

ARTIFICIAL INTELLIGENCE IN PRACTICE

Human Purpose, Legal Responsibility, Governance and the Responsible Exploitation of AI-Related Products



An Isebia-style analytical paper

Updated 1 August 2026

Prepared for boards, product owners, investors, professional advisers, public bodies, SMEs and organisations that develop, procure, deploy, finance or commercialise artificial intelligence.

Document status and legal currency

This paper is a general external guidance document. It is not a legal opinion, investment recommendation, technical design or substitute for jurisdiction-specific advice. Artificial intelligence law changes quickly. The legal overview reflects official materials available on 1 August 2026 and should be reverified before any material decision. ^{[10][11]}

The paper deliberately combines law, governance, technology, commercial design and infrastructure. An AI product is not merely a model. It is a socio-technical arrangement made from data, software, interfaces, people, contracts, energy, hardware, suppliers and decisions. A weakness in any one of those elements can become a weakness in the product as a whole.

Tables, checklists and forms have been confined to the appendices. The main paper is written as a continuous argument: first to explain what AI is, then to identify how value is created, then to establish who remains responsible when the machine becomes capable of doing more.

Generative AI tools assisted with research organisation and drafting. Human direction, source selection, legal judgement, editing and responsibility for the final text remain human. That disclosure is not a formality. It expresses the central proposition of this paper.

How to read this paper

A reader concerned with strategy may begin with the Foreword, Executive Summary and Parts I and II. A legal or compliance reader may proceed to Parts III and IV. A product owner or investor should also read Parts V and VI, because an AI proposition can be legally elegant and commercially useless, or commercially attractive and physically impossible. The appendices convert the argument into working tools.

The Value of a Will

The more answers can be generated artificially, the more important it becomes for a human being to decide what is worth asking, why it is asked, how it is asked, and what should be done with the answer. As artificial output becomes abundant, authenticity increases in value. The artificial must remain a servant, not a master.

Adapted from Dwight P. Isebia, 2024. [2]

Artificial intelligence is often described through its answers. That is understandable, because answers are what the user sees. Yet the real human work begins before an answer appears. Someone chooses the objective. Someone selects the data. Someone decides which uncertainty is acceptable. Someone determines whether an answer is advice, evidence, entertainment, instruction or merely a possibility.

An answer without a will is only output. A system may rank, predict, compose, recommend or act, but it does not thereby acquire a legitimate purpose. Purpose remains a human assignment. So do responsibility and restraint.

This is the dividing line between using AI and surrendering to it. A responsible organisation does not ask only whether a system can perform a task. It asks whether the task should be performed, whether the result can be tested, whether affected persons can be treated fairly, whether the data may lawfully be used, and whether someone has the authority and courage to stop the system when it is wrong.

The value of human intelligence will not disappear because artificial capacity becomes cheap. Human intelligence changes its task. It becomes less valuable as a producer of routine first drafts and more valuable as a source of intention, judgement, context, empathy, contradiction and accountability. This is not a defensive argument against technology. It is the condition under which technology can be used well.

Governance begins before the prompt and continues after the output.

The paper that follows develops this idea in practical terms. It explains how organisations can exploit AI-related products without treating legality as an afterthought, governance as paperwork or human oversight as a decorative phrase. The aim is not to slow innovation. The aim is to distinguish acceleration from loss of control.

Executive summary

Artificial intelligence can reduce cost, increase speed, broaden access to knowledge, improve predictions and support forms of creation that were previously expensive or impossible. It can also amplify error, discrimination, surveillance, manipulation, cybersecurity risk and dependence on opaque suppliers. The same system may create benefit in one context and unacceptable harm in another. That is why AI cannot be governed by the name of the technology alone.

The proper unit of analysis is the use case. Who is using the system, for what purpose, on which data, with what degree of autonomy, affecting whom, under which law, and with what route of correction? A chatbot that drafts a marketing headline and a system that recommends whether a person receives credit may use related technology, but they do not present the same risk or demand the same controls.

A mature organisation therefore needs two connected structures. The first is an organisation-wide management system that sets policy, responsibilities, inventory, training, risk appetite and assurance. The second is a use-case lifecycle that examines each application from intention and data through testing, deployment, monitoring and retirement. ISO/IEC 42001 provides a management-system structure; the NIST AI Risk Management Framework provides a flexible process for governing, mapping, measuring and managing risk; ISO/IEC 42005 adds structured impact assessment. [6][7][8]

Legal compliance is layered. AI-specific regulation does not replace data protection, consumer protection, intellectual-property law, product liability, employment law, sectoral regulation, cybersecurity duties, competition law or contract. The European Union's AI Act is the most comprehensive AI-specific framework, but even inside the EU it operates alongside the GDPR, the Data Act, the Cyber Resilience Act, the revised Product Liability Directive, NIS2, DORA and sector rules. Outside the EU, existing laws often apply even where no single AI Act exists.

Commercial exploitation requires the same discipline. A product is not ready because it produces impressive output in a demonstration. It must have a lawful data position, a defined customer, evidence of performance, safe operating limits, service levels, monitoring, support, incident handling, intellectual-property allocation, cybersecurity, change control, pricing, unit economics and a credible exit path. Claims such as 'accurate', 'unbiased', 'compliant', 'autonomous' or 'human-like' must be treated as factual representations, not as marketing decoration.

The physical infrastructure of AI is equally important. Models consume chips, electricity, cooling, water, networks and capital. A statement that an AI data centre has ten megawatts of capacity says nothing by itself about the number of models it can run. It is an electrical statement. The figure remains ambiguous until it is clear whether ten megawatts refers to total facility power or net IT load, what PUE is assumed, how much capacity is contracted and occupied, what rack density is supported, and how energy-price risk is allocated.

A model can distribute capability. It cannot distribute accountability.

The recommended approach is practical. Inventory AI use. Classify each use case. Appoint accountable people. Carry out impact and legal assessments where risk is material. Define human decision rights. Record data lineage and intellectual-property rights. Test performance and misuse. Contract for access, audit, security, updates and exit. Monitor drift and incidents. Stop or redesign systems whose residual risk is not justified by their value.

Responsible AI is not the opposite of profitable AI. It is the means by which a product survives contact with customers, regulators, courts, investors, employees and reality.

Contents

Foreword	The Value of a Will
Executive summary	
Part I Human purpose and the nature of AI	
1	Capacity without intention
2	What AI is - and what it is not
3	How organisations create value with AI
4	The AI product is more than a model
Part II From capability to a commercial product	
5	Begin with a real problem
6	Business models and the AI value chain
7	Evidence before claims
8	Data as raw material and legal responsibility
9	Intellectual property, training data and human authorship
10	Contracting and commercial exploitation
Part III The legal landscape	
11	Law follows the function
12	The European Union AI Act
13	Privacy, profiling and automated decisions
14	Liability, consumer protection and competition
15	International and sectoral overlays
Part IV Governance through the lifecycle	
16	Who is responsible
17	A controlled lifecycle
18	Impact assessment and risk classification
19	Testing, explanation, monitoring and drift
20	Cybersecurity and agentic AI
21	Third parties, open source and supply chains
Part V The physical body of AI	
22	Compute is a material resource
23	What ten megawatts really means
24	Sustainability, society and the digital divide
Part VI Implementation, assurance and future direction	
25	A proportionate approach for SMEs
26	Boards, regulated firms and investors
27	The next stage: hybrid and agentic systems
28	Conclusion: artificial capacity under human direction
Appendices Forms, registers, checklists and sources	

Human purpose and the nature of AI

The first question is not what the system can do. It is what human purpose the system is being asked to serve.

1. Capacity without intention

AI enlarges capacity. It can search more documents than a person can read, detect statistical relations that no unaided analyst would see, draft variations at negligible marginal cost and operate at a speed that changes the economics of many services. Capacity, however, is not intention. A system can optimise only within the objective and information architecture given to it, whether those are explicit, hidden or inherited.

This distinction matters because modern AI can make an organisation feel that a conclusion has become objective merely because it was generated through computation. Yet the choice of variables, labels, examples, thresholds and success criteria is already a set of judgements. When an AI system predicts default, selects a candidate, recommends a treatment or prioritises a fraud alert, it embodies assumptions about what counts as risk, merit, benefit and error.

Human direction therefore has two meanings. The first is operational: a person can intervene, correct or stop the system. The second is normative: a person or accountable body has decided why the system exists, which values govern it and which consequences are not acceptable. Human-in-the-loop without human purpose is only manual supervision of an undefined objective.

The speed of a system does not shorten the distance between a fact and a justified conclusion.

Good governance preserves that distance. It creates time and evidence between output and action. The more serious the consequence, the stronger the need for verification, explanation, appeal and human judgement.

2. What AI is - and what it is not

The term artificial intelligence is used broadly. It covers systems that infer from inputs how to generate predictions, content, recommendations or decisions. Machine learning is one family within that field. Generative AI creates new text, images, audio, video or code. Natural-language processing deals with human language. Computer vision interprets images and video. Robotics connects computation to physical action. Agentic systems plan and execute multiple steps using tools and external systems. [1][3][9]

Not every automated system is AI. A fixed rule that rejects a transaction above a pre-set amount may be conventional software. A model that learns patterns from transaction history and adjusts its

risk score is more clearly machine learning. The distinction is not always legally decisive. Technology-neutral laws may regulate the effect of automation whether or not a vendor calls the system AI.

The word AI is also used as a marketing device. A product may contain a narrow predictive model, a third-party language-model API and a conventional workflow engine, yet be sold as an intelligent platform. A responsible purchaser looks beneath the label. It identifies the actual model, the data flow, the party controlling updates, the permissions granted to the system, and the decisions that remain with people.

Generative systems deserve particular care because fluency can be mistaken for knowledge. A language model is designed to produce plausible sequences, not to guarantee truth. Retrieval, rules, calculations and human review can improve reliability, but no interface should hide the difference between a generated statement and verified evidence.

2.1 Predictive, generative, hybrid and agentic systems

Predictive AI estimates what may happen or how an observation should be classified. It supports forecasting, maintenance, fraud detection, demand planning and risk scoring. Generative AI creates content and can translate, summarise, draft, code and simulate dialogue. Hybrid AI combines these functions: a statistical or rules-based component performs the analysis, while a language model explains the result. This often produces a more reliable architecture than asking one generative model both to calculate and to communicate.

Agentic AI adds a further step. Instead of producing an answer and waiting, an agent can plan a sequence, call tools, access data, send messages, execute transactions and revise its own route. The relevant risk is therefore not only whether the output is correct. It is whether the system has authority to act, which resources it can reach, how its identity is managed, and what happens when one mistaken step becomes the input for the next.

3. How organisations create value with AI

Organisations use AI to automate routine work, support decisions, personalise services, detect anomalies, optimise operations and create new products. The supplied business-development materials identify predictive AI, generative AI, machine learning, natural-language processing and robotics as distinct business tools, each with advantages and failure modes. [3]

The most common early gains come from activities where the cost of a first draft is high but the cost of verification is manageable. Examples include document summarisation, customer-service assistance, translation, coding support, marketing variants, meeting notes, research triage and extraction of structured information from unstructured files.

More advanced uses move from assistance to decision support. Financial institutions detect fraud and assess risk. Manufacturers predict equipment failure and optimise inventory. Healthcare organisations analyse images and patient data. Professional firms search precedents, compare

contracts and organise evidence. Public bodies classify applications and direct cases to the appropriate team.

The final category changes the product itself. AI may become the service: an agent, diagnostic tool, recommendation engine, digital twin, autonomous component, personalised educational system or generative design platform. At that point governance cannot remain an internal IT policy. It becomes part of product quality, customer protection, contracting and valuation.

3.1 The difference between efficiency and substitution

An organisation should be explicit about whether AI assists a person, replaces a task, or replaces a decision. Those are different interventions. A drafting assistant may reduce preparation time while leaving authorship and responsibility with a professional. A fully automated eligibility decision shifts legal and ethical risk because the affected person may be excluded without meaningful human consideration.

The strongest use cases often preserve a deliberate division of labour. Machines perform scale, pattern detection, repetition and simulation. People define goals, test assumptions, interpret context, communicate uncertainty and take responsibility. This is not a compromise between old and new. It is a design principle.

4. The AI product is more than a model

A model is only one component in an AI product. The product includes the user interface, prompts, retrieval sources, data pipelines, model weights, fine-tunes, rules, APIs, logging, security, hosting, human review, customer support, contracts and updates. A change in any of those elements can change the system's legal status, performance and risk.

The same general-purpose model may be used to draft a poem, advise a patient, rank job candidates or control industrial equipment. The model provider may be the same, but the product provider and deployer create the context in which risk becomes concrete. This is why responsibility cannot be outsourced merely by buying an API.

A useful product description should answer four questions. What does the system do? What does it not do? What evidence supports the claim? Who is responsible when the system is used outside its intended conditions? If those questions cannot be answered, the product is not yet mature, regardless of the quality of its demonstration.

A product is not complete when it can generate output. It is complete when someone can explain how it is controlled, supported, corrected, paid for and stopped.

From capability to a commercial product

Commercial value begins when a technical possibility is converted into a reliable, lawful and supportable service.

5. Begin with a real problem

The correct starting point is not 'Where can we use AI?' It is 'Which problem is important enough to solve, and why is AI an appropriate method?' This protects the organisation from technology theatre: projects selected because they are fashionable rather than useful.

A credible use case identifies the user, the decision or task, the existing process, the pain point, the expected benefit and the cost of error. It also identifies the baseline. Without a baseline there is no meaningful claim that the AI system is faster, cheaper, fairer or more accurate.

The problem must be defined with enough precision to test causality. If a chatbot is expected to reduce service cost, the organisation must know whether cost is driven by repetitive questions, poor internal information, complex policy or inadequate staffing. A fluent interface will not repair a broken process. It may merely make the broken process respond more confidently.

The Isebia consistency method is useful here. It asks whether intention and justification align, whether facts and language are verifiable, whether the causal chain is sound, whether the arithmetic is consistent, whether the proposal is lawful and whether it works in practice. It then adds governance, feedback and an exit route. [4]

5.1 Define success and failure before development

Success criteria should be defined before the team becomes attached to the solution. They may include accuracy, response time, human time saved, complaint rate, conversion, false positives, energy cost, fairness measures or revenue. Failure criteria matter equally. A project should state what must not happen: unlawful discrimination, disclosure of confidential data, unsupervised financial commitments, unsafe instructions or dependence on a supplier that cannot be replaced.

A product that has no stop condition has no real risk appetite. It has optimism.

6. Business models and the AI value chain

AI products are exploited through several commercial forms. A provider may sell software as a service, license a model, charge per API call, embed AI in equipment, provide a managed service, operate a marketplace, license data, distribute an open-source core with paid enterprise controls, or combine cloud capacity with advisory and implementation.

Each model allocates risk differently. In a subscription service, the provider controls updates and availability but may restrict audit. In a customer-hosted licence, the customer gains control but carries more infrastructure responsibility. In an API model, output cost and dependence on a third party may vary with usage. In a managed service, the supplier may become involved in operational decisions and personal-data processing.

The AI value chain may include a foundation-model developer, model host, fine-tuning provider, data supplier, system integrator, application provider, cloud or data-centre operator, distributor, reseller and deployer. Contracts must not assume that one party controls the whole system. The party making the customer promise must understand which promises depend on others.

6.1 Unit economics

A product may be technically valuable but commercially weak if inference cost rises with every customer interaction, if support requires constant expert intervention, or if the model supplier can change price and functionality without compensation. Unit economics should include compute, storage, data licensing, retrieval, monitoring, red-team testing, human review, customer support, insurance, incident response and regulatory maintenance.

Revenue should be matched to the service actually delivered. A per-user price may suit a productivity tool. A per-transaction price may suit fraud detection. Reserved capacity may be necessary for regulated customers. Outcome-based pricing can align incentives, but it creates disputes unless outcomes are measurable and causation is clear.

7. Evidence before claims

AI marketing often moves faster than evidence. Words such as intelligent, autonomous, unbiased, compliant, secure, explainable and human-like can carry legal consequences. A claim about accuracy is testable. A claim that a system complies with the law may imply a level of legal assurance that the seller is not competent or able to provide.

Consumer-protection authorities have repeatedly stressed that existing law applies to AI. In the United States, the Federal Trade Commission has stated that there is no AI exemption from laws against unfair or deceptive practices. It has challenged unsupported accuracy and earnings claims and warned companies not to change data-use promises quietly. [26]

Evidence should match the context of use. A model that performs well on a public benchmark may fail on a customer's language, population, document quality or operating conditions. A product that is accurate on average may still be unacceptable if its errors concentrate on a protected or vulnerable group.

The evidence file should preserve test data, methodology, exclusions, limitations, versions and dates. A model update can invalidate an earlier claim. Monitoring should therefore be connected to marketing: when evidence changes, claims and customer instructions must change too.

A convincing answer is not the same as a supported answer. A successful demonstration is not the same as a reliable product.

8. Data as raw material and legal responsibility

AI learns, retrieves and acts through data. Data may be technically available and still be unlawful, inaccurate, unrepresentative, confidential, biased or commercially restricted. The legal right to use data and the technical ability to use data are different questions.

A data strategy should distinguish training data, fine-tuning data, retrieval sources, prompts, customer inputs, telemetry, labels, evaluation sets and outputs. Each category may have a different owner, legal basis, retention period and confidentiality status.

Data lineage is the record of where data came from, what happened to it, which system used it and where it went. Without lineage, an organisation cannot answer basic questions about consent, deletion, error correction, contamination or intellectual-property rights. It cannot reliably reproduce an output or investigate harm.

Data quality is contextual. A dataset may be large but irrelevant, accurate but outdated, representative of a national population but unsuitable for a local community. Synthetic data can reduce privacy risk but may reproduce the assumptions of the generator. Anonymisation can help, but re-identification risk must be assessed realistically.

8.1 The temptation of secondary use

AI creates a powerful incentive to reuse data for new purposes. Customer conversations become training material. Employee documents become a knowledge base. Product telemetry becomes a behavioural model. That reuse may create value, but it may conflict with consent, purpose limitation, confidentiality, trade secrets or contractual promises.

A lawful and trustworthy product makes secondary use explicit. It does not hide a new business model inside revised terms or a broad statement that data may be used to improve services.

9. Intellectual property, training data and human authorship

AI raises separate questions about inputs, software, models and outputs. The supplied copyright paper places human creative will at the centre of authorship and warns that the use of protected material for training or generation must be analysed separately from ownership of the resulting output. [2]

Copyright generally protects original expression rather than ideas, methods or functionality. Source code can be protected even where the underlying function is not. Model weights, fine-tunes, prompts, taxonomies, retrieval databases, evaluation sets and confidential workflows may also be

protected through copyright, database rights, contract or trade-secret law, depending on jurisdiction and facts.

Training data presents the most contested issue. A provider may rely on licences, public-domain material, statutory text-and-data-mining exceptions, fair use or another legal basis. Those routes differ between jurisdictions and remain fact-sensitive. A customer should not assume that a vendor's ability to offer a model proves that every training input was lawfully used.

In the United States, the Copyright Office's 2025 report on copyrightability maintained the human-authorship requirement and explained that human selection, arrangement or modification may be protected, while purely AI-generated material is not. Its separate training report addresses the use of copyrighted works in generative-AI development. [17][18]

Contractual ownership of output is not the same as copyright ownership. A supplier can assign whatever rights it has, but a contract cannot manufacture copyright where the law does not recognise sufficient human authorship. The practical solution is to preserve evidence of human creative contribution and to allocate rights in prompts, drafts, modifications, selections and final arrangements.

9.1 A rights architecture

A mature product separates background IP from project IP. Background IP includes pre-existing models, software, methods and know-how. Project IP includes customer-specific configurations, fine-tunes, integrations and documentation. Customer data and output should be separately defined. The contract should also address feedback, telemetry, model improvement and whether a customer's confidential material may influence services supplied to others.

Open-source components require the same attention. 'Open' does not mean unrestricted. Model licences may limit commercial use, regulated use, redistribution, scale or prohibited applications. Software, weights and datasets may each carry different licences.

10. Contracting and commercial exploitation

The AI contract is the place where technological uncertainty becomes allocated risk. It should describe the intended use, excluded use, roles, data, intellectual property, service levels, security, human oversight, changes, audit, liability, support, exit and regulatory cooperation.

A supplier may resist a warranty that every output is accurate because probabilistic systems cannot guarantee that result. The customer should then seek architecture and process commitments: agreed testing, documented limitations, error correction, monitoring, escalation and the right to stop or revert after material deterioration.

Service levels must fit the product. Availability alone is insufficient for AI. The contract may need response latency, throughput, model version, context-window limits, rate limits, retrieval freshness, false-positive thresholds, content-filter behaviour or human-review times.

Change control is central. A cloud provider can update a model, safety filter or API without the customer changing any code. That update may improve capability but alter output, compliance or cost. Material changes should trigger notice, regression testing and, for regulated or high-impact uses, approval.

10.1 Liability and indemnity

Liability should follow control. The provider is best placed to manage defects in its model or service, the integrator controls configuration, the deployer controls the use context, and the customer may control data and human decisions. Indemnities can address third-party IP claims, confidentiality, security breach and unlawful processing, but they are meaningful only if backed by financial capacity and insurance.

An exclusion of all indirect loss may remove the very losses most likely to arise from AI, including business interruption, regulatory investigation, reputational repair and data restoration. The parties should negotiate with a realistic view of the harm, not merely copy a conventional software template.

10.2 Exit and substitutability

An organisation should know how it leaves before it enters. Exit requires export of data, prompts, logs, fine-tunes, embeddings, configurations and documentation in a usable format. It may require a transition period, model replacement, deletion evidence and continued access to audit records. A system that cannot be replaced is not merely a vendor relationship. It is a strategic dependency.

The legal landscape

AI law is not one law. It is a layered field in which old duties meet new capabilities.

11. Law follows the function

The legal analysis begins with what the system does, not with whether it is called AI. A recommendation engine may engage consumer law. A hiring tool engages employment and equality law. A medical system may be a regulated product. A credit score may engage financial regulation and rules on automated decisions. A generative platform may raise copyright, privacy and platform issues at the same time.

The relevant role also matters. A foundation-model provider, product provider, deployer, importer, distributor, integrator and infrastructure operator may each have different obligations. A company can occupy several roles at once.

Jurisdiction follows markets, users, data and effects. A company outside the European Union may be within the AI Act or GDPR because it offers a system into the EU or its output is used there. A service operating globally may face different rules on data transfers, content, discrimination, consumer protection and government access.

The practical response is not to seek one universal compliance label. It is to create a jurisdiction and role map for each product and to design a common governance baseline strong enough to support local adaptations.

12. The European Union AI Act

The EU AI Act entered into force on 1 August 2024. Prohibited practices and AI-literacy duties began to apply on 2 February 2025. Governance and general-purpose-AI obligations began to apply on 2 August 2025. The general application date is 2 August 2026. Following the AI Omnibus, which entered into force on 27 July 2026, Annex III high-risk rules apply from 2 December 2027 and rules for high-risk systems embedded in regulated products apply from 2 August 2028. [10][11]

The Act is risk-based. Some practices are prohibited. High-risk systems face lifecycle duties concerning risk management, data governance, technical documentation, records, transparency, human oversight, accuracy, robustness, cybersecurity and post-market monitoring. Certain systems that interact with people or generate synthetic content face transparency obligations. Many lower-risk uses remain largely outside mandatory product controls but are still subject to other law.

The classification exercise should not be reduced to a label. A system may become high-risk because of its intended purpose, sector or integration into regulated equipment. A general-purpose

model may not itself be a high-risk application, but a downstream provider can create one by embedding it in employment, education, credit, critical infrastructure or another sensitive use.

12.1 Prohibited and sensitive uses

Organisations should maintain a prohibited-use policy that reflects both the Act and their own values. A use may be legally permissible and still outside risk appetite. Emotion recognition, biometric categorisation, manipulation, social scoring and certain surveillance functions require particular care. The absence of a clear prohibition does not remove duties under privacy, equality, consumer or human-rights law.

12.2 General-purpose AI and downstream responsibility

General-purpose AI regulation recognises that foundation models serve many applications. Model providers have transparency, documentation and copyright-related duties, with additional obligations for models presenting systemic risk. Downstream providers must obtain enough information to understand and control the product they build. A contractual statement that the customer is responsible for compliance cannot replace the technical information needed to comply.

12.3 Transparency

The Act's transparency duties include informing people when they interact with certain AI systems and identifying or labelling certain synthetic content. The Commission's Article 50 guidance and code of practice materials are intended to support implementation. [10]

Transparency should be useful, not ceremonial. A notice that 'AI may be used' tells a person little. A meaningful notice explains what the system does, what data it uses, how the output affects the person, whether a human reviews it, and how the person can challenge or correct the result.

13. Privacy, profiling and automated decisions

The GDPR remains central where AI processes personal data. Its principles require lawfulness, fairness, transparency, purpose limitation, data minimisation, accuracy, storage limitation, security and accountability. Article 22 provides protection against certain decisions based solely on automated processing that produce legal or similarly significant effects. [12]

AI creates pressure against several of those principles. Developers want more data, while minimisation requires restraint. Models may discover new uses, while purpose limitation demands clarity. Complex systems may be difficult to explain, while transparency requires meaningful information. Continuous learning can change performance, while accuracy and accountability require control.

The legal basis must be analysed for each processing purpose. Consent is not a universal answer. It must be informed, specific, freely given and withdrawable, and it may be unsuitable where there is a power imbalance. Legitimate interests require balancing. Contract necessity must be interpreted narrowly. Special-category data requires an additional condition.

13.1 Data protection by design

Privacy controls should be built into architecture: separation of training and production data, minimisation of prompts and logs, role-based access, retention controls, encryption, local or regional processing where appropriate, and means to honour access, correction, objection and deletion.

A Data Protection Impact Assessment and an AI impact assessment overlap but are not identical. The DPIA focuses on risks to people from personal-data processing. The AI assessment also considers safety, fairness, autonomy, social impact, cybersecurity, IP, operational dependency and environmental effects.

13.2 Human intervention

Human intervention is meaningful only if the person has competence, time, information and authority. A reviewer who approves hundreds of algorithmic recommendations without understanding the model is not exercising oversight. The system may be formally human-assisted and substantively automated.

14. Liability, consumer protection and competition

When AI causes loss, liability may arise through contract, negligence or other civil wrongs, professional duties, product liability, data protection, discrimination, consumer law or sector regulation. The path depends on the parties, product and jurisdiction.

The revised EU Product Liability Directive expressly brings software within the product-liability framework and updates the regime for new technologies, including AI. This matters where defective software, updates or related services cause covered damage. [15]

Professional users remain responsible for professional judgement. A lawyer, doctor, accountant or engineer cannot rely on a disclaimer from an AI supplier to discharge duties owed to a client or patient. AI may change the standard of care over time, but it does not erase the duty to verify.

Consumer law controls presentation as well as performance. Dark patterns, anthropomorphic design and emotional simulation can create undue trust. A system that sounds certain may cause a user to rely on it more than a conventional disclaimer can correct. Product design is therefore part of legal risk.

14.1 Competition and concentration

AI markets may concentrate around chips, cloud capacity, data and foundation models. Organisations should examine exclusivity, switching barriers, access to output and data, self-preferencing, tying and restrictions on interoperability. An apparently convenient integrated stack can become a long-term dependency.

Competition concerns also arise from algorithmic pricing, coordinated market behaviour and use of sensitive competitor information. The fact that an algorithm identifies a pricing strategy does not make an unlawful result lawful.

14.2 The Data Act and connected products

The EU Data Act has applied since 12 September 2025. It addresses access to and use of data generated by connected products and related services, as well as cloud switching and certain contractual issues. AI products that depend on industrial or device data should treat data-access rights as part of product architecture and contracting. [13]

14.3 Cybersecurity duties

The Cyber Resilience Act entered into force in December 2024. Reporting obligations begin on 11 September 2026 and the main obligations apply from 11 December 2027. Products with digital elements may need secure development, vulnerability handling and conformity measures. NIS2 and DORA may impose additional organisational and sectoral duties. [14][23][24]

15. International and sectoral overlays

There is no single international AI code. The United States relies on a mixture of federal and state law, sector regulators, civil-rights rules, consumer protection and voluntary frameworks such as NIST. China combines cybersecurity, data-security and personal-information laws with specific rules on algorithms, generative AI and deep synthesis. The United Kingdom uses existing law and regulator-led guidance while developing its framework. Singapore combines a model governance framework with practical testing initiatives such as AI Verify.

The OECD AI Principles, adopted in 2019 and updated in 2024, provide an international reference for innovative and trustworthy AI that respects human rights and democratic values. In 2026 the OECD also published due-diligence guidance for responsible AI. [9]

A global product should identify the strictest material requirements and design a common baseline, but it should not assume that compliance in one jurisdiction guarantees compliance elsewhere. Data localisation, government access, content controls, employment rules, consumer rights and liability can differ sharply.

15.1 Finance

Financial institutions must integrate AI into existing governance, model risk, outsourcing, operational resilience, conduct and consumer-protection frameworks. A regulator may ask not only what model is used but how the board understands its role, how third-party dependencies are controlled, how unfair outcomes are tested and how decisions can be explained.

15.2 Healthcare and life sciences

Clinical AI can become regulated software or a medical device. Safety, efficacy, validation, change control and post-market surveillance are central. The use of generative AI to explain a result does not remove the need for validated diagnostic logic and professional oversight.

15.3 Employment and education

AI used to screen, rank, monitor or assess people can reproduce historical exclusion and affect fundamental opportunities. Organisations should test for group impacts, provide notice and review, avoid irrelevant proxies and preserve an appeal route.

15.4 Government and public services

Public-sector AI engages legality, equality, procedural fairness, transparency and democratic accountability. Procurement should secure access to documentation, audit, data, source or escrow where necessary, and long-term continuity. A public authority should not become unable to explain a decision because a supplier claims trade secrecy.

15.5 Gaming, media and personalised systems

AI can detect fraud and harmful behaviour, personalise content and create immersive services. It can also intensify addiction, manipulation, impersonation and deepfake risk. Product design should distinguish assistance from behavioural exploitation and should protect children and vulnerable users.

Governance through the lifecycle

Governance is the discipline of keeping purpose, evidence and authority connected as the system changes.

16. Who is responsible

The supplied governance guide recommends an AI Governance Committee or, for a smaller organisation, an AI Governance Officer, supported by cross-functional collaboration among legal, compliance, IT, engineering and business functions. [1]

Responsibility begins with the board or governing body. It sets risk appetite, approves material uses and receives assurance. Management converts that appetite into policy and resources. Product owners are accountable for value and operation. Technical teams build and monitor. Legal and compliance interpret obligations. Security tests resilience. Internal audit or independent assurance challenges the system.

An AI committee should not become a place where responsibility disappears into collective discussion. Every use case needs a named owner. Every material risk needs a control owner. Every decision to deploy or continue needs an accountable approver.

16.1 The AI inventory

The organisation cannot govern systems it does not know it uses. The inventory should include purchased tools, embedded features, employee experimentation, models inside business software, customer-facing products and automated decisions. Shadow AI is not a minor policy issue; it is an information-security and confidentiality risk.

The inventory should record purpose, owner, provider, model, data, affected persons, autonomy, jurisdiction, risk class, approval status, monitoring and retirement date. Appendix B provides a practical intake form.

16.2 AI literacy

Training should match role. All staff need to understand confidentiality, verification and approved tools. Product teams need lifecycle and data knowledge. Reviewers need to understand limitations and automation bias. Boards need enough technical and legal literacy to challenge claims and allocate resources.

17. A controlled lifecycle

Governance is strongest when integrated into the project lifecycle rather than applied at the end. A useful lifecycle begins with intention and problem definition, then moves through data and rights,

design, supplier selection, development or configuration, testing, approval, deployment, monitoring, incident response and retirement.

Each phase should have evidence and a decision gate. A pilot should not drift into production because users find it convenient. A production system should not expand to a new country or sensitive purpose without review. A material model update should not be treated as routine maintenance if it changes behaviour.

17.1 Gate one: justified purpose

The first gate asks whether the problem is real, the expected benefit is material and AI is proportionate. It identifies affected persons and unacceptable outcomes. Projects without a clear purpose should stop here.

17.2 Gate two: lawful and suitable data

The second gate confirms rights, legal basis, quality, representativeness, lineage, security and retention. If the data cannot support the intended claim, more sophisticated modelling will not repair the foundation.

17.3 Gate three: validated system

The third gate requires technical, human-factors, fairness, security and misuse testing. The benchmark must reflect actual operating conditions. Residual risks and limitations must be documented.

17.4 Gate four: controlled deployment

Deployment requires trained users, access control, monitoring, incident response, change control, customer information and an accountable decision. The system should have a rollback or safe-state route.

17.5 Gate five: continued justification

A system should continue only while its value and risk remain justified. Monitoring may show drift, cost increase, supplier deterioration, new law, changed user behaviour or a better alternative. Retirement is a normal lifecycle decision, not an admission of failure.

18. Impact assessment and risk classification

Risk classification should consider consequence and context, not only model type. Relevant factors include impact on rights or safety, vulnerability of affected people, scale, autonomy, reversibility, opacity, data sensitivity, novelty, dependency and the possibility of misuse.

An impact assessment should describe benefit as well as harm. A system can reduce waiting time, improve access or detect fraud while also increasing surveillance or exclusion. The assessment should make those trade-offs visible and identify who receives the benefit and who bears the risk.

ISO/IEC 42005 provides guidance for AI system impact assessments, while the NIST framework encourages organisations to govern, map, measure and manage risk throughout the lifecycle. The supplied Isebia consistency test adds intention, factual clarity, causality, arithmetic, legality, practical execution and feedback. [7][8][4]

18.1 Risk is not a colour

A red, amber or green label can assist communication, but it must not replace analysis. Some risks are showstoppers; others can be mitigated. A high-impact system may be acceptable with strong controls and clear public benefit. A low-impact tool may still be unacceptable if it systematically discloses confidential data.

The risk register should state cause, event, impact, pre-control likelihood and severity, control, owner, residual risk, indicator and escalation route. Appendix D provides a model.

19. Testing, explanation, monitoring and drift

Testing should ask whether the system works, for whom it works, when it fails and how failure is detected. Accuracy is only one measure. Others include robustness, calibration, fairness, latency, cost, security, refusal behaviour, hallucination rate, recovery and usability.

Generative systems require scenario-based evaluation. Test sets should include ordinary requests, ambiguous prompts, adversarial inputs, confidential information, prohibited content, multilingual use and attempts to manipulate tools. Human evaluation remains necessary where quality depends on meaning, tone or professional judgement.

19.1 Explanation

Different people need different explanations. A user may need to know why a recommendation affected them. A regulator may need documentation and logs. A developer needs technical diagnostics. A board needs assurance that risks are controlled. One generic explanation cannot serve all audiences.

An explanation should not pretend that a complex model is simpler than it is. It can describe the purpose, important factors, data sources, confidence, limitations, human role and contest route. Where the underlying process cannot be explained adequately for a high-impact use, the organisation should consider a more interpretable architecture.

19.2 Drift and continuous change

Model drift occurs when the world changes and the relationship learned from earlier data no longer holds. Data drift, concept drift, user adaptation, policy changes and supplier updates can all alter performance. Monitoring should therefore compare current behaviour with the validated baseline.

A system may remain technically available while becoming substantively unfit. Availability dashboards do not reveal whether an eligibility model has become unfair or whether a language model has started producing unsupported citations after an update.

20. Cybersecurity and agentic AI

AI strengthens cybersecurity and creates new attack surfaces. OWASP's guidance for large-language-model applications identifies prompt injection, insecure output handling, data poisoning, denial of service and supply-chain vulnerabilities among the major risks. Its 2026 agentic guidance addresses autonomous systems that plan, use tools and act across workflows. [19][20]

Prompt injection exploits the system's difficulty in separating instruction from content. A malicious document, webpage or user message may cause the model to reveal data or call a tool. The defence is architectural: isolate data from instructions, restrict permissions, validate outputs, apply least privilege and require approval for sensitive actions.

Agents add identity and authority risk. An agent may have access to email, payment, code, customer records or industrial controls. Its credentials should be limited, short-lived and attributable. Tool calls should be logged. High-risk actions should require step-up authentication or human approval. Budgets and rate limits should constrain financial and computational damage.

Memory is another attack surface. Poisoned or incorrect information can persist and influence later actions. Organisations should define what an agent may remember, who can alter memory, how sources are verified and how memory is corrected or deleted.

Autonomy is not an absence of control. It is control designed before the moment of action.

20.1 Incident response

An AI incident may involve data disclosure, harmful output, discriminatory effect, model theft, supply interruption, malicious agency, runaway cost or incorrect automated action. The incident plan should connect technical containment with legal notification, customer communication, evidence preservation and model rollback.

21. Third parties, open source and supply chains

Most organisations do not build every component. They depend on models, clouds, datasets, libraries, vector databases, evaluation tools and consultants. Third-party risk is therefore part of the product, not a procurement side issue.

Due diligence should establish who owns and operates the model, where data is processed, whether customer data is used for training, how updates are controlled, what security and assurance evidence exists, how incidents are reported, what subcontractors are used, and whether the service can be replaced.

Open-source models can improve transparency and control, but they shift responsibility. The deployer may become responsible for hardening, content safety, updates, licensing, export controls and monitoring. Model cards and repository popularity are not substitutes for independent testing.

21.1 Supply-chain concentration

A product may depend on one chip vendor, one cloud, one model API and one dataset. Each dependency can fail commercially, technically or politically. Resilience requires alternatives, inventories, escrow or export rights where appropriate, and a plan for degraded operation.

The physical body of AI

Artificial intelligence may appear immaterial on a screen. In reality it is built from energy, minerals, water, networks, land, labour and capital.

22. Compute is a material resource

The public conversation often treats AI as software alone. Yet every training run and inference request is executed on hardware in a physical location. The product depends on electricity, cooling, fiber, storage, transformers, batteries, generators, semiconductor supply, maintenance and skilled operators.

This physical layer affects law and economics. Data location influences privacy and government access. Energy price affects unit cost. Grid availability limits scale. Cooling and water affect environmental permission and community acceptance. Chip origin and customer use may engage export controls. A cyber-secure model hosted in an insecure facility is not a secure product.

The supplied universal data-centre framework treats land, power, planning, environment, fiber, market demand, finance, cybersecurity, insurance and community impact as parallel diligence streams. No single attractive feature makes a site investable. [22]

23. What ten megawatts really means

A ten-megawatt AI data centre is not a statement about the amount of data stored, the number of models available or the intelligence of the system. It is a statement about electrical power capacity. The supplied technical note distinguishes power from energy, gross facility capacity from net IT load, installed capacity from actual use and redundancy from consumption. [5]

23.1 Power is not energy

A megawatt measures the rate at which electricity is consumed or delivered. One megawatt equals one thousand kilowatts. Energy is power used over time and is expressed in megawatt-hours or gigawatt-hours. A facility drawing ten megawatts continuously uses ten megawatt-hours in one hour, 240 megawatt-hours in a day and 87.6 gigawatt-hours in a year.

The simplest analogy is a water pipe. Megawatts describe the flow rate. Megawatt-hours describe the volume that has passed through the pipe.

23.2 Gross facility capacity and net IT load

The expression 'ten megawatts' is incomplete unless the project states what the number measures. If it is total utility or facility capacity, that power must serve servers, GPUs, networking, cooling,

pumps, fans, lighting, security, UPS losses, transformers and offices. If it is net IT capacity, the grid connection must be larger because the supporting facility also consumes power.

Power Usage Effectiveness, or PUE, expresses the relationship:

PUE = total facility power / IT equipment power

At a PUE of 1.25, a ten-megawatt gross facility can provide approximately eight megawatts to IT equipment. Conversely, ten megawatts of net IT capacity requires approximately 12.5 megawatts of total facility power. Those are materially different projects.

23.3 Installed, contracted, occupied and actual

Capacity has several meanings. The technical system may be designed for ten megawatts. The utility may contract eight. Customers may reserve six. Hardware may be deployed for five. The facility may draw four at a particular moment. Investors and customers should never allow those figures to be presented as though they are the same.

Saleable capacity is also lower than nameplate capacity where power is reserved for networks, storage, safety margin, maintenance or future growth. Revenue depends on saleable and occupied IT capacity, not on the number printed on the substation.

23.4 Rack density and hardware

Conventional enterprise racks may use roughly five to fifteen kilowatts. High-density cloud racks may use twenty to forty. Modern AI racks may require sixty to 150 kilowatts, and advanced liquid-cooled systems can exceed that range. If ten megawatts were entirely available for IT, it could support about 200 racks at fifty kilowatts, 100 racks at 100 kilowatts or about sixty-seven racks at 150 kilowatts.

Rack count still does not reveal the number of GPUs or the useful compute. That depends on GPU model, servers, network fabric, storage, redundancy and power reserved for non-GPU equipment.

23.5 Redundancy

N, N+1 and 2N describe the amount of backup equipment. A 2N electrical design may have two complete UPS or transformer paths, but it does not normally consume twice the operational power. Installed equipment rating, resilient capacity and actual load must remain separate in the financial and technical model.

23.6 Energy cost

A facility drawing ten megawatts continuously uses 87.6 gigawatt-hours per year. At eight euro cents per kilowatt-hour, the annual electricity cost is about EUR 7.0 million. At ten cents it is EUR 8.76 million. At fifteen cents it is EUR 13.14 million.

If ten megawatts refers to net IT load and PUE is 1.25, the total demand is 12.5 megawatts and annual consumption is 109.5 gigawatt-hours. At ten euro cents per kilowatt-hour, the annual

electricity cost is approximately EUR 10.95 million. Energy price and PUE therefore affect both profitability and environmental impact.

A ten-megawatt promise is a power statement, not an intelligence statement.

23.7 Investment meaning

An investable specification should state net critical IT load, design and measured PUE, contracted grid capacity, delivery date, redundancy, supported density, liquid-cooling capability, energised capacity, customer reservations, actual utilisation and the party carrying electricity-price risk. Appendix H converts those questions into a diligence form.

24. Sustainability, society and the digital divide

AI can improve energy systems, health, education and public services, but it can also concentrate resources and increase inequality. Access to advanced models depends on compute, data, skills and capital. Communities that supply energy, water or land should receive a credible account of the public value created.

Environmental claims require measurement. PUE measures facility overhead but does not measure carbon, water or hardware impact. WUE, carbon intensity, embodied emissions, renewable matching, refrigerants, e-waste and hardware utilisation may all matter. A low PUE facility powered by high-carbon electricity is not necessarily sustainable.

Twenty-four-seven carbon-free claims require hourly evidence and contractual attribution, not annual certificates alone. Water claims should distinguish withdrawal from consumption and should address drought and local competition. Circularity should include repair, reuse, component recovery and safe end-of-life treatment.

Social impact includes employment transition. AI may remove repetitive tasks and create new roles, but benefits and losses will not be evenly distributed. A responsible implementation identifies affected workers, redesigns jobs, provides training and measures whether human capability is strengthened or merely displaced.

Artificial abundance makes human responsibility scarce, and therefore more valuable.

Implementation, assurance and future direction

The organisation does not need perfect certainty. It needs a disciplined way to decide under uncertainty.

25. A proportionate approach for SMEs

Small and medium-sized organisations may not have the resources for a large governance department. The supplied governance guide recommends a proportionate core: an accountable AI officer, an inventory, applicable-law identification, risk assessment, a risk register, approved-tool rules and reporting to the governing body. [1]

The first objective is visibility. Identify which tools staff use, what information enters them and which outputs affect customers or decisions. Ban the use of unapproved public tools for confidential or personal data. Provide a secure approved alternative.

The second objective is concentration. Focus resources on the few use cases with material effect. A low-risk drafting assistant may require policy, training and verification. A system that determines pricing, eligibility or employment requires a deeper assessment, testing, documentation and review.

The third objective is contractual leverage. SMEs often rely on third parties and may have limited negotiating power. They can still compare vendors, refuse undisclosed training on customer data, demand deletion, export and incident terms, and avoid products that cannot explain their basic controls.

25.1 A ninety-day start

In the first thirty days, appoint an owner, issue an interim acceptable-use rule and create an inventory. In the next thirty, classify use cases, review the most important suppliers and establish a risk register. In the final thirty, approve priority systems, train users, implement monitoring and present the first governance report to the board.

26. Boards, regulated firms and investors

Boards should treat AI as a strategic and control issue, not as an IT feature. They should understand which use cases are material, where the organisation depends on third parties, what data and IP positions exist, how incidents are handled, and whether investment claims are supported by evidence.

Regulated firms should connect AI governance to existing risk, outsourcing, model, conduct, operational-resilience and audit frameworks. Creating a separate AI committee should not duplicate or weaken existing accountability. The AI layer should make those frameworks more coherent.

Investors should distinguish a technical demonstration from a repeatable business. Due diligence should test customer demand, revenue quality, compute and energy cost, supplier concentration, data rights, model performance, security, legal classification, insurance, team capability and the cost of compliance.

Claims about proprietary technology deserve particular scrutiny. A company may own an interface while depending entirely on a third-party model. It may call a dataset proprietary while lacking rights to use it for training. It may report gross model usage as revenue without accounting for compute cost. Evidence levels should separate identified, tested, contracted and operational facts.

26.1 Assurance

ISO/IEC 42001 can support a certifiable AI management system, while NIST provides a voluntary risk-management structure. Certification can strengthen trust, but it does not certify that every use case is lawful or safe. Assurance must still examine the product and context. [6][7]

Internal audit, independent technical review, red teaming, legal review and customer assurance should be coordinated. Repeated evidence should be stored in an AI assurance file rather than recreated for every questionnaire.

27. The next stage: hybrid and agentic systems

The next phase of AI adoption will not be defined only by larger models. It will be defined by systems that combine models, rules, retrieval, tools and memory to perform sustained work. Hybrid systems can improve reliability by separating calculation from communication. Agentic systems can reduce coordination cost by executing multi-step processes.

This increases the importance of identity, authority and boundaries. The key question shifts from 'What can the model say?' to 'What can the system do?' An agent that drafts an email is different from one that sends it. An agent that recommends a payment is different from one that authorises it.

OWASP's 2026 agentic guidance identifies threats including goal hijacking, tool misuse, identity and privilege abuse, memory poisoning, insecure inter-agent communication, cascading failures and rogue agents. The ICO has similarly emphasised that organisations remain responsible for data-protection compliance when they develop, deploy or integrate agentic AI. [20][25]

The design response is explicit authority. Give agents the minimum permissions, verify tool outputs, separate planning from execution, require approval for irreversible acts, constrain budgets, log actions and maintain a safe shutdown route.

27.1 Personalised AI

Personalised systems may become persistent advisers that know a person's habits, finances, health, relationships and ambitions. Their usefulness will be inseparable from their intimacy. Consent, dependency, manipulation, memory, inheritance, impersonation and the right to disengage will become major legal and ethical questions.

27.2 AI for social good

AI can help address climate, health, education and infrastructure challenges. Good intention is not enough. Social-good systems still require evidence, community participation, rights protection and sustainable operation. The people affected should help define both the problem and the measure of success.

28. Conclusion: artificial capacity under human direction

Artificial intelligence changes the amount of work that can be done, the speed at which it can be done and the number of people who can access sophisticated capability. It does not remove the need to decide what work deserves to be done.

The organisation that succeeds will not be the one that uses AI everywhere. It will be the one that knows where AI creates real value, where it requires restraint, where human judgement must remain decisive and where the evidence no longer justifies continuation.

Law provides boundaries. Governance connects those boundaries to daily decisions. Contracts allocate control and consequence. Technology creates capability. Infrastructure gives that capability a physical body. Human intelligence gives it purpose.

The artificial should extend human capacity without dissolving human responsibility.

That is the standard by which AI-related products should be developed, purchased, exploited and judged.

Forms, registers, checklists and sources

The appendices convert the principles in the paper into working documents. They should be adapted to the organisation, sector and jurisdiction. Red flags should not disappear inside an average score.

Appendix A - Glossary

Term	Working definition
AI system	A machine-based system that infers from inputs how to generate predictions, content, recommendations or decisions.
Agentic AI	An AI system that plans and executes multi-step tasks, often by using tools, memory and external services.
AI impact assessment	A structured evaluation of intended benefit, affected persons, potential harm, controls and residual risk.
AI inventory	The organisation's record of AI systems, tools, owners, data, risks, approvals and lifecycle status.
AI literacy	Knowledge and practical competence needed to use, supervise or govern AI appropriately.
Data lineage	The recorded origin, transformation, use, movement and destination of data through its lifecycle.
Deployer	An organisation or person using an AI system under its authority, excluding purely personal non-professional use.
Drift	A change in data, environment or relationships that causes model performance to deteriorate or change.
Foundation model / GPAI model	A broadly capable model that can be adapted or integrated into many downstream applications.
Generative AI	AI that creates new content such as text, images, audio, video or code.
Human oversight	A designed ability for competent people to understand, intervene, override, correct or stop a system.
Inference	The operation in which a trained model produces an output from new input.
Model	A mathematical or computational representation learned or configured to produce outputs from inputs.
PUE	Power Usage Effectiveness: total facility power divided by IT equipment power.
Provider	A party that develops an AI system or model, or has it developed, and places it on the market or puts it into service under its name.
Retrieval-augmented generation	An architecture in which a generative model receives relevant information from an external source before producing an answer.
Residual risk	Risk remaining after controls have been applied.
Training	The process of adjusting a model using data so that it learns patterns or behaviour.

Appendix B - AI use-case intake form

Field	Response
Use-case title	
Business owner	
Technical owner	
Version and date	
Problem to be solved	
Why AI is appropriate	
Current process and baseline	
Expected benefit	
Affected users and persons	
AI type / model / supplier	
Inputs and data sources	
Outputs and decisions	
Degree of autonomy	☐ Advisory ☐ Decision support ☐ Automated decision ☐ Agentic action
Human review	
Jurisdictions and sectors	
Personal / sensitive data	
IP and confidentiality	
Security classification	
Unacceptable outcomes	
Success KPIs	
Stop / exit criteria	
Initial risk classification	☐ Low ☐ Moderate ☐ High ☐ Prohibited / no-go
Required approvals	

Appendix C - AI impact assessment form

Assessment area	Core question	Evidence / finding	Action / owner
Purpose and necessity	What legitimate objective is served? Is AI necessary and proportionate?		
Stakeholders	Who benefits, who is affected and who may be vulnerable?		
Rights and fairness	Could the system discriminate, exclude, manipulate or impair autonomy?		
Data	Are data lawful, relevant, representative, accurate and traceable?		
Privacy	What personal data is processed? Is a DPIA required?		
Human oversight	Can a competent person review, intervene and reverse?		
Transparency	What must be explained to users, affected persons and regulators?		
Accuracy and robustness	What performance is required and under which conditions?		
Cybersecurity	How could the system be attacked, poisoned, extracted or misused?		
IP and confidentiality	Are training, inputs and outputs lawfully and contractually controlled?		
Operational dependency	What suppliers, clouds, chips, data or staff are critical?		
Environmental impact	What compute, energy, water and hardware impacts arise?		
Social and workforce impact	How are jobs, skills, accessibility and the digital divide affected?		
Misuse and dual use	What foreseeable harmful or prohibited uses exist?		
Residual risk	Is the remaining risk justified by the benefit?		

Appendix D - AI risk register

ID	Risk	Cause	Impact	I	L	Controls	Residual	Owner	Indicator
R-01	Hallucinated or inaccurate output	Reliance on generative output without verification	Wrong advice, loss or harm	4	4	Grounding, verification, confidence and human review	2	Product Owner	Error rate / complaints
R-02	Confidential data disclosure	Unapproved tool or prompt leakage	Breach of confidence / privacy	5	3	Approved tools, DLP, access and training	2	CISO	DLP alert / incident
R-03	Bias or unfair impact	Unrepresentative data or proxy variables	Discrimination and legal liability	5	3	Impact assessment, subgroup testing, appeal	2	Risk Owner	Fairness metrics
R-04	Model or data drift	World or user behaviour changes	Performance deterioration	4	3	Monitoring, thresholds and revalidation	2	Model Owner	Drift threshold
R-05	Prompt injection / tool misuse	Malicious or indirect instruction	Unauthorised action or data loss	5	3	Isolation, least privilege, validation, approvals	2	Security Lead	Tool-call anomaly
R-06	Supplier dependency	Sole model or cloud provider	Outage, price increase or lock-in	4	4	Exit rights, export, alternatives and resilience	3	Procurement	SLA / price change
R-07	IP infringement	Unlawful data, output or component use	Claims, injunction or redesign	4	3	Rights diligence, licences, indemnity and records	2	Legal	Claim / takedown
R-08	Regulatory misclassification	Incorrect role or risk analysis	Non-compliance and market delay	5	2	Legal assessment and change review	2	Compliance	Law / use change
R-09	Runaway cost	Uncontrolled inference or agent action	Financial loss or service denial	3	4	Budgets, quotas, rate limits and alerts	1	Operations	Cost threshold
R-10	Inadequate human oversight	Reviewer lacks time, authority or competence	Rubber-stamping and hidden automation	5	3	Work design, training, authority and sampling	2	Business Owner	Override / review data

Appendix E - Data and intellectual-property rights matrix

Asset / data	Owner	Source	Legal basis / licence	Permitted use	Retention	Exit / deletion
Pre-training data						
Fine-tuning data						
Evaluation data						
Retrieval / knowledge base						
Customer prompts and inputs						
Model weights						
Fine-tuned weights / adapters						
System prompts						
Generated outputs						
Telemetry and logs						
Feedback and corrections						
Embeddings / vector store						

Appendix F - Vendor and model due-diligence questionnaire

Domain	Question	Evidence / response
Identity and control	Who develops, owns, hosts and updates the model or service?	
Model and version	Which model, architecture and version are supplied?	
Intended and prohibited use	What uses are supported, restricted or excluded?	
Training data	What categories and legal bases support training? Are opt-outs respected?	
Customer data	Is customer data retained, reviewed or used for training?	
Location and transfers	Where are data, logs and backups processed?	
Performance evidence	Which benchmarks, populations and limitations apply?	
Fairness	What subgroup testing and mitigation are performed?	
Security	What assurance, penetration tests and AI-specific controls exist?	
Incident response	What constitutes an incident and when is the customer notified?	
Subprocessors and supply chain	Which critical third parties are used?	
Updates	How are material model, API and safety changes notified and tested?	
Availability and support	What SLA, recovery and support commitments apply?	
IP	What rights, warranties and indemnities apply to inputs and outputs?	
Regulatory role	What provider/deployer or sector obligations does the vendor support?	
Audit and evidence	What reports, logs, documentation and audit rights are available?	
Exit	Can data, prompts, embeddings, logs and configurations be exported?	
Financial resilience	Can the vendor support warranty, indemnity and continuity obligations?	

Appendix G - AI contract review checklist

Clause area	Checked	Review point
Scope and intended purpose	☐	Functions, users, territories, limitations and prohibited uses
Roles and compliance	☐	Provider, deployer, processor/controller and regulatory cooperation
Data use	☐	Inputs, logs, retention, training, secondary use and deletion
Intellectual property	☐	Background IP, project IP, outputs, feedback and open-source components
Performance	☐	Benchmarks, acceptance, limitations and remediation
Service levels	☐	Availability, latency, throughput, support and recovery
Human oversight	☐	Review, override, escalation and allocation of decision responsibility
Security	☐	Controls, testing, vulnerability handling and incident notice
Changes and updates	☐	Notice, regression testing, approval and rollback
Audit and documentation	☐	Logs, model/data information, assurance reports and customer audit
Liability	☐	Caps, exclusions, professional risk, regulatory loss and business interruption
Indemnities	☐	IP, privacy, confidentiality, security and third-party claims
Insurance	☐	Cyber, professional, product and technology cover
Subcontractors	☐	Approval, flow-down, location and replacement
Price and cost control	☐	Usage, model changes, token limits, energy/compute and indexation
Exit and transition	☐	Export, migration, deletion, continued access and assistance
Suspension and termination	☐	Safety, law, non-performance, insolvency and transition

Appendix H - AI data-centre capacity clarification form

No.	Question	Response / evidence
1	Does the quoted MW refer to gross site capacity or net critical IT load?	
2	Is capacity installed, contracted, deliverable, energised or merely planned?	
3	What design PUE is assumed?	
4	Is PUE a target or a measured operating figure, and at what load?	
5	How much capacity is currently energised?	
6	How much IT capacity is leased or reserved by customers?	
7	What standard and high-density rack loads are supported?	
8	Is the cooling system designed for direct-to-chip or immersion-cooled AI racks?	
9	What electrical and cooling redundancy standard is used?	
10	What is the expected actual average load and utilisation ramp?	
11	What expansion capacity is legally, physically and electrically available?	
12	Who bears electricity-price, PUE and curtailment risk?	
13	What is the contracted grid capacity, voltage, delivery date and curtailment status?	
14	How much of nameplate power is saleable IT capacity after reserves and overhead?	
15	What annual MWh/GWh and energy cost arise in base and downside cases?	
16	What fiber diversity, latency and carrier commitments support the compute product?	
17	What GPU, server, network and storage architecture is assumed?	
18	What environmental, water, carbon, insurance and community evidence exists?	

Appendix I - Board AI dashboard

Dashboard area	Measure	Current	Trend	Board action
AI inventory	Total systems / approved / pilot / shadow / retired			
Risk profile	Low / moderate / high / prohibited use cases			
High-impact decisions	Systems affecting rights, safety, employment, credit or access			
Incidents	Security, privacy, harm, outage, cost and regulatory events			
Performance	Accuracy, drift, complaints, overrides and service levels			
Data and IP	Open rights issues, DPIAs, deletion and licence risks			
Third parties	Critical providers, concentration and exit readiness			
Compliance	AI Act, sector, privacy, CRA, NIS2/DORA and local actions			
People	AI literacy, reviewer competence and workforce impact			
Economics	Revenue, savings, compute cost, human review and variance			
Sustainability	Energy, PUE/WUE, carbon, hardware utilisation and e-waste			
Decisions required	Approve, redesign, pause, invest or retire			

Appendix J - AI product exploitation readiness checklist

Readiness domain	Minimum evidence	Complete
Purpose and customer	Defined customer problem, intended use and baseline	☐
Market evidence	Interviews, pilots, pricing and credible demand	☐
Product evidence	Context-specific performance, limitations and acceptance	☐
Data rights	Lawful, documented, relevant and exportable data position	☐
IP position	Background/project/output rights and open-source compliance	☐
Legal classification	Role, jurisdiction, risk class and sector analysis	☐
Governance	Owner, committee/officer, inventory, impact assessment and risk register	☐
Human oversight	Competence, authority, time, override and appeal	☐
Security	Threat model, testing, least privilege and incident response	☐
Third parties	Due diligence, SLAs, changes, audit and exit	☐
Operations	Monitoring, drift, support, documentation and retirement	☐
Commercial model	Pricing, unit economics, cost limits and concentration	☐
Contracts	Data, IP, SLA, liability, indemnity, change and exit	☐
Insurance and finance	Coverage, capital, warranty and continuity capacity	☐
Infrastructure	Compute, power, network, location and resilience	☐
Sustainability	Energy, water, carbon, hardware and social impact	☐
Claims	Every material marketing claim is evidenced and current	☐
Go/no-go	Residual risk is justified and decision is recorded	☐

Appendix K - Source register

No.	Source	Type	Access
1	Ramparts, Navigating the Complexities of AI Governance: A Comprehensive Guide, updated March 2026	Supplied source	Supplied file
2	Dwight P. Isebia, Copyright Aspects of Artificial Intelligence, 8 September 2024	Supplied source	Supplied file
3	AI for Business Development: Prompt and Overview	Supplied source	Supplied file
4	Isebia Consistency Test - Prompt and One-Page Form	Supplied source	Supplied file
5	Technical note: What a 10 MW AI data centre means	Supplied source	Supplied file
6	ISO/IEC 42001:2023 - AI management systems	Official standard overview	Official link
7	NIST Artificial Intelligence Risk Management Framework	Official guidance	Official link
8	ISO/IEC 42005 - AI system impact assessment	Official standard overview	Official link
9	OECD AI Principles and Due Diligence Guidance for Responsible AI	Official guidance	Official link
10	European Commission, AI Act and implementation timeline	Official EU source	Official link
11	European Commission, AI Omnibus enters into force, 27 July 2026	Official EU source	Official link
12	Regulation (EU) 2016/679 - General Data Protection Regulation	Official EU law	Official link
13	European Commission, Data Act	Official EU source	Official link
14	European Commission, Cyber Resilience Act	Official EU source	Official link
15	Directive (EU) 2024/2853 on liability for defective products	Official EU law	Official link
16	Directive (EU) 2019/790 on copyright and related rights in the Digital Single Market	Official EU law	Official link
17	U.S. Copyright Office, Copyright and Artificial Intelligence, Part 2: Copyrightability	Official US source	Official link
18	U.S. Copyright Office, Copyright and Artificial Intelligence, Part 3: Generative AI Training	Official US source	Official link
19	OWASP Top 10 for LLM Applications 2025	Industry security guidance	Official link
20	OWASP Top 10 for Agentic Applications 2026	Industry security guidance	Official link
21	MediteranAI / ExpertAigroup Human-AI and infrastructure concept materials	Supplied source	Supplied file
22	Universal Data Centre Site Development Information Requirements and Due Diligence Questionnaire	Supplied source	Supplied file

No.	Source	Type	Access
23	Regulation (EU) 2022/2554 - Digital Operational Resilience Act	Official EU law	Official link
24	Directive (EU) 2022/2555 - NIS2	Official EU law	Official link
25	UK Information Commissioner's Office, AI and data protection / agentic AI materials	Official UK source	Official link
26	US Federal Trade Commission, AI privacy, confidentiality and deceptive claims materials	Official US source	Official link

Publication note

This paper should be reviewed whenever a material law, guidance document, model, supplier or use case changes. Source references support general propositions only; legal conclusions depend on facts and jurisdiction.

mr Dwight P. Isebia LL.M

Author and responsible for the final editorial review

1 August 2026