



SCHULUNGSKATALOG 2026 • KOMPLETT-EDITION

KI GIBT IHNEN TEMPO. IHR KÖNNEN GIBT DIE RICHTUNG.

20 rollenbasierte KI-Tracks vom Berufsstart bis zum Plattform-Betrieb, dazu Programme und Masterclasses. Mit Agenden, Festpreisen je Gruppe und Artefakten, die am Montag weiterarbeiten.

7.900 € je Trainingstag und Gruppe • ab 4 bis 12 Teilnehmende • Zertifikat inklusive

Stand: September 2026 • bordonaro.ai

WAS SIE IN DIESEM KATALOG FINDEN.

Lernmechanik und Trainer	4
Szenarien und Entscheidungshilfe	5
Die 20 Tracks	7
Preisübersicht	54
Konditionen und Kontakt	55

Jede Position trägt eine Artikelnummer und einen QR-Code: Er führt direkt zur passenden Seite mit Details, Agenda und Anfrageweg.

KEIN PAPIER ÜBER MÖGLICHKEITEN. EIN ANGEBOT MIT PREISEN.

Ich habe fast 30 Jahre in Rechenzentren, Projekten und Trainingsräumen verbracht und dabei eine einfache Überzeugung entwickelt: Wer etwas verkauft, sollte sagen können, was es kostet, für wen es gedacht ist und was am Ende existiert. Genau so ist dieser Katalog gebaut. Jede Position hat einen Preis oder einen ehrlichen Weg dorthin, eine Zielgruppe und ein überprüfbares Ergebnis.

Nicht alles muss offline. Alles muss unter Kontrolle. Wenn Sie nach dem Blättern nur eine Frage mitnehmen, dann diese: Läuft bei Ihnen jeder Workload dort, wo er hingehört?

Dino Bordonaro

Gründer und CEO · 8x Microsoft MVP in Folge



DIE LERNMECHANIK

NICHT MEHR TOOLS. EINE ARBEITSWEISE, DIE SIE BEHERRSCHEN.

01

Verstehen

Was kann KI hier leisten, was nicht?

02

Anwenden

Eine passende Aufgabe mit freigegebenen Daten bearbeiten.

03

Prüfen

Quellen, Vollständigkeit, Fehler und notwendige Freigaben untersuchen.

04

Übertragen

Eine geprüfte Vorlage oder Checkliste in den eigenen Alltag übernehmen.

05

Vertiefen

Wiederkehrende Fragen im Rollen-Track oder Workflow-Lab bearbeiten.

06

Dranbleiben

Echte Fragen in regelmäßigen Fragerunden und Updates lösen.

IHR TRAINER

Wer KI erklärt, sollte auch verstehen, worauf sie läuft.

Dino verbindet langjährige IT- und Trainerfahrung mit der Frage, wie KI im wirklichen Arbeitsalltag eingesetzt und verantwortet werden kann. Für technische Vertiefungen zählen Architektur, Datenflüsse, Berechtigungen und Betrieb. Für Ihren Einstieg zählt zunächst eine Aufgabe, die Sie nachvollziehbar bearbeiten können.

8x Microsoft MVP in Folge • Microsoft Certified Trainer • über 800 trainierte Professionals

SO FÜHLT SICH DAS AN

EINE ÜBERZEUGENDE ANTWORT KANN TROTZDEM FALSCH SEIN.

BEISPIELSZENARIO. FIKTIVER ABLAUF, KEINE KUNDENREFERENZ.

Montag, 08:17 Uhr. Die Anfrage liegt seit Freitag da.

Eine erfahrene Vertriebsmitarbeiterin kennt die Produkte und ihre Kunden. Trotzdem beginnt sie die Angebotsvorbereitung mit Suchen: alte E-Mails, Produktunterlagen, offene Fragen. Im Beispiel bekommt ein freigegebener KI-Assistent ausgewählte Unterlagen und eine klare Aufgabe. Er sammelt Anforderungen, markiert Lücken und entwirft eine Gliederung. Einen vorgeschlagenen Liefertermin verwirft sie, denn dafür gibt es keine bestätigte Quelle. Sie beginnt nicht mehr beim leeren Blatt. Die Zusage an den Kunden bleibt ihre Entscheidung.

BEISPIELSZENARIO. FIKTIVER ABLAUF, KEINE KUNDENREFERENZ.

Der erste Supportfall. Die Antwort klingt erstaunlich sicher.

Ein Berufseinsteiger lässt sich einen Fehler erklären. Die KI schlägt einen Befehl vor, den er noch nicht kennt. Einfach ausführen? Genau hier beginnt die eigentliche Kompetenz. Im Übungsfall prüft er die freigegebene Dokumentation, erklärt die mögliche Wirkung in eigenen Worten und testet nur in einer isolierten Umgebung. Für den Eingriff ins Kundensystem braucht er die zuständige Freigabe. KI hilft ihm beim Lernen. Sie erteilt ihm keine Berechtigung.

BEISPIELSZENARIO. FIKTIVER ABLAUF, KEINE KUNDENREFERENZ.

Die Präsentation ist fertig. Die wichtigste Zahl stimmt nicht.

Ein erfahrener Controller lässt einen Entwurf für den Monatsbericht vorbereiten. Der Text klingt schlüssig, doch die KI vergleicht zwei unterschiedliche Zeiträume. Im Szenario bleiben Berechnungen im geprüften Rechenwerk. KI unterstützt bei Struktur, Erklärungsvarianten und Rückfragen. Der Controller kontrolliert Zahlenbezug und Annahmen, bevor der Bericht an die Geschäftsführung geht. Der Fortschritt ist nicht blindes Tempo. Es ist weniger Routinearbeit bei weiterhin klarer Verantwortung.

Genau deshalb beginnen wir nicht mit einer Liste von Tools. Wir beginnen mit Ihrer Aufgabe und damit, woran Sie ein gutes Ergebnis erkennen.

FÜR WEN IST WAS? IHRE ENTSCHEIDUNGSHILFE

SECHS ROLLEN, SECHS LERNWEGE.

Niemand braucht alle 20 Tracks. Finden Sie Ihre Rolle, starten Sie mit dem empfohlenen Weg und lassen Sie Vorwissen im Skill-Gespräch anerkennen. Sichere Nutzung und Ergebnisprüfung gehören in jeden Weg.

Berufseinsteiger und Umschulende

„Ich nutze KI längst, aber weiß nie, ob ich ihr trauen kann.“

Sie bekommen einen sicheren Lernweg: erklären lassen, rückfragen, Quellen prüfen, selbst erklären. Aus der heimlichen Abkürzung wird ein Lernbeschleuniger.

AI-START



AI-LIT-LAB



AI-WORK

Erfahrene Fachkräfte

„Zu viel meiner Zeit steckt im Suchen, Zusammenfassen und ersten Entwurf.“

Sie bringen Ihren fachlichen Kontext in die Arbeit mit KI ein und prüfen deren Vorschläge an Ihren Maßstäben. Ihr Wissen bleibt der Maßstab, nicht die glatteste Antwort.

AI-WORK



AI-CONTEXT



AI-CHECK

Geschäftsführung und Vorstand

„Was ist KI uns wert, was riskieren wir und wer verantwortet es?“

Vom Entscheidungsbild über belastbare Business Cases bis zur Führungsfrage: Sie entscheiden auf Basis eines eigenen Urteils statt fremder Folien.

AI-EXEC



AI-VALUE



AI-LEAD

Governance, Datenschutz, Security

„Wie machen wir KI-Einsatz nachweisbar und klassifiziert?“

EU AI Act rollenklar übersetzt, Anwendungen sauber eingestuft, Bedrohungen systematisch modelliert. Schulung und Entwurf, keine Rechtsberatung.

AI-GOV



AI-RISK



AI-SEC

Architektur und Entwicklung

„Wie baue ich RAG, Agenten und Evaluation, die Produktion aushalten?“

Vom Wissensfundament über RAG und Evaluation bis zu begrenzten, beaufsichtigten Agenten: gebaut und gehärtet auf echter Lab-Hardware.

AI-DATA



AI-RAG



AI-EVAL



AI-AGENT

Plattform und Betrieb

„Wie betreibe ich Modelle und GPUs, verbunden bis getrennt?“

Souveräne Plattformen planen, betreiben und aktualisieren, bis zum vollständig getrennten Betriebsmodell. Ihr Heimspiel wird unser Kursraum.

AI-PLATFORM



AI-OPS



AI-ARC



AI-RESIDENCY

Der Online-Check hilft in 7 Fragen: bordonaro.ai/de/academy

AIA-AI-LIT-LAB

PFAD 1 · VERSTEHEN

k/aia-ai-lit-lab

AI Literacy Praxislabor

Ein halber Tag, nach dem jede Person einen eigenen Assistenten gebaut hat, statt nur davon gehört zu haben.

Sie richten im Labor einen eigenen Assistenten ein, üben die Datenregeln an realen Fällen und schließen mit einem Transfer-Check ab, in Gruppen von 4 bis 12 Personen.

Für wen: Teams und Abteilungen, die vom Zuhören ins Tun kommen sollen: Sachbearbeitung, Vertrieb, Verwaltung, Service.

Dauer: ½ Tag**Gruppe:** 4-12 Personen**Vorlauf:** ab 2 Wochen**3.950 €** pauschal je Durchführung

AGENDA

09:00

Vom Wissen zum Handgriff

Jede Person formuliert den ersten Prompt zur mitgebrachten Aufgabe aus dem Steckbrief. Die Ergebnisse werden verglichen: Warum liefert dieselbe Frage zwanzig verschiedene Antworten? Daraus entsteht die erste Verbesserungsrunde.

09:50

Datenregeln im Ernstfall

Übung mit Fallkarten: Kundenbeschwerde, Inhaltsliste, Pressemitteilung, Angebotskalkulation. Jede Person entscheidet pro Fall, ob und in welche KI die Daten dürfen, und begründet die Entscheidung. Strittige Fälle werden im Plenum ausgefochten.

10:50

Prompts, die wiederholbar funktionieren

Aus den Versuchen des Vormittags wird ein Muster: Rolle, Aufgabe, Kontext, Format. Jede Person überarbeitet den eigenen Prompt nach dem Muster und misst am Ergebnis, was sich verbessert hat.

11:30

Ihr erster eigener Assistent

Jede Person richtet einen Assistenten mit eigener Rollenbeschreibung, Tonfall und klaren Grenzen ein, auf unserer Lab-Umgebung oder Ihrem Tenant. Am Ende beantwortet der Assistent die Steckbrief-Aufgabe besser als der erste Prompt vom Morgen.

12:30

Transfer-Check und Montagsplan

Jede Person führt den eigenen Assistenten an einem neuen Fall vor und entscheidet drei Datenregel-Fälle im Check. Abschluss: Was baue ich am Montag nach, was brauche ich dafür von wem?

Das nehmen Sie mit

- Jede Person hat eigene Prompts an einer echten Aufgabe aus dem eigenen Arbeitsalltag geschrieben, verbessert und bewertet
- Die Datenregeln sind an konkreten Fällen geübt: jede Person hat mindestens zehn Einordnungen selbst entschieden und begründet
- Ein erster eigener Assistent mit Rolle, Tonfall und Grenzen, eingerichtet auf unserer Lab-Umgebung oder Ihrem Tenant
- Ein bestandener Transfer-Check je Person als dokumentierter Baustein Ihres Art.-4-Kompetenznachweises

Voraussetzungen

Grundbegriffe wie Prompt und Halluzination sollten schon einmal gehört sein, etwa aus dem AI Literacy Praxislabor (AI-LIT-LAB). Ein Laptop mit Browser genügt, KI-Vorerfahrung ist nicht nötig.

Abschluss

Transfer-Check am Ende: jede Person demonstriert den eigenen Assistenten an einem neuen Fall und entscheidet drei Datenregel-Fälle. Ergebnis und Teilnahme werden nachvollziehbar und prüffähig dokumentiert.

Enthalten

½ Tag Praxislabor mit Dino Bordonaro, vor Ort oder remote · Vorbereitete Übungs-Accounts auf unserer Lab-Umgebung oder Einrichtung auf Ihrem Tenant nach Absprache · Übungsset mit Fallkarten zu den Datenregeln und Prompt-Startvorlagen · Schritt-für-Schritt-Anleitung zum Nachbauen des Assistenten im eigenen Haus · Teilnahmezertifikate und prüffähige Schulungsdokumentation inklusive Transfer-Check-Ergebnis

AIA-AI-WORK

PFAD 1 · VERSTEHEN



k/aia-ai-work

Sicher produktiv mit KI arbeiten

Ein Tag, an dem drei Ihrer echten Arbeitsabläufe messbar schneller werden, nicht drei Folienbeispiele.

Drei real verbesserte Arbeitsabläufe mit Vorher-Nachher-Messung, persönliche Prompt- und Prüfvorlagen und verabschiedete Teamregeln.

Für wen: Teams bis 12 Personen, die KI schon einmal angefasst haben und jetzt ihre tatsächliche Arbeit damit erledigen wollen: Berichte, Angebote, Protokolle, Kundenkorrespondenz, Recherche.

Dauer: 1 Tag**Gruppe:** 4-12 Personen**Vorlauf:** ab 2 Wochen**7.900 €** pauschal je Durchführung

AGENDA

09:00

Ihre Abläufe auf dem Tisch

Die drei ausgewählten Abläufe werden am Whiteboard zerlegt: Schritte, Zeitaufwand, Fehlerquellen. Für jeden Ablauf legt das Team die Vorher-Messlatte fest: Wie lange dauert er heute, was ist ein gutes Ergebnis?

10:00

Ablauf 1 umbauen

Der erste Ablauf wird live umgebaut: gemeinsame Prompt-Vorlage entwickeln, an drei echten Fällen ausführen, Ergebnisse gegen die Messlatte halten. Erste Frage des Tages: Wo spart KI wirklich, wo kostet die Kontrolle den Gewinn wieder auf?

11:30

Die Prüfroutine

An den Ergebnissen von Ablauf 1 wird die Prüfvorlage gebaut: Faktencheck, Quellenfrage, Zahlenprüfung, Tonalität. Übung mit präparierten Ausgaben: das Team findet eingebaute Halluzinationen und entscheidet, welche Prüfschritte Pflicht werden.

13:30

Ablauf 2 und 3 in Werkstattarbeit

In zwei Gruppen bauen die Teilnehmenden Ablauf 2 und 3 selbst um: Prompt-Vorlage, Testlauf an echten Fällen, Prüfschritte. Der Trainer rotiert, korrigiert und stellt die unbequemen Fragen: Würden Sie dieses Ergebnis unterschreiben?

15:00

Vorher-Nachher-Messung

Beide Gruppen präsentieren am Live-Beispiel. Für jeden der drei Abläufe wird festgehalten: Zeit vorher, Zeit nachher inklusive Prüfung, verbleibende Risiken. Ehrliches Ergebnis: mindestens ein Kandidat fällt oft durch, auch das wird dokumentiert.

16:00

Teamregeln und Montagsplan

Das Team entscheidet und verschriftlicht seine Regeln: welche Aufgaben mit KI, welche Daten erlaubt, wer prüft was, was bleibt tabu. Jede Person benennt den einen Ablauf, den sie ab Montag mit Vorlage und Prüfroutine fährt.

Das nehmen Sie mit

- Drei reale Arbeitsabläufe des Teams sind mit KI umgebaut und mit Vorher-Nachher-Messung dokumentiert
- Jede Person hat persönliche Prompt-Vorlagen für die eigenen wiederkehrenden Aufgaben
- Eine Prüfvorlage je Ablauf: was wird vor der Weiterverwendung kontrolliert, von wem, woran erkennt man Halluzinationen
- Verabschiedete Teamregeln: welche Aufgaben mit KI, welche Daten, wer prüft, was tabu bleibt
- Ein dokumentierter Trainingstag als Baustein Ihres Art.-4-Kompetenznachweises

Voraussetzungen

Jede Person sollte schon eigene Prompts geschrieben haben, etwa aus AI-LIT-LAB oder eigener Praxis. Zugang zu einem im Haus freigegebenen KI-Werkzeug am eigenen Laptop.

Abschluss

Transfernachweis über die Vorher-Nachher-Messung der drei Abläufe und die von jeder Person selbst erstellten Vorlagen. Teamregeln und Messergebnisse werden im Ergebnisdokument nachvollziehbar und prüffähig festgehalten.

Enthalten

1 Trainingstag mit Dino Bordonaro, vor Ort oder remote · Auswertung der eingereichten Ablauf-Steckbriefe und Auswahl der drei Trainings-Abläufe vorab · Vorlagen-Set: Prompt-Vorlage, Prüfvorlage und Teamregel-Vorlage zum Weiterbenutzen · Fotoprotokoll beziehungsweise Ergebnisdokument mit allen erarbeiteten Vorlagen und Messungen · Teilnahmezertifikate und prüffähige Schulungsdokumentation · 30-minütiges Follow-up-Gespräch nach vier Wochen: Was ist im Alltag angekommen, wo hakt es?

AIA-AI-CONTEXT

PFAD 1 · VERSTEHEN

k/aia-ai-context

Prompt- und Context-Engineering

Zwei Tage, nach denen gute KI-Ergebnisse in Ihrem Team keine Glückssache mehr sind, sondern wiederholbare Bausteine.

Wiederverwendbare Kontextmuster, strukturierte Outputs mit Qualitätskontrollen und eine getestete Prompt-Bibliothek für Ihr Team.

Für wen: Vielnutzer, Fachexperten und Entwickler, die täglich mit KI arbeiten und deren Ergebnisse andere weiterverwenden: Analysen, Berichte, Datenextraktion, Textbausteine, Vorstufen für Automatisierung.

Dauer: 2 Tage **Gruppe:** 4-12 Personen **Vorlauf:** ab 4 Wochen **15.800 €** pauschal je Durchführung

AGENDA

Tag 1 · 09:00	Warum derselbe Prompt zehnmal anders antwortet Die eingereichten Prompts im Stresstest: derselbe Prompt, mehrere Läufe, gemessene Streuung. Daran wird sichtbar, welche Formulierungen tragen und welche dem Zufall überlassen. Jede Person wählt ihren wichtigsten Prompt als Arbeitsobjekt für beide Tage.
Tag 1 · 10:30	Kontextmuster statt Bauchgefühl Aus den besten Einreichungen werden Muster destilliert: Rolle, Fachkontext, Beispiele, Grenzen, Negativanweisungen. Übung: jede Person zerlegt den eigenen Prompt in Bausteine und benennt, welcher Baustein welchen Effekt hat, geprüft durch Weglassen.
Tag 1 · 13:30	Kontext aus Unternehmenswissen Wie viel Dokument, Tabelle oder Vorwissen gehört in den Kontext, und in welcher Form? Übung an echten Unterlagen: dasselbe Dokument als Rohtext, als Auszug und als strukturierte Stichpunkte in den Kontext geben und die Ergebnisqualität vergleichen. Dazu die Datenregel-Frage: Was darf überhaupt hinein?
Tag 1 · 15:30	Tagesergebnis: Ihr Muster-Set besteht den Test Jede Person baut den eigenen Arbeitsprompt vollständig auf Kontextmuster um und führt ihn gegen drei vorher definierte Testfälle aus. Bestanden ist, was in allen drei Fällen brauchbare und begründbar bessere Ergebnisse liefert als die Ursprungsfassung. Die Muster wandern in den gemeinsamen Katalog.
Tag 2 · 09:00	Strukturierte Outputs: JSON und Tabellen Vom Fließtext zum weiterverwendbaren Ergebnis: Output-Schemata definieren, Pflichtfelder erzwingen, Tabellen für Excel sauber formatieren. Übung: eine Datenextraktion aus unstrukturierten Dokumenten, deren Ergebnis ohne Nacharbeit in eine Tabelle fällt. Entwickler validieren das JSON zusätzlich maschinell.
Tag 2 · 11:00	Qualitäts- und Halluzinationskontrollen Jeder Bibliotheks-Prompt bekommt Testfälle und Grenzfälle: leere Eingaben, widersprüchliche Angaben, fehlende Informationen. Übung an präparierten Dokumenten mit eingebauten Fehlern: Welche Kontrollen fangen die erfundene Zahl, welche lassen sie durch? Daraus entsteht die Prüfroutine je Prompt.

DIE 20 TRACKS · TRACK 3/20 · FORTSETZUNG

Tag 2 · 13:30

Die Prompt-Bibliothek des Teams

Aus den Einzelergebnissen wird ein Fundus: Namenskonvention, Beschreibung, Einsatzgrenzen, Testfälle, Version. Werkstattarbeit: die zehn wichtigsten Prompts des Teams werden bibliotheksfähig gemacht und gegenseitig getestet. Entscheidung im Team: Wo liegt die Bibliothek, wer pflegt sie, wie kommen neue Prompts hinein?

Tag 2 · 15:30

Tagesergebnis: Bibliothek Version 1 im Abnahmetest

Der Ernstfall: jede Person löst eine fremde Aufgabe nur mit den Bibliotheks-Prompts der anderen. Was ohne Rückfragen funktioniert, ist abgenommen, der Rest bekommt konkrete Nacharbeit mit Verantwortlichem und Termin. Abschluss: Pflegeplan und die ersten drei Ausbau-Kandidaten.

Das nehmen Sie mit

- Ein Satz wiederverwendbarer Kontextmuster: Rollen, Fachkontext, Beispiele und Grenzen als Bausteine, die sich je Aufgabe kombinieren lassen
- Prompts mit strukturierten Outputs: JSON und Tabellen, die sich direkt in Excel, Weiterverarbeitung oder Folgesysteme übernehmen lassen
- Qualitäts- und Halluzinationskontrollen je Prompt: Testfälle, Grenzfälle und definierte Prüfschritte vor der Weiterverwendung
- Eine Prompt-Bibliothek Version 1 für Ihr Team: benannt, dokumentiert, getestet und mit geklärter Pflege-Verantwortung
- Zwei dokumentierte Trainingstage als Baustein Ihres Art.-4-Kompetenznachweises

Voraussetzungen

Regelmäßige eigene KI-Nutzung im Arbeitsalltag wird vorausgesetzt, etwa auf dem Stand nach AI-WORK. Für die JSON-Übungen sind keine Programmierkenntnisse nötig, Entwickler bekommen an denselben Übungen zusätzliche Tiefe.

Abschluss

Zwei überprüfbare Tagesergebnisse: das Muster-Set jeder Person besteht drei definierte Testfälle, die Team-Bibliothek besteht den gegenseitigen Abnahmetest. Beide Ergebnisse werden nachvollziehbar und prüffähig dokumentiert.

Enthalten

2 Trainingstage mit Dino Bordonaro, inhouse oder remote · Vorab-Auswertung der eingereichten Prompts als Startbild des Teams · Musterkatalog Kontextmuster und Output-Schemata als editierbare Vorlagen · Prompt-Bibliothek-Vorlage inklusive Namenskonvention, Versionierung und Testfall-Struktur · Teilnahmezertifikate und prüffähige Schulungsdokumentation · 45-minütige Remote-Sprechstunde vier Wochen später zur Weiterentwicklung der Bibliothek

AIA-AI-START

PFAD 1 · VERSTEHEN



k/aia-ai-start

KI-Einstieg für den Berufsstart

Ein Tag, an dem Berufseinsteiger lernen, KI als Lernwerkzeug zu nutzen, ohne ihr blind zu vertrauen: Aufgaben erklären lassen, gute Rückfragen stellen, Quellen prüfen und die eigene Verständnislücke erkennen, bevor sie zum Fehler wird.

Sie üben, sich Aufgaben von KI erklären zu lassen, die Erklärung zu hinterfragen und einen synthetischen Übungsfall am Ende selbst in eigenen Worten zu erklären, ohne die KI daneben.

Für wen: Auszubildende, Berufseinsteiger, Umschulende und Quereinsteiger in den ersten Berufsjahren.

Dauer: 1 Trainingstag**Gruppe:** 4-12 Personen**Vorlauf:** ab 2 Wochen**7.900 €** pauschal je Durchführung

AGENDA

09:00	Was KI kann und woher die Antworten kommen Funktionsweise in verständlichen Bildern: Sprachmodelle erzeugen plausible Fortsetzungen, keine geprüften Wahrheiten. Erste Übung: dieselbe Frage dreimal stellen und die Unterschiede erklären.
10:00	Aufgaben erklären lassen, richtig Ein Fachbegriff, ein Prozess, eine Fehlermeldung: Wie man eine Erklärung anfordert, die zum eigenen Wissensstand passt, und welche Rückfragen die Erklärung testbar machen.
11:15	Quellen prüfen statt glauben Übung an präparierten Ausgaben: Welche Aussage hat eine prüfbare Quelle, welche klingt nur so? Jeder Teilnehmer findet mindestens eine eingebaute plausible Falschaussage.
13:00	Datenregeln am eigenen Arbeitsplatz Was darf in welche KI: Fallkarten mit typischen Situationen aus Ausbildung und Einstieg. Kundendaten, interne Unterlagen, eigene Notizen. Die Hausregeln werden konkret statt abstrakt.
14:00	Der eigene Übungsfall Jeder bearbeitet einen mitgebrachten Begriff oder eine Aufgabe mit der gelernten Routine: erklären lassen, rückfragen, Quelle prüfen, Verständnis dokumentieren. Peer-Austausch zu zweit.
15:30	Selbst erklären, ohne KI daneben Die Nagelprobe: Jeder erklärt seinen Fall in eigenen Worten vor der Kleingruppe. Was noch wackelt, wird sichtbar und bekommt eine konkrete Lernaufgabe.
16:30	Eskalation und nächster Schritt Wann muss eine erfahrene Person ran: die persönliche Grenze schriftlich festlegen. Jeder geht mit einer Lernroutine-Karte und einem konkreten nächsten Übungsschritt.

Das nehmen Sie mit

- Eine persönliche Lernroutine: erklären lassen, rückfragen, Quelle prüfen, selbst erklären
- Ein geübter Umgang mit Unsicherheit: erkennen, wann eine Antwort nur plausibel klingt
- Eine klare Regel, wann eine erfahrene Person einbezogen werden muss
- Ein selbst bearbeiteter synthetischer Übungsfall mit dokumentierter Prüfung
- Die Datenregeln des Hauses am eigenen Arbeitsplatz angewendet

Voraussetzungen

Keine. Ein Laptop mit Browser genügt. Der Track setzt bewusst keine KI-Vorerfahrung voraus.

Abschluss

Bestanden hat, wer seinen Übungsfall am Ende in eigenen Worten erklären kann und die Prüfroutine daran nachweisbar angewendet hat. Das Ergebnis wird in der Schulungsdokumentation festgehalten.

Enthalten

1 Trainingstag mit Dino Bordonaro, inhouse oder remote · Übungsfälle auf synthetischen Daten, keine echten Kundendaten nötig · Lernroutine-Karte und Prüffragen-Vorlage zur Weiterverwendung · Teilnahmezertifikat und Schulungsdokumentation · 4 Wochen später eine 60-Minuten-Fragerunde remote für offene Fälle

AIA-AI-CHECK

PFAD 1 · VERSTEHEN



k/aia-ai-check

KI-Ergebnisse prüfen und verbessern

Ein Tag nur für die Fähigkeit, die über Nutzen oder Schaden entscheidet: KI-Ergebnisse systematisch prüfen. Quellen und Zahlen, plausible Irrtümer, Gegenproben und die Frage, wann ein Ergebnis eine Freigabe braucht.

Sie üben eine belastbare Prüfliste an mehreren synthetischen Fällen und passen sie an die Arbeit Ihres Teams an. Am Ende erkennt jeder die typischen plausiblen Irrtümer und weiß, wann eskaliert wird.

Für wen: Teams, die KI-Entwürfe bereits nutzen und nun die Qualitätssicherung dazu brauchen: Sachbearbeitung, Assistenz, Fachabteilungen, Controlling-nahe Rollen.

Dauer: 1 Trainingstag **Gruppe:** 4-12 Personen **Vorlauf:** ab 2 Wochen

7.900 € pauschal je Durchführung

AGENDA

09:00	Warum plausibel nicht richtig ist Die Fehlerklassen von Sprachmodellen an echten Mustern: erfundene Quellen, falsche Zahlenbezüge, stille Annahmen, überzeugender Ton. Übung: In drei vorbereiteten Ausgaben steckt je ein Fehler, wer findet ihn ohne Hilfsmittel?
10:15	Die Prüfliste, Schritt für Schritt Quellen, Zahlen, Vollständigkeit, Auftragsstreue, Tonalität, Datenherkunft. Jeder Prüfschritt bekommt eine konkrete Handbewegung und ein Zeitbudget. Prüfen muss schneller sein als selbst schreiben, sonst macht es niemand.
11:30	Gegenproben, die wirken Rückrechnung bei Zahlen, Quellabgleich bei Behauptungen, Gegenfrage an das Modell bei Unsicherheit. Übung an synthetischen Fällen: Welche Gegenprobe entlarvt welchen Fehlertyp?
13:00	Ihre Fälle, Runde eins Die mitgebrachten anonymisierten Beispiele werden mit der Prüfliste bearbeitet. Was hätte die Liste gefangen, was nicht? Die Liste wird dort nachgeschärft, wo sie am eigenen Material versagt.
14:15	Eskalation und Freigabe Nicht jedes Ergebnis braucht vier Augen, manche brauchen sechs. Entscheidungsbaum: Was gebe ich selbst frei, was prüft ein Kollege, was geht an die Führungskraft? Die Regel wird schriftlich und teamtauglich.
15:30	Prüfprotokoll für den Nachweis Für nachweispflichtige Fälle: das schlanke Protokoll, das dokumentiert, was geprüft wurde. Übung am eigenen Fall, danach ist das Muster einsatzbereit.
16:30	Teamregel und nächster Schritt Die angepasste Prüfliste wird zur Teamvereinbarung: Welche Prüfschritte sind ab morgen Pflicht, wer pflegt die Liste, wann wird nachgeschärft. Jeder geht mit einem konkreten Anwendungsvorsatz.

Das nehmen Sie mit

- Eine Prüfliste für KI-Ergebnisse, an mehreren Fallklassen geübt und aufs Team angepasst
- Ein Katalog der plausiblen Irrtümer: erfundene Quellen, verschobene Zeiträume, stille Annahmen
- Geübte Gegenproben: Rückrechnung, Quellabgleich, Gegenfrage an das Modell
- Eine dokumentierte Eskalations- und Freigaberegeln für kritische Ergebnisse
- Ein Prüfprotokoll-Muster für nachweispflichtige Fälle

Voraussetzungen

Erste eigene Erfahrung mit KI-Entwürfen im Arbeitsalltag, etwa aus AI-WORK, AI-LIT-LAB oder eigener Praxis. Ein Laptop mit Browser genügt.

Abschluss

Bestanden hat, wer im Abschlussfall die eingebauten Fehler mit der Prüfliste findet und die Eskalationsentscheidung begründet. Das Ergebnis wird in der Schulungsdokumentation festgehalten.

Enthalten

1 Trainingstag mit Dino Bordonaro, inhouse oder remote · Synthetische Prüffälle mit eingebauten, dokumentierten Fehlern · Prüflisten-Vorlage und Prüfprotokoll-Muster zur Weiterverwendung · Teilnahmezertifikat und Schulungsdokumentation · 4 Wochen später eine 60-Minuten-Fragerunde remote für strittige Fälle

AIA-AI-EXEC

PFAD 2 · FÜHREN

k/aia-ai-exec

Sovereign AI Executive Briefing

Ein Nachmittag, nach dem Ihre Geschäftsführung weiß, ob KI in Ihrem Haus Miete oder Eigentum wird.

Ein dokumentiertes Entscheidungsbild zu Wert, Risiko, Souveränität und Betriebsmodell, plus die offenen Entscheidungsfragen der nächsten 90 Tage als strukturierte Vorlage für Ihre Gremien.

Für wen: Geschäftsführung, Vorstand, Beirat und Gesellschafter, maximal acht Personen.

Dauer: ½ Tag

Gruppe: 4-8 Personen

Vorlauf: ab 2 Wochen

4.900 € pauschal je Durchführung

AGENDA

14:00

Lagebild: der Druck von drei Seiten

Auswertung Ihres Vorab-Fragebogens: Kundenerwartungen, Wettbewerb, Schatten-KI im eigenen Haus. Leitfrage der Runde: Welche drei KI-Entscheidungen stehen bei Ihnen in den nächsten 24 Monaten wirklich an?

14:45

Der Live-Beweis

Demonstration einer vom Internet getrennten KI-Umgebung aus unserem Enterprise-Lab, wenn technisch und organisatorisch möglich: der Sovereign Assistant beantwortet Fragen zu eingespielten Dokumenten, sichtbar ohne Cloud-Verbindung. Entscheiderfrage danach: Was bedeutet das für unsere kritischen Datenbestände?

15:45

Ökonomie: Miete oder Eigentum

Das Rechenmodell: Token- und Lizenzkosten pro Kopf und Jahr gegen Invest und Betrieb einer eigenen Umgebung, mit konservativen Annahmen. Übung: Ihr Nutzungsprofil wird live eingesetzt, der Break-even wird für Ihr Haus sichtbar.

16:30

Strategie-Session: Ihr Entscheidungsbild

Wert, Risiko, Souveränität und Betriebsmodell auf einer Seite. Entscheidung im Raum: Welche zwei Fragen will die Geschäftsleitung in 90 Tagen belastbar beantwortet haben, und wer trägt sie?

17:10

Die nächsten 90 Tage

Die Runde formuliert ihren 90-Tage-Entwurf mit Auftragskandidaten und Zuständigkeitsvorschlägen. Ob und wie er verabschiedet wird, entscheiden Ihre Gremien. Abgleich mit passenden nächsten Schritten wie AI-VALUE oder AI-LEAD. Ende 17:30.

Das nehmen Sie mit

- Ein Entscheidungsbild auf einer Seite: Wert, Risiko, Souveränität und Betriebsmodell für Ihr Haus, im Raum gemeinsam erarbeitet
- Eine Miete-oder-Eigentum-Rechnung mit Ihrem Nutzungsprofil und konservativen Annahmen, inklusive sichtbarem Break-even
- Ein priorisierter 90-Tage-Entwurf mit zwei bis drei Auftragskandidaten, Verantwortliche und Termine schlägt die Runde selbst vor
- Ein gemeinsames Vokabular auf Entscheiderebene, damit die nächste KI-Diskussion eine Entscheidung wird statt einer Vertagung

Voraussetzungen

Keine. Technisches Vorwissen ist nicht erforderlich, mitzubringen sind Entscheidungskompetenz und 3,5 ungestörte Stunden.

Abschluss

Kein Quiz auf Entscheiderebene. Der Transfernachweis ist das dokumentierte Entscheidungsbild mit dem 90-Tage-Entwurf, das als Protokoll an alle Teilnehmenden geht und im Nachgespräch nach vier Wochen besprochen wird.

Enthalten

Halbtägige Session, gehalten von Dino Bordonaro, vor Ort, remote oder in unserem Enterprise-Lab · Auswertung des Vorab-Fragebogens als individuelles Lagebild Ihres Hauses · Live-Demonstration einer vom Internet getrennten KI-Umgebung aus dem Enterprise-Lab, wenn technisch und organisatorisch möglich · Miete-oder-Eigentum-Rechenmodell als weiterverwendbare Vorlage · Dokumentiertes Entscheidungsbild und 90-Tage-Entwurf als Protokoll binnen drei Werktagen · 30-minütiges Nachgespräch mit dem Auftraggeber nach vier Wochen

AIA-AI-LEAD

PFAD 2 · FÜHREN



k/aia-ai-lead

Führen in der Human-Agent-Organisation

Zwei Tage, nach denen Ihre Führungskräfte wissen, was sie an einen Assistenten delegieren und was niemals.

Jede Führungskraft erarbeitet eine Delegationsmatrix, ein Eskalationsmodell und den Entwurf eines 90-Tage-Adoption-Plans für das eigene Team.

Für wen: Führungskräfte mit Personalverantwortung, vom Teamleiter bis zum Bereichsleiter, maximal zwölf Personen.

Dauer: 2 Tage **Gruppe:** 4-12 Personen **Vorlauf:** ab 4 Wochen

15.800 € pauschal je Durchführung

AGENDA

Tag 1 · 09:00	Ihr Führungsalltag im Röntgenbild Arbeit an den mitgebrachten Wochenlisten: Welche der zehn Aufgaben sind Entwurfsarbeit, welche sind echte Entscheidungen? Erste Sortierung an der Wand, die Muster über alle Teilnehmenden hinweg werden sichtbar.
Tag 1 · 10:30	Was Assistenten und Agenten heute wirklich leisten Live-Arbeit am System: Zielvereinbarungs-Entwurf, Vorbereitung eines schwierigen Feedbackgesprächs, Meeting-Zusammenfassung. Inklusive eines sichtbaren Fehlers und der Frage: Woran hätte ich den erkannt, wenn ich nur überflogen hätte?
Tag 1 · 13:00	Delegationsregeln Aufbau der Aufgabenklassen-Matrix: geht an den Assistenten, geht an den Agenten mit Freigabe, bleibt beim Menschen. Übung an den eigenen Wochenlisten, Grenzfälle werden in der Gruppe entschieden. Leitprinzip: der Assistent entwirft, der Mensch gibt frei.
Tag 1 · 15:00	Entscheidungs- und Eskalationsmodell Freigabestufen definieren: Wann genügt Stichprobe, wann braucht es Vier-Augen, wann ist KI tabu? Drei kritische Fälle werden durchgespielt, darunter ein fehlerhafter Entwurf, der beinahe rausgegangen wäre.
Tag 1 · 16:15	Tagesergebnis: Ihre Delegationsmatrix steht Jede Führungskraft dokumentiert Matrix und Freigabestufen für das eigene Team und stellt sie einem Partner im Peer-Review vor. Offene Punkte werden für Tag 2 notiert.
Tag 2 · 09:00	Ängste ernst nehmen Die echten Einwände auf dem Tisch: Ersetzt mich das? Wird meine Arbeit jetzt gemessen? Übung: Antwortlinien formulieren, die ehrlich sind statt beschwichtigend, und benennen, was die Führungskraft zusagen kann und was nicht.
Tag 2 · 10:30	Betriebsrat und Mitbestimmung Welche Einführungsschritte mitbestimmungsrelevant sein können, wie frühe Einbindung konkret aussieht und welche Zusagen tragfähig sind. Keine Rechtsberatung, aber eine Gesprächsstruktur, die im Termin mit dem Gremium trägt. Übung: Eröffnung des Betriebsratsgesprächs im Rollenspiel.

DIE 20 TRACKS · TRACK 7/20 · FORTSETZUNG

Tag 2 · 13:00

Der Adoption-Plan

Die ersten 90 Tage im eigenen Team: Vorbild-Nutzung der Führungskraft, drei Team-Routinen, Erfolgsmessung ohne Überwachungsgefühl. Jede Führungskraft entscheidet: Mit welcher Aufgabe startet mein Team in Woche 1?

Tag 2 · 14:30

Führungsroutinen live

Simulation eines Weeklys mit Assistenten-Entwürfen: Entwurf prüfen, Freigabe erteilen oder verweigern, Feedback an das Team formulieren. Die Freigaberegeln von Tag 1 werden unter Zeitdruck angewendet.

Tag 2 · 16:00

Tagesergebnis: Ihr Plan vor der Gruppe

Jede Führungskraft stellt Delegationsmatrix und Adoption-Plan in fünf Minuten vor, die Gruppe hält mit den härtesten Einwänden dagegen. Fixierte Fassung geht in die Schulungsdokumentation.

Das nehmen Sie mit

- Eine Delegationsmatrix je Führungskraft: welche Aufgabenklassen an Assistenten gehen, welche an Agenten mit Freigabe, welche beim Menschen bleiben
- Ein Entscheidungs- und Eskalationsmodell mit Freigabestufen nach dem Prinzip: der Assistent entwirft, der Mensch gibt frei
- Ein 90-Tage-Adoption-Plan für das eigene Team mit den ersten drei Führungsroutinen
- Eine belastbare Gesprächsline für Betriebsrat und Team, die Ängste ernst nimmt statt sie wegzumoderieren
- Ein dokumentiertes Rollenbild der Führungskraft in der Human-Agent-Organisation

Voraussetzungen

Personalverantwortung und erste eigene Berührung mit einem KI-Assistenten. Wer noch nie mit einem Assistenten gearbeitet hat, holt das mit dem Prework in einer Stunde nach.

Abschluss

Der TransfERNachweis ist die Präsentation von Delegationsmatrix und Adoption-Plan vor der Gruppe am Ende von Tag 2. Beide Artefakte und die Teilnahme werden nachvollziehbar und prüffähig dokumentiert.

Enthalten

Zwei Präsenz- oder Remote-Tage, gehalten von Dino Bordonaro · Vorlagen für Delegationsmatrix, Eskalationsmodell und Adoption-Plan zur Weiterverwendung im Haus · Gesprächsleitfaden für Team- und Betriebsratsgespräche · Teilnahmezertifikate und prüffähige Schulungsdokumentation als Baustein Ihres Art.-4-Nachweises · 60-minütige Remote-Sprechstunde für die Gruppe vier Wochen nach dem Training

AIA-AI-VALUE

PFAD 2 · FÜHREN



k/aia-ai-value

Use-Case- und Business-Case-Sprint

Zwei Tage, nach denen Ihre KI-Ideenliste ein priorisiertes Portfolio mit Eigentümern und Zahlen ist.

Ein priorisiertes Use-Case-Portfolio mit Wert-Risiko-Matrix, konservativ gerechneter Wirtschaftlichkeit, einem Eigentümer je Use Case und einem 90-Tage-Plan für die Top-Kandidaten.

Für wen: Gemischte Runden aus Führungskräften, Fachbereichsvertretern und Geschäftsleitung, maximal zwölf Personen.

Dauer: 2 Tage**Gruppe:** 4-12 Personen**Vorlauf:** ab 4 Wochen**15.800 €** pauschal je Durchführung

AGENDA

Tag 1 · 09:00

Der Ideen-Stapel auf dem Tisch

Alle eingereichten Steckbriefe hängen an der Wand. Je Steckbrief drei Prüffragen: Welches Problem, wessen Zeit, welche Daten? Duplikate werden zusammengeführt, Unklares wird vom Einreicher in 60 Sekunden nachgeschärft.

Tag 1 · 10:30

Das Bewertungsraster

Wert-Dimensionen (Zeitgewinn, Qualität, Risikominderung) und Risiko-Dimensionen (Datenklasse, Prozesskritikalität, Akzeptanz) werden festgelegt. Kalibrierungs-Übung: zwei Fälle bewertet die ganze Runde gemeinsam, bis das Raster für alle gleich tickt.

Tag 1 · 13:00

Die Wert-Risiko-Matrix entsteht

Alle Use Cases werden in Kleingruppen bewertet und in die Matrix eingeordnet. Die Grenzfälle werden im Plenum ausgefochten. Entscheidung des Tages: Welche fünf bis sieben Kandidaten bilden die Top-Gruppe?

Tag 1 · 15:00

Machbarkeits-Check der Top-Gruppe

Je Top-Kandidat: Wie ist die Datenlage, welche Systeme sind beteiligt, welcher Reifegrad ist nötig? Wo es passt, zeigen wir am Sovereign Assistant in unserer Lab-Umgebung, wie nah ein vergleichbarer Fall heute schon ist.

Tag 1 · 16:15

Tagesergebnis: das priorisierte Portfolio

Die Matrix steht, die Top-Gruppe ist benannt und begründet, jeder Grenzfall hat eine dokumentierte Entscheidung. Foto-Protokoll geht noch am Abend an alle.

Tag 2 · 09:00

Konservativ rechnen

Die Rechenlogik: eingesparte Stunden mal Häufigkeit mal Vollkostensatz, abzüglich Adoption-Abschlag und Prüfaufwand. Übung am gemeinsamen Beispiel: dieselbe Idee einmal optimistisch, einmal konservativ gerechnet, die Differenz ist die Diskussion.

Tag 2 · 10:30

Business Cases in Kleingruppen

Jede Gruppe rechnet ein bis zwei Top-Kandidaten durch und dokumentiert jede Annahme. Regel: Jede Zahl muss einer kritischen Nachfrage der Geschäftsführung standhalten.

Tag 2 · 13:00

Eigentümer und Hindernisse

Je Use Case werden festgelegt: ein Eigentümer mit Namen, ein Erfolgsmaß, das größte Hindernis. Entscheidung im Raum: Wer trägt was, und welcher Use Case bleibt ohne Eigentümer liegen, bis sich einer findet?

DIE 20 TRACKS · TRACK 8/20 · FORTSETZUNG

Tag 2 · 14:30

Der 90-Tage-Plan

Für die zwei bis drei besten Kandidaten wird der Pilot zugeschnitten: Scope klein genug für 90 Tage, Beteiligte, Messpunkte, erster Meilenstein in Woche 2. Abgleich mit Ihrem Reifegrad: was braucht erst Daten- oder Governance-Arbeit?

Tag 2 · 16:00

Pitch und Tagesergebnis

Jede Gruppe pitcht ihren Case in fünf Minuten, als säße die Geschäftsführung im Raum. Feedback nach Raster, letzte Korrekturen, das Portfolio-Dokument wird fixiert und der Review-Termin in acht Wochen vereinbart.

Das nehmen Sie mit

- Ein bereinigtes Use-Case-Inventar aus den vorab eingereichten Steckbriefen, Duplikate zusammengeführt, Luftschlösser aussortiert
- Eine Wert-Risiko-Matrix mit allen Kandidaten und einer begründeten Top-Gruppe
- Eine grobe Wirtschaftlichkeitsrechnung je Top-Use-Case mit konservativen, dokumentierten Annahmen
- Ein benannter Eigentümer, ein Erfolgsmaß und das größte Hindernis je Top-Use-Case
- Ein 90-Tage-Plan für die zwei bis drei besten Kandidaten mit Scope, Beteiligten und Messpunkten

Voraussetzungen

Grundverständnis, was Assistenten und Sprachmodelle leisten, etwa aus AI-LIT-LAB oder AI-EXEC. Entscheidend ist, dass die Runde ihre Prozesse kennt und Prioritäten setzen darf.

Abschluss

Der TransfERNachweis ist der Abschluss-Pitch mit dokumentiertem Business Case je Kleingruppe. Portfolio, Annahmen und Eigentümer werden im Portfolio-Dokument festgehalten und im Review nach acht Wochen gegen den Ist-Stand geprüft.

Enthalten

Zwei Präsenz- oder Remote-Tage, gehalten von Dino Bordonaro · Steckbrief-Vorlage und Bewertungsraster als wiederverwendbare Werkzeuge für künftige Portfolio-Runden · Rechenmodell für die Wirtschaftlichkeitsschätzung mit dokumentierten Annahmen · Portfolio-Dokument mit Matrix, Business Cases und 90-Tage-Plan binnen fünf Werktagen · Teilnahmezertifikate und Schulungsdokumentation · 60-minütiges Portfolio-Review remote nach acht Wochen

AIA-AI-CHAMP

PFAD 2 · FÜHREN



k/aia-ai-champ

AI Champions Program

Vier Module und sechs begleitende Fragerunden über zwölf Wochen, nach denen KI in Ihrem Haus ein Team hat statt eines einsamen Beauftragten.

Bis zu zehn befähigte Champions mit Use-Case-Pipeline, interner Community, fertigen Leitfäden und einem messbaren Transfer nach 90 Tagen.

Für wen: Motivierte Mitarbeitende und Teamleiter aus verschiedenen Bereichen, die KI-Wissen in die Fläche tragen sollen, maximal zehn Personen.

Dauer: 4 Präsenztage + 6 Fragerunden über 12 Wochen

Gruppe: 4-10 Personen

Vorlauf: ab 6 Wochen

34.900 € pauschal je Durchführung

AGENDA

Modul 1 · 09:00	Standortbestimmung und Rollenklärung Auswertung der Steckbriefe: Wer bringt was mit? Klärung der Champion-Rolle an konkreten Fällen: vormachen, anleiten, einsammeln, eskalieren. Und ebenso klar, was ein Champion nicht ist: kein Ersatz-Admin, kein Kontrolleur, kein Alleinunterhalter.
Modul 1 · 10:30	Werkzeuge im Griff Hands-on am Assistenten: Prompting-Grundmuster, Kontext mitgeben, Ergebnisse prüfen. Die Datenregeln des Hauses werden an zehn Alltagsfällen entschieden: darf das in die KI oder nicht?
Modul 1 · 13:00	Der erste eigene Fall Jeder Champion löst einen Zeitfresser aus dem eigenen Steckbrief live am System und dokumentiert den Lösungsweg auf der Leitfaden-Vorlage. Die Dokumentation ist der Prototyp der späteren internen Arbeitshilfen.
Modul 1 · 15:00	Erklären üben Übung vor der Gruppe: jeder Champion erklärt ein Konzept, etwa Halluzination oder Kontextfenster, in drei Minuten so, dass es die Kollegin am Empfang versteht. Feedback nach festem Raster.
Modul 1 · 16:15	Modulergebnis: Lernlandkarte und erste Arbeitshilfe Jeder Champion hat eine dokumentierte Arbeitshilfe und eine persönliche Lernlandkarte bis Modul 2. Hausaufgabe: die Arbeitshilfe mit zwei Kollegen testen und Reaktionen notieren.
Modul 2 · 09:00	Vom Flurfunk zur Pipeline Auswertung der Hausaufgabe: Was haben die Kollegen gesagt? Dann die Interview-Technik: mit welchen fünf Fragen ein Champion aus einem Flurgespräch einen brauchbaren Use Case macht.
Modul 2 · 10:30	Steckbrief-Werkstatt Jeder Champion schreibt zwei Use-Case-Steckbriefe aus dem eigenen Bereich: welches Problem, wessen Zeit, welche Daten. Gegenseitiges Review in Paaren, unklare Steckbriefe werden nachgeschärft.
Modul 2 · 13:00	Bewerten mit dem Wert-Risiko-Raster Die kompakte Fassung des Bewertungsrasters aus dem Use-Case-Sprint: Wert und Risiko je Steckbrief einschätzen, Grenzfälle in der Gruppe entscheiden. Übung: zwei bewusst strittige Fälle bis zur Entscheidung bringen.

DIE 20 TRACKS · TRACK 9/20 · FORTSETZUNG

Modul 2 · 15:00	Das Pipeline-Board Aufbau der Pipeline im Werkzeug Ihrer Wahl: Spalten, Verantwortliche, Pflege-Rhythmus. Entscheidung: Wer kuratiert, wie kommen neue Einträge hinein, wann fliegt ein Eintrag raus?
Modul 2 · 16:15	Modulergebnis: die Pipeline lebt Mindestens zehn bewertete Use Cases stehen im Board, jede Spalte hat einen Verantwortlichen. Hausaufgabe: je Champion zwei Interviews im eigenen Bereich führen und die Pipeline füttern.
Modul 3 · 09:00	Die interne Lernsession Pipeline-Review aus der Hausaufgabe, dann das Handwerk: Aufbau einer 20-Minuten-Session, Dramaturgie, die drei häufigsten Fehler beim internen Schulen. Frage an jeden Champion: Welche Session braucht Ihr Bereich zuerst?
Modul 3 · 10:30	Session-Bau aus dem Materialpaket Jeder Champion baut aus den mitgelieferten Vorlagen seine erste eigene Session für den Heimatbereich, mit einem Live-Beispiel aus dem dortigen Alltag.
Modul 3 · 13:00	Generalprobe Die Sessions werden vor der Gruppe gehalten, gnadenlos mit Uhr. Feedback nach Raster: Einstieg, Beispielqualität, Umgang mit Fragen. Jede Session bekommt eine konkrete Verbesserung mit auf den Weg.
Modul 3 · 15:00	Community-Aufbau Kanal, Rhythmus, Spielregeln: wie Fragen ankommen, wie schnell geantwortet wird, was ein Champion tut, wenn er die Antwort nicht kennt. Eskalationswege zu IT, Datenschutz und Führung werden mit Namen hinterlegt.
Modul 3 · 16:15	Modulergebnis: Termine stehen Jeder Champion hat eine erprobte Session und einen fixierten Termin im eigenen Bereich. Das Community-Konzept ist beschlossen, der Kanal ist angelegt. Hausaufgabe: die Session vor echtem Publikum halten.
Modul 4 · 09:00	Pilot-Zuschnitt Erfahrungsbericht aus den ersten echten Sessions. Dann: aus der Pipeline werden zwei bis drei Piloten geschnitten, klein genug für 90 Tage, je mit Champion als Kümmerer und einem Erfolgsmaß.
Modul 4 · 10:30	Messbar machen Vorher-Nachher-Messung ohne Überwachungsgefühl: einfache Kennzahlen wie Durchlaufzeit oder Nachfragequote, sauber erhoben. Übung: je Pilot die Messung definieren und den Nullpunkt festlegen.
Modul 4 · 13:00	Das Transfer-Dashboard Jeder Champion legt das Dashboard für seinen Bereich an: Piloten, Sessions, Community-Aktivität, Messwerte. Regel: alles, was im Abschlussbericht stehen soll, wird ab heute erfasst.
Modul 4 · 14:30	Der Auftritt vor der Führung Vorbereitung der Ergebnispräsentation für den Sponsor: Was ist entstanden, was wird in 90 Tagen gemessen, was brauchen die Champions von der Führung? Probelauf mit hartem Feedback.
Modul 4 · 16:00	Modulergebnis und Startschuss Präsentation vor dem Sponsor, wenn organisatorisch möglich direkt im Anschluss. Jeder Champion hat seinen 90-Tage-Plan, die Office-Hours-Termine stehen im Kalender.

DIE 20 TRACKS · TRACK 9/20 · FORTSETZUNG

Danach · Woche 1-
12

Sechs Fragerunden und Transfer

Sechs Termine zu je 60 Minuten remote, alle zwei Wochen: echte Fälle aus den Bereichen, Pipeline-Review, Nachschärfen der Sessions. Am Ende steht das Abschlussgespräch mit Sponsor und Auftraggeber inklusive Transfer-Bericht mit den Messwerten je Bereich.

Das nehmen Sie mit

- Bis zu zehn Champions, die Assistenten sicher nutzen, Kollegen anleiten und wissen, wann sie eskalieren
- Eine gefüllte und gepflegte Use-Case-Pipeline mit Bewertungsraster und Pflege-Routine
- Eine interne Community mit Kanal, Rhythmus und Eskalationswegen, die nach dem Programm weiterläuft
- Ein Materialpaket aus Leitfäden und Schulungsunterlagen mit Nutzungsrecht zur internen Weiterverwendung
- Ein Transfer-Bericht nach 90 Tagen mit Vorher-Nachher-Messung je Champion-Bereich

Voraussetzungen

Neugier, Standing im eigenen Bereich und die Freistellung für vier Präsenztage plus etwa zwei Stunden pro Woche in den 90 Tagen danach. Technisches Vorwissen ist nicht erforderlich.

Abschluss

Drei überprüfbare Nachweise: die vor echtem Publikum gehaltene Session, die gepflegte Pipeline und der Transfer-Bericht mit Vorher-Nachher-Messung nach 90 Tagen. Alles wird nachvollziehbar und prüffähig dokumentiert.

Enthalten

Vier modulare Präsenztage im Abstand von zwei bis drei Wochen, gehalten von Dino Bordonaro, auf Wunsch ein Modul in unserem Enterprise-Lab · Sechs Fragerunden zu je 60 Minuten remote, alle zwei Wochen über zwölf Wochen, für Fälle aus der Praxis · Materialpaket mit Leitfäden, Session-Vorlagen und Schulungsfolien inklusive Nutzungsrecht zur internen Weiterverwendung · Use-Case-Pipeline-Vorlage und Transfer-Dashboard · Zertifikate und prüffähige Schulungsdokumentation als Baustein Ihres Art.-4-Nachweises · Abschlussgespräch mit Sponsor und Auftraggeber inklusive Transfer-Bericht

AIA-AI-GOV

PFAD 3 · REGELN



k/aia-ai-gov

EU AI Act und AI Governance

Zwei Tage, nach denen Sie ein KI-Register, einen Literacy-Plan und einen Policy-Entwurf haben statt einer PowerPoint mit Paragraphen.

Ein befülltes AI Inventory, ein Rollenmodell, ein Literacy-Plan, ein Policy-Entwurf und ein Kontrollkalender, alle an Ihren echten Systemen erarbeitet.

Für wen: Für Verantwortliche aus Recht, Datenschutz, Informationssicherheit, Compliance und IT-Governance, die den EU AI Act vom Text in einen Arbeitsstand überführen sollen.

Dauer: 2 Tage**Gruppe:** 4-12 Personen**Vorlauf:** ab 4 Wochen**15.800 €** pauschal je Durchführung

AGENDA

Tag 1 · 09:00

Was der EU AI Act wirklich verlangt

Der risikobasierte Ansatz ohne Panik: Art. 4 gilt seit Februar 2025, die Aufsicht seit August 2026. Kernfrage der Runde: Wo endet die Bemühenspflicht und wo beginnt Theater? Kein vorgeschriebenes Zertifikat, kein direktes Bußgeld für Art. 4.

Tag 1 · 10:30

Digital-Omnibus-Sachstand

Was das Vereinfachungspaket bewegt und was nicht. Entscheidung je Teilnehmer: Welche Vorbereitung ist heute schon belastbar, welche warten wir bewusst ab, ohne die Bemühenspflicht zu verletzen?

Tag 1 · 13:30

Rollenmodell an Ihrer Organisation

Anbieter, Betreiber, Importeur, Händler. Übung: Jeder Teilnehmer ordnet seine drei wichtigsten Anwendungen einer Rolle zu und benennt, wer im Haus die Pflicht trägt. Hochrisiko richtet sich nach Use Case und Rolle, nicht pauschal nach Branche.

Tag 1 · 15:30

AI Inventory befüllen

Wir starten Ihr Verzeichnis. Je Anwendung: Zweck, Datenklasse, Rolle, Verantwortlicher. Tagesergebnis: eine erste befüllte Inventory-Tabelle mit mindestens fünf echten Einträgen je Teilnehmer.

Tag 2 · 09:00

Literacy-Plan nach Art. 4

Kompetenz rollen- und kontextangemessen statt Einheitsunterweisung für alle. Übung: Wir leiten aus Ihren Rollen ab, wer welches Niveau braucht, und tragen es in das Literacy-Raster ein.

Tag 2 · 11:00

Policy-Werkstatt an Ihren Dokumenten

Wir schreiben am Policy-Entwurf mit Ihren echten Richtlinien als Grundlage. Entscheidung je Haus: Was regeln wir neu, was verweist auf bestehende Datenschutz- und Sicherheitsrichtlinien?

Tag 2 · 14:00

Freigabeprozess und Kontrollkalender

Wer gibt eine neue KI-Anwendung frei, in welchen Schritten, mit welcher Dokumentation? Wir bauen den Freigabepfad und tragen wiederkehrende Kontrollen mit Terminen und Auslösern in den Kalender ein.

DIE 20 TRACKS · TRACK 10/20 · FORTSETZUNG

Tag 2 · 16:00

Zusammenführen und Lücken benennen

Wir legen Inventory, Rollenmodell, Literacy-Plan, Policy und Kontrollkalender nebeneinander. Tagesergebnis: eine ehrliche Lückenliste mit Verantwortlichen und Fristen für die nächsten vier Wochen.

Das nehmen Sie mit

- Ein befülltes AI Inventory Ihrer real eingesetzten Anwendungen mit Rolle, Zweck und Datenklasse je Eintrag
- Ein Rollenmodell, das Anbieter, Betreiber, Importeur und Händler auf Ihre Organisation abbildet, mit benannten Verantwortlichen
- Ein Literacy-Plan nach Art. 4, der Kompetenz rollen- und kontextangemessen zuordnet statt pauschal alle zu schulen
- Ein Policy-Entwurf für den KI-Einsatz, abgestimmt auf Ihre bestehenden Richtlinien zu Datenschutz und Informationssicherheit
- Ein Kontrollkalender mit Terminen, Auslösern und Zuständigkeiten für die laufende Aufsicht

Voraussetzungen

Grundkenntnisse Ihrer eigenen Organisation und Zugriff auf vorhandene Richtlinien. Juristische Vorbildung ist nicht nötig, Entscheidungsbefugnis in Ihrem Bereich hilft.

Abschluss

Geprüft wird der Stand der fünf Artefakte am Ende von Tag 2 und die Lückenliste. Teilnahme und Inhalte werden nachvollziehbar und prüffähig dokumentiert.

Enthalten

Zwei Präsenz- oder Remote-Tage mit Dino Bordonaro und maximal zwölf Teilnehmenden · Vorlagenpaket: AI-Inventory-Tabelle, Rollenmodell-Matrix, Literacy-Plan-Raster, Policy-Gerüst und Kontrollkalender · Aktueller Sachstand zum Digital-Omnibus-Paket mit Einordnung, was heute belastbar ist und was noch in Bewegung ist · Revisionssichere Schulungsdokumentation und Teilnahmezertifikate als Baustein Ihres Art.-4-Nachweises · 30-minütiges Nachgespräch nach vier Wochen zum Stand Ihrer fünf Artefakte

AIA-AI-RISK

PFAD 3 · REGELN

k/aia-ai-risk

AI Risk Classification Lab

Ein Tag, an dem Ihre echten Anwendungen eingestuft werden statt allgemeiner Fallbeispiele aus einer Folie.

Ihre realen Use Cases sind in verboten, hochrisiko, transparenzpflichtig oder minimal eingestuft, mit dokumentierten Kriterien und Kontrollbedarf je Fall.

Für wen: Für Governance-Verantwortliche und Führungskräfte, die eine Liste eingesetzter oder geplanter KI-Anwendungen belastbar einstufen müssen.

Dauer: 1 Tag**Gruppe:** 4-12 Personen**Vorlauf:** ab 2 Wochen**7.900 €** pauschal je Durchführung

AGENDA

09:00

Die vier Klassen sauber getrennt

Verboten, hochrisiko, transparenzpflichtig, minimal. Woran man jede erkennt und wo die Graubereiche liegen. Kernfrage: Warum richtet sich Hochrisiko nach Use Case und Rolle und nicht pauschal nach Branche?

10:30

Der Entscheidungsbaum in der Anwendung

Wir gehen den Kriterienkatalog Schritt für Schritt durch. Übung an einem gemeinsamen Beispielfall, bis jeder Teilnehmer den Baum sicher liest.

13:00

Ihre echten Fälle, Runde eins

Jeder Teilnehmer klassifiziert zwei eigene Use Cases live. Entscheidung: In welche Klasse fällt der Fall und welche zwei Kriterien tragen die Entscheidung?

14:30

Ihre echten Fälle, Runde zwei und Streitfälle

Wir nehmen die schwierigen Fälle vor die Gruppe. Übung: Wo Meinungen auseinandergehen, dokumentieren wir beide Lesarten und die Bedingung, die den Ausschlag gibt.

15:45

Kontrollbedarf und Priorisierung

Je Fall benennen wir, welche Kontrolle nötig ist und wie dringend. Tagesergebnis: eine dokumentierte, priorisierte Klassifikationsliste Ihrer eingebrachten Anwendungen.

Das nehmen Sie mit

- Jeder eingebrachte Use Case ist einer Klasse zugeordnet: verboten, hochrisiko, transparenzpflichtig oder minimal
- Dokumentierte Entscheidungskriterien, die zeigen, warum ein Fall in seine Klasse fällt
- Ein benannter Kontrollbedarf je Fall statt einer pauschalen Aussage für alles
- Eine priorisierte Liste, welche Anwendungen zuerst Maßnahmen brauchen

Voraussetzungen

Eine Liste Ihrer eingesetzten oder geplanten KI-Anwendungen und Grundkenntnisse ihres jeweiligen Zwecks. Vorwissen zum EU AI Act hilft, ist aber nicht Bedingung.

Abschluss

Geprüft wird die dokumentierte Klassifikationsliste am Ende des Tages: Klasse, zwei tragende Kriterien und Kontrollbedarf je Fall. Teilnahme und Inhalte werden nachvollziehbar und prüffähig dokumentiert.

Enthalten

Ein Präsenz- oder Remote-Tag mit Dino Bordonaro und maximal zwölf Teilnehmenden · Klassifikations-Entscheidungsbaum und Kriterienkatalog als Arbeitsvorlage · Dokumentationsvorlage je Fall mit Klasse, Begründung und Kontrollbedarf · Revisionssichere Schulungsdokumentation und Teilnahmezertifikate · 15-minütiges Nachgespräch zur Priorisierung Ihrer Maßnahmenliste

AIA-AI-SEC

PFAD 3 · REGELN



k/aia-ai-sec

Secure GenAI und Agent Threat Modeling

Zwei Tage, nach denen Sie ein Threat Model, ein Datenflussbild und einen Testplan für eine echte Anwendung haben statt einer Liste theoretischer Risiken.

Ein Threat Model, ein Datenflussbild, ein Katalog konkreter Schutzmaßnahmen und ein Testplan für eine reale GenAI- oder Agenten-Anwendung Ihres Hauses.

Für wen: Für Sicherheits- und Governance-Verantwortliche zusammen mit Engineering, die eine GenAI- oder Agenten-Anwendung absichern müssen.

Dauer: 2 Tage**Gruppe:** 4-12 Personen**Vorlauf:** ab 4 Wochen**15.800 €** pauschal je Durchführung

AGENDA

Tag 1 · 09:00	Warum GenAI andere Bedrohungen hat Prompt Injection, unsichere Ausgaben, Datenabfluss über Wissensquellen und Werkzeuge. Kernfrage: Wo unterscheidet sich das von klassischer Anwendungssicherheit und wo nicht?
Tag 1 · 10:30	Datenflussbild Ihrer Anwendung Wir zeichnen den vollständigen Fluss: Eingabe, Modell, RAG-Quellen, Tool-Aufrufe, Ausgabe. Übung: Jede Vertrauensgrenze wird markiert und benannt.
Tag 1 · 13:30	STRIDE-artige Systematik auf GenAI übertragen Wir gehen die Bedrohungskategorien durch und übersetzen sie auf Prompt, Kontext, Werkzeuge und Ausgabe. Der OWASP-LLM-Referenzrahmen dient als Checkliste, damit keine Kategorie vergessen wird.
Tag 1 · 15:30	Bedrohungen an Ihrer Anwendung sammeln Übung an Ihrem Datenflussbild: Je Vertrauensgrenze benennen wir plausible Angriffe. Tagesergebnis: ein erstes Threat Model mit priorisierten Bedrohungen.
Tag 2 · 09:00	Live-Demonstration im Lab Wir zeigen an einer präparierten Umgebung, wie eine Prompt Injection über eine RAG-Quelle wirkt und was ein zu weit berechtigtes Werkzeug preisgibt. Keine Kundendaten, nur zum Verständnis der Wirkung.
Tag 2 · 11:00	Schutzmaßnahmen je Bedrohung Guardrails, Berechtigungen nach geringstem Recht, Ausgabefilter, Quellenprüfung. Entscheidung je Bedrohung: welche Maßnahme, welcher Aufwand, welches Restrisiko bleibt?
Tag 2 · 14:00	Testplan bauen Zu jeder Maßnahme ein Test, der ihre Wirkung nachweist. Übung: Wir formulieren pro wichtiger Bedrohung einen konkreten Testfall mit erwartetem Ergebnis.
Tag 2 · 16:00	Restrisiken und Übergabe Wir tragen zusammen, was bleibt, und benennen Verantwortliche. Tagesergebnis: Threat Model, Maßnahmenkatalog, Testplan und priorisierte Restrisikoliste für Ihre Anwendung.

Das nehmen Sie mit

- Ein Datenflussbild Ihrer Anwendung von der Eingabe über Modell, Wissensquellen und Werkzeuge bis zur Ausgabe
- Ein Threat Model, das eine STRIDE-artige Systematik auf GenAI und Agenten überträgt
- Ein Katalog konkreter Schutzmaßnahmen: Guardrails, Berechtigungen und Ausgabefilter, je Bedrohung zugeordnet
- Ein Testplan, der prüfbar macht, ob eine Maßnahme wirkt
- Eine priorisierte Restrisikoliste mit Verantwortlichen

Voraussetzungen

Grundkenntnisse in Anwendungsarchitektur und Informationssicherheit. Sie brauchen eine reale Anwendung, die Sie betrachten wollen, mindestens als Architekturskizze.

Abschluss

Geprüft wird die Vollständigkeit von Threat Model, Maßnahmenkatalog und Testplan am Ende von Tag 2, jeweils an Ihrer Anwendung. Teilnahme und Inhalte werden nachvollziehbar und prüffähig dokumentiert.

Enthalten

Zwei Präsenz- oder Remote-Tage mit Dino Bordonaro und maximal zwölf Teilnehmenden · Threat-Modeling-Vorlage für GenAI und Agenten mit Datenfluss-Notation · Referenzrahmen entlang bekannter OWASP-LLM-Risikofelder als Checkliste, in eigener Formulierung · Maßnahmen- und Testplan-Vorlage zum Weiterführen im Haus · Zugang zur Lab-Umgebung für die Demonstrationen und 30-minütiges Nachgespräch nach vier Wochen

AIA-AI-DATA

PFAD 4 · BAUEN



k/aia-ai-data

AI-ready Data und Knowledge Architecture

RAG-Projekte scheitern selten am Modell, fast immer an Quellen ohne Besitzer und Berechtigungen ohne Konzept.

Eine priorisierte Wissenslandkarte Ihres Hauses mit Ownership-Matrix, Quellenbewertung und einem Retrieval- und Berechtigungskonzept, das ein RAG-Projekt tragen kann.

Für wen: Architekten, Data Engineers und Entwickler, die vor der Aufgabe stehen, Unternehmenswissen für Assistenten nutzbar zu machen.

Dauer: 2 Tage**Gruppe:** 4-12 Personen**Vorlauf:** ab 4 Wochen**15.800 €** pauschal je Durchführung

AGENDA

Tag 1 · 09:00	Das Lagebild: Ihre Quellen auf dem Tisch Auswertung der anonymisierten Quellensteckbriefe. Erste Sortierfrage: Welche dieser Quellen würde ein Assistent heute zitieren und wem würden Sie die Antwort glauben?
Tag 1 · 10:30	Warum Indexieren ohne Landkarte teuer wird Anatomie gescheiterter RAG-Projekte: veraltete Duplikate, verwaiste Laufwerke, ACL-Verlust beim Crawlen. Live-Demo am Sovereign Assistant im Lab: dieselbe Frage gegen kuratierte und unkuratierte Quellen.
Tag 1 · 13:00	Werkstatt Wissenslandkarte Jede Organisation kartiert ihre Quellen entlang dreier Achsen: fachlicher Wert, Datenqualität, Zugriffsstruktur. Ergebnis ist eine Vier-Quadranten-Karte mit klaren Kandidaten und klaren Ausschlüssen.
Tag 1 · 15:00	Ownership oder Datenfriedhof Übung an den eigenen Karten: Wer besitzt welche Quelle fachlich, wer pflegt sie, was passiert bei Widerspruch zwischen zwei Versionen? Die Ownership-Matrix entsteht am eigenen Fall.
Tag 1 · 16:30	Tagesergebnis: Landkarte und Matrix stehen Jede Organisation präsentiert ihre Wissenslandkarte mit Ownership-Matrix in 5 Minuten. Prüffrage der Runde: Ist jede Top-Quelle einem Namen zugeordnet und jede Ausschluss-Entscheidung begründet?
Tag 2 · 09:00	Priorisierung: die erste Iteration schneiden Aus der Landkarte wird eine geordnete Liste: Welche 3-5 Quellen tragen die erste RAG-Iteration? Kriterien sind Antwortwert, Datenreife und Berechtigungsklarheit, nicht politische Sichtbarkeit.
Tag 2 · 10:30	Retrieval-Anforderungen statt Wunschliste Aus realen Fragetypen der Teilnehmenden entsteht ein Anforderungsprofil: Faktenfragen, Prozessfragen, Vergleichsfragen. Welche Metadaten, Strukturen und Aktualitätsgarantien braucht jede Klasse?
Tag 2 · 13:00	Berechtigungen: von der Quelle bis zur Antwort Werkstatt am Referenzkonzept: Wie kommen ACLs aus SharePoint und Fileserver in den Index und wie werden sie bei jeder Abfrage geprüft? Entscheidungsbaum für document-level security am eigenen Zugriffsmodell.

DIE 20 TRACKS · TRACK 13/20 · FORTSETZUNG

Tag 2 · 15:00

Der Pflegeprozess: aktuell bleiben statt einmal aufräumen

Übung: Für die priorisierten Quellen wird der Aktualisierungspfad definiert. Wer löscht, wer versioniert, wie erfährt der Index davon? Ohne diesen Prozess veraltet jede Landkarte in Monaten.

Tag 2 · 16:15

Tagesergebnis: das übergabefähige Paket

Jede Organisation legt Landkarte, Ownership-Matrix, priorisierte Quellenliste und Berechtigungskonzept-Entwurf als Paket ab. Gegenseitiges Review anhand einer Checkliste: Könnte ein RAG-Team am Montag damit starten?

Das nehmen Sie mit

- Eine Wissenslandkarte Ihrer realen Quellenlandschaft, anonymisiert erarbeitet, mit Bewertung nach Aktualität, Qualität und Zugriffsstruktur
- Eine Ownership-Matrix: je priorisierter Quelle ein benannter fachlicher Besitzer und ein definierter Aktualisierungsprozess
- Eine priorisierte Quellenliste für die erste RAG-Iteration mit dokumentierten Ausschlusskriterien
- Ein Retrieval-Anforderungsprofil: welche Fragetypen, welche Dokumentarten, welche Metadaten das System beantworten muss
- Ein Berechtigungskonzept-Entwurf, der Quell-ACLs bis zur Antwort durchreicht statt sie beim Indexieren zu verlieren

Voraussetzungen

Kenntnis der eigenen Systemlandschaft und Grundverständnis von Datenstrukturen. Programmierkenntnisse sind nicht erforderlich, Architekturdenken schon.

Abschluss

Das Abschluss-Review am Tag 2 prüft jedes Paket gegen eine dokumentierte Checkliste: Ownership benannt, Priorisierung begründet, Berechtigungspfad beschrieben. Ergebnis und Teilnahme werden nachvollziehbar und prüffähig dokumentiert.

Enthalten

2 Präsenztage mit Dino Bordonaro, inhouse oder remote · Vorlagen für Wissenslandkarte, Ownership-Matrix und Quellenbewertung als wiederverwendbare Dokumente · Referenz-Berechtigungskonzept mit Entscheidungsbaum für document-level security · Anonymisierte Auswertung der Quellensteckbriefe als Lagebild · Teilnahmezertifikate und Schulungsdokumentation für Ihr Compliance-Archiv · 60-minütige Remote-Sprechstunde 4 Wochen nach dem Training zur Überprüfung der Landkarte

AIA-AI-RAG

PFAD 4 · BAUEN



k/aia-ai-rag

Enterprise RAG Engineering

Ein RAG-Demo ist ein Nachmittag, ein RAG-System mit Berechtigungen und messbarer Qualität ist Engineering. Hier bauen Sie das zweite.

Ein funktionierender RAG-Prototyp auf der Lab-Umgebung mit Ingestion-Pipeline, Hybrid-Retrieval, Reranking, Evaluationsset, Berechtigungsmodell und dokumentierten Grenzen.

Für wen: Entwickler, Data Engineers und Architekten, die ein RAG-System für den Unternehmenseinsatz bauen oder verantworten.

Dauer: 3 Tage**Gruppe:** 4-12 Personen**Vorlauf:** ab 4 Wochen**23.700 €** pauschal je Durchführung

AGENDA

Tag 1 · 09:00

Warum das Tutorial-RAG an Ihren Dokumenten scheitert

Live-Befund am Lab: ein naives RAG gegen reale Dokumenttypen mit Tabellen, Scans und Versionen. Jeder Fehler wird einer Komponente zugeordnet, daraus entsteht die Bauliste der drei Tage.

Tag 1 · 10:30

Ingestion: aus Dokumenten werden verwertbare Einheiten

Hands-on: Parser für PDF, Office und HTML, Umgang mit Tabellen und Kopfzeilen, Metadaten-Extraktion. Jeder baut die Pipeline gegen einen präparierten Dokumentensatz mit bekannten Fallen.

Tag 1 · 13:00

Chunking ist eine Entscheidung, kein Default

Übung mit Messvergleich: feste Fenster, strukturbasiertes und semantisches Chunking gegen denselben Fragenkatalog. Wann welche Strategie trägt und wie man die Entscheidung dokumentiert.

Tag 1 · 15:00

Embeddings und Index in Betrieb nehmen

Embedding-Modelle im Vergleich, Dimensionen, Sprachen, Kosten. Aufbau des Vektorindex in der eigenen Lab-Umgebung, erste Abfragen gegen den eigenen Bestand.

Tag 1 · 16:30

Tagesergebnis: die Pipeline läuft

Jede Person zeigt: Dokumentensatz ingestiert, Chunks inspiziert, Index abfragbar. Prüfkriterium ist eine dokumentierte Chunking-Entscheidung mit Begründung, nicht nur laufender Code.

Tag 2 · 09:00

Die Grenzen der reinen Vektorsuche

Messübung am eigenen Index: Produktnummern, Eigennamen und Abkürzungen, an denen semantische Suche vorbeigreift. Die Fehlerliste motiviert den Hybrid-Ansatz.

Tag 2 · 10:30

Hybrid-Retrieval bauen

Hands-on: Volltextsuche neben dem Vektorindex, Score-Fusion und Gewichtung. Jeder misst den Effekt am eigenen Fragenkatalog statt ihn zu glauben.

Tag 2 · 13:00

Reranking: Präzision auf den letzten Metern

Einbau eines Rerankers in die eigene Pipeline, Latenz- und Kostenmessung inklusive. Entscheidung am Messwert: Wann lohnt der zusätzliche Schritt und wann nicht?

DIE 20 TRACKS · TRACK 14/20 · FORTSETZUNG

Tag 2 · 15:00	Das Evaluationsset entsteht Aus den mitgebrachten Fragenkatalogen wird ein strukturiertes Evaluationsset mit Erwartungsantworten und Quellenangaben. Automatisierter Lauf gegen die eigene Pipeline liefert die erste Baseline.
Tag 2 · 16:30	Tagesergebnis: gemessene Qualität statt Bauchgefühl Jede Person legt Zahlen vor: Trefferquote und Antwortqualität für naive Suche, Hybrid und Hybrid mit Reranking im Vergleich. Die Messtabelle ist das überprüfbare Ergebnis des Tages.
Tag 3 · 09:00	Berechtigungen: die Frage, die Projekte stoppt Demonstration am Lab: ein RAG ohne Berechtigungsprüfung beantwortet Gehaltsfragen aus einem HR-Dokument. Danach Architekturen für document-level security: ACL-Übernahme, Filterung zur Abfragezeit, Trusted-Identity-Durchreichung.
Tag 3 · 10:30	Document-level security implementieren Hands-on: Berechtigungsmetadaten im Index, Sicherheitsfilter in der Retrieval-Schicht, Test mit zwei Nutzerrollen gegen denselben Bestand. Verifikation: Die gesperrte Rolle bekommt weder Antwort noch Zitat.
Tag 3 · 13:00	Antwortqualität: Grounding, Zitate, Ich-weiß-es-nicht Prompting der Generierungsschicht: Antworten mit Quellenzitat, Verhalten bei fehlendem Kontext, Umgang mit widersprüchlichen Dokumenten. Messlauf gegen das Evaluationsset zeigt den Effekt jeder Änderung.
Tag 3 · 15:00	Grenzdokument und Betriebsübergabe Jedes Team schreibt die ehrliche Grenzdokumentation: Welche Fragetypen beantwortet der Prototyp belastbar, welche nicht, was fehlt bis zur Produktion? Ausblick auf Observability und Release-Gates als Brücke zu AI-EVAL.
Tag 3 · 16:15	Tagesergebnis: der Prototyp im Review Abschluss-Review je Person: Live-Abfrage mit zwei Rollen, Evaluationsergebnis, Grenzdokument. Die Runde prüft gegen eine Checkliste: Läuft es, misst es, schützt es, und steht ehrlich drin, was fehlt?

Das nehmen Sie mit

- Ein lauffähiger RAG-Prototyp auf der Lab-Umgebung, gebaut vom eigenen Team, mit nachvollziehbarer Architekturentscheidung je Komponente
- Eine Ingestion-Pipeline mit dokumentierten Chunking-Strategien für Fließtext, Tabellen und strukturierte Dokumente
- Hybrid-Retrieval aus Vektor- und Volltextsuche mit Reranking, im direkten Messvergleich zur naiven Vektorsuche
- Ein Evaluationsset mit 30-50 Frage-Antwort-Paaren und ein erster gemessener Qualitätsstand als Baseline
- Ein implementiertes Berechtigungsmodell mit document-level security und ein Grenzdokument, das ehrlich festhält, was der Prototyp nicht kann

Voraussetzungen

Solide Programmierkenntnisse in Python, Grundverständnis von APIs und Datenbanken. Vorerfahrung mit Embeddings ist hilfreich, aber nicht Bedingung.

Abschluss

Das Abschluss-Review am Tag 3 dokumentiert je Person den lauffähigen Prototyp, die Evaluationsergebnisse und das Grenzdokument gegen eine feste Checkliste. Teilnahme und Ergebnisse werden nachvollziehbar und prüffähig dokumentiert.

Enthalten

3 Trainingstage mit Dino Bordonaro, wahlweise inhouse oder im Lab · Persönliche Umgebung je Teilnehmendem auf unserem Enterprise-Lab für die Dauer des Trainings plus 14 Tage Nachlauf · Vollständiges Code-Repository mit Referenzimplementierung, Übungsständen und Musterlösungen · Vorlagen für Evaluationsset, Berechtigungsmodell und Grenzdokumentation · Teilnahmezertifikate und Schulungsdokumentation für Ihr Compliance-Archiv · 90-minütige Remote-Sprechstunde 4 Wochen nach dem Training für Fragen aus der eigenen Umsetzung

AIA-AI-EVAL

PFAD 4 · BAUEN



k/aia-ai-eval

Evaluation, Observability und Guardrails

Ohne Golden Dataset und Release-Gates ist jedes KI-Release ein Experiment an Ihren Nutzern.

Ein Golden Dataset, automatisierte Qualitätsmetriken, Tracing, Kosten- und Latenzbudgets sowie definierte Release-Gates, die Regressionen erkennen, bevor Nutzer sie erleben.

Für wen: Entwickler, Platform Engineers und technische Verantwortliche, die ein KI-System in Richtung Produktion bringen oder bereits betreiben.

Dauer: 2 Tage**Gruppe:** 4-12 Personen**Vorlauf:** ab 4 Wochen**15.800 €** pauschal je Durchführung

AGENDA

Tag 1 · 09:00	Der Blindflug-Befund Live-Experiment am Referenzsystem: ein scheinbar harmloser Prompt-Fix verschlechtert messbar die Hälfte der Antworten. Frage an die Runde: Wie viele Ihrer letzten Änderungen gingen ohne Messung live?
Tag 1 · 10:30	Das Golden Dataset bauen Werkstatt an den mitgebrachten Fällen: Erwartungsergebnisse formulieren, Randfälle und Fangfragen ergänzen, Repräsentativität prüfen. Ergebnis ist ein versioniertes Dataset mit 30-50 Fällen je Team.
Tag 1 · 13:00	Metriken implementieren: Groundedness, Relevanz, Vollständigkeit Hands-on: automatisierte Bewertung mit LLM-as-Judge, Kalibrierung gegen menschliche Stichproben-Urteile, Umgang mit Judge-Fehlurteilen. Jede Metrik läuft am Ende gegen das eigene Dataset.
Tag 1 · 15:00	Der erste Evaluationslauf Vollständiger Lauf des eigenen Datasets gegen das System, Auswertung der Ausreißer: Ist der Fehler im Retrieval, im Prompt oder im Judge? Die Fehlerklassifikation entscheidet über die richtige Korrektur.
Tag 1 · 16:30	Tagesergebnis: eine belastbare Baseline Jedes Team legt seine Baseline vor: Dataset-Umfang, Metrikwerte, klassifizierte Ausreißer. Prüffrage: Würden Sie auf Basis dieser Zahlen eine Änderung freigeben oder ablehnen können?
Tag 2 · 09:00	Tracing: die Kette sichtbar machen Hands-on: Instrumentierung der Pipeline mit Traces je Anfrage, von Retrieval-Treffern über Modellaufrufe bis zur Antwort. Übung: einen realen Schlechtfall allein anhand des Trace diagnostizieren.
Tag 2 · 10:30	Kosten- und Latenzbudget Auswertung der Trace-Daten: Was kostet welcher Anfragetyp in Token und Millisekunden? Jedes Team definiert Budgets und Alarmschwellen und verankert sie im Monitoring.
Tag 2 · 13:00	Guardrails vor und nach dem Modell Implementierung: Eingabefilter gegen Prompt Injection und Datenabfluss, Ausgabeprüfung auf Grounding und Formatverstöße. Messübung: Was fangen die Guardrails, was kostet das an Latenz?

DIE 20 TRACKS · TRACK 15/20 · FORTSETZUNG

Tag 2 · 15:00

Release-Gates und Regression-Erkennung

Werkstatt: Evaluationslauf als Pflicht-Stage in der Deployment-Pipeline, Schwellwerte je Metrik, Rollback-Kriterien. Test am eigenen System: eine absichtlich eingebaute Regression muss vom Gate gestoppt werden.

Tag 2 · 16:15

Tagesergebnis: das Gate hält

Abnahme je Team: Die präparierte Regression wird vom eigenen Release-Gate erkannt und blockiert, die saubere Version passiert. Das dokumentierte Gate-Regelwerk ist das überprüfbare Abschlussartefakt.

Das nehmen Sie mit

- Ein Golden Dataset für das eigene System: kuratierte Fälle mit Erwartungsergebnissen, Randfällen und bewusst schwierigen Anfragen
- Implementierte Qualitätsmetriken für Groundedness, Relevanz und Vollständigkeit, inklusive kalibriertem LLM-as-Judge mit Stichproben-Gegenprüfung
- Tracing über die gesamte Kette von der Anfrage über Retrieval und Modellaufruf bis zur Antwort, mit Kosten und Latenz je Schritt
- Ein Kosten- und Latenzbudget je Anfragetyp mit Alarmschwellen
- Definierte Release-Gates: welche Messwerte ein Deployment passieren muss und wann ein Rollback ausgelöst wird

Voraussetzungen

Programmierkenntnisse in Python und ein reales oder im Training bereitgestelltes KI-System als Messobjekt. Erfahrung mit CI/CD-Pipelines ist hilfreich.

Abschluss

Transfernachweis ist der bestandene Gate-Test am Tag 2: Regression erkannt, saubere Version freigegeben, Regelwerk dokumentiert. Teilnahme und Ergebnisse werden nachvollziehbar und prüffähig dokumentiert.

Enthalten

2 Trainingstage mit Dino Bordonaro, inhouse, remote oder im Lab · Persönliche Lab-Umgebung mit vorbereitetem Referenzsystem und Evaluations-Stack, plus 14 Tage Nachlauf · Code-Repository mit Metrik-Implementierungen, Judge-Prompts und Pipeline-Vorlagen · Vorlagen für Golden Dataset, Budgetdefinition und Release-Gate-Checkliste · Teilnahmezertifikate und Schulungsdokumentation für Ihr Compliance-Archiv · 60-minütige Remote-Sprechstunde 4 Wochen nach dem Training zum Review der eigenen Gates

AIA-AI-AGENT

PFAD 4 · BAUEN



k/aia-ai-agent

Agentic Systems Engineering

Ein Agent, der handeln darf, bevor er Gates bestanden hat, ist keine Automatisierung, sondern ein Betriebsrisiko mit API-Zugriff.

Ein selbst gebauter agentischer Workflow mit Tool-Anbindung, explizitem Zustandsmodell, Human-in-the-Loop-Gates, Evaluation sowie Abbruch- und Eskalationslogik, abgenommen im Härtestest.

Für wen: Erfahrene Entwickler und Architekten, die Agenten an reale Geschäftsprozesse und Systeme anbinden sollen.

Dauer: 4 Tage**Gruppe:** 4-12 Personen**Vorlauf:** ab 6 Wochen**31.600 €** pauschal je Durchführung

AGENDA

Tag 1 · 09:00

Autopsie eines Agenten-Fehlschlags

Live am Lab: ein naiver Agent gerät in eine Tool-Schleife und verbrennt sichtbar Budget. Analyse der Ursachen und Ableitung der Architekturprinzipien: expliziter Zustand, begrenzte Autonomie, prüfbare Schritte.

Tag 1 · 10:30

Anatomie eines belastbaren Agenten

Architektur-Durchstich: Planner, Tool-Schicht, Speicher, Prüfschicht. Abgrenzungsfrage je mitgebrachtem Prozess-Steckbrief: Braucht das einen Agenten oder reicht ein Workflow mit einem LLM-Schritt?

Tag 1 · 13:00

Tools entwerfen, die ein Modell sicher bedienen kann

Hands-on: Tool-Definitionen mit engen Schemata, Idempotenz, Fehlerrückgaben, Least-Privilege-Zugriff. Jeder baut die ersten zwei Tools des Übungsprozesses und testet sie isoliert.

Tag 1 · 15:00

Der erste Durchstich

Der eigene Agent führt den Übungsprozess erstmals Ende-zu-Ende gegen die Mock-Systeme aus, noch ohne Gates. Beobachtungsauftrag: Wo trifft er Annahmen, die niemand geprüft hat?

Tag 1 · 16:30

Tagesergebnis: laufender Agent mit Befundliste

Jede Person zeigt den laufenden Durchstich und legt eine Befundliste vor: mindestens drei beobachtete ungeprüfte Annahmen oder Risikostellen. Diese Liste wird an Tag 2 und 3 systematisch abgearbeitet.

Tag 2 · 09:00

Zustand explizit machen

Vom Prompt-Gedächtnis zum Zustandsmodell: Zustände, Übergänge, erlaubte Aktionen je Zustand. Jeder modelliert seinen Übungsprozess als Graph, bevor wieder Code entsteht.

Tag 2 · 10:30

Das Zustandsmodell implementieren

Hands-on: Orchestrierung mit persistiertem Zustand, Wiederaufnahme nach Absturz, deterministische Übergänge. Test: Prozess in der Mitte abbrechen und korrekt fortsetzen.

DIE 20 TRACKS · TRACK 16/20 · FORTSETZUNG

Tag 2 · 13:00	Kontext- und Speicherstrategie Was der Agent je Schritt wissen muss und was nicht: Arbeitskontext, Langzeitwissen über RAG-Anbindung, Kontextbudget. Messübung: Kontextdisziplin gegen Vollkontext im Kosten- und Qualitätsvergleich.
Tag 2 · 15:00	Fehlerpfade zuerst Übung an den Mock-Systemen mit eingebauten Störungen: Timeouts, widersprüchliche Daten, halbfertige Transaktionen. Der Agent muss jeden Fehlerpfad in einen definierten Zustand überführen statt zu improvisieren.
Tag 2 · 16:30	Tagesergebnis: kontrollierter Umgang mit Störungen Abnahme je Person: Der Agent durchläuft ein Störungs-Szenario und landet nachweisbar in definierten Zuständen, dokumentiert im Zustandsprotokoll. Improvisierte Weiterläufe gelten als nicht bestanden.
Tag 3 · 09:00	Verantwortung erst nach Gates Das Prinzip des Tracks im Detail: Aktionsklassen von lesend bis irreversibel, je Klasse ein Gate-Typ. Werkstatt am eigenen Prozess-Steckbrief: Welche Aktion bekommt welches Gate und wer gibt frei?
Tag 3 · 10:30	Human-in-the-Loop implementieren Hands-on: Freigabe-Schritt mit Vorschau der geplanten Aktion, Begründung des Agenten und explizitem Approve oder Reject. Der Agent wartet nachweisbar, statt bei ausbleibender Antwort weiterzulaufen.
Tag 3 · 13:00	Abbruch- und Eskalationslogik Implementierung von Budgetgrenzen für Schritte, Kosten und Laufzeit, Erkennung von Schleifen und Zielabweichung, Eskalation an definierte Empfänger mit Zustandsübergabe. Test: der Schleifen-Agent von Tag 1 gegen die neuen Grenzen.
Tag 3 · 15:00	Sicherheit an der Tool-Grenze Übung: Prompt Injection über Tool-Rückgaben und Dokumenteninhalte, die den Agenten zu ungeplanten Aktionen verleiten soll. Härtung durch Eingabesäuberung, Aktionsvalidierung und Gate-Pflicht für sensible Klassen.
Tag 3 · 16:30	Tagesergebnis: Gates unter Beschuss Abnahme je Person: Drei vorbereitete Stör- und Fehlerszenarien laufen gegen den eigenen Agenten. Bestanden ist, wer jede kritische Aktion nachweisbar vor dem Gate stoppt und die Eskalation korrekt auslöst.
Tag 4 · 09:00	Agenten evaluieren: Trajektorien statt Antworten Evaluationsdesign für mehrschrittige Systeme: Erfolgskriterien je Prozessziel, Bewertung von Handlungspfaden, Kosten je erfolgreichem Durchlauf. Aufbau des Evaluationssets für den eigenen Übungsagenten.
Tag 4 · 10:30	Der Evaluationslauf Hands-on: automatisierter Lauf mit Erfolgsquote, Gate-Auslösungen, Abbruchgründen und Kosten je Durchlauf. Auswertung: Welche Fehlklasse dominiert und welche Architekturänderung adressiert sie?
Tag 4 · 13:00	Das Abnahmeprotokoll Werkstatt: Jedes Team schreibt das Abnahmeprotokoll seines Agenten: bestandene Gates, Evaluationswerte, Restrisiken, Bedingungen für Verantwortungsübernahme und für den Entzug. Ehrlichkeit ist Prüfkriterium.
Tag 4 · 15:00	Übertrag auf den eigenen Prozess Rückkehr zum mitgebrachten Steckbrief: Architektur-Skizze, Gate-Zuordnung und Aufwandsschätzung für den realen Fall im Haus. Peer-Review der Skizzen in Zweiergruppen mit dokumentierten Einwänden.

DIE 20 TRACKS · TRACK 16/20 · FORTSETZUNG

Tag 4 · 16:15

Tagesergebnis: Abnahme im Härtetest

Finale je Person: Live-Durchlauf des Agenten vor der Gruppe inklusive eines nicht vorab bekannten Störszenarios, danach Übergabe von Abnahmeprotokoll und Übertrags-Skizze. Die Checkliste der Runde entscheidet über die Abnahme.

Das nehmen Sie mit

- Ein lauffähiger agentischer Workflow auf der Lab-Umgebung: Planung, Tool-Aufrufe, Ergebnisprüfung, gebaut vom eigenen Team
- Ein explizites Zustandsmodell mit definierten Übergängen statt impliziter Schleifen, inklusive persistierbarem Zustand für Wiederaufnahme
- Implementierte Human-in-the-Loop-Gates: welche Aktionen der Agent vorschlagen darf und welche erst nach Freigabe ausgeführt werden
- Abbruch- und Eskalationslogik mit Budgetgrenzen für Schritte, Kosten und Zeit sowie definierten Eskalationsempfängern
- Ein Evaluations- und Abnahmeprotokoll, das dokumentiert, unter welchen Bedingungen der Agent Verantwortung übernehmen darf und unter welchen nicht

Voraussetzungen

Sichere Programmierpraxis in Python und API-Erfahrung. Vorkenntnisse aus AI-RAG oder AI-EVAL sind hilfreich, insbesondere ein Grundverständnis von Evaluation.

Abschluss

Transfernachweis ist der bestandene Härtetest am Tag 4 samt Abnahmeprotokoll: Gates halten unter Störung, Eskalation greift, Restrisiken sind ehrlich dokumentiert. Teilnahme und Ergebnisse werden nachvollziehbar und prüffähig dokumentiert.

Enthalten

4 Trainingstage mit Dino Bordonaro, wahlweise inhouse oder im Lab · Persönliche Lab-Umgebung mit Übungsprozess, Mock-Systemen und lokal betriebenen Modellen, plus 14 Tage Nachlauf · Code-Repository mit Referenzarchitektur, Zustandsmodell-Vorlagen und Musterlösungen je Tagesstand · Gate-Katalog und Abnahmeprotokoll-Vorlage für die Übertragung auf eigene Prozesse · Teilnahmezertifikate und Schulungsdokumentation für Ihr Compliance-Archiv · 90-minütige Remote-Sprechstunde 4 Wochen nach dem Training zum Review des eigenen Agenten-Designs

AIA-AI-PLATFORM

PFAD 4 · BAUEN

k/aia-ai-platform

Sovereign AI Platform Engineering

Fünf Tage, nach denen Ihr Team eine Inferenz-Plattform nicht nur beschreiben kann, sondern eine gebaut, vermessen und übergeben hat.

Eine referenzierbare Plattformarchitektur von GPU-Sizing bis Betriebsübergabe, im Lab durchgebaut, vermessen und als Blaupause dokumentiert.

Für wen: Plattform-Teams und Infrastruktur-Architekten, die LLM-Inferenz im eigenen Rechenzentrum oder auf Azure Local verantworten sollen.

Dauer: 5 Tage **Gruppe:** 4-12 Personen **Vorlauf:** ab 6 Wochen

39.500 € pauschal je Durchführung

AGENDA

Tag 1 · 09:00	Ihre Lastannahmen auf dem Prüfstand Auswertung der Umgebungs-Steckbriefe: Wie viele gleichzeitige Nutzer, welche Modellklassen, welche Antwortzeiten? Aus vagen Erwartungen werden messbare Anforderungen in Tokens pro Sekunde.
Tag 1 · 10:30	Die Referenzarchitektur Schicht für Schicht Compute, Netzwerk, Storage, Identity, Registry, Inference, Gateway. Für jede Schicht die eine Entscheidung, die später teuer zu korrigieren ist, und woran BORDONARO sie in der Sovereign AI Appliance festgemacht hat.
Tag 1 · 13:00	GPU-Kapazitätsmodell von L40S bis H100 VRAM-Bedarf, Quantisierung, Batching, Concurrency: was die Wahl der Karte wirklich bestimmt. Live-Benchmark im Lab, dasselbe Modell auf zwei GPU-Klassen im direkten Vergleich.
Tag 1 · 15:00	Übung: Sizing-Rechnung für Ihr Szenario Jedes Team rechnet sein eigenes Szenario mit dem Sizing-Blatt durch und verteidigt die Rechnung vor der Gruppe. Die häufigste Erkenntnis: erst falsch dimensioniert, dann korrigiert, dann verstanden.
Tag 1 · 16:15	Tagesergebnis: dokumentiertes Kapazitätsmodell Jedes Team hat ein begründetes Kapazitätsmodell für sein Szenario, geprüft gegen die Benchmark-Werte des Tages. Es wird im Laufe der Woche gegen die echte Plattform validiert.
Tag 2 · 09:00	Azure Local als Fundament Cluster-Anatomie, Storage Spaces Direct, Rollen und Grenzen der Plattform in der jeweils verfügbaren Version. Was Azure Local kann, was es nicht kann und wo ein reiner Kubernetes-Unterbau die ehrlichere Antwort ist.
Tag 2 · 10:45	Identity durchgängig gedacht Entra ID, lokale Identitäten, Workload-Identity und Secrets-Verwaltung ohne Klartext in Configs. Konkrete Frage: Wer darf welches Modell mit welchen Daten aufrufen, und wo wird das durchgesetzt?
Tag 2 · 13:00	Netzwerkdesign und Egress-Kontrolle Segmentierung der Inference-Zone, TLS im Innenverhältnis, kontrollierte Ausgänge. Übung am Whiteboard: Welche sieben Verbindungen braucht die Plattform wirklich, und welche davon lassen sich schließen?

DIE 20 TRACKS · TRACK 17/20 · FORTSETZUNG

Tag 2 · 15:00	Übung: Fundament-Design im Lab umsetzen Jedes Team richtet Identity-Anbindung und Netzsegmentierung für seine Lab-Umgebung ein und weist mit einem Testzugriff nach, dass die Grenzen halten.
Tag 2 · 16:15	Tagesergebnis: freigegebenes Fundament Identity und Netz stehen und sind getestet. Jedes Team dokumentiert sein Fundament-Design in der Blaupause, offene Streitfragen wandern begründet ins Entscheidungslog.
Tag 3 · 09:00	Registry: Herkunft und Signatur von Modellen Modell- und Container-Registry als eine Disziplin: Versionierung, Signaturprüfung, Herkunftsnachweis. Warum ein unsigniertes Modell aus dem Internet in keiner souveränen Plattform landen darf.
Tag 3 · 10:30	Inference-Server im ehrlichen Vergleich vLLM und Alternativen an denselben Messpunkten: Durchsatz, Latenz, Betriebsaufwand. Dazu der Blick auf Foundry Local auf Azure Local in der jeweils verfügbaren Version, derzeit Preview, mit klarer Einordnung, was das für Produktionsentscheidungen bedeutet.
Tag 3 · 13:00	Aufbau: vom Modell zum API-Endpunkt Jedes Team baut den Stack durch: Modell aus der Registry ziehen, Signatur prüfen, servieren, hinter das API-Gateway legen. Am Ende antwortet die eigene Plattform auf den ersten authentifizierten Request.
Tag 3 · 15:00	Lasttest und Tuning Lastgenerator gegen den eigenen Stack: Batching-Parameter, KV-Cache und Quantisierung verändern und die Wirkung messen. Das Kapazitätsmodell von Tag 1 trifft auf die Realität.
Tag 3 · 16:15	Tagesergebnis: lauffähiger Stack mit Messprotokoll Jedes Team hat einen funktionierenden Inference-Stack und ein Messprotokoll, das die Abweichung zwischen Sizing-Rechnung und gemessener Leistung ausweist und erklärt.
Tag 4 · 09:00	Observability für GPU-Plattformen GPU-Auslastung, Token-Durchsatz, Warteschlangen, Fehlerraten: welche Metriken ein Betriebsteam wirklich braucht und welche Alarme nachts wecken dürfen. Aufbau des Dashboards auf der eigenen Lab-Plattform.
Tag 4 · 10:45	Update-Pfad: Modelle, Treiber, Firmware, Container Vier Update-Ströme mit vier Risikoprofilen. Konkrete Entscheidung je Strom: Wartungsfenster, Canary oder rollierend, und was vor jedem Modellwechsel zwingend gemessen wird.
Tag 4 · 13:00	Mandanten, Quoten und interne Verrechnung Mehrere Teams auf einer GPU-Plattform: Quotenmodell, Priorisierung, Kostenzuordnung. Übung: ein Quotenmodell für drei konkurrierende Abteilungen entwerfen und begründen.
Tag 4 · 15:00	Störungsübung: der Node fällt aus Wir ziehen einem Team kontrolliert einen GPU-Node weg. Erkennen, entscheiden, umleiten, kommunizieren: die Störung wird nach Runbook abgearbeitet und anschließend seziert.
Tag 4 · 16:15	Tagesergebnis: Betriebs-Dashboard und Störungsprotokoll Jedes Team hat ein funktionierendes Monitoring-Dashboard und ein ausgefülltes Störungsprotokoll mit Zeitleiste, Entscheidungen und Verbesserungen fürs eigene Runbook.

DIE 20 TRACKS · TRACK 17/20 · FORTSETZUNG

Tag 5 · 09:00	Das Runbook wird vollständig Startprozeduren, Update-Abläufe, Störungsszenarien, Eskalationswege: die Fragmente der Woche werden zu einem Runbook-Gerüst zusammengeführt, das ein Betriebsteam übernehmen kann.
Tag 5 · 10:30	Architektur-Review gegen Ihre Anforderungen Die Blaupause jedes Teams im Review gegen Datenschutz-Anforderungen, interne Sicherheitsvorgaben und Anforderungen bis VS-NfD, für die die Architektur vorbereitbar sein soll. Lücken werden benannt, nicht wedgdiskutiert.
Tag 5 · 13:00	Abnahme-Simulation: Übergabe an den Betrieb Rollenspiel mit vertauschten Teams: ein Team übergibt seine Plattform, das andere prüft nach Abnahme-Checkliste und darf alles fragen. Was die Übergabe nicht überlebt, wird nachgebessert.
Tag 5 · 15:00	Transferplan für Montag Jedes Team priorisiert die ersten drei Schritte im eigenen Haus: Was wird übernommen, was wird anders entschieden, was braucht eine Beschaffung? Terminierung des Follow-ups.
Tag 5 · 16:15	Abschlussergebnis: referenzierbare Blaupause Jedes Team verlässt den Kurs mit einer dokumentierten Plattform-Blaupause, validiertem Kapazitätsmodell, Runbook-Gerüst und Entscheidungslog. Kurzes Review der Ergebnisse durch den Trainer als Abschluss.

Das nehmen Sie mit

- Ein Referenzarchitektur-Dokument über alle Schichten: Compute, Netzwerk, Storage, Identity, Registry, Inference-Stack
- Ein GPU- und Kapazitätsmodell mit nachvollziehbarer Sizing-Rechnung von L40S bis H100 für Ihre Lastannahmen, gestützt auf eigene Messwerte aus dem Lab
- Ein lauffähiger Inference-Stack, den Ihr Team selbst aufgebaut hat: Modell-Registry, Inference-Server, API-Gateway, Monitoring
- Ein Betriebsübergabe-Paket mit Runbook-Gerüst, Update-Pfad und Störungsprotokollen aus den Übungen
- Ein Entscheidungslog zu den strittigen Architekturfragen Ihres Hauses mit dokumentierter Begründung

Voraussetzungen

Solide Infrastruktur-Grundlagen: Virtualisierung, Netzwerk, Container und Linux-Kommandozeile. Azure-Grundkenntnisse helfen, sind aber keine Bedingung.

Abschluss

Die Abnahme-Simulation am Tag 5 ist der Transfernachweis: jedes Team übergibt seine Plattform nach Checkliste und besteht das Review des Partner-Teams. Blaupause, Messprotokolle und Störungsprotokoll werden dokumentiert und den Teilnehmenden übergeben.

Enthalten

5 Trainingstage, gehalten von Dino Bordonaro · Dedizierte Lab-Umgebung mit NVIDIA-GPUs je Zweierteam für die gesamte Woche, aus unserem Enterprise-Lab mit 2.516 Cores und 24 TB RAM · Referenzarchitektur-Vorlagen und Infrastructure-as-Code-Bausteine zur Weiterverwendung · Sizing-Rechenblatt mit den im Kurs erhobenen Benchmark-Werten · Teilnahmezertifikate und prüffähige Schulungsdokumentation · 60-minütiges Follow-up mit dem Team 4 Wochen nach dem Kurs

AIA-AI-OPS

PFAD 4 · BAUEN



k/aia-ai-ops

LLMOps für verbundene, getrennte und Air-Gap-Umgebungen

Drei Tage, nach denen ein Modell-Update auch dann sicher ankommt, wenn kein Kabel nach draußen führt.

Ein getesteter Promotion-Prozess von dev über test nach prod, ein Offline-Updatekonzept mit Signaturprüfung und ein im Lab geprobtes Incident- und Rollback-Runbook.

Für wen: Plattform- und Betriebsteams, die eine LLM-Plattform in Umgebungen mit eingeschränkter oder fehlender Internetanbindung verantworten: Behörden, kritische Infrastruktur, Verteidigung, produzierende Industrie mit getrennten Netzen.

Dauer: 3 Tage**Gruppe:** 4-12 Personen**Vorlauf:** ab 4 Wochen**23.700 €** pauschal je Durchführung

AGENDA

Tag 1 · 09:00	Drei Betriebsmodelle, drei Wahrheiten Verbunden, zeitweise getrennt, dauerhaft getrennt: was sich je Modell an Update-Frequenz, Monitoring und Supportfähigkeit wirklich ändert. Ehrliche Einordnung von Azure Local Disconnected Operations: zugangsbeschränkt, mit Eligibility-Prüfung, kein frei bestellbares Produkt.
Tag 1 · 10:45	Alles ist ein Artefakt Modelle, Container, Prompts, Konfigurationen und Evaluationsdaten als versionierte, signierte Artefakte. Konkrete Frage je Artefakttyp: Wo entsteht es, wer signiert es, was beweist die Signatur?
Tag 1 · 13:00	Promotion-Gates, die etwas aufhalten Der Weg von dev über test nach prod mit Gates, die scharf gestellt sind: Evaluations-Suite, Signaturprüfung, Vier-Augen-Freigabe. Diskussion am Fall: Welches Gate hätte den letzten schlechten Modellwechsel Ihres Hauses gestoppt?
Tag 1 · 15:00	Übung: Promotion-Pipeline im Lab Jedes Team baut die Pipeline nach: ein Modellartefakt durchläuft dev und test, ein absichtlich verschlechtertes Modell muss am Evaluations-Gate hängen bleiben. Tut es das nicht, wird das Gate nachgeschärft.
Tag 1 · 16:15	Tagesergebnis: dokumentierter Promotion-Prozess Jedes Team hat einen Promotion-Prozess mit definierten Gates als Dokument und als lauffähige Pipeline, nachgewiesen durch ein Modell, das durchkam, und eines, das gestoppt wurde.
Tag 2 · 09:00	Anatomie eines Transferpakets Was über die Schleuse geht: Modellgewichte, Container-Images, Signaturen, Manifeste, Evaluationsdaten. Aufbau eines Manifests, das Vollständigkeit und Unversehrtheit prüfbar macht, bevor irgendetwas installiert wird.
Tag 2 · 10:45	Die Schleuse als Prozess, nicht als USB-Stick Rollen, Prüfschritte und Protokollierung beim Übergang in die getrennte Zone. Welche Prüfungen laufen außen, welche innen, und wer trägt die Freigabe? Abgleich mit den mitgebrachten Netzskizzen.

DIE 20 TRACKS · TRACK 18/20 · FORTSETZUNG

Tag 2 · 13:00	Übung: Offline-Update unter Realbedingungen Die Lab-Zone jedes Teams ist vom Internet getrennt. Ein neues Modell samt Container wird als Transferpaket gebaut, eingeschleust, gegen das Manifest geprüft und im Staging aktiviert. Ein manipuliertes Paket ist im Umlauf und muss an der Signaturprüfung scheitern.
Tag 2 · 15:30	Nachbereitung: wo es hakte Auswertung der Übung entlang der Protokolle: Welche Prüfung hat gefehlt, welche war doppelt, wie lange hat der Durchlauf gedauert? Die Erkenntnisse fließen direkt ins eigene Updatekonzept.
Tag 2 · 16:15	Tagesergebnis: durchgeführtes Offline-Update mit Protokoll Jedes Team hat ein Update ohne Internetzugang vollständig durchgeführt und das manipulierte Paket nachweislich abgewiesen. Das Schleusenprotokoll dokumentiert jeden Schritt.
Tag 3 · 09:00	Monitoring ohne Cloud-Backend Modellqualität, Drift und Plattformgesundheit sichtbar machen, wenn kein SaaS-Dashboard erlaubt ist. Aufbau der Metrik-Pipeline in der getrennten Zone, inklusive der Frage, welche Auswertungen periodisch nach außen dürfen und welche nie.
Tag 3 · 10:45	Incident-Szenarien für LLM-Plattformen Modellregression nach Update, Verdacht auf kompromittiertes Artefakt, GPU-Ausfall, Prompt-Injection-Welle: für jedes Szenario Erkennungsweg, Sofortmaßnahme und Kommunikationskette im Runbook festlegen.
Tag 3 · 13:00	Übung: Rollback unter Zeitdruck Das am Vortag eingespielte Modell zeigt in der Übung eine massive Regression. Jedes Team rollt nach Runbook auf den letzten guten Stand zurück, misst die Zeit und weist per Evaluation nach, dass der alte Zustand wiederhergestellt ist.
Tag 3 · 15:00	Ihr 90-Tage-Plan für den echten Betrieb Priorisierung im Team: welche drei Lücken zwischen Lab-Übung und eigener Umgebung zuerst geschlossen werden, wer die Signaturkette im Haus verantwortet, wann der erste eigene Schleusen-Testlauf stattfindet.
Tag 3 · 16:15	Tagesergebnis: getestetes Incident- und Rollback-Runbook Jedes Team hat ein Runbook, das eine echte Rollback-Übung überstanden hat, samt gemessener Wiederherstellungszeit und dokumentiertem 90-Tage-Plan für die eigene Umgebung.

Das nehmen Sie mit

- Ein dokumentierter Artefaktfluss für Modelle, Container, Prompts und Konfigurationen mit Versionierung und Signaturkette
- Ein Promotion-Prozess dev-test-prod mit definierten Gates, inklusive Evaluations-Gate vor jedem Modellwechsel
- Ein Offline-Updatekonzept: Transferpaket, Schleusenprozess, Signaturprüfung und Staging ohne Internetzugang, im Lab einmal vollständig durchgespielt
- Ein Incident- und Rollback-Runbook für Modellregression, kompromittierte Artefakte und Plattformausfall, in einer Übung getestet
- Ein Monitoring-Konzept, das ohne Cloud-Backend auskommt und trotzdem Modellqualität sichtbar macht

Voraussetzungen

Betriebserfahrung mit Containern und CI/CD-Grundverständnis. Ideal ist der vorherige Besuch von AI-PLATFORM oder eine bestehende Inference-Umgebung im eigenen Haus.

Abschluss

Drei dokumentierte Nachweise: die Pipeline, die ein schlechtes Modell stoppt, das Schleusenprotokoll des Offline-Updates mit abgewiesenem Manipulationspaket und die gemessene Rollback-Übung. Alle drei werden den Teilnehmenden als Dokumentation übergeben.

Enthalten

3 Trainingstage, gehalten von Dino Bordonaro · Lab-Umgebung mit einer bewusst abgeschotteten Netzzone je Team, in der die Offline-Übungen real stattfinden · Vorlagen für Transferpaket-Manifest, Promotion-Gates und Rollback-Runbook · Teilnahmezertifikate und prüffähige Schulungsdokumentation · 60-minütiges Follow-up mit dem Team 4 Wochen nach dem Kurs

AIA-AI-ARC

PFAD 4 · BAUEN



k/aia-ai-arc

Azure Arc und Azure Local als Hybrid-Control-Plane

Drei Tage, nach denen Arc nicht mehr nur ein Agent auf Ihren Servern ist, sondern eine Control-Plane mit Regeln, die Sie selbst gesetzt haben.

Ein durchdachtes Arc-Ressourcenmodell mit Policy- und GitOps-Grundlage und ein begründetes Entscheidungsbild, wann Connected und wann Disconnected das richtige Betriebsmodell ist.

Für wen: Plattform-Teams, die Azure Local betreiben oder einführen und die Verwaltung über Azure Arc sauber aufsetzen wollen, statt sie zu erben.

Dauer: 3 Tage **Gruppe:** 4-12 Personen **Vorlauf:** ab 4 Wochen

23.700 € pauschal je Durchführung

AGENDA

Tag 1 · 09:00	Was Arc wirklich projiziert Das Arc-Ressourcenmodell von innen: Server, Kubernetes, Datenservices und Azure Local als projizierte Ressourcen. Welche Daten fließen dabei in welche Richtung, und was bleibt lokal? Klärung am Datenfluss-Diagramm, nicht am Marketing-Bild.
Tag 1 · 10:45	Ressourcen-Topologie für Ihr Haus Management-Groups, Subscriptions, Ressourcengruppen, Tags und RBAC: Übung am eigenen Fall. Konkrete Entscheidung: Wo trennen Sie Produktion von Test, und wer darf auf Arc-Ebene was?
Tag 1 · 13:00	Onboarding im Lab Jedes Team verbindet Azure-Local-Instanzen und weitere Systeme mit Arc, in der jeweils verfügbaren Version. Dabei wird protokolliert, welche Endpunkte, Konten und Berechtigungen das Onboarding tatsächlich verlangt.
Tag 1 · 15:15	RBAC scharf stellen Rollenzuschnitt in der Praxis: Betriebsteam, Sicherheitsteam, Auditor. Test mit drei Identitäten, ob die Grenzen halten. Wer zu viel sieht oder zu viel darf, wird nachjustiert.
Tag 1 · 16:15	Tagesergebnis: dokumentiertes Ressourcenmodell Jedes Team hat seine Systeme in Arc onboarded und ein Ressourcenmodell mit RBAC-Zuschnitt als Diagramm dokumentiert, geprüft mit echten Testzugriffen.
Tag 2 · 09:00	Policy als Leitplanke, nicht als Bremse Azure Policy auf Arc-Ressourcen: Audit gegen Deny, Initiativen, Ausnahmen mit Ablaufdatum. Auswahlübung: Welche zehn Policies gehören in Ihre Startmenge, und welche einzige davon steht auf Deny?
Tag 2 · 10:45	GitOps-Grundlage legen Konfiguration als Code mit Flux auf Arc-enabled Kubernetes: Repository-Struktur, Umgebungs-Trennung, Freigabewege. Die Grundsatzfrage wird entschieden: Was darf nie wieder per Portal geändert werden?
Tag 2 · 13:00	Übung: Änderung per Pull-Request Jedes Team ändert eine Workload-Konfiguration ausschließlich über einen Pull-Request und beobachtet die Auslieferung durch die GitOps-Kette. Danach die Gegenprobe: eine manuelle Änderung am System wird als Drift erkannt und zurückgerollt.

DIE 20 TRACKS · TRACK 19/20 · FORTSETZUNG

Tag 2 · 15:15	Azure Local unter Arc im Alltag Updates, Monitoring und Kapazitätssicht auf Azure Local über die Arc-Ebene, in der jeweils verfügbaren Version. Ehrliche Bilanz: was die Control-Plane heute gut abdeckt und wo weiterhin lokale Werkzeuge nötig sind.
Tag 2 · 16:15	Tagesergebnis: laufende Policy- und GitOps-Kette Jedes Team hat zehn Policies zugewiesen, eine Änderung per Pull-Request ausgeliefert und einen Konfigurationsdrift nachweislich erkannt und korrigiert. Repositories und Policy-Definitionen gehen als Artefakte mit.
Tag 3 · 09:00	Connected und Disconnected ohne Ideologie Die beiden Betriebsmodelle im nüchternen Vergleich: Datenflüsse, Update-Wege, Funktionsumfang, Betriebsaufwand. Azure Local Disconnected Operations wird ehrlich eingeordnet: zugangsbeschränkt, mit Eligibility-Prüfung, nicht frei bestellbar, Funktionsumfang in der jeweils verfügbaren Version.
Tag 3 · 10:45	Der Kriterienkatalog Erarbeitung eines Entscheidungsrasters: regulatorische Auflagen, Latenz- und Autonomie-Anforderungen, Personal, Supportfähigkeit, Kosten. Jedes Kriterium bekommt eine Gewichtung und eine messbare Prüffrage.
Tag 3 · 13:00	Fallarbeit: Ihr Haus im Entscheidungsbild Jedes Team wendet den Katalog auf die eigene Ausgangslage aus der Vorarbeit an und erarbeitet eine Empfehlung: Connected, Disconnected-Prüfung anstoßen oder Mischmodell. Die Empfehlung wird vor der Gruppe verteidigt.
Tag 3 · 15:15	Die Eligibility-Realität und der Plan B Durchgehen des realen Wegs zu Disconnected Operations: Voraussetzungen, Prüfprozess, Zeithorizont. Für jedes Team wird ein Plan B festgehalten, falls die Prüfung scheitert oder zu lange dauert.
Tag 3 · 16:15	Tagesergebnis: begründete Betriebsmodell-Empfehlung Jedes Team hat ein dokumentiertes Entscheidungsbild mit gewichtetem Kriterienkatalog, einer Empfehlung für das eigene Haus und einer Liste der offenen Prüfpunkte samt Verantwortlichen.

Das nehmen Sie mit

- Ein Arc-Ressourcenmodell für Ihr Haus: Management-Groups, Subscriptions, Ressourcengruppen, Tags und RBAC-Zuschnitt, als Diagramm dokumentiert
- Eine Policy-Grundlage mit den zehn wichtigsten Regeln für den Start, als Code versioniert
- Eine funktionierende GitOps-Kette im Lab: Konfigurationsänderung per Pull-Request statt per Handarbeit
- Ein Entscheidungsbild Connected gegen Disconnected mit Kriterienkatalog, Datenfluss-Übersicht und ehrlicher Eligibility-Einordnung
- Eine dokumentierte Empfehlung für das Betriebsmodell Ihres Hauses mit den offenen Prüfpunkten

Voraussetzungen

Grundkenntnisse in Azure-Konzepten wie Subscriptions, Ressourcengruppen und RBAC sowie Betriebserfahrung mit Servern oder Kubernetes. Azure-Local-Praxis ist hilfreich, wird aber im Lab gestellt.

Abschluss

Drei dokumentierte Nachweise: das Ressourcenmodell mit bestandenen RBAC-Tests, die GitOps-Kette mit erkanntem Drift und die vor der Gruppe verteidigte Betriebsmodell-Empfehlung. Alles wird als Teil der Schulungsdokumentation übergeben.

Enthalten

3 Trainingstage, gehalten von Dino Bordonaro · Lab-Zugang mit Azure-Local-Instanzen und eigenem Azure-Tenant-Bereich je Team · Policy- und GitOps-Vorlagen als versionierte Repositories zur Weiterverwendung · Kriterienkatalog Connected gegen Disconnected als Entscheidungsvorlage · Teilnahmezertifikate und prüffähige Schulungsdokumentation · 60-minütiges Follow-up mit dem Team 4 Wochen nach dem Kurs

AIA-AI-RESIDENCY

PFAD 4 · BAUEN

k/aia-ai-residency

Sovereign AI Architecture Residency

Fünf Tage im Lab, an deren Ende keine Folien stehen, sondern Ihr eigener Architekturentwurf, validiert auf echter Hardware.

Ein kundeneigener Architektur- und Betriebsentwurf, dessen kritische Annahmen auf echter Hardware validiert wurden, mit getestetem Slice, priorisiertem Backlog und Executive Readout.

Für wen: Das Kernteam eines Hauses, das eine souveräne KI-Plattform konkret plant: Architekten, Plattform-Verantwortliche und die Person, die die Investition vertreten muss.

Dauer: 5 Tage im Lab **Gruppe:** 4-8 Personen **Vorlauf:** ab 6 Wochen

44.500 € pauschal je Durchführung

AGENDA

Tag 1 · 09:00	Zielbild und die drei kritischen Annahmen Ihr Vorhaben auf einer Wand: Workloads, Nutzer, Datenklassen, regulatorischer Rahmen. Die im Vorgespräch benannten Annahmen werden geschärft und in messbare Prüfkriterien übersetzt.
Tag 1 · 10:45	Anforderungen mit Vetorecht Datenschutz, Sicherheit und Betrieb formulieren ihre harten Grenzen, bevor die Architektur entsteht: Was darf das System nie tun, welche Nachweise bis VS-NfD sollen vorbereitbar sein, wer trägt den Betrieb?
Tag 1 · 13:00	Slice-Auswahl: ein Fall, der etwas beweist Aus Ihren Kandidaten-Workloads wird der Anwendungsfall gewählt, der die riskanteste Annahme testet, nicht der bequemste. Festlegung der Messlatte: Woran erkennen wir am Freitag, dass der Slice trägt?
Tag 1 · 15:00	Erste Messreihe: Baseline auf Lab-Hardware Ihre Ziel-Modelle laufen erstmals auf den reservierten GPU-Systemen. Baseline-Messung von Durchsatz und Latenz als Referenz für alle Entscheidungen der Woche.
Tag 1 · 16:15	Tagesergebnis: geprüftes Zielbild und Messplan Zielbild, harte Grenzen, gewählter Slice und ein Messplan für die drei kritischen Annahmen liegen schriftlich vor und sind vom ganzen Team getragen.
Tag 2 · 09:00	Architekturentwurf: die tragenden Entscheidungen Compute, Netz, Storage, Identity, Registry, Inference-Stack: der Entwurf entsteht an der Wand, jede Schicht mit Begründung und Alternative. Azure Local und Azure Arc werden dort eingeplant, wo sie tragen, in der jeweils verfügbaren Version.
Tag 2 · 10:45	Sizing-Rechnung gegen Baseline Das Kapazitätsmodell wird mit den Baseline-Werten von Tag 1 durchgerechnet: Wie viele Karten welcher Klasse, von L40S bis H100, für Ihre Lastspitzen? Der Unterschied zwischen Herstellerangabe und eigener Messung wird beziffert.
Tag 2 · 13:00	Validierung: Annahme eins und zwei im Test Die beiden ersten kritischen Annahmen werden auf der Hardware geprüft, etwa Quantisierungseffekte auf Ihre Qualitätsanforderung oder Concurrency-Verhalten unter realistischer Last. Ergebnisse fließen sofort in den Entwurf zurück.

DIE 20 TRACKS · TRACK 20/20 · FORTSETZUNG

Tag 2 · 15:30	Betriebsmodell-Weiche Verbunden, zeitweise getrennt oder Disconnected-Prüfung anstoßen: Entscheidung am Kriterienkatalog, mit ehrlicher Einordnung der Eligibility-Voraussetzungen bei Azure Local Disconnected Operations.
Tag 2 · 16:15	Tagesergebnis: belastbarer Architekturentwurf v1 Der Entwurf steht in Version 1, mit Sizing-Rechnung auf Basis eigener Messwerte und einer dokumentierten Betriebsmodell-Entscheidung samt offenen Prüfpunkten.
Tag 3 · 09:00	Slice-Bau: das Fundament Aufbau des Slice auf der Lab-Umgebung entlang des Entwurfs: Identity-Anbindung, Netzsegment, Registry-Anbindung mit Signaturprüfung. Ihr Team baut, der Trainer hält den Entwurf und die Realität zusammen.
Tag 3 · 13:00	Slice-Bau: Daten und Anwendungslogik Ihre anonymisierten Beispieldaten treffen auf den Stack: Aufbereitung, Indexierung oder Anbindung je nach Fall, dazu die Anwendungsschicht. Wo inhaltlich passend, kommen Muster aus Sovereign Assistant oder VOXA als Referenz zum Einsatz.
Tag 3 · 15:30	Erster Durchstich Der Slice antwortet zum ersten Mal durchgängig: von der authentifizierten Anfrage bis zur belegten Antwort. Was klemmt, wird protokolliert und priorisiert statt schöngeredet.
Tag 3 · 16:15	Tagesergebnis: durchgängiger Slice Der Anwendungsfall läuft Ende-zu-Ende auf der Lab-Umgebung. Die Abweichungen vom Entwurf sind protokolliert und terminiert für Tag 4.
Tag 4 · 09:00	Annahme drei und der Lasttest Die verbliebene kritische Annahme wird geprüft, danach Lasttest des Slice gegen die Messlatte von Tag 1: Durchsatz, Latenz, Verhalten an der Kapazitätsgrenze. Jede Abweichung bekommt eine Ursache oder ein Backlog-Item.
Tag 4 · 10:45	Betriebs- und Sicherheitsprobe Der Slice wird behandelt wie Produktion: Monitoring-Sicht aufgebaut, ein Modell-Rollback geprobt, ein Prompt-Injection-Versuch gegen die Schutzschicht gefahren. Ergebnisse fließen in den Betriebsentwurf.
Tag 4 · 13:30	Backlog-Priorisierung Alles Offene wird zu einem Backlog mit Aufwandsklassen, Abhängigkeiten und Risiken: der Weg vom validierten Slice zur produktiven Plattform, priorisiert vom Team, nicht vom Trainer.
Tag 4 · 15:30	Investitionsbild und Make-or-Buy Sizing, Betriebsmodell und Backlog werden zu einem Investitionsbild verdichtet: Eigenbau, Sovereign AI Appliance oder Mischweg, mit Zahlen und benannten Unsicherheiten für das Readout.
Tag 4 · 16:15	Tagesergebnis: validierter Entwurf mit Backlog Architektur- und Betriebsentwurf liegen in geprüfter Fassung vor, alle drei kritischen Annahmen sind mit Messwerten beantwortet, das Backlog ist priorisiert.
Tag 5 · 09:00	Härtung des Entwurfs Letzter Review-Durchgang mit vertauschten Rollen: das Team greift den eigenen Entwurf an, der Trainer verteidigt ihn und umgekehrt. Was den Angriff nicht übersteht, wird geändert oder als Risiko dokumentiert.

DIE 20 TRACKS · TRACK 20/20 · FORTSETZUNG

Tag 5 · 10:30	Readout-Probe Das Executive Readout wird einmal komplett geprobt: Kernaussagen, Messwerte, Investitionsbild, Empfehlung. Das Team entscheidet, wer welchen Teil vor der eigenen Entscheiderebene vertritt.
Tag 5 · 13:00	Dokumentenübergabe Entwurf, Messprotokolle, Backlog und Readout-Unterlage werden final zusammengeführt und übergeben. Alle Artefakte gehören Ihnen und sind ohne BORDONARO weiterverwendbar.
Tag 5 · 14:30	Executive Readout 90 Minuten vor Ihrer Entscheiderebene, vor Ort im Lab oder remote zugeschaltet: der validierte Entwurf, die Live-Demonstration des Slice, das Investitionsbild und die Empfehlung mit offenen Risiken. Fragen werden vom Team beantwortet, nicht nur vom Trainer.
Tag 5 · 16:15	Abschlussergebnis: Entscheidungsgrundlage übergeben Ihre Entscheiderebene hat einen validierten Entwurf mit Messwerten, getestetem Slice und priorisiertem Backlog gesehen. Die nächsten Schritte und die beiden Follow-up-Termine sind terminiert.

Das nehmen Sie mit

- Ein Architektur- und Betriebsentwurf für Ihre Plattform, in Ihrer Struktur und mit Ihren Anforderungen, nicht als generische Vorlage
- Validierte Kernannahmen: Sizing, Durchsatz und Antwortzeiten Ihrer Ziel-Workloads, gemessen auf GPU-Systemen unseres Labs von L40S bis H100
- Ein getesteter Slice: ein priorisierter Anwendungsfall Ihres Hauses läuft Ende der Woche durchgängig auf der Lab-Umgebung
- Ein priorisiertes Backlog für die Umsetzung mit Aufwandsklassen, Abhängigkeiten und benannten Risiken
- Ein Executive Readout am Tag 5: Entwurf, Messwerte und Investitionsbild in 90 Minuten für Ihre Entscheiderebene, mit Dokument zum Mitnehmen

Voraussetzungen

Ein entscheidungsnahes Vorhaben und ein Team, das es trägt: mindestens ein Architekt oder Plattform-Verantwortlicher mit Infrastruktur-Tiefe. Vorbesuch von AI-PLATFORM ist hilfreich, aber nicht Bedingung.

Abschluss

Der Transfernachweis ist das Executive Readout selbst: Ihr Team präsentiert Entwurf, Messwerte und Empfehlung vor der eigenen Entscheiderebene und beantwortet deren Fragen. Alle Validierungen sind als Messprotokolle dokumentiert und Teil der Übergabe.

Enthalten

5 Tage exklusive Arbeit im Enterprise-Lab mit Dino Bordonaro, kein Parallelbetrieb mit anderen Kunden · Reservierte GPU-Systeme mehrerer Klassen für Vergleichsmessungen, aus dem Verbund mit 2.516 Cores, 24 TB RAM und 3.400 TB Storage · 90-minütiges Vorgespräch und Aufbereitung Ihrer Unterlagen vor der Woche · Architektur- und Betriebsentwurf, Messprotokolle und Backlog als übergabefähige Dokumente · Executive-Readout-Unterlage für Ihre Entscheiderebene · Zwei 60-minütige Follow-up-Termine in den 8 Wochen nach der Residency

PREISMATRIX

ALLE PREISE AUF EINEN BLICK

ARTIKEL-NR.	LEISTUNG	UMFANG	PREIS NETTO
AIA-AI-LIT-LAB	AI Literacy Praxislabor	½ Tag	3.950 € pauschal je Durchführung
AIA-AI-WORK	Sicher produktiv mit KI arbeiten	1 Tag	7.900 € pauschal je Durchführung
AIA-AI-CONTEXT	Prompt- und Context-Engineering	2 Tage	15.800 € pauschal je Durchführung
AIA-AI-START	KI-Einstieg für den Berufsstart	1 Trainingstag	7.900 € pauschal je Durchführung
AIA-AI-CHECK	KI-Ergebnisse prüfen und verbessern	1 Trainingstag	7.900 € pauschal je Durchführung
AIA-AI-EXEC	Sovereign AI Executive Briefing	½ Tag	4.900 € pauschal je Durchführung
AIA-AI-LEAD	Führen in der Human-Agent-Organisation	2 Tage	15.800 € pauschal je Durchführung
AIA-AI-VALUE	Use-Case- und Business-Case-Sprint	2 Tage	15.800 € pauschal je Durchführung
AIA-AI-CHAMP	AI Champions Program	4 Präsenztage + 6 Fragerunden über 12 Wochen	34.900 € pauschal je Durchführung
AIA-AI-GOV	EU AI Act und AI Governance	2 Tage	15.800 € pauschal je Durchführung
AIA-AI-RISK	AI Risk Classification Lab	1 Tag	7.900 € pauschal je Durchführung
AIA-AI-SEC	Secure GenAI und Agent Threat Modeling	2 Tage	15.800 € pauschal je Durchführung
AIA-AI-DATA	AI-ready Data und Knowledge Architecture	2 Tage	15.800 € pauschal je Durchführung
bordonaro.ai AIA-AI-RAG	Enterprise RAG Engineering	3 Tage	Preise 23.700 € pauschal je Durchführung

KONDITIONEN & KONTAKT

KLARE REGELN. KURZE WEGE.

Preisanker Trainings	7.900 € netto pro Trainingstag und Gruppe, Halbtagsformate anteilig, die Lab-Residency mit 8.900 €.
Mindestgruppe	Durchführung ab vier tatsächlich angemeldeten Teilnehmenden, bis zur genannten Gruppengröße. Kleinere Runden wechseln kostenfrei auf den nächsten Termin.
Terminplanung	Je nach Umfang 2 bis 6 Wochen Vorlauf ab schriftlicher Bestätigung von Termin und Umfang.
Abschluss	Teilnahmezertifikat und prüffähige Schulungsdokumentation bei jedem Track inklusive, als Baustein Ihres Kompetenznachweises nach Art. 4 EU AI Act.
Alle Preise	Netto zzgl. 19 % USt. Reisekosten nur bei Vor-Ort-Terminen außerhalb der Region, vorab ausgewiesen.

BORDONARO IT GmbH & Co. KG

Wormser Landstraße 83a · 67346 Speyer
+49 6232 298 53 - 0 · info@bordonaro-it.de

[bordonaro.ai](#) · [www.bordonaro.de](#) · [dino.bordonaro.de](#)

AG Ludwigshafen HRA 61166 · Komplementärin: BORDONARO Beteiligungs GmbH, HRB 64031 · USt-IdNr. DE815511519 · Vertreten durch Dino Bordonaro



BORDONARO

NICHT ALLES MUSS OFFLINE. ALLES MUSS UNTER KONTROLLE.

bordonaro.ai · www.bordonaro.de · dino.bordonaro.de

+49 6232 298 53 - 0 · info@bordonaro-it.de · Stand: September 2026