

AI LITERACY: SAFE, PRODUCTIVE AND RESPONSIBLE USE

4-Hour Applied Workshop | Complete Training Kit for Individuals and Organizations



INDIGONIX SYSTEM INTELLIGENCE

Facilitator Guide + Participant Workbook + Templates + Assessment | Version 1.0 — July 2026

1. Program Positioning and Value Proposition

This program does more than teach tool usage. It enables participants to decide when, why and with what data to use AI—and what verification and human oversight are required.

Item	Individual Track	Enterprise Track
Audience	Professionals, entrepreneurs, students, educators and everyday users	Employees, managers, specialists, transformation teams, public-sector/university/NGO staff
Core promise	Think, create and verify better with AI	Make AI use safe, productive, traceable and policy-aligned
Recommended cohort	12-30 participants	15-40 participants
Format	In-person or live online	Preferably in-person; adaptable to live online
Output	Personal AI Use Plan + Prompt Portfolio	Department AI Use Map + Safe-Use Commitment + 30-Day Action Plan

1.1 Learning Outcomes

- Explain AI, generative AI, large language models, tokens, prompts, hallucination and bias accurately.
- Assess whether a task is suitable for AI based on risk and human oversight.
- Write strong prompts using Role + Context + Task + Data + Constraints + Format + Quality Criteria.
- Distinguish personal, sensitive and confidential organizational data and apply data minimization.
- Verify AI outputs for sources, logic, freshness, bias and fitness for purpose.
- Build responsible workflows for text, image, presentation, summary, ideation and code generation.
- Create a 30-day personal or organizational AI adoption plan.

1.2 EU AI Act Alignment

Article 4 of the EU AI Act requires providers and deployers to take measures to ensure a sufficient level of AI literacy among staff and other persons operating AI systems on their behalf, considering technical knowledge, experience, education, training and the context of use. AI literacy provisions have applied since 2 February 2025. This workshop can serve as one piece of evidence in a role- and risk-based literacy program; it is not, by itself, a guarantee of legal compliance.

2. Four-Hour Workshop Agenda

Time	Module	Focus	Method	Tangible Output
00-15	Opening and AI Pulse Check	Expectations, usage level, safety perception	Poll + pair share	Personal/team goal
15-50	M1 — How AI Works and Where It Fails	AI/GenAI/LLM, probabilistic generation, hallucination, bias	Briefing + live demo + mini game	Risk awareness card
50-80	M2 — Safe and Responsible Use	Data classes, privacy, copyright, human oversight, shadow AI	Scenario cards + group decision	Use / do-not-use decision matrix
80-90	Break	—	—	—
90-140	M3 — Prompt Design and Productivity	7-part prompt, iteration, multimodal creation	Demo + individual + pair practice	3 optimized prompts + quality rubric
140-170	M4 — Research and Verification with AI	Source hierarchy, fact-checking, freshness, multi-source checks	Fact-check lab	Verified mini brief
170-180	Break	—	—	—
180-225	M5 — Real-World Application Sprint	Individual or departmental case	Team sprint + presentation	End-to-end AI workflow
225-240	Close, Post-Test and 30-Day Plan	Recap, commitment, transfer	Quiz + action board	30-day plan + certificate criteria

Recommended interactive dose: 45-55%. Every 10-15 minutes includes a decision, trial, discussion or short creation task.

3. Module 1 — How AI Works and Where It Fails

Objective

Position AI correctly as a probabilistic content-generation system rather than an all-knowing authority.

Learning Content

- Differences among AI, machine learning, generative AI and LLMs
- Input → tokenization → context/attention → probabilistic output
- Hallucination, bias, freshness limits and overconfidence
- Text, image, audio, video, code and summary generation

Activity

“Next token” game: The facilitator gives a half sentence; participants suggest continuations, then compare with AI. The aim is to make probabilistic generation visible.

Facilitator Note

Live demo 1: Run the same prompt twice and mark differences. Live demo 2: Ask for sources on a non-existent report and verify whether the references are real.

4. Module 2 — Safe and Responsible Use

Objective

Enable participants to decide what data can be shared, where human approval is required and when AI should not be used.

Learning Content

- Data classes: public, internal, confidential, personal/sensitive
- Data minimization, anonymization, masking and approved tools
- Copyright and attribution; do not present AI output as unreviewed human work
- Bias, discrimination, accessibility and inclusion
- Human oversight: draft, review, approval and record

Activity

Scenario cards: (1) CV screening, (2) content creation with customer data, (3) pasting company code into a chatbot, (4) health advice, (5) board report. Groups decide “Allowed / Conditional / Prohibited / Expert Approval”.

Facilitator Note

Facilitator emphasis: “Classify the data before the prompt; classify the risk before the output; identify the human before use.”

5. Module 3 — Prompt Design and Productivity

Objective

Enable participants to treat a prompt not as a one-shot command but as an iterative design governed by quality criteria.

Learning Content

- 7 components: Role + Context + Task + Data + Constraints + Format + Quality Criteria
- Examples, counterexamples, option generation and critique loops
- Output formats: table, checklist, executive summary, slide flow, JSON
- Multimodal use: image description, document summary, presentation outline, narration script
- Second-pass quality prompt: “Critique and improve this answer against the rubric.”

Activity

Prompt clinic: Each participant writes a weak prompt from their work, improves it using the 7-part framework and produces a second version based on peer rubric feedback.

Facilitator Note

Individual examples: weekly planning, CV, learning plan, travel/spend comparison. Enterprise examples: report summary, risk analysis, customer communication, process improvement, technical documentation.

6. Module 4 — Research and Verification with AI

Objective

Enable participants to treat AI responses as hypotheses or drafts requiring verification—not final truth.

Learning Content

- Source hierarchy: law/official body, academic publication, standard, reputable news, industry report, blog/social media
- SIFT approach: Stop, Investigate the source, Find better coverage, Trace claims to the original
- Check date, version, author, institution and methodology
- Numerical verification: calculator/code/second method
- Contradiction log and confidence level

Activity

Fact-check lab: Groups receive a short AI response, separate its key claims, search sources, label each “Verified / Partial / Unverified / False” and produce a reliable brief.

Facilitator Note

Rule: Asking for a source link is not verification. Opening it, reading the context and matching it to the claim is verification.

7. Module 5 — Real-World Application Sprint

Objective

Apply the workshop learning end-to-end to a real personal or organizational case.

Learning Content

- Define the problem and success criterion
- Check AI suitability and risk
- Classify and, where needed, anonymize data
- Design, run, critique and improve the prompt
- Verify the output and identify human approval
- Share the result in a 2-minute presentation

Activity

Individual cases: learning plan, career development, content creation, budget/purchase comparison, personal project plan. Enterprise cases: meeting notes to action plan, complaint analysis, policy summary, process improvement, risk-and-control table.

Facilitator Note

Presentation rubric: Business value 25%, prompt design 20%, safety 20%, verification 20%, human oversight and feasibility 15%.

8. Prompt Design Framework and Sample Library

Component	Question	Example
Role	What expert perspective should AI use?	Senior process analyst
Context	What is the situation, audience and purpose?	Monthly performance meeting for the leadership team
Task	What exactly should it produce?	Extract decisions and actions from notes
Data	What safe data will be used?	Anonymized meeting notes
Constraints	What must it avoid; what are the limits?	Do not invent facts; flag ambiguity
Format	How should the output be structured?	Decision / Owner / Due date / Risk table
Quality Criteria	What defines a good answer?	Complete, traceable, no more than 300 words
Use Case		Ready-to-Use Prompt
Summarization		Summarize the following text under Decisions, Risks, Open Questions and Actions. Do not add information absent from the text. For each action, write “not specified” where owner or due date is missing.
Comparison		Compare X and Y by cost, benefit, risk, feasibility and strength of evidence. Show assumptions separately and leave the final decision to the human decision-maker.
Content Creation		Audience [audience], objective [objective], channel [channel]. Create three content options. Use a professional, warm tone. Avoid fabricated claims, unverified statistics and long copyrighted quotations.
Research		First refine the research question. Then recommend only verifiable source types. For each claim, provide source date and confidence level. Do not pretend to have read a source you cannot access.
Quality Review		Score the following AI output from 1-5 for accuracy, completeness, source quality, bias, privacy, copyright and usability. Provide the three most critical issues and a corrected version.
Meeting to Action		Analyze these anonymized meeting notes. Extract only explicitly stated decisions and actions. Output a table with Decision Action Owner Due Date Dependency Evidence Sentence.

9. Participant Worksheets

9.1 AI Use Decision Card

Question	If Yes	If No
Does the task affect health, law, finance, hiring, credit, security or fundamental rights?	Expert/organizational approval and stronger verification are required.	Continue to the next question.
Does it contain personal, sensitive, confidential or trade-secret data?	Do not use raw data; check authorization, anonymization and approved tools.	Continue to the next question.
Will the output be published externally or used for a decision?	Source, copyright, bias, accuracy and human approval checks are mandatory.	May be used as a low-risk draft; still review it.

9.2 Prompt Canvas

Field	Complete
Role	
Context	
Task	
Safe data	
Constraints	
Format	
Quality criteria	
Human approval	

9.3 AI Output Verification Checklist

- I separated the key claims.
- I checked date and version.
- I opened the sources and read the context.
- I cross-checked with at least one independent source.
- I checked numbers using a second method.
- I reviewed for bias, missing perspectives and inclusion.
- I checked privacy/confidentiality and copyright risks.
- I obtained final human approval or identified the accountable person.

9.4 30-Day Action Plan

Week	Goal	Action	Success Indicator	Risk/Control
1				
2				
3				
4				

10. Pre-Test / Post-Test and Certificate

No.	Question	Expected Answer
1	What is the core generation mechanism of an LLM?	Generating likely next tokens based on context
2	What is hallucination?	AI confidently presenting false or non-existent information
3	What is the first step before using AI with text containing personal data?	Check authorization, data minimization and anonymization needs
4	Name three of the seven prompt components.	Any three of Role, Context, Task, Data, Constraints, Format, Quality Criteria
5	Why is asking for a source link not sufficient fact-checking?	The source must be opened, read in context and matched to the claim
6	What should AI's role be in critical decisions?	Supportive draft/analysis; final decision remains with an authorized human
7	How should a numerical AI output be verified?	With a calculator, code or an independent second method
8	What is shadow AI?	Use of AI tools for work outside organizational approval
9	What three checks should be performed before publishing AI-generated content?	Any three of accuracy/source, copyright, bias/privacy and human approval
10	What is the main subject of Article 4 of the EU AI Act?	Measures by providers and deployers to ensure sufficient AI literacy for relevant persons

Assessment

- Pre-test: 10 questions, no score pressure; used to identify baseline.
- Post-test: 10 parallel questions measuring the same outcomes; recommended pass mark 80%.
- Recommended certificate conditions: at least 80% attendance, 80% post-test and completion of the 30-day plan.

11. Facilitator Guide

11.1 Preparation Checklist

- For the enterprise track, confirm approved AI tools, data policy and cases with the organization in advance.
- Log in to demo accounts; never use confidential or real personal data.
- Test projection, internet, power, QR codes and backup demo screenshots.
- Prepare participant workbook, scenario cards, prompt canvas and rubric.
- Prepare pre-test, post-test, survey and certificate workflow.
- Plan an interaction point every 10-15 minutes; avoid long theory blocks.

11.2 Concise Answers to Difficult Questions

Question	Suggested Response
Will AI take my job?	It will transform parts of jobs. The strongest position belongs to people combining domain expertise, AI literacy and human judgment.
Which tool is best?	There is no single best tool. Selection depends on task, data class, security, cost, integration and verification needs.
Should we never give data to AI?	Public and authorized data may be used. Personal, sensitive, confidential or contract-protected data requires policy and approval checks.
Who owns AI output?	Copyright and usage rights vary by tool terms, human contribution and jurisdiction. Organizational/legal review is needed before publication.

11.3 Post-Workshop Follow-Up

Timing	Action	Evidence / Output
Same day	Send resources, workbook and prompt set	Distribution record
7 days	Short usage survey and request for a success example	Survey result
30 days	Review action plan; in organizations, assess with team lead/AI owner	Completed plan + improvement areas
3 months	Optional impact measurement: time saved, quality, incidents, adoption rate	AI literacy impact note

12. Enterprise Customization Pack

- The organization's approved-tool list and data-classification scheme are integrated into the workshop.
- Cases are customized for management, HR, finance, legal, marketing, sales, R&D, software and operations.
- Outputs may be retained as an AI literacy evidence pack: attendance, test result, content, role-based plan, survey and actions.
- At the end, a department-level "Allowed / Conditional / Prohibited" use map is produced.
- As a next step, the workshop connects to an AI use policy, AI inventory, role-based training matrix and governance roadmap.

12.1 Deliverables

Deliverable	Individual	Enterprise
Facilitator slide flow and speaker notes	✓	✓
Participant workbook	✓	✓
Prompt canvas and library	✓	✓
AI use decision card	✓	✓
Pre-test / post-test / survey	✓	✓
30-day action plan	✓	✓
Department use map	—	✓
Role-based AI literacy evidence pack	—	✓
Starter organizational AI use policy template	—	✓

12.2 Indigonix and Workshop Contact

Indigonix System Intelligence

indigonix.ai

info@indigonix.ai

Dr. Damla Sivrioğlu Aslan — Founder & Chief Architect

Oya Demir, MBA — Head of Global Business Development

Workshop Booking

Contact the Indigonix team for tailored enterprise delivery, open-enrollment individual sessions, online or in-person formats, and Turkish or English facilitation.

13. References and Currency Note

- Regulation (EU) 2024/1689 (Artificial Intelligence Act), especially Article 4 — AI literacy.
- European Commission, "AI Literacy — Questions & Answers" and AI Act application timeline.
- For organizational policy, KVKK/GDPR, copyright and sector regulation, current review by legal/compliance teams should prevail.
- Tool capabilities, model names, retention settings and knowledge cutoffs change rapidly and should be updated from official product pages before delivery.