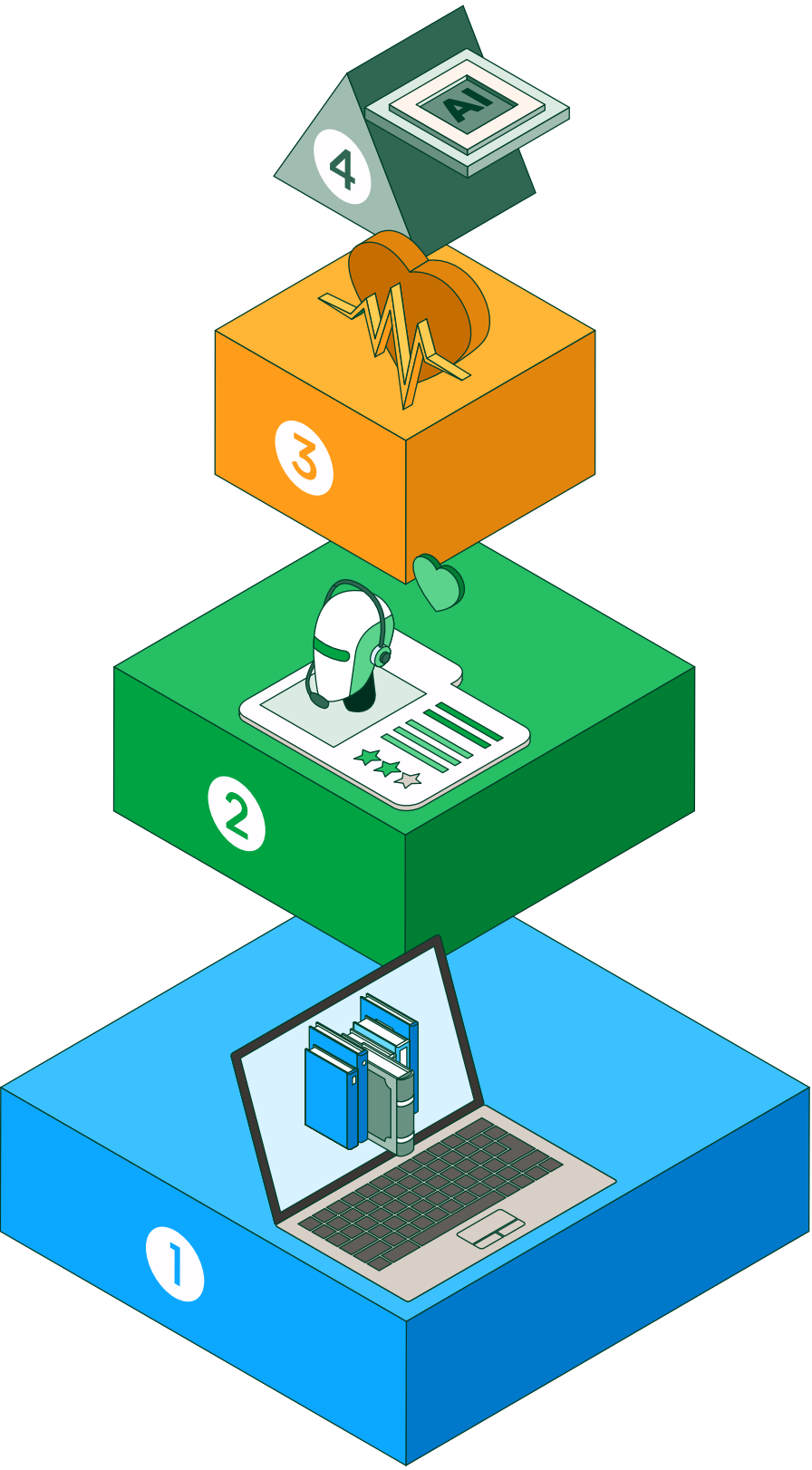


Risk classification under the European AI Act:

An impact-based approach

The AI Act establishes **four levels of risk** for AI systems:



4 Unacceptable risk

Types of use
AI systems considered a clear threat to people’s health, safety and fundamental rights.

Legal implications
Prohibited within the European Union.

Example
The use of AI through subliminal or manipulative techniques to influence human behaviour, or AI systems designed to evaluate or classify individuals based on their behaviour in a way that produces a “social scoring” effect resulting in unfair or adverse treatment.

3 High risk

Types of use
AI systems that may pose serious risks to people’s health, safety or fundamental rights.

Legal implications
Subject to strict requirements regarding transparency, human oversight, accuracy, robustness and cybersecurity, among others.

Example
Biometric AI systems for emotion recognition, or systems used for hiring or selecting individuals – in particular, for publishing targeted job advertisements, analysing and filtering applications, and assessing candidates.

2 Limited risk

Types of use
AI systems not classified as high risk but whose use may pose certain risks to health and safety or to individuals’ fundamental rights.

Legal implications
Transparency obligations to ensure that users are informed, clearly know they are interacting with an AI system, and can identify content that has been artificially generated or manipulated, among other requirements.

Example
Chatbots, virtual assistants or image and text generators.

1 Riesgo mínimo o nulo

Types of use
AI systems with no impact on fundamental rights, health or safety.

Legal implications
No specific requirements, although voluntary compliance measures (codes of conduct) may be adopted.

Example
Email spam or junk filters, or AI-based or AI-enabled video games.

Source: Regulation (EU) 2024/1689 of the European Parliament and of the Council; European Commission, ‘AI Act enters into force’