

PROGRAMME OVERVIEW

AI & Generative AI Practitioner Programme

Eighteen weeks, one day a week, five hours a day. A mixed non-technical cohort taken from no prior knowledge to designing, governing and defending generative AI in a real organisation.

18 weeks

90 guided learning hours

6 blocks

3 checkpoints + capstone

Instructor-led, online

No prior technical knowledge

THE SIX BLOCKS

Block A**Foundations**

Weeks 1-4

What generative AI is, what it can do, and how it works

Block B**Prompting Craft**

Weeks 5-8

Getting reliable, repeatable, high-quality output

Block C**Microsoft Copilot & M365**

Weeks 9-10

Hands-on Copilot across M365, Copilot Studio agents, and the governance to roll it out safely

Block D**Risk & Governance**

Weeks 11-13

What goes wrong, what the law requires, and how to control it

Block E**Agents & Automation**

Weeks 14-15

From question-and-answer to goal-and-result

Block F**Business Value & Career**

Weeks 16-18

Where the value is, why pilots fail, and where you go next

Prefer to start smaller?

Not ready for the full eighteen weeks? Individual blocks can be taken as standalone modules, a lower-cost way to focus on one area such as [Microsoft Copilot & M365](#), [Prompting Craft](#) or [Risk & Governance](#). Ask about single-module options and we'll tell you what is currently available.

THE TEACHING DAY, 300 MINUTES OF LEARNING TIME

TIME	SESSION	LENGTH
09:00	Arrival, recap and outcomes	20 min
09:20	Session 1: Core concepts	60 min
10:20	Break	15 min
10:35	Session 2: Guided lab	55 min
11:30	Session 3: Deep dive	40 min
12:10	Lunch	40 min
12:50	Session 4: Applied workshop	70 min
14:00	Break	10 min
14:10	Session 5: Discussion and share-back	40 min
14:50	Consolidation, homework and exit ticket	15 min

Balance. Taught input 100 minutes (33%). Lab 55, workshop 70, discussion 40-55% of learning time is active. Week 18 runs to a different timetable for capstone presentations and the final assessment, and still totals 300 minutes.

THE EIGHTEEN WEEKS

Expand all

Collapse all

Welcome & The AI Landscape

What generative AI actually is, where it came from, and what it can do

LEARNING OUTCOMES

- Explain the difference between discriminative AI and generative AI using a workplace example of your own
- Describe the seven capabilities of generative AI and give a use case for each
- Place the key milestones of AI development on a timeline from the 1950s to 2026
- Set up and safely use at least two frontier AI assistants on the free tier
- State the programme's ground rules for confidential data and academic honesty

KEY TAKEAWAYS

- Discriminative AI classifies; generative AI creates. Both are deep learning underneath.
- There are seven capability families: text, image, audio, video, code, data and virtual worlds.
- The technology is older than the hype; the interface is what changed in November 2022.
- Near-universal adoption is not the same as transformation, 37% of organisations use AI at surface level with minimal process change (Deloitte, State of AI in the Enterprise, 2026).
- Nothing confidential goes in a prompt, and you are accountable for everything you submit.

LAB · 55 MIN

Lab 1: First contact with frontier assistants

Output: Completed comparison grid, minimum one documented error per tool

WORKSHOP · 70 MIN

Workshop: Mapping the seven capabilities onto your own work

A ranked seven-row capability map in Workbook W1.5. Two groups present.

ASSESSMENT

Formative: exit ticket self-assessment against the five outcomes. Trainer records a baseline confidence score in the tracker.

DIRECTED STUDY

- Use one AI assistant daily for a real task; log what worked and what failed (45 min)
- Read the Week 1 glossary, Workbook W1.7 (30 min)
- Bring one real piece of work you would like to speed up

SOURCE MATERIAL

- C1 L1 Course Introduction; Specialisation Introduction
- C1 L2 Introduction to Generative AI; History and Evolution
- C1 L3 Capabilities of Generative AI

LEARNING OUTCOMES

- Select an appropriate generative AI tool for a given task and justify the choice
- Describe the four image transformation capabilities and demonstrate two of them
- Explain why an organisation might choose a locally-hosted model over a hosted one
- Identify the data-handling risk in a proposed tool and state a mitigation
- Produce a first-draft business asset using at least two different modalities

KEY TAKEAWAYS

- Choose tools by task, not by vendor loyalty. The frontier assistants overlap; the specialists earn narrow wins.
- Assume anything you type into a consumer tool may be retained and reviewed.
- Image models do four things beyond generation: translate, restyle, inpaint and outpaint.
- Training-data provenance determines what you can safely ship commercially.
- 'Technically correct but does the wrong thing' is the dominant failure mode across every modality.

LAB · 55 MIN**Lab 2: Four modalities in one hour**

Output: One draft communication plus one generated image, with an honest time log

WORKSHOP · 70 MIN**Workshop: The tool selection matrix**

A completed and peer-challenged five-row selection matrix in Workbook W2.5.

ASSESSMENT

Formative: peer challenge of the tool selection matrix. Trainer records participation and quality of risk identification.

DIRECTED STUDY

- Complete the five-task tool selection matrix in Workbook W2.6 (60 min)
- Read your employer's IT acceptable use policy and note where it is silent on AI (30 min)

SOURCE MATERIAL

- C1 L4 Applications of Generative AI
- C1 L5-L8 Tools for Text, Image, Audio and Video, Code Generation

LEARNING OUTCOMES

- Describe the three layers of a neural network and what parameters, weights and biases are
- Distinguish supervised from unsupervised learning and give a business example of each
- Explain the four generative model families and match each to a suitable task
- Explain in two sentences, to a non-technical colleague, why a transformer beat what came before
- Justify a model-family choice for a real brief and defend it under challenge

KEY TAKEAWAYS

- A neural network is input, hidden and output layers; weights sit between neurons, biases sit inside them.
- Supervised learning needs labels and is auditable; unsupervised learning scales and surprises you.
- Transformers replaced recurrence with attention, which is why they scale to long sequences.
- Four generative families: VAEs compress, GANs compete, transformers attend, diffusion denoises.
- Published parameter counts for frontier models are largely unverified. Cite sources and say 'reportedly'.

LAB · 55 MIN**Lab 3: Watching a model think (temperature and reasoning)**

Output: Temperature comparison table plus one documented reasoning failure

WORKSHOP · 70 MIN**Workshop: Model selection clinic**

A defended model-family recommendation, recorded in Workbook W3.5.

ASSESSMENT

Formative: self-check quiz (15 questions) plus the written explainer, peer-marked in Week 4 against a clarity rubric.

DIRECTED STUDY

- Write a one-page plain-English explainer of how an LLM works (45 min)
- Complete the Workbook W3.6 self-check quiz, 15 questions (30 min)
- Revise Weeks 1-3 for the Block A checkpoint

SOURCE MATERIAL

- C4 L1 Deep Learning and Large Language Models
- C4 L2 Generative AI Models

Foundation Models, Platforms & the 2026 Landscape

What a foundation model is, who makes them now, and where you will actually run them

LEARNING OUTCOMES

- Define a foundation model and place generative models, foundation models and LLMs in the correct nested relationship
- Name the current frontier and open-weight model families as of August 2026 and state which are open
- Compare open-weight and closed-weight deployment on cost, control, capability and risk
- Navigate at least one enterprise AI platform and run a prompt against two different models
- Complete the Block A checkpoint assessment

KEY TAKEAWAYS

- A foundation model is trained once on vast data and adapted to many applications.
- All LLMs are foundation models; not all foundation models are LLMs. Same pattern one level up.
- The model landscape turns over roughly every six months. Cite dates and verify before teaching.
- Open weights buy you control and privacy; closed weights buy you capability and convenience.
- Google AI Studio is the most reliable free classroom platform; test every free tier before you rely on it.

LAB · 55 MIN

Lab 4: Same prompt, different models

Output: Three-model comparison scored against pre-declared criteria

WORKSHOP · 70 MIN

Block A checkpoint assessment + capstone briefing

Completed checkpoint paper; draft capstone problem statement in Workbook Appendix C.1.

ASSESSMENT

SUMMATIVE CHECKPOINT 1: Block A. 30 minutes, closed book, 50 marks. MCQ, short answer and applied scenario. Recorded in the tracker. 60% is the recorded pass threshold; below 50% triggers a support conversation.

DIRECTED STUDY

- Draft your capstone problem statement using the Workbook Appendix C.1 template (60 min)
- Audit which AI tools your employer has already licensed (30 min)

SOURCE MATERIAL

- C4 L3 Foundation Models; IBM Granite Foundation Models
- C4 Platforms for Generative AI (watsonx.ai, Hugging Face)

The Anatomy of a Prompt

Four building blocks, six steps, and the difference between asking and engineering

LEARNING OUTCOMES

- Identify the four building blocks of a well-constructed prompt in any example given to you
- Rewrite a naive prompt into an engineered prompt and evidence the improvement
- Apply the six-step prompt engineering process to a real task from your own work
- Explain why prompt engineering is iterative rather than a single correct answer
- Start a personal prompt library with at least five reusable templates

KEY TAKEAWAYS

- Four building blocks: instruction, context, input data, output indicator.
- The output indicator is the most neglected and the highest-leverage of the four.
- Prompt engineering is a six-step loop. Most good prompts are on their third or fourth version.
- Change one thing at a time, or you will not know what worked.
- When a task recurs, stop prompting and start building a reusable template.

LAB · 55 MIN

Lab 5: Rebuild your worst prompt

Output: Four-version prompt log with scores and an annualised time saving

WORKSHOP · 70 MIN

Workshop: The prompt clinic

A diagnosed and rebuilt prompt from another learner, with before-and-after outputs.

ASSESSMENT

Formative: prompt log reviewed against the four-building-blocks rubric. Prompt library progress recorded in the tracker.

DIRECTED STUDY

- Add five templates to your personal prompt library (60 min)
- Apply the six-step process to a second real task and log all versions (30 min)

SOURCE MATERIAL

- C3 L1 What is a Prompt
- C3 L2 What is Prompt Engineering

Prompting Craft: Best Practice & Core Techniques

Clarity, context, precision, persona, plus zero-shot, few-shot and bias mitigation

LEARNING OUTCOMES

- Apply the four best-practice dimensions to improve any prompt
- Use the persona pattern deliberately and explain what it changes
- Distinguish zero-shot from few-shot prompting and choose correctly between them
- Write a prompt that actively mitigates a foreseeable bias in the output
- Use framing and domain terminology to constrain a model to a required boundary

KEY TAKEAWAYS

- Four quality dimensions sit on top of the four structural blocks: clarity, context, precision, persona.
- A persona changes vocabulary, priorities and, most importantly, what gets left out.
- Few-shot prompting is the answer when you cannot describe what you want but you can show it.
- Explicit neutrality instructions measurably reduce stereotyped output.
- Prompt-level mitigation is a patch. It fails silently when nobody remembers to apply it.

LAB · 55 MIN

Lab 6: The four-dimension audit

Output: Five-version dimension audit plus a zero-shot / few-shot comparison

WORKSHOP · 70 MIN

Workshop: Persona swap and bias stress test

A persona comparison grid and a tested bias-mitigation instruction, in Workbook W6.5.

ASSESSMENT

Formative: dimension audit reviewed. Bias-mitigation instruction peer-tested. Both recorded in the tracker.

DIRECTED STUDY

- Expand your prompt library to ten technique-tagged templates (60 min)
- Rewrite the ten supplied weak prompts in Workbook W6.6 (45 min)

SOURCE MATERIAL

- C3 L3 Best Practices for Prompt Creation
- C3 L5 Text-to-Text Prompt Techniques

07

Advanced Prompting Patterns

Interview, chain-of-thought, tree-of-thought, and how to stop a model agreeing with you

LEARNING OUTCOMES

- Run the interview pattern to extract requirements you had not articulated
- Construct a chain-of-thought prompt with a worked exemplar and explain why it helps

- Apply the tree-of-thought pattern to a decision with genuinely competing options
- Design an adversarial review that forces a model to argue against its own output
- Select the right pattern for a given task rather than defaulting to one

KEY TAKEAWAYS

- Interview for discovery, chain for procedure, tree for choice, adversarial for confidence.
- 'One at a time' is what makes the interview pattern work at all.
- A chain-of-thought exemplar must cover every kind of reasoning the real question needs.
- Tree-of-thought explores several branches at once and prunes the weak ones.
- Models argue equally well in both directions. Adversarial review is how you stop that being a problem.

LAB · 55 MIN

Lab 7: Four patterns, one problem

Output: Four pattern outputs plus a written reflection on which changed your thinking

WORKSHOP · 70 MIN

Workshop: Build a decision council

A four-persona decision council with a chairman verdict, in Workbook W7.5.

ASSESSMENT

Formative: decision council reviewed for genuine persona diversity. Capstone statement revision recorded.

DIRECTED STUDY

- Add all four advanced patterns to your prompt library as templates (60 min)
- Run an adversarial review on your capstone problem statement and revise it (45 min)

SOURCE MATERIAL

- C3 L6 Interview Pattern Approach
- C3 L7 Chain-of-Thought Approach
- C3 L8 Tree-of-Thought Approach

08

Multimodal Prompting & Your Prompt Library

Image prompting techniques, negative prompts, and turning practice into reusable assets

LEARNING OUTCOMES

- Apply the five image prompting techniques and evidence the effect of each
- Use weighted terms and negative prompts to control an image model's output
- Structure a personal prompt library so a colleague could pick it up and use it
- Explain when a reusable template beats an ad-hoc prompt

- Complete the Block B checkpoint assessment

KEY TAKEAWAYS

- Five image techniques: style modifiers, quality boosters, repetition, weighted terms, negative prompts.
- A negative weight actively pushes away from a concept; omission does not.
- Negative instruction is well known in image prompting and badly under-used in text prompting.
- A prompt library needs naming, tagging, versioning and documented failure modes.
- A library only its author can use has failed.

LAB · 55 MIN

Lab 8: Stack the five image techniques

Output: Five-stage image progression with all prompts, plus a text transfer test

WORKSHOP · 70 MIN

Block B checkpoint assessment + prompt library peer review

Completed checkpoint paper; peer-review feedback recorded in Workbook W8.6.

ASSESSMENT

SUMMATIVE CHECKPOINT 2: Block B. 45 minutes, 50 marks. Practical prompt rebuild, prompt library review, and one written explanation. Recorded in the tracker. 60% recorded pass threshold.

DIRECTED STUDY

- Finalise, back up and share your prompt library; record a colleague's feedback (60 min)
- Pre-read the Week 9 outcomes and glossary in your workbook (30 min)

SOURCE MATERIAL

- C3 L9 Text-to-Image Prompt Techniques
- C3 L4 Common Prompt Engineering Tools

09

Copilot across M365

Word, Excel, PowerPoint, Outlook and Teams, applied to real work

LEARNING OUTCOMES

- Use Copilot in Word, Excel, PowerPoint, Outlook and Teams for real tasks: drafting, summarising, analysis and meeting notes
- Apply the prompting techniques from Block B to Copilot, which works from your own files, emails and meetings
- Judge where Copilot adds value and where it does not, and check its output before you rely on it

KEY TAKEAWAYS

- Copilot works from your organisation's own content, so context and permissions decide the quality.
- The prompting craft from Block B transfers directly to Copilot.
- A confident answer is still only a draft. Verify everything.

LAB · 55 MIN

Copilot across three M365 apps on a real task

Output: completed task with an honest before-and-after time log

WORKSHOP · 70 MIN

Build a reusable Copilot prompt set for your role

Peer-reviewed prompt set in Workbook CP.1

ASSESSMENT

Formative: exit ticket and the completed Copilot task log, recorded in the tracker.

DIRECTED STUDY

- Use Copilot daily on real work and log time saved and errors (45 min)
- Read your organisation's Copilot and data policies (30 min)

SOURCE MATERIAL

- Microsoft Copilot for M365 documentation
- Flowline Robotics Copilot lab guide CP.1

10

Extending and governing Copilot

Copilot Studio agents, data handling, licensing and rollout

LEARNING OUTCOMES

- Build a simple, no-code agent in Copilot Studio for a routine task
- Explain how Copilot handles M365 data: permissions, sensitivity labels, data residency, and what Copilot can and cannot see
- Assess licensing and a practical rollout, including spotting and managing shadow AI

KEY TAKEAWAYS

- Copilot inherits each user's permissions, so oversharing is the main risk.
- Sensitivity labels and governance come first, before a wide rollout.
- Licensing and adoption decide whether the value is real.

LAB · 55 MIN

Build and test a small Copilot Studio agent

Output: a working no-code agent for a repeatable task in your area

WORKSHOP · 70 MIN

One-page Copilot rollout and governance plan

Completed plan in Workbook CP.2

ASSESSMENT

Applied checkpoint: Copilot task log plus the one-page rollout and governance plan. Recorded in the tracker.

DIRECTED STUDY

- Draft a Copilot rollout note for your own team (45 min)
- Review Microsoft Purview sensitivity labels in your tenant, if available (30 min)

SOURCE MATERIAL

- Microsoft Copilot Studio and Purview documentation
- Flowline Robotics Copilot governance guide CP.2

11

Limitations, Hallucination & Grounding

Why models get things confidently wrong, and what actually fixes it

LEARNING OUTCOMES

- Name the principal limitations of generative AI and identify which are engineering problems and which are inherent
- Explain hallucination in text, image and code, with a worked example of each
- Describe how retrieval-augmented generation grounds a model, and what it does not fix
- Design a verification step appropriate to the risk of a given task
- Detect a hallucination in output relating to your own area of expertise

KEY TAKEAWAYS

- A generative model is not a repository of information. It generates plausible output, not verified output.
- Hallucination has one root cause, likelihood maximisation, and three visible flavours: text, image and code.
- Nothing in the output tells you which parts are fabricated. The verification step must be external.
- RAG fixes staleness and sourcing. It does not fix bad retrieval, and a wrong answer with a citation is worse than one without.
- Verification effort should be proportionate to risk, not applied uniformly.

LAB · 55 MIN

Lab 9: Hunt for hallucinations in your own field

Output: Twenty-question accuracy audit with a grounded / ungrounded comparison

WORKSHOP · 70 MIN

Workshop: Design a proportionate verification step

A six-task risk-proportionate verification design in Workbook W9.5.

ASSESSMENT

Formative: verification design reviewed for proportionality. Hallucination audit quality recorded in the tracker.

DIRECTED STUDY

- Design proportionate verification steps for three tasks in your role (60 min)
- Find and summarise one published case of AI hallucination causing real harm (30 min)

SOURCE MATERIAL

- C2 L1 Limitations of Generative AI
- C2 L3 Hallucinations of Text and Image LLMs; L4 Hallucinations of Code LLMs
- C2 L6 Enhancing LLM Accuracy with RAG (Marina Danilevsky, IBM Research)

12

Bias, Privacy, Copyright & the Law

The four ethical concerns, deepfakes, and the regulatory position as it stands in August 2026

LEARNING OUTCOMES

- Identify the four principal ethical concerns and give a concrete example of each from your own sector
- Explain how bias enters a model and where in the lifecycle it can be addressed
- State the current EU AI Act obligations and the dates they apply from
- Describe the UK's regulatory approach and how it differs from the EU's
- Assess a proposed AI use case against privacy, copyright and disclosure requirements

KEY TAKEAWAYS

- Four ethical concerns: inaccuracy, bias, privacy and security, copyright.
- You can act on bias at prompt level and review level. Only process controls survive contact with scale.
- Copyright in AI output is genuinely unsettled. Read every tool's terms, some assign you liability as well as ownership.
- Deepfake detection is an arms race; disclosure and provenance are more durable controls.
- EU AI literacy obligations have applied since February 2025. High-risk deadlines moved to Dec 2027 and Aug 2028.

LAB · 55 MIN

Lab 10: Assess a real use case against the risk screen

Output: Completed use case risk assessment with named control and owner

WORKSHOP · 70 MIN

Workshop: The regulatory readiness review

A three-deployment regulatory readiness review and a single board recommendation.

ASSESSMENT

Formative: use case risk assessment reviewed against the eight-question screen. Recorded in the tracker.

DIRECTED STUDY

- Complete a use case risk assessment for a colleague's use case, not your own (60 min)
- Review the acceptable use policy draft and mark what is missing (45 min)

SOURCE MATERIAL

- C2 L2 Issues and Concerns About Generative AI
- C2 L5 AI Portraits and Deepfakes
- C2 L7 Legal Issues and Implications of Generative AI

13

Responsible AI in Practice

Four considerations, seven guardrails, and governance you can actually run on Monday

LEARNING OUTCOMES

- Apply the four considerations for responsible generative AI to a real deployment
- Audit an organisation against the seven safety guardrails and identify the biggest gap
- Compare the responsible AI frameworks published by IBM, Google, OpenAI and Microsoft
- Draft an acceptable use policy fit for your own organisation
- Complete the Block D checkpoint assessment

KEY TAKEAWAYS

- Four considerations: transparency, accountability, privacy, safety guardrails.
- A model cannot be accountable. The organisation deploying it is, every time.
- Seven guardrails. Monitoring and user education are the two most commonly missing.
- All four vendor frameworks name transparency; all are voluntary.
- A policy that is comprehensive but unenforceable is worse than a short one people follow.

LAB · 55 MIN

Lab 11: The seven-guardrail audit

Output: Scored guardrail audit with named owners and one written proposal

WORKSHOP · 70 MIN

Block D checkpoint, acceptable use policy and defence

A defended one-page acceptable use policy. Assessed and recorded.

ASSESSMENT

SUMMATIVE CHECKPOINT 3: Block C. Applied. 50 marks: one-page acceptable use policy (35) plus five-minute defence (15). Recorded in the tracker. 60% recorded pass threshold.

DIRECTED STUDY

- Revise your policy against the challenge and send it to a real colleague (45 min)
- Finalise your capstone problem statement (30 min)

SOURCE MATERIAL

- C2 L8 Considerations for Responsible Generative AI
- C2 L9 Implementing Responsible Generative AI Across Domains
- C2 L10 AI Ethics: Perspective of Key Players; Trustworthy AI: An IBM Perspective

14

From Chat to Agents

The shift from question-and-answer to goal-and-result

LEARNING OUTCOMES

- Explain the difference between a chat interaction, a workflow and an agent
- Describe the observe-think-act loop and identify what makes it stop
- Name the four levers that tune an agent and map each to the loop
- Judge correctly when a task warrants an agent and when chat is sufficient
- Write a goal specification precise enough for an agent to act on

KEY TAKEAWAYS

- Chat is question to answer. An agent is goal to result.
- Observe, think, act, repeat, until the goal is met. The goal definition controls everything.
- Four levers: context, skills, model, tools. Model choice matters least.
- An agent is a really capable stranger. Onboard it like one.
- If the same steps happen in the same order every time, you want a workflow, not an agent.

LAB · 55 MIN

Lab 12: Watch an agent work, then improve the goal

Output: Two annotated agent traces with analysis of what the clearer goal changed

WORKSHOP · 70 MIN

Workshop: The goal specification clinic

A goal specification that survived adversarial literal execution, in Workbook W12.5.

ASSESSMENT

Formative: goal specification tested by adversarial literal execution. Task audit reviewed and recorded.

DIRECTED STUDY

- Write three agent-ready goal specifications (60 min)
- Audit one week of your work; classify every recurring task as chat, workflow or agent (45 min)

SOURCE MATERIAL

- Supplementary: AI Agents Explained (Open Residency, 2026)
- Supplementary: Model Context Protocol and the Agentic AI Foundation
- C5 L1 Generative AI Trends and Impacts (custom GPTs, Auto-GPT)

15

Building Your Personal AI Operating System

Context files, tools and MCP, and turning your expertise into reusable skills

LEARNING OUTCOMES

- Build a structured context set describing you, your role and your organisation
- Explain what the Model Context Protocol is and why a neutral standard matters
- Connect at least one tool to an assistant and describe the permission model you would apply
- Convert a process you already perform into a documented, reusable skill
- Explain why skills, not prompts, are where organisational value accumulates

KEY TAKEAWAYS

- Context is what the agent should already know. Plain text, owned by you, portable between tools.
- MCP is a neutral standard under the Linux Foundation, backed by every major lab.
- Start every tool connection read-only. Earn each permission stage with observed behaviour.
- A skill is a standard operating procedure written for a machine. It is where your expertise lives.
- Test skills cold, on someone else. Almost none survive the first attempt.

LAB · 55 MIN

Lab 13: Build your personal AI operating system

Output: Four reviewed context files and one twice-revised working skill

WORKSHOP · 70 MIN

Workshop: Skill build and cold test

A skill that has survived two independent cold tests, with the defect log.

ASSESSMENT

Formative: skill assessed by two independent cold tests. Context file quality reviewed. Recorded in the tracker.

DIRECTED STUDY

- Build and cold-test two more skills (90 min)
- Draft the automation section of your capstone (30 min)

SOURCE MATERIAL

- Supplementary: AI Agents Explained, context, tools, skills
- Supplementary: MCP donated to the Linux Foundation's Agentic AI Foundation, Dec 2025
- C5 L2.6 Demo: Creating Custom GPTs (Bradley Steinfeld, IBM)

16

The Business Case

Where the value actually is, why most pilots fail, and how to run a feasibility assessment

LEARNING OUTCOMES

- Identify where generative AI value concentrates and explain the 75% rule
- Apply the five-pillar adoption framework to a real proposal
- Build a simple dashboard of your own operational data and identify where an AI assistant misreads it
- Explain the principal reasons enterprise AI pilots fail to reach production
- Run a feasibility-versus-impact assessment and build a defensible business case for one use case

KEY TAKEAWAYS

- Roughly 75% of estimated value sits in customer operations, marketing and sales, software engineering and R&D (McKinsey, 2023).
- Five pillars: feasibility, prioritisation, talent, risk, alignment. Alignment is the one that gets skipped.
- An AI assistant reading your data does not know what happened in your business. It will over-read a trend it has no context for.
- Pilots fail on undefined success measures, unready data, unredesigned processes and absent owners.
- If you cannot write a success measure your sponsor would accept, the case is not ready.

LAB · 55 MIN

Lab 14: Read your own numbers

Output: Three dashboard views with AI interpretations and one documented misreading

WORKSHOP · 70 MIN

Workshop: Feasibility clinic and investment committee

A plotted use case with five-pillar analysis and a testable success measure, plus a funding decision with written reasons, in Workbook W14.4.

ASSESSMENT

Formative: business case pitched to a peer investment committee, with written feedback on the weakest element.

DIRECTED STUDY

- Complete the business case section of your capstone (90 min)
- Rebuild one of your dashboard views for a real audience and note what you changed (30 min)

SOURCE MATERIAL

- C5 L1.3 Elevating Businesses Through Generative AI
- C5 L1.4 An Ecosystem of Generative AI Providers
- C5 L1.5 Impact of Generative AI on Different Industries

17

Careers, Roles & Building Your Showcase

How roles are changing, where you fit, and evidencing what you can now do

LEARNING OUTCOMES

- Describe the principal AI-related roles and the skills each requires
- Explain how generative AI is changing your own profession specifically
- Identify a credible progression route and the certification that supports it
- Build a portfolio artefact evidencing your capability to someone who was not here
- Submit a complete capstone draft for feedback

KEY TAKEAWAYS

- Prompt engineering and AI ethics are the two roles most accessible from a non-technical background.
- Across every profession the pattern is the same: production gets cheap, judgement gets expensive.
- AI-900 retired in June 2026. AWS AI Practitioner is the most accessible first certification for this audience.
- Six or more months of sustained practice measurably improves success rates (Anthropic Economic Index, March 2026).
- Evidence beats assertion. Five artefacts and a one-page summary are worth more than a certificate alone.

LAB · 55 MIN

Lab 15: Build your portfolio one-pager

Output: One-page portfolio summary validated by a peer cold read

WORKSHOP · 70 MIN

Workshop: Capstone draft clinic and presentation rehearsal

Draft capstone submitted; rehearsal feedback captured in Workbook W15.5.

ASSESSMENT

Formative: draft capstone submitted and marked against the full rubric, with written feedback returned within 72 hours.

DIRECTED STUDY

- Revise your capstone against the feedback received (120 min)
- Rehearse and time your ten-minute presentation (30 min)

SOURCE MATERIAL

- C5 L2.1 Career Opportunities in Generative AI
- C5 L2.2 Enhancing Your Career with Generative AI
- C5 L2.3-2.5 GenAI for Content Creators, IT Professionals, Leaders and Managers

18

Capstone Presentations, Assessment & Next Steps

Present, defend, be assessed, and leave with a plan

LEARNING OUTCOMES

- Present a complete capstone: problem, approach, evidence, risks and recommendation
- Defend your recommendation against structured adversarial challenge
- Assess a peer's work against a published rubric and give useful feedback
- Complete the final written assessment
- Leave with a written 90-day plan and an identified progression route

KEY TAKEAWAYS

- You can define a problem, choose an approach, assess the risk and defend a recommendation.
- Evidence beats assertion. You now have six artefacts.
- The tools will change every six months. The disciplines will not.
- Verify before you trust. Define what done looks like. Never put confidential data in a prompt.
- Sustained practice is what separates people who used a course from people who changed how they work.

LAB · 55 MIN

Capstone presentations, ten minutes plus five minutes of challenge

Output: Delivered and defended capstone presentation, assessed against the published rubric

WORKSHOP · 70 MIN

Final written assessment, peer assessment and 90-day planning

Completed final assessment; peer assessment returned; written 90-day plan; certificate presented.

ASSESSMENT

SUMMATIVE: Capstone (100 marks) plus Final Written Assessment (60 marks). Programme pass requires 15/18 attendance, all block checkpoints attempted, and 60% or above in both capstone and final assessment. Distinction at 80% or above combined.

DIRECTED STUDY

- Execute your 90-day plan
- Book your next certification
- Return for a quarterly alumni session

SOURCE MATERIAL

- Synthesis across all five IBM courses plus supplementary material

ASSESSMENT ACROSS THE PROGRAMME

ASSESSMENT	WEEK	MARKS	FORMAT
Block A checkpoint	Week 4	50 marks	MCQ, short answer, applied scenario. 30 minutes, closed book.
Block B checkpoint	Week 8	50 marks	Practical prompt rebuild, prompt library review, one written explanation. 45 minutes.
Block C checkpoint	Week 10	Applied	Copilot task log plus a one-page rollout and governance plan.
Block D checkpoint	Week 13	50 marks	Acceptable use policy drafted and defended under challenge. 70 minutes.
Capstone project	Week 18	100 marks	Real problem taken end to end, presented in 10 minutes and defended for 5.
Final assessment	Week 18	60 marks	MCQ, short answer and applied scenario across the whole programme. 60 minutes.

To pass. Attendance at 15 of 18 sessions · all block checkpoints attempted · capstone at 60% or above · final assessment at 60% or above. Distinction at 80% or above across capstone and final combined.

THE DOCUMENT SET

	DOCUMENT	FORMAT	PURPOSE
01	Programme Handbook	DOCX · 34 pp	Specification, assessment framework, delivery guidance, quality assurance
02	Trainer Lesson Plans	DOCX · 52 pp	Sixteen minute-by-minute five-hour plans with facilitation notes
03	Learner Workbook	DOCX · 81 pp	Every exercise, lab template, glossary and the capstone workspace
04	Assessment Bank	DOCX · 21 pp	Checkpoint papers, final assessment, model answers, trainer only
05	Onboarding Pack	DOCX · 14 pp	Enrolment, skills audit, goals, learning agreement, acceptable use, IT readiness
06	Feedback Forms	DOCX · 10 pp	Exit ticket, mid-point review, end-of-programme evaluation, observation form
07	Candidate Progress Tracker	XLSX · 11 sheets	Attendance, assessment, portfolio, skills audit, dashboard, at-risk register
08	Certificate & Programme Record	DOCX · 5 pp	Pass and distinction certificates plus the full programme record
09	Competency Mapping & Evidence Pack	DOCX · 18 pp	Skills England, EU AI Act Article 4, Turing framework, Skills Bootcamps and Ofqual, with a verification log
10	Employer Design Panel	DOCX · 14 pp	Half-day panel pack with a signed outcome record evidencing employer co-design
11	Skills England Enquiry	DOCX · 11 pp	Ready-to-send benchmark enquiry, response log, decision tree and FOI fallback
-	Week 1-16 Presentation Decks	16 × PPTX	188 branded slides with complete speaker notes

BEFORE EVERY COHORT

This programme has perishable content. Two weeks before a cohort starts, verify and update the model landscape table in Week 4, the regulatory dates table in Week 10, every free-tier tool used in a lab, and the certification costs in Week 15. Free tiers change monthly, test each one on

the venue network in the week before the session. Where a fact carries a date, teach the date alongside the fact.

Flowline Robotics · AI & Generative AI Practitioner Programme · Version 1.0 · August 2026 · Review due February 2027
Conceptual content derives in part from the IBM Skills Network Generative AI Fundamentals specialisation. Statistics and regulatory positions are attributed at the point of use and were verified in August 2026. Course design, assessment instruments, workshop formats, forms and tracking tools are original to Flowline Robotics.