

Module 1: Introduction to AI Safety

Module Overview

Artificial intelligence is transforming industries, improving service delivery, and creating new opportunities for innovation. However, alongside these benefits come significant risks that must be understood and managed. AI systems can unintentionally produce harmful outcomes through bias, misinformation, privacy violations, inaccurate predictions, or misuse by malicious actors. Understanding these risks is the first step toward building AI systems that are safe, responsible, and trustworthy.

This module introduces the fundamental concepts of AI safety and explores how safety differs from ethics and security. Learners will examine real-world examples from Kenya and around the world to understand the consequences of unsafe AI deployment and the importance of designing AI systems that protect people, respect human rights, and promote public trust.

What is AI Safety?

AI Safety refers to the practice of designing, developing, deploying, and maintaining artificial intelligence systems in ways that minimize risks to individuals, communities, and society while maximizing their benefits. Safe AI systems are reliable, transparent, ethical, fair, and accountable throughout their lifecycle.

AI safety encompasses several important dimensions:

Functional Safety focuses on preventing accidents, errors, and unintended consequences. For example, an AI system used in healthcare should provide accurate diagnoses and minimize the risk of harming patients through incorrect recommendations.

Ethical Safety ensures that AI systems operate fairly, respect human rights, avoid discrimination, and include diverse populations in their design and implementation.

Societal Safety addresses the broader impacts of AI on society, including misinformation, public trust, democratic processes, economic inequality, and long-term social risks.

Historical and Kenyan Case Studies of Unsafe AI Deployments

Examining real-world examples helps illustrate why AI safety is essential.

Kenyan Context

One of the most common examples involves **AI-powered credit scoring** used by digital lenders and fintech platforms. Some mobile lending applications rely on algorithms that may unfairly classify low-income earners, informal workers, or individuals with limited digital histories as high-risk borrowers. This can disproportionately affect women, youth, and rural communities by limiting their access to affordable financial services.

Another example is the growing use of **facial recognition technologies** in public spaces. AI-powered surveillance systems have been piloted in parts of Nairobi to improve public security and traffic management. While these technologies can support law enforcement, they also present risks such as privacy violations, inaccurate identification, false positives, and the potential targeting of marginalized communities if appropriate safeguards are not in place.

Kenya has also experienced the impact of **AI-assisted misinformation campaigns**, particularly during the 2017 and 2022 general elections, as well as during the 2024 and 2025 nationwide protests. AI-generated content, manipulated images, automated social media accounts, and deepfake technologies have contributed to the rapid spread of false information, undermining public trust and influencing public discourse.

Global Examples

Unsafe AI deployment has also been observed internationally.

In 2018, Amazon discontinued an AI recruitment system after discovering that it consistently favored male applicants because it had been trained using historical hiring data that reflected gender bias.

The COMPAS Risk Assessment System, widely used in parts of the United States criminal justice system, attracted criticism after studies found that its predictions disproportionately affected certain racial groups, raising concerns about fairness and accountability.

Microsoft's Tay chatbot, launched in 2016, was manipulated by users into generating offensive and harmful content within hours of its release, demonstrating the importance of robust content moderation and safety guardrails.

Tesla's Autopilot technology continues to highlight the risks associated with overreliance on autonomous systems, where incorrect identification of road conditions or driver overconfidence has contributed to several accidents.

Understanding the Difference Between Safety, Ethics, and Security

Although these concepts are closely related, they address different aspects of responsible AI.

AI Safety is concerned with preventing accidental harm and ensuring AI systems perform reliably. For example, an AI-powered health chatbot in Kenya should provide accurate medical information and avoid recommending harmful treatments based on poor-quality training data.

AI Ethics focuses on fairness, accountability, transparency, and respect for human rights. An example would be ensuring that AI-based loan approval systems do not unfairly disadvantage informal sector workers simply because they lack conventional financial records.

AI Security focuses on protecting AI systems from intentional attacks and malicious use. Examples include defending AI systems against hacking, preventing prompt injection attacks, and identifying AI-generated deepfakes that spread false political information during election campaigns.

A useful analogy is to compare these concepts with road transport:

- Safety is like wearing seatbelts and having airbags that protect passengers during accidents.
- Ethics is like ensuring that everyone has an equal opportunity to obtain a driver's licence through a fair process.
- Security is like protecting vehicles from theft, vandalism, or cyberattacks.

Together, safety, ethics, and security create trustworthy AI systems.

Key Risks in Artificial Intelligence

AI systems face several categories of risk that developers and organizations should understand.

Misuse occurs when AI technologies are intentionally used to spread misinformation, generate fake content, manipulate public opinion, or facilitate cybercrime.

Bias arises when AI systems produce unfair outcomes because of unrepresentative datasets or flawed design decisions. In Kenya, biased credit scoring and insurance algorithms may disadvantage people from low-income communities or those working in the informal sector.

Accidents happen when AI systems make unintended errors. Examples include diagnostic systems providing incorrect medical recommendations or automated systems making unsafe operational decisions.

Adversarial Manipulation involves deliberately deceiving AI systems. Examples include modified vehicle licence plates designed to confuse AI-powered traffic cameras or carefully crafted prompts intended to bypass chatbot safety controls.

Learning Outcomes

By the end of this module, learners should be able to:

- Explain the concept of AI safety and its importance in responsible AI development.
- Distinguish between AI safety, AI ethics, and AI security.
- Analyze Kenyan and international case studies involving unsafe AI deployments.
- Identify and categorize common AI risks and their potential impacts.
- Recommend practical strategies for reducing AI-related risks.

Learning Activities

1. Case Study Analysis

Learners will examine the use of AI-based credit scoring in Kenya's digital lending sector.

Working individually or in groups, learners should:

1. Review a summary describing how mobile lending platforms use AI to assess creditworthiness.
2. Identify the factors that contributed to unfair outcomes, such as biased datasets, limited transparency, or the over-penalization of informal workers.
3. Determine whether the issue primarily relates to AI safety, ethics, or security, while recognizing that multiple dimensions may overlap.
4. Recommend at least two practical measures that could reduce or prevent similar risks in future AI systems.

2. Discussion Questions

Learners should reflect on and discuss the following questions:

- Should fintech companies be required to demonstrate fairness before deploying AI-based credit scoring systems?
- Is algorithmic bias primarily a technical problem, or does it reflect broader inequalities within society?
- How can Kenya encourage AI innovation while ensuring that vulnerable populations are not excluded from essential services?

Teaching Materials

The instructor should provide presentation slides covering key AI safety concepts, comparisons between safety, ethics, and security, AI risk categories, and both Kenyan and international case studies.

Recommended readings include *The Malicious Use of Artificial Intelligence: Forecasting, Prevention, and Mitigation* by Brundage et al. (2018), together with articles, policy papers, or reports discussing AI bias and access to financial services in Kenya.

Suggested multimedia resources include news reports on AI-powered surveillance technologies in Nairobi and documentaries or news segments examining misinformation during Kenyan elections.

“

Assessment

Learners' understanding will be assessed through multiple approaches.

A short quiz will test their ability to distinguish between AI safety, ethics, and security using practical Kenyan examples.

Learners will also complete a reflection essay discussing the AI risk they believe poses the greatest challenge to Kenya and explaining their reasoning.

Participation during case study discussions will contribute to the final assessment, with emphasis placed on critical thinking, evidence-based reasoning, and the ability to propose practical solutions for safer AI systems.

Revision #2

Created 2026-08-06 07:16:37 UTC by Angela

Updated 2026-08-06 07:20:48 UTC by Angela