

Trustworthy AI: Auditing, Bias Detection, and Interpretability

COURSE OVERVIEW

This course explores how to build AI systems that are transparent, fair, and accountable. Participants will learn the principles and techniques for auditing AI models, detecting and mitigating bias, and improving model interpretability. By using real-world datasets and tools like SHAP, LIME, and fairness libraries, participants will gain practical experience in evaluating AI systems for ethical integrity and regulatory compliance. The course emphasizes responsible AI development to support trustworthy decision-making across sectors.

WHO SHOULD ATTEND?

This course is designed for data scientists, machine learning engineers, compliance officers, AI researchers, policy makers, risk managers, and anyone involved in developing or deploying AI solutions. It is especially relevant for professionals working in high-stakes environments such as finance, healthcare, human resources, and public services where AI outcomes must be explainable and fair.

COURSE OUTCOMES

Delegates will gain the knowledge and skills to:

- Learn the ethical and legal foundations of trustworthy AI.
- Conduct AI audits to evaluate fairness, accountability, and transparency.
- Detect and mitigate data and algorithmic bias.
- Use interpretability tools to explain AI predictions.
- Design AI systems aligned with regulatory frameworks and ethical guidelines.
- Communicate model decisions clearly to technical and non-technical audiences.
- Apply best practices for documenting and validating AI models.

KEY COURSE HIGHLIGHTS

At the end of the course, you will understand;

- The foundations of responsible and ethical AI.
- Types and sources of bias in datasets and models.
- Auditing frameworks and governance for AI systems.
- Hands-on tools: SHAP, LIME, Fairlearn, and AI Fairness 360.
- Techniques for improving model transparency and accountability.
- Regulatory considerations (e.g., GDPR, AI Act, CCPA).
- Case studies in finance, healthcare, HR, and law enforcement.
- Risk mitigation strategies and human-in-the-loop design.
- Communicating AI outcomes responsibly.
- Developing internal policies for ethical AI deployment.

All our courses are dual-certificate courses. At the end of the training, the delegates will receive two certificates.

1. A GTC end-of-course certificate
2. Continuing Professional Development (CPD) Certificate of completion with earned credits awarded