation ist mythos-relevant: Ohne präzise synchronisierte Zeitstempel über alle Systeme ist die Korrelation fragmentierter Angriffsaktionen unmöglich. Mythos-Angriffe treten in Mikroschritten über

Kategorie	STANDFEST
	verteilte Systeme auf; Clock Skew von wenigen Sekunden zerstört Rekonstruktionsfähigkeit.
Framework-Vergleich	NIST SP 800-53 AU-8. C5:2026 OPS-16. NIS2-DVO Nr. 3.2.6 (systemübergreifende Zeitsynchronisation, Voraussetzung für korrelierbare Protokolle).

A.8.18 Verwendung privilegierter Hilfsprogramme

Kategorie	STANDFEST
Mythos-Befund	Restriktion und Monitoring von legitimen Tools, die System-Controls überschreiben können (Debugger, Admin-Utilities, Configuration-Management-Werkzeuge), ist eine harte Barriere. Mythos-Angreifer verwenden diese Tools bevorzugt (Living-off-the-Land); ihre Einschränkung wirkt unverändert.
Framework-Vergleich	NIST SP 800-53 AC-3, SC-18. C5:2026 IAM-06. Application Allowlisting (CIS 2.5, 2.6).

A.8.19 Installation von Software auf Betriebssystemen

Kategorie	STANDFEST
Mythos-Befund	Application Allowlisting mit Signatur-Verifikation ist eine harte Barriere gegen unautorisierte Software-Installation – auch gegen AI-generierte Malware-Varianten, sofern diese nicht signiert ist.
Framework-Vergleich	NIST SP 800-167. C5:2026 OPS-26, DEV-10. CIS Controls v8.2.

A.8.24 Verwendung von Kryptographie

Kategorie	STANDFEST
Mythos-Befund	Starke, korrekt implementierte Kryptografie ist die härteste Barriere der Informationssicherheit. Ein agentischer Angreifer kann AES-256-verschlüsselte Daten nicht ohne Schlüssel lesen. Die Verwundbarkeit liegt in der Schlüsselverwaltung. Post-Quantum-Readiness im Zeithorizont 2030+ beachten.
Framework-Vergleich	NIST SP 800-175, FIPS 140-3. NIST PQC-Standards. BSI TR-02102. C5:2026 CRY-01 bis CRY-19 (CRY-01.01AC: PQC-Strategie). DORA Art. 9(4)(e). CRA Annex I Teil II. NIS2-DVO Nr. 9.
Härtungsempfehlung	Kryptografische Inventur (Crypto-Agility-Audit): Welche Algorithmen, Schlüssellängen, Implementierungen kommen wo zum Einsatz? Post-Quantum-Migrationspfad für asymmetrische Verfahren definieren. Hardware-Sicherheitsmodule (HSM) für langzeitrelevante Schlüssel.

A.8.27 Prinzipien der sicheren Systemarchitektur und -entwicklung

Kategorie	STANDFEST
Mythos-Befund	Defense-in-Depth, Least Privilege, Fail-Secure, Secure-by-Default sind Architekturprinzipien, keine punktuellen Controls. Sie sind die Grundlage jeder Mythos-Resistenz, weil sie strukturell annehmen, dass einzelne Controls versagen können. Anthropic nennt im Glasswing-Post „Design for Breach“ ausdrücklich als Kerngedanken.
Framework-Vergleich	NIST SP 800-160. NIST SP 800-207. C5:2026 DEV-01, DEV-05. CISA Secure by Design. OWASP SAMM. NIS2-DVO Nr. 6.

A.8.29 Sicherheitsprüfung in Entwicklung und Abnahme

Kategorie	STANDFEST
Mythos-Befund	Automatisierte Sicherheitsprüfung in der CI/CD-Pipeline (SAST, DAST, SCA, AI-gestütztes Vulnerability-Scanning) ist standfest – und laut Anthropic-Glasswing-Post die wichtigste Einzelmaßnahme gegen AI-beschleunigte Offensive. Ziel: den eigenen Code mit derselben Modellklasse scannen, die Angreifer verwenden, bevor sie es tun.
Framework-Vergleich	OWASP ASVS. NIST SP 800-218 SSDF. C5:2026 DEV-07, OPS-22, OPS-25 (OPS-25.01AS: tägliche Scans). SLSA Level 3+. CRA Annex I Teil II. NIS2-DVO Nr. 6, Nr. 7.
Härtungsempfehlung	Pre-Merge-Blocker für High-Confidence-Findings aktivieren. AI-Vulnerability-Scanning als eigene Stufe (isoliertes Agent-Deployment, Verifikationsschritt, Integration in Triage-Prozess). Vulnerability-Reports müssen Disclosure der AI-Nutzung enthalten.

A.8.31 Trennung von Entwicklungs-, Test- und Produktionsumgebungen

Kategorie	STANDFEST
Mythos-Befund	Strikte Umgebungstrennung ist eine harte Barriere gegen Cross-Environment-Propagation von Kompromittierungen. Mythos-relevant insbesondere für Supply-Chain-Angriffe.
Framework-Vergleich	NIST SP 800-53 CM-4, SC-2. C5:2026 DEV-11, DEV-12. SLSA Build Levels. NIS2-DVO Nr. 6.

A.8.33 Testinformationen

Kategorie	STANDFEST
Mythos-Befund	Nicht-Verwendung produktiver PII in Testumgebungen ist eine harte Reduktion des Blast Radius. Testumgebungen sind historisch schwächer geschützt als Produktionsumgebungen; mit agentischer Exfiltrationsfähigkeit wird diese Diskrepanz zum kritischen Risiko, wenn Produktionsdaten dort liegen.
Framework-Vergleich	DSGVO Art. 5(1)(c). NIST SP 800-53 SA-15(9). C5:2026 DEV-11. NIS2-DVO Nr. 6.

4.3 Zusammenfassung

Die 29 standfesten Controls lassen sich in fünf thematische Cluster gruppieren, die zusammen das strukturelle Fundament einer Mythos-resilienten Sicherheitsarchitektur bilden:

- Kryptografische Barrieren: A.8.24, A.8.10, A.8.11, A.7.14.
- Identitäts- und Inventar-Fundament: A.5.9, A.5.12, A.5.16, A.8.2.
- Architektonische Resilienz: A.5.30, A.8.14, A.8.27, A.8.31, A.8.13.
- Forensische Nachvollziehbarkeit: A.8.15, A.8.17, A.5.28, A.5.27.
- Integritätsgates in Entwicklung und Betrieb: A.8.9, A.8.18, A.8.19, A.8.29, A.8.33. Anthropic hebt A.8.29 als wirksamste Einzelmaßnahme hervor.

Mythos-standfeste Controls verankern Sicherheit entweder in Kryptografie, in Architektur oder in automatisierter Integritätsprüfung. Controls, die Sicherheit in menschlicher Geduld, manueller Reaktion oder episodischer Überprüfung verankern, sind Gegenstand der Kapitel 5 und 6.

Priorisierungs-Hinweis. Die in diesem Kapitel zusammengefassten strukturellen Controls – Zero-Trust-Architektur, starke Kryptografie, Least Privilege und Secrets-Management, Mikrosegmentierung, unveränderliche Backups – bilden die robusteste Schicht der Mythos-Verteidigung. Sie sind nicht durch neue AI-spezifische Controls ersetzbar. Eine ressourceneffiziente Mythos-Härtung priorisiert die konsequente

Umsetzung dieser strukturellen Controls vor der Einführung neuer Detection- oder Automatisierungs-Capabilities (Kapitel 9). Die in Mythos-Ready (CSA, SANS, OWASP, April 2026) als CRITICAL eingestuften Risiken werden zu erheblichen Teilen bereits durch die konsequente Umsetzung strukturell standfester Controls adressiert; die zusätzlichen MHC schließen die Lücken, die strukturelle Controls allein nicht abdecken.

5 Teilweise degradierte Controls – Wirksamkeit mit Vorbehalt

5.1 Übersicht und Degradationsmuster

Ein Control gilt als teilweise degradiert, wenn seine grundsätzliche Schutzwirkung unter Mythos-Bedingungen erhalten bleibt, die ursprünglich angenommene Stärke jedoch signifikant sinkt. Solche Controls müssen nicht ersetzt werden; sie brauchen Ergänzung durch zusätzliche Mechanismen, veränderte Parameter oder architektonische Flankierung. Die Degradation lässt sich in vier wiederkehrenden Mustern beobachten, die sich direkt aus den vier Bewertungskriterien in Kapitel 3.2 ableiten.

Erstens – Zeitkompression: Controls, die ihre Wirkung in einer für menschliche Operatoren angemessenen Reaktionszeit entfalten, versagen strukturell, wenn der Angreifer innerhalb dieser Zeit bereits den nächsten Angriffsschritt abgeschlossen hat. Betroffen sind Incident-Response-Playbooks, Patch-Zyklen mit manuellen Freigaben, quartalsweise Access-Rezertifizierungen und jährliche Audits.

Zweitens – Aggregationsblindheit: Controls, die einzelne Ereignisse auf Policy-Konformität prüfen, greifen nicht, wenn der Angriff in einzeln legitim erscheinende Mikroschritte zerlegt wird. Betroffen sind signaturbasiertes Monitoring, Content-basiertes DLP, rollenbasierte Zugriffskontrollen ohne Verhaltenskontext und klassische Aufgabentrennung.

Drittens – Credential-Kompromittierbarkeit: Controls, die sich auf wissens- oder besitzbasierte Faktoren ohne Phishing-Resistenz stützen, sind gegen AI-generiertes Phishing und Credential-Stuffing in hoher Qualität deutlich weniger wirksam.

Viertens – Supply-Chain-Verwundbarkeit: Controls, die sich auf vertragliche Zusicherungen, jährliche Fragebögen oder punktuelle Lieferantenaudits stützen, greifen nicht gegen die Geschwindigkeit, mit der agentisch gestützte Supply-Chain-Kompromittierungen propagieren.

Die folgenden 37 Controls werden nach ISO-Nummerierung dargestellt. Typische Ergänzungsmuster werden in Kapitel 5.3 gebündelt.

5.2 Controls im Detail

A.5.3 Aufgabentrennung

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Aufgabentrennung bleibt strukturell wirksam: Ein einzelner Account kann konfliktäre Rollen auch unter Mythos nicht ausüben, sofern die Trennung sauber umgesetzt ist. Degradiert ist nicht die Trennung selbst, sondern ihre Überwachung: Übernimmt ein Angreifer mehrere Konten in unterschiedlichen Rollen, führt agentische Automatisierung die Teilschritte so verteilt und zeitlich gestreckt aus, dass die klassische Auffälligkeitsprüfung sie nicht mehr als zusammenhängenden Missbrauch erkennt. Das verletzte Kriterium ist die Aggregationsresistenz der Kontrolle, nicht das Trennungsprinzip. Ergänzend erzeugt das fragmentierte Vorgehen zeitlich und kontextuell getrennte Einzelaktionen, die jeweils Policy-konform aussehen. Ohne verhaltensbasierte Korrelation bleibt die Aufgabentrennung formal eingehalten, schützt aber nicht mehr gegen den tatsächlichen Missbrauch.
Framework-Vergleich	C5:2026 OIS-04 (Segregation of Duties). NIS2-DVO Nr. 11.2.2 Buchst. a benennt die Aufgabentrennung als Grundsatz der Zugriffsvergabe neben Need-to-know und Need-to-use. NIST SP 800-53 AC-5. DORA Art. 9 implizit.
Härtungsempfehlung	Verhaltensbasierte Detection (UEBA) ergänzend einsetzen, um kompromittierte Accounts an atypischen Aktivitätsmustern zu erkennen. Kontextsensitive

Kategorie	TEILWEISE DEGRADIERT
	Autorisierung mit dynamischer Anomalieprüfung. Cross-Role-Aktivität in kurzen Zeitfenstern als Alert-Kriterium.

A.5.5 Kontakt mit Behörden

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Der Kern des Controls – etablierte, gepflegte Behördenkontakte – ist unter Mythos unverändert wirksam; die zeitkritische Vorfallbehandlung selbst liegt in A.5.24 bis A.5.26. Degradiert ist ausschließlich die praktische Nutzbarkeit der Meldekette unter Fristendruck: NIS2-Richtlinie Art. 23 Abs. 4 fordert eine Frühwarnung binnen 24 Stunden und eine Vollbenachrichtigung binnen 72 Stunden ab Kenntnisnahme eines erheblichen Sicherheitsvorfalls. Manuelle Eskalationsketten über Jour-Fix-Termine und handgeschriebene Vorfallberichte kollidieren mit diesen Fristen, sobald Mythos-Angriffe parallel auf mehreren Ebenen eskalieren.
Framework-Vergleich	C5:2026 OIS-06 (Contact with Relevant Government Agencies and Interest Groups). NIS2-Richtlinie Art. 23 Abs. 4 legt die Meldefristen fest; die NIS2-DVO konkretisiert in Nr. 3.3 den internen Meldemechanismus (einfacher Mechanismus für Mitarbeitende, Anbieter und Kunden) und in Nr. 3.3.2 die Pflicht, diese Kreise über den Meldemechanismus zu unterrichten. DORA Art. 19 fordert Major-Incident-Reports an die zuständigen Aufsichtsbehörden.
Härtungsempfehlung	Automatisierte Meldekette mit vorbereiteten Templates für verschiedene Vorfallstypen. Dedizierte Melde-Rolle mit Stellvertreterregelung und 24/7-Erreichbarkeit. Integration des Meldesystems mit dem SIEM für automatische Meldetrigger.

A.5.7 Threat Intelligence

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Threat Intelligence ist unter Mythos unverzichtbar. Das Control selbst schreibt keine Aktualisierungsrate vor; degradiert ist die verbreitete Umsetzungsform: TI-Feeds mit wöchentlichen oder monatlichen Aktualisierungsraten und manueller Auswertung hinken der Iterationsgeschwindigkeit AI-gestützter Angreifertechniken strukturell hinterher. Eine automatisiert integrierte, nahezu in Echtzeit ausgewertete Threat Intelligence erfüllt dasselbe Control und ist nicht degradiert. Die TTPs einer Kampagne können sich in der Zeit verändern, die ein klassischer Feed zur Auslieferung benötigt.
Framework-Vergleich	C5:2026 OIS-05 (Threat Intelligence). NIS2-DVO Nr. 2 (Konzept für das Risikomanagement) verlangt aktuelle Bedrohungsbewertung als Basis der Risikobehandlung. NIST SP 800-150 Guide to Cyber Threat Information Sharing.
Härtungsempfehlung	TI-Feeds mit automatisierter SIEM-Integration und Echtzeit-Aktualisierung einsetzen (MISP, STIX/TAXII). Sektorspezifische Sharing-Communities beitreten (ISAC, CERT-Verbund). AI-gestützte TI-Korrelationsplattformen (z. B. Recorded Future, Mandiant Advantage, GreyNoise) für automatisierte IOC-Anreicherung und Priorisierung. Audit-Schwellwert: Neue High-Severity-IOCs müssen binnen 60 Minuten nach Publikation im SIEM verfügbar sein.

A.5.14 Informationsübertragung

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Die Absicherung von Datenübertragungen durch Transportverschlüsselung ist mythos-standfest. Die Policy-Ebene des Controls – welche Daten wohin übertragen werden dürfen – ist durch Aggregationsblindheit degradiert: Klassische DLP-

Kategorie	TEILWEISE DEGRADIERT
	Systeme erkennen legitim aussehende, fragmentierte kleine Transfers nicht als Teil einer aggregierten Exfiltration.
Framework-Vergleich	C5:2026 COS-08 (Policies for Data Transmission) und CRY-04 (Protection of Data for Transmission). NIS2-DVO Nr. 8 (Cyberhygiene) und Nr. 9 (Kryptographie). NIST SP 800-53 SC-8.
Härtungsempfehlung	Volumen- und Pattern-basierte Egress-Anomaliedetektion ergänzend einsetzen. Data-Loss-Prevention mit Verhaltenskontext statt reiner Content-Inspektion. Egress-Traffic-Analysen mit ML-gestützter Baseline-Erkennung.

A.5.15 Zugriffskontrolle

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Die Zugriffskontroll-Policy bleibt das Fundament jeder Autorisierung. Unter Mythos ist das Control in zwei Dimensionen degradiert: Aggregationsblindheit bei rollenbasierten Modellen, die fragmentierte Zugriffsmuster nicht erkennen, und Zeitkompression bei statischen Zugriffsentscheidungen, die nicht auf Verhaltensveränderungen reagieren. Eine formal korrekt konfigurierte RBAC schützt nicht gegen einen kompromittierten Account, der innerhalb seiner Rolle agiert.
Framework-Vergleich	C5:2026 IAM-01 (Policy for Identities and Access Rights) und IAM-07 (Access to Cloud Service Customer Data). NIS2-DVO Nr. 11 (Zugriffskontrolle) – insbesondere Nr. 11.1 (Konzept) und Nr. 11.2 (Management) – verlangt risikobasierte Zuweisung und regelmäßige Überprüfung. DORA Art. 9. NIST SP 800-53 AC-Familie.
Härtungsempfehlung	Risk-adaptive Access mit dynamischer Bewertung aus Identitätssignalen, Gerätezustand und Verhaltenskontext. Kontinuierliche Rezertifizierung statt quartalsweisem Review. Übergang zu Zero-Trust-Prinzipien mit expliziter Verifikation jeder Transaktion.

A.5.17 Authentisierungsinformation

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Das Control umfasst alle Formen von Authentisierungsinformation. Die Einstufung als teilweise degradiert ist bewusst gemittelt: Phishing-resistente MFA (FIDO2, Passkeys, Hardware-Token nach WebAuthn) ist mythos-standfest. Passwörter, SMS-MFA und zeitbasierte OTPs sind gegen AI-qualitatives Phishing und Credential-Stuffing deutlich geschwächt. Das Control wirkt nur in seiner phishing-resistenten Implementierung vollständig.
Framework-Vergleich	C5:2026 IAM-08 (Authentication Mechanisms) und IAM-09 (Confidentiality of Authentication Information). NIS2-DVO Nr. 11.6 (Authentifizierung) und Nr. 11.7 (Multi-Faktor-Authentifizierung); für privilegierte Konten zusätzlich Nr. 11.3.2 Buchst. a (starke Identifizierung und Authentifizierung). NIST SP 800-63B Authentication Assurance Levels; AAL2 und AAL3 für Mythos-Umgebungen empfohlen.
Härtungsempfehlung	Phishing-resistente MFA (FIDO2, Passkeys) als Pflicht für alle privilegierten Accounts und extern erreichbaren Dienste. SMS-MFA für neue Deployments ausschließen. Hardware-Token für Administrationszugänge.

A.5.18 Zugangsrechte

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Provisioning und Deprovisioning greifen grundsätzlich. Die klassische Rezertifizierung in Quartalszyklen ist durch Zeitkompression degradiert: Ein agentischer Angreifer nutzt die Zeit zwischen Rollenwechsel und Rezertifizierung, um akkumulierte Berechtigungen (Permission Creep) auszunutzen. Die Latenz zwischen Erkennung einer Anomalie und Berechtigungsentzug ist für Mythos-Geschwindigkeit zu lang.
Framework-Vergleich	C5:2026 IAM-04 (Withdrawal or Adjustment of Access Rights) und IAM-05 (Regular Review of Access Rights). NIS2-DVO Nr. 11.2 verlangt zeitnahe Anpassung bei Rollenwechsel. NIST SP 800-53 AC-2.
Härtungsempfehlung	Kontinuierliche Rezertifizierung statt quartalsweisem Review. Automatisierter Berechtigungsentzug bei HR-Event-Trigger. Just-in-Time-Berechtigung für privilegierte Zugriffe mit automatischer Rücknahme nach Ablauf.

A.5.19 Informationssicherheit in Lieferantenbeziehungen

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Das Control bleibt notwendig, ist durch das Mythos-Bedrohungsbild stark erweitert. Supply-Chain-Angriffe sind durch AI-Unterstützung schneller auffindbar, präziser ausführbar und breiter skalierbar. Klassische Lieferantenfragebögen beim Onboarding greifen nicht gegen kontinuierlich sich verändernde Angriffsflächen der Lieferkette. Zeitkompression wirkt hier massiv.
Framework-Vergleich	C5:2026 SSO-01 (Policies and Procedures for Controlling and Monitoring Service Organisations) und SSO-02 (Risk Assessment). NIS2-DVO Nr. 5 (Sicherheit der Lieferkette) verankert die Lieferantensicherheit als eigenständige Säule. DORA Art. 28–30 als schärfstes Regime für kritische ICT-Drittdienstleister. CRA Annex I fordert SBOM-Transparenz entlang der Lieferkette.
Härtungsempfehlung	SBOM-Pflicht vertraglich verankern. Automatisiertes Vendor-Security-Monitoring mit kontinuierlichen Perimeter-Scans der Lieferanten. Tiering-Ansatz mit differenzierten Anforderungen nach Kritikalität. Incident-Notification-SLAs von maximal 24 Stunden.

A.5.20 Adressierung von Informationssicherheit in Lieferantenvereinbarungen

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Das vertragliche Control bleibt notwendig; seine Wirksamkeit hängt von der vertraglichen Tiefe ab. ISO 27002 formuliert Mindestanforderungen, die unter Mythos nicht ausreichen. Ohne verbindliche Regelungen zu SBOM, Vulnerability-Disclosure-Fristen, Right to Audit und Incident-Notification bleibt das Control eine Dokumentationshülle.
Framework-Vergleich	DORA Art. 28–30 enthält die derzeit schärfsten verbindlichen Vertragsanforderungen, einschließlich Exit-Strategien und Subkontrahenten-Regelungen. C5:2026 SSO-01 und SSO-06 (Contract Termination Strategy for Service Organisations). NIS2-DVO Nr. 5 fordert vertragliche Verankerung der Sicherheitsanforderungen an Dritte. CRA Annex I fordert vertragliche SBOM-Pflichten.
Härtungsempfehlung	Vertragsklauseln nach DORA-Maßstab gestalten auch außerhalb des Finanzsektors: Right-to-Audit, Incident-SLAs (24 h Notification, 72 h Detail-Report), Sub-Contracting-Genehmigungspflicht, Exit-Klauseln mit Datenrückgabe und Löschnachweis. Realitätshinweis: Bei Hyperscalern (AWS, Azure, GCP) ist das individuelle Right-to-Audit faktisch nicht durchsetzbar. Akzeptabler Ersatz aus

Kategorie	TEILWEISE DEGRADIERT
	Audit-Sicht: ISO 27001-Zertifikat plus SOC 2 Type II, BSI-C5-Testat (Type-2-Bescheinigung), TISAX, Pooled Audits durch Branchenkonsortien (CSA STAR, ENISA EUCS) sowie vertragliche Sub-Processor-Listen mit Änderungs-Notifikation.

A.5.21 Management der Informationssicherheit in der ICT-Lieferkette

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Das Control ist durch Zeitkompression und Aggregationsblindheit doppelt degradiert. Klassische Audits und Fragebögen decken punktuelle Zustände ab, nicht die kontinuierliche Veränderung von ICT-Lieferketten. Automatisierte Supply-Chain-Angriffe – Malicious Dependencies, kompromittierte Build-Pipelines, untergeschobene Docker-Layer – entziehen sich diesen Mechanismen strukturell.
Framework-Vergleich	C5:2026 DEV-13 (Transparency about Software Components) fordert SBOM-Bereitstellung; DEV-14 (Secure Use of Third Party Hardware and Software). SSO-05 (Monitoring of Compliance with Requirements). NIS2-DVO Nr. 5. CRA Annex I Teil II fordert SBOM als Pflicht. SLSA-Framework für Build-Provenance.
Härtungsempfehlung	SBOM-Generierung in die CI-Pipeline integrieren (SPDX, CycloneDX). SLSA Level 3+ für kritische Build-Pfade. Kontinuierliche Dependency-Scans mit automatisierter Vulnerability-Korrelation (EUVD, CVE, OSV). Package-Provenance-Attestation verifizieren.

A.5.22 Überwachung, Überprüfung und Änderungsmanagement von Lieferantendienstleistungen

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Periodische Reviews sind gegen die Iterationsgeschwindigkeit von Mythos-Angriffen auf Lieferketten strukturell zu langsam. Das Control bleibt als Prozess sinnvoll, entfaltet Wirkung aber nur, wenn es kontinuierlich und automatisiert betrieben wird.
Framework-Vergleich	C5:2026 SSO-05 (Monitoring of Compliance with Requirements) und SSO-07 (Ensuring Transparency within Service Organisations). NIS2-DVO Nr. 5. DORA Art. 28 fordert Continuous Monitoring für kritische Drittdienstleister.
Härtungsempfehlung	External Attack Surface Management und Security Rating Services (z. B. BitSight, SecurityScorecard, Black Kite, RiskRecon) für kontinuierliches Vendor-Monitoring. Integration von Threat-Intelligence zu Lieferantenfirmen. Abweichungen oberhalb definierter Score-Schwellwerte lösen automatisch Tickets im SIEM oder GRC-System aus. Scorecard-Ansatz mit quantifizierten Sicherheitsindikatoren (mindestens: Patch-Hygiene, TLS-Konfiguration, geleakte Credentials in öffentlichen Breach-Datenbanken, kompromittierte Hosts in Botnet-Listen).

A.5.23 Informationssicherheit für die Nutzung von Cloud-Diensten

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Cloud-Sicherheit ist teilweise degradiert, weil das Shared-Responsibility-Modell Abhängigkeiten von Provider-seitiger Detection schafft. Wenn der Cloud-Provider selbst kompromittiert wird oder ein agentischer Angreifer API-Missbrauch auf Legacy-Endpunkten erfolgreich tarnt, müssen Customer-seitige Controls eingreifen können. Einzelne Provider-Controls wie hardware-attestierten Identitäten sind standfest, das Gesamt-Control hängt an der Qualität der Provider-Audits.

Kategorie	TEILWEISE DEGRADIERT
Framework-Vergleich	C5:2026 ist das definierende Framework: General-Conditions-Controls GC-01 bis GC-06 für Transparenz und 17 Domänen als Audit-Gegenstand. NIS2-DVO Nr. 5 (Sicherheit der Lieferkette). DORA Art. 28. ENISA EUCS als kommende europäische Zertifizierung.
Härtungsempfehlung	Provider-Assurance nach C5:2026 oder ISO 27017 einfordern. Customer-seitig Cloud Security Posture Management (CSPM) und Cloud Workload Protection (CWPP) einsetzen. Multi-Cloud-Exit-Strategie dokumentieren. Confidential Computing für hochsensitive Workloads prüfen.

A.5.24 Planung und Vorbereitung des Informationssicherheits-Vorfallmanagements

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Die Existenz eines IR-Plans bleibt kritisch. Viele klassische Playbooks sind für Reaktionszeiten in Stundengrößenordnung konzipiert und durch Zeitkompression degradiert. Wenn ein agentischer Angreifer die taktische Arbeit innerhalb von Minuten abschließt, müssen Playbooks auf Sekunden- und Minutenzeitfenster ausgelegt sein. Zusätzlich unterschätzen viele Pläne die Möglichkeit paralleler Incidents.
Framework-Vergleich	C5:2026 SIM-01 (Policy for Security Incident Management) und SIM-02 (Security Incident Response Plans). NIS2-DVO Nr. 3.1 (Konzept für die Bewältigung von Sicherheitsvorfällen). DORA Art. 17. NIST SP 800-61 Revision 3.
Härtungsempfehlung	Playbook-Revision für Minuten-Zeitfenster mit vordefinierten automatisierten Containment-Aktionen. SOAR-Integration. Tabletop-Übungen für drei bis fünf parallele Incidents. Vordefinierte Kommunikationsvorlagen für Kunden, Behörden, Öffentlichkeit.

A.5.26 Reaktion auf Informationssicherheitsvorfälle

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Die IR-Ausführung selbst ist teilweise degradiert, weil die menschliche Entscheidungskette zwischen Detection und Reaktion die Reaktionslatenz dominiert. Wo ein Angreifer 80 bis 90 Prozent autonom operiert und Reaktionen in Sekunden auslöst, bleibt einem menschlichen SOC-Analysten oft zu wenig Zeit für eine informierte Entscheidung. Standfest wird das Control nur, wenn definierte Eskalationsstufen vollautomatisch ausgeführt werden.
Framework-Vergleich	C5:2026 SIM-03 (Processing of Security Incidents) und SIM-04 (Documentation and Reporting). NIS2-DVO Nr. 3.5 (Reaktion auf Sicherheitsvorfälle) nennt in 3.5.2 drei Reaktionsphasen (Eindämmung, Beseitigung, erforderlichenfalls Wiederherstellung); die nachträgliche Überprüfung und Lessons Learned sind in Nr. 3.6 separat geregelt. DORA Art. 17–18. NIST SP 800-61 Revision 3.
Härtungsempfehlung	SOAR-Plattform für automatisierte Containment-Aktionen (Account-Sperre, Netzwerkisolation, Schlüsselrotation). Mythos-spezifische Playbooks: Massen-Datenexfiltration, paralleler Multi-Vector-Angriff, Supply-Chain-Kompromittierung. Dedizierte Threat-Hunter-Funktion.

A.5.29 Informationssicherheit während einer Störung

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Das Control adressiert die Aufrechterhaltung von Sicherheitsmaßnahmen im Ausnahmebetrieb. Unter Mythos ist es doppelt degradiert: Erstens parallele Angriffsszenarien, die klassische BCM-Pläne nicht annehmen. Zweitens die

Kategorie	TEILWEISE DEGRADIERT
	Möglichkeit, dass der Angreifer gerade den BCM-Fall provoziert hat und die Notfallinfrastruktur selbst zum Ziel macht. DR-Systeme sind oft schwächer geschützt als Produktionssysteme.
Framework-Vergleich	C5:2026 BCM-01 bis BCM-04. CRY-16 (Operational Continuity for Key Management). NIS2-DVO Nr. 4 (Betriebskontinuitäts- und Krisenmanagement). DORA Art. 11–12.
Härtungsempfehlung	Parallele-Incidents-Szenarien in Tabletops einbauen. DR-Infrastruktur mit gleichwertigem Sicherheitsniveau wie Produktion. Separate IAM-Pfade für DR-Administration. Regelmäßige DR-Failover-Tests einschließlich Sicherheitscontrols.

A.5.33 Schutz von Aufzeichnungen

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Die reine Aufbewahrung ist mythos-orthogonal; die Integritätssicherung ist durch agentische Manipulationsmöglichkeiten bedroht. Ein Angreifer mit Infrastruktur-Zugriff kann Aufzeichnungen nachträglich verändern oder gezielt verfälschen, wenn diese nicht kryptografisch integritätsgesichert sind.
Framework-Vergleich	C5:2026 OPS-12 (Logging and Monitoring – Access, Retention and Deletion) und OPS-14 (Retention of the Logging Data). COM-03 (Internal Audits of the ISMS). NIS2-DVO Nr. 3.2.5 (Aufbewahrung und Schutz der Protokolle). DORA Art. 10.
Härtungsempfehlung	Append-only Storage mit kryptografischer Integritätssicherung (WORM, Object Lock). Hashing-Ketten zur Manipulationsdetektion. Getrennte IAM-Pfade für lesenden und schreibenden Zugriff auf Archive. Separate Aufbewahrung bei externen Dritten für kritische Aufzeichnungen.

A.5.34 Datenschutz und Schutz personenbezogener Daten

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Das Policy-Control bleibt notwendig; die Wirksamkeit hängt an den technischen Controls A.8.11 (Maskierung) und A.8.24 (Kryptografie). Unter Mythos wird die DSGVO-konforme Datenminimierung zum wichtigsten Hebel, weil sie den Blast Radius bei erfolgreicher Exfiltration strukturell begrenzt. Nicht gesammelte Daten können nicht exfiltriert werden.
Framework-Vergleich	DSGVO Art. 5 und Art. 32. C5:2026 PI-03 (Secure Deletion of Data) und OPS-30/31 (Separation of Datasets). ISO/IEC 27701 als Privacy-Erweiterung des ISMS. NIS2-DVO Nr. 8 und Nr. 9 implizit.
Härtungsempfehlung	Datenminimierung technisch durchsetzen (Schema-Validation, ORM-Regeln, API-Response-Filterung). Purpose Limitation als zwingende Attribut-Prüfung in Datenzugriffen. Pseudonymisierung als Default in Analytics- und ML-Pipelines. Regelmäßige Privacy-Impact-Assessments mit Mythos-Bedrohungsszenarien.

A.5.35 Unabhängige Überprüfung der Informationssicherheit

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Jährliche oder zweijährliche Audits sind gegen die Iterationsgeschwindigkeit von Mythos-Bedrohungen strukturell zu langsam. Das Control bleibt als Governance-Mechanismus sinnvoll, liefert aber keine zeitnahe Wirksamkeitsbewertung mehr. Zwischen zwei Audits können sich Angriffsfläche, TTPs und Control-Wirksamkeit mehrfach signifikant verändert haben.

Kategorie	TEILWEISE DEGRADIERT
Framework-Vergleich	C5:2026 COM-03 (Internal Audits of the ISMS) und COM-04 (Information on Information Security Performance). NIS2-DVO Nr. 2.3 (Unabhängige Überprüfung der Netz- und Informationssicherheit). DORA Art. 6(5) für Finanzsektor deutlich schärfer.
Härtungsempfehlung	Continuous-Audit-Mechanismen über automatisiertes Control-Monitoring. Kennzahlenbasierte Wirksamkeitsmessung mit Echtzeit-Dashboards. Ergänzende Penetrationstests und Red-Teaming zwischen formellen Audits. Purple-Team-Übungen zur fortlaufenden Validierung.

A.6.3 Informationssicherheitsbewusstsein, Aus- und Weiterbildung

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Klassische Awareness-Trainings bereiten Mitarbeitende nicht auf Mythos-generierte Social-Engineering-Qualität vor. AI-generierte Phishing-Nachrichten sind grammatikalisch, kontextuell und stilistisch von legitimer Kommunikation kaum mehr zu unterscheiden. Deepfake-basierte Vishing- und Impersonation-Angriffe umgehen die Erkennungsmuster klassischer Trainings. Das Control bleibt notwendig, muss inhaltlich nachgeschärft werden.
Framework-Vergleich	C5:2026 HR-03 (Security Training and Awareness Programme) und DEV-04 (Safety Training Regarding Continuous Software Delivery). NIS2-DVO Nr. 8 (Cyberhygiene und Schulungen im Bereich der Cybersicherheit). DORA Art. 13(6).
Härtungsempfehlung	Mythos-spezifische Trainingsmodule: AI-Phishing-Erkennung, Deepfake-Vishing, Impersonation-Szenarien. Regelmäßige AI-generierte Phishing-Simulationen statt statischer Templates. Stärkere Betonung von Prozess-Verifikation statt Sender-Erkennung (Out-of-Band-Verifikation bei Finanzanweisungen, Credential-Resets).

A.6.7 Remote-Arbeit

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Remote-Work-Policies bleiben notwendig, sind aber durch die erweiterte Angriffsfläche privater Netze, unverwalteter Geräte und schwächerer physischer Sicherung degradiert. Mythos-Angreifer können diese Schwächen automatisiert ausnutzen. Standfest wird das Control nur mit hardware-attestierter Access-Steuerung und strikter Device-Compliance-Prüfung.
Framework-Vergleich	C5:2026 HR-07 (Remote Working – Policy) und HR-08 (Remote Working – Implementation). NIS2-DVO Nr. 11 (Zugriffskontrolle) verlangt differenzierte Absicherung abhängig vom Zugriffskontext. NIST SP 800-46.
Härtungsempfehlung	Hardware-attestierte Endpunkte für privilegierte Remote-Zugänge. Zero-Trust-Network-Access (ZTNA) statt klassisches VPN. Conditional Access mit Device-Posture-Prüfung. Verschlüsselte, verwaltete Arbeitsstationen mit striktem Application-Allowlisting.

A.6.8 Meldung von Informationssicherheitsereignissen

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	User-Reporting ist als Kanal notwendig, verliert aber unter Mythos an Detection-Bedeutung. Agentische Angriffe laufen oft vollständig unterhalb der User-Wahrnehmungsschwelle ab: API-Missbrauch, Credential-Stuffing auf Service-Accounts, fragmentierte Daten-Exfiltration. Der Nutzer sieht nichts, was er melden könnte.

Kategorie	TEILWEISE DEGRADIERT
Framework-Vergleich	C5:2026 SIM-05 (Duty of the Personnel to Report Security Incidents to a Central Body). NIS2-DVO Nr. 3.3 (Meldung von Ereignissen) verlangt einen einfachen Mechanismus für Mitarbeitende, Anbieter und Kunden; Nr. 3.3.2 verpflichtet zur Unterrichtung dieser Kreise über den Mechanismus.
Härtungsempfehlung	Fokus auf automatisierte Detection statt User-Reporting. Einfache Ein-Klick-Meldung im Mail-Client für verdächtige Nachrichten. Feedback-Kanal für den Nutzer, damit Meldungen ernst genommen werden und Meldeverhalten gefördert wird. Mythos-spezifische Schulungsschwerpunkte: Mitarbeitende explizit schulen, ungewöhnliche Authentisierungs-Indikatoren zu melden – unerwartete MFA-Push-Notifications ohne eigenen Login-Versuch, fremde aktive Sitzungen im Account-Übersichtsbildschirm, unerklärbare Account-Lockouts, plötzliche Push-Notification-Stürme (MFA Fatigue Attacks). Diese Indikatoren sind oft das einzige sichtbare Zeichen einer agentischen Credential-Stuffing-Kampagne.

A.7.9 Sicherheit von Werten außerhalb der Räumlichkeiten

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Das Control adressiert primär physischen Diebstahlsschutz für mobile Geräte; in dieser Dimension bleibt es wirksam. Unter Mythos ist die Remote-Kompromittierung mobiler Geräte das größere Risiko, die nur teilweise von A.7.9 abgedeckt wird (ansonsten A.8.1). Die physische Ebene allein ist nicht ausreichend.
Framework-Vergleich	C5:2026 AM-12 (Removable Media and Endpoint Devices). NIS2-DVO Nr. 12 (Anlagen- und Wertemanagement) und Nr. 13 (physische Sicherheit). NIST SP 800-124.
Härtungsempfehlung	Full-Disk-Encryption mit HSM-gesicherter Schlüsselverwaltung. Remote-Wipe-Fähigkeit über MDM. Tracking-Fähigkeit (Find-My-Device). DLP auf Endpunktebene mit Volumen-Anomalieerkennung: Egress-Schwellwerte je Geräteklasse (Alert bei mehr als 2σ Abweichung von der 30-Tage-Baseline der ausgehenden Datenmenge, automatischer Block bei mehr als 3σ). Geofencing für hochsensitive Geräteklassen.

A.8.1 Benutzer-Endpunktgeräte

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Klassische signaturbasierte Endpoint-Protection ist gegen AI-generierte Malware-Varianten stark degradiert: Die Fähigkeit, funktional äquivalente, aber signaturresistente mutierte Malware automatisiert zu erzeugen, entwertet das Signatur-Paradigma. EDR mit Verhaltensanalyse bleibt standfest. Das Control als Ganzes ist teilweise degradiert, weil die Implementierungstiefe über Wirksamkeit entscheidet.
Framework-Vergleich	C5:2026 OPS-04/05 (Protection Against Malware) und OPS-26 (System Hardening). AM-12 (Removable Media and Endpoint Devices). NIS2-DVO Nr. 11 (Zugriffskontrolle) und Nr. 8 (Cyberhygiene). NIST SP 800-53 SI-3 und SI-4.
Härtungsempfehlung	EDR mit Verhaltensanalyse und AI-gestützter Detektion. Application Allowlisting statt Blocklisting. Hardware-basierte Identität (TPM, Secure Enclave) als Voraussetzung für Zugriff. Automatisierte Netzwerk-Isolation bei High-Confidence-Verdacht (klare Definition: EDR-Confidence-Score von 80 % oder höher, oder Korrelation von drei oder mehr ATT&CK-Techniken im 60-Minuten-Fenster). Zero-Trust-Endpoint-Architektur.

A.8.3 Einschränkung des Informationszugangs

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Rollenbasierte Datenzugriffskontrolle greift gegen Einzelzugriffe konform der Policy, ist aber aggregationsblind: Ein kompromittierter Account, der in seiner Rolle agiert und Daten in Mikroschritten abrufen, erzeugt keinen Alarm. Die Degradation betrifft insbesondere Read-Zugriffe auf große Datenbestände – klassisches Symptom einer Massen-Exfiltration durch legitim authentifizierte Accounts.
Framework-Vergleich	C5:2026 IAM-07 (Access to Cloud Service Customer Data) und IAM-05 (Regular Review of Access Rights). NIS2-DVO Nr. 11 (Zugriffskontrolle). DORA Art. 9. NIST SP 800-53 AC-3 und AC-4.
Härtungsempfehlung	Attribute-Based Access Control (ABAC) mit Verhaltenskontext und Volumengrenzen. Data Access Governance mit ML-Anomalie-Detection für Read-Pattern. Bloat-Detection für ungewöhnliche Datenmengen innerhalb erlaubter Rollen. Automatische Stufung von Abfragen basierend auf Zugriffs-Historie.

A.8.4 Zugang zu Quellcode

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Source-Code-Zugriffe sind bei starkem IAM-Schutz formal abgesichert. Unter Mythos ist die Mass-Exfiltration durch kompromittierte Developer-Accounts ein signifikantes Risiko: Wholesale-Clone ganzer Repositories über agentisch automatisierte Git-Operationen hinterlässt keine anomalen Einzelereignisse, sondern nur im Aggregat erkennbare Muster.
Framework-Vergleich	C5:2026 DEV-11 (Protection of Development and Test Environments) und IAM-06 (Privileged Access Rights). NIS2-DVO Nr. 6 (Sicherheitsmaßnahmen bei Entwicklung) verlangt angemessenen Zugriffsschutz auf Entwicklungsartefakte.
Härtungsempfehlung	Repository Anomaly Detection für Clone-Velocity und ungewöhnliche Access-Patterns. SSO mit FIDO2-MFA-Pflicht für Code-Repositories. Dedicated Developer-Workstations mit Device-Attestation. Code-Signing mit Hardware-Keys. Audit-Trails aller Clone- und Fork-Operationen.

A.8.5 Sichere Authentisierung

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Wie A.5.17 adressiert dieses Control die technische Umsetzung der Authentisierung. Phishing-resistente MFA nach WebAuthn-Standard ist mythosstandfest. SMS-basierte MFA, zeitbasierte OTPs und Push-Notifications ohne Number-Matching sind gegen AI-qualitatives Phishing degradiert.
Framework-Vergleich	C5:2026 IAM-08 (Authentication Mechanisms) und PSS-05 (Authentication Mechanisms) mit produktseitigen Anforderungen an Cloud-Dienste. NIS2-DVO Nr. 11.6 und Nr. 11.7. NIST SP 800-63B; AAL2 oder AAL3 für Mythos-relevante Systeme. FIDO Alliance Standards.
Härtungsempfehlung	Phishing-resistente MFA (FIDO2, Passkeys, Hardware-Token nach WebAuthn) als Pflicht für privilegierte und extern erreichbare Accounts. Zertifikatbasierte Authentisierung für Service-to-Service-Kommunikation. Kontinuierliche Authentisierung mit Verhaltensbiometrie für sensitive Sitzungen.

A.8.7 Schutz vor Malware

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Signaturbasierter Malware-Schutz ist unter Mythos stark degradiert, weil AI-Assistenz die Erzeugung funktional äquivalenter Malware-Varianten dramatisch verbilligt hat. Verhaltensbasierte Detektion und Sandbox-Analyse sind standfest. Die Einstufung des Gesamt-C Controls hängt von der gewählten Implementierungstiefe ab.
Framework-Vergleich	C5:2026 OPS-04 (Protection Against Malware – Policies and Procedures) und OPS-05 (Implementation). NIS2-DVO Nr. 8 (Cyberhygiene). NIST SP 800-53 SI-3. CIS Control 10.
Härtungsempfehlung	EDR mit verhaltensbasierter Detection und ML-Engine (z. B. CrowdStrike Falcon, SentinelOne Singularity, Microsoft Defender for Endpoint mit Defender XDR, Palo Alto Cortex XDR). Sandboxing für unbekannte Dateien und Prozesse mit ML-basierter Detonation-Analyse (z. B. Joe Sandbox, ANY.RUN, Hatching Triage). Application Allowlisting als komplementärer Schutz. Deep-Inspection auf Netzwerkebene für Command-and-Control-Traffic.

A.8.12 Verhinderung von Datenabfluss

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Content-basiertes DLP greift gegen klassische Exfiltration-Patterns, ist aber aggregationsblind: Fragmentierte Exfiltration in kleinen, unauffälligen Transfers über längere Zeiträume bleibt unerkannt. Verhaltensbasiertes DLP mit Volumen- und Pattern-Anomalie-Analyse ist standfest.
Framework-Vergleich	C5:2026 COS-08 (Policies for Data Transmission) und PI-01 (Safety of Input and Output Interfaces). NIS2-DVO Nr. 8 (Cyberhygiene). NIST SP 800-53 AC-4 und SI-4.
Härtungsempfehlung	DLP mit Verhaltenskontext und Volumen-Anomalie-Detection. UEBA-Integration. Egress-Traffic-Analyse mit ML-gestützter Baseline. Kombination von content-based, context-based und behavior-based DLP. Automatisierte Quarantäne bei hohem Risk-Score.

A.8.16 Überwachungsaktivitäten

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Signaturbasiertes Monitoring ist durch Aggregationsblindheit gegen Mikroschritt-Angriffe degradiert. Verhaltensbasiertes Monitoring mit AI-Unterstützung ist standfest. Die Einstufung reflektiert, dass viele ISMS das Control klassisch implementieren und damit gegen Mythos-TTPs unwirksam ausführen.
Framework-Vergleich	C5:2026 OPS-13 (Security Information and Event Management) fordert SIEM-basierte Korrelation. OPS-18 bis OPS-25 für Vulnerability-Management-Integration. NIS2-DVO Nr. 3.2 (Überwachung und Protokollierung) mit expliziten Anforderungen an automatisierte, kontinuierliche Überwachung. DORA Art. 10.
Härtungsempfehlung	UEBA mit ML-gestützter Baseline. Kill-Chain-orientierte Detection statt isoliertem Alert-Fokus. Korrelation fragmentierter Einzelereignisse über längere Zeiträume. Integration von MITRE ATT&CK als Detection-Framework mit messbarer Coverage (Validierung via DeTT&CT oder Atomic Red Team): Zielwert mindestens 60 % Coverage der für das dokumentierte organisationsspezifische Threat Profile relevanten Techniken der eigenen Branche. Detection-Schwerpunkte gegen Mythos-TTPs: Living-off-the-Land-Binaries (PowerShell mit Base64-Encoding, schtasks-Persistenz, Office-App spawn cmd oder powershell, ungewöhnliche curl-/wget-Aufrufe von Service-Accounts), Beacon-less C2 (DNS-Tunneling, HTTPS-C2 mit langen Sleep-Intervallen, Cloud-API-basierte C2-Kanäle wie missbrauchte

Kategorie	TEILWEISE DEGRADIERT
	SaaS-Integrationen). Dedizierte Threat-Hunting-Funktion mit dokumentierter Methodik (PEAK-Framework oder TaHiTI), Mindestkapazität 0,5 FTE bis 500 Mitarbeitende, 1 oder mehr FTE darüber; mindestens zwölf Hypothesen-basierte Hunts pro Jahr mit ATT&CK-Mapping. Praktikabilitätshinweis: Für Organisationen ohne dediziertes SOC sind Managed Detection & Response (MDR)-Services (z. B. CrowdStrike Falcon Complete, Arctic Wolf, SentinelOne Vigilance, Sophos MDR) eine valide Alternative mit vertraglich verankerter MTTC-SLA.

A.8.20 Netzwerksicherheit

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Klassische Zonen-Architekturen mit Perimeter-Firewalls und Pivot-Routen zwischen vertrauenswürdigen internen Zonen bieten einem Mythos-Angreifer Lateral-Movement-Räume, die automatisiert ausgekundschaftet werden können. Zero-Trust-Architektur mit Workload-Identity ist standfest.
Framework-Vergleich	C5:2026 COS-01 (Technical Safeguards) bis COS-07 (Documentation of the Network Topology). NIS2-DVO Nr. 6.7 (Netzicherheit) verlangt dokumentierte Netzarchitektur, Fernzugriffskontrolle und Deaktivierung nicht benötigter Dienste; Nr. 11 (Zugriffskontrolle) und Nr. 8 (Cyberhygiene) ergänzen. NIST SP 800-207 Zero Trust Architecture.
Härtungsempfehlung	Reifegrad-Pfad zur Zero-Trust-Netzwerkarchitektur in drei Stufen, statt Big-Bang-Migration. Stufe 1 (sofort, Brückenkonzept für klassische Rechenzentren): Interne Segment-Firewalls ohne Trust-Pivot zwischen internen Zonen, Identity-Aware Proxy (z. B. Cloudflare Access, Tailscale, Zscaler ZPA) für privilegierte Pfade. Stufe 2 (mittelfristig): SASE für Remote-Zugriffe, Service Mesh mit mTLS für die kritischsten Service-zu-Service-Pfade. Stufe 3 (langfristig): Identitätsbasierte Mikrosegmentierung auf Workload-Ebene flächendeckend.

A.8.21 Sicherheit von Netzwerkdiensten

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Perimeter-orientierter Schutz von Netzwerkdiensten ist degradiert. Identitätsbasierter Schutz über mTLS, SPIFFE-Identitäten und Service-Mesh ist standfest. Das Control hängt an der gewählten Architekturdizziplin – klassische VPN- und Jump-Host-Modelle sind teilweise degradiert.
Framework-Vergleich	C5:2026 COS-02 (Security Requirements for Connections in the Cloud Service Provider Network) und COS-05 (Networks for Administration). NIS2-DVO Nr. 6.7 (Netzicherheit) und Nr. 8. NIST SP 800-207 und SP 800-53 SC-Familie.
Härtungsempfehlung	mTLS für alle Service-to-Service-Verbindungen. SPIFFE/SPIRE für Workload-Identity. Service Mesh (Istio, Linkerd) zur zentralen Policy-Durchsetzung. API-Gateway mit Token-basierter Authentisierung statt IP-Allowlisting.

A.8.22 Trennung in Netzwerken

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	VLAN-basierte Segmentierung ist gegen Lateral Movement eines agentischen Gegners zu grob: Innerhalb einer Zone sind klassische Netzwerk-Trennmechanismen oft nicht wirksam. Mikrosegmentierung auf Identitäts- statt Netzwerkebene ist standfest.

Kategorie	TEILWEISE DEGRADIERT
Framework-Vergleich	C5:2026 COS-05 (Networks for Administration) und COS-06 (Separation of Data Traffic in Jointly Used Network Environments). NIS2-DVO Nr. 6.7 (Netzsicherheit) und Nr. 11. NIST SP 800-53 AC-4 und SC-7.
Härtungsempfehlung	Mikrosegmentierung auf Workload-Ebene (z. B. Calico, Cilium). Identity-Aware Firewalls. East-West-Traffic-Inspection mit ML-basierter Anomalieerkennung. Getrennte Admin-Netzwerke mit dedizierter Infrastruktur.

A.8.25 Sicherer Entwicklungslebenszyklus

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Das klassische Secure-SDLC-Konzept mit gelegentlichem Security-Review vor Releases ist teilweise degradiert. Die Geschwindigkeit AI-gestützter Angriffsvektor-Entwicklung erfordert, dass Sicherheitsprüfung in jeder CI-Iteration erfolgt.
Framework-Vergleich	C5:2026 DEV-01 bis DEV-15 decken den gesamten SDLC ab, inkl. DEV-04 (Safety Training), DEV-05 (Design Documentation for Security Features) und DEV-06 (Risk Assessment). NIS2-DVO Nr. 6 (Sicherheitsmaßnahmen bei Erwerb, Entwicklung und Wartung). NIST SP 800-218 SSDF. OWASP SAMM.
Härtungsempfehlung	AI-gestütztes Pre-Commit-Scanning mit aktueller Frontier-Modellgeneration. Konkrete Tool-Klassen: SAST (Semgrep, SonarQube, Checkmarx), DAST (OWASP ZAP, Burp Suite Enterprise), SCA (Snyk, Dependency-Track, OWASP Dependency-Check), AI-gestütztes Code-Scanning der nächsten Generation (GitHub Copilot Autofix, Snyk Code DeepCode AI, Semgrep Pro AI-Assist, Endor Labs Code, Socket.dev). Secure-by-Default-Frameworks und Libraries. Threat Modeling als Pflichtphase in Design-Reviews mit dokumentierter Methodik (STRIDE, PASTA, LINDDUN für Privacy-Threats).

A.8.26 Sicherheitsanforderungen für Anwendungen

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Formalisierte Sicherheitsanforderungen bleiben notwendig, entfalten Wirkung aber erst durch automatisierte Verifikation. Ohne Verknüpfung mit dem Security-Testing in Abnahme (A.8.29) bleibt das Control Dokumentation ohne Wirkung.
Framework-Vergleich	C5:2026 DEV-05 (Design Documentation for Security Features) und DEV-07 (Testing Changes). PSS-01 bis PSS-12 für Cloud-Produktsicherheitsanforderungen. NIS2-DVO Nr. 6. OWASP ASVS als etablierter Requirements-Katalog.
Härtungsempfehlung	Security Requirements als ausführbare Tests formalisieren (Security as Code). Automatisierte Compliance-Checks in CI. Threat Modeling als Requirements-Quelle mit STRIDE oder PASTA. OWASP ASVS als Basis-Requirements-Katalog.

A.8.28 Sicheres Programmieren

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Schulungen und Guidelines allein sind gegen Mythos degradiert; automatisiertes SAST mit AI-gestützter Verifikation ist standfest. Manuelle Code-Reviews ohne Werkzeugunterstützung sind zu langsam und übersehen systematisch die Fehlerklassen, die Mythos-Angreifer automatisiert finden.
Framework-Vergleich	C5:2026 DEV-04 (Safety Training Regarding Continuous Software Delivery). NIS2-DVO Nr. 6 (Sicherheitsmaßnahmen bei Entwicklung). NIST SP 800-218 SSDF. OWASP Secure Coding Practices. CERT Secure Coding Standards.

Kategorie	TEILWEISE DEGRADIERT
Härtungsempfehlung	SAST in der IDE (Shift-Left) mit Entwicklerfeedback vor Commit. AI-Code-Review als Pflicht-Stufe vor Merge. Secure Coding Guidelines via Pre-Commit-Hooks enforced. Dependency-Pinning und Supply-Chain-Scanning.

A.8.30 Ausgelagerte Entwicklung

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Vertragliche Kontrolle allein ist unter Mythos degradiert. Externe Entwicklungsteams können selbst zum Angriffsvektor werden: durch Kompromittierung, Social Engineering oder unbeabsichtigte Introdution unsicherer Dependencies. Automatisierte Artefakt-Prüfung und SBOM-Pflichten härten das Control.
Framework-Vergleich	C5:2026 DEV-02 (Outsourcing of the Development). NIS2-DVO Nr. 5 (Sicherheit der Lieferkette) und Nr. 6. CRA Annex I Teil II. SLSA für Build-Attestation.
Härtungsempfehlung	SBOM-Pflicht für alle gelieferten Artefakte. SLSA Level 3+ für Build-Provenance. Code-Audit-Rechte vertraglich festschreiben. Automatische Artefakt-Prüfung in der eigenen Pipeline. Identische Security-Standards wie für interne Entwicklung.

A.8.32 Änderungsmanagement

Kategorie	TEILWEISE DEGRADIERT
Mythos-Befund	Klassische Change-Management-Prozesse mit mehrwöchigen Freigabezyklen sind gegen Mythos-Zeitkompression selbst zum Risiko geworden. Security-Patches können nicht Wochen warten, wenn Exploits binnen Stunden öffentlich werden. Anthropic empfiehlt explizit die Vorbereitung von Emergency-Change-Procedures.
Framework-Vergleich	C5:2026 DEV-03 (Policies for Changes to System Components), DEV-06 (Risk Assessment, Categorisation and Prioritisation of Changes) und DEV-15 (Exceptions to the Change Management Process). NIS2-DVO Nr. 6. ITIL als klassische Referenz.
Härtungsempfehlung	Emergency-Change-Procedures im Voraus definieren und auditierbar machen. Automated Change Validation mit CI-gestützten Security-Regression-Tests. GitOps als Standard-Change-Pfad mit automatisierter Rollback-Fähigkeit. Separate SLAs für Security-kritische Patches (maximal 24 bis 48 Stunden für KEV-Listings).

5.3 Zusammenfassung

Die 37 teilweise degradierten Controls teilen sich in fünf wiederkehrende Ergänzungsmuster auf, die in Kombination den Übergang von teilweiser Degradation zu Mythos-Standfestigkeit ermöglichen:

- Verhaltensbasierte Detection statt Signaturen: UEBA, ML-gestützte Anomalieerkennung, Kill-Chain-Korrelation. Betrifft A.8.7, A.8.12, A.8.16 und ergänzend A.5.3, A.5.14, A.5.15, A.8.3.
- Phishing-resistente und hardware-gebundene Authentisierung: FIDO2, Passkeys, Hardware-Token, mTLS mit SPIFFE. Betrifft A.5.17, A.8.5, A.8.21 und ergänzend A.6.7, A.8.1, A.8.4.
- Identitätsbasierte Netzwerkarchitektur (Zero Trust): Mikrosegmentierung, Identity-Aware Firewalls, Service Mesh. Betrifft A.8.20, A.8.21, A.8.22 und ergänzend A.5.15, A.8.3, A.6.7.
- Automatisiertes, kontinuierliches Security-Testing: AI-gestütztes Pre-Commit-Scanning, SAST/DAST in CI, SBOM-Pipelines. Betrifft A.8.25, A.8.26, A.8.28, A.8.30, A.8.32 und ergänzend A.5.19, A.5.21.
- Supply-Chain-Transparenz und kontinuierliches Vendor-Monitoring: SBOM, SLSA, automatisiertes Vendor-Security-Monitoring. Betrifft A.5.19, A.5.20, A.5.21, A.5.22, A.5.23.

Diese fünf Ergänzungsmuster sind nicht redundant, sondern komplementär. Sie werden in den Härtungsempfehlungen der Einzelcontrols wiederkehrend referenziert und in Kapitel 9 zur Synthese der Mythos-Härtungs-Controls zusammengeführt.

6 Reine Reibung – Controls, deren klassische Implementierungsform ersetzt oder fundamental gehärtet werden muss

6.1 Übersicht und Risikoprofil

Ein Control gilt als reine Reibung, wenn seine Kernwirkung gegenüber Mythos-Angreifern strukturell entfällt – entweder weil sie auf begrenzter Angreiferkapazität beruht oder weil sie menschliche Reaktionszeit in einer automatisiert durchlaufenen Kill-Chain voraussetzt (siehe Definition in Kap. 3.1.3). Die im Glasswing-Defensivpost (April 2026) formulierte Kernaussage des Anthropic-Security-Engineering-Teams – dass Mitigationen, deren Wert in der Erzeugung von Reibung liegt, gegen einen Gegner mit unbegrenzter Geduld deutlich an Wirksamkeit verlieren – trifft auf die folgenden vier Controls unmittelbar zu.

Die vier Controls eint, dass ihre Schutzwirkung nicht aus struktureller Unmöglichkeit oder kryptografischer Härte stammt, sondern entweder aus der Annahme begrenzter Angreiferkapazität oder aus der Annahme menschlich handhabbarer Reaktionszeiten. Beide Annahmen sind unter Mythos-Bedingungen empirisch widerlegt. Die Controls müssen in ihrer Grundlogik ersetzt oder fundamental umgestaltet werden.

Wichtig zur Einordnung: Die Klassifikation als reine Reibung bedeutet nicht, dass diese Controls obsolet sind. Sie bleiben nützlich gegen bestimmte Angreifertypen und in bestimmten Kontexten. Unter der spezifischen Mythos-Bedrohungslage gilt jedoch: Die klassische beziehungsweise reibungsbasierte Implementierungsform darf nicht als alleinige wirksame Mitigation für Mythos-relevante Risiken angesetzt werden. Eine automatisierte oder anderweitig gehärtete Umsetzung desselben Controls kann sehr wohl wirksame Risikoreduktion herstellen; ihre Wirksamkeit ist organisationsspezifisch nachzuweisen (siehe Kap. 3.1.3 und Anhang F).

6.2 Controls im Detail

A.5.25 Beurteilung und Entscheidung zu Informationssicherheits-Ereignissen

Kategorie	REINE REIBUNG
Mythos-Befund	Das Control adressiert die menschliche Triage-Ebene: einen SOC-Analysten oder Security-Engineer, der eingehende Alerts sichtet, klassifiziert und über Eskalation entscheidet. Unter Mythos überholt die Angreifergeschwindigkeit die menschliche Triage-Zeit systematisch. Im dokumentierten GTG-1002-Fall führte Claude Code 80 bis 90 Prozent der taktischen Angriffsschritte bei Anfrageraten aus, die physisch nicht menschlich erreichbar sind. Eine Triage-Kette, die Entscheidungen in Minuten bis Stunden trifft, steht einem Angreifer gegenüber, der in derselben Zeit bereits mehrere Angriffsstufen abgeschlossen hat. Das Control bleibt formal intakt, entfaltet aber keine wirksame Reaktionsbarriere mehr – die menschliche Reaktionszeit in der automatisch durchlaufenen Kill-Chain ist die strukturelle Schwäche gemäß Kap. 3.1.3.
Framework-Vergleich	C5:2026 SIM-03 (Processing of Security Incidents) und OPS-13 (Security Information and Event Management). NIS2-DVO Nr. 3.4 (Bewertung und Klassifizierung von Ereignissen) verpflichtet zu dokumentierten Triage-Kriterien und zum Korrelations-/Analyseverfahren für Protokolle. NIST SP 800-61 Revision 3. DORA Art. 17.
Ersatz-Empfehlung	Menschliche Erst-Triage durch AI-gestützte Tier-1-Automation ersetzen. SOAR-Plattform mit vordefinierten Playbooks für autonome Containment-Aktionen bei eindeutigen Signalen (bekannte IOCs, Volumen-Anomalien über Schwellwerten, Credential-Stuffing-Muster). Menschliche Eskalation auf ambige Fälle und hohe Schadensszenarien konzentrieren. SOC als Threat-Hunting- und Incident-Commander-Funktion neu aufstellen, nicht als Alert-Sichtung. Measurable SLA: Mean Time to Containment unter 10 Minuten für High-Confidence-Alerts. Praktikabilitätshinweis: SOAR-Plattformen (Splunk SOAR, Palo Alto XSOAR,

Kategorie	REINE REIBUNG
	Microsoft Sentinel mit Automation Rules) erfordern dediziertes Engineering für Playbook-Pflege. Für Organisationen ohne 24/7-SOC ist Managed Detection & Response (z. B. CrowdStrike Falcon Complete, Arctic Wolf, SentinelOne Vigilance, Sophos MDR) mit vertraglich verankerter MTTC-SLA die gangbarere Alternative.

A.5.36 Einhaltung von Richtlinien und Standards für die Informationssicherheit

Kategorie	REINE REIBUNG
Mythos-Befund	Das Control prüft Aktivitäten und Systemzustände gegen dokumentierte Richtlinien. Die grundlegende Schwäche: Die Prüfung basiert auf der Annahme, dass abweichende Angriffsaktionen als solche erkennbar sind. Ein agentischer Angreifer, der einen Angriff in Mikroschritte zerlegt, die jeweils Policy-konform aussehen, umgeht diese Prüfung strukturell. Die Aggregationsblindheit klassischer Compliance-Prüfungen – quartalsweise Reviews, Policy-Abgleiche gegen Konfigurations-Baselines – macht sie zur reinen Reibung: Sie schaffen Aufwand für den Angreifer, aber keinen Schutz, wenn dieser Aufwand automatisiert werden kann.
Framework-Vergleich	C5:2026 COM-03 (Internal Audits of the Information Security Management System) und COM-01 (Identification of Applicable Legal, Regulatory, Self-imposed or Contractual Requirements). NIS2-DVO Nr. 2.2 (Überwachung der Einhaltung) verlangt wirksame Berichterstattung und regelmäßige Konformitätsprüfung. DORA Art. 6(5) und Art. 24.
Ersatz-Empfehlung	Continuous Control Monitoring zusätzlich zu periodischen Audits. Automatisierte, Echtzeit-fähige Compliance-Checks gegen Baselines. Policy-as-Code (OPA/Rego, Sentinel) mit CI-Integration und Runtime-Enforcement. Integration mit UEBA und SIEM zur Erkennung richtlinienkonformer, aber verhaltensanomalier Aktivitäten. Compliance-Dashboards mit quantifizierten Abweichungsindikatoren anstelle von PDF-Auditreports. Verknüpfung mit automatisiertem Incident-Ticket bei Abweichung.

A.8.8 Verwaltung technischer Schwachstellen

Kategorie	REINE REIBUNG
Mythos-Befund	Klassisches Vulnerability-Management mit monatlichen oder quartalsweisen Patch-Zyklen, Change-Advisory-Board-Terminen und dreifach gestaffelten Testphasen ist unter Mythos reine Reibung gegen den in Kapitel 2.1 beschriebenen Kollaps des Patch-Gap-Fensters. Wenn zwischen der Veröffentlichung eines Patches und einem funktionierenden Exploit nur noch Stunden liegen, ist ein Prozess mit mehrwöchigen Freigabezyklen keine Schutzmaßnahme, sondern eine strukturelle Schwäche. Die Existenz eines ordnungsgemäßen, dokumentierten Vulnerability-Management-Prozesses erzeugt Compliance-Nachweise, aber keine Mythos-Resistenz. Präzisierung: Die Reibungs-Einstufung betrifft den zyklischen Patch-Mechanismus. Die organisatorische Prozesshülle des Controls – Schwachstellen-Ownership, SLA-Verantwortung, Eskalationswege – bleibt wirksam und ist Voraussetzung der Automatisierung nach MHC-09 (vgl. Kap. 3.1.3).
Framework-Vergleich	C5:2026 bietet den granularsten Vergleichsmaßstab: OPS-18 (Managing Vulnerabilities – Policies), OPS-25 (Vulnerability Scans) mit der Additional-Sharpening-Stufe OPS-25.01AS (Scan-Frequenz täglich), OPS-27 (Patch Management Policies) und OPS-28 (Patch Management Implementation). NIS2-DVO Nr. 6 (Sicherheitsmaßnahmen bei Erwerb, Entwicklung und Wartung). CRA Art. 13/14 fordert für Software-Hersteller Vulnerability-Handling mit dokumentierten Reaktionszeiten. NIST SP 800-40.

Kategorie	REINE REIBUNG
Ersatz-Empfehlung	EPSS-basierte Priorisierung (Exploit Prediction Scoring System) statt reiner CVSS-Gewichtung. 24-Stunden-SLA für KEV-Listings (CISA Known Exploited Vulnerabilities). Automatisierte Patch-Pipelines mit Canary-Deployments, automatischem Rollback und Security-Regression-Tests. Separater Hotfix-Channel außerhalb des regulären Change-Management für KEV-Einträge. AI-gestütztes Pre-Ship-Vulnerability-Scanning (siehe A.8.29) als Shift-Left-Ergänzung. Emergency-Change-Procedures im Voraus genehmigt und auditierbar (siehe A.8.32).

A.8.23 Webfilterung

Kategorie	REINE REIBUNG
Mythos-Befund	URL- und Reputations-basierte Webfilter versagen unter Mythos strukturell. Die Mythos-generierte Phishing-Infrastruktur rotiert schneller, als Reputation-Feeds aktualisiert werden können: Domain-Generation-Algorithmen, kurzlebige Lookalike-Domains mit automatisch generierten, legitim wirkenden Landing-Pages, auf kompromittierter Legacy-Infrastruktur gehostete Payloads. Die Reibung, die Webfilter erzeugen – dass der Angreifer mehr Phishing-Domains registrieren muss –, ist für einen automatisierten Gegner kein signifikanter Aufwand.
Framework-Vergleich	C5:2026 COS-04 (Cross-Network Access). NIS2-DVO Nr. 8 (Cyberhygiene). NIST SP 800-53 SC-7 (Boundary Protection). CIS Control 9 (Email and Web Browser Protections).
Ersatz-Empfehlung	DNS-Security als Protective-DNS-Layer mit AI-gestützter Erkennung (Cloudflare Gateway, Cisco Umbrella, Quad9) anstelle reiner URL-Blocklisten. Remote Browser Isolation (RBI) für alle externen Links, insbesondere aus Mail. DNS-over-HTTPS ausschließlich zu kontrollierten, vertrauenswürdigen Resolvern. Strukturell: phishing-resistente MFA (A.5.17, A.8.5) als harte Barriere, die auch dann wirkt, wenn der Nutzer eine Phishing-Seite erreicht.

6.3 Zusammenfassung

Die vier Reibungs-Controls teilen eine Logik: Ihre Schutzwirkung beruht auf der Annahme unverhältnismäßigen Angreiferaufwands (A.8.8, A.8.23, A.5.36) oder menschlich handhabbarer Reaktionszeit (A.5.25) – beides hebt Mythos-Klasse-AI aus, weil die Grenzkosten einer weiteren Angriffs-Iteration nahe null liegen. Konsequenz für das SoA: Diese Controls sollten nicht als alleinige Mitigation für Mythos-relevante Risiken verbucht werden; in der Risikobewertung ist zu dokumentieren, welche Ergänzungen aus den Kategorien Standfest und Teilweise degradiert die eigentliche Schutzwirkung übernehmen.

7 Nicht betroffene Controls – neutrale Basis

7.1 Controls ohne Mythos-Relevanz

23 Controls der ISO/IEC 27002:2022 sind gegenüber der Mythos-Bedrohungslage neutral – überwiegend organisatorische, dokumentatorische, Governance- oder physisch gebundene Controls. Sie bleiben für ISMS-Aufbau und Compliance wichtig, erhalten aber keine gesonderte Mythos-Härtungsempfehlung, weil ihre Einstufung unverändert bestehen bleibt.

7.2 Übersicht der 23 nicht betroffenen Controls

ID	Titel	Kurzbegründung
A.5.1	Informationssicherheitsrichtlinien	Policy-Artefakt. Die Wirksamkeit ergibt sich aus der Umsetzung in nachgelagerten Controls.
A.5.2	Informationssicherheitsrollen und -verantwortlichkeiten	Governance-Artefakt. Rollenzuordnung als solche wird durch Mythos nicht betroffen.
A.5.4	Verantwortlichkeiten der Leitung	Führungsverantwortung und Management-Commitment. Nicht mythos-sensitiv.
A.5.6	Kontakt mit speziellen Interessengruppen	Community-Austausch mit fachlichen Gremien. Mythos-neutral.
A.5.8	Informationssicherheit im Projektmanagement	Prozessuales Control für die Integration von Sicherheit in Projektlebenszyklen.
A.5.10	Akzeptable Nutzung von Informationen und zugehörigen Werten	Nutzungspolicy. Die Wirksamkeit ergibt sich aus technischer Durchsetzung in anderen Controls.
A.5.11	Rückgabe von Werten	Offboarding-Artefakt. Mythos-neutral.
A.5.13	Kennzeichnung von Informationen	Operationales Artefakt der Klassifizierung aus A.5.12.
A.5.31	Identifikation gesetzlicher und regulatorischer Anforderungen	Compliance-Inventar.
A.5.32	Rechte an geistigem Eigentum	Lizenz- und IP-Verwaltung.
A.5.37	Dokumentierte Betriebsprozesse	Operationales Artefakt.
A.6.1	Überprüfung (Screening)	Pre-Employment-Check. Hintergrundprüfungen adressieren Insider-Risiko, nicht AI-beschleunigte Externangriffe.
A.6.2	Beschäftigungsbedingungen	Vertragliches Artefakt.
A.6.4	Disziplinarverfahren	HR-Prozess.
A.6.6	Vertraulichkeits- oder Geheimhaltungsvereinbarungen	Rechtliches Artefakt.
A.7.5	Schutz vor physischen und umweltbedingten Bedrohungen	Umweltschutz gegen Feuer, Wasser, Strom-Ausfall.
A.7.7	Aufgeräumter Arbeitsplatz und Bildschirmsperre	Adressiert physische Präsenz-Angriffe (Shoulder Surfing, unbeaufsichtigte Unterlagen); Mythos-Angriffe sind per Definition remote und agentisch.
A.7.8	Aufstellung und Schutz von Geräten	Physische Aufstellungsfrage.
A.7.11	Versorgungseinrichtungen	Unterbrechungsfreie Stromversorgung, Klimatisierung.
A.7.12	Verkabelungssicherheit	Physisches Layer der Infrastruktur.
A.7.13	Wartung von Geräten	Operative Wartung mit physischem Zugang.
A.8.6	Kapazitätsmanagement	Kapazitätsplanung gegen Überlastung; die Abwehr gezielter Erschöpfungsangriffe wird durch A.8.14 (Redundanz) und A.8.20/A.8.21 (Netzwerksicherheit) abgedeckt. Hinweis: Automatisierte Erschöpfungsmuster – API-Abuse, Autoscaling- und Kostenerschöpfung – sind unter Mythos-Bedingungen relevant, betreffen aber die genannten Controls und nicht die Kapazitätsplanung als solche.
A.8.34	Schutz von Informationssystemen während Audittests	Operatives Audit-Artefakt zur Vermeidung von Audit-Seiteneffekten.

Teil III – Konvergenzanalyse: Lücken und Synthese

Die Einzelbewertung aller 93 Controls in Teil II zeigt, welche Controls standhalten, welche ergänzt werden müssen und welche zu ersetzen sind. Teil III identifiziert systematisch, welche Anforderungen andere regulatorisch relevante Frameworks explizit stellen, die die ISO 27002:2022 so nicht kennt oder nur implizit über die Risikobetrachtung abdeckt, und die unter Mythos-Bedingungen eine belastbare Schutzwirkung entfalten. Es geht dabei nicht um Lücken der ISO-Methodik selbst, sondern um Wirksamkeits-Aspekte, die andere Frameworks konkreter benennen. Ergebnis ist ein priorisierter Zusatzkatalog konkreter Controls.

Die Konvergenzanalyse folgt zwei Schritten: Kapitel 8 gruppiert die Lücken thematisch in sieben Clustern und begründet für jeden, warum die jeweilige Anforderung mythos-standfest ist. Kapitel 9 konsolidiert die Ergebnisse in einem nummerierten Katalog von Mythos-Härtungs-Controls.

8 Lückenanalyse – was andere Frameworks explizit fordern, das die ISO 27002 nur implizit abdeckt

8.1 Methodik der Lückenanalyse

Die Lückenanalyse erfolgt dreistufig. Erstens werden alle Anforderungen der fünf primären Vergleichsframeworks – C5:2026, NIST, DORA, CRA und NIS2 mit Durchführungsverordnung – gegen den ISO-27002:2022-Katalog abgeglichen. Zweitens werden jene Anforderungen extrahiert, die ISO 27002 nicht oder nur sehr abstrakt kennt. Drittens wird für jede extrahierte Anforderung geprüft, ob sie unter Mythos-Bedingungen eine harte Schutzwirkung entfaltet – also eines der Kriterien aus Kapitel 3.2 erfüllt: kryptografische Unmöglichkeit, strukturelle Reduktion des Blast Radius, Aggregationsresistenz oder harte Zeitbarriere.

Anforderungen, die lediglich eine andere Formulierung für bereits in ISO 27002 vorhandene Controls sind, werden nicht als Lücke gezählt. Anforderungen, die zwar in einem anderen Framework detaillierter, aber unter Mythos-Bedingungen selbst degradiert sind, ebenfalls nicht. Die folgende Analyse beschränkt sich auf substantielle Differenzen zum ISO-Text bei gleichzeitiger Mythos-Standfestigkeit.

Die Lücken gruppieren sich zu sieben thematischen Clustern, geordnet nach operativer Nähe zur Mythos-Kernbedrohung.

8.2 Cluster 1: Post-Quantum-Strategie und Crypto-Agility

ISO-27002-Bezug: A.8.24 (Verwendung von Kryptographie) fordert einen dokumentierten, risikobasierten Umgang mit kryptografischen Mechanismen, bleibt aber auf abstrakter Ebene. Keine explizite Anforderung an Post-Quantum-Vorbereitung, kryptografisches Inventar oder Crypto-Agility.

Forderung in anderen Frameworks: C5:2026 CRY-01.01AC fordert verbindlich eine dokumentierte Post-Quantum-Strategie mit vier Elementen: einem kryptografischen Inventar mit Prioritätsstufen, laufender Beobachtung des Stands der Technik, Einsatz von Hybrid-Cryptography-Modellen und definierten Triggern, Ressourcen, Transitionsplänen und Erfolgskriterien für die PQC-Umstellung. CRY-01.03AC fordert mindestens jährliche Überprüfung. CRY-02 (Cryptographic Change Management) verankert Crypto-Agility operativ.

Mythos-Standfestigkeit: Die Lücke ist mythos-standfest aus zwei Gründen. Erstens: Das Bedrohungsmodell Store-now-decrypt-later ist konkret und heute wirksam. Exfiltrierte Daten können Jahre gelagert werden, bis Quantum-Computing operationalisiert ist. Zweitens: Crypto-Agility ist eine strukturelle Voraussetzung, um auf plötzlich entwertete Algorithmen reagieren zu können, ohne in wochenlange Reengineering-Zyklen zu geraten. Einschränkend gilt: Der Kern dieser Lücke – die Quantenbedrohung – resultiert nicht aus Mythos, sondern besteht unabhängig davon; die ISO fordert über die Risikobetrachtung implizit bereits angemessene Verfahren. Mythos-spezifisch ist ein dritter Aspekt: PQC-Verfahren sind vergleichsweise neu und weniger erprobt als RSA oder ECDH, sodass AI-gestützte Kryptoanalyse unreife Schemata schneller schwächen könnte. Das erhöht die Anforderung an Krypto-Agilität zusätzlich – etwa mehrere Trust-Anchor mit unterschiedlichen Algorithmen und die Möglichkeit, einzelne kurzfristig zu deaktivieren.

Schlüsselmanagement-Facetten. Eine Analyse der Cloud Security Alliance zum PQC-Schlüsselmanagement (Cloud Key Management Working Group, Entwurf, August 2026) konkretisiert vier Aspekte, die eine reine Algorithmus-Umstellung übersieht. Erstens die Trust-Now-Forge-Later-Variante des Ernte-Angriffs: Nicht nur verschlüsselte Daten, auch heute gültige Signaturen – Zertifikate, Code-Signaturen, langlebige signierte

Dokumente – können geerntet und später mit einem kryptografisch relevanten Quantencomputer rückwirkend gefälscht oder abgestritten werden; langlebige Signaturen gehören daher in Inventar und Migrationsscope. Zweitens das Downgrade-Risiko hybrider Protokolle: Ein Angreifer kann eine hybride Aushandlung (etwa Hybrid-TLS) auf den klassischen Anteil zurückzwingen; erforderlich sind erzwungene quantensichere Aushandlung, Ablehnung rein klassischer Sitzungen und protokollierte Downgrade-Erkennung. Drittens die Migrationsreihenfolge in der Schlüsselhierarchie: Zuerst muss die Wurzel (Master- bzw. Key-Encryption-Key, Stamm-CA) auf PQC umgestellt werden, dann die abgeleiteten Schlüssel – PQC nur auf unteren Ebenen bei klassischer Wurzel lässt den Trust-Anchor verwundbar, besonders bei Archivdaten mit langem Schutzbedarf. Viertens quantensichere Zeitstempel: Ein klassischer Zeitstempeldienst (TSA) schwächt die Non-Repudiation-Kette unabhängig vom Signaturalgorithmus; Signatur- und TSA-Zertifikate sind auf PQC zu migrieren, für Langzeitarchive mit periodischer Zeitstempel-Erneuerung.

8.3 Cluster 2: Supply-Chain-Transparenz und SBOM-Pflicht

ISO-27002-Bezug: A.5.19 bis A.5.23 adressieren die Lieferkette auf Policy- und Vertragsebene. Keine Anforderung an technische Transparenz über eingesetzte Softwarekomponenten, keine SBOM-Pflicht und keine verbindliche Build-Provenance-Anforderung.

Forderung in anderen Frameworks: C5:2026 DEV-13 (Transparency about Software Components). Die Erstellung einer SBOM gehört nach EU CRA Annex I Teil II zu den grundlegenden Anforderungen an das Schwachstellenmanagement; die Erfüllung der anwendbaren Anforderungen ist ihrerseits Voraussetzung der Konformität, auf deren Basis die CE-Kennzeichnung erfolgt. Verbindlich ist dies nur für Hersteller von Produkten mit digitalen Elementen, die in der EU in Verkehr gebracht werden; für andere Organisationen bleibt die SBOM eine risikobasierte Entscheidung. DORA Art. 28 fordert für ICT-Drittdienstleister detaillierte Transparenz. NIS2-DVO Nr. 5 (Sicherheit der Lieferkette) verpflichtet zu dokumentierten Abhängigkeiten. SLSA-Framework definiert Build-Provenance-Levels.

Mythos-Standfestigkeit: SBOM ermöglicht strukturelle Detektion, die nicht auf Angreiferreißung beruht. Der Angriff wird erkennbar, sobald die kompromittierte Komponente in EUVD, CVE oder OSV publiziert wird. Die Reduktion der Detektionslatenz ist strukturell, nicht zeitlich.

8.4 Cluster 3: Container, Confidential Computing und Multi-Tenancy

C5:2026 hat mit der Revision von März 2026 drei strukturell neue Kontrolldomänen eingeführt, die ISO 27002 nicht kennt und die explizit auf moderne Cloud-Workload-Architekturen zugeschnitten sind.

Container-Sicherheit

ISO-27002-Bezug: Keine container-spezifischen Controls. Allgemeine Umgebungs- und Konfigurations-Controls (A.8.9, A.8.19, A.8.31) sind anwendbar, adressieren aber nicht die spezifischen Container-Angriffsflächen.

Forderung: C5:2026 OPS-34 (Container Management – Policies and Procedures) und OPS-35 (Implementation) sowie PSS-11 (Images for Virtual Machines and Containers) fordern dokumentierte Image-Quellen, signierte Images, Runtime-Protection und Netzwerk-Segmentierung auf Container-Ebene. NIST SP 800-190 als Referenz.

Mythos-Standfestigkeit: Container-Signatur schafft eine kryptografisch-strukturelle Barriere gegen kompromittierte Basis-Images aus Public Registries.

Confidential Computing

ISO-27002-Bezug: Keine Anforderungen an Trusted Execution Environments oder Remote Attestation.

Forderung: C5:2026 OPS-32 (Confidential Computing – Policies and Procedures) und OPS-33 (Remote Attestation). NIST-Konzepte zu Trusted Computing als Referenz.

Mythos-Standfestigkeit: Confidential Computing ist die derzeit härteste Barriere gegen Angreifer mit Infrastruktur-Zugriff, weil sie die Vertraulichkeit von Daten in der Verarbeitung schützt – nicht nur in Transit und at Rest.

Multi-Tenancy-Isolation

ISO-27002-Bezug: Keine spezifischen Multi-Tenancy-Controls. Generische Trennungs-Controls (A.8.22, A.8.31) reichen nicht aus.

Forderung: C5:2026 OPS-30 (Separation of Datasets – Policies and Procedures) und OPS-31 (Implementation) fordern explizit die Trennung von Kundendaten in Multi-Tenant-Umgebungen mit nachweisbarer Implementierung. PSS-10 (Software Defined Networking) ergänzt die Netzwerkseite.

Mythos-Standfestigkeit: Strukturelle Begrenzung des Blast Radius: Ein Angreifer, der in einem Tenant Fuß fasst, wird nachweislich daran gehindert, auf andere Tenants zuzugreifen.

8.5 Cluster 4: Verbindliche Meldefristen und Root-Cause-Analyse

ISO-27002-Bezug: A.5.5 (Kontakt mit Behörden) und A.5.25 (Beurteilung von Ereignissen) adressieren Meldewege ohne konkrete Fristen. A.5.27 (Lernen aus Vorfällen) fordert keine Root-Cause-Analysis-Pflicht.

Forderung in anderen Frameworks: NIS2-Richtlinie Art. 23 Abs. 4 fordert binnen 24 Stunden eine Frühwarnung, binnen 72 Stunden eine Vollbenachrichtigung bei erheblichen Sicherheitsvorfällen sowie einen Abschlussbericht nach einem Monat. DORA Art. 19 fordert Major-Incident-Reports an Aufsichtsbehörden mit detaillierten Inhaltsanforderungen. DORA Art. 17 macht Root-Cause-Analyse zur verpflichtenden Phase im Incident-Management-Lifecycle. Der Cyber Resilience Act (CRA) verpflichtet Inverkehrbringer von Produkten mit digitalen Elementen zusätzlich zu Meldepflichten bei aktiv ausgenutzten Schwachstellen (u. a. 24-Stunden-Frühwarnung an ENISA bzw. das zuständige CSIRT). NIS2-DVO Nr. 3.6 verpflichtet zu nachträglichen Überprüfungen einschließlich Ursachenermittlung.

Mythos-Standfestigkeit: Zeitbasierte Meldefristen sind mythos-standfest, weil sie die Zeitkompression, die Mythos-Angreifern zugutekommt, auf der regulatorischen Seite spiegeln: Wer in 24 Stunden melden muss, braucht detection- und triage-seitig entsprechende Automatisierung. Root-Cause-Analyse schafft die empirische Grundlage für Detection-Verbesserung. Die Fristen gelten allerdings nicht für jede Organisation gleichermaßen – maßgeblich sind die Rolle (etwa Inverkehrbringer nach CRA) und die Einstufung unter NIS2; für nicht erfasste Organisationen fällt die Lücke entsprechend kleiner aus. Über die Compliance-Pflichten eines ISMS (Betrachtung der Erwartungen interessierter Parteien, einschließlich rechtlicher Vorgaben) fließen die jeweils anwendbaren Fristen ohnehin in die Sicherheitsziele ein.

8.6 Cluster 5: Kontinuierliche Prüfung zusätzlich zu periodischen Audits

ISO-27002-Bezug: A.5.35 (Unabhängige Überprüfung) beschreibt periodische Audits. A.5.36 (Einhaltung von Richtlinien) und A.8.8 (Vulnerability Management) sind in Kapitel 6 als reine Reibung eingestuft.

Forderung in anderen Frameworks: C5:2026 OPS-25.01AS fordert als Sharpening tägliche Vulnerability-Scans. COM-03 und COM-04 fordern kontinuierliche ISMS-Wirksamkeitsprüfung. NIS2-DVO Nr. 2.2 (Überwachung der Einhaltung) verlangt ein wirksames System der Berichterstattung. DORA Art. 24–26 fordert ein Digital Operational Resilience Testing Programme; Art. 26–27 das Threat-Led Penetration Testing (TLPT) für kritische Funktionen.

Mythos-Standfestigkeit: Kontinuierliche Prüfung gleicht die Zeitkompression auf Verteidigungsseite aus. Die Detektionslatenz wird von Monaten und Quartalen auf Stunden und Tage reduziert. TLPT mit Mythos-spezifischen Szenarien ist das bekannteste regulatorische Regime, das die Empirie-Ebene explizit adressiert: Es prüft, ob die Controls wirken, nicht ob sie dokumentiert sind.

8.7 Cluster 6: Phishing-resistente Identität und Zero-Trust-Architektur

ISO-27002-Bezug: A.5.17 und A.8.5 fordern dem Schutzbedarf angemessene Authentisierung, ohne konkrete Verfahren oder Assurance-Levels zu nennen. A.5.16 fordert Identitätsmanagement ohne explizite Workload-Identity-Anforderung.

Forderung in anderen Frameworks: NIST SP 800-63B definiert drei Authentication Assurance Levels; für Mythos-Umgebungen sind AAL2 und AAL3 relevant. FIDO-Alliance-Standards (WebAuthn, Passkeys). NIST SP 800-207 Zero Trust Architecture definiert Workload-Identity als strukturelles Prinzip. NIST NCCoE-Arbeiten zu Agent Identity. SPIFFE/SPIRE als De-facto-Standard.

Mythos-Standfestigkeit: Phishing-resistente MFA ist die direkte Antwort auf AI-qualitatives Phishing: Selbst wenn ein Nutzer auf einer perfekt gestalteten Phishing-Seite landet, kann der Angreifer keinen verwertbaren Authentisierungsnachweis erlangen, weil FIDO2 eine Origin-Bindung vornimmt. Workload-Identity ist strukturell: Sie ersetzt Shared Secrets durch hardware-verankerte, kurzlebige Identitäten.

8.8 Cluster 7: Automatisierungspflicht und Resilienz-Testing

ISO-27002-Bezug: Keine ausdrückliche Automatisierungspflicht. A.5.30 (ICT-BCM) und A.8.16 (Monitoring) erwähnen keine konkreten Testszenarien für parallele Incidents.

Forderung in anderen Frameworks: NIS2-DVO Nr. 3.2.2 verpflichtet zu automatisierter Überwachung, soweit durchführbar. NIS2-DVO Nr. 3.5.5 fordert die Testung der Reaktionsverfahren. DORA Art. 24–26 fordert ein umfassendes Digital Operational Resilience Testing Programme. DORA Art. 26–27 fordert TLPT mit Szenarien, die reale Bedrohungen abbilden. Anthropic empfiehlt im Glasswing-Defensivpost explizit Tabletop-Übungen für drei bis fünf parallele Incidents.

Hinweis zur Schärfe der NIS2-DVO-Formulierung: Die DVO formuliert die Automatisierungspflicht als Soll-Vorschrift („soweit durchführbar“, „vorbehaltlich der betrieblichen Kapazitäten“) und nicht als absolute Pflicht. Für Mythos-exponierte Organisationen sollte die Durchführbarkeit restriktiv ausgelegt werden, da manuelle Reaktionszeiten gegenüber agentischen Angreifern strukturell versagen.

Mythos-Standfestigkeit: Die Automatisierungspflicht ist die einzig realistische Antwort auf die Asymmetrie zwischen agentischem Angreifer und menschlicher Verteidigungskette. Parallel-Incident-Testing adressiert direkt die Fragmentierungs-Strategie.

8.9 Zusammenfassung: Das Lücken-Delta der ISO 27002

Die sieben Cluster beschreiben gemeinsam das systematische Delta zwischen ISO/IEC 27002:2022 und den regulatorischen Anforderungen der letzten drei Jahre. Dieses Delta wurde in der Konzeption der ISO 27002 nicht als Mangel angelegt – es reflektiert die Framework-weise Spezialisierung der Normenlandschaft.

Unter Mythos-Bedingungen wird dieses Delta zum operativ relevanten Problem. Eine Organisation, die sich auf die reine ISO-27002-Basis stützt, verfügt nicht über die konkreten Anforderungen, die C5:2026, NIST, DORA, CRA und NIS2 mit Durchführungsverordnung an die aktuelle Bedrohungslage stellen. Das Delta zu schließen ist keine Compliance-Übung, sondern die Kernaufgabe der operativen Mythos-Härtung.

Die sieben Cluster lassen sich zu zwei Bewegungen verdichten: einer strukturellen Modernisierung der Sicherheitsarchitektur (Cluster 1 Kryptografie, 3 Cloud-Infrastruktur, 6 Identität) und einer zeitlichen Verkürzung der Verteidigungsschleifen (Cluster 2 Lieferkette, 4 Meldepflichten, 5 kontinuierliche Prüfung, 7 Automatisierung). Kapitel 9 führt diese Bewegungen zum priorisierten Katalog von Mythos-Härtungs-Controls zusammen.

9 Synthese – Katalog der Mythos-Härtungs-Controls

9.1 Zweck und Anwendungsweise des Katalogs

Dieses Kapitel konsolidiert kompakt die Erkenntnisse aus Kapitel 8 zu einem nummerierten Katalog von dreizehn Mythos-Härtungs-Controls (MHC).

Hinweis: Nutzen Sie den MRIS IMPLEMENTATION GUIDE als Umsetzungsleitlinie zu den dreizehn Mythos-Härtungs-Controls.

Jedes MHC ist so formuliert, dass es direkt in das Statement of Applicability eines bestehenden ISO-27001-zertifizierten ISMS aufgenommen werden kann.

Die dreizehn MHC sind ein priorisierter Zusatzkatalog für die Mythos-Bedrohungslage. Sie kommen zusätzlich zu den bestehenden ISO-Controls zum Einsatz, adressieren jedoch Schutzziele, die ISO nicht explizit kennt oder nur abstrakt beschreibt. Jedes MHC referenziert eine ISO-27002-Anbindung.

Die Nummerierung MHC-01 bis MHC-13 folgt der operativen Nähe zur Mythos-Kernbedrohung, nicht einer Implementierungsreihenfolge. Jedes MHC steht für sich und kann unabhängig eingeführt werden; die Gesamtwirkung ergibt sich aus der kombinierten Anwendung.

Der Katalog erhebt keinen Anspruch auf Vollständigkeit. Er beschränkt sich auf die derzeit robustesten und operativ klar beschreibbaren Ergänzungen. Die Versionshinweise dokumentieren die Entwicklung dieses Katalogs.

Priorisierung gegenüber bestehenden Controls. Die dreizehn MHC ergänzen, sie ersetzen nicht. Vor der Investition in neue Mythos-Härtungs-Controls steht die konsequente Umsetzung der bestehenden strukturellen

Controls aus Kapitel 4 (insbesondere Zero-Trust-Architektur, starke Kryptografie, Least Privilege, Secrets-Management, Mikrosegmentierung, unveränderliche Backups). Die strukturellen Controls bleiben unter Mythos die robusteste Verteidigungsschicht. Eine Organisation, die diese Controls nicht konsequent umgesetzt hat, gewinnt durch die Einführung neuer MHC weniger Mythos-Resilienz als durch die Härtung des Bestehenden. Die Bewertungssprache in den Kapiteln 5 und 6 ist daher zu lesen als „selektiv degradiert“ – nicht als „obsolet“. Klassische Controls werden durch Mythos in spezifischen Wirksamkeits-Dimensionen geschwächt, behalten aber in anderen Dimensionen ihre volle Schutzwirkung. Die korrekte Reaktion ist Härtung, nicht Ersatz.

9.2 Dreizehn Mythos-Härtungs-Controls

MHC-01 Post-Quantum-Strategie und kryptografisches Inventar

Einordnung: übergreifendes Resilienz-Control mit erhöhter Relevanz unter AI-beschleunigter Crypto-Discovery und Exploit-Ökonomie. Die Quantenbedrohung selbst besteht unabhängig von Mythos; AI beschleunigt Auffindung schwacher Kryptografie, Migrationsdruck und Ausnutzungsökonomie.

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Flankiert A.8.24 (Verwendung von Kryptographie).
Framework-Vergleich	Framework-Basis: C5:2026 CRY-01.01AC bis CRY-01.03AC. BSI TR-02102. FIPS 203 (ML-KEM), FIPS 204 (ML-DSA) und FIPS 205 (SLH-DSA) als finalisierte PQC-Standards (13.08.2024); HQC als Backup-KEM (ausgewählt 11.03.2025, Standardisierung läuft). NIST SP 1800-38 als NCCoE-Migrationsleitfaden. Roadmap der EU-Kommission zur Post-Quantum-Migration.
Härtungsempfehlung	Kernidee: Dokumentierte PQC-Strategie aus vier Elementen — kryptografisches Inventar aller Algorithmen und Schlüssel, laufende Beobachtung des Stands der Technik, hybride Kryptografie-Modelle, definierte Trigger-Events und Transitionspfade (Wurzel der Schlüsselhierarchie zuerst, Downgrade-resistente Hybrid-Aushandlung). Mythos-Wirkung: Adressiert Store-now-decrypt-later und schafft strukturelle Crypto-Agility; Schutzwirkung beruht auf kryptografischer Unmöglichkeit, nicht auf Angreiferreibung. → Umsetzung, Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-01.

MHC-02 SBOM und Build-Provenance

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Flankiert primär A.5.21 (ICT-Lieferkettenmanagement) und A.8.30 (Ausgelagerte Entwicklung); vollständige Zuordnung siehe Anhang A.
Framework-Vergleich	Framework-Basis: C5:2026 DEV-13. CRA Annex I Teil II als rechtsverbindliche SBOM-Pflicht für Hersteller von Produkten mit digitalen Elementen (Schwachstellen-Meldepflichten ab September 2026, Hauptpflichten ab 11.12.2027). DORA Art. 28. NIS2-DVO Nr. 5. SLSA-Framework. BSI TR-03183-2 als SBOM-Qualitätsmaßstab. Marktlage 2026: Die einsetzenden Offenlegungswellen bestätigen den Volumenbedarf der SBOM-Abfragefähigkeit (Kap. 2.5).
Härtungsempfehlung	Kernidee: Automatisierte SBOM-Generierung in der CI-Pipeline (SPDX/CycloneDX, CRA-konform) für alle selbst entwickelten und weitergegebenen Artefakte, vertragliche SBOM-Pflicht für kritische Lieferanten, automatisierte Vulnerability-Korrelation (CVE/EUVD/OSV) mit getrennter Ausnutzbarkeits-Kommunikation über VEX/CSAF, SLSA Level 3+ für Build-Provenance. Mythos-Wirkung: Schafft strukturelle Detektionsfähigkeit für Supply-Chain-Angriffe; Detektionslatenz sinkt auf die Publikationszeit in Vulnerability-Datenbanken. → Umsetzung, Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-02.

MHC-03 Phishing-resistente Multi-Faktor-Authentisierung

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Flankiert primär A.5.17 (Authentisierungsinformation) und A.8.5 (Sichere Authentisierung); vollständige Zuordnung siehe Anhang A.
Framework-Vergleich	Framework-Basis: NIST SP 800-63B Revision 4 (final 31.07.2025), AAL2/AAL3 – AAL3 verlangt gerätegebundene, nicht exportierbare Schlüssel; synchronisierbare Passkeys nur AAL2. Begriffliche Präzisierung: Ein gerätegebundener Passkey erfüllt das Mehrfaktor-Kriterium erst zusammen mit lokaler Nutzerverifikation (PIN oder Biometrie); ohne sie liegt ein reiner Besitzfaktor vor. Maßgeblich für die Schutzwirkung gegen Mythos-Phishing ist die Origin-Bindung, nicht die Zahl der Faktoren. C5:2026 IAM-08. FIDO-Alliance-Standards (WebAuthn, Passkeys, CTAP2). NIS2-DVO Nr. 11.6 und Nr. 11.7; für privilegierte Konten zusätzlich Nr. 11.3.2 Buchst. a.
Härtungsempfehlung	Kernidee: Origin-gebundene Verfahren (WebAuthn/FIDO2) für alle privilegierten und extern erreichbaren Zugänge; für höchsten Schutzbedarf nur attestierte, gerätegebundene Authenticator (Attestation-Enforcement), synchronisierbare und Software-Passkeys (Passkey-Vaults) ausgeschlossen; SMS-MFA und Push ohne Number-Matching für neue Deployments ausgeschlossen. Mythos-Wirkung: Origin-Bindung macht Authentisierungsnachweise auf Phishing-Seiten wertlos; strukturelle Antwort auf AI-qualitatives Phishing. Wirksamkeitsgrenze: phishing-resistent ist nicht phishing-sicher — Endpoint-Kompromittierung und Sitzungsdiebstahl nach der Anmeldung bleiben über MHC-05 und A.8.7 zu adressieren. → Umsetzung, Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-03.

MHC-04 Workload-Identität und Zero-Trust-Netzwerkarchitektur

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Flankiert primär A.5.16 (Identitätsmanagement), A.8.20 (Netzwerksicherheit) und A.8.21 (Sicherheit von Netzwerkdiensten); vollständige Zuordnung siehe Anhang A.
Framework-Vergleich	Framework-Basis: NIST SP 800-207 Zero Trust Architecture. NIST NCCoE-Arbeiten zu Agent Identity (2026, angekündigt). SPIFFE/SPIRE. Service-Mesh-Architekturen (Istio, Linkerd).
Härtungsempfehlung	Kernidee: Workload-Identität – jede Maschine, jeder Dienst und jeder Container erhält eine eindeutige, kryptografisch verifizierbare Identität mit kurzlebigen, automatisch rotierenden Credentials statt gemeinsam genutzter Zugangsdaten – (SPIFFE/SPIRE, mTLS) in drei Reifegrad-Stufen statt Big-Bang — privilegierte Workloads, Production-Mainstream, Flächendeckung; AI-Agenten als Sonderfall mit eigenen, zeitlich begrenzten, capability-scoped Identitäten pro Tool-Aufruf statt persönlicher Developer-Credentials, vollständiger Audit-Trail. Mythos-Wirkung: Reduziert klassische credential-basierte Lateral-Movement-Pfade strukturell und begrenzt den erreichbaren Schadensradius; Shared Secrets werden durch hardware-verankerte, kurzlebige Credentials ersetzt. Eine kompromittierte legitime Workload kann weiterhin innerhalb ihrer eigenen Berechtigungen handeln – deshalb ist MHC-04 mit Detection (MHC-05) und Berechtigungsminimierung zu kombinieren. → Umsetzung, Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-04.

MHC-05 Verhaltensbasierte Detection und Kill-Chain-Korrelation

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Flankiert primär A.8.16 (Überwachungsaktivitäten) und A.8.12 (Verhinderung von Datenabfluss); vollständige Zuordnung siehe Anhang A.
Framework-Vergleich	Framework-Basis: C5:2026 OPS-13 (SIEM). NIS2-DVO Nr. 3.2. DORA Art. 10. MITRE ATT&CK als operatives Detection-Framework (laufend aktualisiert, Inhaltsversion v19.2). NIST SP 800-94 nur als datierter Hintergrund (Fassung 2007; der Entwurf einer Revision 1 wurde zurückgezogen, ein aktueller Nachfolger liegt nicht vor).
Härtungsempfehlung	Kernidee: UEBA mit ML-Baseline und ATT&CK-basierten Detection-Regeln (messbare Coverage), Kill-Chain-Korrelation über Ereignisse und Zeitfenster; Schwerpunkte gegen Mythos-TTPs (Living-off-the-Land, Beacon-less C2, WebAuthn-/Passkey-Anomalien); Threat-Hunting als eigenständige Funktion; Adversarial Robustness der eigenen Detection-AL. Mythos-Wirkung: Adressiert die Aggregationsblindheit klassischer Detection; fragmentierte Mikroschritt-Angriffe werden durch Muster-Korrelation erkennbar. Realismus-Hinweis: Detection ist nicht Prevention — Low-and-slow-Angriffe, getarnte Agenten in legitimen Baselines und Adversarial-ML-Evasion bleiben Restrisiko; immer in Kombination mit den strukturellen Controls aus Kapitel 4 zu bewerten. → Umsetzung (inkl. MDR-Alternative), Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-05.

MHC-06 Container-Sicherheit und Confidential Computing

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: A.5.23 (Cloud-Dienste, flankiert) und A.8.31 (Trennung von Umgebungen).
Framework-Vergleich	Framework-Basis: C5:2026 OPS-34/35 (Container Management), OPS-32/33 (Confidential Computing) und PSS-11. NIST SP 800-190 (Fassung 2017, weiterhin maßgebliche Container-Baseline). Confidential Computing Consortium.
Härtungsempfehlung	Kernidee: Für Container signierte Base-Images aus kontrollierten Registries (Sigstore/cosign), Admission-Controller-Policies (OPA Gatekeeper, Kyverno) und Image-Scans vor dem Deployment; für Confidential Computing TEE/Secure Enclaves (Intel TDX, AMD SEV-SNP, ARM CCA) mit Remote-Attestation für erhöhten Schutzbedarf. Mythos-Wirkung: Container-Signatur schafft eine kryptografische Hartbarriere gegen manipulierte Images; Confidential Computing ist die derzeit härteste Barriere gegen Angreifer mit Provider-Infrastruktur-Zugriff. → Umsetzung, Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-06.

MHC-07 Multi-Tenancy-Isolation mit nachweisbarer Trennung

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: A.5.23 (Cloud-Dienste, flankiert) und A.8.22 (Trennung von Netzwerken); erweitert um daten- und rechenzentrumsspezifische Aspekte der Mandantentrennung.
Framework-Vergleich	Framework-Basis: C5:2026 OPS-30 und OPS-31. PSS-10 (Software Defined Networking). CSA Cloud Controls Matrix. ISO/IEC 27017 (Informationssicherheit für Cloud-Dienste).
Härtungsempfehlung	Kernidee: Dokumentierte Trennungskonzepte für Daten (tenant-spezifische Schlüssel, logische oder physische Trennung), Netzwerke (VPCs, Namespaces, SDN-Regeln) und Compute (tenant-spezifische Scheduling-Policies), mit nachweisbarer Trennung durch regelmäßige, idealerweise automatisierte Tests. Mythos-Wirkung: Strukturelle Begrenzung des Blast Radius — ein Angreifer, der in einem Tenant Fuß

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
	fasst, wird nachweislich am Zugriff auf andere gehindert. → Umsetzung, Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-07.

MHC-08 Unveränderliche Backups und Recovery-Validierung

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Vertieft primär A.8.13 (Informationssicherung), das im Grundsatz standfest ist, aber konkrete Anforderungen an Unveränderlichkeit nicht explizit formuliert; vollständige Zuordnung siehe Anhang A.
Framework-Vergleich	Framework-Basis: C5:2026 OPS-06 bis OPS-09 mit separater Adressierung von Policies, Monitoring, Testing und Storage. DORA Art. 12. NIS2-DVO Nr. 4.1. Industriepraxis 3-2-1-1-0-Modell.
Härtungsempfehlung	Kernidee: 3-2-1-1-0-Modell (drei Kopien, zwei Medientypen, mindestens eine offsite und eine unveränderlich oder offline, null Fehler beim Restore-Test), getrennte IAM-Pfade für die Backup-Infrastruktur, quartalsweise dokumentierte Restore-Tests. Mythos-Wirkung: Unveränderliche Backups sind eine der wirksamsten Hartbarrieren gegen Ransomware und agentische Datenintegritätsangriffe. → Umsetzung, Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-08.

MHC-09 AI-gestütztes Security-Testing in der Pipeline

Abgrenzung: Das Control setzt automatisiertes Security-Testing (SAST, DAST, SCA) als Basis voraus und beschreibt die Mythos-spezifische Ergänzung darauf – AI-gestützte Scan-Verfahren der nächsten Generation sowie die Absicherung AI-generierten Codes. Die Kennzahl zur Behebungslatenz (Anhang G.4) misst nicht das Testverfahren selbst, sondern dessen Wirkung im Betrieb.

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Flankiert primär A.8.29 und A.8.25; vollständige Zuordnung siehe Anhang A.
Framework-Vergleich	Framework-Basis: C5:2026 DEV-07, OPS-22 und OPS-25 mit der Verschärfung OPS-25.01AS (tägliche Scans). NIST SP 800-218 SSDF, ergänzt um NIST SP 800-218A (Secure Software Development für Generative AI, 26.07.2024); SSDF v1.2 (SP 800-218r1) als Entwurf (17.12.2025). OWASP ASVS. CRA Annex I Teil II. Multi-System-Lage: Vergleichbare agentische Testfähigkeit ist bei mehreren Anbietern verfügbar (u. a. Daybreak, MDASH; Kap. 2.5); die Werkzeugwahl ist anbieterunabhängig zu treffen.
Härtungsempfehlung	Kernidee: SAST/DAST/SCA plus AI-gestütztes Code-Scanning der nächsten Generation in der Pipeline, Pre-Merge-Block bei High/Critical ab 80 % Confidence (KEV sofort blockierend), tägliche Scans, Mythos-spezifische Pentests; sekundäre Bedrohung adressieren — AI-Coding-Assistenten erzeugen unsichere Patterns, daher Pre-Commit-AI-Code-Audit und Ausweis AI-generierten Codes als eigener Risikofaktor im Secure-SDLC-Standard und in der Risikobewertung (Risikoregister bzw. Anwendungs-Risikoklassifikation), nicht im Statement of Applicability – dieses dokumentiert die Anwendbarkeit von Controls, nicht Risiken. Mythos-Wirkung: Shift-Left entzieht dem Angreifer das Fenster zwischen Exploit-Verfügbarkeit und Produktions-Patch; strukturelle Antwort auf den Kollaps des Patch-Gap-Fensters. → Umsetzung, Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-09.

MHC-10 Continuous Control Monitoring und Policy-as-Code

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Ersetzt in weiten Teilen die Wirksamkeit von A.5.35 und A.5.36 unter Mythos-Bedingungen; vollständige Zuordnung siehe Anhang A.
Framework-Vergleich	Framework-Basis: C5:2026 COM-03 und COM-04. NIS2-DVO Nr. 2.2. DORA Art. 6(5). OWASP SAMM. NIST SP 800-137/137A (Information Security Continuous Monitoring). NIST OSCAL (maschinenlesbare Control-Kataloge und Nachweise). Industriepraxis Open Policy Agent (OPA/Rego), HashiCorp Sentinel.
Härtungsempfehlung	Kernidee: Policy-as-Code (OPA, Sentinel) als ausführbare Regeln, integriert in CI (Pre-Deploy) und Runtime-Enforcement; kontinuierliche Compliance-Checks gegen Baselines, Echtzeit-Dashboards mit quantifizierten Abweichungsindikatoren, automatische Incident-Tickets bei Abweichung. Mythos-Wirkung: Verkürzt die Detektionslatenz für Compliance-Abweichungen von Monaten oder Quartalen auf Sekunden bis Minuten. → Umsetzung, Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-10.

MHC-11 SOAR-basierte Tier-1-Automation und parallele Reaktions-Playbooks

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Ersetzt in weiten Teilen die Wirksamkeit von A.5.25, ergänzt A.5.24 und A.5.26; vollständige Zuordnung siehe Anhang A.
Framework-Vergleich	Framework-Basis: C5:2026 SIM-02 und SIM-03. NIS2-DVO Nr. 3.5 (Reaktion auf Sicherheitsvorfälle, drei Phasen nach 3.5.2). DORA Art. 17. NIST SP 800-61 Revision 3 (Vorfallsbehandlung als CSF-2.0-Profil, 03.04.2025; löst die Fassung von 2012 ab). Anthropic-Empfehlung aus dem Glasswing-Defensivpost zu parallelen Incident-Szenarien.
Härtungsempfehlung	Kernidee: SOAR-Plattform mit vordefinierten Playbooks für vollautomatisches Containment bei eindeutigen Signalen und menschlicher Eskalation für ambige Fälle; Mythos-spezifische Playbooks (Massen-Exfiltration, Multi-Vector, Supply-Chain, MFA-Fatigue), Tabletop für drei bis fünf parallele Incidents; Härtung der Automation-Schicht selbst (signierte Trigger, Audit-Trail mit Ausführungs-Identität, Rate-Limits, Human-in-the-Loop bei hohem Blast-Radius, Playbook-Schutz wie Quellcode). Mythos-Wirkung: Schließt den Zeitversatz zwischen agentischer Angreifergeschwindigkeit und menschlicher Reaktionskette; MTTC von Stunden auf Minuten. Wirksamkeitsgrenze: Die SOAR-Logik selbst kann Ziel werden — ohne die genannte Härtung der Automation-Schicht ist MHC-11 nicht wirksam. → Umsetzung (inkl. MDR-Alternative), Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-11.

MHC-12 Threat-Led Penetration Testing mit Mythos-Szenarien

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Ergänzt A.5.35 und A.8.29.
Framework-Vergleich	Framework-Basis: DORA Art. 26 und 27 als derzeit einziges verbindliches regulatorisches TLPT-Regime, konkretisiert durch die TLPT-RTS (Delegierte Verordnung (EU) 2025/1190, anwendbar seit 08.07.2025). TIBER-EU-Rahmenwerk der EZB (seit 11.02.2025 auf DORA ausgerichtet, mit verpflichtendem Purple Teaming; in Deutschland TIBER-DE über Bundesbank/BaFin). C5:2026 OPS-22 (Penetration Tests). NIS2-DVO Nr. 3.5.5 mit Anforderung an regelmäßige Tests der Reaktionsverfahren.
Härtungsempfehlung	Kernidee: Threat-Led Penetration Testing mit externen Red-Teamern nach realen Szenarien, Mythos-spezifisch erweitert (AI-Spear-Phishing mit Deepfake-Voice,

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
	fragmentierte Angriffsketten, Multi-Vector, Supply-Chain, Cloud-API-Missbrauch); mindestens jährlich für kritische Funktionen, Purple Teaming zwischen Zyklen mit ATT&CK-Coverage; Szenario-Pool multi-systemisch anlegen (mindestens zwei unabhängige agentische Systeme). Mythos-Wirkung: TLPT testet die Wirksamkeit gegen reale Mythos-TTPs statt gegen Dokumentation; reduziert falsch-positive Compliance-Zusicherungen. → Umsetzung (inkl. kostengünstigerer Formate für Nicht-DORA-Adressaten), Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-12.

MHC-13 AI-Agent-Governance und Harness-Sicherheit

Kategorie	MYTHOS-HÄRTUNGS-CONTROL
Mythos-Befund	ISO-27002-Anbindung: Erweitert A.5.9 (Inventar, Shadow-AI-Erfassung), A.5.16 (Identitätsmanagement), A.8.27 (sichere Systemarchitektur) und MHC-04 um die spezifischen Governance- und Sicherheitsanforderungen agentischer AI-Systeme im eigenen Betrieb.
Framework-Vergleich	Framework-Basis: OWASP Top 10 for LLM Applications 2026 (insbesondere LLM01 Prompt Injection, LLM03 Excessive Agency). OWASP Agentic Security Initiative (ASI01 bis ASI10). MITRE ATLAS (AML.T0051 LLM Prompt Injection, AML.T0053 AI Agent Tool Invocation (vormals LLM Plugin Compromise; ATLAS-Stand August 2026), AML.T0086 Exfiltration via AI Agent Tool Invocation, AML.T0110 AI Agent Tool Poisoning). Anthropic, OpenAI und Google Responsible-Use-Guidelines. NIST AI RMF 1.0. ISO/IEC 42001 Annex A.
Härtungsempfehlung	Kernidee: Governance der AI-Agenten in fünf Dimensionen — Harness-Audit (Prompts, Tools, Retrieval wie Produktivcode, versioniert, signiert, Pre-Production-Pen-Tests), Blast-Radius-Limits (Capability-Scoping pro Tool-Aufruf, Rate-Limits, Circuit-Breaker, Human-Approval für irreversible Aktionen), Human-Override (Owner mit Kill-Switch, Audit-Trail mit menschlicher Korrelation), Supply-Chain-Inventar agentischer Komponenten (MCP-Server, IDE-Extensions, Agentic Skills) und Pre-Production-Checks (Scope, Bedrohungsmodell, Rollback, Monitoring). Mythos-Wirkung: Schließt die in Mythos-Ready als CRITICAL eingestufte Lücke „Unmanaged AI Agent Attack Surface“ — defensive Risikoseite und Supply-Chain-Risikoseite zugleich. Reifegrad-Realismus: konzeptionell ausgereift, operative Umsetzungsreife derzeit meist auf Stufe Initial (12 bis 24 Monate Übergang); zuerst Inventarisierung, Audit-Trail und Capability-Scoping priorisieren. → Umsetzung, Reifegrad-Pfad und Nachweisführung: Implementation Guide, MHC-13.

9.3 Ableitung aus den Lückenclustern

Zwölf der dreizehn MHC lassen sich primär auf die sieben Lückencluster aus Kapitel 8 zurückführen; MHC-08 entsteht primär aus der Control-Degradation-Analyse und wird durch externe Framework-Anforderungen zusätzlich gestützt:

- Cluster 1 (Kryptografie und Post-Quantum): MHC-01.
- Cluster 2 (Supply-Chain-Transparenz): MHC-02; MHC-13 ergänzt um agentische Supply-Chain-Komponenten (MCP-Server, IDE-Extensions, Agentic Skills).
- Cluster 3 (Container, Confidential Computing, Multi-Tenancy): MHC-06 und MHC-07.
- Cluster 4 (Meldepflichten und Root-Cause): MHC-11 (Reaktionsseite); die Meldefrist-Compliance wird nicht als eigenes MHC gelistet, weil sie als direkte Rechtspflicht aus NIS2 und DORA ohnehin umzusetzen ist.
- Cluster 5 (Kontinuierliche Prüfung): MHC-10 und MHC-12.

- Cluster 6 (Phishing-resistente Identität und Zero Trust): MHC-03, MHC-04 und MHC-13 (Identitäts-Lifecycle der AI-Agenten).
- Cluster 7 (Automatisierung und Resilienz-Testing): MHC-05, MHC-09 und MHC-11.

MHC-08 (unveränderliche Backups) resultiert nicht primär aus der externen Framework-Lückenanalyse, sondern aus der vertieften Bewertung von A.8.13, A.5.29, A.5.30 und der angrenzenden Resilienz-Controls in den Kapiteln 4 und 5: Die ISO fordert Sicherung und Wiederherstellbarkeit, aber keine Unveränderlichkeit und keine getrennten Zugriffspfade für die Backup-Infrastruktur – genau diese beiden Eigenschaften entscheiden unter Mythos-Bedingungen über die Wirksamkeit. MHC-05 ist die operative Grundlage der Automatisierungspflicht (Cluster 7) und zugleich Antwort auf die Aggregationsblindheit klassischer Detection.

9.4 Kompakt-Übersicht des MHC-Katalogs

Die folgende Übersicht fasst die dreizehn MHC zusammen und ordnet jedem MHC seine primäre Framework-Basis und seine ISO-27002-Anbindung zu. Die Anbindungen sind konsistent mit der Bewertungsmatrix in Anhang A und den Einzelbewertungen in den Kapiteln 4 bis 6.

MHC	Titel	Framework-Basis	ISO-Anbindung
MHC-01	Post-Quantum-Strategie und kryptografisches Inventar	C5:2026 CRY-01.01AC	A.8.24
MHC-02	SBOM und Build-Provenance	C5:2026 DEV-13, CRA Annex I, SLSA	A.5.19, A.5.20, A.5.21, A.5.22, A.8.4, A.8.30
MHC-03	Phishing-resistente MFA	NIST SP 800-63B Rev. 4 (AAL2/3), FIDO2	A.5.17, A.6.7, A.8.1, A.8.5, A.8.23
MHC-04	Workload-Identität und Zero Trust	NIST SP 800-207, SPIFFE	A.5.15, A.5.16, A.6.7, A.8.2, A.8.20, A.8.21, A.8.22
MHC-05	Verhaltensbasierte Detection	C5:2026 OPS-13, MITRE ATT&CK	A.5.3, A.5.7, A.5.14, A.5.15, A.8.1, A.8.3, A.8.4, A.8.7, A.8.12, A.8.15, A.8.16
MHC-06	Container und Confidential Computing	C5:2026 OPS-32 bis OPS-35	A.5.23, A.8.31
MHC-07	Multi-Tenancy-Isolation	C5:2026 OPS-30/31, PSS-10	A.8.22, A.5.23
MHC-08	Unveränderliche Backups	C5:2026 OPS-06 bis OPS-09	A.8.13, A.5.29, A.5.30, A.5.33, A.8.14
MHC-09	AI-gestütztes Security-Testing	C5 OPS-25.01AS, CRA, SSDF	A.8.8, A.8.25, A.8.26, A.8.28, A.8.29, A.8.32
MHC-10	Continuous Control Monitoring	C5 COM-03/04, NIS2-DVO 2.2	A.5.18, A.5.35, A.5.36, A.8.9
MHC-11	SOAR und Tier-1-Automation	C5 SIM-02/03, NIS2-DVO 3.5	A.5.5, A.5.24, A.5.25, A.5.26
MHC-12	Threat-Led Penetration Testing	DORA Art. 26/27, TIBER-EU	A.5.35, A.8.29
MHC-13	AI-Agent-Governance und Harness-Sicherheit	OWASP LLM Top 10 2026, ASI01–ASI10, MITRE ATLAS, ISO 42001 Annex A	A.5.9, A.5.16, A.8.27, ergänzt MHC-04

Der Katalog ist Gegenstand der Handlungsempfehlungen in Kapitel 10. Diese Hinweise sind keine Roadmap, da die konkrete Einführungsreihenfolge von ISMS-Reifephase, vorhandener Control-Landschaft und priorisierten Risikoszenarien abhängt.

9.5 Reifegrad-Stufen pro MHC

Die dreizehn MHC sind keine binären Ja-/Nein-Entscheidungen. Für jeden MHC sind drei Reifegrad-Stufen definiert, die als Soll-Ist-Vergleich im Statement of Applicability geführt werden. Eine Organisation kann je MHC unterschiedliche Stufen erreichen; die Stufen sind kumulativ (Stufe 2 setzt die Anforderungen aus Stufe 1 voraus).

Bewertungsraster: Initial = begonnener Implementierungspfad, ad hoc oder dokumentiert; Defined = standardisiert und eingeführt; Managed = automatisiert, gemessen und kontinuierlich verbessert.

Verbindliche Lesart der Stufen. Die Stufe Initial dokumentiert den begonnenen Implementierungspfad und stellt nicht zwingend die vollständige Erfüllung des jeweiligen MHC-Control-Statements dar. Erst ab Defined gilt ein MHC grundsätzlich als umgesetzt; Managed beschreibt die fortgeschrittene beziehungsweise kontinuierlich wirksame Umsetzung. Eine Organisation kann daher nicht mit Verweis auf die Stufe Initial die Umsetzung eines MHC als erreicht ausweisen – im Statement of Applicability ist in diesem Fall der Umsetzungsstand „teilweise“ mit Zielstufe und Zieltermin zu führen. Die Control-Statements der MHC-Bausteine im Implementation Guide beschreiben die Mindestanforderungen der Stufe Defined; darüber hinausgehende Anforderungen sind dort ausdrücklich als weiterführende Härting (Stufe Managed) gekennzeichnet.

MHC	Initial	Defined	Managed
MHC-01	Kryptografisches Inventar dokumentiert	PQC-Strategie verabschiedet, Hybrid-Cryptography in Pilot	Hybrid-Cryptography produktiv, jährliche Review automatisiert
MHC-02	SBOM-Generierung in einzelnen Pipelines	SBOM für alle eigenen Artefakte, Vulnerability-Korrelation	SBOM von kritischen Lieferanten, SLSA Level 3+, automatisierte Provenance-Verifikation
MHC-03	FIDO2/Passkeys für privilegierte Accounts	FIDO2-Pflicht für alle privilegierten Zugänge und alle extern erreichbaren Dienste	Passwortlose Authentisierung organisationsweit, SMS-MFA abgeschaltet
MHC-04	SPIFFE/SPIRE für privilegierte Workloads, mTLS auf 3-5 kritischsten Pfaden	Workload-Identity für alle produktiven Microservices und für produktive, unter MHC-13 fallende AI-Agenten und agentische Workflows, ZTNA für privilegierten Zugriff	Flächendeckend Workload-Identity inkl. Dev/Test, AI-Agenten mit capability-scoped Identitäten
MHC-05	MITRE ATT&CK als Detection-Framework, Coverage nicht systematisch erhoben	Abdeckung der für das eigene Threat-Profile relevanten ATT&CK-Techniken dokumentiert und priorisiert, dokumentiertes Threat-Hunting	Abdeckung durch Tests validiert (Purple Teaming, BAS), adversarial-robuste ML-Detection, kontinuierliche Nachführung
MHC-06	Container-Images signiert, Vulnerability-Scan vor Deployment	Runtime-Protection mit Admission-Controllern, dokumentierte TEE-Use-Cases	Confidential Computing produktiv für hochsensitive Workloads, Remote-Attestation automatisiert
MHC-07	Logische Tenant-Trennung dokumentiert	Tenant-spezifische Schlüssel, Netzwerk-/Compute-Isolation	Nachweisbare Trennung durch automatisierte Tests, regelmäßige Audits
MHC-08	Backups vorhanden, Restore-Tests jährlich	3-2-1-1-0-Modell umgesetzt (mindestens eine unveränderliche oder netzgetrennte Kopie), getrennte Backup-Administrationszugänge,	Eigene Administrations- und Identitätsebene für die Backup-Infrastruktur, automatisierte Integritäts- und Wiederherstellungs-Validierung, adversariale

MHC	Initial	Defined	Managed
		quartalsweise dokumentierte Restore-Tests	Löschtests, kontinuierliche Unveränderlichkeits-Überwachung
MHC-09	SAST/DAST/SCA in CI	Mindestens ein AI-gestütztes Testverfahren für geeignete Anwendungen nachweislich produktiv + Pre-Merge-Block bei High-Findings	Tägliche Vulnerability-Scans, AI-Code-Audit für Vibe-Coded Code, automatisierter KEV-Hotfix-Channel
MHC-10	Compliance-Dashboards mit manueller Pflege	Policy-as-Code in CI, Echtzeit-Compliance-Checks gegen Baselines, Alarmierung und Nachverfolgung von Abweichungen	Runtime-Enforcement, automatisierte Incident-Tickets bei Abweichung
MHC-11	Reaktions-Playbooks dokumentiert, manuelle Ausführung	SOAR mit teilautomatisierten Playbooks, MTTC < 60 Minuten	Vollautomatische Tier-1-Containment, MTTC < 10 Minuten, halbjährliche Multi-Incident-Tabletops; alternativ MDR-Service mit vertraglicher SLA
MHC-12	Jährliche Penetrationstests mit Standard-Scope	Threat-Led Tests mit Mythos-Szenarien, Purple-Team zwischen Zyklen	Kontinuierliche Adversary-Emulation, BAS-Plattform produktiv, Closure-Rate ≥ 90 % im SLA-Zeitfenster
MHC-13	Produktive AI-Agenten und agentische Workflows inventarisiert (einschließlich SaaS-, Low-Code- und No-Code-Agenten), manuelles Inventar von MCP-Servern und Extensions	Capability-scoped Identitäten + Audit-Trail produktiv, benannter Verantwortlicher und erprobter Not-Aus je Agent, dokumentiertes Bedrohungsmodell pro Agent-Use-Case, Allowlist für Extensions und MCP-Server	Vollständiger Harness-Code-Review-Prozess, automatisierte Pre-Production-Checks, Mean Time to Revoke < 5 Minuten, regelmäßige Prompt-Injection-Pen-Tests

Die Stufen sind als Audit-Hilfsmittel gedacht: Ein externer Auditor sieht auf einen Blick, wo die Organisation pro MHC steht und welcher Verbesserungspfad geplant ist. Der Soll-Stand pro MHC sollte sich aus der Risikobewertung der jeweiligen Organisation ableiten – nicht jede Organisation braucht überall Stufe Managed.

Teil IV – Handlungsempfehlungen und Ausblick

Die Teile I bis III haben die Mythos-Bedrohungslage, die Bewertung aller 93 ISO-Controls, die Lücken gegenüber anderen Frameworks und den Katalog der dreizehn Mythos-Härtungs-Controls dargelegt. Teil IV formuliert daraus abgeleitete Handlungsempfehlungen und schließt mit einer Reflexion über die Grenzen dieser Arbeit sowie mit einem Ausblick.

Die Empfehlungen sind bewusst nicht als Arbeitsplan oder Projektstruktur formuliert. Jedes ISMS hat eine eigene Reifephase und eigene Prioritäten. Ein konkreter Einführungspfad muss in der jeweiligen Organisation auf Basis der eigenen Risikobewertung, Ressourcenlage und Stakeholder-Struktur entwickelt werden.

10 Handlungsempfehlungen für die ISMS-Anpassung

10.1 Geltungsbereich und Grenzen der Empfehlungen

Die in diesem Kapitel formulierten Handlungsempfehlungen sind eine priorisierte Auswahl, kein abschließender Katalog. Sie decken fünf Themenfelder ab: Sofort-Neubewertung bestehender Controls, strukturelle Härtung, Reifung des ISMS-Prozesses, Vorstandskommunikation und Kennzahlenarbeit. Andere Themenfelder (sektorale

Spezifika, regulatorische Sonderfragen, konkrete Lieferantenbeziehungen) bleiben bewusst außen vor, weil sie nicht verallgemeinerbar sind.

Hinweis zur Nicht-Vollständigkeit: Die Empfehlungen enthalten bewusst keine Zeitangaben. Die realistische Umsetzungsdauer hängt von ISMS-Reife, Team-Stärke, Budgetlage und Tool-Landschaft ab. Die Reihenfolge innerhalb der Empfehlungen ist keine Implementierungssequenz, sondern eine thematische Gliederung.

10.2 Sofort-Neubewertung bestehender Controls

Aus der Analyse in Teil II folgt, dass die vier Controls aus Kapitel 6 (reine Reibung) und die 37 Controls aus Kapitel 5 (teilweise degradiert) in der Risikobewertung eines Mythos-exponierten ISMS neu bewertet werden müssen. Die Neubewertung erfolgt anhand der vier Bewertungskriterien aus Kapitel 3.2 und hat konkrete Folgen für das Statement of Applicability:

- Für die vier Controls A.5.25, A.5.36, A.8.8 und A.8.23 wird geprüft, ob sie in der bisherigen Form als alleinige Mitigation für Mythos-relevante Risiken im SoA geführt werden. Ist dies der Fall, wird die Eintragung um die jeweilige Ersatz-Empfehlung aus Kapitel 6 ergänzt oder korrigiert.
- Für die 37 teilweise degradierten Controls wird geprüft, welches der fünf Ergänzungsmuster aus Kapitel 5.3 jeweils bereits umgesetzt ist und wo Lücken bestehen. Ergebnis ist eine Gap-Liste je Control als Eingabe für die strukturelle Härtung.
- Für die 29 standfesten Controls aus Kapitel 4 wird die konsequente Umsetzung geprüft. Ein formal im SoA eingetragenes, aber nicht durchgängig implementiertes standfestes Control bietet unter Mythos-Bedingungen keinen Schutz. Beispiel: Ein Asset-Inventar ohne automatisierte Aktualisierung ist schnell veraltet, weshalb Mythos-spezifische Anomaliedetektion blinde Flecken hat.

Die Neubewertung erzeugt kein neues ISMS-Dokument, sondern eine Aktualisierung des bestehenden SoA mit zwei Ergänzungen: erstens Vermerken zur Mythos-Resilienz pro Control, zweitens Zuordnungen zu den dreizehn MHC aus Kapitel 9. Die Aktualisierung ist Gegenstand der nächsten Management-Review.

10.3 Strukturelle Härtung: Übernahme der MHC ins SoA

Die dreizehn MHC konkretisieren Mythos-relevante technische und organisatorische Härtungsmuster, die die ISO/IEC 27002:2022 nicht in dieser technischen Tiefe oder Verbindlichkeit vorgibt. Mehrere MHC sind dabei ausdrücklich gehärtete Ausprägungen bestehender ISO-Controls – etwa MHC-08 zu A.8.13, MHC-03 zu A.5.17 und A.8.5 oder MHC-10 zu A.5.35, A.5.36 und A.8.9; andere adressieren Anforderungen, die die ISO nur implizit über die Risikobetrachtung abdeckt. Die Übernahme in das SoA folgt einem einfachen Muster:

- Jeder MHC wird als zusätzlicher Eintrag ins SoA aufgenommen, mit Verweis auf die flankierenden ISO-Controls aus der MHC-Kompakt-Tabelle (Kap. 9.4).
- Der Umsetzungsstand wird ehrlich dokumentiert: bereits umgesetzt, teilweise umgesetzt, geplant oder nicht angegangen. Die Kategorie „geplant“ sollte mit expliziter Begründung versehen sein, warum sofortige Umsetzung nicht möglich ist.
- Für jeden MHC wird geprüft, ob er im organisatorischen Kontext überhaupt anwendbar ist. MHC-07 (Multi-Tenancy-Isolation) ist beispielsweise für Organisationen ohne Multi-Tenant-Architektur nicht relevant. Solche Nicht-Anwendbarkeit wird ebenfalls explizit vermerkt.

Die Übernahme erzeugt keine Compliance-Pflicht, weil die MHC keine normative Basis im ISO-27001-Sinne haben. Sie dokumentiert die bewusste Auseinandersetzung mit der Mythos-Bedrohungslage und gibt internen und externen Prüfern eine nachvollziehbare Grundlage. Für Organisationen, die einer Regulierung unterliegen, die einen MHC-Inhalt bereits verlangt, deckt sich die Übernahme mit einer bestehenden Pflicht. Das gilt allerdings nur im jeweiligen Anwendungsbereich: MHC-12 entspricht den TLPT-Anforderungen für diejenigen DORA-regulierten Finanzunternehmen, die gemäß Art. 26 Abs. 8 für TLPT identifiziert wurden – nicht für alle DORA-Adressaten. MHC-11 operationalisiert die Anforderungen aus NIS2 und der Durchführungsverordnung an Erkennung, Reaktion und Wiederherstellung; ein SOAR-System oder automatisiertes Tier-1-Containment wird dort nicht pauschal vorgeschrieben.

10.4 Reifung des ISMS-Prozesses

Mehrere Beobachtungen aus der Mythos-Bedrohungslage haben direkte Implikationen für die Betriebsweise des ISMS:

- Die Risikobewertung nach ISO 27005 bleibt die Basis, muss aber in kürzeren Zyklen aktualisiert werden. Was vor einem Jahr als geringe Bedrohung eingestuft war, kann heute akut sein. Eine jährliche formale Risikobewertung mit unterjährig ergänzender Überprüfung spezifischer Risiken ist angemessen.
- Das Management-Review nach ISO 27001 Kapitel 9.3 sollte die Mythos-Bedrohungslage als wiederkehrenden Agenda-Punkt aufnehmen. Grundlage können die Versionshinweise dieses Schnellhefts und einschlägige Threat-Intelligence-Berichte sein.
- Die interne Audit-Funktion sollte um Continuous-Audit-Elemente ergänzt werden (inhaltlich deckungsgleich mit MHC-10). Die Ergebnisse des automatisierten Monitorings fließen in die reduzierten, weiterhin jährlichen formalen Audits ein, ersetzen sie aber nicht.
- Die KPI-Landschaft des ISMS wird um Mythos-relevante Indikatoren ergänzt (siehe Kap. 10.6).
- Innovation-Acceleration-Governance etablieren: ein cross-funktionales Gremium aus Security, Legal, Engineering, Datenschutz und ggf. Beschaffung mit dem expliziten Mandat, defensive AI-Tools und neue Sicherheits-Controls beschleunigt zu evaluieren und produktiv zu nehmen. Ohne dieses Gremium läuft jede Mythos-Härtungsmaßnahme in die übliche Approval-Friction (Beschaffungszyklen, Legal-Review, Datenschutz-Freigabe), die unter Mythos-Bedingungen zur strukturellen Schwäche wird – die NIS2-DVO erkennt dies in Nr. 1.1 als Governance-Pflicht implizit an. Empfohlene Kadenz: zweiwöchentlich, mit definiertem Time-to-Decision-Ziel von maximal 30 Tagen für Tools mit gültigem ISO-27001/SOC-2/C5-Testat.
- Permanente VulnOps-Funktion analog DevOps aufbauen, statt Vulnerability-Management als Teilzeit-Aufgabe innerhalb des SOC zu führen. Verantwortung: kontinuierliche Discovery von Zero-Days über die gesamte Software-Estate (eigener Code, Drittsoftware, Cloud-Konfiguration), automatisierte Remediations-Pipelines, EPSS- und KEV-basierte Priorisierung. Mindestkapazität ab etwa 1.000 Mitarbeitenden oder ab dem Zeitpunkt, an dem die Patch-Last 100 oder mehr Critical-Findings pro Monat erreicht. Diese Funktion adressiert das in Mythos-Ready (CSA, SANS, OWASP) als CRITICAL eingestufte Risiko der Continuous-Vulnerability-Management-Reife.

10.5 Vorstandskommunikation und Risikodialog

Die Herausforderung besteht darin, die Bedrohung ohne Alarmismus, aber auch ohne Verharmlosung darzustellen. Folgende Punkte haben sich als nützlich erwiesen:

- Die vier Feststellungen aus Kapitel 2 eignen sich als Struktur für einen Board-Vortrag: Kollaps des Patch-Gap-Fensters, Zeitachsen-Kompression, Fragmentierung, Fähigkeits-Entkopplung. Alle vier Punkte sind empirisch belegt.
- Die Unterscheidung zwischen standfesten, teilweise degradierten und reibungsbasierten Controls vermittelt, dass nicht alles an Sicherheit versagt, sondern dass gezielt anzupassen ist.
- Die MHC aus Kapitel 9 sind als Diskussionsgrundlage für Investitionsentscheidungen geeignet. Sie liefern konkrete Bezugspunkte, die mit Framework-Referenzen belegt sind.
- Die regulatorische Einbettung sollte transparent gemacht werden. Mythos-Härtung ist in erheblichen Teilen eine Rechtspflicht, nicht nur Best Practice – abhängig jedoch von Sektor und Adressatenkreis. Für CRA-pflichtige Software-Hersteller, DORA-regulierte Finanzinstitute und NIS2-DVO-Adressaten (bestimmte Digitalanbieter, siehe Kap. 3.3) gilt sie unmittelbar; für andere NIS2-Entities ist die jeweilige nationale Umsetzung maßgeblich.
- Mögliche Entwicklung der Sorgfaltserwartung thematisieren – als strategische Einschätzung des Autors, nicht als geltende Rechtspflicht. Der EU AI Act begründet keine Pflicht zum Einsatz AI-gestützter Threat-Hunting-, Code-Scanning- oder Vulnerability-Discovery-Werkzeuge. Mit der breiten Verfügbarkeit defensiver AI-Werkzeuge ist jedoch damit zu rechnen, dass sich der Maßstab dessen verschiebt, was als angemessene Sicherheitsmaßnahme gilt. Leitungsorgane könnten künftig gegenüber Aufsicht, Versicherern und Gerichten zunehmend belegen können, welche AI-Tools für defensive Zwecke (Code-Audit, Threat-Hunting, Vulnerability-Discovery) eingesetzt werden – und die Nicht-Nutzung verfügbarer Tools begründen. Empfehlung als wiederkehrendes Board-Briefing-Element: „Welche defensiven AI-Tools

haben wir im Berichtszeitraum eingeführt? Welche evaluieren wir? Welche nutzen wir bewusst nicht – und mit welcher Begründung?" Dies adressiert das in Mythos-Ready (CSA, SANS, OWASP) explizit als HIGH eingestufte Risiko regulatorischer und Haftungsexposition.

10.6 Mythos-relevante Kennzahlen

Die folgenden Kennzahlen sind eine Auswahl, kein abschließendes KPI-Set. Sie orientieren sich an den dreizehn MHC und den vier Bewertungskriterien aus Kapitel 3.2.

Kennzahl	Beschreibung	Bezug
Mean Time to Containment (MTTC)	Durchschnittliche Zeit von Detection bis zur automatischen oder manuellen Eindämmung eines Vorfalls. Zielwert für High-Confidence-Alerts unter 10 Minuten.	MHC-11, MHC-05
Phishing-Resistant Account Coverage	Anteil aller privilegierten und extern erreichbaren In-Scope-Accounts mit phishing-resistenter MFA (Account-Coverage, nicht Anmeldungen). Zielwert privilegiert: 100 % oder dokumentierte, befristete und risikogenehmigte Ausnahme; extern erreichbar typisch mindestens 95 % für privilegierte Accounts.	MHC-03
SBOM-Coverage	Anteil der eigenen Software-Artefakte mit automatisch generiertem SBOM und Anteil kritischer Lieferanten mit vertraglich zugesicherter SBOM-Pflicht.	MHC-02
Behebungslatenz für KEV-Listings	Zeit zwischen der Verfügbarkeit einer belastbaren Behebung oder wirksamen Kompensationsmöglichkeit und deren verifizierter Umsetzung. Als wirksame Maßnahme gelten Patch, Workaround, kompensierendes Control, Isolierung oder Ersatz des betroffenen Assets. Zielwert unter 24 Stunden (Mittel). Definition in Anhang G.4.	MHC-09
Pipeline-Gate-Abdeckung und Escape Rate	Anteil der produktiven Pipelines mit Sicherheitsprüfung und Pre-Merge-Block; ergänzend Anteil der Schwachstellen, die trotz Pipeline-Prüfung in die Produktion gelangen, sowie Anteil geeigneter Anwendungen mit produktivem AI-gestütztem Testverfahren.	MHC-09
Restore-Test-Erfolgsquote	Anteil erfolgreicher quartalsweiser Restore-Tests aus immutable Backups. Zielwert 100 %.	MHC-08
Continuous-Monitoring-Coverage	Anteil der als automatisierbar klassifizierten SoA-Anforderungen, deren Umsetzung durch automatisiertes Continuous Monitoring ergänzend zu periodischen Audits geprüft wird. Bezugsgröße ist die Automatisierungs-Eignung, nicht die Gesamtzahl der Controls (Definition in Anhang G.6a/G.6b).	MHC-10
Kryptografisches-Inventar-Abdeckung	Anteil der eingesetzten kryptografischen Verfahren, die im zentralen Inventar erfasst sind und eine Prioritätseinstufung für PQC-Migration haben.	MHC-01
TLPT-Findings-Closure-Rate	Anteil der Findings aus Threat-Led Penetration Tests, die innerhalb des definierten – vertraglich, intern oder regulatorisch vorgegebenen – Behebungszeitraums geschlossen wurden.	MHC-12
Container-Policy-Konformität	Anteil produktiver Container-Workloads, die beim Deployment Signaturprüfung und Admission-Control-Richtlinien bestehen. Zielwert: 100 % in kritischen bzw. vollständig in-scope produktiven Namespaces, organisationsweit mindestens 95 % technische Coverage; verbleibende Ausnahmen sind dokumentiert, risikobewertet und genehmigt. Ursprünglicher Zielwert 100 % für kritische Namespaces.	MHC-06

Kennzahl	Beschreibung	Bezug
Agent-Inventar-Abdeckung	Anteil produktiver AI-Agenten und agentischer Workflows im Verantwortungsbereich der Organisation – einschließlich SaaS-, Low-Code- und No-Code-Agenten –, die vollständig inventarisiert sind. Zielwert 100 %. Entspricht der Reifegradstufe Initial.	MHC-13
Agent-Governance-Abdeckung	Anteil der inventarisierten produktiven AI-Agenten und agentischen Workflows, für die mindestens benannter Verantwortlicher, capability-scoped Identität, dokumentierte Berechtigungsgrenzen und manipulationssicheres Ausführungsprotokoll umgesetzt und dokumentiert sind. Zielwert 100 %. Entspricht der Reifegradstufe Defined.	MHC-13

Die zehn primären Kennzahlen decken nicht alle dreizehn MHC ab (ergänzend werden für MHC-10 die Bezugskennzahl Automatisierungs-Eignung und für MHC-13 die Governance-Abdeckung erhoben). Für die strukturellen Controls MHC-04 (Workload-Identität) und MHC-07 (Multi-Tenancy-Isolation) ist das eine bewusste Designentscheidung: Der Umsetzungsnachweis erfolgt über Implementierungs-Stage-Gates im Projektcontrolling, nicht über kontinuierliche Kennzahlen.

Gate-Kriterien (Richtwerte, je Reifegradstufe): MHC-04 — Initial: SPIFFE/SPIRE-Identitäten für ≥ 80 % der privilegierten Service-Workloads; Defined: mTLS-Pflicht für Ost-West-Verkehr — 100 % der In-Scope-Verbindungen oder dokumentierte, befristete und risikogenehmigte Ausnahmen; ≥ 95 % gilt ausdrücklich nur als Übergangsschwelle und nicht als vollständige Defined-Erfüllung; Managed: attestierungsgestützte Zugriffspolicies produktiv. MHC-07 — je Stufe ein bestandener, dokumentierter Mandantentrennungstest (Initial: jährlich intern; Defined: je Major-Release; Managed: zusätzlich durch unabhängige Dritte). Die Gate-Erfüllung wird im Projektcontrolling als Ja/Nein-Nachweis geführt.

10.7 Zusammenfassung der Empfehlungen

Die fünf Empfehlungsfelder – Sofort-Neubewertung, strukturelle Härtung, ISMS-Reifung, Vorstandskommunikation und Kennzahlenarbeit – sind voneinander unabhängig adressierbar; Organisationen mit begrenzten Ressourcen können einzelne Felder priorisieren. Ihre verbindende Klammer ist die Rücknahme der Annahme, dass Angreiferaufwand eine verlässliche Schutzgröße sei: Schutzwirkung ist stattdessen in harten Barrieren zu verankern – kryptografisch, architektonisch, aggregationsresistent oder automatisierungsbasiert.

11 Reflexion und Ausblick

11.1 Grenzen dieser Arbeit

Die Grenzen von MRIS gelten wie in Kap. 1.4 (Position im ISMS-Stack) und Kap. 3.5 (Bedrohungs-Scope) beschrieben: eine Momentaufnahme statt eines Dauerzustands, kein Ersatz für eine organisationsspezifische Risikoanalyse und kein Vollständigkeitsanspruch des MHC-Katalogs. Die Analyse stützt sich ausschließlich auf öffentlich verfügbare Primärquellen; interne Incident-Daten einzelner Organisationen sind nicht Teil der Bewertungsgrundlage.

11.2 Was dieses Schnellheft nicht leistet

MRIS ist kein Werkzeug zur Zertifizierung oder Compliance-Prüfung, ersetzt keine Hersteller-, Produktauswahl- oder Architektur-Entscheidung, ist kein juristisches Dokument und keine Threat-Intelligence-Quelle; es arbeitet auf der Ebene von Kategorien und Prinzipien.

11.3 Erwartete Weiterentwicklungen

Steigende Agenten-Fähigkeit: Die Fortschrittsgeschwindigkeit impliziert, dass der Anteil taktischer Angriffsarbeit, den Modelle autonom übernehmen können, weiter zunehmen wird. Einstufungen als teilweise degradiert können sich in Richtung reine Reibung verschieben.

Aufbau von Verteidigungs-AI: Die gleichen Fähigkeiten sind zunehmend auf der Verteidigungsseite verfügbar (MHC-05, MHC-09, MHC-11). Das wird die Asymmetrie teilweise wieder schließen, bringt aber neue

Herausforderungen: Adversarial-Robustness der Verteidigungs-AI, Governance der AI-gestützten Sicherheits-Funktionen, neue Schulungs-Anforderungen.

Regulatorische Verdichtung: Die jüngsten Entwicklungen – C5:2026, NIS2-DVO, CRA, DORA – sind Teil einer erkennbaren Verdichtung der Cybersecurity-Regulierung in Europa. Diese Verdichtung wird sich fortsetzen, insbesondere im Bereich der AI-System-Governance (AI Act), Produktsicherheit (CRA) und sektorspezifischer Regulierung.

Pflege- und Review-Zusage. MRIS wird quartalsweise überprüft. Außerplanmäßige Reviews werden ausgelöst durch: (a) General Availability oder wesentliche Fähigkeitssprünge der Mythos-Modellklasse; (b) neue Releases von MITRE ATT&CK oder ATLAS; (c) Veröffentlichung der angekündigten BSI-Kreuzreferenztafel zu C5:2026; (d) neue regulatorische Vorgaben mit Control-Bezug (z. B. NIS2-DVO-Ergänzungen, CRA-Durchführungsrechtsakte).

Trigger-Status Juli 2026: Trigger (a) ist mit der Multi-Vendor-Entwicklung und der allgemeinen Verfügbarkeit der Mythos-Klasse ausgelöst und mit Version 1.8 abgearbeitet (Kap. 2.5). Trigger (b) bis (d) sind zum 19. Juli 2026 nicht ausgelöst bzw. in laufender Beobachtung (ATT&CK v19 ist Referenzstand; die BSI-Kreuzreferenztafel zu C5:2026 ist angekündigt, aber nicht veröffentlicht).

11.4 Bedrohungshorizont

Vier Entwicklungslinien sind für kommende Versionen kartiert, aber noch nicht control-wirksam bewertet:

- Agent-zu-Agent-Protokolle (A2A): Delegationsketten zwischen Agenten schaffen neue Vertrauens- und Angriffsflächen (Task-Injection über Zwischenagenten, transitive Berechtigungen); betrifft perspektivisch MHC-04 und MHC-13.
- MCP-Ökosystem als Lieferkette: Tool-Server, Manifeste und Registries werden zur Software-Lieferkette der Agenten; die Herkunfts- und Signaturlogik aus MHC-02 ist auf Tool-Server zu übertragen.
- Agentische Angriffe auf OT/ICS: Die Übertragung der Exploit-Ökonomie auf cyber-physische Systeme hat Safety-Folgen und erfordert eine eigene Bewertungsachse jenseits von Annex A.
- Konvergenz PQC × agentische Kryptoanalyse: Automatisierte Analyse von Protokoll- und Implementierungsfehlern verkürzt die Zeit, in der „harvest now, decrypt later“ wirtschaftlich ist; das erhöht den Druck auf MHC-01.

Diese Punkte sind Kandidaten für eine Version 2.0; ihre Bewertung folgt der Review-Kadenz aus Kap. 11.3.

11.5 Schlussbemerkung

MRIS liefert keinen neuen theoretischen Rahmen, keine eigene Methodik und keine Compliance-Aussage, sondern eine Arbeitshilfe: Es stellt vorhandene Frameworks gegen die aktuelle Bedrohungslage, macht Lücken sichtbar und benennt konkrete Anhaltspunkte für deren Schließung. Klassische Controls werden dabei nicht obsolet – sie werden in spezifischen Wirksamkeits-Dimensionen geschwächt und sind gezielt zu härten, nicht zu ersetzen; die MHC ergänzen sie, sie treten nicht an ihre Stelle.

Teil V – Ergänzungs-Layer: Weiterführende Frameworks

Teil V ergänzt die primären Frameworks aus Kap. 3.3 um fünf weitere, die für eine vertiefte Mythos-Härtungs-Arbeit nützlich sind. Sie werden in Teil II und III dort zitiert, wo sie greifen; die folgende Tabelle fasst zusammen, was jedes liefert und wann es relevant ist. Wer sich auf die primären Frameworks konzentriert, kann Teil V überspringen.

Framework	Was es liefert	Wann relevant
BSI IT-Grundschutz-Kompendium	Technisch konkrete ISMS-Bausteine (Basis-/Standard-/Hoch-Anforderungen), breiter als die abstrakte ISO 27002	Öffentlicher Sektor, KRITIS, Grundschutz-Zertifizierung; als zweite Primärreferenz neben C5:2026
MITRE ATT&CK und D3FEND	Operatives Vokabular für Detection (ATT&CK) und Verteidigung (D3FEND); Kill-Chain-Struktur bildet fragmentierte Angriffe zusammenhängend ab	Konkretisierung von SIEM/Detection (C5 OPS-13, NIS2-DVO 3.2) und DORA-TLPT-

Framework	Was es liefert	Wann relevant
		Szenarien; Grundlage von MHC-05 und MHC-12
ISO/IEC 42001	Management-System für AI mit Annex-A-Controls (Risikomanagement, Datenqualität, Lifecycle, Transparenz, Robustheit)	Governance defensiver AI-Komponenten und Absicherung eigener Modelle; besonders für SaaS-Anbieter mit hohem AI-Anteil
CIS Controls v8	18 priorisierte Control-Kategorien mit Implementation-Groups (IG1–IG3) und konkreten Safeguards	Pragmatische Priorisierungshilfe für Organisationen in früher Reifephase (IG1 als Einstieg)
OWASP ASVS und SAMM	ASVS: operativer Verifikationskatalog für sichere Software (L1–L3); SAMM: Reifegradmodell für die Secure-SDLC-Organisation	Direkt in CI integrierbare Grundlage für SAST/DAST (MHC-09); Messlatte für Software-Hersteller mit CRA-Pflicht

Keines dieser Frameworks ersetzt die Primärbasis aus Kap. 3.3; die Auswahl richtet sich nach Branche, Größe und strategischer Ausrichtung der Organisation.

C5:2026 ↔ IT-Grundschutz-Korrespondenz (Orientierungshilfe). Zur Abbildung der in Teil II und III genannten C5-Referenzen auf typische Grundschutz-Bausteine:

C5:2026-Domäne	IT-Grundschutz-Bausteine (typisch)
C5 OIS (Organisation of Information Security)	ISMS-Schicht, ORP
C5 HR (Personnel)	ORP.2
C5 AM (Asset Management)	CON.10, OPS.1.1.3
C5 PS (Physical Security)	INF.1 bis INF.7
C5 OPS (Operations)	OPS.1.1.1 bis OPS.1.2.5, DER
C5 IAM (Identity and Access Management)	ORP.4
C5 CRY (Cryptography and Key Management)	CON.1, CON.6
C5 COS (Communication Security)	NET.1, NET.3
C5 DEV (Development)	CON.8
C5 SSO (Service Providers and Suppliers)	OPS.2, CON.14
C5 SIM (Security Incident Management)	DER.2.1 bis DER.2.3
C5 BCM (Business Continuity Management)	DER.4
C5 COM (Compliance)	ISMS-Schicht
C5 PSS (Product Safety and Security)	APP-Bausteine je nach System

Die Liste ist Orientierungshilfe, kein verbindliches Eins-zu-eins-Mapping.

Anhänge – Arbeitsmaterialien

Die Anhänge bündeln die in den Teilen I bis V erarbeiteten Inhalte in einer für die praktische ISMS-Arbeit direkt nutzbaren Form. Anhang A liefert die vollständige Bewertungsmatrix aller 93 Controls mit Kategorie-Einstufung und MHC-Zuordnung. Anhang B ergänzt diese Matrix um eine Cross-Reference zwischen den primären Frameworks. Anhang C stellt den MHC-Katalog als Standalone-Arbeitsblatt bereit. Anhang D listet Indikatoren- und Monitoring-Quellen. Anhang E liefert das RACI-Modell für die MHC-Umsetzung. Anhang F beschreibt die Risikobewertungs-Brücke nach ISO/IEC 27005. Anhang G standardisiert die KPI-Definitionen aus Kap. 10.6 für audit-fähige, organisationsübergreifend reproduzierbare Messung.

Anhang A – Bewertungsmatrix aller 93 Controls

Die folgende Tabelle listet alle 93 Controls aus ISO/IEC 27002:2022 in ISO-Nummerierung. Für jedes Control werden angegeben: ID, Titel (gekürzt), Mythos-Kategorie (S = Standfest, T = Teilweise degradiert, R = Reine Reibung, N = Nicht betroffen) und – sofern zutreffend – die flankierenden Mythos-Härtungs-Controls aus Kapitel 9.

Die MHC-Zuordnungen in dieser Matrix sind mit der Kompakt-Übersicht in Kapitel 9.4 konsistent.

ID	Control-Titel	Kat.	Flankierend MHC
A.5.1	Informationssicherheitsrichtlinien	N	-
A.5.2	Informationssicherheitsrollen und -verantwortlichkeiten	N	-
A.5.3	Aufgabentrennung	T	MHC-05
A.5.4	Verantwortlichkeiten der Leitung	N	-
A.5.5	Kontakt mit Behörden	T	MHC-11
A.5.6	Kontakt mit speziellen Interessengruppen	N	-
A.5.7	Threat Intelligence	T	MHC-05
A.5.8	Informationssicherheit im Projektmanagement	N	-
A.5.9	Inventar von Informationen und zugehörigen Werten	S	MHC-13
A.5.10	Akzeptable Nutzung von Informationen	N	-
A.5.11	Rückgabe von Werten	N	-
A.5.12	Klassifizierung von Informationen	S	-
A.5.13	Kennzeichnung von Informationen	N	-
A.5.14	Informationsübertragung	T	MHC-05
A.5.15	Zugriffskontrolle	T	MHC-04, MHC-05
A.5.16	Identitätsmanagement	S	MHC-04, MHC-13
A.5.17	Authentisierungsinformation	T	MHC-03
A.5.18	Zugangsrechte	T	MHC-10
A.5.19	Informationssicherheit in Lieferantenbeziehungen	T	MHC-02
A.5.20	Adressierung in Lieferantenvereinbarungen	T	MHC-02
A.5.21	Management der ICT-Lieferkette	T	MHC-02
A.5.22	Überwachung Lieferantendienstleistungen	T	MHC-02
A.5.23	Cloud-Dienste	T	MHC-06, MHC-07
A.5.24	Planung IR-Management	T	MHC-11
A.5.25	Beurteilung und Entscheidung zu Ereignissen	R	MHC-11
A.5.26	Reaktion auf Vorfälle	T	MHC-11
A.5.27	Lernen aus Vorfällen	S	-
A.5.28	Sammlung von Beweismaterial	S	-
A.5.29	Informationssicherheit während einer Störung	T	MHC-08
A.5.30	ICT-Bereitschaft für BCM	S	MHC-08
A.5.31	Identifikation rechtlicher Anforderungen	N	-
A.5.32	Rechte an geistigem Eigentum	N	-
A.5.33	Schutz von Aufzeichnungen	T	MHC-08
A.5.34	Datenschutz und PII	T	-
A.5.35	Unabhängige Überprüfung	T	MHC-10, MHC-12

ID	Control-Titel	Kat.	Flankierend MHC
A.5.36	Einhaltung von Richtlinien	R	MHC-10
A.5.37	Dokumentierte Betriebsprozesse	N	-
A.6.1	Überprüfung (Screening)	N	-
A.6.2	Beschäftigungsbedingungen	N	-
A.6.3	Bewusstsein, Aus- und Weiterbildung	T	-
A.6.4	Disziplinarverfahren	N	-
A.6.5	Verantwortlichkeiten nach Beendigung	S	-
A.6.6	Vertraulichkeitsvereinbarungen	N	-
A.6.7	Remote-Arbeit	T	MHC-03, MHC-04
A.6.8	Meldung von Sicherheitsereignissen	T	-
A.7.1	Physische Sicherheitsperimeter	S	-
A.7.2	Physische Zutrittskontrolle	S	-
A.7.3	Sicherung Büros und Einrichtungen	S	-
A.7.4	Physische Sicherheitsüberwachung	S	-
A.7.5	Schutz vor physischen und Umwelt-Bedrohungen	N	-
A.7.6	Arbeiten in Sicherheitsbereichen	S	-
A.7.7	Clear Desk und Screen Lock	N	-
A.7.8	Aufstellung und Schutz von Geräten	N	-
A.7.9	Sicherheit von Werten außerhalb der Räumlichkeiten	T	-
A.7.10	Speichermedien	S	-
A.7.11	Versorgungseinrichtungen	N	-
A.7.12	Verkabelungssicherheit	N	-
A.7.13	Wartung von Geräten	N	-
A.7.14	Sichere Entsorgung oder Wiederverwendung	S	-
A.8.1	Benutzer-Endpunktgeräte	T	MHC-03, MHC-05
A.8.2	Privilegierte Zugriffsrechte	S	MHC-04
A.8.3	Einschränkung des Informationszugangs	T	MHC-05
A.8.4	Zugang zu Quellcode	T	MHC-02, MHC-05
A.8.5	Sichere Authentisierung	T	MHC-03
A.8.6	Kapazitätsmanagement	N	-
A.8.7	Schutz vor Malware	T	MHC-05
A.8.8	Verwaltung technischer Schwachstellen	R	MHC-09
A.8.9	Konfigurationsmanagement	S	MHC-10
A.8.10	Informationslöschung	S	-
A.8.11	Datenmaskierung	S	-
A.8.12	Verhinderung von Datenabfluss	T	MHC-05
A.8.13	Informationssicherung (Backup)	S	MHC-08
A.8.14	Redundanz	S	MHC-08
A.8.15	Protokollierung (Logging)	S	MHC-05
A.8.16	Überwachungsaktivitäten	T	MHC-05

ID	Control-Titel	Kat.	Flankierend MHC
A.8.17	Uhrenzeit-Synchronisation	S	-
A.8.18	Privilegierte Hilfsprogramme	S	-
A.8.19	Software-Installation	S	-
A.8.20	Netzwerksicherheit	T	MHC-04
A.8.21	Sicherheit von Netzwerkdiensten	T	MHC-04
A.8.22	Trennung in Netzwerken	T	MHC-04, MHC-07
A.8.23	Webfilterung	R	MHC-03
A.8.24	Verwendung von Kryptographie	S	MHC-01
A.8.25	Sicherer Entwicklungslebenszyklus	T	MHC-09
A.8.26	Sicherheitsanforderungen für Anwendungen	T	MHC-09
A.8.27	Prinzipien der sicheren Systemarchitektur	S	MHC-13
A.8.28	Sicheres Programmieren	T	MHC-09
A.8.29	Sicherheitsprüfung in Entwicklung und Abnahme	S	MHC-09, MHC-12
A.8.30	Ausgelagerte Entwicklung	T	MHC-02
A.8.31	Trennung von Umgebungen	S	MHC-06
A.8.32	Änderungsmanagement	T	MHC-09
A.8.33	Testinformationen	S	-
A.8.34	Schutz während Audittests	N	-

Legende: S = Standfest (Kapitel 4), T = Teilweise degradiert (Kapitel 5), R = Reine Reibung (Kapitel 6), N = Nicht betroffen (Kapitel 7). Die MHC-Bezüge verweisen auf Mythos-Härtungs-Controls aus Kapitel 9.

Anhang B – Framework-Mapping für Kernthemen

Die folgende Tabelle fasst für zwölf zentrale Themenfelder der Informationssicherheit zusammen, wie die primären Frameworks sie jeweils adressieren. Sie ist Orientierung, nicht vollständiges Mapping.

Themenfeld	ISO 27002	C5:2026	NIST CSF	DORA	CRA	NIS2-DVO
Asset-Management	A.5.9, A.5.12	AM-02, AM-03, AM-04, AM-09	ID.AM	Art. 8	Annex I	Nr. 12
Identität und Zugriff	A.5.15, A.5.16, A.8.2	IAM-01, IAM-06, IAM-08	PR.AA	Art. 9	Annex I	Nr. 11
Kryptografie	A.8.24	CRY-01 bis CRY-19	PR.DS	Art. 9(4)(e)	Annex I II	Nr. 9
Logging und Monitoring	A.8.15, A.8.16	OPS-10 bis OPS-17	DE.CM	Art. 10	-	Nr. 3.2
Incident Response	A.5.24, A.5.26	SIM-01 bis SIM-06	RS.	Art. 17, 19	-	Nr. 3.1 bis 3.6
Business Continuity	A.5.29, A.5.30, A.8.13	BCM-01 bis BCM-04, OPS-06 bis OPS-09	RC.	Art. 11, 12	-	Nr. 4

Themenfeld	ISO 27002	C5:2026	NIST CSF	DORA	CRA	NIS2-DVO
Supply Chain	A.5.19 bis A.5.23	SSO-01 bis SSO-08, DEV-13	GV.SC	Art. 28-30	Annex I II	Nr. 5
Entwicklung	A.8.25 bis A.8.29	DEV-01 bis DEV-15	PR.IR	Art. 8(3)	Annex I II	Nr. 6
Physische Sicherheit	A.7.1 bis A.7.14	PS-01 bis PS-08	PR.IR	Art. 9(4)	-	Nr. 13
Risikomanagement	A.5.1 bis A.5.8	OIS-07, OIS-08, OIS-09	GV.RM, ID.RA	Art. 6, 7	-	Nr. 2
Compliance	A.5.31 bis A.5.36	COM-01 bis COM-04	GV.OC	Art. 6(5)	Art. 13, 14	Nr. 2.2, 2.3
Cloud-Nutzung	A.5.23	GC-01 bis GC-06, OPS-30 bis OPS-35	GV.SC-07	Art. 28	-	Nr. 5

Leere Zellen (-) bedeuten, dass das Framework für das jeweilige Themenfeld keine eigene spezifische Anforderung formuliert. Die NIST-CSF-Spalte nennt die jeweiligen Kategorien aus CSF 2.0; die konkrete Umsetzung erfolgt über SP 800-53.

Hinweis zum Geltungsbereich der NIS2-DVO: Die in der Spalte „NIS2-DVO“ zitierten Nummern entstammen der Durchführungsverordnung (EU) 2024/2690, die laut ihrem Artikel 1 ausschließlich für bestimmte Digitalanbieter gilt (siehe Kap. 3.3). Für andere NIS2-Adressaten sind die Meldefristen aus der NIS2-Richtlinie Art. 23 Abs. 4 in Verbindung mit der jeweiligen nationalen Umsetzung maßgeblich.

Anhang C – MHC-Katalog als Standalone-Arbeitsblatt

Die folgende Übersicht stellt den Mythos-Härtungs-Control-Katalog aus Kapitel 9 als Standalone-Arbeitsblatt bereit. Anhang C kann aus dem Gesamtdokument extrahiert und als eigenständiges Arbeitsblatt für das Statement of Applicability verwendet werden. Die Spalte „Status“ ist bewusst leer gelassen und kann organisationsspezifisch ausgefüllt werden (z. B. bereits umgesetzt, teilweise umgesetzt, geplant, nicht anwendbar).

MHC	Titel	Framework-Basis	ISO-Anbindung	Status
MHC-01	Post-Quantum-Strategie und kryptografisches Inventar	C5:2026 CRY-01.01AC	A.8.24	
MHC-02	SBOM und Build-Provenance	CRA, C5:2026 DEV-13, SLSA	A.5.19, A.5.20, A.5.21, A.5.22, A.8.4, A.8.30	
MHC-03	Phishing-resistente MFA	NIST 800-63B, FIDO2	A.5.17, A.6.7, A.8.1, A.8.5, A.8.23	
MHC-04	Workload-Identität und Zero Trust	NIST 800-207, SPIFFE	A.5.15, A.5.16, A.6.7, A.8.2, A.8.20, A.8.21, A.8.22	
MHC-05	Verhaltensbasierte Detection	MITRE ATT&CK, C5 OPS-13	A.5.3, A.5.7, A.5.14, A.5.15, A.8.1, A.8.3, A.8.4, A.8.7, A.8.12, A.8.15, A.8.16	
MHC-06	Container und Confidential Computing	C5:2026 OPS-32 bis OPS-35	A.5.23, A.8.31	
MHC-07	Multi-Tenancy-Isolation	C5:2026 OPS-30/31	A.5.23, A.8.22	

MHC	Titel	Framework-Basis	ISO-Anbindung	Status
MHC-08	Unveränderliche Backups	C5:2026 OPS-06 bis OPS-09	A.5.29, A.5.30, A.5.33, A.8.13, A.8.14	
MHC-09	AI-gestütztes Security-Testing	C5 OPS-25.01AS, CRA, SSDF	A.8.8, A.8.25, A.8.26, A.8.28, A.8.29, A.8.32	
MHC-10	Continuous Control Monitoring	C5 COM-03/04, NIS2-DVO 2.2	A.5.18, A.5.35, A.5.36, A.8.9	
MHC-11	SOAR und Tier-1-Automation	C5 SIM-02/03, NIS2-DVO 3.5	A.5.5, A.5.24, A.5.25, A.5.26	
MHC-12	Threat-Led Penetration Testing	DORA Art. 26/27, TIBER-EU	A.5.35, A.8.29	
MHC-13	AI-Agent-Governance und Harness-Sicherheit	OWASP LLM Top 10 2026, ASI01–ASI10, MITRE ATLAS, ISO 42001 Annex A	A.5.9, A.5.16, A.8.27, ergänzt MHC-04	

Anhang D – Indikatoren und Monitoring-Quellen

Die folgende Tabelle listet öffentliche Indikatoren-Quellen auf, die für die laufende Beobachtung der Mythos-Bedrohungslage und für die Detektion Mythos-relevanter Ereignisse nützlich sind. Die Liste ist eine Auswahl, kein vollständiger Katalog.

Quelle	Beschreibung	URL / Zugang
CVE – Common Vulnerabilities and Exposures	Zentrale Schwachstellenbibliothek von MITRE, Publikation über NVD.	cve.org / nvd.nist.gov
EUVD – European Union Vulnerability Database	ENISA-geführte Schwachstellendatenbank für die EU.	euvd.enisa.europa.eu
KEV – Known Exploited Vulnerabilities Catalog	CISA-Liste aktiv ausgenutzter Schwachstellen, Patch-Priorisierung.	cisa.gov/known-exploited-vulnerabilities-catalog
EPSS – Exploit Prediction Scoring System	Wahrscheinlichkeit der Exploit-Nutzung einer CVE in 30 Tagen.	first.org/epss
OSV – Open Source Vulnerabilities	Vulnerability-Feed für Open-Source-Abhängigkeiten.	osv.dev
MITRE ATT&CK	Taxonomie von Angreifertaktiken und -techniken.	attack.mitre.org
CISA Advisories	US-Behördenbenachrichtigungen zu akuten Bedrohungen.	cisa.gov/news-events/cybersecurity-advisories
BSI-Warnungen und Lageberichte	Deutsche Sicherheitswarnungen und jährlicher Lagebericht.	bsi.bund.de
CERT-EU Advisories	Warnungen für EU-Institutionen mit breiter Relevanz.	cert.europa.eu
Anthropic Threat Intelligence	Veröffentlichungen zu AI-gestützten Angriffskampagnen.	anthropic.com/news
GitHub Security Advisories	Zentrale Stelle für Open-Source-Vulnerability-Reports.	github.com/advisories
FIRST.org	Forum of Incident Response and Security Teams.	first.org

Anhang E – RACI-Modell für die MHC-Umsetzung

Die folgende Tabelle ordnet jedem der dreizehn Mythos-Härtungs-Controls die typischen Verantwortlichkeiten in einem mittleren Unternehmen zu. Die Zuordnung ist als Vorschlag zu verstehen und ist organisationsspezifisch zu adjustieren – insbesondere in Organisationen ohne dedizierten AI Governance Lead oder ohne Vorstands-Reporting-Linie zu Cybersecurity.

Die folgende RACI ist eine Beispielbelegung und keine normative Governance-Vorgabe. Sie zeigt eine in mittleren Unternehmen verbreitete Zuordnung. Je nach Organisationsform kann die Accountability für einzelne MHC bei anderen Rollen liegen – etwa beim CTO für Workload-Architektur (MHC-04), beim Head of Engineering für Secure Development (MHC-09) oder bei einer Produkt-Security-Funktion für AI-Agent-Governance (MHC-13).

Bewertungsraster: R = Responsible (führt aus), A = Accountable (rechenschaftspflichtig, eine Person je Zeile), C = Consulted (wird vorher befragt), I = Informed (wird nachträglich informiert).

Rollen-Spalten: CISO = Chief Information Security Officer, ISMS = ISMS-Manager / Beauftragter, AI-Gov = AI Governance Lead, IT = IT-Betrieb / Plattform-Engineering, Dev = Entwicklung, SOC = Security Operations Center / IR-Funktion, Recht = Rechts- und Compliance-Funktion, Vorstand = Geschäftsleitung / Aufsichtsorgan.

MHC	Titel (kurz)	CISO	ISMS	AI-Gov	IT	Dev	SOC	Recht	Vorstand
MHC-01	Post-Quantum-Strategie / Crypto-Inventar	A	R	C	R	C	I	C	I
MHC-02	SBOM und Build-Provenance	A	C	I	C	R	I	C	I
MHC-03	Phishing-resistente MFA	A	C	I	R	C	I	I	I
MHC-04	Workload-Identität / Zero Trust	A	C	C	R	R	C	I	I
MHC-05	Verhaltensbasierte Detection	A	I	C	C	I	R	I	I
MHC-06	Container / Confidential Computing	A	I	C	R	R	C	I	I
MHC-07	Multi-Tenancy-Isolation	A	C	I	R	R	C	C	I
MHC-08	Unveränderliche Backups	A	C	I	R	I	C	I	I
MHC-09	AI-gestütztes Security-Testing	A	I	C	C	R	C	I	I
MHC-10	Continuous Control Monitoring	A	R	C	R	C	C	I	I
MHC-11	SOAR / Tier-1-Automation	A	C	C	C	I	R	C	I
MHC-12	Threat-Led Penetration Testing	A	C	C	C	C	R	C	I
MHC-13	AI-Agent-Governance	A	C	R	C	R	C	C	C

Hinweise zur Anwendung

- Pro Zeile genau ein A: Accountability ist nicht delegierbar und wird nicht aufgeteilt. In der Vorlage ist überall der CISO als Accountable hinterlegt; in größeren Organisationen kann die Accountability für einzelne MHC an Domänen-Verantwortliche delegiert werden (z. B. CTO für MHC-04, Head of Engineering für MHC-09).
- Mehrere R sind möglich, wenn die Ausführung mehrere Domänen betrifft. Beispiel MHC-04: IT setzt SPIFFE/SPIRE auf, Dev migriert die Service-Identitäten – beide sind Responsible.
- MHC-13 ist die einzige Zeile, in der der Vorstand als Consulted geführt wird. Begründung: Die Governance produktivsetzender AI-Agenten betrifft Risiko-Akzeptanzfragen oberhalb der CISO-Mandatsschwelle (Standard of Care, siehe Kap. 10.5).
- Externe Dienstleister (MDR, MSSP, TLPT-Provider, GRC-Beratung) erscheinen nicht als eigene Spalte. Sie sind über die jeweils Responsible-Rolle steuerungsmäßig eingebunden, nicht über eine eigene RACI-Spalte.

Anhang F – Risikobewertungs-Brücke nach ISO/IEC 27005

Dieser Anhang beschreibt, wie eine MRIS-Bewertung in den Risikomanagement-Prozess nach ISO/IEC 27005:2022 einfließt. Er schließt die in Kap. 1.4 explizit gemachte Lücke „kein eigener Risikomanagement-Prozess“ durch eine konkrete Schnittstellen-Definition zwischen MRIS-Wirksamkeits-Layer und ISMS-Risikoprozess.

F.1 Eingabe: MRIS-Bewertungsergebnis je Control

Aus der Anwendung von MRIS auf das Statement of Applicability ergeben sich pro Control vier Datenpunkte:

- Mythos-Kategorie (Standfest / Teilweise degradiert / Reine Reibung / Nicht betroffen) – aus den Kapiteln 4 bis 7.
- Identifizierte Schwächen entlang der vier Bewertungskriterien aus Kap. 3.2 (Angreifergeduld, Zeitkompression, Fähigkeitsentkopplung, Aggregationsresistenz).
- Flankierende Mythos-Härtungs-Controls aus Anhang A.
- Aktueller Reifegrad pro flankierendem MHC nach Kap. 9.5 (Initial / Defined / Managed).

F.2 Verarbeitungsschritte im Risikoprozess

Schritt 1: Risiko-Reassessment (ISO/IEC 27005:2022 Kap. 7.3)

Für jedes Control mit Mythos-Kategorie „Teilweise degradiert“ oder „Reine Reibung“ wird die zugeordnete Risiko-Position im Risikoregister erneut bewertet. Konkret:

Der methodische Ansatzpunkt ist die bisher angenommene Wirksamkeit der betroffenen Controls (Control Effectiveness), nicht unmittelbar die inhärente Eintrittswahrscheinlichkeit. Die Residual-Likelihood beziehungsweise das Restrisiko wird anschließend anhand der organisationsspezifischen Risikomethodik erneut bestimmt.

MRIS-Befund	Auswirkung auf die Risikobewertung
Standfest	Keine automatische Anpassung.
Teilweise degradiert	Angenommene Control Effectiveness neu bewerten; Restrisiko nach eigener Methodik neu bestimmen.
Reine Reibung	Nicht als alleinige wirksame Mitigation ansetzen; ergänzende Maßnahmen ausweisen oder Restrisiko explizit akzeptieren.
Nicht betroffen	Keine Mythos-spezifische Anpassung.

Die Consequence-Achse bleibt unverändert, sofern nicht das Schadensszenario selbst (z. B. Massen-Exfiltration durch agentische Tools) eine höhere Consequence-Stufe nahelegt.

Bewusst wird keine pauschale Stufen-Anhebung („+1“, „+2“) vorgegeben: Organisationen arbeiten mit unterschiedlichen Modellen (3×3, 4×4, 5×5, 10×10, quantitativ nach FAIR), in denen eine feste Stufenzahl keine übertragbare Bedeutung hätte. Zu vermeiden ist außerdem Doppelzählung, wenn mehrere degradierte Controls dasselbe Risiko behandeln – in diesem Fall wird die Wirksamkeit des Mitigationsbündels gemeinsam bewertet, nicht je Control einzeln.

Schritt 2: Behandlungsoptionen (ISO/IEC 27005:2022 Kap. 8.1)

Für jedes neu bewertete Risiko werden die vier ISO-27005-Behandlungsoptionen geprüft:

- **Modifikation:** Übernahme der flankierenden MHC ins SoA (siehe Kap. 10.3) – die Standardantwort für die meisten Mythos-Risiken.
- **Übernahme:** Akzeptanz erhöhter Restrisiken durch das Risk-Owner-Gremium, sofern Behandlungskosten unverhältnismäßig sind. Begründung muss schriftlich dokumentiert sein.
- **Vermeidung:** Außerbetriebnahme der betroffenen Funktion oder des Service. In Mythos-Kontexten selten umsetzbar, kann aber bei Legacy-Systemen ohne Patch-Pfad relevant sein.
- **Teilung:** Versicherungslösungen (Cyber Risk Insurance), Outsourcing an spezialisierte Provider mit eigener Mythos-Härtung.

Schritt 3: Restrisiko-Bewertung und Akzeptanz (ISO/IEC 27005:2022 Kap. 8.6)

Nach Behandlungsentscheidung wird das Restrisiko bewertet und vom Risk Owner formal akzeptiert. Bei Akzeptanz von Restrisiken oberhalb der Akzeptanzschwelle ist eine zeitlich befristete Akzeptanzentscheidung mit Re-Review-Datum (typisch sechs Monate) zu treffen.

F.3 Beispielhafte Anwendung

Beispiel: Risiko R-042 „Massen-Exfiltration aus Code-Repositories durch kompromittierten Developer-Account“. MRIS-Bewertung der relevanten Controls:

Control	MRIS-Befund	Konsequenz
A.5.17 Authentisierung	Teilweise degradiert (Aggregationsblindheit, Credential-Kompromittierbarkeit). Flankierend MHC-03.	Angenommene Wirksamkeit gegen Credential-Phishing deutlich reduziert; ohne MHC-03 (phishing-resistente MFA, FIDO2) mindestens auf Reifegrad Defined wirkt das Control nur eingeschränkt.
A.8.4 Quellcode-Zugang	Teilweise degradiert (Aggregationsblindheit). Flankierend MHC-02 und MHC-05.	Angenommene Wirksamkeit gegen fragmentierten Massenzugriff reduziert; erst in Kombination mit Repository-Anomaly-Detection (MHC-05) und SBOM-Pflicht (MHC-02) auf Reifegrad Managed wird sie wiederhergestellt.
A.8.16 Monitoring	Teilweise degradiert (Aggregationsblindheit). Flankierend MHC-05.	Angenommene Detektionswirksamkeit reduziert; erst mit verhaltensbasierter Detection auf Reifegrad Managed werden fragmentierte Mass-Clones erkannt.

Resultierende Behandlungsentscheidung: Nach Neubewertung der Control Effectiveness liegt das Restrisiko oberhalb der Akzeptanzschwelle. Übernahme von MHC-02, MHC-03 und MHC-05 ins SoA mit Ziel-Reifegrad Managed innerhalb von zwölf Monaten. Bis zur Erreichung wird ein temporäres Restrisiko vom Risk Owner mit Re-Review nach sechs Monaten akzeptiert.

F.4 Ausgabe in das ISMS

Die Risikobewertungs-Brücke produziert vier Artefakte für das ISMS:

- Aktualisierte Einträge im Risikoregister mit MRIS-Bezug pro Risikoposition.
- Aktualisiertes Statement of Applicability mit zusätzlichen MHC-Einträgen und Verweis auf die jeweilige Mythos-Kategorie der flankierten ISO-Controls.
- Behandlungsplan mit Reifegrad-Zielwerten je MHC und definierten Meilensteinen.
- Restrisiko-Akzeptanzentscheidungen mit Re-Review-Daten.

Diese vier Artefakte schließen die Brücke zwischen MRIS-Bewertungs-Layer und ISO-27005-Risikoprozess. Sie ersetzen weder das Risikoregister noch den ISMS-Prozess, sondern liefern dem Risikoprozess die unter Mythos-Bedingungen notwendigen Eingaben.

Anhang G – KPI-Definitionen und Mess-Standardisierung

Dieser Anhang standardisiert die zehn primären Mythos-Kennzahlen aus Kap. 10.6 so, – für MHC-10 wird zusätzlich die Bezugskennzahl Automatisierungs-Eignung (G.6a) und für MHC-13 die ergänzende Governance-Kennzahl (G.10b) erhoben – dass sie zwischen Organisationen und über Zeit reproduzierbar erhoben werden können. Ohne diese Standardisierung führen zwei Organisationen denselben KPI mit identischem Zielwert und produzieren dennoch nicht vergleichbare Ergebnisse.

KPI	Definition, Messpunkte und Datenquelle	Zielwert	Frequenz · Verantwortung	Bezug
G.1 Mean Time to Containment (MTTC)	Zeit vom ersten High-Confidence-Alert (SIEM/EDR, Confidence $\geq 80\%$ oder High-Severity-Flag) bis zur bestätigten ersten Containment-Aktion; Mittel und Median. Quelle: SOAR-Audit-Log (primär), SIEM, IR-Ticket-System.	Mittel < 10 Min; Median < 5 Min	Monatlich · SOC-Lead / CISO	MHC-11, MHC-05
G.2 Phishing-Resistant Account Coverage	Anzahl phishing-resistent abgesicherter In-Scope-Accounts geteilt durch die Anzahl aller In-Scope-Accounts, in Prozent; erfasst werden privilegierte und extern erreichbare Accounts mit AAL2+ phishing-resistenter MFA (FIDO2/WebAuthn/Passkey/CTAP2). Quelle: IdP (Entra ID, Okta, Ping) + CMDB der privilegierten Accounts.	Privilegiert 100 % oder dokumentierte, befristete und risikogenehmigte Ausnahme; extern erreichbar $\geq 95\%$; SMS-Neuanlage = 0 %	Monatlich · IAM-Lead / ISMS	MHC-03
G.3 SBOM-Coverage	(a) eigene Release-Artefakte mit automatischem SBOM in CI; (b) kritische Lieferanten (Tier 1/2) mit vertraglicher SBOM-Pflicht. Quelle: CI-Reports (Snyk, Dependency-Track); Vertragsdatenbank.	(a) $\geq 95\%$ (12 M); (b) $\geq 80\%$ (18 M)	Quartalsweise · Head of Eng. / CISO	MHC-02
G.4 Behebungslatenz für KEV-Listings	Zeit von der Verfügbarkeit einer belastbaren Behebung bis zur nachgewiesenen Behebung oder wirksamen Kompensation auf allen betroffenen Hosts; Mittel und 95-Perzentil. Als Behebung zählen Patch, Workaround, kompensierendes Control, Trennung vom Netz oder Ablösung.	Mittel < 24 h; P95 < 72 h (exponiert)	Pro Eintrag / monatlich · VulnOps-Lead / CISO	MHC-09, A.8.8

KPI	Definition, Messpunkte und Datenquelle	Zielwert	Frequenz · Verantwortung	Bezug
	Quelle: KEV-Katalog und Herstellerinformation (Start), Patch-/Vulnerability-Management (Stopp).			
G.5 Restore-Test-Erfolgsquote	Anteil erfolgreicher Restores aus immutable Backups innerhalb RTO mit fehlerfreien Daten. Quelle: Backup-Audit-Log + Test-Protokolle IT-Betrieb. Mindestfrequenz quartalsweise je kritischem System.	100 %; jeder Fehlschlag ist ein Control-Failure und wird als Finding mit Ursachenanalyse und Frist behandelt, Eskalation risikobasiert	Quartalsweise · IT-Betrieb / ISMS	MHC-08, A.8.13
G.6a Automatisierungseignung	Anzahl der als automatisiert prüfbar klassifizierten Anforderungen geteilt durch die Anzahl der bewerteten SoA-Anforderungen, in Prozent. Die Klassifikation wird dokumentiert und jährlich überprüft. Quelle: GRC-Plattform.	Kein Zielwert; Bezugsgröße für G.6b	Jährlich · ISMS-Manager	MHC-10
G.6b Continuous-Monitoring-Coverage	Anzahl der aktiv automatisiert überwachten Anforderungen geteilt durch die Anzahl der als automatisierbar klassifizierten Anforderungen, in Prozent (Prüfung mindestens täglich, mit Alarmierung bei Abweichung, ergänzend zu periodischen Audits). Quelle: OPA/Sentinel-Repo, GRC-Plattform.	≥ 60 % (12 M); ≥ 80 % (24 M)	Quartalsweise · ISMS-Manager	MHC-10
G.7 Kryptografisches-Inventar-Abdeckung	Anteil produktiver Krypto-Verfahren im Crypto-Inventar mit PQC-Prioritätsstufe; Mindestattribute Algorithmus, Schlüssellänge, Library+Version, Einsatzkontext. Quelle: CMDB-Crypto-Tabelle, Crypto-Discovery-Scanner.	≥ 80 % (12 M); ≥ 95 % (24 M)	Halbjährlich · Architecture-Lead / CISO	MHC-01
G.8 TLPT-Findings-Closure-Rate	Anteil TLPT-Findings, im SLA-Fenster per Re-Test bestätigt geschlossen; Standard-SLA Critical 30 / High 60 / Medium 90 / Low 180 Tage; die Standard-SLA-Fristen sind MRIS-Richtwerte und gelten, soweit keine strengere interne, vertragliche oder regulatorische Frist vorgeht. Quelle: TLPT-Report, Findings-Tracker, Re-Test.	≥ 90 % aggregiert	Quartalsweise · CISO / Risk Owner	MHC-12
G.9 Container-Policy-Konformität	Anteil produktiver Container mit gültiger Image-Signatur, bestandenem Scan und ohne	Krit. Namespaces 100 %; gesamt ≥	Monatlich · Plattform-Lead / ISMS	MHC-06

KPI	Definition, Messpunkte und Datenquelle	Zielwert	Frequenz · Verantwortung	Bezug
	unbefristete Privileg-Ausnahme. Quelle: Admission-Controller (Kyverno, OPA Gatekeeper), Sigstore/Cosign, Registry-Scans.	95 %; ungeprüfte Ausnahmen 0		
G.10a Agent-Inventar-Abdeckung	Anzahl vollständig inventarisierter produktiver AI-Agenten und agentischer Workflows im Verantwortungsbereich der Organisation (selbst entwickelte, SaaS-, Low-Code- und No-Code-Agenten) geteilt durch deren Gesamtzahl, in Prozent; Shadow-AI-Funde (A.5.9) zählen zum Nenner. Quelle: Agent-Inventar, Discovery-Scans, CMDB.	100 % (nicht inventarisierte werden deaktiviert)	Monatlich · AI-Gov-Lead / ISMS	MHC-13
G.10b Agent-Governance-Abdeckung	Anteil der inventarisierten Agenten und Workflows mit benanntem Verantwortlichen, capability-scoped Identität, dokumentierten Berechtigungsgrenzen und manipulationssicherem Protokoll (A.8.15). Quelle: Agent-Inventar, IdP/Workload-Identitäten, Ausführungsprotokolle.	100 % (nicht inventarisierte werden deaktiviert)	Monatlich · AI-Gov-Lead / ISMS	MHC-13

Standardisierungs-Hinweise. Die KPIs sind organisationsübergreifend reproduzierbar, sofern die genannten Datenquellen verfügbar sind und die Scope-Definitionen konsistent angewendet werden. Für Vergleichbarkeit über Zeit werden KPI-Definitionen bei wesentlichen Änderungen versioniert und Berechnungsänderungen als Bruch in der Zeitreihe ausgewiesen. Bei Branchen-Benchmarks oder ISAC-Sharing wird die Definitions-Version mit dem Wert übermittelt und sektorspezifische Anpassungen (z. B. höhere Schwellwerte für DORA-regulierte Institute) werden deklariert.

Glossar

Das folgende Glossar erklärt zentrale Begriffe des Schnellhefts in kompakter Form. Begriffe sind alphabetisch sortiert.

Begriff	Erklärung
Aggregationsresistenz	Eigenschaft eines Controls, auch dann zu greifen, wenn ein Angriff in viele einzeln legitim erscheinende Mikroschritte zerlegt wird. Eines der vier Bewertungskriterien (Kap. 3.2).
Angreifergeduld	Bewertungskriterium: Beruht die Schutzwirkung darauf, dass ein Angriff für einen Gegner zu mühsam ist? Unter Mythos strukturell geschwächt.
Crypto-Agility	Fähigkeit eines Systems, kryptografische Mechanismen kurzfristig zu ändern, etwa bei Kompromittierung eines Algorithmus oder bei PQC-Migration.
Fähigkeitsentkopplung	Bewertungskriterium: Hängt die Likelihood-Annahme von einer historischen Korrelation zwischen Akteur und Fähigkeit ab, die durch Mythos aufgehoben wurde?
FIDO2 / WebAuthn	Authentisierungsstandards der FIDO-Alliance für phishing-resistente MFA. Verankern die Authentisierung kryptografisch an der Origin-Domain.
Mythos	In diesem Schnellheft Klassenbegriff: agentische Frontier-Modelle mit autonomer Schwachstellen-Entdeckung und Exploit-Kettenbau, unabhängig vom Anbieter (u. a. Claude Mythos/Fable 5, OpenAI Daybreak mit GPT-5.5-Cyber, Microsoft MDASH). Ursprünglich Bezeichnung des Referenzmodells Claude Mythos Preview (Kap. 2.5).
MHC – Mythos-Härtungs-Control	In diesem Schnellheft eingeführte Kategorie eines Zusatz-Ccontrols, das eine Mythos-relevante Wirksamkeitslücke schließt, die die ISO 27002:2022 nur implizit abdeckt. Dreizehn MHC in Kapitel 9.
AI – Artificial Intelligence	Im deutschen Sprachgebrauch auch „KI“. Dieses Werk verwendet durchgängig „AI“, um an die MHC-Bezeichnungen und die zitierten Frameworks anzuschließen.
PQC – Post-Quantum-Cryptography	Kryptografische Verfahren, die gegen Angreifer mit Quantum-Computer-Fähigkeiten resistent sind. NIST-Standardisierung (ML-KEM, ML-DSA, SLH-DSA).
Reine Reibung	Kategorie für Controls, deren klassische bzw. verbreitete Implementierungsform gegenüber Mythos-Angreifern keine hinreichende eigenständige Schutzwirkung mehr bietet, weil sie auf begrenzter Angreiferkapazität oder auf menschlich handhabbaren Reaktionszeiten beruht. Vier Controls in Kapitel 6.
SBOM – Software Bill of Materials	Maschinenlesbare Liste aller Softwarekomponenten eines Artefakts. Formate: SPDX, CycloneDX. Pflicht durch CRA Annex I Teil II.
SLSA – Supply-chain Levels for Software Artifacts	Rahmenwerk für Build-Provenance mit vier Reifestufen. Level 3+ für kritische Build-Pfade empfohlen.
SOAR	Security Orchestration, Automation and Response. Plattform-Kategorie für automatisierte Incident-Response-Workflows.
SPIFFE / SPIRE	Standard und Referenz-Implementierung für Workload-Identitäten. Basis für Zero-Trust-Architekturen in Cloud-native-Umgebungen.
Standfest	Kategorie für Controls, deren Schutzwirkung aus einer harten Barriere stammt, die auch unter Mythos fortbesteht. 29 Controls in Kapitel 4.
Store-now-decrypt-later	Bedrohungsmodell, bei dem heute exfiltrierte verschlüsselte Daten aufbewahrt werden, um sie nach Verfügbarkeit von Quantum-Computing zu entschlüsseln.

Begriff	Erklärung
Teilweise degradiert	Kategorie für Controls, deren Schutzwirkung erhalten bleibt, deren ursprünglich angenommene Stärke aber unter Mythos sinkt. 37 Controls in Kapitel 5.
TEE – Trusted Execution Environment	Hardware-gestützter, isolierter Ausführungsbereich für vertrauliche Datenverarbeitung. Basis für Confidential Computing.
TLPT – Threat-Led Penetration Testing	Penetrationstest nach realen Bedrohungsszenarien, verankert in DORA Art. 26/27 und im TIBER-EU-Framework.
UEBA – User and Entity Behavior Analytics	Detektionskategorie auf Basis von Verhaltens-Baselines und ML-gestützter Anomalieerkennung.
Zeitkompression	Bewertungskriterium: Setzt das Control menschliche Reaktionszeit in einer Kette voraus, in der der Angreifer autonom operiert?
Zero Trust	Architekturprinzip, nach dem keinem Akteur, keinem Gerät und keinem Netz implizit vertraut wird. Jede Transaktion wird explizit verifiziert. NIST SP 800-207.

Quellenverzeichnis

Normen und regulatorische Quellen

- ISO/IEC 27001:2022 – Information security management systems – Requirements.
- ISO/IEC 27002:2022 – Information security controls.
- ISO/IEC 27005:2022 – Information security risk management.
- ISO/IEC 42001:2023 – Artificial intelligence management system.
- ISO/IEC 20889 – Privacy enhancing data de-identification terminology.
- ISO/IEC 27017 – Code of practice for information security controls for cloud services.
- ISO/IEC 27037 – Guidelines for identification, collection, acquisition and preservation of digital evidence.
- ISO/IEC 27701 – Privacy information management system.
- ISO 22301 – Business continuity management systems.
- BSI Cloud Computing Compliance Criteria Catalogue (C5:2026), März 2026.
- BSI IT-Grundschutz-Kompendium (aktuelle Fassung).
- BSI TR-02102 Kryptographische Verfahren: Empfehlungen und Schlüssellängen.
- BSI TR-03183 Cyber-Resilienz-Anforderungen an Hersteller und Produkte, Teil 1 (Allgemeine Anforderungen) und Teil 2 (SBOM). bsi.bund.de (Abruf: 19.07.2026).
- Richtlinie (EU) 2022/2555 (NIS2-Richtlinie) über Maßnahmen für ein hohes gemeinsames Cybersicherheitsniveau.
- Durchführungsverordnung (EU) 2024/2690 der Kommission vom 17. Oktober 2024 (NIS2-DVO).
- Verordnung (EU) 2022/2554 (DORA) über die digitale operationelle Resilienz im Finanzsektor.
- Verordnung über horizontale Cybersicherheitsanforderungen (EU Cyber Resilience Act, CRA).
- DSGVO – Verordnung (EU) 2016/679.
- AI Act – Verordnung (EU) 2024/1689 zur Festlegung harmonisierter Vorschriften für künstliche Intelligenz.

NIST-Publikationen

- NIST Cybersecurity Framework 2.0 (CSF 2.0), Februar 2024.
- NIST SP 800-34 Rev. 1 – Contingency Planning Guide for Federal Information Systems.
- NIST SP 800-40 Rev. 4 – Guide to Enterprise Patch Management Planning.
- NIST SP 800-46 Rev. 2 – Guide to Enterprise Telework, Remote Access, and BYOD Security.

- NIST SP 800-53 Rev. 5 – Security and Privacy Controls for Information Systems and Organizations.
- NIST SP 800-60 – Guide for Mapping Types of Information and Information Systems to Security Categories.
- NIST SP 800-61 Revision 3 – Incident Response Recommendations and Considerations for Cybersecurity Risk Management.
- NIST SP 800-63B-4 (Juli 2025) – Digital Identity Guidelines: Authentication and Authenticator Management.
- NIST SP 800-86 – Guide to Integrating Forensic Techniques into Incident Response.
- NIST SP 800-88 Rev. 1 – Guidelines for Media Sanitization.
- NIST SP 800-92 – Guide to Computer Security Log Management.
- NIST SP 800-94 – Guide to Intrusion Detection and Prevention Systems (2007; Entwurf Revision 1 zurückgezogen).
- NIST SP 800-124 Rev. 2 – Guidelines for Managing the Security of Mobile Devices.
- NIST SP 800-150 – Guide to Cyber Threat Information Sharing.
- NIST SP 800-160 Vol. 1 – Systems Security Engineering.
- NIST SP 800-167 – Guide to Application Whitelisting.
- NIST SP 800-175B – Guideline for Using Cryptographic Standards.
- NIST SP 800-188 – De-Identifying Government Datasets.
- NIST SP 800-190 – Application Container Security Guide.
- NIST SP 800-207 – Zero Trust Architecture.
- NIST SP 800-218 – Secure Software Development Framework (SSDF) v1.1.
- NIST SP 1800-38 (Draft) – Migration to Post-Quantum Cryptography.
- FIPS 140-3 – Security Requirements for Cryptographic Modules.
- NIST NCCoE Project – Agent Identity Framework (2026; angekündigt, Primärquelle bei Redaktionsschluss nicht öffentlich verifiziert).

Threat Intelligence und Industrie-Publikationen

- Anthropic – Detecting and countering misuse of AI: August 2025 Threat Intelligence Report.
- Anthropic – Disrupting the first reported AI-orchestrated cyber espionage campaign (GTG-1002), November 2025.
- Anthropic – Preparing your security program for AI-accelerated offense (Glasswing-Post), April 2026. claude.com/blog/preparing-your-security-program-for-ai-accelerated-offense (Abruf: 19.07.2026).
- Anthropic – Claude Mythos Preview, April 2026. anthropic.com/news (Abruf: 19.07.2026).
- OpenAI – Daybreak: AI-gestützte Schwachstellen-Entdeckung und Patch-Validierung (GPT-5.5, Codex Security), 11. Mai 2026. Sekundäranalysen: thehackernews.com, futurumgroup.com (Abruf: 19.07.2026).
- Microsoft – MDASH (Multi-Model Agentic Scanning Harness), 12. Mai 2026: 16 Windows-Schwachstellen im Mai-Patchday, CyberGym 88,45 %. Sekundäranalysen: csoonline.com, thehackernews.com (Abruf: 19.07.2026).
- Anthropic – Claude Fable 5 und Claude Mythos 5, 9. Juni 2026. anthropic.com/news/claude-fable-5-mythos-5 (Abruf: 19.07.2026).
- Anthropic – Redeploying Claude Fable 5, 1. Juli 2026. anthropic.com/news/redeploying-fable-5 (Abruf: 19.07.2026).
- Anthropic – Project Glasswing: An initial update, 22. Mai 2026. anthropic.com/research/glasswing-initial-update (Abruf: 19.07.2026).
- Cloud Security Alliance – Project Glasswing: AI Discovery Outpaces Open Source Patching Capacity, Juni 2026. cloudsecurityalliance.org (Abruf: 19.07.2026).
- Cloud Security Alliance, Cloud Key Management Working Group – Post-Quantum Cryptography Key Management. Entwurf, 13. August 2026. cloudsecurityalliance.org (Abruf: 24.08.2026).
- Cloud Security Alliance, SANS Institute, [un]prompted, OWASP Gen AI Security Project (Hrsg.) – The „AI Vulnerability Storm": Building a „Mythos-ready" Security Program. Expedited Strategy Briefing. Öffentliche

Freigabe 14. April 2026; zitierte Fassung v0.95 vom 18. April 2026. cloudsecurityalliance.org (Abruf: 19.07.2026).

- OWASP Gen AI Security Project – Top 10 for LLM Applications 2026 (LLM01–LLM10) und Agentic Security Initiative (ASI01–ASI10).
- MITRE ATLAS – Adversarial Threat Landscape for Artificial-Intelligence Systems, atlas.mitre.org (Abruf: 19.07.2026).
- NIST AI Risk Management Framework 1.0, AI RMF 1.0.
- MITRE ATT&CK Enterprise Matrix – attack.mitre.org (Abruf: 19.07.2026).
- MITRE D3FEND – d3fend.mitre.org (Abruf: 19.07.2026).
- OWASP Application Security Verification Standard (ASVS) – owasp.org/www-project-application-security-verification-standard (Abruf: 19.07.2026).
- OWASP Software Assurance Maturity Model (SAMM) – owaspsamm.org (Abruf: 19.07.2026).
- Center for Internet Security (CIS) – CIS Critical Security Controls v8.
- CSA Cloud Controls Matrix v4.
- SLSA – Supply-chain Levels for Software Artifacts, slsa.dev.
- ENISA – European Cybersecurity Certification Scheme for Cloud Services (EUCS).
- ENISA – EU Vulnerability Database (EUVD), euvd.enisa.europa.eu (Abruf: 19.07.2026).
- FIDO Alliance – WebAuthn Specifications.
- TIBER-EU – European framework for Threat Intelligence-based Ethical Red Teaming, Europäische Zentralbank.

Online-Quellen zuletzt geprüft: 19. Juli 2026 (PQC-Schlüsselmanagement: 24. August 2026; MITRE ATT&CK und ATLAS: 27. August 2026).

Versionshinweise

Dieses Schnellheft wird als lebendes Dokument geführt. Die folgende Tabelle dient als Änderungsprotokoll für Folgeversionen. Die vorliegende Fassung ist Version 2.0.

Version	Datum	Änderungen
1.0	April 2026	Erstveröffentlichung. Bewertung aller 93 Controls der ISO/IEC 27002:2022 gegen die Mythos-Bedrohungslage. Framework-Quervergleich zu BSI C5:2026, NIST, DORA, CRA und NIS2 mit Durchführungsverordnung. Katalog der zwölf Mythos-Härtungs-Controls (MHC-01 bis MHC-12).
1.1	April 2026	Überarbeitete Fassung zur Audit-Standfestigkeit. Korrektur falscher NIS2-DVO-Verweise (Nr. 1.2.5 existiert nicht, korrigiert auf 11.2.2 Buchst. a; Fristen nicht in DVO 3.3, sondern in NIS2-Richtlinie Art. 23 Abs. 4; Fehlinterpretation 3.3.2; Phasenzahl 3.5 von vier auf drei korrigiert). Ergänzung Nr. 6.7 (Netzsicherheit) für A.8.20–A.8.22. Präzisierung Geltungsbereich der NIS2-DVO. Umklassifizierung A.7.7 von „Reine Reibung“ nach „Nicht betroffen“ (neue Zählung 29/37/4/23). Schärfung Kap. 3.1.3 zur Auflösung des Widerspruchs zwischen Definition und Anwendung. Kennzeichnung der Zeitfaktor-Ergänzung als Autoren-Interpretation zur ISO/IEC 27005. Einordnung BSI C5 als Prüfkatalog. Konsistenz Kap. 9.4 vs. Anhang A. Ergänzung Anhänge A–G. Vereinheitlichung Rechtschreibung (ß) und Typografie.
1.2	April 2026	Konkretisierung der Härtungsempfehlungen für Audit-Standfestigkeit. Tool-Klassen explizit benannt (EDR, SAST/DAST/SCA, MDR, SOAR, ZTNA, EASM). Schwellwerte und Mengengerüste ergänzt (MTTC < 10 Min, ATT&CK-Coverage ≥ 60 %, Egress-Anomalie $2\sigma/3\sigma$, EDR-Confidence ≥ 80 %, mindestens zwölf Threat-Hunts pro Jahr). Reifegrad-Pfade in drei Stufen (Initial/Defined/Managed) für alle zwölf MHC in Kap. 9.5. AI-Agenten als Insider-Risiko in MHC-04 adressiert (capability-scoped Identitäten). Vibe-Coded vulnerable Code in MHC-09 als sekundäre Bedrohung adressiert. Living-off-the-

Version	Datum	Änderungen
		Land und Beacon-less C2 als Detection-Schwerpunkte in MHC-05. Praktikabilitätshinweise (MDR-Alternative, Hyperscaler-Realität, BAS-Plattformen) ergänzt. Methodik-Verweise (PEAK-Framework, TaHITI, DeTT&CT, STRIDE) konkretisiert. Mythos-spezifische User-Reporting-Indikatoren in A.6.8.
1.3	April 2026	Abdeckungserweiterung gegenüber dem CSA/SANS/OWASP-Bericht „Building a Mythos-ready Security Program“ (April 2026). Neues MHC-13 (AI-Agent-Governance und Harness-Sicherheit) schließt das in Mythos-Ready als CRITICAL eingestufte Risiko der unmanagten AI-Agent-Angriffsfläche: Harness-Audit, Blast-Radius-Limits, Human-Override, Supply-Chain-Inventar für MCP-Server/Extensions/Agentic Skills, Pre-Production-Checks. A.5.9 um Shadow-AI-Inventar erweitert (Browser-Plug-ins, IDE-Extensions, MCP-Server). Kap. 10.4 um Innovation-Acceleration-Governance (cross-funktionales Gremium mit 30-Tage-Decision-Ziel) und permanente VulnOps-Funktion erweitert. Kap. 10.5 um Standard-of-Care-Verschiebung durch EU AI Act und Board-Briefing-Element ergänzt. Reifegrad-Tabelle Kap. 9.5 um MHC-13 erweitert; Bewertungsmatrix (Anhang A) und MHC-Standalone-Arbeitsblatt (Anhang C) entsprechend aktualisiert. Mythos-Ready-Strategiepapier im Vorwort und Kap. 3.4 als zentrale Industrie-Konsens-Referenz verortet (Status zwischen Threat-Intelligence-Bericht und normativer Grundlage). Damit sind 12 von 13 in Mythos-Ready genannten Risiken auf Reifegrad „Managed“ und 1 auf „Defined“ abgebildet; keine Lücken mehr.
1.4	April 2026	Schließung der durch Audit-Bewertung („Vollständigkeit = mittel“) aufgezeigten Strukturlücken. Neues Kap. 1.4 „Position im ISMS-Stack und Grenzen des Schnellhefts“ verortet MRIS explizit als Wirksamkeits- und Realitäts-Layer auf einem ISMS-Fundament; macht klar, was MRIS leistet (Wirksamkeits-Bewertung, Lückenanalyse, MHC-Katalog, Operationalisierung) und was bewusst nicht (vollständiges ISMS, eigener Risikomanagement-Prozess, PDCA-System, Compliance-Mapping-Werkzeug, organisationsspezifische Policy-Hierarchie). Neuer Anhang E „RACI-Modell für die MHC-Umsetzung“ mit acht Rollen (CISO, ISMS, AI-Gov, IT, Dev, SOC, Recht, Vorstand) für alle dreizehn MHC. Neuer Anhang F „Risikobewertungs-Brücke nach ISO/IEC 27005“ als Schnittstellen-Definition zwischen MRIS und ISMS-Risikoprozess (Likelihood-Anhebung, Behandlungsoptionen, Restrisiko-Akzeptanz). Neuer Anhang G „KPI-Definitionen und Mess-Standardisierung“ mit vollständiger Standardisierung der acht Mythos-Kennzahlen (Datenquelle, Uhr-Start/-Stopp, Berechnungsformel, Reporting-Frequenz, Messverantwortung). Vorwort kompakt auf eine Seite reduziert; Faktenfokus, weniger Prosa.
1.5	April 2026	Schärfung der Klassifikations-Sprache und der Wirksamkeitsaussagen nach kritischer Audit-Bewertung. Klassische Controls werden klar als „selektiv degradiert“ bezeichnet, nicht als „obsolet“ – Korrekturabsätze in Kap. 9.1, Kap. 4.3 und Kap. 11.4 stellen die Priorität struktureller Controls vor neuen MHC heraus. Neues Kap. 3.1.5 zur Kontextabhängigkeit der Klassifikation (gleiches Control kann gegen unterschiedliche Angriffsvektoren in unterschiedliche Kategorien fallen, Beispiel A.5.17). Aggregationsresistenz in Kap. 3.2 operationalisiert (SIEM-Korrelation, UEBA, Graph-Analysen, Kill-Chain-Tracking oder strukturelle Unteilbarkeit). Bedrohungs-Scope in Kap. 3.5 explizit auf Gen-AI-beschleunigte Angriffe begrenzt. MHC-05 um Realismus-Hinweis ergänzt (Detection ist nicht Prevention; Low-and-slow, getarnte Agenten, Adversarial-ML-Evasion als Restrisiko). MHC-11 um Härtung der Automation-Schicht erweitert (signierte Trigger, Audit-Trail, Rate-Limits, Human-in-the-Loop für Aktionen mit hohem Blast-Radius). MHC-13 um Reifegrad-Realismus ergänzt (operative Umsetzungsreife liegt aktuell bei den meisten Organisationen auf Stufe Initial; 12 bis 24 Monate Übergangszeit; phasenweise Priorisierung Inventarisierung → Audit-Trail → Capability-Scoping).

Version	Datum	Änderungen
1.6	Juni 2026	Verweis-Aktualisierung ohne inhaltliche Neubewertung der Controls. Kapitel 9.2: kryptografische Standards auf die finalisierten FIPS 203/204/205 sowie das zur Standardisierung ausgewählte HQC umgestellt (MHC-01); NIST SP 800-94 als Fassung 2007 mit zurückgezogenem Revisions-Entwurf eingeordnet (MHC-05); MITRE-ATLAS-IDs in MHC-13 berichtigt (AML.T0047 → AML.T0053 für LLM Plugin Compromise, Ergänzung der agentischen Techniken AML.T0086/T0110) und OWASP-LLM-Nummerierung bereinigt (LLM08 „Excessive Permissions“ entfällt; abgedeckt durch LLM06 Excessive Agency). Aktualisierte Referenzen: NIST SP 800-63B Rev. 4 (MHC-03), CycloneDX/SPDX-Standardisierung, BSI TR-03183-2, CRA-Zeitleiste und VEX/CSAF (MHC-02), Sigstore/Admission-Controller/TEE-Beispiele (MHC-06), ISO/IEC 27017 (MHC-07), NIST SP 800-218A (MHC-09), NIST SP 800-137/137A und OSCAL (MHC-10), NIST SP 800-61 Rev. 3 (MHC-11), DORA-TLPT-RTS (EU) 2025/1190 und TIBER-EU/TIBER-DE (MHC-12). Quellenstand der Verweise: Juni 2026.
1.7	Juli 2026	Konsistenz-Release nach externer Doppelbewertung (Auditor/CISO): Mapping Kap. 9.4 ↔ Anhänge A/C vereinheitlicht (Anhang A führend; A.8.4, A.8.27, A.5.9); Anhangsverweise auf E/G berichtigt; MHC-Zählung durchgängig „dreizehn“; zwei neue Kennzahlen (Container-Policy-Konformität, Agent-Inventar-Abdeckung) in Kap. 10.6 und Anhang G; Prozesshüllen-Präzisierung zu Reibungs-Controls (Kap. 3.1.3, A.8.8); Review-Trigger in Kap. 11.3; Richtwert-Kennzeichnung unbelegter Zahlen; Lizenz auf CC BY-SA 4.0.
1.8	Juli 2026	Multi-Vendor-Update (Reaktion auf Review-Trigger a): neues Kap. 2.5 (Daybreak, MDASH, GA der Mythos-Klasse, Offenlegungswellen); „Mythos“ als Klassenbegriff; MHC-02/-09/-12 multi-systemisch geschärft; neues Kap. 11.4 Bedrohungshorizont; Gate-Kriterien MHC-04/-07 in Kap. 10.6; Referenzbereinigung (Mythos-Ready-Datierung, TR-03183, Abrufdatum je Webquelle); maschinenlesbare Fassung (YAML) referenziert.
1.9	August 2026	Externe Evidenz-Aktualisierung (Pass-the-Passkey-Angriffsklasse, SpecterOps, August 2026; klassische Implementierungs-/Endpoint-Schwäche, keine Control-Umstufung). MHC-03 um Attestation-Enforcement und Wirksamkeitsgrenze ergänzt; A.8.15 um Vertraulichkeitsdimension der Protokolle (CVE-2026-34348); MHC-05 um Detection-Signale für WebAuthn-/Passkey-Missbrauch. MHC-13-Referenzen auf OWASP Top 10 for LLM Applications 2026 gepinnt und Nummerierung korrigiert (Excessive Agency LLM06→LLM03; 9.2, 9.4, Anhang C, Vorwort, Quellenliste).
1.10	August 2026	Straffung für Handlichkeit (Umfang deutlich reduziert) ohne Substanzverlust am Kern. Kap. 9.2: Härtungsempfehlungen der 13 MHC auf Kernidee und Mythos-Wirkung reduziert, operative Umsetzung per Verweis (MHC-ID) an den Implementation Guide abgegeben; Mythos-Befund und Framework-Vergleich unverändert. Teil V (Ergänzungs-Frameworks) und Anhang G (KPI) von Prosa auf Tabellen verdichtet. Grenzen- und These-Dopplungen konsolidiert (Kap. 3.5, 11.1/11.2), Schlussbemerkung 11.5 entkernt, Zusammenfassungen 6.3/10.7 und Kap. 7 gestrafft, Zitatdichte der Standfest-Framework-Vergleiche reduziert (Framework-Mapping erhalten, Glossen und Erläuterungen entfernt). Kern unverändert: Kap. 5/6, Anhang A, MHC-Befunde.
1.11	August 2026	Umsetzung des externen Fachreviews (Daniel Heldt) zu Version 1.8. Korrektur des Risiko-Begriffs in Kap. 1.2 gegen ISO/IEC 27005:2022 (Risiko = „Auswirkung von Unsicherheit auf Ziele“, 3.1.3; Wahrscheinlichkeit × Schaden = Risikoniveau, 3.1.15; Motivation/Fähigkeit/Ausnutzbarkeit präziser in Kap. 7 bzw. Anhang A.2.3 verortet). Zentraler Claim und Lückenanalyse weicher gefasst (explizit gefordert vs. implizit von ISO abgedeckt); Kapazitätsbarriere-Aussage entschärft (physische/Hardware-Angriffe ausgenommen). PQC-Lücke ehrlich eingeordnet (quantengetrieben, nicht Mythos-beschleunigt; Krypto-

Version	Datum	Änderungen
		Agilität mit mehreren Trust-Anchor). Continuous Monitoring „zusätzlich zu" statt „statt" periodischer Audits. CRA als Quelle der Meldepflichten und deren Anwendbarkeit ergänzt. Rahmung „Wirksamkeit der Maßnahme, nicht abstraktes Control" (Kap. 3.2); Begriffsklärung Rahmenwerk MRIS vs. Schnellheft. Tortendiagramm der Control-Verteilung. Dank an Daniel Heldt.
1.12	August 2026	PQC-Schlüsselmanagement vertieft (Quelle: CSA Cloud Key Management Working Group, Entwurf 08/2026; eigene Formulierung mit Attribution). Kap. 8.2 um vier schlüsselmanagement-spezifische Facetten ergänzt: Trust-Now-Forge-Later (Ernte signierter Artefakte), Downgrade-Resistenz hybrider Aushandlung, Migrationsreihenfolge in der Schlüsselhierarchie (Wurzel/KEK zuerst), quantensichere Zeitstempel. MHC-01-Kernidee entsprechend präzisiert; CSA-Quelle im Quellenverzeichnis ergänzt. Operative Umsetzung im Implementation Guide, MHC-01.
1.13	August 2026	Präzisierung der ISO-Anbindung in Kap. 9.2: Wo die Befund-Zelle nur die primären Ankercontrols nennt, ist dies jetzt als Auswahl gekennzeichnet („flankiert primär") mit Verweis auf die vollständige Zuordnung in Anhang A. Sprachliche Korrekturen aus dem Fachreview (Heldt): Rollenbezug auf die ISMS-Verantwortung geöffnet, AI-Systeme nicht mehr als Urheber von Phishing und Schadsoftware bezeichnet, Geltungsbereich der CRA-SBOM-Pflicht auf Hersteller von Produkten mit digitalen Elementen eingegrenzt, Mehrfaktor-Kriterium bei gerätegebundenen Passkeys präzisiert, Workload-Identität bei Erstnennung erläutert, Kap. 2.5 wartungsarm betitelt.
2.0	August 2026	Qualitäts-Release nach externem Qualitäts-, CISO- und Audit-Review. Methodik: Trennung von struktureller Schutzfähigkeit und organisationsspezifischer Implementierungsreife (Kap. 3.2); Kategorie „Reine Reibung" implementierungsbezogen definiert (Kap. 3.1.3); Reifegrad-Lesart verbindlich geregelt — Initial ist kein Umsetzungsnachweis (Kap. 9.5); Risikobewertungs-Brücke ohne pauschale Stufen-Anhebung, stattdessen Neubewertung der Control Effectiveness nach eigener Methodik (Anhang F). Korrekturen: Anteil der härtingsbedürftigen Controls auf 44 % (41 von 93) berichtigt; NIS2-DVO-Verweise für Authentisierung auf Nr. 11.6/11.7 und 11.3.2 Buchst. a umgestellt; DORA-TLPT auf die nach Art. 26 Abs. 8 identifizierten Unternehmen begrenzt; SoA nicht mehr als Ort der Risikoführung genannt (MHC-09); CRA-SBOM und AI-Act-Sorgfaltserwartung regulatorisch präzisiert. Weitere Schärfungen: MHC-01 als übergreifendes Resilienz-Control eingeordnet, MHC-04 Identität/Credential und Lateral-Movement-Wirkung präzisiert, MHC-08 aus der Control-Degradation-Analyse hergeleitet, MHC-09 gegen die Testbasis abgegrenzt, RACI als Beispielbelegung gekennzeichnet, Aussage zur ISO-Lückenabdeckung entschärft, Begründungen zu A.5.3, A.5.5, A.5.7, A.8.6 und A.8.11 nachgeschärft, Kennzahlen G.2, G.4, G.5 und ATT&CK-Abdeckung überarbeitet, Schreibweise „AI" vereinheitlicht.

Aktualisierungsanlässe für künftige Versionen werden insbesondere sein: Veröffentlichung neuer BSI-C5-Fassungen, ISO-27002-Revisionen, Änderungen in NIS2-Durchführungsakten, neue DORA-RTS, Updates der CRA-Umsetzungsakte, signifikant neue Erkenntnisse zur Mythos-Bedrohungslage sowie Lessons Learned aus der Anwendung des Schnellhefts in der Praxis.