

Course No. 1779

When to Trust the Algorithm? Understanding Model Limitations Training Course

DATA SCIENCE & ADVANCED ANALYTICS
DATA, ANALYTICS & DIGITAL TRANSFORMATION

DURATION

5 Days

DELIVERY

Classroom



Course Overview

The When to Trust the Algorithm? Understanding Model Limitations Training Course is a practical, non-technical programme designed to help managers, executives, and decision-makers understand when algorithmic and artificial intelligence outputs can be trusted, when they require validation, and when human judgment must take precedence.

The course focuses on the limitations and risks behind analytical models and artificial intelligence systems, including data quality, bias, uncertainty, inaccurate predictions, changing environments, overfitting, lack of context, and misleading correlations. Participants will learn how these limitations can affect business decisions and organizational outcomes.

Rather than teaching programming or model development, the programme provides participants with a management-oriented framework for critically evaluating algorithmic outputs. It enables them to ask the right questions about data, assumptions, model performance, reliability, explainability, and applicability before relying on automated recommendations.

Through practical scenarios and case studies, participants will assess situations in which algorithms perform effectively, identify warning signs that indicate unreliable outputs, and develop appropriate human oversight, governance, and escalation mechanisms for responsible algorithm-supported decision-making.

Objectives

- By the end of the course, participants will be able to:
- Understand how algorithms and predictive models support organizational decision-making.
- Identify the key assumptions and limitations underlying algorithmic outputs.
- Evaluate the reliability and relevance of model-generated recommendations.
- Assess how data quality affects algorithmic performance.
- Recognize different forms of bias in data, models, and automated decisions.
- Interpret uncertainty, confidence, prediction errors, and model performance indicators.
- Distinguish between correlation, prediction, and causation.
- Identify situations where algorithmic outputs should not be accepted without human review.

- Evaluate model performance when business conditions or data patterns change.
- Recognize warning signs of unreliable or inappropriate algorithmic recommendations.
- Assess explainability, transparency, and accountability requirements.
- Establish appropriate levels of human oversight and intervention.
- Apply governance principles to algorithm-supported business decisions.
- Improve communication between business leaders, data professionals, and technical teams.
- Develop a practical framework for deciding when to trust, challenge, or reject algorithmic outputs.

Who Should Attend

This course is designed for executives, senior managers, department heads, business leaders, decision-makers, risk professionals, strategy leaders, and managers who use or oversee decisions supported by algorithms, predictive analytics, or artificial intelligence.

It is particularly relevant to professionals working in strategy, finance, risk management, operations, compliance, audit, digital transformation, data analytics, artificial intelligence, business intelligence, human resources, marketing, customer experience, and performance management.

The programme is suitable for government and public-sector organizations, banks and financial institutions, oil and gas companies, energy organizations, engineering and industrial companies, telecommunications, healthcare, logistics, and large corporations adopting algorithmic and AI-supported decision-making.

Learning Outcomes

- By the end of the course, participants will be able to:
- Explain the role of algorithms and models in modern organizational decision-making.
- Assess whether an algorithmic output is relevant to a specific business context.
- Identify the main assumptions behind model-generated recommendations.
- Evaluate the impact of data quality and data limitations on model reliability.
- Recognize potential sources of bias in algorithmic systems.
- Interpret model accuracy, confidence, uncertainty, and error indicators.

- Distinguish reliable predictions from potentially misleading outputs.
- Identify model limitations caused by changing markets, behaviors, or operating conditions.
- Detect situations where human judgment should override or challenge an algorithm.
- Evaluate the transparency and explainability of automated recommendations.
- Assess risks associated with excessive reliance on automated decision-making.
- Establish appropriate validation, monitoring, and escalation mechanisms.
- Apply responsible AI and algorithm governance principles.
- Communicate effectively with technical teams when questioning model outputs.
- Develop a practical Trust, Challenge & Escalation Framework for algorithm-supported decisions.

Course Outline

DAY 01

Understanding Algorithms & Model-Based Decision-Making

- What is an algorithm and how does it support business decisions?
- From traditional analytics to predictive and AI-powered models.
- How data is transformed into algorithmic recommendations.
- Model inputs, assumptions, outputs, and decisions.
- Automated recommendations versus human judgment.
- Common types of models used in business environments.
- Where algorithms add value and where they can create risk.
- Understanding the difference between automation and intelligent decision support.
- The role of managers in overseeing algorithm-supported decisions.

PRACTICAL APPLICATION

Mapping how an algorithm could influence a real organizational decision.

DAY 02

Why Algorithms Fail — Data, Bias & Uncertainty

- The relationship between data quality and model performance.
- Missing, inaccurate, outdated, and inconsistent data.
- Sampling and representation problems.
- Understanding algorithmic and data bias.
- Sources of bias throughout the model lifecycle.
- Prediction errors and false conclusions.
- Understanding uncertainty and confidence.
- Correlation versus causation.
- Overfitting, underfitting, and model generalization from a managerial perspective.

PRACTICAL APPLICATION

Diagnosing why a seemingly accurate algorithm produced an unreliable business recommendation.

DAY 03

Evaluating Model Reliability & Performance

- How non-technical professionals can assess model performance.
- Understanding accuracy, precision, recall, and error rates.
- Interpreting performance indicators without technical complexity.
- Model validation and testing.
- Comparing predictions with actual outcomes.
- Identifying performance deterioration over time.
- Model drift and changing business environments.
- Testing whether a model remains relevant to current conditions.
- Recognizing red flags in model-generated recommendations.

PRACTICAL APPLICATION

Reviewing model performance information and determining whether an algorithm should be trusted, challenged, or investigated further.

DAY 04

Human Judgment, Governance & Responsible AI

- The role of human judgment in algorithm-supported decisions.
- Defining human oversight and intervention points.
- When human expertise should override algorithmic recommendations.
- Explainability and transparency in automated decision-making.
- Accountability for algorithm-supported decisions.
- Managing operational, financial, regulatory, and reputational risks.
- Governance frameworks for algorithms and AI systems.
- Monitoring and reviewing automated decisions.
- Establishing escalation and exception-management mechanisms.

PRACTICAL APPLICATION

Designing a human oversight and governance framework for an AI-supported decision process.

DAY 05

When to Trust, Challenge or Reject the Algorithm

- Building a structured framework for evaluating algorithmic outputs.
- Trust criteria: data quality, performance, relevance, stability, and explainability.
- Challenge criteria: uncertainty, anomalies, bias, changing conditions, and conflicting evidence.
- Rejection criteria and high-risk decision situations.
- Establishing decision thresholds and escalation rules.
- Continuous monitoring of model performance.
- Communicating model limitations to executives and stakeholders.
- Building organizational capability for responsible algorithm use.
- Integrating human judgment with AI-supported decision-making.

FINAL WORKSHOP

Developing an integrated Trust, Challenge & Escalation Framework to determine when an algorithmic output can be accepted, when it requires further validation, and when human judgment should override the recommendation.



Register or Request More Information

Course page: <https://quest-training.com/courses/when-to-trust-the-algorithm-understanding-model-limitations-training-course>

Website: <https://quest-training.com>

Email: info@quest-training.com

Phone: +905016771440

