



TRAINING CATALOG 2026 • FULL EDITION

AI BRINGS SPEED. YOUR SKILLS SET THE DIRECTION.

20 role-based AI tracks from career start to platform operations, plus programs and masterclasses. With agendas, fixed group prices and artifacts that keep working on Monday.

€7,900 per training day and group • 4 to 12 participants • certificate included

As of: September 2026 • bordonaro.ai

WHAT YOU WILL FIND IN THIS CATALOG.

Learning mechanic and trainer	4
Scenarios and decision guide	5
The 20 tracks	7
Price overview	54
Terms and contact	55

Every item carries an item number and a QR code leading directly to the matching page with details, agenda and inquiry path.

WHY THIS CATALOG

NOT A PAPER ABOUT POSSIBILITIES. AN OFFER WITH PRICES.

I have spent almost 30 years in datacenters, projects and training rooms, and developed one simple conviction: if you sell something, you should be able to say what it costs, who it is for and what exists at the end. That is exactly how this catalog is built. Every item has a price or an honest path to one, a target group and a verifiable result.

Not everything needs to be offline. Everything needs to stay under control. If you take just one question away from these pages, let it be this: does every workload in your organisation run where it belongs?

Dino Bordonaro

Founder and CEO · 8x Microsoft MVP in a row



THE LEARNING MECHANIC

NOT MORE TOOLS. A WAY OF WORKING YOU CAN RELY ON.

01

Understand

What can AI do here, and what can it not?

02

Apply

Work on a suitable task with approved data.

03

Check

Examine sources, completeness, errors and required approvals.

04

Put it to work

Adopt a verified template or checklist into your daily routine.

05

Deepen

Work through recurring questions in a role track or workflow lab.

06

Keep learning

Solve real questions in regular Q&A sessions and updates.

YOUR TRAINER

Teaching AI also means understanding what it runs on.

Dino connects long-standing IT and teaching experience with the practical question of how AI should be used and managed at work. Technical tracks address architecture, data flows, permissions and operations. For your first step, the focus is a task you can work through and understand.

8x Microsoft MVP in a row • Microsoft Certified Trainer • more than 800 professionals trained

WHAT IT FEELS LIKE

A CONVINCING ANSWER CAN STILL BE WRONG.

ILLUSTRATIVE SCENARIO. A FICTIONAL WORKFLOW, NOT A CLIENT CASE STUDY.

Monday, 8:17 am. The request has been waiting since Friday.

An experienced sales specialist knows her products and customers. Preparing a proposal still starts with searching through emails, product documents and open questions. In this scenario, an approved AI assistant receives selected documents and a clear task. It gathers requirements, flags gaps and drafts an outline. She rejects a suggested delivery date because there is no confirmed source. She no longer starts from a blank page. The commitment to the customer remains her decision.

ILLUSTRATIVE SCENARIO. A FICTIONAL WORKFLOW, NOT A CLIENT CASE STUDY.

The first support case. The answer sounds remarkably confident.

A new team member asks AI to explain an error. It suggests a command they have never used. Run it straight away? This is where real competence begins. In the exercise, they check approved documentation, explain the possible effect in their own words and test only in an isolated environment. A change to the customer system still requires the responsible person's approval. AI helps them learn. It does not grant them permission.

ILLUSTRATIVE SCENARIO. A FICTIONAL WORKFLOW, NOT A CLIENT CASE STUDY.

The presentation is ready. The key figure is wrong.

An experienced controller uses AI to help draft a monthly report. The text sounds convincing, but it compares two different reporting periods. In this scenario, calculations stay in the verified reporting system. AI helps with structure, alternative explanations and follow-up questions. The controller checks the figures and assumptions before the report goes to management. Progress is not speed at any cost. It is less routine work with responsibility still clearly assigned.

That is why we do not begin with a list of tools. We begin with your task and how you recognise a good result.

WHO IS THIS FOR? YOUR DECISION GUIDE

SIX ROLES, SIX LEARNING PATHS.

Nobody needs all 20 tracks. Find your role, start with the recommended path and have prior knowledge recognised in a skills conversation. Safe use and result checking belong in every path.

Career starters and career changers

„I already use AI, but I never know whether I can trust it.“

You get a safe learning path: have it explained, ask back, check sources, explain it yourself. The secret shortcut becomes a learning accelerator.

AI-START



AI-LIT-LAB



AI-WORK

Experienced professionals

„Too much of my day goes into searching, summarising and drafting.“

You bring your professional context into working with AI and test its suggestions against your standards. Your knowledge sets the bar, not the smoothest answer.

AI-WORK



AI-CONTEXT



AI-CHECK

Executive management and board

„What is AI worth to us, what do we risk and who answers for it?“

From decision picture to resilient business cases to the leadership questions: you decide based on your own judgement instead of borrowed slides.

AI-EXEC



AI-VALUE



AI-LEAD

Governance, privacy, security

„How do we make AI use provable and properly classified?“

The EU AI Act translated per role, applications classified cleanly, threats modelled systematically. Training and drafts, not legal advice.

AI-GOV



AI-RISK



AI-SEC

Architecture and engineering

„How do I build RAG, agents and evaluation that survive production?“

From the knowledge foundation to RAG, evaluation and bounded, supervised agents: built and hardened on real lab hardware.

AI-DATA



AI-RAG



AI-EVAL



AI-AGENT

Platform and operations

„How do I run models and GPUs, connected to disconnected?“

Plan, operate and update sovereign platforms, up to fully disconnected operating models. Your home turf becomes our classroom.

AI-PLATFORM



AI-OPS



AI-ARC



AI-RESIDENCY

The online check helps in 7 questions: bordonaro.ai/en/academy

AIA-AI-LIT-LAB

PATH 1 · UNDERSTAND

k/aia-ai-lit-lab

AI Literacy Hands-on Lab

Half a day after which every participant has built their own assistant instead of just hearing about one.

You set up your own assistant in the lab, practise the data rules on real cases and finish with a transfer check, in groups of 4 to 12 people.

Who for: Teams and departments that need to move from listening to doing: case handling, sales, administration, service.

Duration: ½ day **Group:** 4-12 people **Lead time:** from 2 weeks

€3,950 flat per delivery

AGENDA

09:00

From knowing to doing

Every participant writes the first prompt for the task from their profile. The results are compared: Why does the same question produce twenty different answers? This starts the first improvement round.

09:50

Data rules under fire

Case card exercise: customer complaint, salary list, press release, quote calculation. Each participant decides per case whether and into which AI the data may go, and justifies the decision. Contested cases are argued out in plenary.

10:50

Prompts that work repeatably

The morning's attempts are distilled into a pattern: role, task, context, format. Each participant reworks their own prompt along the pattern and measures against the output what actually improved.

11:30

Your first own assistant

Each participant sets up an assistant with their own role description, tone and clear limits, on our lab environment or your tenant. By the end, the assistant answers the profile task better than the first prompt of the morning.

12:30

Transfer check and Monday plan

Each participant demonstrates their assistant on a new case and decides three data rule cases in the check. Closing question: What do I rebuild on Monday, and what do I need from whom?

What you take away

- Every participant has written, improved and evaluated their own prompts on a real task from their own daily work
- The data rules are practised on concrete cases: every participant has decided and justified at least ten classifications themselves
- A first own assistant with role, tone and limits, set up on our lab environment or your tenant
- A passed transfer check per person as a documented building block of your Art. 4 competence evidence

Prerequisites

Basic terms such as prompt and hallucination should have been heard before, for example in the AI Literacy Praxislabor (AI-LIT-LAB). A laptop with a browser is enough, prior AI experience is not required.

Completion

Transfer check at the end: each participant demonstrates their own assistant on a new case and decides three data rule cases. Results and attendance are documented in a traceable, audit-ready form.

Included

Half-day hands-on lab with Dino Bordonaro, on site or remote · Prepared practice accounts on our lab environment, or setup on your tenant by prior arrangement · Exercise set with data rule case cards and prompt starter templates · Step-by-step guide for rebuilding the assistant in your own organization · Attendance certificates and audit-ready training documentation including transfer check results

AIA-AI-WORK

PATH 1 • UNDERSTAND



k/aia-ai-work

Working Safely and Productively with AI

One day on which three of your real workflows become measurably faster, not three slide examples.

Three genuinely improved workflows with before-and-after measurements, personal prompt and review templates, and team rules the team has agreed on.

Who for: Teams of up to 12 people who have already touched AI and now want to get their actual work done with it: reports, quotes, minutes, customer correspondence, research.

Duration: 1 day **Group:** 4-12 people **Lead time:** from 2 weeks

€7,900 flat per delivery

AGENDA

09:00

Your workflows on the table

The three selected workflows are dissected on the whiteboard: steps, time spent, failure points. For each workflow the team sets the before benchmark: How long does it take today, and what counts as a good result?

10:00

Rebuilding workflow 1

The first workflow is rebuilt live: develop a shared prompt template, run it on three real cases, hold the results against the benchmark. First question of the day: Where does AI really save time, and where does checking eat the gain back up?

11:30

The review routine

The review template is built on the outputs of workflow 1: fact check, source question, number verification, tone. Exercise with prepared outputs: the team hunts down planted hallucinations and decides which review steps become mandatory.

13:30

Workflows 2 and 3 in workshop mode

In two groups, participants rebuild workflows 2 and 3 themselves: prompt template, test run on real cases, review steps. The trainer rotates, corrects and asks the uncomfortable questions: Would you sign off on this result?

15:00

Before-and-after measurement

Both groups present on a live example. For each of the three workflows the record shows: time before, time after including review, remaining risks. Honest outcome: at least one candidate often fails, and that gets documented too.

16:00

Team rules and Monday plan

The team decides and writes down its rules: which tasks with AI, which data allowed, who reviews what, what stays off limits. Each participant names the one workflow they will run with template and review routine starting Monday.

What you take away

- Three real workflows of the team are rebuilt with AI and documented with before-and-after measurements
- Every participant has personal prompt templates for their own recurring tasks
- A review template per workflow: what gets checked before reuse, by whom, and how to spot hallucinations
- Agreed team rules: which tasks with AI, which data, who reviews, what stays off limits
- A documented training day as a building block of your Art. 4 competence evidence

Prerequisites

Every participant should have written their own prompts before, for example in AI-LIT-LAB or through their own practice. Access to an AI tool approved in your organization on their own laptop.

Completion

Transfer evidence via the before-and-after measurements of the three workflows and the templates each participant created themselves. Team rules and measurement results are recorded in the results document in a traceable, audit-ready form.

Included

1 training day with Dino Bordonaro, on site or remote · Evaluation of the submitted workflow profiles and selection of the three training workflows in advance · Template set: prompt template, review template and team rules template for continued use · Results document with all templates and measurements created during the day · Attendance certificates and audit-ready training documentation · 30-minute follow-up call after four weeks: What has landed in daily work, where is it stuck?

AIA-AI-CONTEXT

PATH 1 · UNDERSTAND

k/aia-ai-context

Prompt and Context Engineering

Two days after which good AI results in your team are no longer a matter of luck, but repeatable building blocks.

Reusable context patterns, structured outputs with quality controls, and a tested prompt library for your team.

Who for: Heavy users, subject-matter experts and developers who work with AI daily and whose results are reused by others: analyses, reports, data extraction, text modules, precursors to automation.

Duration: 2 days **Group:** 4-12 people **Lead time:** from 4 weeks

€15,800 flat per delivery

AGENDA

Day 1 · 09:00

Why the same prompt answers differently ten times

The submitted prompts under stress test: the same prompt, multiple runs, measured variance. This makes visible which formulations carry weight and which are left to chance. Each participant picks their most important prompt as their working object for both days.

Day 1 · 10:30

Context patterns instead of gut feeling

Patterns are distilled from the best submissions: role, domain context, examples, limits, negative instructions. Exercise: each participant dissects their own prompt into building blocks and names which block has which effect, verified by leaving blocks out.

Day 1 · 13:30

Context from company knowledge

How much document, table or background knowledge belongs in the context, and in what form? Exercise on real material: feed the same document as raw text, as an excerpt and as structured bullet points into the context and compare result quality. Plus the data rule question: What is allowed in at all?

Day 1 · 15:30

Day result: your pattern set passes the test

Each participant rebuilds their working prompt fully on context patterns and runs it against three predefined test cases. It passes if all three cases produce usable results that are demonstrably better than the original version. The patterns go into the shared catalogue.

Day 2 · 09:00

Structured outputs: JSON and tables

From prose to reusable results: define output schemas, enforce mandatory fields, format tables cleanly for Excel. Exercise: a data extraction from unstructured documents whose result drops into a table without rework. Developers additionally validate the JSON by machine.

Day 2 · 11:00

Quality and hallucination controls

Every library prompt gets test cases and edge cases: empty inputs, contradictory information, missing data. Exercise on prepared documents with planted errors: Which controls catch the invented number, which let it through? This produces the review routine per prompt.

THE 20 TRACKS · TRACK 3/20 · CONTINUED

Day 2 · 13:30

The team's prompt library

The individual results become a shared repertoire: naming convention, description, usage limits, test cases, version. Workshop: the team's ten most important prompts are made library-ready and cross-tested. Team decision: Where does the library live, who maintains it, how do new prompts get in?

Day 2 · 15:30

Day result: library version 1 passes the acceptance test

The real test: each participant solves someone else's task using only the other participants' library prompts. Whatever works without questions is accepted, the rest gets concrete rework with an owner and a deadline. Closing: maintenance plan and the first three expansion candidates.

What you take away

- A set of reusable context patterns: roles, domain context, examples and limits as building blocks that can be combined per task
- Prompts with structured outputs: JSON and tables that can be taken straight into Excel, further processing or downstream systems
- Quality and hallucination controls per prompt: test cases, edge cases and defined checks before reuse
- A team prompt library version 1: named, documented, tested, with maintenance responsibility clarified
- Two documented training days as a building block of your Art. 4 competence evidence

Prerequisites

Regular hands-on AI use in daily work is assumed, roughly at the level reached after AI-WORK. The JSON exercises require no programming skills, developers get additional depth from the same exercises.

Completion

Two verifiable day results: each participant's pattern set passes three defined test cases, and the team library passes the mutual acceptance test. Both results are documented in a traceable, audit-ready form.

Included

2 training days with Dino Bordonaro, in-house or remote · Advance evaluation of the submitted prompts as the team's starting picture · Pattern catalogue of context patterns and output schemas as editable templates · Prompt library template including naming convention, versioning and test case structure · Attendance certificates and audit-ready training documentation · 45-minute remote office hour four weeks later on evolving the library

AIA-AI-START

PATH 1 • UNDERSTAND



k/aia-ai-start

AI Onboarding for Career Starters

One day teaching career starters to use AI as a learning tool without trusting it blindly: have tasks explained, ask good follow-up questions, check sources and spot your own knowledge gap before it becomes a mistake.

You practise having tasks explained by AI, questioning the explanation and, at the end of the day, explaining a synthetic practice case in your own words, without the AI at your side.

Who for: Apprentices, career starters, career changers and juniors in their first years on the job.

Duration: 1 training day**Group:** 4-12 people**Lead time:** from 2 weeks**€7,900** flat per delivery

AGENDA

09:00	What AI can do and where answers come from How it works, in understandable images: language models produce plausible continuations, not verified truths. First exercise: ask the same question three times and explain the differences.
10:00	Having tasks explained, properly A technical term, a process, an error message: how to request an explanation that matches your level, and which follow-up questions make the explanation testable.
11:15	Checking sources instead of believing Exercise on prepared outputs: which statement has a verifiable source, which one only sounds like it? Every participant finds at least one built-in plausible false statement.
13:00	Data rules at your own workplace What may go into which AI: case cards with typical situations from training and early career. Customer data, internal documents, your own notes. House rules become concrete instead of abstract.
14:00	Your own practice case Everyone works on a term or task they brought along using the routine: have it explained, ask back, check the source, document understanding. Peer exchange in pairs.
15:30	Explaining it yourself, without AI The litmus test: everyone explains their case in their own words to the small group. Whatever still wobbles becomes visible and gets a concrete learning task.
16:30	Escalation and next step When an experienced person must step in: writing down your personal boundary. Everyone leaves with a learning-routine card and a concrete next practice step.

What you take away

- A personal learning routine: have it explained, ask back, check the source, explain it yourself
- A practised feel for uncertainty: recognising when an answer only sounds plausible
- A clear rule for when to involve an experienced colleague
- A self-completed synthetic practice case with a documented check
- Your organisation's data rules applied to your own workplace

Prerequisites

None. A laptop with a browser is enough. The track deliberately assumes no prior AI experience.

Completion

You pass if you can explain your practice case in your own words at the end and have demonstrably applied the check routine to it. The result is recorded in the training records.

Included

1 training day with Dino Bordonaro, in-house or remote · Practice cases on synthetic data, no real customer data required · Learning-routine card and check-question template for continued use · Certificate of participation and training records · One 60-minute remote Q&A session four weeks later for open cases

AIA-AI-CHECK

PATH 1 • UNDERSTAND



k/aia-ai-check

Checking and Improving AI Results

One day dedicated to the skill that decides between benefit and damage: systematically checking AI results. Sources and figures, plausible errors, counter-checks and the question of when a result needs sign-off.

You practise a robust checklist on several synthetic cases and adapt it to your team's work. By the end, everyone recognises the typical plausible errors and knows when to escalate.

Who for: Teams that already use AI drafts and now need the quality assurance to go with them: case handling, assistants, business departments, controlling-adjacent roles.

Duration: 1 training day**Group:** 4-12 people**Lead time:** from 2 weeks**€7,900** flat per delivery

AGENDA

09:00

Why plausible is not correct

The error classes of language models shown on real patterns: invented sources, wrong figure references, silent assumptions, a convincing tone. Exercise: three prepared outputs each contain one error, who finds it unaided?

10:15

The checklist, step by step

Sources, figures, completeness, faithfulness to the task, tone, data provenance. Every check step gets a concrete hand movement and a time budget. Checking must be faster than writing it yourself, or nobody will do it.

11:30

Counter-checks that work

Recalculating figures, comparing claims against sources, counter-questioning the model when unsure. Exercise on synthetic cases: which counter-check exposes which error type?

13:00

Your cases, round one

The anonymised examples you brought are worked through with the checklist. What would the list have caught, what not? The list is sharpened exactly where it fails on your own material.

14:15

Escalation and sign-off

Not every result needs four eyes, some need six. Decision tree: what do I sign off myself, what does a colleague check, what goes to the manager? The rule becomes written and team-ready.

15:30

Check protocol for the record

For cases requiring evidence: the lean protocol documenting what was checked. Exercise on your own case, after which the template is ready to use.

16:30

Team rule and next step

The adapted checklist becomes a team agreement: which check steps are mandatory from tomorrow, who maintains the list, when is it sharpened. Everyone leaves with a concrete application plan.

What you take away

- A checklist for AI results, practised on several case classes and adapted to your team
- A catalogue of plausible errors: invented sources, shifted time periods, silent assumptions
- Practised counter-checks: recalculation, source comparison, counter-questions to the model
- A documented escalation and sign-off rule for critical results
- A check-protocol template for cases requiring evidence

Prerequisites

First experience with AI drafts in daily work, for example from AI-WORK, AI-LIT-LAB or your own practice. A laptop with a browser is enough.

Completion

You pass if you find the built-in errors in the final case using the checklist and justify the escalation decision. The result is recorded in the training records.

Included

1 training day with Dino Bordonaro, in-house or remote · Synthetic check cases with built-in, documented errors · Checklist template and check-protocol template for continued use · Certificate of participation and training records · One 60-minute remote Q&A session four weeks later for contested cases

AIA-AI-EXEC

PATH 2 · LEAD

k/aia-ai-exec

Sovereign AI Executive Briefing

One afternoon after which your executive team knows whether AI in your company will be rented or owned.

A documented decision picture covering value, risk, sovereignty and operating model, plus the open decision questions for the next 90 days as a structured template for your governing bodies.

Who for: Managing directors, board members, advisory boards and shareholders, up to eight people.

Duration: ½ day **Group:** 4-8 people **Lead time:** from 2 weeks

€4,900 flat per delivery

AGENDA

14:00	Situation picture: pressure from three sides Analysis of your questionnaire: customer expectations, competition, shadow AI inside your own company. Guiding question for the room: Which three AI decisions are genuinely pending in your company over the next 24 months?
14:45	The live proof Demonstration of an internet-disconnected AI environment from our enterprise lab, where technically and organizationally feasible: the Sovereign Assistant answers questions about loaded documents, visibly without any cloud connection. Executive question afterwards: What does this mean for our critical data assets?
15:45	Economics: rent or own The calculation model: token and license cost per head and year against investment and operation of an own environment, with conservative assumptions. Exercise: your usage profile is entered live and the break-even for your company becomes visible.
16:30	Strategy session: your decision picture Value, risk, sovereignty and operating model on a single page. Decision in the room: Which two questions does the executive team want answered reliably within 90 days, and who owns them?
17:10	The next 90 days Assignments are fixed with owners and dates. Alignment with the fitting next steps, such as the use case sprint (AI-VALUE) or the leadership track (AI-LEAD). Ends at 17:30.

What you take away

- A one-page decision picture: value, risk, sovereignty and operating model for your company, developed together in the room
- A rent-or-own calculation using your usage profile and conservative assumptions, with a visible break-even
- A prioritised 90-day draft with two to three assignment candidates, owners and dates are proposed by the group itself
- A shared vocabulary at executive level, so the next AI discussion ends in a decision instead of another postponement

Prerequisites

None. No technical background is required, bring decision-making authority and 3.5 undisturbed hours.

Completion

No quiz at executive level. The transfer evidence is the documented decision picture with the 90-day draft, sent to all participants as minutes and discussed in the follow-up call after four weeks.

Included

Half-day session delivered by Dino Bordonaro, on site, remote or in our enterprise lab · Analysis of the pre-briefing questionnaire as an individual situation picture of your company · Live demonstration of an internet-disconnected AI environment from the enterprise lab, where technically and organizationally feasible · Rent-or-own calculation model as a reusable template · Documented decision picture and 90-day draft delivered as minutes within three working days · 30-minute follow-up call with the sponsor after four weeks

AIA-AI-LEAD

PATH 2 · LEAD

k/aia-ai-lead

Leading in the Human-Agent Organization

Two days after which your managers know what to delegate to an assistant and what never to.

Every manager works out a delegation matrix, an escalation model and the draft of a 90-day adoption plan for their own team.

Who for: Managers with people responsibility, from team lead to department head, up to twelve people.

Duration: 2 days **Group:** 4-12 people **Lead time:** from 4 weeks

€15,800 flat per delivery

AGENDA

Day 1 · 09:00	Your leadership week under the X-ray Work on the week lists brought along: Which of the ten tasks are drafting work, which are genuine decisions? First sorting on the wall, and the patterns across all participants become visible.
Day 1 · 10:30	What assistants and agents genuinely deliver today Live work on the system: drafting a target agreement, preparing a difficult feedback conversation, summarizing a meeting. Including one visible error and the question: How would I have caught it if I had only skimmed?
Day 1 · 13:00	Delegation rules Building the task-class matrix: goes to the assistant, goes to the agent with approval, stays with the human. Exercise on the own week lists, edge cases are decided by the group. Guiding principle: the assistant drafts, the human approves.
Day 1 · 15:00	Decision and escalation model Defining approval levels: When is a spot check enough, when does it take four eyes, when is AI off limits? Three critical cases are played through, including a flawed draft that almost went out.
Day 1 · 16:15	Day result: your delegation matrix stands Each manager documents matrix and approval levels for their own team and presents them to a partner in a peer review. Open points are noted for day 2.
Day 2 · 09:00	Taking fears seriously The real objections on the table: Will this replace me? Is my work being measured now? Exercise: formulating answers that are honest instead of soothing, and naming what a manager can promise and what not.
Day 2 · 10:30	Works council and co-determination Which rollout steps can be subject to co-determination, what early involvement looks like in practice and which commitments hold. No legal advice, but a conversation structure that holds up in the meeting with the council. Exercise: opening the works council conversation in a role play.
Day 2 · 13:00	The adoption plan The first 90 days in the own team: the manager's visible own usage, three team routines, measuring success without a feeling of surveillance. Each manager decides: Which task does my team start with in week 1?

THE 20 TRACKS · TRACK 7/20 · CONTINUED

Day 2 · 14:30

Leadership routines live

Simulation of a weekly with assistant drafts: review the draft, grant or refuse approval, formulate feedback to the team. The approval rules from day 1 are applied under time pressure.

Day 2 · 16:00

Day result: your plan in front of the group

Each manager presents delegation matrix and adoption plan in five minutes while the group pushes back with the hardest objections. The fixed version goes into the training documentation.

What you take away

- A delegation matrix per manager: which task classes go to assistants, which to agents with approval, which stay with humans
- A decision and escalation model with approval levels built on the principle: the assistant drafts, the human approves
- A 90-day adoption plan for the own team with the first three leadership routines
- A robust line of conversation for works council and team that takes fears seriously instead of talking them away
- A documented role profile of the manager in the human-agent organization

Prerequisites

People responsibility and some first-hand contact with an AI assistant. Anyone who has never worked with an assistant catches up in one hour with the prework.

Completion

The transfer evidence is the presentation of delegation matrix and adoption plan in front of the group at the end of day 2. Both artifacts and attendance are documented in a traceable, audit-ready form.

Included

Two on-site or remote days delivered by Dino Bordonaro · Templates for delegation matrix, escalation model and adoption plan for further use in your company · Conversation guide for team and works council discussions · Certificates and audit-ready training documentation as a building block of your Article 4 evidence · 60-minute remote group consultation four weeks after the training

AIA-AI-VALUE

PATH 2 • LEAD



k/aia-ai-value

Use Case and Business Case Sprint

Two days after which your AI idea list is a prioritized portfolio with owners and numbers.

A prioritized use case portfolio with a value-risk matrix, conservatively calculated economics, one owner per use case and a 90-day plan for the top candidates.

Who for: Mixed groups of managers, department representatives and executives, up to twelve people.

Duration: 2 days **Group:** 4-12 people **Lead time:** from 4 weeks

€15,800 flat per delivery

AGENDA

Day 1 • 09:00	The idea pile on the table All submitted one-pagers go up on the wall. Three test questions per one-pager: Which problem, whose time, which data? Duplicates are merged, anything unclear is sharpened by its submitter in 60 seconds.
Day 1 • 10:30	The evaluation grid Value dimensions (time saved, quality, risk reduction) and risk dimensions (data class, process criticality, acceptance) are agreed. Calibration exercise: the whole group scores two cases together until the grid ticks the same for everyone.
Day 1 • 13:00	The value-risk matrix takes shape All use cases are scored in small groups and placed in the matrix. The edge cases are fought out in plenary. Decision of the day: Which five to seven candidates form the top group?
Day 1 • 15:00	Feasibility check of the top group Per top candidate: What is the data situation, which systems are involved, which maturity level is required? Where it fits, we show on the Sovereign Assistant in our lab environment how close a comparable case already is today.
Day 1 • 16:15	Day result: the prioritized portfolio The matrix stands, the top group is named and justified, every edge case has a documented decision. A photo protocol goes out to everyone the same evening.
Day 2 • 09:00	Calculating conservatively The calculation logic: hours saved times frequency times fully loaded cost rate, minus an adoption discount and review effort. Exercise on a shared example: the same idea calculated once optimistically, once conservatively, and the difference is the discussion.
Day 2 • 10:30	Business cases in small groups Each group calculates one or two top candidates and documents every assumption. Rule: every number must survive a critical follow-up question from the executive team.
Day 2 • 13:00	Owners and obstacles Per use case the group fixes: one owner by name, one success measure, the biggest obstacle. Decision in the room: Who carries what, and which use case stays parked without an owner until one is found?

THE 20 TRACKS · TRACK 8/20 · CONTINUED

Day 2 · 14:30

The 90-day plan

For the two to three best candidates, the pilot is cut to size: scope small enough for 90 days, participants, measuring points, first milestone in week 2. Cross-check against your maturity level: what needs data or governance work first?

Day 2 · 16:00

Pitch and day result

Each group pitches its case in five minutes as if the executive team were in the room. Feedback against the grid, final corrections, the portfolio document is fixed and the review date in eight weeks is agreed.

What you take away

- A cleaned use case inventory from the pre-submitted one-pagers, duplicates merged, castles in the air sorted out
- A value-risk matrix with all candidates and a well-reasoned top group
- A rough economics calculation per top use case with conservative, documented assumptions
- A named owner, a success measure and the biggest obstacle per top use case
- A 90-day plan for the two to three best candidates with scope, participants and measuring points

Prerequisites

A basic understanding of what assistants and language models deliver, for example from AI-LIT-LAB or AI-EXEC. What matters is that the group knows its processes and is allowed to set priorities.

Completion

The transfer evidence is the closing pitch with a documented business case per small group. Portfolio, assumptions and owners are recorded in the portfolio document and checked against reality in the review after eight weeks.

Included

Two on-site or remote days delivered by Dino Bordonaro · One-pager template and evaluation grid as reusable tools for future portfolio rounds · Calculation model for the economics estimate with documented assumptions · Portfolio document with matrix, business cases and 90-day plan within five working days · Certificates and training documentation · 60-minute remote portfolio review after eight weeks

AIA-AI-CHAMP

PATH 2 · LEAD



k/aia-ai-champ

AI Champions Program

Four modules and six accompanying Q&A sessions over twelve weeks, after which AI in your company has a team instead of one lonely delegate.

Up to ten enabled champions with a use case pipeline, an internal community, ready-to-use guides and measurable transfer after 90 days.

Who for: Motivated employees and team leads from different departments who are meant to carry AI knowledge into the organization, up to ten people.

Duration: 4 on-site days + 6 Q&A sessions over 12 weeks

Group: 4-10 people

Lead time: from 6 weeks

€34,900 flat per delivery

AGENDA

Module 1 · 09:00	Baseline and role clarity Analysis of the profiles: Who brings what? The champion role is clarified on concrete cases: demonstrate, guide, collect, escalate. And just as clearly what a champion is not: no substitute admin, no supervisor, no one-person show.
Module 1 · 10:30	Mastering the tools Hands-on with the assistant: core prompting patterns, providing context, checking results. The company's data rules are decided on ten everyday cases: may this go into the AI or not?
Module 1 · 13:00	The first own case Each champion solves one time sink from their own profile live on the system and documents the solution path on the guide template. That documentation is the prototype of the later internal job aids.
Module 1 · 15:00	Practicing explanation Exercise in front of the group: each champion explains one concept, for example hallucination or context window, in three minutes so that the colleague at the front desk understands it. Feedback against a fixed grid.
Module 1 · 16:15	Module result: learning map and first job aid Each champion has one documented job aid and a personal learning map until module 2. Homework: test the job aid with two colleagues and note their reactions.
Module 2 · 09:00	From hallway talk to pipeline Review of the homework: What did the colleagues say? Then the interview technique: the five questions with which a champion turns a hallway conversation into a usable use case.
Module 2 · 10:30	One-pager workshop Each champion writes two use case one-pagers from their own area: which problem, whose time, which data. Mutual review in pairs, unclear one-pagers get sharpened.
Module 2 · 13:00	Scoring with the value-risk grid The compact version of the evaluation grid from the use case sprint: estimating value and risk per one-pager, edge cases decided by the group. Exercise: driving two deliberately contentious cases all the way to a decision.

THE 20 TRACKS • TRACK 9/20 • CONTINUED

Module 2 • 15:00	The pipeline board Setting up the pipeline in the tool of your choice: columns, owners, maintenance rhythm. Decision: Who curates, how do new entries come in, when does an entry get thrown out?
Module 2 • 16:15	Module result: the pipeline is alive At least ten scored use cases are on the board, and every column has an owner. Homework: each champion runs two interviews in their own area and feeds the pipeline.
Module 3 • 09:00	The internal learning session Pipeline review from the homework, then the craft: structure of a 20-minute session, dramaturgy, the three most common mistakes in internal training. Question to each champion: Which session does your area need first?
Module 3 • 10:30	Building sessions from the material package Each champion builds their first own session for their home area from the supplied templates, with a live example from that area's daily work.
Module 3 • 13:00	Dress rehearsal The sessions are delivered in front of the group, mercilessly against the clock. Feedback against a grid: opening, quality of examples, handling of questions. Every session leaves with one concrete improvement.
Module 3 • 15:00	Building the community Channel, rhythm, ground rules: how questions arrive, how fast they are answered, what a champion does when they do not know the answer. Escalation paths to IT, data protection and leadership are recorded with names.
Module 3 • 16:15	Module result: dates are set Each champion has a rehearsed session and a fixed date in their own area. The community concept is agreed and the channel is set up. Homework: deliver the session in front of a real audience.
Module 4 • 09:00	Cutting the pilots Experience reports from the first real sessions. Then: two to three pilots are cut from the pipeline, small enough for 90 days, each with a champion as caretaker and one success measure.
Module 4 • 10:30	Making it measurable Before-and-after measurement without a feeling of surveillance: simple indicators such as lead time or follow-up-question rate, collected cleanly. Exercise: defining the measurement per pilot and fixing the baseline.
Module 4 • 13:00	The transfer dashboard Each champion sets up the dashboard for their area: pilots, sessions, community activity, measurements. Rule: everything that should appear in the final report is tracked from today.
Module 4 • 14:30	The appearance before leadership Preparing the results presentation for the sponsor: What has been created, what will be measured in 90 days, what do the champions need from leadership? Trial run with hard feedback.
Module 4 • 16:00	Module result and starting signal Presentation to the sponsor, directly afterwards where organizationally possible. Each champion has their 90-day plan, and the six Q&A session dates are in the calendar.

THE 20 TRACKS · TRACK 9/20 · CONTINUED

Afterwards ·
Weeks 1-12

Six Q&A sessions and transfer

Six 60-minute remote sessions, every two weeks: real cases from the areas, pipeline review, sharpening the sessions. It ends with the closing meeting with sponsor and client, including the transfer report with the measurements per area.

What you take away

- Up to ten champions who use assistants confidently, guide colleagues and know when to escalate
- A filled and maintained use case pipeline with an evaluation grid and a maintenance routine
- An internal community with channel, rhythm and escalation paths that keeps running after the program
- A material package of guides and training materials with a license for internal reuse
- A transfer report after 90 days with a before-and-after measurement per champion area

Prerequisites

Curiosity, standing in the own department and release from duties for four on-site days plus about two hours per week during the 90 days after. No technical background is required.

Completion

Three verifiable proofs: the session delivered in front of a real audience, the maintained pipeline, and the transfer report with a before-and-after measurement after 90 days. Everything is documented in a traceable, audit-ready form.

Included

Four modular on-site days at intervals of two to three weeks, delivered by Dino Bordonaro, with one module in our enterprise lab on request · Six 60-minute remote Q&A sessions, every two weeks over twelve weeks, for cases from daily practice · Material package with guides, session templates and training slides including a license for internal reuse · Use case pipeline template and transfer dashboard · Certificates and audit-ready training documentation as a building block of your Article 4 evidence · Closing meeting with sponsor and client including the transfer report

AIA-AI-GOV

PATH 3 · GOVERN



k/aia-ai-gov

EU AI Act and AI Governance

Two days after which you hold an AI register, a literacy plan and a policy draft instead of a slide deck full of paragraphs.

A populated AI inventory, a role model, a literacy plan, a policy draft and a control calendar, all built on your real systems.

Who for: For owners in legal, data protection, information security, compliance and IT governance who need to turn the EU AI Act from text into a working state.

Duration: 2 days **Group:** 4-12 people **Lead time:** from 4 weeks

€15,800 flat per delivery

AGENDA

Day 1 · 09:00	What the EU AI Act actually requires The risk-based approach without panic: Art. 4 has applied since February 2025, oversight since August 2026. Core question for the room: where does the duty of best effort end and where does theatre begin? No prescribed certificate, no direct fine for Art. 4.
Day 1 · 10:30	Digital Omnibus status What the simplification package moves and what it does not. Decision per participant: which preparation is already solid today, which do we deliberately hold off on without breaching the duty of best effort?
Day 1 · 13:30	Role model on your organisation Provider, deployer, importer, distributor. Exercise: each participant assigns their three most important applications to a role and names who in the house carries the obligation. High risk follows the use case and role, not the industry as a blanket rule.
Day 1 · 15:30	Populating the AI inventory We start your register. Per application: purpose, data class, role, owner. Day result: a first populated inventory table with at least five real entries per participant.
Day 2 · 09:00	Literacy plan under Art. 4 Competence by role and context instead of mandatory training for all. Exercise: from your roles we derive who needs which level and enter it into the literacy grid.
Day 2 · 11:00	Policy workshop on your documents We write the policy draft using your real policies as the basis. Decision per house: what do we regulate anew, what points back to existing data protection and security rules?
Day 2 · 14:00	Approval process and control calendar Who approves a new AI application, in which steps, with which documentation? We build the approval path and enter recurring controls with dates and triggers into the calendar.
Day 2 · 16:00	Bringing it together and naming gaps We line up inventory, role model, literacy plan, policy and control calendar. Day result: an honest gap list with owners and deadlines for the next four weeks.

What you take away

- A populated AI inventory of your real applications with role, purpose and data class per entry
- A role model that maps provider, deployer, importer and distributor onto your organisation, with named owners
- A literacy plan under Art. 4 that assigns competence by role and context instead of training everyone across the board
- A policy draft for AI use, aligned with your existing data protection and information security rules
- A control calendar with dates, triggers and responsibilities for ongoing oversight

Prerequisites

Basic knowledge of your own organisation and access to existing policies. Legal training is not required, decision authority in your area helps.

Completion

The state of the five artefacts at the end of day two and the gap list are reviewed. Attendance and content are documented in a traceable, audit-ready form.

Included

Two on-site or remote days with Dino Bordonaro and up to twelve participants · Template pack: AI inventory table, role model matrix, literacy plan grid, policy skeleton and control calendar · Current status of the Digital Omnibus package with a clear read on what is solid today and what is still moving · Audit-proof training documentation and certificates of attendance as a building block for your Art. 4 evidence · A 30-minute follow-up call after four weeks on the state of your five artefacts

AIA-AI-RISK

PATH 3 • GOVERN

k/aia-ai-risk

AI Risk Classification Lab

One day on which your real applications get classified instead of generic example cases from a slide.

Your real use cases are sorted into prohibited, high risk, transparency obligation or minimal, with documented criteria and control need per case.

Who for: For governance owners and leaders who need to classify a list of deployed or planned AI applications with confidence.

Duration: 1 day **Group:** 4-12 people **Lead time:** from 2 weeks

€7,900 flat per delivery

AGENDA

09:00	The four classes cleanly separated Prohibited, high risk, transparency obligation, minimal. How to recognise each and where the grey zones sit. Core question: why does high risk follow use case and role rather than the industry as a blanket?
10:30	The decision tree in practice We work through the criteria catalogue step by step. Exercise on a shared example case until every participant reads the tree with confidence.
13:00	Your real cases, round one Each participant classifies two of their own use cases live. Decision: which class does the case fall into and which two criteria carry the decision?
14:30	Your real cases, round two and disputes We bring the hard cases before the group. Exercise: where opinions diverge, we document both readings and the condition that tips the balance.
15:45	Control need and prioritisation Per case we name which control is needed and how urgent it is. Day result: a documented, prioritised classification list of your submitted applications.

What you take away

- Every submitted use case is assigned to a class: prohibited, high risk, transparency obligation or minimal
- Documented decision criteria that show why a case falls into its class
- A named control need per case instead of one blanket statement for everything
- A prioritised list of which applications need measures first

Prerequisites

A list of your deployed or planned AI applications and basic knowledge of their purpose. Prior knowledge of the EU AI Act helps but is not required.

Completion

The documented classification list at the end of the day is reviewed: class, two carrying criteria and control need per case. Attendance and content are documented in a traceable, audit-ready form.

Included

One on-site or remote day with Dino Bordonaro and up to twelve participants · Classification decision tree and criteria catalogue as a working template · Documentation template per case with class, rationale and control need · Audit-proof training documentation and certificates of attendance · A 15-minute follow-up call to prioritise your list of measures

AIA-AI-SEC

PATH 3 · GOVERN



k/aia-ai-sec

Secure GenAI and Agent Threat Modeling

Two days after which you hold a threat model, a data flow diagram and a test plan for a real application instead of a list of theoretical risks.

A threat model, a data flow diagram, a catalogue of concrete protective measures and a test plan for a real GenAI or agent application of yours.

Who for: For security and governance owners together with engineering who need to secure a GenAI or agent application.

Duration: 2 days **Group:** 4-12 people **Lead time:** from 4 weeks

€15,800 flat per delivery

AGENDA

Day 1 · 09:00	Why GenAI carries different threats Prompt injection, unsafe outputs, data leakage through knowledge sources and tools. Core question: where does this differ from classic application security and where does it not?
Day 1 · 10:30	Data flow diagram of your application We draw the full flow: input, model, RAG sources, tool calls, output. Exercise: every trust boundary is marked and named.
Day 1 · 13:30	Carrying a STRIDE-style method over to GenAI We work through the threat categories and translate them onto prompt, context, tools and output. The OWASP LLM reference frame serves as a checklist so no category is forgotten.
Day 1 · 15:30	Collecting threats on your application Exercise on your data flow diagram: per trust boundary we name plausible attacks. Day result: a first threat model with prioritised threats.
Day 2 · 09:00	Live demonstration in the lab On a prepared environment we show how a prompt injection works through a RAG source and what an over-permissioned tool gives away. No customer data, only to understand the effect.
Day 2 · 11:00	Protective measures per threat Guardrails, least-privilege permissions, output filters, source checks. Decision per threat: which measure, which effort, which residual risk remains?
Day 2 · 14:00	Building the test plan For each measure a test that proves its effect. Exercise: for each important threat we formulate a concrete test case with an expected result.
Day 2 · 16:00	Residual risks and handover We gather what remains and name owners. Day result: threat model, measures catalogue, test plan and prioritised residual risk list for your application.

What you take away

- A data flow diagram of your application from input through model, knowledge sources and tools to output
- A threat model that carries a STRIDE-style method over to GenAI and agents
- A catalogue of concrete protective measures: guardrails, permissions and output filters, assigned per threat
- A test plan that makes it verifiable whether a measure works
- A prioritised residual risk list with owners

Prerequisites

Basic knowledge of application architecture and information security. You need a real application to examine, at least as an architecture sketch.

Completion

The completeness of threat model, measures catalogue and test plan at the end of day two is reviewed, each on your application. Attendance and content are documented in a traceable, audit-ready form.

Included

Two on-site or remote days with Dino Bordonaro and up to twelve participants · Threat modeling template for GenAI and agents with data flow notation · A reference frame along known OWASP LLM risk areas as a checklist, in our own wording · Measures and test plan template to continue in-house · Access to the lab environment for the demonstrations and a 30-minute follow-up call after four weeks

AIA-AI-DATA

PATH 4 · BUILD

k/aia-ai-data

AI-ready Data and Knowledge Architecture

RAG projects rarely fail at the model, they fail at sources without owners and permissions without a concept.

A prioritized knowledge map of your organization with an ownership matrix, source assessment and a retrieval and permission concept strong enough to carry a RAG project.

Who for: Architects, data engineers and developers tasked with making company knowledge usable for assistants.

Duration: 2 days **Group:** 4-12 people **Lead time:** from 4 weeks

€15,800 flat per delivery

AGENDA

Day 1 · 09:00	The situation report: your sources on the table Review of the anonymized source profiles. First sorting question: which of these sources would an assistant cite today, and whom would you trust with that answer?
Day 1 · 10:30	Why indexing without a map gets expensive Anatomy of failed RAG projects: stale duplicates, orphaned drives, ACL loss during crawling. Live demo on the Sovereign Assistant in our lab: the same question against curated and uncurated sources.
Day 1 · 13:00	Knowledge map workshop Each organization maps its sources along three axes: business value, data quality, access structure. The result is a four-quadrant map with clear candidates and clear exclusions.
Day 1 · 15:00	Ownership or data graveyard Exercise on your own maps: who owns each source, who maintains it, what happens when two versions contradict each other? The ownership matrix is built on your own case.
Day 1 · 16:30	Day result: map and matrix stand Each organization presents its knowledge map with ownership matrix in 5 minutes. The group's test question: is every top source assigned to a name, and is every exclusion justified?
Day 2 · 09:00	Prioritization: cutting the first iteration The map becomes an ordered list: which 3-5 sources carry the first RAG iteration? The criteria are answer value, data readiness and permission clarity, not political visibility.
Day 2 · 10:30	Retrieval requirements instead of a wish list Real question types from the participants become a requirements profile: fact questions, process questions, comparison questions. Which metadata, structures and freshness guarantees does each class need?
Day 2 · 13:00	Permissions: from source to answer Workshop on the reference concept: how do ACLs travel from SharePoint and file servers into the index, and how are they enforced on every query? Decision tree for document-level security applied to your own access model.

THE 20 TRACKS · TRACK 13/20 · CONTINUED

Day 2 · 15:00

The maintenance process: staying current instead of cleaning up once

Exercise: the update path is defined for each prioritized source. Who deletes, who versions, how does the index learn about it? Without this process, every map goes stale within months.

Day 2 · 16:15

Day result: the handover-ready package

Each organization files its map, ownership matrix, prioritized source list and draft permission concept as one package. Peer review against a checklist: could a RAG team start with this on Monday?

What you take away

- A knowledge map of your real source landscape, developed in anonymized form, rated by freshness, quality and access structure
- An ownership matrix: a named business owner and a defined update process for every prioritized source
- A prioritized source list for the first RAG iteration with documented exclusion criteria
- A retrieval requirements profile: which question types, document types and metadata the system must serve
- A draft permission concept that carries source ACLs all the way to the answer instead of losing them at indexing time

Prerequisites

Familiarity with your own system landscape and a basic understanding of data structures. Programming skills are not required, architectural thinking is.

Completion

The final review on Day 2 checks every package against a documented checklist: ownership named, prioritization justified, permission path described. Attendance and results are documented in a traceable, audit-ready form.

Included

2 training days with Dino Bordonaro, on site or remote · Reusable templates for knowledge map, ownership matrix and source assessment · A reference permission concept with a decision tree for document-level security · An anonymized evaluation of the source profiles as a situation report · Certificates of attendance and training documentation for your compliance archive · A 60-minute remote office hour 4 weeks after the training to review your map

AIA-AI-RAG

PATH 4 · BUILD



k/aia-ai-rag

Enterprise RAG Engineering

A RAG demo is an afternoon, a RAG system with permissions and measurable quality is engineering. Here you build the latter.

A working RAG prototype on our lab environment with an ingestion pipeline, hybrid retrieval, reranking, an evaluation set, a permission model and honestly documented limits.

Who for: Developers, data engineers and architects who build or own a RAG system for enterprise use.

Duration: 3 days **Group:** 4-12 people **Lead time:** from 4 weeks

€23,700 flat per delivery

AGENDA

Day 1 · 09:00	Why the tutorial RAG fails on your documents Live findings in the lab: a naive RAG against real document types with tables, scans and versions. Every failure is attributed to a component, which produces the build list for the three days.
Day 1 · 10:30	Ingestion: turning documents into usable units Hands-on: parsers for PDF, Office and HTML, handling tables and headers, metadata extraction. Everyone builds the pipeline against a prepared document set with known traps.
Day 1 · 13:00	Chunking is a decision, not a default Exercise with measured comparison: fixed windows, structure-based and semantic chunking against the same question catalog. When each strategy holds and how to document the decision.
Day 1 · 15:00	Embeddings and index in operation Embedding models compared: dimensions, languages, cost. Building the vector index in your own lab environment, first queries against your own corpus.
Day 1 · 16:30	Day result: the pipeline runs Each person demonstrates: document set ingested, chunks inspected, index queryable. The pass criterion is a documented chunking decision with rationale, not just running code.
Day 2 · 09:00	The limits of pure vector search Measurement exercise on your own index: product numbers, proper names and abbreviations that semantic search misses. The error list motivates the hybrid approach.
Day 2 · 10:30	Building hybrid retrieval Hands-on: full-text search next to the vector index, score fusion and weighting. Everyone measures the effect on their own question catalog instead of taking it on faith.
Day 2 · 13:00	Reranking: precision on the last mile Adding a reranker to your own pipeline, including latency and cost measurement. Decision by the numbers: when is the extra step worth it and when is it not?

THE 20 TRACKS • TRACK 14/20 • CONTINUED

Day 2 • 15:00	The evaluation set takes shape The question catalogs you brought become a structured evaluation set with expected answers and source references. An automated run against your own pipeline delivers the first baseline.
Day 2 • 16:30	Day result: measured quality instead of gut feeling Each person presents numbers: hit rate and answer quality for naive search, hybrid, and hybrid with reranking side by side. The measurement table is the verifiable result of the day.
Day 3 • 09:00	Permissions: the question that stops projects Lab demonstration: a RAG without permission checks answers salary questions from an HR document. Then the architectures for document-level security: ACL ingestion, query-time filtering, trusted identity passthrough.
Day 3 • 10:30	Implementing document-level security Hands-on: permission metadata in the index, a security filter in the retrieval layer, testing with two user roles against the same corpus. Verification: the restricted role gets neither answer nor citation.
Day 3 • 13:00	Answer quality: grounding, citations, saying I do not know Prompting the generation layer: answers with source citations, behavior on missing context, handling contradicting documents. A run against the evaluation set shows the effect of every change.
Day 3 • 15:00	Limits document and operations handover Each team writes the honest limits documentation: which question types the prototype answers reliably, which it does not, and what is missing before production. Outlook on observability and release gates as the bridge to AI-EVAL.
Day 3 • 16:15	Day result: the prototype in review Final review per person: live query with two roles, evaluation results, limits document. The group checks against a checklist: does it run, does it measure, does it protect, and does it honestly state what is missing?

What you take away

- A running RAG prototype on the lab environment, built by your own team, with a traceable architecture decision per component
- An ingestion pipeline with documented chunking strategies for prose, tables and structured documents
- Hybrid retrieval combining vector and full-text search with reranking, measured head-to-head against naive vector search
- An evaluation set of 30-50 question-answer pairs and a first measured quality baseline
- An implemented permission model with document-level security and a limits document that honestly states what the prototype cannot do

Prerequisites

Solid programming skills in Python, a basic understanding of APIs and databases. Prior experience with embeddings helps but is not required.

Completion

The final review on Day 3 documents each person's running prototype, evaluation results and limits document against a fixed checklist. Attendance and results are documented in a traceable, audit-ready form.

Included

3 training days with Dino Bordonaro, on site or in our lab · A personal environment per participant on our enterprise lab for the training plus 14 days afterwards · A complete code repository with reference implementation, exercise stages and sample solutions · Templates for evaluation set, permission model and limits documentation · Certificates of attendance and training documentation for your compliance archive · A 90-minute remote office hour 4 weeks after the training for questions from your own implementation

AIA-AI-EVAL

PATH 4 • BUILD

k/aia-ai-eval

Evaluation, Observability and Guardrails

Without a golden dataset and release gates, every AI release is an experiment on your users.

A golden dataset, automated quality metrics, tracing, cost and latency budgets and defined release gates that catch regressions before your users experience them.

Who for: Developers, platform engineers and technical owners taking an AI system towards production or already operating one.

Duration: 2 days **Group:** 4-12 people **Lead time:** from 4 weeks

€15,800 flat per delivery

AGENDA

Day 1 • 09:00	The blind flight diagnosis Live experiment on the reference system: a seemingly harmless prompt fix measurably degrades half the answers. Question to the group: how many of your recent changes went live without measurement?
Day 1 • 10:30	Building the golden dataset Workshop on the cases you brought: formulating expected results, adding edge cases and trap questions, checking representativeness. The result is a versioned dataset of 30-50 cases per team.
Day 1 • 13:00	Implementing metrics: groundedness, relevance, completeness Hands-on: automated scoring with LLM-as-judge, calibration against human spot-check judgments, handling judge errors. By the end, every metric runs against your own dataset.
Day 1 • 15:00	The first evaluation run A full run of your dataset against the system, then outlier analysis: is the failure in retrieval, in the prompt or in the judge? The error classification determines the right fix.
Day 1 • 16:30	Day result: a reliable baseline Each team presents its baseline: dataset size, metric values, classified outliers. Test question: could you approve or reject a change based on these numbers?
Day 2 • 09:00	Tracing: making the chain visible Hands-on: instrumenting the pipeline with traces per request, from retrieval hits through model calls to the answer. Exercise: diagnose a real bad case from the trace alone.
Day 2 • 10:30	Cost and latency budget Analyzing the trace data: what does each request type cost in tokens and milliseconds? Each team defines budgets and alert thresholds and wires them into monitoring.
Day 2 • 13:00	Guardrails before and after the model Implementation: input filters against prompt injection and data exfiltration, output checks for grounding and format violations. Measurement exercise: what do the guardrails catch and what do they cost in latency?

THE 20 TRACKS · TRACK 15/20 · CONTINUED

Day 2 · 15:00

Release gates and regression detection

Workshop: the evaluation run as a mandatory stage in the deployment pipeline, thresholds per metric, rollback criteria. Test on your own system: a deliberately introduced regression must be stopped by the gate.

Day 2 · 16:15

Day result: the gate holds

Sign-off per team: the prepared regression is detected and blocked by your own release gate, the clean version passes. The documented gate ruleset is the verifiable closing artifact.

What you take away

- A golden dataset for your own system: curated cases with expected results, edge cases and deliberately hard queries
- Implemented quality metrics for groundedness, relevance and completeness, including a calibrated LLM-as-judge with human spot-check validation
- Tracing across the full chain from request through retrieval and model calls to the answer, with cost and latency per step
- A cost and latency budget per request type with alert thresholds
- Defined release gates: which measurements a deployment must pass and when a rollback is triggered

Prerequisites

Programming skills in Python and a real AI system as the measurement object, or the one we provide in the training. Experience with CI/CD pipelines helps.

Completion

The transfer proof is the passed gate test on Day 2: regression caught, clean version released, ruleset documented. Attendance and results are documented in a traceable, audit-ready form.

Included

2 training days with Dino Bordonaro, on site, remote or in our lab · A personal lab environment with a prepared reference system and evaluation stack, plus 14 days of continued access · A code repository with metric implementations, judge prompts and pipeline templates · Templates for golden dataset, budget definition and release gate checklist · Certificates of attendance and training documentation for your compliance archive · A 60-minute remote office hour 4 weeks after the training to review your own gates

AIA-AI-AGENT

PATH 4 · BUILD



k/aia-ai-agent

Agentic Systems Engineering

An agent allowed to act before passing its gates is not automation, it is an operational risk with API access.

A self-built agentic workflow with tool integration, an explicit state model, human-in-the-loop gates, evaluation and abort and escalation logic, signed off in a live stress test.

Who for: Experienced developers and architects tasked with connecting agents to real business processes and systems.

Duration: 4 days **Group:** 4-12 people **Lead time:** from 6 weeks

€31,600 flat per delivery

AGENDA

Day 1 · 09:00	Autopsy of an agent failure Live in the lab: a naive agent gets stuck in a tool loop and visibly burns budget. Root cause analysis and derivation of the architecture principles: explicit state, bounded autonomy, verifiable steps.
Day 1 · 10:30	Anatomy of a dependable agent Architecture walkthrough: planner, tool layer, memory, verification layer. Scoping question for every process profile brought along: does this need an agent, or is a workflow with one LLM step enough?
Day 1 · 13:00	Designing tools a model can operate safely Hands-on: tool definitions with tight schemas, idempotency, error returns, least-privilege access. Everyone builds the first two tools of the training process and tests them in isolation.
Day 1 · 15:00	The first end-to-end run Your agent executes the training process end to end against the mock systems for the first time, still without gates. Observation task: where does it make assumptions nobody has verified?
Day 1 · 16:30	Day result: running agent with a findings list Each person demonstrates the running walkthrough and presents a findings list: at least three observed unverified assumptions or risk points. This list is worked off systematically on Days 2 and 3.
Day 2 · 09:00	Making state explicit From prompt memory to a state model: states, transitions, allowed actions per state. Everyone models their training process as a graph before writing more code.
Day 2 · 10:30	Implementing the state model Hands-on: orchestration with persisted state, resumption after a crash, deterministic transitions. Test: abort the process midway and resume it correctly.
Day 2 · 13:00	Context and memory strategy What the agent must know per step and what it must not: working context, long-term knowledge via RAG integration, context budget. Measurement exercise: context discipline versus full context in a cost and quality comparison.

THE 20 TRACKS · TRACK 16/20 · CONTINUED

Day 2 · 15:00	Failure paths first Exercise on the mock systems with injected faults: timeouts, contradictory data, half-finished transactions. The agent must transfer every failure path into a defined state instead of improvising.
Day 2 · 16:30	Day result: controlled handling of faults Sign-off per person: the agent runs through a fault scenario and demonstrably lands in defined states, documented in the state log. Improvised continuation counts as failed.
Day 3 · 09:00	Responsibility only after gates The core principle of this track in detail: action classes from read-only to irreversible, one gate type per class. Workshop on your own process profile: which action gets which gate, and who approves?
Day 3 · 10:30	Implementing human-in-the-loop Hands-on: an approval step with a preview of the planned action, the agent's reasoning and an explicit approve or reject. The agent demonstrably waits instead of proceeding when no answer comes.
Day 3 · 13:00	Abort and escalation logic Implementing budget limits for steps, cost and runtime, detection of loops and goal drift, escalation to defined recipients with state handover. Test: the looping agent from Day 1 against the new limits.
Day 3 · 15:00	Security at the tool boundary Exercise: prompt injection via tool returns and document content designed to lure the agent into unplanned actions. Hardening through input sanitization, action validation and mandatory gates for sensitive classes.
Day 3 · 16:30	Day result: gates under fire Sign-off per person: three prepared fault and disturbance scenarios run against your own agent. You pass if every critical action demonstrably stops at the gate and escalation triggers correctly.
Day 4 · 09:00	Evaluating agents: trajectories instead of answers Evaluation design for multi-step systems: success criteria per process goal, scoring of action paths, cost per successful run. Building the evaluation set for your own training agent.
Day 4 · 10:30	The evaluation run Hands-on: an automated run with success rate, gate triggers, abort reasons and cost per run. Analysis: which failure class dominates, and which architecture change addresses it?
Day 4 · 13:00	The sign-off protocol Workshop: each team writes the sign-off protocol for its agent: gates passed, evaluation results, residual risks, conditions for granting responsibility and for withdrawing it. Honesty is the grading criterion.
Day 4 · 15:00	Transfer to your own process Back to the profile you brought: architecture sketch, gate mapping and effort estimate for the real case in your organization. Peer review of the sketches in pairs with documented objections.
Day 4 · 16:15	Day result: sign-off in the stress test Finale per person: a live run of the agent in front of the group including a fault scenario not known in advance, then handover of sign-off protocol and transfer sketch. The group's checklist decides the sign-off.

What you take away

- A running agentic workflow on the lab environment: planning, tool calls, result verification, built by your own team
- An explicit state model with defined transitions instead of implicit loops, including persistable state for resumption
- Implemented human-in-the-loop gates: which actions the agent may propose and which execute only after approval
- Abort and escalation logic with budget limits for steps, cost and time, and defined escalation recipients
- An evaluation and sign-off protocol documenting under which conditions the agent may take responsibility and under which it may not

Prerequisites

Confident programming practice in Python and API experience. Prior knowledge from AI-RAG or AI-EVAL helps, especially a basic understanding of evaluation.

Completion

The transfer proof is the passed stress test on Day 4 together with the sign-off protocol: gates hold under fault conditions, escalation triggers, residual risks are honestly documented. Attendance and results are documented in a traceable, audit-ready form.

Included

4 training days with Dino Bordonaro, on site or in our lab · A personal lab environment with a training process, mock systems and locally hosted models, plus 14 days of continued access · A code repository with reference architecture, state model templates and sample solutions per day stage · A gate catalog and sign-off protocol template for transfer to your own processes · Certificates of attendance and training documentation for your compliance archive · A 90-minute remote office hour 4 weeks after the training to review your own agent design

AIA-AI-PLATFORM

PATH 4 · BUILD

k/aia-ai-platform

Sovereign AI Platform Engineering

Five days after which your team has not just described an inference platform but built, measured and handed one over.

A referenceable platform architecture from GPU sizing to operations handover, built and measured in the lab and documented as a blueprint.

Who for: Platform teams and infrastructure architects who are expected to own LLM inference in their own data center or on Azure Local.

Duration: 5 days **Group:** 4-12 people **Lead time:** from 6 weeks

€39,500 flat per delivery

AGENDA

Day 1 · 09:00	Your load assumptions under scrutiny Review of the environment profiles: how many concurrent users, which model classes, which response times? Vague expectations become measurable requirements in tokens per second.
Day 1 · 10:30	The reference architecture layer by layer Compute, network, storage, identity, registry, inference, gateway. For each layer the one decision that is expensive to correct later, and how BORDONARO settled it in the Sovereign AI Appliance.
Day 1 · 13:00	GPU capacity model from L40S to H100 VRAM demand, quantization, batching, concurrency: what really drives the choice of card. Live benchmark in the lab, the same model on two GPU classes in direct comparison.
Day 1 · 15:00	Exercise: sizing calculation for your scenario Each team works through its own scenario with the sizing sheet and defends the calculation in front of the group. The most common insight: first sized wrong, then corrected, then understood.
Day 1 · 16:15	Day result: documented capacity model Each team has a justified capacity model for its scenario, checked against the day's benchmark values. It will be validated against the real platform over the course of the week.
Day 2 · 09:00	Azure Local as the foundation Cluster anatomy, Storage Spaces Direct, roles and limits of the platform in the version available at the time. What Azure Local can do, what it cannot, and where a plain Kubernetes base is the more honest answer.
Day 2 · 10:45	Identity thought through end to end Entra ID, local identities, workload identity and secrets management without plaintext in configs. The concrete question: who may call which model with which data, and where is that enforced?
Day 2 · 13:00	Network design and egress control Segmenting the inference zone, internal TLS, controlled egress. Whiteboard exercise: which seven connections does the platform really need, and which of them can be closed?

THE 20 TRACKS · TRACK 17/20 · CONTINUED

Day 2 · 15:00	Exercise: implementing the foundation in the lab Each team sets up identity integration and network segmentation for its lab environment and proves with a test access that the boundaries hold.
Day 2 · 16:15	Day result: approved foundation Identity and network are in place and tested. Each team documents its foundation design in the blueprint, open disputes go into the decision log with reasoning.
Day 3 · 09:00	Registry: provenance and signatures for models Model and container registry as one discipline: versioning, signature verification, provenance. Why an unsigned model from the internet must never land in a sovereign platform.
Day 3 · 10:30	Inference servers in an honest comparison vLLM and alternatives against the same measuring points: throughput, latency, operational effort. Plus a look at Foundry Local on Azure Local in the version available at the time, currently preview, with a clear framing of what that means for production decisions.
Day 3 · 13:00	Build: from model to API endpoint Each team builds the stack end to end: pull the model from the registry, verify the signature, serve it, put it behind the API gateway. At the end the team's own platform answers its first authenticated request.
Day 3 · 15:00	Load test and tuning A load generator against your own stack: change batching parameters, KV cache and quantization and measure the effect. The capacity model from day one meets reality.
Day 3 · 16:15	Day result: running stack with measurement protocol Each team has a working inference stack and a measurement protocol that states and explains the gap between the sizing calculation and the measured performance.
Day 4 · 09:00	Observability for GPU platforms GPU utilization, token throughput, queues, error rates: which metrics an operations team really needs and which alerts are allowed to wake someone at night. Building the dashboard on your own lab platform.
Day 4 · 10:45	Update path: models, drivers, firmware, containers Four update streams with four risk profiles. A concrete decision per stream: maintenance window, canary or rolling, and what must be measured before every model change.
Day 4 · 13:00	Tenants, quotas and internal chargeback Several teams on one GPU platform: quota model, prioritization, cost allocation. Exercise: design and defend a quota model for three competing departments.
Day 4 · 15:00	Incident drill: the node goes down We pull a GPU node from one team in a controlled way. Detect, decide, reroute, communicate: the incident is worked through by runbook and then dissected.
Day 4 · 16:15	Day result: operations dashboard and incident protocol Each team has a working monitoring dashboard and a completed incident protocol with timeline, decisions and improvements for its own runbook.
Day 5 · 09:00	The runbook becomes complete Startup procedures, update flows, incident scenarios, escalation paths: the fragments of the week are consolidated into a runbook skeleton an operations team can take over.

THE 20 TRACKS · TRACK 17/20 · CONTINUED

Day 5 · 10:30

Architecture review against your requirements

Each team's blueprint reviewed against data protection requirements, internal security standards and requirements up to VS-NfD the architecture should be preparable for. Gaps are named, not argued away.

Day 5 · 13:00

Acceptance simulation: handover to operations

Role play with swapped teams: one team hands over its platform, the other checks against the acceptance checklist and may ask anything. Whatever does not survive the handover gets fixed.

Day 5 · 15:00

Transfer plan for Monday

Each team prioritizes its first three steps back home: what gets adopted, what gets decided differently, what needs procurement? The follow-up is scheduled.

Day 5 · 16:15

Final result: referenceable blueprint

Each team leaves with a documented platform blueprint, a validated capacity model, a runbook skeleton and a decision log. A short trainer review of the results closes the week.

What you take away

- A reference architecture document across all layers: compute, network, storage, identity, registry, inference stack
- A GPU and capacity model with a traceable sizing calculation from L40S to H100 for your load assumptions, backed by your own lab measurements
- A running inference stack your team built itself: model registry, inference server, API gateway, monitoring
- An operations handover package with runbook skeleton, update path and incident protocols from the exercises
- A decision log covering the contested architecture questions of your organization with documented reasoning

Prerequisites

Solid infrastructure fundamentals: virtualization, networking, containers and the Linux command line. Azure basics help but are not required.

Completion

The acceptance simulation on day 5 is the transfer evidence: each team hands over its platform against a checklist and passes the partner team's review. Blueprint, measurement protocols and the incident protocol are documented and handed to the participants.

Included

5 training days delivered by Dino Bordonaro · A dedicated lab environment with NVIDIA GPUs per team of two for the entire week, drawn from our enterprise lab with 2,516 cores and 24 TB RAM · Reference architecture templates and infrastructure-as-code building blocks for reuse · Sizing spreadsheet with the benchmark values collected during the course · Certificates of attendance and audit-ready training documentation · A 60-minute follow-up with the team 4 weeks after the course

AIA-AI-OPS

PATH 4 • BUILD



k/aia-ai-ops

LLMOps for Connected, Disconnected and Air-Gapped Environments

Three days after which a model update arrives safely even when no cable leads to the outside world.

A tested promotion process from dev through test to prod, an offline update concept with signature verification and an incident and rollback runbook rehearsed in the lab.

Who for: Platform and operations teams responsible for an LLM platform in environments with restricted or no internet connectivity: public authorities, critical infrastructure, defense, manufacturing with segregated networks.

Duration: 3 days **Group:** 4-12 people **Lead time:** from 4 weeks

€23,700 flat per delivery

AGENDA

Day 1 • 09:00

Three operating models, three truths

Connected, intermittently disconnected, permanently disconnected: what really changes per model in update frequency, monitoring and supportability. Honest framing of Azure Local Disconnected Operations: access-restricted, subject to an eligibility check, not a freely orderable product.

Day 1 • 10:45

Everything is an artifact

Models, containers, prompts, configurations and evaluation data as versioned, signed artifacts. A concrete question per artifact type: where is it created, who signs it, what does the signature prove?

Day 1 • 13:00

Promotion gates that actually stop things

The path from dev through test to prod with gates that are armed: evaluation suite, signature verification, four-eyes approval. Case discussion: which gate would have stopped your organization's last bad model change?

Day 1 • 15:00

Exercise: promotion pipeline in the lab

Each team builds the pipeline: a model artifact passes dev and test, a deliberately degraded model must get stuck at the evaluation gate. If it does not, the gate gets sharpened.

Day 1 • 16:15

Day result: documented promotion process

Each team has a promotion process with defined gates as a document and as a running pipeline, proven by one model that passed and one that was stopped.

Day 2 • 09:00

Anatomy of a transfer package

What passes through the gateway: model weights, container images, signatures, manifests, evaluation data. Building a manifest that makes completeness and integrity verifiable before anything gets installed.

Day 2 • 10:45

The gateway as a process, not a USB stick

Roles, verification steps and logging at the transition into the disconnected zone. Which checks run outside, which inside, and who owns the approval? Cross-checked against the network sketches you brought.

THE 20 TRACKS • TRACK 18/20 • CONTINUED

Day 2 • 13:00	Exercise: offline update under real conditions Each team's lab zone is cut off from the internet. A new model and its container are built into a transfer package, brought through the gateway, verified against the manifest and activated in staging. A tampered package is in circulation and must fail signature verification.
Day 2 • 15:30	Debrief: where it jammed Review of the exercise along the logs: which check was missing, which was redundant, how long did the run take? The findings flow straight into your own update concept.
Day 2 • 16:15	Day result: completed offline update with protocol Each team has performed a full update without internet access and demonstrably rejected the tampered package. The gateway protocol documents every step.
Day 3 • 09:00	Monitoring without a cloud backend Making model quality, drift and platform health visible when no SaaS dashboard is allowed. Building the metrics pipeline in the disconnected zone, including the question of which evaluations may leave periodically and which never.
Day 3 • 10:45	Incident scenarios for LLM platforms Model regression after an update, suspected compromised artifact, GPU failure, a wave of prompt injection: for each scenario the runbook records detection path, immediate action and communication chain.
Day 3 • 13:00	Exercise: rollback under time pressure The model installed the day before shows a massive regression in the exercise. Each team rolls back to the last known good state by runbook, measures the time and proves via evaluation that the old state is restored.
Day 3 • 15:00	Your 90-day plan for real operations Team prioritization: which three gaps between lab exercise and your own environment get closed first, who owns the signature chain in your organization, when does the first in-house gateway test run take place.
Day 3 • 16:15	Day result: tested incident and rollback runbook Each team has a runbook that survived a real rollback exercise, including measured recovery time and a documented 90-day plan for its own environment.

What you take away

- A documented artifact flow for models, containers, prompts and configurations with versioning and a signature chain
- A dev-test-prod promotion process with defined gates, including an evaluation gate before every model change
- An offline update concept: transfer package, gateway process, signature verification and staging without internet access, played through once in full in the lab
- An incident and rollback runbook for model regression, compromised artifacts and platform outage, tested in an exercise
- A monitoring concept that works without a cloud backend and still makes model quality visible

Prerequisites

Operational experience with containers and a basic understanding of CI/CD. Having attended AI-PLATFORM or running an existing inference environment in-house is ideal.

Completion

Three documented proofs: the pipeline that stops a bad model, the gateway protocol of the offline update with the rejected tampered package and the timed rollback exercise. All three are handed to the participants as documentation.

Included

3 training days delivered by Dino Bordonaro · A lab environment with a deliberately isolated network zone per team where the offline exercises take place for real · Templates for the transfer package manifest, promotion gates and the rollback runbook · Certificates of attendance and audit-ready training documentation · A 60-minute follow-up with the team 4 weeks after the course

AIA-AI-ARC

PATH 4 · BUILD

k/aia-ai-arc

Azure Arc and Azure Local as a Hybrid Control Plane

Three days after which Arc is no longer just an agent on your servers but a control plane with rules you set yourself.

A well-founded Arc resource model with a policy and GitOps foundation, plus a reasoned decision picture of when connected and when disconnected is the right operating model.

Who for: Platform teams that run or introduce Azure Local and want to set up management via Azure Arc deliberately instead of inheriting it.

Duration: 3 days **Group:** 4-12 people **Lead time:** from 4 weeks

€23,700 flat per delivery

AGENDA

Day 1 · 09:00	What Arc actually projects The Arc resource model from the inside: servers, Kubernetes, data services and Azure Local as projected resources. Which data flows in which direction, and what stays local? Clarified on the data flow diagram, not the marketing picture.
Day 1 · 10:45	Resource topology for your organization Management groups, subscriptions, resource groups, tags and RBAC: an exercise on your own case. The concrete decision: where do you separate production from test, and who may do what at the Arc level?
Day 1 · 13:00	Onboarding in the lab Each team connects Azure Local instances and further systems to Arc, in the version available at the time. Along the way the team logs which endpoints, accounts and permissions the onboarding actually demands.
Day 1 · 15:15	Arming RBAC Role design in practice: operations team, security team, auditor. Testing with three identities whether the boundaries hold. Whoever sees or can do too much gets adjusted.
Day 1 · 16:15	Day result: documented resource model Each team has onboarded its systems into Arc and documented a resource model with RBAC layout as a diagram, verified with real test accesses.
Day 2 · 09:00	Policy as guardrail, not brake Azure Policy on Arc resources: audit versus deny, initiatives, exemptions with expiry dates. Selection exercise: which ten policies belong in your starting set, and which single one of them is set to deny?
Day 2 · 10:45	Laying the GitOps foundation Configuration as code with Flux on Arc-enabled Kubernetes: repository structure, environment separation, approval paths. The fundamental question gets decided: what must never again be changed via the portal?
Day 2 · 13:00	Exercise: change via pull request Each team changes a workload configuration exclusively through a pull request and watches the GitOps chain deliver it. Then the counter-test: a manual change on the system is detected as drift and rolled back.

THE 20 TRACKS · TRACK 19/20 · CONTINUED

Day 2 · 15:15	Azure Local under Arc in daily operations Updates, monitoring and capacity views on Azure Local through the Arc layer, in the version available at the time. An honest balance: what the control plane covers well today and where local tooling is still required.
Day 2 · 16:15	Day result: running policy and GitOps chain Each team has assigned ten policies, delivered a change via pull request and demonstrably detected and corrected a configuration drift. Repositories and policy definitions leave with the team as artifacts.
Day 3 · 09:00	Connected and disconnected without ideology The two operating models in sober comparison: data flows, update paths, feature scope, operational effort. Azure Local Disconnected Operations framed honestly: access-restricted, subject to an eligibility check, not freely orderable, feature scope in the version available at the time.
Day 3 · 10:45	The criteria catalog Building a decision grid: regulatory constraints, latency and autonomy requirements, staffing, supportability, cost. Every criterion gets a weight and a measurable verification question.
Day 3 · 13:00	Case work: your organization in the decision picture Each team applies the catalog to its own starting position from the prework and produces a recommendation: connected, initiate the disconnected eligibility check, or a mixed model. The recommendation is defended in front of the group.
Day 3 · 15:15	The eligibility reality and plan B Walking the real path towards Disconnected Operations: prerequisites, review process, time horizon. For each team a plan B is recorded in case the check fails or takes too long.
Day 3 · 16:15	Day result: reasoned operating model recommendation Each team has a documented decision picture with a weighted criteria catalog, a recommendation for its own organization and a list of open verification points with owners.

What you take away

- An Arc resource model for your organization: management groups, subscriptions, resource groups, tags and RBAC layout, documented as a diagram
- A policy foundation with the ten most important rules for the start, versioned as code
- A working GitOps chain in the lab: configuration changes via pull request instead of manual work
- A connected versus disconnected decision picture with a criteria catalog, a data flow overview and an honest eligibility assessment
- A documented operating model recommendation for your organization with the open verification points

Prerequisites

Basic knowledge of Azure concepts such as subscriptions, resource groups and RBAC, plus operational experience with servers or Kubernetes. Azure Local practice helps but is provided in the lab.

Completion

Three documented proofs: the resource model with passed RBAC tests, the GitOps chain with detected drift and the operating model recommendation defended in front of the group. Everything is handed over as part of the training documentation.

Included

3 training days delivered by Dino Bordonaro · Lab access with Azure Local instances and a dedicated Azure tenant area per team · Policy and GitOps templates as versioned repositories for reuse · Connected versus disconnected criteria catalog as a decision template · Certificates of attendance and audit-ready training documentation · A 60-minute follow-up with the team 4 weeks after the course

AIA-AI-RESIDENCY

PATH 4 · BUILD

k/aia-ai-residency

Sovereign AI Architecture Residency

Five days in the lab that end not with slides but with your own architecture design, validated on real hardware.

A customer-owned architecture and operations design whose critical assumptions were validated on real hardware, with a tested slice, a prioritized backlog and an executive readout.

Who for: The core team of an organization concretely planning a sovereign AI platform: architects, platform owners and the person who has to defend the investment.

Duration: 5 days in the lab **Group:** 4-8 people **Lead time:** from 6 weeks

€44,500 flat per delivery

AGENDA

Day 1 · 09:00	Target picture and the three critical assumptions Your initiative on one wall: workloads, users, data classes, regulatory framework. The assumptions named in the preliminary call are sharpened and translated into measurable verification criteria.
Day 1 · 10:45	Requirements with veto power Data protection, security and operations state their hard limits before the architecture is drawn: what must the system never do, which evidence up to VS-NfD should be preparable, who will carry operations?
Day 1 · 13:00	Slice selection: a case that proves something From your candidate workloads the use case is chosen that tests the riskiest assumption, not the most convenient one. The bar is set: how do we recognize on Friday that the slice holds?
Day 1 · 15:00	First measurement run: baseline on lab hardware Your target models run on the reserved GPU systems for the first time. Baseline measurement of throughput and latency as the reference for every decision of the week.
Day 1 · 16:15	Day result: verified target picture and measurement plan Target picture, hard limits, the selected slice and a measurement plan for the three critical assumptions exist in writing and are carried by the whole team.
Day 2 · 09:00	Architecture design: the load-bearing decisions Compute, network, storage, identity, registry, inference stack: the design takes shape on the wall, every layer with reasoning and an alternative. Azure Local and Azure Arc are planned in where they carry weight, in the version available at the time.
Day 2 · 10:45	Sizing calculation against the baseline The capacity model is calculated with the baseline values from day 1: how many cards of which class, from L40S to H100, for your peak loads? The gap between vendor claims and your own measurements is quantified.
Day 2 · 13:00	Validation: assumptions one and two under test The first two critical assumptions are checked on the hardware, for example quantization effects on your quality requirement or concurrency behavior under realistic load. Results flow straight back into the design.

THE 20 TRACKS · TRACK 20/20 · CONTINUED

Day 2 · 15:30	The operating model switch Connected, intermittently disconnected, or initiate the disconnected eligibility check: decided against the criteria catalog, with an honest framing of the eligibility prerequisites of Azure Local Disconnected Operations.
Day 2 · 16:15	Day result: robust architecture design v1 The design exists in version 1, with a sizing calculation based on your own measurements and a documented operating model decision including open verification points.
Day 3 · 09:00	Slice build: the foundation Building the slice on the lab environment along the design: identity integration, network segment, registry connection with signature verification. Your team builds, the trainer keeps design and reality aligned.
Day 3 · 13:00	Slice build: data and application logic Your anonymized sample data meets the stack: preparation, indexing or integration depending on the case, plus the application layer. Where it honestly fits, patterns from Sovereign Assistant or VOXA serve as references.
Day 3 · 15:30	First end-to-end run The slice answers end to end for the first time: from authenticated request to a substantiated answer. Whatever jams is logged and prioritized instead of talked up.
Day 3 · 16:15	Day result: end-to-end slice The use case runs end to end on the lab environment. Deviations from the design are logged and scheduled for day 4.
Day 4 · 09:00	Assumption three and the load test The remaining critical assumption is tested, then the slice faces a load test against the bar from day 1: throughput, latency, behavior at the capacity limit. Every deviation gets a cause or a backlog item.
Day 4 · 10:45	Operations and security drill The slice is treated like production: a monitoring view is built, a model rollback is rehearsed, a prompt injection attempt is run against the protection layer. Results flow into the operations design.
Day 4 · 13:30	Backlog prioritization Everything open becomes a backlog with effort classes, dependencies and risks: the path from validated slice to productive platform, prioritized by the team, not the trainer.
Day 4 · 15:30	Investment picture and make or buy Sizing, operating model and backlog are condensed into an investment picture: self-build, Sovereign AI Appliance or a mixed path, with numbers and named uncertainties for the readout.
Day 4 · 16:15	Day result: validated design with backlog Architecture and operations design exist in a verified version, all three critical assumptions are answered with measurements, the backlog is prioritized.
Day 5 · 09:00	Hardening the design A final review round with swapped roles: the team attacks its own design, the trainer defends it, and vice versa. Whatever does not survive the attack is changed or documented as a risk.
Day 5 · 10:30	Readout rehearsal The executive readout is rehearsed once in full: key messages, measurements, investment picture, recommendation. The team decides who presents which part to its own leadership.

THE 20 TRACKS · TRACK 20/20 · CONTINUED

Day 5 · 13:00

Document handover

Design, measurement protocols, backlog and readout deck are finalized and handed over. All artifacts belong to you and remain usable without BORDONARO.

Day 5 · 14:30

Executive readout

90 minutes in front of your leadership level, on site in the lab or joining remotely: the validated design, a live demonstration of the slice, the investment picture and the recommendation with open risks. Questions are answered by the team, not just the trainer.

Day 5 · 16:15

Final result: decision basis handed over

Your leadership has seen a validated design with measurements, a tested slice and a prioritized backlog. Next steps and the two follow-up sessions are scheduled.

What you take away

- An architecture and operations design for your platform, in your structure and with your requirements, not a generic template
- Validated core assumptions: sizing, throughput and response times of your target workloads, measured on GPU systems of our lab from L40S to H100
- A tested slice: one prioritized use case of your organization runs end to end on the lab environment by the end of the week
- A prioritized implementation backlog with effort classes, dependencies and named risks
- An executive readout on day 5: design, measurements and investment picture in 90 minutes for your leadership level, with a takeaway document

Prerequisites

A decision-ready initiative and a team that carries it: at least one architect or platform owner with infrastructure depth. Prior attendance of AI-PLATFORM helps but is not required.

Completion

The transfer evidence is the executive readout itself: your team presents design, measurements and recommendation to its own leadership and answers their questions. All validations are documented as measurement protocols and are part of the handover.

Included

5 days of exclusive work in the enterprise lab with Dino Bordonaro, no parallel operation with other customers · Reserved GPU systems of several classes for comparative measurements, from the cluster with 2,516 cores, 24 TB RAM and 3,400 TB storage · A 90-minute preliminary call and preparation of your documents before the week · Architecture and operations design, measurement protocols and backlog as handover-ready documents · Executive readout deck for your leadership level · Two 60-minute follow-up sessions in the 8 weeks after the residency

PRICE MATRIX

ALL PRICES AT A GLANCE

ITEM NO.	SERVICE	SCOPE	NET PRICE
AIA-AI-LIT-LAB	AI Literacy Hands-on Lab	½ day	€3,950 flat per delivery
AIA-AI-WORK	Working Safely and Productively with AI	1 day	€7,900 flat per delivery
AIA-AI-CONTEXT	Prompt and Context Engineering	2 days	€15,800 flat per delivery
AIA-AI-START	AI Onboarding for Career Starters	1 training day	€7,900 flat per delivery
AIA-AI-CHECK	Checking and Improving AI Results	1 training day	€7,900 flat per delivery
AIA-AI-EXEC	Sovereign AI Executive Briefing	½ day	€4,900 flat per delivery
AIA-AI-LEAD	Leading in the Human-Agent Organization	2 days	€15,800 flat per delivery
AIA-AI-VALUE	Use Case and Business Case Sprint	2 days	€15,800 flat per delivery
AIA-AI-CHAMP	AI Champions Program	4 on-site days + 6 Q&A sessions over 12 weeks	€34,900 flat per delivery
AIA-AI-GOV	EU AI Act and AI Governance	2 days	€15,800 flat per delivery
AIA-AI-RISK	AI Risk Classification Lab	1 day	€7,900 flat per delivery
AIA-AI-SEC	Secure GenAI and Agent Threat Modeling	2 days	€15,800 flat per delivery
AIA-AI-DATA	AI-ready Data and Knowledge Architecture	2 days	€15,800 flat per delivery
AIA-AI-RAG	Enterprise RAG Engineering	3 days	€23,700 flat per delivery
AIA-AI-EVAL	Evaluation, Observability and Guardrails	2 days	€15,800 flat per delivery
AIA-AI-AGENT	Agentic Systems Engineering	4 days	€31,600 flat per delivery
AIA-AI-PLATFORM	Sovereign AI Platform Engineering	5 days	€39,500 flat per delivery
AIA-AI-OPS	LLMOps for Connected, Disconnected and Air-Gapped Environments	3 days	€23,700 flat per delivery
bordonaro.ai AIA-AI-ARC	Azure Arc and Azure Local as a Hybrid Control Plane	3 days	Prices €23,700 flat per delivery

TERMS & CONTACT

CLEAR RULES. SHORT PATHS.

Training price anchor	€7,900 net per training day and group, half-day formats pro rata, the lab residency at €8,900.
Minimum group	Delivery from four actually registered participants, up to the stated group size. Smaller groups move to the next date free of charge.
Scheduling	Depending on scope, 2 to 6 weeks lead time counted from written confirmation of date and scope.
Completion	Certificate of participation and audit-ready training records included with every track, as a building block of your Art. 4 EU AI Act evidence.
All prices	Net plus 19 % VAT. Travel costs only for on-site dates outside the region, stated up front.

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

Ludwigshafen registry HRA 61166 · General partner: BORDONARO Beteiligungs GmbH, HRB 64031 · VAT ID DE815511519 · Represented by Dino Bordonaro



BORDONARO

**NOT EVERYTHING NEEDS TO BE OFFLINE. EVERYTHING NEEDS TO STAY
UNDER CONTROL.**

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

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