



Enghouse
Interactive



PRECHEL CONSULTING

VISION 20XX · COLOGNE · 24–25
JUNE 2026

DAY 1 · CERTIFICATION TRAINING

The EU AI Act

Foundations & Practical Compliance for Customer Service

Prof. Dr. Kerstin Prechel

Professor of Business Ethics & Digital Transformation · Prechel Consulting

Regulation (EU) 2024/1689 · legal status as of June 2026



PRECHEL CONSULTING

WELCOME

A warm welcome to Vision 20XX

Today is a working session, not a lecture.

We will turn 144 pages of EU regulation into a handful of decisions you can actually make about the AI in your contact centre – what you may use, what you must disclose, and what is now simply off the table.

You don't need a legal or technical background. We'll build the AI basics together first, then the law.



Built for customer service



**Certification-style: you leave
able to classify**



A compass, not a cage



Prof. Dr. Kerstin Prechel

Business ethics & AI · EU AI Act · responsible AI

WHY THIS MATTERS

AI is already running your contact centre

The question is no longer whether you use AI – it's whether you understand it well enough to use it responsibly.



Chatbots & voicebots

Front-line conversational AI handling routine contacts



Agent assist

Real-time suggestions, summaries, next-best-action



Speech & sentiment analytics

Analysing tone and content of calls



Live translation

Cross-lingual chat and call handling



Quality & workforce mgmt

Scoring, forecasting, scheduling



Generative knowledge

Drafting replies, summarising, knowledge search

ROADMAP

Four modules – from the basics to your action plan



1 · Foundations

What AI, ML, LLMs and GPTs really are – plus why the law exists and the AI-literacy duty



2 · The risk classes

The four tiers, the articles, and where your CX systems land



3 · The roles

Provider vs deployer – and the duties that fall on you



4 · Governance

Operationalising responsible AI and your 90-day plan

MODULE 1

Foundations

First the technology, then the law – so the rules actually make sense.

By the end you can...

- explain what AI, ML, LLMs and GPTs are
- say which AI is inside your CX tools
- explain why the EU regulates AI
- describe the AI-literacy duty (Art. 4)
- spot bias, black-box and hallucination risks

MODULE 1 · FOUNDATIONS

What the EU AI Act is

Regulation (EU) 2024/1689

The world's first comprehensive, horizontal law on artificial intelligence. In force since 1 August 2024; rolling out in phases to August 2026 and beyond.

Its purpose (Art. 1): a well-functioning single market with human-centric, trustworthy AI – while protecting health, safety and fundamental rights, and supporting innovation.

It is risk-based: the obligations scale with the potential to cause harm – not with the technology itself.



Horizontal

Applies across every sector – including customer service



Risk-based

Four tiers, from banned to minimal



Extraterritorial

Applies if the output affects people in the EU

MODULE 1 · FOUNDATIONS

AI literacy is already mandatory



Article 4 – in force since 2 February 2025. Providers and deployers must ensure staff who operate or use AI systems have a sufficient level of AI literacy.



Know the technology

What AI is, how ML models work, what your systems actually do – and where they fail



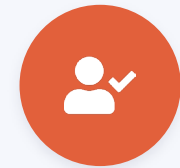
See chances & risks

Recognise error, discrimination and hallucination risks in your use cases



Know the rules & ethics

Grasp the AI Act, GDPR and the ethical principles that frame responsible use



Apply it in practice

Check AI outputs critically; keep a human in the loop where it matters

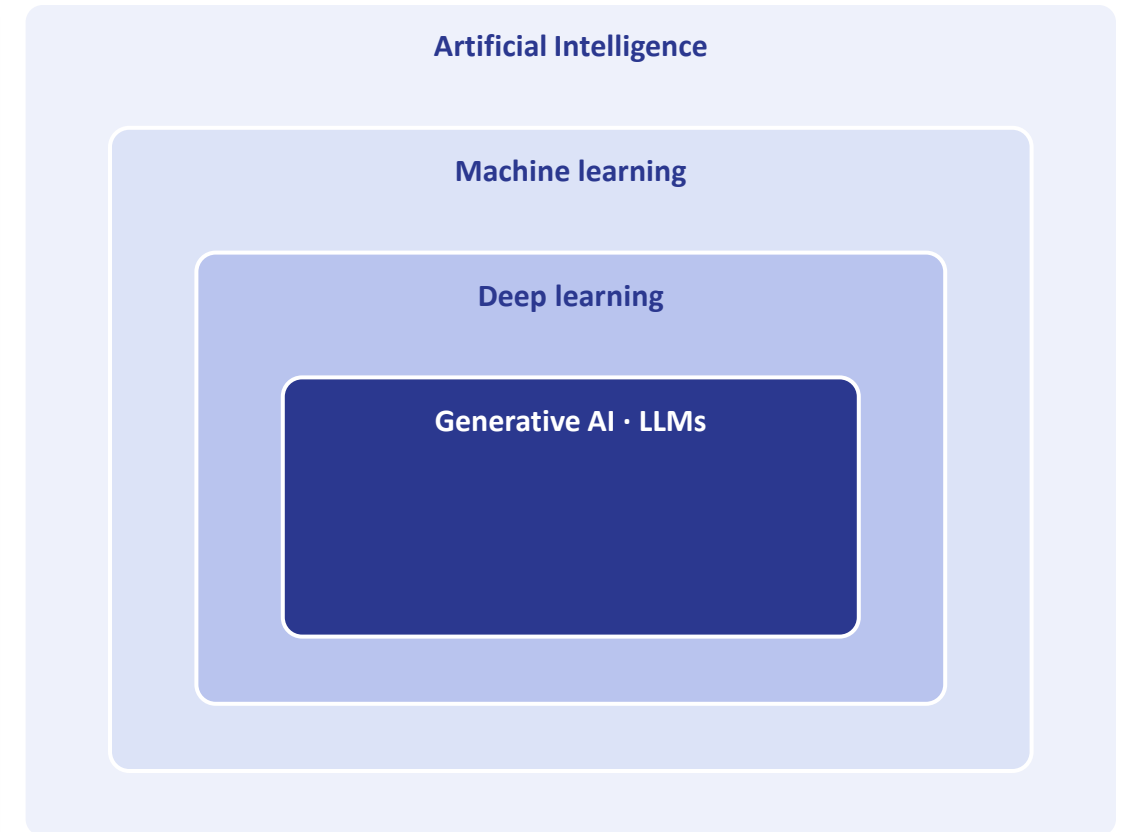
The next slides build dimension one – “know the technology.”

What is AI, really?

AI = systems that perform tasks we'd normally call “intelligent”

understanding language, recognising images, predicting, deciding. Today almost all of it is **machine learning**: systems that learn patterns from data instead of following rules a human wrote.

Narrow, not general: every AI in use today is good at one narrow task. The bot that books appointments cannot drive your car. “General” human-like AI does not exist yet.



Each sits inside the last.

MODULE 1 · AI LITERACY

How machines actually “learn”



Traditional software

A human writes the rules. “IF the balance is below zero, THEN flag the account.” Predictable, explainable – but it can't handle what it wasn't told.



Machine learning

You show it thousands of examples and it infers the pattern itself. Nobody writes the rule – the model discovers it. Powerful, but only as good and as fair as the data.

Two words you'll hear from vendors

Training = the model learns from data. Slow, expensive, done once by the provider.

Inference = the trained model answers a new input. Fast – this is every customer interaction.

What is an LLM? (Large Language Model)

A model trained on enormous amounts of text to do one deceptively simple thing: predict the next word.

From that single objective – guessing the most likely next “token” – emerges fluent writing, summarising, answering and drafting. Scale the data and the model, and the behaviour gets remarkably capable.

The catch: it predicts, it doesn't “know.” An LLM has no understanding and no built-in fact-check. That's why it can be confident and wrong – a hallucination.

In your contact centre

- Powers most modern chatbots & voicebots
- Drafts replies, emails and summaries
- Answers from your knowledge base
- Translates and rephrases in real time

If it writes like a person, it's probably an LLM.

MODULE 1 · AI LITERACY

GPT – decoded

“GPT” is just three words. Knowing them tells you exactly what the tool is doing.



G – Generative

It produces new content – text, and increasingly images, audio and code – rather than just classifying or scoring.



P – Pre-trained

It already learned from a huge body of text before you ever typed anything. You use a finished model.



T – Transformer

The neural-network architecture (from 2017) that made this work – it weighs which words matter to each other (“attention”).



So: a GPT is one kind of LLM. ChatGPT, Claude, Gemini and Copilot are all LLM-based assistants. But not everything labelled “AI” is a GPT – a spam filter or a forecast model is not.

Which AI is probably in your tools?

Most customer-facing tools in 2026 are LLM-powered – exactly the type that can hallucinate and needs oversight.

Chatbots & voicebots

→ LLM + intent understanding (NLU)

Agent assist · generative replies · summaries

→ LLM (generative)

Knowledge search / “ask the docs”

→ LLM + retrieval (RAG)

Speech analytics & transcription

→ Speech-to-text (ASR) + NLP

Sentiment analysis

→ NLP classifier or LLM

Routing, triage & ticket tagging

→ Classic ML classifier

Forecasting & scheduling (WFM)

→ Statistical / time-series ML

Voice output (the bot's voice)

→ Generative speech (TTS)

Neuronal Network – a form of AI

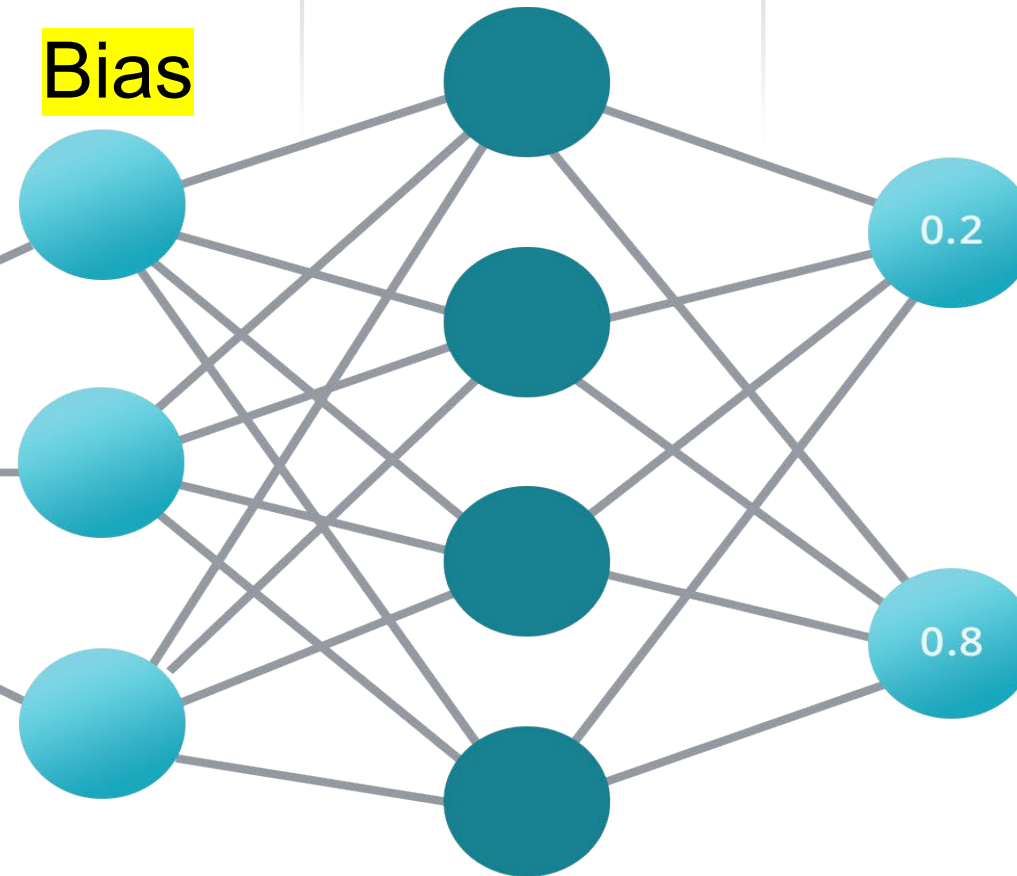
Black Box Problem

Input Layer

Hidden Layer

Output Layer

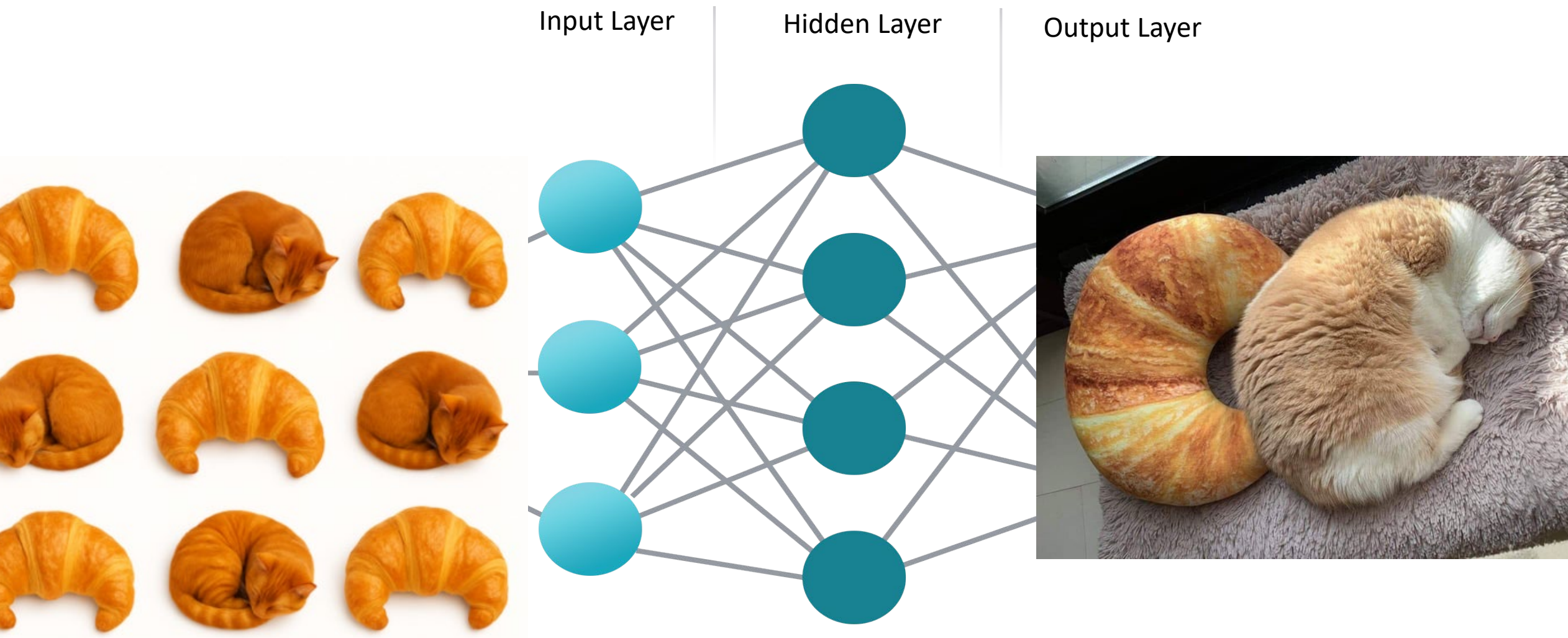
Bias



Judging the
output



Cat or Croissant?



MODULE 1 · AI LITERACY

A quick vocabulary

Token

A chunk of text (roughly a word-part) the model reads and predicts.

Parameters

The billions of tunable values that store what a model “learned.”

Training data

The examples a model learned from – its strengths and biases come from here.

Prompt

The instruction or question you give the model.

Inference

Running the trained model on a new input – every live interaction.

Hallucination

A fluent, confident output that is simply false.

Fine-tuning

Extra training to specialise a model for your domain or tone.

RAG

Retrieval-augmented generation: the model answers from your documents, reducing made-up answers.

Three things to watch with modern AI



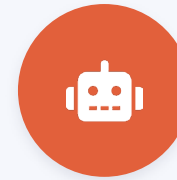
The black box

Neural networks learn patterns we can't fully read back. We see the output, not the reasoning – which is why explainability and oversight matter.



Bias

Models inherit the patterns in their data. Amazon scrapped a hiring tool that learned to down-rank CVs containing the word “women's” (Dastin, 2018).



Hallucination

Generative systems can produce fluent, confident answers that are simply wrong. In CX that becomes a wrong promise to a customer – in your name.



KNOWLEDGE CHECK · MODULE 1

Test yourself: the basics

1

Which statement about today's AI is correct?

- A It mostly follows rules written by programmers
- B It mostly learns patterns from data
- C We now have general, human-like AI

2

An LLM gives a customer a confident but false answer. This is called...

- A a routing error
- B a hallucination
- C a data breach

3

What does the “P” in GPT stand for?

- A Predictive
- B Pre-trained
- C Programmed

Module 1 – how did you do?

B

Q1 – It mostly learns patterns from data Modern AI is machine learning: it infers patterns from examples rather than following hand-written rules. General human-like AI does not exist.

B

Q2 – A hallucination An LLM predicts plausible text, not verified truth, so it can be fluent and wrong. That failure mode is called a hallucination.

B

Q3 – Pre-trained GPT = Generative, Pre-trained, Transformer. “Pre-trained” means it already learned from large text data before you use it.



REFLECTION

Which AI systems are live in your contact centre right now?

- List the ones you bought versus the ones you built
- For each: roughly which AI type is inside it?
- Which are LLM-based – and therefore can hallucinate?

MODULE 2

The Risk Classes

Four tiers. Know where each of your systems lands.

By the end you can...

- name the four risk classes and their articles
- give CX examples for each tier
- identify a prohibited practice
- recognise a high-risk CX use case
- read the implementation timeline

MODULE 2 · RISK CLASSES

A risk-based pyramid



Obligations increase as you move up the pyramid.

CX examples

Emotion AI scoring your agents

Agent quality scoring · WFM that allocates work

Customer chatbots · voicebots · deepfakes

Spam filters · skill-based routing

MODULE 2 · RISK CLASSES

Unacceptable risk – simply banned



Article 5 · prohibited since 2 Feb 2025

Practices judged to violate fundamental rights – no compliance route, no exceptions for business convenience.

- Social scoring of individuals
- Manipulative or deceptive techniques causing harm
- Exploiting age, disability or vulnerability
- Untargeted face-scraping for recognition databases
- **Emotion recognition in the workplace**



Why a contact centre must care

One Article 5 prohibition lands squarely in customer service:

inferring the emotions of your agents

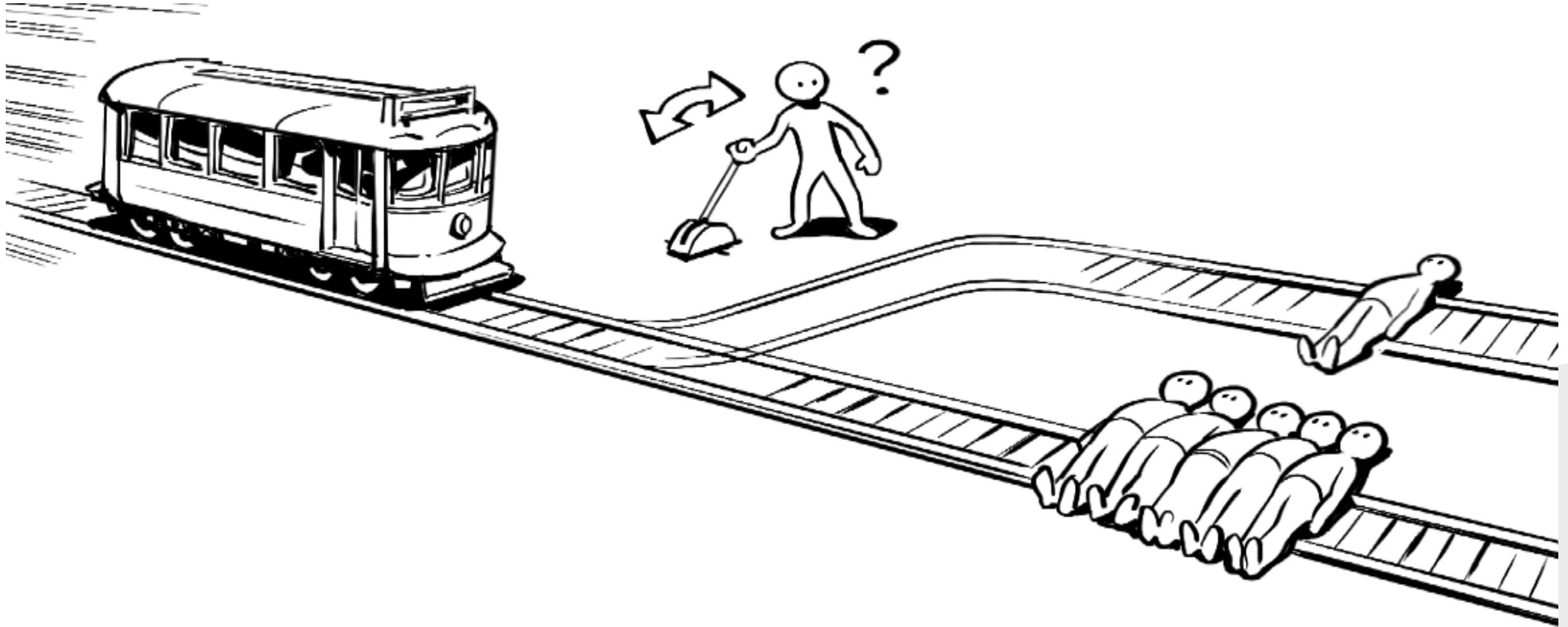
Any AI that reads an employee's emotional state from voice or face in the workplace is prohibited – and breaches here carry the highest fines (up to €35M / 7%).

The Ethics behind the Law

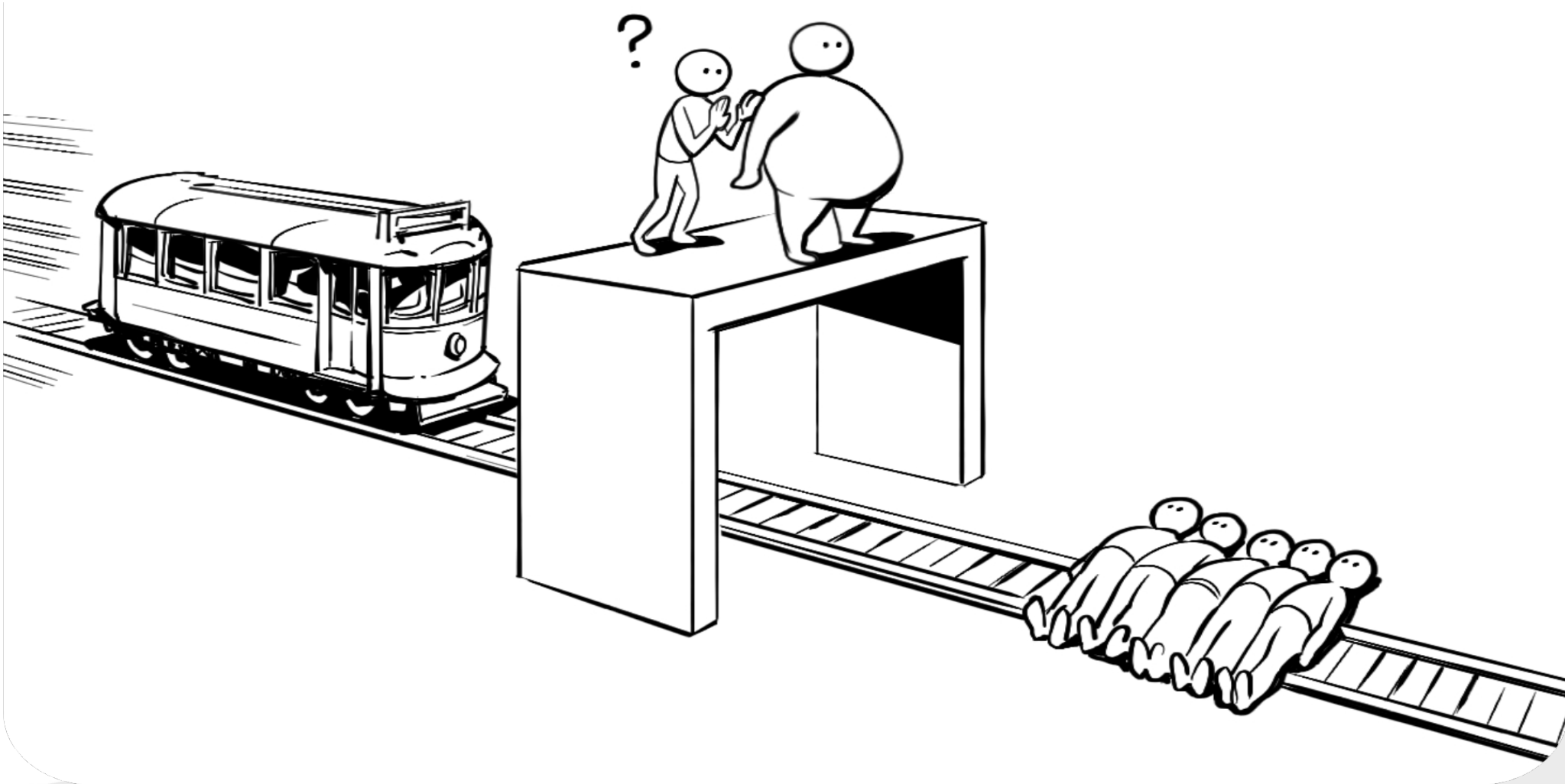
In most studies the ethics of AI are pretty clear. Worldwide scientists are in agreement.

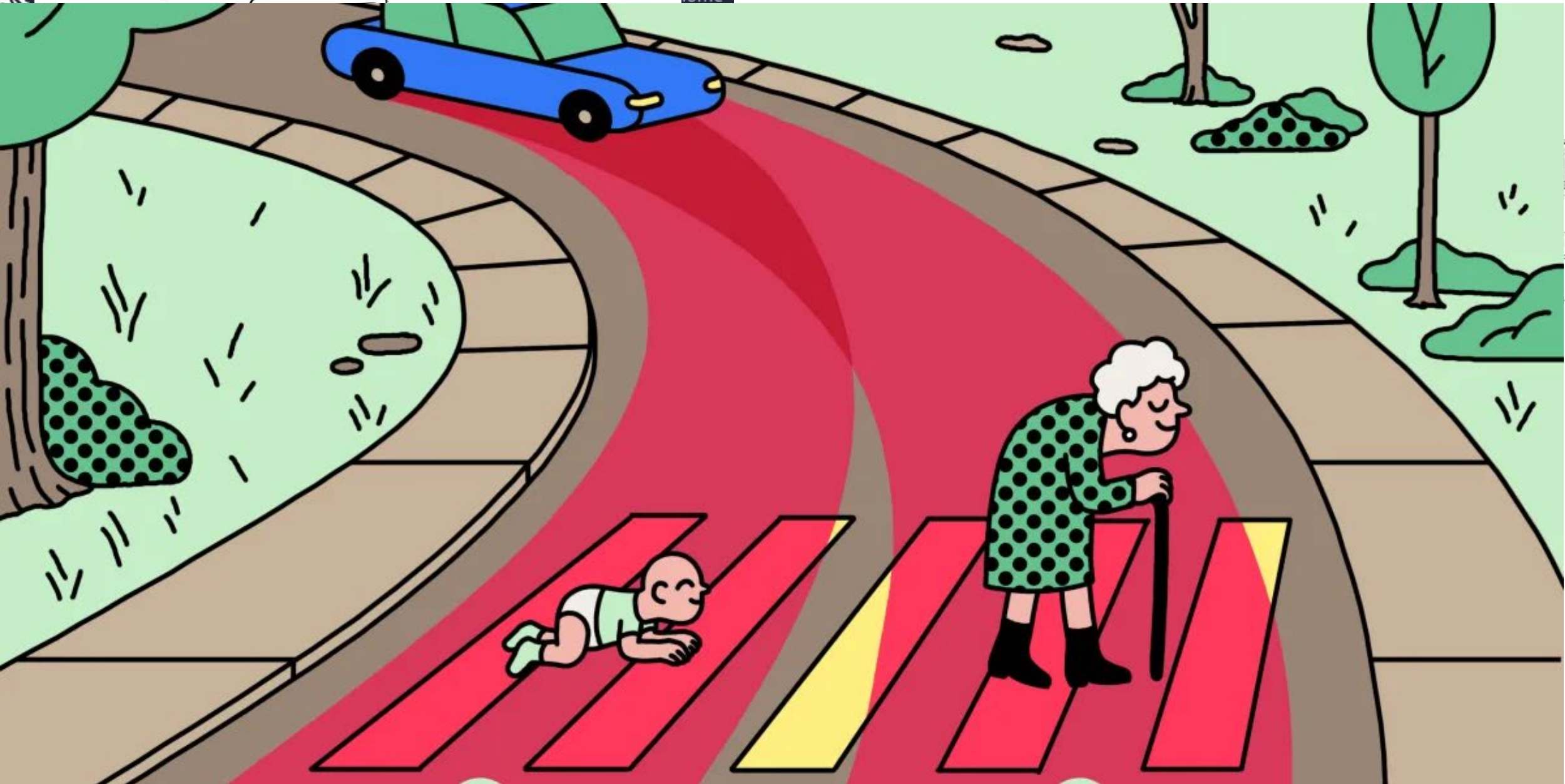


The Trolley Problem



Fat Man Problem





Emotion AI: agents vs. customers

Art. 5(1)(f) bans inferring emotions from biometric data in the workplace. Who the system reads decides everything.



PROHIBITED

Reading your agents' emotions

Voice- or face-based AI that scores an employee's mood, stress or engagement in the workplace.

Narrow exceptions only: genuine medical or safety reasons (read strictly).



NOT banned by Art. 5

Reading the customer's emotion

A call-centre voice system tracking caller sentiment is outside the Art. 5 ban (Commission Guidelines, 2025).

But still regulated: transparency (Art. 50(3)), likely high-risk, and GDPR Art. 9 biometric data.



Watch the overlap: many systems read both at once. If a tool scores caller sentiment but also infers the agent's emotion, the prohibition bites. Pure text sentiment is generally not “biometric” – but tread carefully if it profiles employees.

MODULE 2 · RISK CLASSES

High risk – allowed, but heavily governed

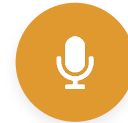
**Art. 6 + Annex III**

Permitted only with a full compliance package: risk management, data governance, technical documentation, logging, transparency, human oversight, accuracy & robustness – plus conformity assessment and EU-database registration.

Which CX systems can be high-risk?

**Worker management (Annex III §4)**

AI that monitors, scores or allocates work to agents, or feeds promotion / termination decisions

**Emotion recognition (Annex III §1)**

Permitted emotion systems (e.g. caller sentiment) sit here – not minimal risk

**Access to essential services**

If your AI helps decide eligibility for a service or benefit

MODULE 2 · THE TIER YOU'LL LIVE IN MOST

Limited risk – the duty to be transparent

Most customer-facing CX AI lands here. The duty is simple: don't let people be fooled. Applies from 2 August 2026 (Art. 50).

**50(1) · Chatbots & voicebots**

Tell people they're talking to a machine – unless it's obvious

**50(2) · Generative output**

Mark AI-generated content as machine-readable (grace period to 2 Dec 2026)

**50(3) · Emotion / biometrics**

Inform people exposed to emotion-recognition or biometric categorisation

**50(4) · Deepfakes & AI text**

Disclose synthetic images, audio, video and AI-generated public text

Minimal risk – the vast majority



No specific obligations under the Act

The largest group – and the safest. Think spam filters, skill-based call routing, ticket categorisation, forecasting.

Two caveats:

AI literacy (Art. 4) still applies; and voluntary codes of conduct are encouraged. Good practice here builds trust at no regulatory cost.

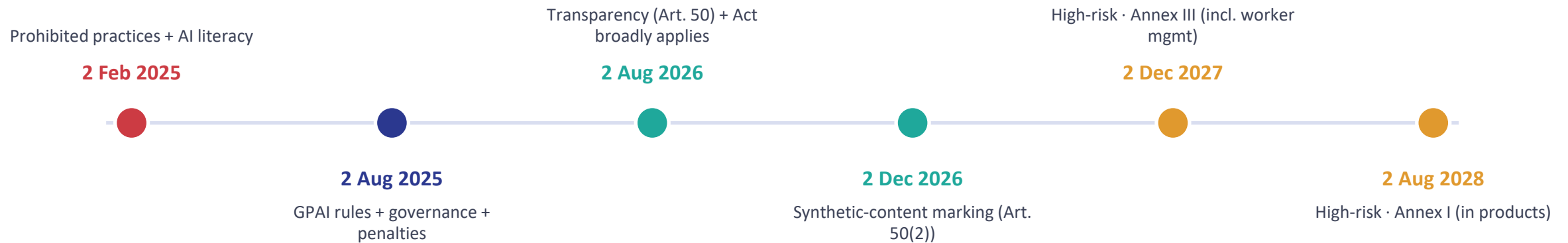
*“Minimal risk” is not
“no responsibility.”*

Classification can change with use. The moment a “simple” tool starts evaluating people or making decisions about them, re-check the tier.

MODULE 2 · RISK CLASSES

The timeline – what applies, and when

In force since 1 Aug 2024. The May 2026 “Digital Omnibus” moved the high-risk deadlines later – the duties already live haven’t moved.



Already in force: literacy, the prohibitions, GPAI and penalties. **Coming next:** transparency in Aug 2026. The high-risk deadlines slipped to 2027/28 under the Omnibus (politically agreed 7 May 2026; formal adoption pending) – but underlying GDPR and liability exposure exists **now**.

MODULE 2 · RISK CLASSES

The AI Act doesn't stand alone



GDPR

Call recordings, transcripts and voiceprints are personal data; emotion/voice biometrics are special-category data (Art. 9). Consent and purpose limits still apply.



Digital Services Act

Governs how platforms moderate and rank content – overlapping on algorithmic transparency and bias.



Product liability & sectoral law

When AI causes harm, liability rules apply regardless of the Act's deadlines.

For contact centres, GDPR is usually the law that bites first.

Who enforces the Act – country by country

Each member state names its own market-surveillance authority. Here's who's watching in six European markets.

	Germany Bundesnetzagentur (BNetzA) Lead market-surveillance authority; BaFin handles financial-sector AI.	IN PROGRESS
	Netherlands AP + RDI Data-protection authority & Digital-Infrastructure inspectorate co-ordinate; sectoral model, RDI is the contact point.	IN PROGRESS
	Denmark Agency for Digital Government Digitaliseringsstyrelsen – single point of contact; plus Datatilsynet and the courts administration.	DESIGNATED
	Italy ACN + AgID Cybersecurity Agency (ACN) for surveillance; AgID as notifying authority. Law 132/2025.	DESIGNATED
	Austria RTR – KI-Servicestelle The AI Service Centre at regulator RTR acts as national contact & co-ordination point.	PARTLY
	Spain AESIA Dedicated AI-supervision agency – the first purpose-built AI regulator in the EU.	DESIGNATED

Above all of them: the EU AI Office (general-purpose AI) and the EDPS (EU institutions). Most states also name sectoral and fundamental-rights authorities. Status: June 2026 – some designations are still pending formal adoption.

MODULE 2 · RISK CLASSES

What's at stake – penalties (Art. 99)

€35M

or 7% of global turnover

Prohibited practices (Art. 5)

€15M

or 3% of global turnover

High-risk & transparency breaches

€7.5M

or 1% of global turnover

Incorrect info to authorities

Whichever is higher (SMEs: whichever is lower). **Enforcement is live:** the EU AI Office and national authorities are already opening cases – and reputational damage often outweighs the fine.



KNOWLEDGE CHECK · MODULE 2

Test yourself: the risk classes

1

AI that scores your agents' call quality and feeds performance reviews is...

- A minimal risk
- B high risk (worker management)
- C prohibited

2

A chatbot answering “where's my order?” mainly must...

- A pass a conformity assessment
- B tell users it's an AI (Art. 50)
- C nothing – it's exempt

3

Inferring your agents' emotions from their voice is...

- A high risk but allowed with safeguards
- B prohibited (Art. 5)
- C limited risk

Module 2 – how did you do?

B

Q1 – High risk – worker management Scoring agents and feeding performance management is Annex III §4; full high-risk duties apply and a human must stay in control.

B

Q2 – Tell users it's an AI (Art. 50) A simple order-status bot makes no decisions about people: limited risk. The duty is transparency – disclose that it's a machine.

B

Q3 – Prohibited (Art. 5) Emotion recognition of employees in the workplace is banned outright; only narrow medical/safety exceptions, and the top penalty tier applies.



REFLECTION

Take your CX system list – can you now tier each one?

- Any candidates for the prohibited line (agent emotion AI)?
- Anything that scores or schedules your people (high-risk)?
- Which need an “you're talking to AI” disclosure by Aug 2026?

MODULE 3

The Roles

Provider or deployer? Your duties depend on the answer.

By the end you can...

- define provider and deployer (Art. 3)
- explain deployer duties for high-risk (Art. 26)
- split Art. 50 duties between the two roles
- locate your own organisation correctly

MODULE 3 · THE ROLES

Provider vs. deployer – where do you sit?



Provider (Art. 3(3))

Develops an AI system (or has it developed) and places it on the market or puts it into service under its own name or brand.

Typically: your software vendor – the company that builds your contact-centre platform.



Deployer (Art. 3(4))

Uses an AI system under its own authority in a professional capacity – without changing its core purpose.

Typically: you – the contact centre running the chatbot, the analytics, the assist tools.



You can become a provider without realising it. Substantially modify a system, put your own brand on it, or change its intended purpose – and provider duties attach to you.

MODULE 3 · THE ROLES

Your duties as a deployer of high-risk AI (Art. 26)

**Human oversight**

Assign trained, empowered people who can interpret and override the AI

**Use as instructed**

Operate the system in line with the provider's instructions for use

**Monitor & flag**

Watch performance; suspend use and inform the provider if risks emerge

**Keep the logs**

Retain system logs that are under your control

**Inform the workforce**

Tell workers & their reps before putting worker-affecting AI to use (Art. 26(7))

**Protect data & rights**

Run a data-protection impact assessment where required; inform affected people

Transparency: who does what (Art. 50)



The provider builds it in

- Design chatbots so users know they're AI (50(1))
- Mark generative output as machine-readable & detectable (50(2))



The deployer does the telling

- Inform people exposed to emotion-recognition or biometric systems (50(3))
- Disclose deepfakes and AI-generated public text (50(4))
- Make the notice clear, distinguishable and accessible

As a deployer you can't outsource compliance to your vendor – you carry your own duties.

MODULE 3 · WORKED SCENARIO

Scenario: “TalkBot” – a customer service chatbot



You buy a chatbot from your platform vendor to answer “Where's my order?” and “How do I reset my password?” It does not make decisions about people. What's the tier – and who does what?

**Tier**

Limited risk (Art. 50). No decisions about people – so not high-risk.

**Provider (vendor)**

Designs the bot to announce it's an AI; documents the system.

**Deployer (you)**

Use as instructed; keep the AI disclosure visible; train staff; handle escalation to a human.

MODULE 3 · WORKED SCENARIO

Scenario: “AgentCoach” – speech analytics on calls



A vendor offers a tool that analyses live calls. The same product can be configured two very different ways – and the configuration decides its legality.



Config A – over the line

Infers each agent's emotional state from their voice to score “engagement.” That's emotion recognition in the workplace – prohibited (Art. 5(1)(f)). Don't deploy it.



Config B – workable

Detects customer sentiment to route and prioritise, with the agent's audio excluded. Allowed – but inform callers (Art. 50(3)), handle voice data under GDPR, and treat as high-risk if it profiles people.



KNOWLEDGE CHECK · MODULE 3

Test yourself: the roles

1

Your vendor builds & sells the chatbot; you run it in your contact centre. You are the...

- A provider
- B deployer
- C neither

2

Your vendor says “we handle all AI Act compliance.” For you that is...

- A true – nothing left to do
- B false – you still owe deployer duties
- C true only for high-risk systems

3

Under Art. 50, who must tell the customer “you're chatting with AI” in practice?

- A only the provider
- B the deployer, using a system the provider designed to disclose
- C nobody – it's optional

Module 3 – how did you do?

B

Q1 – Deployer You use the system under your own authority in a professional capacity without changing its purpose (Art. 3(4)). The vendor is the provider.

B

Q2 – False – you still owe deployer duties Provider compliance covers building the system. You still owe oversight, use-as-instructed, monitoring, logging, worker information and Art. 50 disclosures.

B

Q3 – The deployer, on a system designed to disclose Provider engineers the disclosure capability (50(1)); the deployer operates it so users are actually informed. Shared system, separate duties.

MODULE 4

Governance & Responsible Deployment

From understanding the rules to running them every day.

By the end you can...

- list the steps to operationalise AI governance
- describe an AI governance team
- apply human oversight in CX
- build vendor due-diligence questions
- draft your 90-day plan

MODULE 4 · GOVERNANCE

Operationalising responsible AI – five steps



It's a loop, not a project with an end date – classification is revisited whenever a system or its use changes.

Keep a human in the loop – by design



Oversight is a design choice, not a disclaimer

- The human can understand the system's limits
- The human can interpret the output in context
- The human can decide not to follow it
- The human has the time and authority to intervene

In customer service that means...

- A clean, fast escalation path from bot to human
- Agents who can override AI suggestions without penalty
- No automated decision about a person without review
- Logging what the AI recommended and what the human did

The Tuesday test: if it doesn't make an average Tuesday calmer for customers and agents, it's theatre.

MODULE 4 · GOVERNANCE

Ten questions for any AI vendor

1 What risk tier does this fall in, and on what basis?

2 **Does it ever infer the emotions of our employees?**

3 How does it disclose to users that it's AI?

4 Is generative output marked as AI-generated?

5 What training data was used – and how is bias tested?

6 What can a human override, and how?

7 What logs do we receive and retain?

8 Where is data processed, and how is GDPR met?

9 What's your conformity-assessment & documentation status?

10 Who is liable if the system causes harm?

Take this list to your next procurement conversation – question 2 is the one most vendors won't expect.

MODULE 4 · GOVERNANCE

Your first 90 days

Days 1–30

See it

- Build your AI inventory
- Name an owner for AI governance
- Run AI-literacy basics for the team

Days 31–60

Sort it

- Tier every system
- Flag the prohibited line (agent emotion AI)
- Gap-check disclosures vs Aug 2026

Days 61–90

Run it

- Fix the urgent gaps
- Add vendor questions to procurement
- Set a monitoring & review rhythm



KNOWLEDGE CHECK · MODULE 4

Test yourself: governance

1

What's the single most valuable first step to start AI governance?

- A Buy more AI tools
- B Build an AI inventory
- C Wait for the high-risk deadline

2

Genuine human oversight (Art. 26) requires that the human can...

- A click “approve” quickly
- B understand, interpret and override the AI
- C retrain the model

3

AI governance is best understood as...

- A a one-off certification
- B a continuous loop you revisit
- C the vendor's job

Module 4 – how did you do?

B

Q1 – Build an AI inventory You can't classify, govern or defend what you can't see. The inventory of bought-and-built AI is the foundation of every later step.

B

Q2 – Understand, interpret and override the AI Oversight only counts if the human can grasp the system's limits, read context, and genuinely decline its output – not rubber-stamp it.

B

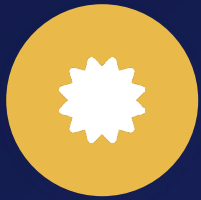
Q3 – A continuous loop you revisit Inventory → classify → gap → controls → monitor, then repeat. Classification changes whenever a system or its use changes.



REFLECTION

What is the one thing you'll do in your first week back?

- Who owns AI governance in your organisation today?
- Where is your biggest disclosure gap before August 2026?
- Is anything you run near the prohibited line?



CERTIFICATION

You can now...



Explain what AI, ML, LLMs and GPTs actually are



Place any CX system in the right risk tier



Recognise the prohibited line – agent emotion AI



Tell provider from deployer, and name your duties



Apply the Art. 50 transparency duties to your bots



Run an AI-governance loop and brief your vendors

WRAP-UP

Five things to carry out the door



Know what's inside the tool Most customer-facing CX AI is an LLM – fluent, fast, and able to be confidently wrong.



Classify by use, not by buzzword The same tech can be minimal or high-risk – impact on people decides.



Know the one hard line in CX Inferring your agents' emotions is prohibited. Full stop.



You're the deployer – own it Vendor compliance doesn't discharge your duties.



A compass, not a cage Done well, the rules make your CX more trustworthy – and that sells.

Thank you

Prof. Dr. Kerstin Prechel

kerstin.prechel@dhsh.de

www.prechel-consulting.de · www.dhsh.de

See you tomorrow for the keynote.



PRECHEL CONSULTING

REFERENCES

Sources & further reading

- Regulation (EU) 2024/1689 (EU AI Act) – esp. Arts. 1, 3, 4, 5, 6, 26, 50, 99; Annex III.
- European Commission, Guidelines on prohibited AI practices, C(2025) 884, 4 Feb 2025.
- European Commission, Draft Guidelines on transparency (Art. 50), 2026.
- Council of the EU & Parliament, Digital Omnibus on AI – political agreement, 7 May 2026.
- Vaswani et al. (2017). Attention Is All You Need (the Transformer). NeurIPS.
- <https://artificialintelligenceact.eu/assessment/eu-ai-act-compliance-checker/>
- European Commission, AI Act – Shaping Europe's digital future (digital-strategy.ec.europa.eu).
- Dastin, J. (2018). Amazon scraps secret AI recruiting tool that showed bias against women. Reuters.
- Bender, Gebru, McMillan-Major & Shmitchell (2021). On the Dangers of Stochastic Parrots. FAccT.
- Lütge, C. (2020). Ethik der Künstlichen Intelligenz. Springer.

Legal status reflects the position as of June 2026; the Digital Omnibus amendments were politically agreed but await formal adoption. This training is educational and not legal advice.