



This course is part of the **AI Literacy** series and is designed to help meet the **AI literacy requirements** set out in **Article 4** of the **EU AI Act**, which **entered into force** on **2 February 2025**.

By building the skills, knowledge, and understanding to interact effectively and responsibly with AI systems, you will not only improve your day-to-day results but also strengthen your organisation's compliance posture. The courses focus on practical, plain-language learning that benefits any role, from operational staff to senior leadership, while ensuring you can oversee, apply, and guide AI use with competence, compliance, and confidence.



MODULE 1: ETHICAL FRAMING FOR PROMPTING

Learning Outcome: Design high-impact prompts for public and enterprise/private GPTs that deliver accurate, efficient results in regulated workflows.

1.1 Setting the Stage

- 1.1.1 Connect prompting to governance, compliance, and accountability.
- 1.1.2 Recognise that the goal is defensible, auditable, and policy-aligned outputs, not just accuracy.
- 1.1.3 Understand why frameworks like the EU AI Act's Article 4, ISO/IEC AI Management, and NIST AI RMF treat ethical prompting as a core workplace skill.

1.2 Why Ethics Shapes Prompting in Practice

- 1.2.1 Understand why prompting is not just about accuracy, but about producing responsible, defensible outputs.
- 1.2.2 Learn how global AI regulations, including the EU AI Act's Article 4, make ethical prompting a core workplace skill.

1.3 Ethical Risk Zones in Prompting

- 1.3.1 Identify the three main risk categories: content, process, and impact.
- 1.3.2 Recognise where unverified outputs, bias, and governance failures can occur.

1.4 Aligning Prompting with Governance Models

- 1.4.1 Connect prompting practices to organisational AI policies and recognised governance standards.
- 1.4.2 Understand how ethical prompting supports compliance with frameworks like ISO/IEC AI Management and NIST AI RMF.

1.5 From Prompting Skills to Ethical Prompt Fluency

- 1.5.1 See how earlier prompting skills evolve into transparent, auditable, and policy-aligned techniques.
- 1.5.2 Understand the value of prompts designed for accountability and audit-readiness.





MODULE 2: GUARDRAILS, GOVERNANCE & RISK

Learning outcome: Use GPT safely, ethically, and legally across deployment types.

2.0 Deployment spectrum (sets the frame for the whole module)

2.1 Responsible inputs across GPT deployments: aligning with risk and governance requirements

- 2.1.1 Company-hosted/private GPT (lowest direct privacy risk)
- 2.1.2 Enterprise ChatGPT / Azure OpenAI (lower risk, not zero)
- 2.1.3 Public ChatGPT: UI controls & “quiet defaults”

2.2 Prompt logs, audit trails & legal discoverability for audit-ready workflows

- 2.2.1 Prompt logging & documentation
- 2.2.2 Can your prompt logs be subpoenaed?
- 2.2.3 Prompt as if it will be reviewed later
- 2.2.4 Redaction & minimization patterns

2.3 What happens to your data under different GPT deployments

- 2.3.1 Company-hosted/private GPT
- 2.3.2 Enterprise GPT
- 2.3.3 Public ChatGPT & OpenAI privacy policy
- 2.3.4 Consent & control — human reality vs. legal text

2.4 Memory & retrieval risks in ChatGPT and enterprise RAG systems

- 2.4.1 Enterprise/private “memory”: RAG & vector stores
- 2.4.2 ChatGPT Memory: use it, control it, trust it?
- 2.4.3 Trust boundaries & version drift

2.5 Human factors: anthropomorphizing, overtrust & misuse

- 2.5.1 Not your colleague: tone ≠ trust
- 2.5.2 Reinforcing tone ≠ agreement or accuracy
- 2.5.3 Keep it strictly professional
- 2.5.4 Regulated decision boundaries

2.6 Designing your AI use playbook for consistent, compliant adoption

- 2.6.1 Core components
- 2.6.2 Ready-to-use templates
- 2.6.3 Mini-playbooks by sector for faster, safer adoption
- 2.6.4 “Just risky enough” without reckless





MODULE 3: AI WORKFLOWS IN REGULATED ROLES

Learning outcome: Putting GPT to work in high-stakes roles, fast, useful, and under control.



3.1 GPT in professional workflows

3.1.1 Identify where GPT delivers measurable value in regulated and operationally critical roles

3.1.2 Understand the workflow stages where GPT can accelerate work while retaining human judgment

3.2 GPT for client-facing communications

3.2.1 Draft professional and compliant external communications

3.2.2 Guide GPT to meet accuracy, tone, and compliance standards

3.2.3 Recognise liability points requiring human review

3.3 Internal training & policy onboarding

3.3.1 Translate complex policies for onboarding and training

3.3.2 Maintain accountability with human oversight

3.4 Document summaries & risk narratives

3.4.1 Convert source content into decision-support summaries or narratives

3.4.2 Control scope, tone, and factual reliability

MODULE 4: CHATGPT IN THE ENTERPRISE CONTEXT: TOOLS, INTEGRATIONS & RISK POINTS

Learning Outcome: This module uses ChatGPT, the most widely adopted public GPT, as the reference point. The principles and governance considerations apply equally when evaluating enterprise/private GPT features, making this essential knowledge for both environments.

4.0 Why ChatGPT still matters in the enterprise conversation

4.0.1 ChatGPT as the most widely used public GPT

4.0.2 Feature pipeline and enterprise adoption

4.0.3 Cross-platform fluency

4.1 ChatGPT tools panel

4.1.1 Overview of built-in tools

4.1.2 Data handling and risk exposure

4.1.3 Enterprise equivalents and governance

4.2 File & image use

4.2.1 Secure, purposeful use

4.2.2 Common mistakes

4.2.3 Enterprise parallel





4.3 Exploring the GPT world

- 4.3.1 Custom GPTs in the public library
- 4.3.2 Reliability assessment checklist
- 4.3.3 Enterprise equivalent: internal GPT catalogues

4.4 ChatGPT vs Microsoft Copilot vs browser extensions

- 4.4.1 Functional differences
- 4.4.2 Data handling & governance changes
- 4.4.3 Decision criteria

4.5 Connecting to external platforms

- 4.5.1 Meaning of “connect”
- 4.5.2 Accessible/storable/shared data
- 4.5.3 Enterprise vs public connectors
- 4.5.4 Role-based decision points

4.6 Agent Mode

- 4.6.1 How it works in ChatGPT
- 4.6.2 Real-world risks
- 4.6.3 Privacy/legal concerns
- 4.6.4 Enterprise parallels

4.7 Applying ChatGPT learning to regulated workflows

- 4.7.1 Mapping features to equivalents/governance
- 4.7.2 Safe prototyping & training uses
- 4.7.3 Pre-enablement checklist

(NEW) MODULE 5: GPT-5 OVERSIGHT & SAFETY PATTERNS

Learning Outcome:

Apply GPT-5's enhanced controls to enforce ethical safeguards, limit AI autonomy in high-risk contexts, and maintain transparency with context-only answering and explicit refusal rules.

5.1 Autonomy Boundaries with the “Eagerness” Dial

- 5.1.1 What eagerness settings do in GPT-5
- 5.1.2 Using low vs high autonomy ethically
- 5.1.3 Setting stop conditions for high-risk tools

5.2 Tool Preambles for Auditability

- 5.2.1 How preambles make AI reasoning visible
- 5.2.2 Using preambles to support compliance reviews
- 5.2.3 Example preamble for safe summarisation





5.3 Context-Only Answering & Explicit Refusals

- 5.3.1 Restricting answers to approved context
- 5.3.2 Writing clear refusal instructions
- 5.3.3 Why refusals protect accuracy and trust

5.4 Prompt Hygiene for Sensitive Outputs

- 5.4.1 Avoiding conflicting ethical instructions
- 5.4.2 Using checklists to pre-screen AI outputs
- 5.4.3 Applying hygiene steps to compliance tasks

