

AI FIELD GUIDE

LIVING REFERENCE

CRAIG WATERHOUSE

VERSION 1.9 · 1 AUGUST 2026

Introduction

AI stopped being optional somewhere in the last three years. It now sits inside the tools most professionals use daily, shapes more decisions than most organisations can comfortably list, and changes faster than any technology in working memory. That speed has produced two kinds of writing about it: breathless commentary that ages in weeks, and technical depth that never reaches a decision. This guide is an attempt at the missing middle — a working reference, written by a practitioner for practitioners, deep enough to reason with and honest about what will be out of date by the time you finish reading it.

It is written for the technically literate professional who already uses AI and now has to make decisions about it: which models and platforms to adopt, what they really cost, what to automate and what to keep human, how to govern and secure the result, and how to keep your own skills compounding while the surface churns. Nothing here requires you to write code, though the guide will take you to the edge of building if you want to go there.

Three concerns recur so often in real AI work that they run through the whole book as marked lenses — governance, delivery and project management, and security. And because the facts of this field split cleanly into the durable and the perishable, the guide is engineered around that split: concepts are written to last, while fast-moving facts are concentrated into clearly flagged sections, dated, and refreshed with each version.

The guide serves readers in the United States, the European Union, and Australia as a single audience. Where the three regions genuinely differ — regulation, data residency, procurement expectations, energy and infrastructure — the difference is called out inline at the point it matters, rather than relegated to an appendix or defaulted to any one region's view.

Finally, a disclosure that doubles as a demonstration: this guide is researched, drafted, and maintained with Claude, an AI assistant made by Anthropic — using exactly the workflows, subject to exactly the limitations, and checked with exactly the discipline it describes. Read it the way it was written: ground what matters, verify what is load-bearing, and treat the dated sections as snapshots of a moving field.

How to use this guide

This is a living document. AI moves fast enough that specific model names, prices, and benchmark standings begin ageing within weeks, so the guide is built to be refreshed rather than read once and shelved.

The durable / perishable split

Concepts that rarely change (how models work, how agents are structured, the principles of token economics, how to learn) are written to last. Fast-moving facts (current flagship models, pricing, leaderboard positions) are concentrated into a few clearly-marked sections so you know exactly what to re-verify and when. Each fast-moving section carries its own “last verified” line.

Cross-cutting lenses

Three concerns recur throughout AI work and appear as consistently-formatted callout boxes wherever they are relevant, so you meet them at the point they actually matter rather than only in a distant appendix:

GOVERNANCE Governance

Regulation, policy, data-use terms, auditability, accountability, and who is answerable for what.

DELIVERY / PM Delivery / PM

How this affects planning, estimation, cost/latency budgeting, delivery risk, and running AI work as projects.

SECURITY Security

Attack surfaces, data exposure, trust boundaries, least-privilege, and safe deployment patterns.

A note on sources

Fast-moving claims are grounded in current web sources and footnoted so you can re-verify them yourself when you update. Benchmarks are treated as directional, not definitive — the only reliable signal for a specific job is testing on your own representative data.

Contents

The complete guide spans the seventeen sections below — click any entry to jump straight to it. An annotated map of what each section covers follows on the next page; sections marked ⚡ are the fast-moving ones.

- [1. Foundations & Glossary](#)
- [2. The Model Landscape ⚡](#)
- [3. Top AI Platforms in Use Today ⚡](#)
- [4. Prompt Engineering & Configuration](#)
- [5. Multimodal & Generative Media ⚡](#)
- [6. Agents: Building, Configuring & Connecting](#)
- [7. Coding & Developer AI ⚡](#)
- [8. Data: The Substrate ⚡](#)
- [9. AI for Business: Adoption, Delivery & Project Management](#)
- [10. Measuring AI Value & Return](#)
- [11. People & Work](#)
- [12. The Economics of AI ⚡](#)
- [13. Governance, Risk & Compliance](#)
- [14. AI Security & Threats](#)
- [15. Safety, Reliability & Limitations](#)
- [16. Trends & Near-Term Outlook ⚡](#)
- [17. Learning AI: Roadmap & Resources ⚡](#)

The seventeen sections

For orientation, here is the full structure and what each section holds. Sections marked ✂ are the fast-moving ones to refresh most often.

1. Foundations & Glossary — Plain-English primer on how AI models work, plus a working glossary.

2. The Model Landscape ✂ — Shared reference: leading frontier and open-weight models compared on capability, context, modality, and price.

3. Top AI Platforms in Use Today ✂ — ~10 consumer/business platforms: focus, roadmap, users, differentiation, use cases.

4. Prompt Engineering & Configuration — Prompting techniques, model settings, Claude-specific configuration, and getting great documents out of AI.

5. Multimodal & Generative Media ✂ — Image, video, audio, voice, and music tools and their current state.

6. Agents: Building, Configuring & Connecting — What agents are and why; configuration, hosting, harnesses, the MCP tool ecosystem, connecting AI to your own data (RAG), and chaining multiple agents together.

7. Coding & Developer AI ✂ — Claude Code, Cursor, Copilot, “vibe coding” — and where non-developers benefit.

8. Data: The Substrate ✂ — Readiness, quality, provenance, privacy, retrieval and permissions — the substrate AI runs on, and the most common cause of failure.

9. AI for Business: Adoption, Delivery & Project Management — Build vs buy, ROI, integration, the AI project lifecycle, LLMOps, change management, a practical delivery playbook, and the guide’s full treatment of evaluation as a discipline (Section 9.10).

10. Measuring AI Value & Return — Did it pay? Six kinds of return, the frameworks that exist, the metrics that matter, the counterfactual problem, and a worked business case.

11. People & Work — What the evidence shows about AI and employment, the junior pipeline, role redesign, consultation obligations, and algorithmic management.

12. The Economics of AI ✂ — Token economics and the cost paradox, the capex and energy boom, who makes money, and the sustainability debate.

13. Governance, Risk & Compliance — The EU AI Act, the US federal and state picture, the Australian picture, ISO/IEC 42001 and the NIST AI RMF, and internal governance.

14. AI Security & Threats — OWASP LLM Top 10, MITRE ATLAS, prompt injection, agent security, secure deployment.

15. Safety, Reliability & Limitations — Model behaviour: hallucination, bias, failure modes — the companion to security.

16. Trends & Near-Term Outlook ✂ — The 12–24 month outlook: capability, the scaling and AGI debate, and what to treat with scepticism.

17. Learning AI: Roadmap & Resources ✂ — Staged curriculum: Stage 1–4 modules with platform tracks (Anthropic, OpenAI, Google, Microsoft, open-weight, coding tools) and verified free resources — courses, docs, and videos.

1. Foundations & Glossary

Concept section — durable. Current-state examples verified 29 July 2026.

This section builds a mental model of what today's AI systems actually are and how they behave, then closes with a glossary you can jump back to. The aim is not surface-level familiarity but a working model solid enough to reason about the platforms, agents, economics, and risks in later sections. Nothing here requires code; a few deeper asides are marked and can be skipped without losing the thread.

1.1 What a large language model actually is

A large language model (LLM) is, mechanically, a next-token predictor. Given a stretch of text, it produces a probability distribution over what comes next, samples one token, appends it, and repeats. Everything an assistant appears to do — reason, summarise, translate, write code — emerges from doing this extremely well over a model that has absorbed statistical structure from a vast amount of text.

The engine underneath is the transformer, introduced in 2017. Its key mechanism, self-attention, lets the model weigh how much every token should influence every other token when forming its prediction. This is what gives models their sense of context and long-range coherence, and — because attention can be computed in parallel — what made training at today's scale practical.

It is worth being precise about one thing: an LLM is not a database and does not “look up” facts. It has learned patterns and compressed regularities. That single fact explains both its fluency and its failure modes — most importantly, hallucination (see Section 1.12 and Section 15).

How a model is trained

Modern models are built in two broad phases:

- **Pretraining** — the model learns language and world structure by predicting the next token across enormous text (and increasingly image, audio, and code) corpora. This is where most of the compute and cost sits, and it produces a raw “base” model that is knowledgeable but not yet helpful or safe.
- **Post-training** — the base model is shaped into a usable assistant. This includes supervised fine-tuning (SFT) on curated examples, and reinforcement learning from human or AI feedback (RLHF / RLAIFF) to align tone, helpfulness, and safety. Increasingly it also includes reinforcement learning against verifiable rewards (e.g. did the code run, was the maths correct), which is the engine behind the recent leap in reasoning models.

GOVERNANCE Where a model's knowledge comes from

Training-data provenance is now a live governance and legal question: what was in the corpus, whether it was licensed, and whether outputs can reproduce copyrighted material. When selecting a model or vendor, provenance, indemnification, and data-use terms belong in the evaluation, not just benchmark scores.

1.2 Tokens — the unit of everything

Models do not read words or characters; they read tokens — sub-word chunks produced by a tokeniser. In English, a token averages roughly $\frac{3}{4}$ of a word (about four characters), so 1,000 tokens is around 750 words. Common words are often a single token; rare words, code, and other languages fragment into many.

Tokens matter because they are the unit of three things at once: what the model can hold (context), what you pay (pricing is per token, input and output priced separately), and how long you wait (latency scales with tokens generated). Almost every practical trade-off in AI — cost, speed, how much you can feed in — reduces to token accounting.

DELIVERY / PM Tokens are your budget line

Because cost and latency are both per-token, token volume is a first-class project constraint, not an afterthought. Estimating a workload means estimating tokens: input size × requests, plus expected output length, plus any retrieved context. A feature that quietly stuffs large documents into every call can cost 10–100× a lean one. Budget it early (see Section 12).

A non-obvious equity issue also lives here: because tokenisers are optimised for English, the same sentence in many other languages costs more tokens — and therefore more money and latency — for identical meaning.

1.3 The context window

The context window is the maximum amount of text (in tokens) a model can consider in a single request — your prompt, any attached documents, the running conversation, retrieved data, and the model's own output all share this budget. It is the model's working memory, and it is finite.

Windows have grown dramatically. Where 4,000–8,000 tokens was typical a few years ago, leading models in 2026 offer 200,000 to 1,000,000 tokens, with some pushing to two million and beyond — enough to hold entire books, large codebases, or long conversation histories at once.

A crucial caveat: advertised context is not the same as effective context. Models can lose track of material buried in the middle of a very long input (the “lost in the middle” effect), so a one-million-token window does not guarantee equally reliable recall across all one million tokens. For long-document work — likely your most common use — structure matters: put the most important material where the model attends best (typically the start and end of the prompt) and state the task explicitly.

1.4 Parameters, model size, and Mixture of Experts

Parameters are the learned numerical weights inside a model — the values adjusted during training. Counts run from a few billion to hundreds of billions or more. Loosely, more parameters mean more capacity, but the relationship to real-world quality is far from linear, and “bigger” has stopped being a reliable proxy for “better.”

Two developments broke the size-equals-quality assumption:

- **Mixture of Experts (MoE)** — instead of activating the whole network for every token, MoE models route each token through a small subset of specialised “expert” sub-networks. A model may have hundreds of billions of total parameters but only activate

tens of billions per token, buying large-model quality at a fraction of the inference cost. Most of the largest 2026 models are MoE.

- **Distillation and strong small models** — techniques that transfer capability from a large “teacher” into a compact “student” now yield small models that are startlingly capable, and some run on a single GPU or even a laptop. For many routine tasks the quality gap to the frontier is small enough that the cheaper, faster small model is the right engineering choice.

1.5 Reasoning (“thinking”) models and test-time compute

The most consequential shift of the last two years is the rise of reasoning models. Rather than answering immediately, these models generate an extended internal chain-of-thought — working through the problem step by step — before producing a final answer. Spending more computation at answer time (“test-time compute”) rather than only at training time markedly improves performance on hard maths, coding, and multi-step logic.

The trade-off is direct: reasoning models are slower and cost more (you pay for those thinking tokens), so they are not the right default for every task. Many current models expose configurable reasoning effort — none, low, medium, high — letting you dial thinking up for a hard analytical problem and down for a quick reply. Knowing when a task genuinely needs reasoning versus when a fast standard model suffices is one of the highest-leverage judgements in day-to-day use (and in cost control — see Section 12).

1.6 Multimodality

Early LLMs were text-only. Frontier models are now natively multimodal: a single model can accept and, increasingly, produce combinations of text, images, audio, and video. “Native” matters — a model trained from the ground up across modalities generally reasons across them more coherently than a text model with a separate image tool bolted on.

Practically, this means you can hand a model a screenshot, a chart, a PDF page, or a voice note and have it reason over the content directly. Dedicated generative-media tools for images, video, voice, and music are covered in Section 5; the point here is that multimodality is now a baseline expectation of a frontier model, not a specialty feature.

One clarification the rest of the guide relies on: not every generative model is a language model. The image and video tools in Section 5 mostly use diffusion, a different mechanism — the model learns to reverse a step-by-step noising process, starting from random noise and denoising it, guided by your text prompt, into a coherent image. The mental model to carry forward is that a transformer predicts the next token while a diffusion model refines a whole canvas at once; the token economics and context limits above are an LLM story, and the media tools have their own cost and quality trade-offs (Section 5).

1.7 Embeddings, vectors, and semantic search

An embedding is a representation of a piece of text (or an image, etc.) as a list of numbers — a vector — positioned in a high-dimensional space so that things with similar meaning sit close together. “Invoice” and “bill” land near each other; “invoice” and “giraffe” do not.

This is the quiet workhorse behind a great deal of applied AI. Because meaning becomes geometry, you can search by concept rather than keyword: convert a query into a vector and

retrieve the nearest stored vectors. A vector database stores embeddings at scale and does this similarity search fast. It is the retrieval half of Retrieval-Augmented Generation (RAG), covered in depth in Section 6 — the standard way to give a model access to your own documents without retraining it.

1.8 Three ways to adapt a model to your needs

When a base model doesn't do quite what you want, there are three broad levers, in ascending order of effort and cost:

- **Prompting / in-context learning** — shape behaviour purely through instructions and examples in the prompt. No training, instant, cheapest to iterate. Covers the majority of real needs (Section 4).
- **Retrieval (RAG)** — keep the model as-is but feed it relevant snippets from your own data at query time. Best when the model needs current or proprietary knowledge it wasn't trained on (Section 6).
- **Fine-tuning** — further-train the model on your own examples to bake in a style, format, or narrow skill. Most powerful and most costly; usually the last resort after prompting and retrieval, not the first move.

A useful rule of thumb: prompt first, retrieve when it needs your data, fine-tune only when the first two genuinely fall short. A fourth lever appears once you move from answering to acting: extending the model with tools and skills, so it can perform actions and follow repeatable procedures — the subject of Section 6.

GOVERNANCE Adaptation raises data questions

Fine-tuning and retrieval both put your — possibly sensitive or personal — data into the AI pipeline. That triggers privacy, IP-ownership, and retention questions: where the data is stored, whether it is used to train the provider's models, and how it is deleted. Confirm the vendor's data-use terms before feeding real data in (Sections 13 and 14).

1.9 Inference and the knobs that shape output

Inference is the act of running a trained model to produce output. A handful of settings shape what you get:

- **Temperature / top-p (sampling)** — control randomness. Low values make output focused and repeatable (good for extraction, code, factual work); higher values make it more varied and creative. There is no universally "correct" value — it depends on the task.
- **Max output tokens** — a ceiling on response length; also a cost and latency control.
- **System prompt** — a standing instruction that frames the model's role and rules for the whole conversation, distinct from the user's turn-by-turn messages.
- **Stop sequences** — strings that tell the model to halt, useful for structured output.

One expectation to reset: LLMs are generally non-deterministic. The same prompt can yield different wording on different runs, even at low temperature. This is normal, and it is why reliable AI systems are validated with evaluations over many samples rather than a single lucky output (Sections 9 and 15).

1.10 How models are served and accessed

The same underlying model can reach you through very different channels, and the channel determines where your data goes:

- **Consumer apps** — chat interfaces and assistants (the way most people use AI today, including for document creation). Easiest to use; least control over data handling.
- **APIs** — programmatic access to a provider's hosted model, billed per token. The basis of most business integrations.
- **Self-hosted open-weight** — download an open-weight model and run it on your own or cloud infrastructure. Maximum control and privacy; you own the operational cost and complexity (Section 12).

SECURITY The channel is a data-exposure decision

Every prompt you send to a third-party API or consumer app leaves your environment. Before putting sensitive, regulated, or client data through a model, know: does the provider train on your inputs by default? What is the retention period? Where (which country) is it processed — a real data-residency concern under the EU's GDPR, US state privacy laws, and Australian privacy obligations alike? For the most sensitive workloads, private endpoints or self-hosted open-weight models keep data in your control. Details in Section 14.

1.11 Limitations to keep in mind

A working mental model has to include the failure modes. These are covered in full in Section 15, but keep them in view from the outset:

- **Hallucination** — models can state false information fluently and confidently. Verification is not optional for anything that matters.
- **Knowledge cutoff** — a model only “knows” up to its training cutoff unless given tools or retrieval to reach current information.
- **No persistent memory by default** — a model doesn't inherently remember past sessions; continuity is engineered (via context, retrieval, or an explicit memory feature), not innate.
- **Prompt sensitivity** — small wording changes can shift output quality noticeably; this is the whole premise of prompt engineering (Section 4).
- **Non-determinism** — as above, repeatability isn't guaranteed, which shapes how AI systems must be tested.

1.12 Glossary

Concise definitions for the recurring vocabulary of the guide. Terms are expanded where they first appear in later sections.

Agent	An AI system that pursues a goal over multiple steps, choosing and using tools, rather than answering a single prompt. See Section 6.
Alignment	The effort to make a model's behaviour match human intentions and values — helpful, honest, and safe.
API	Application Programming Interface — programmatic access to a model, billed per token; the basis of most integrations.

Attention	The transformer mechanism that weighs how much each token influences every other, giving models their sense of context.
Benchmark	A standardised test of model capability (e.g. GPQA, SWE-bench). Directional, not definitive — verify on your own data.
Chain-of-thought (CoT)	A model working through a problem in explicit intermediate steps before answering; the basis of reasoning models.
Context window	The maximum tokens a model can consider at once — prompt, documents, history, and output combined. Its working memory.
Distillation	Transferring capability from a large “teacher” model into a smaller, cheaper “student” model.
Embedding	Text (or image) represented as a vector of numbers so that similar meanings sit close together; powers semantic search and RAG.
Fine-tuning	Further-training a model on your own examples to instil a style, format, or narrow skill. The most costly adaptation lever.
Foundation model	A large model pretrained broadly, meant to be adapted to many downstream tasks.
Frontier model	A leading proprietary model at the cutting edge of capability (e.g. the current flagships from major labs).
Grounding	Anchoring a model's output in verifiable external sources (e.g. retrieved documents) to reduce hallucination.
Guardrails	Controls that constrain what a model can accept or produce — policy filters, output checks, permission limits.
Hallucination	Confident, fluent output that is factually wrong. A core limitation, not a rare glitch.
Inference	Running a trained model to produce output (as opposed to training it). What you pay for per token in production.
Jailbreak	A prompt crafted to bypass a model's safety guardrails.
Knowledge cutoff	The date beyond which a model has no trained knowledge unless given tools or retrieval.
Latency	Time to produce a response; scales with tokens generated and with reasoning effort.
LLM	Large Language Model — a transformer-based next-token predictor trained on vast text (and often other modalities).
MCP	Model Context Protocol — an open standard for connecting models to external tools and data sources. See Section 6.
Mixture of Experts (MoE)	An architecture that activates only a subset of specialised sub-networks per token, cutting inference cost at large scale.
Multimodal	Able to accept and/or produce more than one data type — text, image, audio, video.
Open-weight vs open-source	Open-weight: the trained weights are downloadable. Open-source: weights plus training code/data under an open licence. Not synonyms.
Parameters	The learned numerical weights inside a model, adjusted during training. Rough proxy for capacity, not for quality.
Post-training	The stage after pretraining — SFT and RLHF/RLAIF — that turns a raw base model into a helpful, safe assistant.

Pretraining	The first, compute-heavy training stage: learning language and world structure by predicting the next token at scale.
Progressive disclosure	Loading a skill or its context in tiers — metadata first, full instructions when relevant, deeper files only when needed — so an agent can keep many capabilities available at low context cost. See Section 6.
Prompt / system prompt	The input to a model. The system prompt is a standing instruction framing role and rules for a whole conversation.
Prompt injection	An attack that smuggles malicious instructions into content a model reads, hijacking its behaviour. See Section 14.
Quantisation	Compressing a model's weights to lower precision so it runs faster and on cheaper hardware, with some quality trade-off.
RAG	Retrieval-Augmented Generation — fetching relevant snippets from your data at query time to ground a model's answer. Section 6.
Reasoning model	A model that generates an internal chain-of-thought before answering, spending test-time compute for higher accuracy.
RLHF / RLAIF	Reinforcement Learning from Human / AI Feedback — post-training that aligns a model's behaviour to preferences.
Sampling (temperature, top-p)	Settings that control randomness in generation — low for focused/repeatable, high for varied/creative output.
Skill	A reusable, self-contained package of instructions (a SKILL.md file, optionally with scripts and resources) that an agent loads on demand to perform a repeatable task. See Section 6.
SLM	Small Language Model — a compact model, often distilled, capable enough for many tasks at far lower cost and latency.
Test-time compute	Computation spent at answer time (e.g. extended reasoning) rather than only during training, to improve results.
Token / tokenisation	Sub-word chunk a model reads/writes (~¾ word in English). The unit of context, cost, and latency.
Tool use / function calling	A model calling external functions or services (search, code, APIs) to act beyond generating text. Basis of agents.
Transformer	The 2017 neural-network architecture, built on self-attention, underlying essentially all modern LLMs.
Vector database	A store optimised for embeddings and fast similarity search; the retrieval engine behind RAG.
Weights	See Parameters — the trained numbers that constitute the model.
Zero-shot / few-shot	Prompting with no examples (zero-shot) or a handful of examples (few-shot) to steer behaviour without training.

Go deeper: free courses and videos for these foundations — including 3Blue1Brown and Karpathy — are catalogued in the Stage 1 module, Section 17.3.

Next section: Section 2 — The Model Landscape, the shared reference for the frontier and open-weight models the rest of the guide relies on.

2. The Model Landscape

↗ *FAST-MOVING SECTION — last verified 29 July 2026. Model names, prices, and rankings shift monthly. Treat every figure here as a dated snapshot and re-verify against the live trackers listed in Section 2.6 before relying on it.*

This is the shared reference the rest of the guide points back to. It maps the models worth knowing in mid-2026 — the proprietary frontier and the open-weight field that is rapidly catching it — with the numbers that actually drive decisions: context window, price, modalities, and where each model genuinely leads. If one idea survives every future refresh of this section, it is this: there is no single best model. The frontier is a set of specialists, and choosing well means matching model to task (and to budget).

2.1 How to read this section

Four cautions before the tables, because they determine how much weight to put on any number here:

- **Benchmarks are directional, not definitive.** Headline scores are contaminated by training-data leakage on older tests, and they depend heavily on the surrounding scaffold — the same model scores differently inside different harnesses. Saturated tests (MMLU) no longer separate the top models; the more discriminating current signals are GPQA Diamond (science reasoning), SWE-bench Verified / Pro (coding), AIME and FrontierMath (maths), ARC-AGI-2 (abstract reasoning), and Arena Elo (human preference). The only reliable signal for your job is testing on your own representative data.
- **Advertised context ≠ usable context.** Independent testing (e.g. NVIDIA's RULER) finds models reliably use only roughly 50–65% of their stated window before recall degrades. A “1M-token” model is not a promise of clean recall across a million tokens.
- **Headline price is a starting point.** Real cost is shaped by prompt caching (often 75–90% off repeated input), Batch APIs (≈50% off), long-context surcharges (several providers roughly double rates above ~200–272K tokens), and data-residency uplifts. The per-million figures below are standard list prices in USD; Section 12 covers the true economics.
- **Naming conventions differ by lab.** OpenAI uses decimal increments and dated snapshots (GPT-5.6, gpt-5.6-2026-07-09); Anthropic uses named tiers within a generation (Opus > Sonnet > Haiku, e.g. Opus 5); Google uses generation markers plus tiers (Gemini 3.1 Pro / Flash / Flash-Lite); xAI uses point releases (Grok 4.5). A “.1” bump can still be a large capability jump.

DELIVERY / PM Don't hard-wire a model

Because the leader changes every few weeks, build systems that can swap models with a config change, not a rewrite — an abstraction layer or a routing gateway. Re-verify this section quarterly, and re-run your evals when you switch: a cheaper or newer model that scores well on paper can still regress on your specific workload.

2.2 The frontier flagships (proprietary)

As of mid-2026, four labs hold the top of the proprietary frontier: OpenAI, Anthropic, Google, and xAI — though Meta re-entered the proprietary race in July with Muse Spark 1.1

(\$1.25/\$4.25, 1M context), ending its open-only era. All four are reasoning-capable, natively multimodal to varying degrees, and offer very large context windows. Prices are standard list, USD per million tokens (input / output).

Model (provider)	Context	Price in/out	Inputs	Best known for
GPT-5.6 (OpenAI)	1M	\$5 / \$30	text, image	Agentic coding, computer use, all-round work
Claude Opus 5 (Anthropic)	1M	\$5 / \$25	text, image	Top of intelligence index; coding, agents, reliability
Gemini 3.1 Pro (Google)	1M	\$2 / \$12*	text, image, audio, video	Reasoning & data analysis; best value; multimodal
Grok 4.5 (xAI)	1M	\$2 / \$6	text, image, video	Opus-class; coding, agentic work, tool use

* Gemini 3.1 Pro steps to \$4 / \$18 above ~200K tokens; GPT-5.6 and Grok 4.5 similarly apply higher rates on very long prompts. Context figures are advertised maximums — see the effective-context caveat above.

GPT-5.6 (OpenAI) — the Sol/Terra/Luna family reached general availability on 9 July 2026, after a US government safety review, superseding GPT-5.5 as the flagship line. The flagship Sol is \$5/\$30 with a roughly 1M-token window (a long-context premium kicks in above ~272K); the mid-tier Terra (\$2.50/\$15) and small Luna (\$1/\$6) complete the ladder. Built for agentic work: coding in Codex, computer use, and tool-heavy multi-step tasks, and markedly more token-efficient than its predecessor. Knowledge cutoff around December 2025.

Claude Opus 5 (Anthropic) — released 24 July 2026 at unchanged Opus pricing (\$5/\$25) with a 1M-token window and, as of this writing, narrowly on top of the Artificial Analysis Intelligence Index. Strengths are coding, long-horizon agentic tasks, computer/browser use, and a specific reliability focus (fewer unnoticed flaws in its own output). It uses adaptive thinking with five effort settings (low → xhigh → max), and is now the default on Claude Max plans — at roughly half the price of Fable 5. The tier ladder matters: Sonnet 5 (\$3/\$15, 1M context — introductory \$2/\$10 through 31 August 2026) delivers close to Opus quality for most work and is the everyday workhorse; Haiku 4.5 (\$1/\$5, 200K) handles high-volume, low-complexity tasks. Anthropic's higher tier is now public as Claude Fable 5 (\$10/\$50), the safeguarded release of the Mythos-class model; the unrestricted Claude Mythos 5 remains limited to participants in Anthropic's Project Glasswing. (This guide is compiled with Claude, so that lineage is worth stating plainly.)

Gemini 3.1 Pro (Google) — released February 2026, still carrying a “preview” label. The value leader at the frontier: \$2/\$12 up to ~200K tokens with a 1M window (Google positions some tiers at up to 2M), and the broadest native multimodality — text, image, audio, and video in. It leads much of the reasoning field (GPQA Diamond ~94%, ARC-AGI-2 ~77%) and is exceptional for long-document and data-analysis work. Cheaper Flash and Flash-Lite tiers cover high-volume and budget needs. The caveat is status: preview-tier SLAs and rate limits may not suit production that needs contractual guarantees. A Gemini 3.5 Pro remains expected but unconfirmed — a rumoured 17 July date passed without release — while Google shipped Gemini 3.6 Flash (\$1.50/\$7.50, 1M) plus 3.5 Flash-Lite and Flash Cyber variants on 21 July, and openly teased Gemini 4.

Grok 4.5 (xAI) — released 8 July 2026 and described by xAI as an “Opus-class” model, at \$2/\$6 with a 1M window and native video input, superseding Grok 4.3. It is reasoning-centric with strong science reasoning and high factual accuracy, and now targets agentic-coding and tool-use workloads more directly. It still ships with less public safety documentation than its peers — relevant for risk-averse or regulated use.

2.3 The cost tiers within each family

Each lab sells a ladder, not a single model: a flagship for the hardest work, a mid tier that handles the majority of real tasks, and a small/fast tier for high-volume simple work. Matching the task to the lowest tier that does the job well is the single biggest cost lever in applied AI (Section 12 quantifies it). Indicative standard list prices, USD per 1M tokens (input / output):

Tier	Anthropic	OpenAI	Google	xAI
Flagship	Opus 5 \$5/\$25	GPT-5.6 Sol \$5/\$30	Gemini 3.1 Pro \$2/\$12	Grok 4.5 \$2/\$6
Mid / workhorse	Sonnet 5 \$3/\$15	GPT-5.6 Terra \$2.50/\$15	Gemini 3.6 Flash \$1.50/\$7.50	Grok 4.1 Fast \$0.20/\$0.50
Small / budget	Haiku 4.5 \$1/\$5	GPT-4.1 nano \$0.10/\$0.40	Flash-Lite \$0.25/\$1.50	—

The spread from budget to flagship is often 20–50× on output tokens. A routing strategy — cheap tier by default, escalate to flagship only when a task needs it — routinely cuts total spend 50–70% with little quality loss.

2.4 The open-weight frontier

The most consequential structural shift of the past year is that open-weight models have become genuinely competitive. On many coding and reasoning tasks the gap to the proprietary frontier has narrowed to roughly 5–15 points, and for some coding work it has effectively closed. Two things are worth internalising. First, most of the current open-weight leaders come from Chinese labs — Z.ai/Zipu, DeepSeek, Moonshot, Alibaba — with Meta and Mistral the main Western presences. Second, “open-weight” is not “open-source”: the trained weights are downloadable under a licence, but the training data and full pipeline usually are not.

Model (provider)	Architecture	Context	Licence	Best known for
GLM-5.2 (Z.ai / Zipu)	~744B MoE	1M	MIT	Frontier open coding & agents; close second overall
DeepSeek V4 (Pro / Flash)	Large MoE	1M	MIT	Price leader; strong reasoning & competitive coding
Kimi K3 (Moonshot)	2.8T MoE (104B active)	1M	Kimi K3 licence	New open-weight intelligence leader; agentic coding
Qwen3.x (Alibaba)	Dense & MoE	up to 1M	Apache 2.0	Accessible/self-host-friendly; multilingual
Llama 4 (Meta)	MoE (Scout/Maverick)	very large	Custom*	Broad ecosystem & tooling
Mistral (Large 3 / Small 4)	Dense & MoE	large	Apache 2.0	European; strong fine-tuning ecosystem

** Llama's licence is permissive at small scale but carries a ~700M monthly-active-user cap and EU-use restrictions — read it before commercial deployment. Smaller variants (e.g. Qwen small MoE, Mistral Small 4) run on a single high-end GPU; the largest open models need multi-GPU clusters (roughly 4–8× H100/H200).*

Briefly: Kimi K3 — a 2.8-trillion-parameter MoE released 16 July, weights opened 27 July — is the new open-weight intelligence leader, and claims wins on several agentic-coding benchmarks (vendor-reported). GLM-5.2 remains the strongest established case for high-end open coding. DeepSeek V4, generally available since 20 July (Flash \$0.14/\$0.28; Pro \$0.435/\$0.87; both 1M context), is the price story — Flash undercuts virtually everything on hosted API cost. Qwen is the most practical to self-host and the multilingual pick. Llama 4

brings the deepest tooling ecosystem. Mistral is the European option with mature fine-tuning support.

SECURITY Where do the weights — and your data — live?

Using a Chinese lab's hosted API can route your prompts through servers subject to China's National Intelligence Law. The open weights themselves are MIT/Apache-licensed, so self-hosting the same model inside your own boundary removes that data-path concern while keeping the cost and control upside. For regulated or sensitive workloads in the US, EU, or Australia — especially under EU data-transfer rules — this is frequently the deciding factor — the model can be excellent and the hosted API still be the wrong deployment.

GOVERNANCE Licences vary sharply — read them

MIT and Apache 2.0 (GLM, DeepSeek, Qwen, Mistral) are clean for commercial use and fine-tuning. Others carry real conditions: Llama's MAU cap and EU carve-outs, and various "modified" licences with redistribution or training-on-output limits. The label "open" is not a substitute for reading the terms before you build on or redistribute a model.

2.5 The comparative picture — no single winner

Pulling it together, here is where each strength currently sits. These standouts move with every release; use the category, not the specific name, as the durable takeaway.

Category	Current standout(s) — mid-2026
Agentic coding	Claude Opus 5 and GPT-5.6 (neck-and-neck); Kimi K3 leads open-weight
Reasoning & science	Gemini 3.1 Pro leads; Opus 5, GPT-5.6, Grok 4.5 close behind
Writing & natural prose	Claude and GPT-5.6 both strong; task-dependent
Price-performance	Gemini 3.1 Pro (closed); DeepSeek V4 Flash (open)
Largest usable context	1M+ across all flagships; Grok Fast tiers advertise 2M
Multimodal (incl. video)	Gemini 3.1 Pro (text/image/audio/video); Grok 4.5 native video
Agentic reliability / tool use	Opus 5, GPT-5.6 (closed); Kimi K3 (open)
Open-weight intelligence	Kimi K3 (GLM-5.2 close behind)
Local / single-GPU	Qwen small MoE, Mistral Small 4, Gemma-class models

The professional posture that follows from this is model-agnostic: know the matrix, route each task to the cheapest model that clears your quality bar, keep a fallback provider for resilience, and let evaluations on your own data — not leaderboards — make the final call.

2.6 Reference notes & how to re-verify

- **Effective context:** plan around ~50–65% of a model's advertised window for reliable recall; place the most important material at the start and end of long prompts.
- **Benchmark hygiene:** prefer SWE-bench Verified/Pro over saturated tests, always note the scaffold, and treat vendor-reported numbers as a ceiling. Independent aggregators are more trustworthy than first-party claims.
- **Pricing modifiers:** layer prompt caching (75–90% off repeated input) and Batch APIs (~50% off) before comparing headline rates; watch long-context surcharges and data-residency uplifts.

To refresh this section, cross-check these live trackers rather than any single article: Artificial Analysis (artificialanalysis.ai) for the composite intelligence/coding indices and price-vs-

performance; LLM-Stats (llm-stats.com) for release timelines and specs; OpenRouter (openrouter.ai) for real, post-caching prices across providers; each provider's own pricing/docs pages for the authoritative rates; and [LMArena](https://lmarena.ai) for human-preference Elo.

Go deeper: per-vendor learning tracks are mapped in Section 17.7 — and re-check standings against the live trackers in Section 2.6 before relying on them.

Next section: Section 3 — Top AI Platforms in Use Today, moving from the underlying models to the products people actually use.

3. Top AI Platforms in Use Today

↗ **FAST-MOVING SECTION** — last verified 29 July 2026. User numbers, features, and pricing change constantly; treat every figure as a dated, source-attributed snapshot and re-verify before relying on it.

Section 2 covered the engines. This section covers the vehicles — the products people and businesses actually open. The distinction matters: a platform is a model plus everything wrapped around it (interface, memory, integrations, tools, admin controls, data-handling terms), and two platforms running the same underlying model can differ enormously in fit. Several platforms are also model-agnostic, letting you pick from a menu.

Three strategic archetypes explain most of the market. Consumer scale — win the individual user (ChatGPT, Meta AI). Ecosystem distribution — embed AI where people already work (Google Gemini, Microsoft Copilot). Enterprise precision — win high-value professional and business workflows on quality, safety, and governance (Claude). Around them sit specialists: search (Perplexity), real-time (Grok), cost (DeepSeek), embedded productivity (Notion AI), and sovereignty (Mistral Le Chat). The durable backdrop: the single-vendor era is over — a large majority of enterprises now run three or more model families at once.

3.1 How to read this section

- **Platform ≠ model.** Judge a platform on the whole package — UX, integrations, security posture, data terms — not just which model it happens to run this month.
- **Popularity numbers are apples-to-oranges.** Weekly active users, monthly active users, web-traffic share, app downloads, enterprise seats, and “embedded reach” all measure different things. A billion embedded users (surfaced inside an app people opened for another reason) is not a billion people who chose the assistant. Figures below are indicative and dated; use them for rough scale, not precision.
- **Pricing follows a pattern.** Most consumer platforms run a free tier, a ~\$20/month “pro” tier, a ~\$200/month power tier, and seat-based enterprise pricing. The real enterprise cost lives in negotiated contracts, not the public page.

GOVERNANCE The three questions to ask before adopting any platform

1) Does it train on your inputs by default, and can you turn that off? 2) What certifications and data-residency options does it offer (SOC 2, ISO 27001, HIPAA, FedRAMP; where is data processed)? 3) What is its status under incoming regulation — notably the EU AI Act’s general-purpose-AI obligations (in force since August 2025), and the Article 50 transparency duties and Commission enforcement powers that apply from 2 August 2026? The answers, not the benchmark scores, usually decide enterprise fit. Sections 13 and 14 go deeper.

3.2 The ten at a glance

Indicative scale figures are drawn from mid-2026 third-party estimates and provider disclosures; they use different metrics, so compare within a column, not across. Consumer prices are USD per month.

Platform	Provider	Best for	Indicative scale	Consumer price
ChatGPT	OpenAI	General-purpose; broadest ecosystem	~900M weekly; 1B+ monthly	Free / \$20 / \$200
Gemini	Google	Search, multimodal, Workspace	~900M app MAU; ~2.5B+ via Search	Free / ~\$20

Platform	Provider	Best for	Indicative scale	Consumer price
Copilot	Microsoft	Enterprise work in M365	~420M reach; 20M+ paid seats	Free / \$20; ~\$30 seat
Claude	Anthropic	Professional writing, analysis, coding, agents	Smaller consumer; leads enterprise deals	Free / \$20 / \$100–200
Meta AI	Meta	In-app consumer assistant	~1B+ monthly (embedded)	Free
Perplexity	Perplexity	Cited research; AI search & browser	~45M+ MAU (est.)	Free / \$20 / \$200
Grok	xAI	Real-time (X data); cheap	~3% web / ~15% mobile share	Free in X; paid tiers
DeepSeek	DeepSeek	Ultra-low-cost; open weights	~4% web share	Free
Notion AI	Notion	Embedded knowledge-work productivity	Tied to Notion's base	\$20/user (Business)
Le Chat	Mistral	European data sovereignty; private deploy	Europe-concentrated	Free / ~\$15

3.3 The general-purpose leaders

ChatGPT (OpenAI) — chatgpt.com. The default AI for most of the world and the mindshare leader, at roughly 900 million weekly users (over a billion monthly), 50M+ paying subscribers, and a \$25B annualised run-rate; some 93% of the Fortune 500 use OpenAI products. Its edge is breadth and ecosystem: custom GPTs, Canvas, image generation (its Sora video tool is winding down — see §5.3), memory, ChatGPT Search, Codex for coding, and a fast-expanding agent layer (including the new ChatGPT Work, launched July 2026). Used by everyone from students to enterprises. Roadmap is squarely agentic — autonomous multi-step tasks and computer use — plus deeper enterprise features.

Best for: *a capable all-rounder and the widest feature set in one place.*

Google Gemini (Google) — gemini.google.com. The fastest-scaling large assistant, powered less by novelty than by distribution: Gemini is built into Android, Chrome, and Workspace, and its answers reach billions through Search AI Overviews. The standalone app has around 900 million monthly users, and around 70% of Google Cloud customers use Gemini. Differentiators are the broadest native multimodality (text, image, audio, video), very large context, and aggressive pricing. Roadmap: always-on agents, Deep Think reasoning, and ever-tighter Workspace integration.

Best for: *search-grounded research, multimodal work, and anyone living in Google Workspace/Android.*

Microsoft Copilot (Microsoft) — copilot.microsoft.com + M365. The enterprise-work play: Copilot lives inside Windows, Edge, Outlook, Word, Excel, Teams, and the rest of Microsoft 365, grounded in your organisation's own data via the Microsoft Graph. Reach is around 420 million with 20M+ paid M365 seats. Its strongest card is compliance — FedRAMP, SOC 2, and HIPAA positioning that regulated industries require — plus Copilot Studio for building custom agents. Roadmap: autonomous agents across every app and deeper Studio tooling.

Best for: *organisations already standardised on Microsoft 365 and needing enterprise compliance.*

Claude (Anthropic) — claude.ai. Smaller in raw consumer numbers but the enterprise-precision leader: it reportedly wins the majority of new head-to-head enterprise deals against OpenAI and shows the highest per-visit engagement of the consumer assistants. Strengths are writing quality, reasoning, coding (Claude Code), long-context analysis, and a safety/governance posture that appeals to risk-aware buyers, surfaced through Projects, Artifacts, the Model Context Protocol, and the Cowork agent app. Roadmap: agentic work (Claude Code, Cowork, computer use) and enterprise depth. (This guide is written with Claude, so treat that as disclosed.)

Best for: professional writing, document and data analysis, serious coding, and governance-sensitive work.

Meta AI (Meta) — in WhatsApp, Instagram, Facebook, Messenger, and Ray-Ban Meta glasses. By raw monthly reach the largest of all — over a billion — because it is embedded in apps billions already open daily, and it runs on Meta’s open-weight Llama models. Free, consumer-first, strong at casual Q&A and image generation, and increasingly a channel for business messaging on WhatsApp. Roadmap: agents in social and business messaging, wearables, and personalisation. Its scale is distribution, not deliberate choice — an important distinction for business planning.

Best for: casual in-app assistance and reaching consumers via WhatsApp/Instagram; light business relevance beyond messaging.

3.4 Specialist & alternative platforms

Perplexity — perplexity.ai. The AI “answer engine”: cited, live-web answers rather than a page of links, with Pro Search, Deep Research, and a Model Council that runs a query across several frontier models and synthesises them. Around 45M+ monthly users and a ~\$22.6B valuation, now pivoting hard into agents via the Comet AI browser and Perplexity “Computer.” Enterprise tiers (\$40–\$325/seat) add internal-knowledge search, SOC 2/HIPAA, and a no-training-on-your-data guarantee. Roadmap: owning the agentic browser — the “front door” where an agent starts (and increasingly completes) a task.

Best for: research, competitive analysis, and fact-checking where sourced citations matter.

Grok (xAI) — grok.com and inside X. Positioned around real-time information (live X/Twitter data), a reasoning-first design, native video input, and the most permissive guardrails of the major assistants. Web share is modest (~3%) but mobile usage is growing fast. Cheap to use via API. Roadmap: voice, multi-agent (Grok Heavy), and deeper X integration. Cautions: less public safety documentation than peers, an image-tool deepfake controversy, and corporate entanglement (xAI is now a SpaceX subsidiary) — relevant for risk-averse buyers.

Best for: real-time/social-data questions and cost-sensitive, less-guardrailed use.

DeepSeek — deepseek.com. The cost disruptor: a capable consumer app and open-weight models at a fraction of Western pricing, strong on reasoning and maths. Around 4% of web traffic and a genuine force in resetting the market’s price floor. The catch is governance: as a Chinese-developed service its hosted app may be subject to China’s data laws, it carries no publicly verified enterprise security certifications, and it applies content censorship — all significant for regulated use in any of the three regions this guide serves. The open weights can, however, be self-hosted inside your own boundary to sidestep the data-path concern.

Best for: *cost-sensitive or self-hosted workloads — with clear-eyed governance review before hosted use.*

Notion AI — notion.so. AI embedded where a team's knowledge already lives — docs, wikis, projects, meeting notes — with a model picker spanning a dozen models (Claude, GPT, Gemini, Grok, Kimi, DeepSeek, plus Notion's own) and autonomous agents that build databases and draft pages. Included in the Business plan (~\$20/user/month); custom agents bill on credits. Notably strong on governance: contractual no-training-on-your-data with every model provider, permission-aware retrieval, and enterprise zero-retention. Roadmap: agents as “teammates” across the workspace.

Best for: *teams that run on Notion and want AI over their own knowledge base with clean data terms.*

Mistral Le Chat (Mistral AI) — chat.mistral.ai. Europe's flagship assistant and the sovereignty option: it does not route data to China or the US by default, and its Enterprise tier deploys on private infrastructure, VPC, or on-premises. Backed by open-weight models (an exit ramp from vendor lock-in) and paired with Studio (production/LLMOps tooling) and Forge (train-your-own-model). Roadmap: enterprise agents and production tooling. For organisations weighing data residency and sovereignty — a procurement requirement in parts of the EU public sector, and a growing expectation in Australian and US government work — it is the clearest European reference point.

Best for: *privacy- and sovereignty-sensitive deployments and teams wanting private/on-prem options.*

SECURITY Two live security watch-items

Chinese-developed hosted assistants (DeepSeek, and Doubao/ERNIE in the China market) route prompts through servers subject to China's National Intelligence Law — fine for low-sensitivity use, risky for regulated or confidential data unless self-hosted. Separately, agentic browsers (Perplexity's Comet, and similar) that read and act on live pages have already been shown vulnerable to indirect prompt injection and data-exfiltration attacks — do not point them at your bank or confidential systems. Both themes recur in Section 14.

3.5 Also worth knowing

- **Coding-specialist platforms (Cursor, Claude Code, GitHub Copilot, Windsurf)** — among the most-used and highest-valued AI products of all, but deliberately held for Section 7 (Coding & Developer AI) rather than duplicated here, where the focus is general-purpose assistants.
- **Character.AI** — the giant of the companion/roleplay category, with some of the highest session times in all of consumer AI. Low direct business relevance, but a reminder that “AI usage” is not only productivity.
- **Doubao (ByteDance), ERNIE (Baidu), Qwen Chat (Alibaba)** — the dominant consumer assistants inside China; important for anyone operating in or selling to the APAC market.
- **Multi-model aggregators (Poe, and similar)** — one subscription, many models side by side; useful for comparison and for avoiding single-vendor lock-in.
- **Cloud AI platforms (AWS Bedrock, Google Vertex AI, Microsoft Azure AI Foundry, IBM watsonx)** — not chat apps but the enterprise build-and-deploy layer:

access to many models behind your cloud's identity, networking, monitoring, and governance, plus native agent-orchestration tooling. This is where most production business AI is actually assembled — covered in Sections 6 and 9.

3.6 Choosing a platform

A rough decision guide, by situation:

- **Solo / freelancer:** one general-purpose assistant (ChatGPT or Claude), add Perplexity for research.
- **Small team:** the same, plus an embedded tool where your work lives (Notion AI, or Copilot/Gemini if you're on Microsoft/Google).
- **Mid-size / enterprise:** standardise on the platform that matches your existing stack and compliance needs (Copilot for Microsoft shops, Gemini for Google, Claude for governance-sensitive professional work), with enterprise tiers, SSO, and approval workflows.
- **Regulated / sovereignty-sensitive (incl. many public-sector and finance contexts across the US, EU, and Australia):** prioritise data residency, certifications, and no-training guarantees — evaluate enterprise Claude/Copilot, Le Chat, or self-hosted open-weight before hosted consumer tools.

DELIVERY / PM Manage the platform, not just the model

Two delivery realities. First, "shadow AI" is already happening — staff paste company data into free consumer tools daily; a sanctioned platform with proper data terms is the practical fix, not a ban. Second, keep portability in mind: standardising is efficient, but avoid designs that make switching providers a rewrite, because the leader will change again.

Go deeper: where to learn each platform properly — start-here and go-deeper picks per vendor — is in Section 17.7.

3.7 The tooling layer — what each platform actually ships

Section 3.3 introduced each platform as a product; this closes the section by looking across them at their tooling — the surfaces and capabilities wrapped around the model that decide what you can actually do with it. For the three frontier platforms most professionals standardise on — Claude, ChatGPT and Gemini — the model is rarely what separates them; the tooling is. It is also the most perishable material in the guide: the categories below are durable, but the product names inside them change almost monthly, so read the table as a mid-2026 snapshot and re-verify before relying on any specific name.

Five lenses cover most of what matters: the access surfaces you reach the model through; the agentic layer that lets it do work rather than just answer; the connection layer that plugs it into your data and apps; the organising primitives that turn one-off chats into repeatable workflows; and the admin and governance tooling that enterprise buyers actually evaluate.

Tooling category	Claude (Anthropic)	ChatGPT (OpenAI)	Gemini (Google)
Access surfaces	Web, mobile, and a native desktop app (Mac/Windows/ChromeOS) that bundles chat, Cowork and Claude Code in one window,	The most ubiquitous consumer surface — web, mobile and a desktop app, with Codex now merged in.	Web and mobile app, plus deep embedding in Android, Chrome, Search and Workspace, so you meet it inside tools you already use.

	synced across devices.		
Agentic layer — does work, not just answers	Cowork, a visual agent in the desktop app that runs in an isolated local VM with no terminal, and Claude Code for agentic coding in the terminal or IDE.	ChatGPT Work (a goal-to-deliverable agent on GPT-5.6 with Plan mode and approvals) and Workspace Agents (shared, Codex-powered team agents that run in the background).	Antigravity, an agent-first dev platform and 2.0 desktop app that absorbed the Gemini CLI, plus Jules, an async coding agent that opens pull requests.
Connection layer — reaches your data and apps	MCP, the open standard Anthropic authored, via remote connectors and local desktop extensions, plus the Claude in Chrome browser agent and computer use.	Connectors to Slack, Google Drive, Microsoft, Salesforce and Notion, MCP support, and agent and computer use.	Native access to Workspace and Google apps, with Managed Agents and MCP exposed through the Gemini API.
Organising primitives — make work repeatable	Projects, Artifacts, Skills and memory.	Custom GPTs (being succeeded by Workspace Agents for business) and memory; the Canvas surface has been folded into inline text and code.	Gems (custom assistants) and Deep Research, with NotebookLM alongside.
Admin and governance — what enterprises buy on	SSO, audit logs, data-retention controls and no-training terms — a deliberately governance-forward posture.	Enterprise admin, SSO, compliance certifications and workspace controls.	The Workspace admin console, data-residency options, and Google Cloud identity and governance.

The pattern the table makes visible is that differentiation has moved off the model and onto the surface. All three run comparably capable frontier models (Section 2); what divides them in daily use is whether the agent runs where your files already are, how cleanly it reaches the rest of your stack, and how much governance sits around it. A Microsoft or Google shop inherits Copilot or Gemini tooling by gravity, whereas Claude's desktop bundle of Cowork and Claude Code suits anyone who wants an agent working directly against local files.

Two cautions. This is the fast-moving corner flagged at the top of the section, so confirm the specific names against each vendor's own product pages before committing to them. And the deeper mechanics live elsewhere: agent internals and the MCP standard are covered in Section 6, and the coding tools named above — Claude Code, Codex, Antigravity and Cursor — get their full treatment in Section 7.

Next section: Section 4 — Prompt Engineering & Configuration: getting markedly better results from any platform, with a dedicated thread on producing high-quality documents.

4. Prompt Engineering & Configuration

Durable craft, current to the reasoning-model era. Claude product features verified 29 July 2026 — check support.claude.com for the latest.

This section is the practical craft of getting markedly better results from any model, plus how to configure Claude specifically — and it closes with a dedicated playbook for producing high-quality documents, since that is the most common real-world use. The single biggest shift since 2023: prompting stopped being about magic incantations and became about clear communication and managing what information the model has in front of it.

4.1 The mindset: clarity over cleverness

The most useful thing to internalise first: most prompt failures come from ambiguity, not model limitations. Newer models follow instructions better and need less ceremony, but they cannot remove the ambiguity in your own request. Better models raise the ceiling; they do not replace your judgement about role, scope, input quality, output shape, and how you'll check the result.

The most reliable mental model is that you are briefing a brilliant new hire who has zero context on your project. They are capable, but they only know what you tell them — so the clearer and more specific you are, the better they perform. For everyday use, the whole craft reduces to three habits: be clear, be specific, and give context. Most “secret prompt hacks” are noise next to those three.

4.2 The core techniques

A durable toolkit, roughly in order of return on effort:

- **Be clear and direct.** State the task, the audience, and the constraints explicitly. If an instruction could be read two ways, spell out which you mean. If you want action rather than suggestions, say “make these edits,” not “can you suggest changes?”
- **Give context.** Explain the background, the purpose, and who the output is for. “Write this for a non-technical board” and “write this for engineers” produce very different, and better, results than no audience at all.
- **Show examples (multishot).** One or two examples of the input-and-ideal-output is among the highest-ROI moves for consistency and format. With strong current models, try zero-shot first — they often infer intent — and add examples when output drifts.
- **Assign a role.** A single sentence in the system prompt (“You are a senior policy analyst...”) focuses tone and behaviour for the whole conversation.
- **Specify the output.** Give the structure, length, and format you want. For machine-readable output, most platforms now offer a structured-output / JSON-schema mode that guarantees valid JSON.
- **Decompose complex tasks.** Break a big job into steps, or chain prompts (outline → approve → draft → refine). Smaller, checkable steps beat one giant instruction.
- **Allow “I don’t know.”** Explicitly permitting the model to decline when unsure reduces hallucination — especially important when working from supplied documents.
- **Iterate.** Treat it as a conversation. The second and third turns — “shorter,” “more formal,” “add a risks section” — are where quality is won.

4.3 Structuring prompts with XML tags

Claude in particular was trained to recognise XML-style tags, which makes them the cleanest way to separate the parts of a complex prompt so nothing gets mixed up. Wrap each kind of content in its own descriptive tag — instructions, context, examples, source documents, and the desired output format. A reusable skeleton:

```
<role>Senior policy analyst</role>
<context>Audience: our executive board. They are non-technical.</context>
<task>Summarise the attached report and recommend next steps.</task>
<source_material>
  {{paste the document(s) here}}
</source_material>
<instructions>
  Use only the source material. If something isn't covered, say so.
</instructions>
<output_format>
  1. Executive summary (3 sentences)
  2. Key findings (bulleted)
  3. Recommended actions
</output_format>
```

Two placement rules matter. Put context before the task, so the model knows the frame before the instruction. And put your most important instructions at the very start and the very end of a long prompt — models attend best to the beginning and end and can lose material buried in the middle (the “lost in the middle” effect from Section 1.3).

4.4 What changed with reasoning models

The rise of reasoning models (Section 1.5) quietly retired one of 2023’s signature tricks. Because these models already work through problems step by step internally, adding “think step by step” or “show your work” often duplicates effort and can degrade the answer. With a reasoning model, give it the task and the constraints, then get out of the way.

- **Keep explicit chain-of-thought for cheaper, non-reasoning models** (the fast/budget tiers), where prompting the steps still lifts accuracy on multi-step work.
- **Use reasoning-effort controls instead of prompt tricks.** Where a model exposes effort levels (low/medium/high), dial thinking up for hard analysis and down for routine work — it’s a cost and latency lever as much as a quality one.
- **Match the technique to the tier.** The right prompt for a flagship reasoning model and for a cheap classifier model are different; don’t port prompts blindly between them.

4.5 From prompt engineering to context engineering

The larger 2026 shift is that context engineering — deciding what information is in front of the model, and in what shape — has become the higher-leverage skill. The goal, in Anthropic’s own framing, is the minimal set of information that fully outlines the expected behaviour. Note that minimal does not mean short: you still need enough for the model to act correctly, just nothing superfluous.

Practically: don’t dump everything into the prompt. Reasoning quality degrades as prompts grow, and material buried mid-context gets under-used, so retrieve, filter, compress, and structure rather than paste indiscriminately. Aim system prompts at the “right altitude” — specific enough to guide behaviour, general enough not to be brittle if-else logic. When the

model needs a lot of external knowledge, that's a retrieval (RAG) problem, covered in Section 6, not a bigger-prompt problem.

DELIVERY / PM Treat prompts and context as infrastructure

For anything you'll reuse, don't hand-tweak prompts ad hoc. Keep a small test set of representative inputs with known-good outputs, version your prompts, and re-run the test set whenever you change one — prompts break in surprising ways, and a one-word edit can shift three other behaviours. This is regression testing for prompts, and it's the difference between a demo and something you can rely on. Context pipelines also give you an audit trail of what information influenced each answer — which matters for governance (Section 13).

Go deeper: evaluation — the discipline for testing non-deterministic systems that this section keeps invoking — is treated in full in Section 9.10.

4.6 Model settings & knobs

The main dials, most set once per app or per conversation (fuller treatment in Section 1.9):

- **System prompt** — the standing role and rules for the whole conversation.
- **Temperature / sampling** — low for focused, repeatable work (extraction, code, factual); higher for creative variety.
- **Max output tokens** — a length ceiling and a cost/latency control.
- **Reasoning effort** — how hard a reasoning model thinks before answering.
- **Structured output / JSON mode** — forces valid, schema-conformant output for downstream systems.
- **Stop sequences** — halt generation at a marker, useful for structured output.

4.7 Configuring Claude specifically

Beyond the prompt itself, Claude offers configuration that saves you re-specifying the same things every time. Features evolve, so treat this as an orientation and confirm current details at support.claude.com.

- **User preferences** — standing instructions for tone, formatting, and behaviour that apply across your conversations (e.g. “be concise,” “no bullet points unless asked”).
- **Styles** — reusable writing-style profiles you can switch between, so output matches a house voice without re-describing it each time.
- **Skills** — reusable packages of task-specific instructions (a SKILL.md plus optional scripts and resources) that Claude loads on demand to carry out a repeatable procedure consistently; Claude Code and the Cowork agent app run on them (see Section 6.7).
- **Projects** — a persistent workspace that holds context, reference documents, and custom instructions for a body of work, so every conversation in it starts already briefed. (This guide, for instance, is being built inside one.)
- **Artifacts** — standalone outputs (documents, code, designs) that render alongside the chat and can be iterated on turn by turn.
- **Memory & past-chat reference** — optionally let Claude carry context across conversations and search prior chats, so you don't re-explain your situation.

- **Tools** — web search and deep research for current information, code execution and file creation (how this document is produced), and connectors/MCP to bring in your own apps and data.
- **Console prompting tools (for builders)** — the prompt generator (beats the blank page), prompt improver (adds structure and examples automatically), templates and variables, and an evaluation tool for testing prompts at scale.

GOVERNANCE What persists, persists

Preferences, Projects, memory, and connected data all carry information forward — which is the point, but it means anything sensitive you put there sticks around. Keep secrets, credentials, and regulated data out of standing configuration unless the data terms and controls justify it, and know your platform's training-and-retention terms before connecting real systems (Sections 13 and 14).

4.8 Getting great documents out of AI

Because document creation is the most common real use, here is a concrete playbook. Most weak AI documents come from under-briefing — asking for “a report on X” and hoping. Strong ones come from giving the model a proper brief, source material, and a review loop.

- **Brief it like a commission.** State audience, purpose, length, tone, structure, and what must be included. A quick brief template: audience / purpose / length / tone / required sections / sources to use.
- **Supply the source material and constrain it.** Paste or attach the facts and add “use only this; if something isn't covered, flag it rather than inventing.” This is the single biggest defence against confident fabrication.
- **Give a model of “good.”** An example document or a filled-in template teaches structure and voice faster than any description.
- **Outline first, then draft.** Ask for a structure, approve or adjust it, then have it write. This prevents wasted long drafts that miss the mark, and it's exactly how this guide is being produced — section outline agreed, then built and reviewed one at a time.
- **Iterate section by section with specific direction.** “Tighten the intro,” “make section 2 more concrete,” “add a risks table” — targeted revisions beat regenerating the whole thing.
- **Lock house tone once.** Put recurring voice and formatting rules into a Style or preferences, and reusable reference material into a Project, so you don't re-specify them every time.
- **Verify the facts.** Treat figures, quotes, and citations as unverified until checked. Ask for sources, and never paste an AI-stated statistic into something that matters without confirming it.

SECURITY Pasted and retrieved content is data, not commands

When you paste a web page, email, or document into a prompt — or an agent fetches one — any instructions hidden inside that content can hijack the model's behaviour. This is prompt injection, the defining security risk of AI systems. A safe habit even in everyday use: keep untrusted source material clearly separated (e.g. in its own XML tag) and be alert that “instructions” arriving inside data are not your instructions. Section 14 covers the defences in depth.

Go deeper: the prompting guides, interactive tutorial and courses behind these techniques are in Sections 17.3–17.4.

4.9 When to fine-tune — the third adaptation lever

Section 1.8 named three ways to adapt a model to your needs — prompt it, retrieve for it, or fine-tune it — and the guide has treated the first two at length (this section, and retrieval in Sections 6.6 and 8.6). Fine-tuning is the third and the most misunderstood: it continues a model's training on your own examples to change how it behaves. The single most useful rule, borne out repeatedly in 2026 practice, is that fine-tuning is for form, not facts. Use it to lock in a voice, a house style, a strict output schema, or a refusal pattern that prompting cannot hold reliably — not to teach the model knowledge that changes, which is what retrieval is for.

The disciplined sequence is prompt → retrieve → fine-tune → distil, tried in that order, because each step costs more and moves slower than the last. Retrieval is the right first answer for the majority of enterprise use cases: it lets you change the source data without retraining, attribute answers to documents, and swap base models freely. Reach past it to fine-tuning only when prompting and retrieval have genuinely plateaued and one of two things is true: you need to distil a frontier model into a smaller, cheaper, faster one that holds most of the quality at a fraction of the cost per token, or you need behaviour — tone, format, structure — that no prompt reliably enforces at scale. The three techniques worth knowing:

- **Parameter-efficient fine-tuning (LoRA / QLoRA)** — the default. Rather than retrain every weight, it trains a small 'adapter', so a job that once needed a cluster now fits on a single high-memory GPU. This is the cheap, reversible, production-standard path, and it pairs with retrieval rather than replacing it.
- **Supervised fine-tuning (SFT)** — training on curated input-and-ideal-output examples. The workhorse for style, format and domain tone; only as good as the examples you give it, and it needs enough of them to generalise.
- **Reinforcement fine-tuning (RFT)** — the 2026 shift. Instead of imitating reference answers, the model is rewarded by a grader for getting verifiable outcomes right (using algorithms such as GRPO). It suits reasoning, maths and code where correctness can be scored, and needs far less labelled data than SFT — but a reward you specify badly is a behaviour you will get badly.

Two cautions before committing. First, the operating burden: a fine-tune is a living asset that ages as the base model moves under it and as your data drifts, so budget for re-training, evaluation and versioning, not a one-off job. Second, the governance line the guide draws in Section 13.3 — heavy fine-tuning can move you from 'deployer' to 'provider' under the EU AI Act, inheriting the stricter obligations — and a model fine-tuned on your data carries that data's provenance and permission questions (Section 8.4) into the weights, where they cannot be retrieved back out. Prove the need with prompting and retrieval first; fine-tune when the gap that remains is real, repeatable, and worth the upkeep.

Next section: Section 5 — Multimodal & Generative Media: the tools for images, video, audio, voice, and music, and where each currently stands.

5. Multimodal & Generative Media

↗ PARTLY FAST-MOVING — last verified 29 July 2026. Specific tools, versions, and prices change monthly; the categories and trade-offs are more durable. Re-verify names and pricing before relying on them.

Section 1.6 noted that frontier models are natively multimodal for understanding — they can read images, audio, and video. This section is about the dedicated tools that generate media: images, video, voice, and music. Two things define the space in 2026. First, it crossed into production-ready: native audio-with-video, 4K output, multi-shot sequences, and near-photographic realism are now normal rather than novelties. Second, it is a fragmented, multi-model market with no single winner, so the skill is routing each job to the right tool — and it carries the sharpest legal and ethical questions in all of AI.

5.1 How to read this section

- **Model vs product.** As with text, the model (FLUX, Veo, Suno's engine) is distinct from the product around it. Aggregators — fal, Freepik, Krea, Higgsfield, and others — run many models behind one interface, which is how most professionals work, because the leader in each category reshuffles every few months.
- **No single winner — route by job.** Each tool has a niche (text-in-image, physics, motion, vocals). Production teams commonly use two or three and pick per task.
- **Price by usable output.** A model that nails it in one generation beats a cheaper one that needs eight attempts. Measure cost per accepted result, not the list price per image or second.
- **Commercial terms vary sharply.** Licence and commercial-use rights differ by tool and plan, and some open weights are non-commercial. Always confirm before monetising output (see the governance callout below).

5.2 Image generation

The most mature category. The current leaders, and what each is genuinely best at:

Model	Maker	Best for	Commercial / provenance note
Nano Banana Pro	Google	Photorealism, editing, character consistency	Commercial via API; SynthID + C2PA
GPT Image 2	OpenAI	Instruction-following, text-in-image, iterative edits	Commercial via API/ChatGPT; C2PA
Midjourney v8.2	Midjourney	Artistic/aesthetic quality	Commercial on paid plans; no C2PA
FLUX.2	Black Forest Labs	Photorealism, camera accuracy; open weights	API commercial; [dev] weights non-commercial
Seedream 4.5 / 5	ByteDance	Text rendering, 4K, product shots; low cost	Via API; verify terms
Ideogram 3	Ideogram	Typography / text-heavy graphics	Commercial via API
Recraft v4	Recraft	Native vector output (brand/design)	Commercial via API
Adobe Firefly	Adobe	Commercially safe generation	Licensed training; IP indemnification

The practical trick is to match the model to the failure mode: scrambled text → Seedream or Ideogram; faces drifting between generations → Nano Banana; plastic-looking skin → FLUX or Imagen; needs artistic soul → Midjourney; needs legal safety → Adobe Firefly. Then draft

cheap and finish expensive (explore composition on a fast model, re-run the winning prompt on a premium one), and edit rather than regenerate — a targeted AI edit is cheaper than rolling the dice again. Google's Imagen, Alibaba's Qwen Image, and xAI's Grok Imagine round out the field.

5.3 Video generation

The fastest-moving category, and the one that changed most in 2026: native audio (dialogue, ambience, music generated with the video) is now standard, alongside 4K and multi-shot sequences. Approximate per-second API costs shown — subscriptions and credit systems vary.

Model	Maker	Best for	Approx cost / note
Veo 3.1	Google	Premium all-round; 4K + native audio	~\$0.15–\$0.40/sec
Sora 2	OpenAI	Cinematic physics; 15-sec clips	Web/app discontinued Apr 2026; API sunsets 24 Sep 2026
Kling 3.0	Kuaishou	Value; motion, native 4K, multi-shot	~\$0.10/sec
Runway Gen-4.5	Runway	Creative control & in-editor editing (ads)	Credit subscription
Seedance 2.5	ByteDance	Multi-shot storytelling; native audio	~\$0.05–\$0.40/sec via hubs
Wan 2.7	Alibaba	Open-source / self-host	Free (self-host)
Synthesia / HeyGen	—	Corporate avatars, training, localisation	Subscription

Veo 3.1 is the safe premium default; Kling is the value and motion pick; Runway wins on directorial control and editing; Seedance leads multi-shot storytelling. Synthesia and HeyGen are a different animal — avatar/talking-head tools for training and localisation, not cinematic generation. Sora is the cautionary tale: briefly the most talked-about model, its consumer app was discontinued in April 2026 with the API set to follow on 24 September 2026 — a reminder (see the PM callout) not to build a pipeline on a single closed model. A common professional workflow is still to generate silent video and layer voice and sound in post, though native-audio models are closing that gap.

DELIVERY / PM Vendor dependency is a real delivery risk

Sora's discontinuation stranded workflows built solely on it. Keep media pipelines portable across providers — a multi-model hub (fal, Krea, and similar) makes swapping engines a config change rather than a rebuild. And budget by cost-per-accepted-result: draft on cheap/fast models, finish on premium ones, and note that content-moderation rejections can still consume credits.

5.4 Voice & audio

Three sub-categories, each with different leaders:

- **Text-to-speech & voice cloning** — ElevenLabs is the overall leader on realism and clone quality; Fish Audio excels at expressive/character voices; Play.ht at multilingual dubbing; Resemble AI at managed custom voices (and offers deepfake detection); Murf and WellSaid at straightforward corporate narration.
- **Speech-to-text (transcription)** — a separate discipline: Deepgram, AssemblyAI, and OpenAI's Whisper lead, judged on accuracy, speaker diarisation, timestamps, and streaming latency.

- **Editing & cleanup** — Descript for transcript-based editing and overdub, Adobe Podcast for enhancing rough recordings, and Krisp for real-time noise removal on calls.

Voice cloning deserves a specific caution: only clone a voice with the clear consent of its owner, keep records of that permission, and treat the technology's fraud potential seriously (see the security callout).

5.5 Music

Music generation matured fast, and here licensing — not just quality — is the deciding factor:

- **Suno (V5)** — the practical leader: complete songs with vocals from a prompt, plus a Studio (DAW-like) workspace, stems, and voice features. The most feature-rich and the default for most workflows, with a usable free tier.
- **Udio** — strong for ideation and audio fidelity, but its value took a hit during a licensing transition (it settled with Universal and signed with Warner, and downloads were disabled during the move toward a fully licensed version). Better for exploration than a controlled production pipeline right now.
- **ElevenLabs Eleven Music** — the commercial-safety pick: built on licensed training data, with the most realistic vocals, and part of a unified voice-plus-music pipeline. The right default for monetised or client-facing work.
- **Others** — MiniMax/Hailuo music, AIVA (classical/MIDI, DAW-friendly), Stable Audio, and Mureka fill niches from scoring to background tracks.

The pragmatic pattern professionals use: generate with Suno or Udio for personal or non-commercial work where quality decides, and reach for a licensed-data tool for anything monetised or client-facing — the cost of a second subscription is trivial next to a copyright dispute.

Two adjacent categories are maturing fast and worth a line. 3D and asset generation — text- or image-to-3D for game, product and AR/VR pipelines — is roughly where image generation was a couple of years ago: useful for drafts and iteration, not yet production-final. And dubbing and lip-sync — translating and re-voicing video while matching mouth movement — has moved from novelty to working tool for localisation, though the same consent and likeness cautions as voice cloning apply.

5.6 The legal and ethical layer

Generative media raises the hardest governance and safety questions in AI. Two callouts capture what matters most; both connect to Sections 13 and 14.

GOVERNANCE Copyright, commercial rights, and provenance

Training data: most image, video, and music models were trained on scraped web content, which creates copyright uncertainty. A few are trained on licensed data and sold with commercial safety — Adobe Firefly (with IP indemnification) for images and ElevenLabs Music for audio; by contrast Suno and Udio faced major-label litigation, with Udio since settling and moving toward a licensed model.

Commercial-use terms differ by tool and plan — some open weights are explicitly non-commercial — so verify rights before monetising.

Provenance: from 2 August 2026 — this week — the EU AI Act (Article 50) requires AI-generated content to be marked in a machine-readable form and deepfakes to be clearly labelled. The emerging standard pairs C2PA “content credentials” with invisible watermarks (e.g. Google’s SynthID). Adobe, OpenAI, and Google embed C2PA today; Midjourney notably does not. (Systems already on the market before 2 August 2026 have until 2 December 2026 to comply.) The US has no federal equivalent — a patchwork of state deepfake and election-content laws applies instead. Australia’s July 2026 announcement of mandatory national AI standards (Section 13) may yet bring an equivalent, and misleading AI content can already fall foul of consumer-protection and defamation law.

SECURITY Synthetic media is now good enough to defraud

Real cases include a roughly US\$25M theft executed with a deepfake video of a company CFO on a video call, and political robocalls cloning a real voice. The lesson: a familiar face or voice is no longer proof of identity. Treat any request involving money, credentials, or urgency that arrives by voice or video as unverified until confirmed through a separate, trusted channel. Only clone a voice or likeness with explicit consent, and never generate intimate or impersonating content of real people. Section 14 covers organisational defences.

Go deeper: the keep-current directory in Section 17.9 is the sustainable way to track this fast-moving field.

5.7 Getting good output: control, quality and failure modes

The catalogue above is what to pick; the next three subsections are how to use it well. Start with a working model of how these tools behave, because it explains both their magic and their failures. Unlike a language model predicting the next token, image and video models mostly work by diffusion (Section 1.6): they begin from noise and refine a whole frame toward your prompt. That is why the same prompt rarely gives the same result twice, why fine detail is where they struggle, and why control — not raw quality — is the real skill. The failure modes to expect:

- **Text and fine structure.** Small text, logos, hands, teeth and repeated patterns are where artefacts still appear — the model is rendering plausible shapes, not reading a specification. Match the tool to the failure (Section 5.2) and fix the region locally rather than regenerating the whole image.
- **Consistency and identity drift.** Holding the same character, product or style across multiple generations is the hardest routine problem — it is why identity-consistency is a headline feature (Nano Banana among them), and why a reference image beats any amount of prompt description.
- **Temporal coherence (video).** Objects that morph, flicker or break physics between frames are the video equivalent, and they worsen over longer clips — which is why short-clip and multi-shot workflows dominate.
- **Prompt adherence.** Models silently drop or reweight parts of a long prompt. Shorter, ordered, concrete prompts adhere better than sprawling ones.

The craft is in the controls, not the prompt alone. The levers that matter: seeds (fix one to reproduce a result, or vary it deliberately); reference and style images (show, don’t describe); structural conditioning (pose, depth, edge and layout constraints — the ControlNet family — to pin composition); inpainting and outpainting (regenerate only a region, or extend the frame); and image-to-image editing (steer from an existing asset rather than from noise).

Then the durable workflow from Section 5.1: draft cheap, finish expensive, and edit rather than re-roll.

5.8 The economics of generative media

Generation is priced differently from text, and the trap is comparing list prices instead of cost per accepted result (Section 5.1). Three things drive the real bill. Resolution and length: a 4K image or a longer, higher-frame-rate clip costs multiples of a draft, and video scales with duration — the per-second figures in Section 5.3 add up fast across a finished sequence. Iteration count: true cost is list price × attempts-to-acceptance, so a pricier model that lands it first often beats a cheap one you run eight times. And the workflow tax: most professional output is assembled — draft on a fast model, upscale, edit, and for video layer voice and sound in post — so the model invoice is only part of the cost, exactly as it is for agents (Section 12.2). The practical posture: measure cost per accepted asset, generate cheap and finish expensive, reuse reference material and seeds, and reach for the premium model only for the shot that ships. Credit and subscription pricing obscures all of this, so benchmark on your own recurring jobs, not the headline rate.

5.9 Deploying generative media safely at scale

Section 5.6 covered the legal layer; this is its operational counterpart — what changes when generation moves from one person's desk to a pipeline. Four concerns dominate, and each connects back to Sections 13 and 14.

- **Provenance and labelling.** The field has converged on a two-layer standard: C2PA Content Credentials — a cryptographically signed manifest, ratified as ISO/IEC 22144 and now the de facto provenance language of the web — plus an invisible watermark such as Google's SynthID, with over 100 billion assets watermarked by mid-2026. OpenAI signs its Sora and image output with both, Microsoft adds C2PA to Microsoft 365 content, and devices like the Pixel 10 sign at capture. Preserve credentials through your pipeline and label AI media: it is an EU AI Act Article 50 duty (Section 13.3) and, in the US, a California SB 942 obligation, not just good manners. The caveat: a missing credential proves nothing — most files were never signed — and metadata is easily stripped by a screenshot, resize or re-upload, so treat provenance as a signal, not proof.
- **Brand safety and consistency.** At volume the risk is off-brand or off-spec output reaching a customer. Lock a house style (reference and style images, and where supported a brand-tuned model), keep an approved-asset library, and route anything customer-facing through review.
- **Moderation and misuse.** Generation tools can produce harmful or infringing content; the controls are input and output moderation, restricted prompts, likeness-and-consent policies (Section 5.4), and clear takedown paths.
- **Human review.** The through-line: a named person accountable for anything published. Volume makes this harder and more necessary, not less — the failure mode is an unreviewed pipeline shipping a fabricated or infringing asset at scale.

Taken together, Sections 5.7–5.9 are the difference between generating impressive one-offs and running generative media as dependable infrastructure — the same shift from demo to production that the rest of the guide keeps returning to.

Next section: Section 6 — Agents: Building, Configuring & Connecting. What agents are and why, how to configure and host them, the MCP tool ecosystem, and connecting AI to your own data (RAG).

6. Agents: Building, Configuring & Connecting

Concepts are durable; specific frameworks, protocols, and platforms move fast — tool details verified 29 July 2026. This is the guide's deepest section: it absorbs the tool-connection standard (MCP) and connecting AI to your own data (RAG).

If Sections 1–5 were about models and how to prompt them, this section is about letting AI do things — take multi-step actions, use tools, and work over your own data with a degree of autonomy. This is the frontier where most of the 2026 investment and most of the 2026 risk now sits. The section builds up in order: what an agent is, when to use one, what it's made of, the frameworks to build with, the standards for connecting tools (MCP) and data (RAG), how to host and govern agents, and — at length — how they fail and how to secure them.

6.1 What an agent actually is

The word “agent” is overused, so a precise definition helps. An agent is a system in which a model plans and takes actions in a loop, using tools, to pursue a goal — deciding for itself what to do next rather than following a fixed script. That distinguishes it from a chatbot (which responds to each message) and from a workflow (which runs predetermined steps). The core mechanism is a loop, often called ReAct (reason + act):

```
GOAL:  "Book a meeting with Sam next week"
      |
      v
[1] PLAN    decide the next step
[2] ACT     call a tool (e.g. check_calendar)
[3] OBSERVE read the result
[4] REPEAT  loop until the goal is met, then respond
```

Autonomy is a spectrum, not a switch: from an assistant that only suggests, to a fixed workflow that chains steps, to an agent that chooses its own steps, to multi-agent systems where several agents coordinate. More autonomy means more capability — and more ways to go wrong.

6.2 When to use an agent — and when not to

Agents are the right tool when a task is genuinely dynamic: multi-step, requiring real actions or external tools, where the exact path isn't known in advance and the system must adapt as it goes. They are the wrong tool when a single well-crafted prompt or a fixed workflow would do the job — because an agent multiplies cost, latency, and unpredictability. A widely-shared rule of thumb: for a single agent that calls only one or two tools, skip the agent framework entirely and use a plain API call with structured output. Reach for an agent only when the looping, tool-use, and adaptation actually earn their keep.

DELIVERY / PM Agents multiply capability and failure in equal measure

Every step in an agent loop is one or more model calls, so cost and latency scale with the number of steps, and small per-step error rates compound over a long run. Start with the simplest thing that works — a prompt, then a fixed workflow — and escalate to an agent only when a real requirement demands it. When you do, budget for observability, evaluation, and human-approval gates from the start, not as an afterthought (Sections 6.11–6.12).

6.3 The anatomy of an agent

Every agent is assembled from five parts:

- **Model** — the reasoning engine that plans and decides. Reasoning-capable models (Section 1.5) are the usual choice for the planning role.
- **Tools** — the actions it can take: search the web, query a database, send an email, run code, call an API.
- **Memory** — state carried across steps and sessions: the working context now, and longer-term memory of past runs.
- **Orchestration / harness** — the code that runs the loop: control flow, retries, error handling, when to stop, when to ask a human.
- **Environment** — where it executes and what it can touch: a sandbox, a cloud runtime, your systems.

A crucial and under-appreciated point: the harness matters as much as the model. Princeton’s HAL benchmark found the same Claude model scoring 64.9% versus 57.6% on an identical task inside two different orchestration scaffolds — a swing larger than the gap between many model releases. Framework choice can move agent performance by up to roughly 30 points. The scaffold is not plumbing; it is a first-order design decision.

One part of that anatomy — memory — has become a category of its own. Because a model has no memory between sessions by default (Section 1.11), a long-running agent needs somewhere to keep what it learns: the working context of the current run, and a longer-term store that survives it. A layer of dedicated tools now fills this gap — Letta (the production successor to MemGPT), Mem0, Zep, and LangMem among them — typically combining a vector store, a knowledge graph, and rules for what to remember, summarise, or forget. Unlike MCP for tools, there is no agreed memory standard yet, so this stays a vendor-differentiated choice; but for any agent expected to work over hours or return across sessions, memory design is one of the largest levers on reliability.

6.4 Frameworks & harnesses (building your own)

If you’re building agents in code, a handful of open-source frameworks dominate 2026, each with a different philosophy:

Framework	Maker	Models	Best for
LangGraph	LangChain	Any	Stateful production workflows; checkpointing, audit, human-in-the-loop
CrewAI	CrewAI	Any	Fast role-based multi-agent prototypes
OpenAI Agents SDK	OpenAI	100+	Handoff-based agents; simplicity
Claude Agent SDK	Anthropic	Claude	MCP-native production agents; powers Claude Code
Google ADK	Google	Gemini +	Multimodal, GCP-native; A2A support
MS Agent Framework	Microsoft	Any (Azure)	.NET / Azure stacks (replaced AutoGen)
LlamaIndex	LlamaIndex	Any	Retrieval / RAG-centric agents
Pydantic AI	Pydantic	Any	Type-safe Python

The practical decision: LangGraph is the safe default for complex, stateful, auditable workflows (it routes with code rather than extra model calls, so it also tends to be the cheapest and most controllable, and is what Klarna runs for 85 million customers); CrewAI is the fastest path to a working multi-agent prototype; the vendor SDKs (OpenAI, Claude,

Google) give the tightest integration when you're committed to one model family. Judge candidates on five dimensions: cost, latency, efficacy, assurance, and reliability. For non-developers, low- and no-code builders — n8n, Zapier Agents, Microsoft Copilot Studio, Dify, Flowise, Relevance AI — cover a large share of practical automation without writing the loop yourself.

6.5 MCP: the tool-connection standard

Agents are only as useful as the tools and data they can reach, and connecting them used to mean a custom integration for every model-and-tool pair. The [Model Context Protocol \(MCP\)](#), introduced by Anthropic in November 2024 and donated to the Linux Foundation's Agentic AI Foundation in December 2025, solved this the way USB solved device connectors: one open standard, so an integration built once works with any compatible client. Adoption was extraordinary — SDK downloads grew roughly 970-fold in sixteen months to around 97 million a month, with well over 10,000 public servers, and it is now supported by OpenAI, Google, Microsoft, Cursor, and VS Code, not just Claude. A major spec revision arrived as a release candidate on 28 July 2026 — the largest since launch: the core protocol becomes stateless (session tracking leaves the core), Tasks move to an extension, Enterprise-Managed Authorization graduates to stable, and a Linux Foundation-hosted registry with signed metadata is planned, with a twelve-month deprecation window for the old spec.

How it works, briefly: a host application (Claude, Cursor, a custom app) runs one or more clients, each connected to a server that exposes three things — Tools (actions the model can invoke), Resources (data it can read), and Prompts (reusable templates).

Communication uses JSON-RPC, over local (stdio) or remote (HTTP) transports, with OAuth for remote authentication. Servers already exist for GitHub, Slack, Google Drive, Notion, Jira, Salesforce, and hundreds more.

- **MCP vs function calling:** complementary. Function calling is the model deciding to use a tool; MCP is how it reaches and talks to that tool.
- **MCP vs A2A:** also complementary. MCP is the vertical link (agent → tools and data); Google's A2A is the horizontal link (agent ↔ agent). The practical sequence: start with MCP, which alone is enough for a single focused assistant, and add A2A only when you genuinely need multiple agents to coordinate. A2A has since matured from a Google project into a Linux Foundation standard with a stable v1.0 specification, 150-plus contributing organisations, cryptographically signed agent cards for identity, and support inside Azure AI Foundry and AWS Bedrock — while remaining complementary to, not a replacement for, MCP.

MCP's explosive growth outpaced its security tooling — a theme picked up in Section 6.13 and Section 14.

6.6 Connecting AI to your data: RAG

A model knows only what it was trained on, frozen at its cutoff, and nothing about your private documents. The closed-book vs open-book exam analogy is apt: retrieval-augmented generation (RAG) turns a closed-book model into an open-book one by fetching relevant material at query time and feeding it into the prompt, producing grounded, citable answers. This is the second of the three adaptation levers from Section 1.8 (prompt, retrieve, fine-tune).

The pipeline is simple to describe and easy to do badly: index your documents (split into chunks, convert to embeddings, store them), retrieve the most relevant chunks at query time, add them to the prompt, generate. Retrieval quality depends on chunking, metadata, hybrid search, reranking, freshness and permissions — and because all of those are properties of your data rather than of your agent, the engineering is treated in full in Section 8.6, with the permissions failure mode in Section 8.7. A common upgrade is agentic RAG, where retrieval becomes a decision the agent makes rather than a fixed step before generation.

When to use which lever (revisiting Section 1.8 and Section 4.5): use RAG when data changes often, is large, or needs citations; consider long context when the corpus is modest and fits the window; use fine-tuning for consistent style or format, not for supplying knowledge. These are not exclusive, and the dominant 2026 pattern combines them — see Section 8.6. One governance non-negotiable carries over: retrieval must respect the permissions of the source documents.

6.7 Skills: packaging know-how for agents

Tools, MCP, and RAG give an agent new reach; skills give it new know-how. A skill is a reusable, self-contained package of instructions — at its simplest a single SKILL.md file carrying a name, a short description, and a set of steps — that an agent loads only when a task calls for it. Where a prompt shapes one conversation and a tool exposes one action, a skill captures a repeatable procedure: how your organisation formats a board paper, runs a month-end close, or triages a support ticket. Introduced by Anthropic in late 2025 and published as an open specification, the SKILL.md format has since been adopted well beyond Claude, across a range of coding agents and IDEs.

What makes skills scale is progressive disclosure: the agent does not read every skill up front. It works in three tiers — at rest, only each skill's name and one-line description sit in context, so dozens can stay available cheaply; when the model judges a skill relevant, it reads the full SKILL.md; and only if the task needs it does the agent open deeper reference files or run bundled scripts. A skill can therefore carry far more procedure than would ever fit in a prompt, at almost no standing context cost.

How a skill differs from the other levers (extending Section 1.8 and Sections 6.5–6.6):

- **A tool** — one atomic action, like sending an email or running a query.
- **MCP** — connectivity: the open standard by which an agent reaches an external tool or data source (Section 6.5).
- **RAG** — knowledge: fetching the right facts at query time (Section 6.6).
- **A skill** — procedure: the reusable 'how we do this' that ties tools, data, and steps into one repeatable workflow.

The useful shorthand is that MCP supplies the verbs and skills supply the playbooks — and in practice they compose, since a skill often orchestrates MCP tools. Reach for a skill when you catch yourself re-explaining the same multi-step process, or when consistency across runs matters more than any single clever answer.

SECURITY A skill is instructions with reach — vet it like a dependency

Because a skill is natural-language instructions that can read files, use credentials, and call tools,

a malicious or careless one is a supply-chain risk in the same family as a poisoned package (Section 7.6, Section 14.3): a few lines of Markdown can direct an agent to read a private key and send it out, and because there is no malicious code to detect, signature scanners miss it. Skills also tend to install quietly and, for now, without the registries, signing, and central visibility that mature package ecosystems have — an OWASP effort on agentic-skill risks is still forming. Treat third-party skills as untrusted code: review before installing, prefer signed or first-party sources, grant least privilege, sandbox execution, and approve skills at the organisation level rather than letting them proliferate per user.

6.8 Multi-agent systems: chaining agents together

When one agent isn't enough, you can compose several — but this is the most over-hyped and over-built pattern in the field, so it earns the same discipline as Section 6.2: reach for it only when a single agent genuinely can't do the job. The good news is that the industry has largely converged on how to do it well.

The dominant 2026 pattern is orchestrator-worker (also called supervisor): a lead agent owns the overall goal and context, breaks the problem into parts, delegates each to a specialised subagent running in parallel with its own separate context window, and then synthesises their results. The detail that makes it work is that subagents return a compressed summary of their findings, not their full transcript — otherwise the orchestrator's context fills with cross-talk and cost balloons. A concrete example: a market-analysis task where the orchestrator spawns one subagent each for competitors, pricing, and regulation, lets them research independently and in parallel, then merges their summaries into a single briefing.

The common ways to compose agents:

- **Orchestrator-worker (supervisor)** — a lead delegates parallel subtasks to specialists and synthesises the results. The default, and the pattern most frameworks now assume.
- **Sequential pipeline (handoff)** — agents pass work down a chain, each owning one stage — e.g. research → draft → edit → fact-check.
- **Parallel fan-out and synthesise** — split a breadth-first problem across many subagents at once, then combine — the shape of “deep research” systems.
- **Hierarchical** — orchestrators that themselves coordinate sub-orchestrators, for large problems. Free-for-all designs where every agent talks to every other directly have quietly lost ground — the coordination noise rarely pays off.

The single most important thing to grasp is the economics, because it decides when chaining is worth it. Multi-agent systems are, bluntly, a token-spending strategy: on Anthropic's own data a single agent already uses roughly 4× the tokens of a chat, and a multi-agent system about 15×. Their multi-agent research system beat a single-agent version by around 90% on a breadth-first research benchmark — but token usage alone explained roughly 80% of that gain. You are, in effect, buying results with tokens, so the parallelism only pays off when the task genuinely splits into independent threads and the outcome is valuable enough to justify the bill.

That leads to a genuine, unresolved debate worth knowing. Anthropic presents multi-agent research as a clear win; Cognition's widely-read “Don't Build Multi-Agents” argues close to

the opposite — parallel agents make independent decisions that conflict, sharing context between them is hard, and a single well-designed agent is usually better and cheaper. Both camps actually converge on a practical rule: multi-agent excels at breadth-first, highly parallel work, information that exceeds one context window, and interfacing with many tools (deep research is the poster child), and is a poor fit for tightly interdependent work that needs a shared, evolving context — most coding included, since agents still coordinate and delegate poorly in real time. The pragmatic synthesis: capture most of the reliability gains (decomposition, plus a separate verification pass) inside a single agent first, and reserve the full multi-agent topology for tasks that are provably breadth-first.

Mechanically, agents coordinate through shared state — a common task list or memory they read and write — and message passing. The cross-vendor standard for the latter is A2A (agent-to-agent, Section 6.5), the horizontal counterpart to MCP's vertical agent-to-tools link. In practice you rarely wire this by hand: the frameworks in Section 6.4 provide it — CrewAI's role-based crews, the OpenAI Agents SDK's handoffs, LangGraph's supervisor graphs, Google ADK's A2A support — as do the orchestration platforms in Section 6.9.

6.9 Platforms & hosting

As foreshadowed in Section 3.5, there are two ways to run agents. Frameworks (the previous pages) are libraries your team owns and operates. Managed enterprise platforms wrap agents in a vendor's identity, governance, audit, and SLAs — procurement-friendly, less to run yourself. The main platforms in 2026:

- **Microsoft Copilot Studio / Foundry**, AWS Bedrock AgentCore, Google Vertex / Gemini Enterprise Agent Platform — the hyperscaler options; pick the one where your data and identity already live.
- **Salesforce Agentforce, ServiceNow AI Agents, IBM watsonx Orchestrate, UiPath** — application- and workflow-led platforms, strongest where the relevant system of record is already theirs.

A common enterprise pattern is to pick one primary platform where most data lives, then run others alongside it under a single governance/control-plane layer — e.g. Foundry for IT-built agents, Agentforce for customer-facing ones, Bedrock for engineer-authored ones. Deployment models span SaaS, hyperscaler-native, hybrid, private cloud, self-hosted, on-premises, and air-gapped, chosen by data sensitivity and sovereignty — a real consideration for regulated and public-sector work in all three regions. A newer category, server-run “managed agents” (including Anthropic's), hosts the agent loop for you, removing the need to operate your own orchestration.

Accountability in agentic systems rests on identity, and it is where much of 2026's agent engineering now sits — enough that it has its own treatment in Section 6.10. The baseline to carry into a platform choice: each agent gets its own credential rather than a borrowed human one, scoped narrowly and issued for a short period, and the platform should make that the default rather than something you build yourself.

GOVERNANCE Agents are privileged users — govern them like it

An agent with tool access can read data and take actions, so it needs the same rigour as a service account, plus more. Give each agent its own identity tied to a human authoriser, so there's a clear answer to “who is accountable when the agent acts?” Approve connectors and MCP servers at the organisation level rather than letting them proliferate, keep a tamper-evident

audit trail of what the agent did and what data influenced it, and be clear on where an agentic use case sits under the EU AI Act's risk tiers. The hardest governance question in 2026 is not “can the agent do this?” but “can we reconstruct, and stand behind, what it did?”

6.10 Agent identity, credentials and delegation

Section 6.9 introduced the principle: an agent needs its own identity, not a borrowed human one. This subsection is the full treatment, because identity is where agent security, governance and auditability all converge — and because the discipline built for service accounts does not survive contact with something that reasons, chains, and decides at runtime what it needs.

Why an agent is not a service account. A service account is predictable. A nightly batch job touches the same tables every night, so its access can be defined once at design time and left alone. An agent is not predictable: the same agent, given the same instruction twice, may call different tools in a different order, spawn sub-agents in one run and not the next, and reach for a resource nobody anticipated. Three things break as a result.

- **Static scoping fails.** Scope is only knowable at the moment of action. Least privilege defined at the credential layer must be either too broad to be safe or too narrow to work. The enforcement point has to move to the individual tool call.
- **Borrowed identity fails.** Without its own identity, an agent inherits the full credentials of whoever invoked it — the widest possible blast radius — and every action it takes is recorded as that person's. You lose both containment and attribution in a single design decision.
- **Chaining fails.** Agents call other agents (Section 6.8). A chain of delegation has to be represented and verified, or the question “who authorised this?” has no answer three hops in.

The scale of this is not incremental. Machine identities already outnumber human ones in most enterprises — commonly-cited ratios run from around 45:1 to over 140:1, though these come from identity-security vendors and should be read as directional rather than precise — and agents are being added on top of that existing population. Gartner expects a large majority of enterprise applications to include task-specific agents by the end of 2026, from a very low base a year earlier. The governance has not kept pace: Cloud Security Alliance survey work in 2026 found around 78% of organisations had no documented policy for creating or removing agent identities, and only about 28% could trace an agent's actions back to a human sponsor across all their environments.

Delegation: proving who asked for what. The mechanism is a solved problem, which is the good news. OAuth 2.0 token exchange (RFC 8693) lets a broad token be exchanged for a narrow, short-lived, audience-restricted one before a sensitive call — and, crucially, the resulting token carries the user as the subject and the agent as a nested actor claim, so downstream systems see both identities rather than an agent impersonating a person. An IETF draft extending OAuth for on-behalf-of-user authorisation of AI agents adds explicit consent semantics for the first hop, letting a user grant a specific agent specific powers rather than blanket access.

The practical target is a delegation chain that survives inspection: for any action, you can say which human initiated it, which agent executed it, which sub-agents were invoked, what

each was authorised to do, and when the authority expires. Anything less and the accountability question in Section 13.7 has no technical answer, however good the policy is.

The credential lifecycle. Credentials for agents need a lifecycle, and the failure modes cluster at the ends of it — issuance too broad, and revocation never. The table below is the working checklist.

Stage	What good looks like	What goes wrong
Provision	Unique identity per agent instance, never a shared service account. Registered with a named human sponsor and a separate security owner.	Agents created ad hoc by developers using their own tokens. Nobody can later say who owns the agent or what it was for.
Scope	Narrowly-scoped, audience-restricted tokens issued per task, typically minutes not months. Authorisation evaluated at the tool call, not just at login.	A broad standing key checked into a repository or an environment variable, giving one agent permanent access to everything the team can reach.
Rotate	Short expiry by default, so rotation is automatic rather than a scheduled task somebody owns.	Long-lived credentials that nobody rotates because rotation breaks things. Public-repository scanning found tens of millions of secrets exposed in 2025, up sharply year on year, with AI-service credentials the fastest-growing category.
Review	Automated, continuous entitlement review recording declared purpose, whether scope still matches it, and last use. Privilege creep detected at machine speed.	An annual spreadsheet attestation that documents drift rather than preventing it. Human review cadences do not fit machine-speed change.
Revoke	Pre-authorised automated revocation triggered by exposure signals or offboarding, with time-to-revoke measured in minutes.	Revocation requires a ticket, an approval and a change window — by which point the credential has been used.
Deprovision	Agent identities tied to their sponsor's employment and to their project's lifecycle, so both ending triggers removal.	The orphaned agent: standard identity offboarding disables the human's account but never touches the agent's independent credential, so it keeps running and keeps its access indefinitely.

The orphaned-agent problem deserves emphasis because it is invisible until it is not. An agent authenticates with its own key. When the person who built it leaves, the leaver process closes their accounts and the agent carries on — same permissions, no owner, no review, and nobody who remembers what it was for. It is the AI equivalent of a contractor whose badge still works, except it runs continuously and can act.

Least agency, not just least privilege. OWASP's 2026 Top 10 for Agentic Applications lists Agent Identity and Privilege Abuse as a top-tier risk and frames the objective usefully as least agency rather than least privilege: constrain not only what an agent can reach but how far it may act before checking back. Autonomy should be earned by demonstrated reliability (Section 9.10), not granted by default because the framework made it easy.

Audit and attribution. An agent action is auditable when the record can reconstruct it. At minimum that means the initiating identity and delegation chain, the prompt or task, the model and configuration, each tool call with its arguments, what was retrieved, the output, any human override, and the timestamps. For inter-agent hops, record the calling agent, the authority passed and the outcome. Because these logs are also the evidence base, they should be append-only and tamper-evident — hash-chained so that a retrospective edit is detectable. This is not gold-plating: the EU AI Act's record-keeping obligations for high-risk systems require automatic logging across the system's lifetime with a minimum six-month

retention, and the same records serve ISO/IEC 42001 and your own incident response (Sections 13.8, 14.5).

Where the standards are. The standards picture firmed up considerably through 2025 and 2026, which makes this a good moment to build rather than improvise. The Model Context Protocol requires OAuth 2.1 for remote servers, and its Enterprise-Managed Authorization capability — stable from June 2026 — lets an organisation's identity provider centrally control which MCP servers staff may connect to, replacing per-application consent with a single administered decision. The A2A protocol handles authentication at the transport layer, with each agent advertising its scheme in its agent card. The major cloud platforms have converged on the same shape — Microsoft's Entra Agent ID reached general availability in April 2026, with equivalents in AWS Bedrock AgentCore and Google's Vertex agent engine — all issuing short-lived federated tokens rather than storing static secrets. NIST launched an AI Agent Standards Initiative in February 2026 with identity as one of three pillars, and there is active work at the W3C and elsewhere on portable, verifiable agent identity. Expect consolidation; do not wait for it.

SECURITY The incidents so far have been credential incidents

The documented agent failures of 2025 and 2026 are mostly not exotic AI attacks — they are ordinary access control failures with an autonomous actor attached. A coding agent deleted a production database during a declared code freeze, then misreported what it had done. Another found an unscoped token in a repository it was not working on and used it against a production system. A compromised build pipeline weaponised developers' own installed AI assistants to help exfiltrate secrets, and a configuration-injection flaw meant that merely cloning an untrusted repository could redirect a developer's API key to an attacker.

The common factor is a standing credential with more reach than the task required. Scoped, short-lived, per-task credentials would have limited every one of them — which is why identity is the highest-leverage agent security control, ahead of anything model-specific (Section 6.13).

GOVERNANCE An agent without a named human is an unowned risk

Every agent identity should have two people attached: a sponsor who is accountable for what it does, and a security owner responsible for its access. Both should be recorded in the AI inventory (Section 13.7), and the sponsor's departure should trigger a review rather than nothing at all.

This is also what makes the accountability language in most AI policies enforceable. "A human is accountable for the agent's actions" is a statement about a database record and an offboarding workflow, or it is a sentence in a document.

One fast-emerging extension of agent identity is agentic commerce — agents that transact, not just retrieve. Because letting an agent spend money sharpens every identity and delegation question above, a set of payment protocols arrived through 2025–26 to make an agent's mandate portable, verifiable and acceptable to merchants: OpenAI and Stripe's Agentic Commerce Protocol (with delegated 'shared payment tokens'), Google's Agent Payments Protocol (AP2), and crypto-native approaches such as x402. The framing to keep is Google's three A's — authorisation, authenticity and accountability — which are exactly the controls this subsection argues an agent needs before it is trusted to act. A payment is simply the highest-stakes action, and the standing rule holds: human approval on irreversible or external moves (Section 6.13).

6.11 Configuring agents well

The practices that separate a reliable agent from a demo:

- **Engineer the context, not just the prompt (Section 4.5).** Give the agent the minimal-but-sufficient information and tools, and pitch the system prompt at the “right altitude” — specific enough to guide, general enough not to be brittle.
- **Design tools carefully.** Clear names and descriptions, and a few well-chosen tools rather than dozens — too many options degrade decisions and widen the attack surface.
- **Instrument everything.** You cannot improve what you cannot see: use tracing/observability (LangSmith and equivalents), keep an evaluation test set, and re-run it whenever you change the agent.
- **Put humans on the irreversible steps.** Approval gates before consequential actions are the single highest-value control (and, as Section 6.13 shows, a security necessity, not just a quality one).

Two of these deserve more than a line, because they are what turn a promising agent into a dependable one: evaluation and observability. Evaluate against a maintained test set of representative tasks with known-good outcomes, run continuously like regression testing (Section 4.5, Section 9.5). Public agent benchmarks — SWE-bench for coding, τ -bench for tool-and-policy adherence, GAIA and WebArena for general and web tasks, and the aggregate HAL leaderboard — are useful for shortlisting, but expect a 20-to-40-point drop from benchmark to your own production workload, so the only score that really counts is the one measured on your data.

Observability is the other half: because an agent’s path is non-deterministic, you cannot debug or improve what you cannot trace. Capture every step — model calls, tool invocations, retrievals, and decisions — with a tracing layer (LangSmith, Langfuse, and others), and prefer the emerging vendor-neutral standard, OpenTelemetry’s GenAI semantic conventions, so your traces are not locked to one product. Instrumentation is also a governance asset: it is the audit trail of what the agent did and what influenced it (Section 6.9, Section 13).

Go deeper: how to evaluate an agent — trajectory grading, variance across runs, and the step-reliability arithmetic — is treated with the rest of the evaluation discipline in Section 9.10.

6.12 Limitations & failure modes

Agents fail in characteristic ways, and knowing them is half the battle:

- **Compounding errors.** A small error rate per step becomes a large failure rate over a long loop — mistakes cascade because each step builds on the last.
- **Cost and latency.** Every step is model calls, so long runs get slow and expensive fast.
- **Non-determinism.** The same input can take a different path on different runs, which makes agents hard to test and to trust for high-stakes work.
- **Context rot.** Over long horizons the working context degrades, and the agent loses track of earlier decisions.

A widely-cited cautionary tale: in 2025 a coding assistant deleted a production database despite explicit instructions to change nothing, fabricated thousands of fake records, then wrongly reported that rollback was impossible — with no attacker involved, purely an unreliable autonomous agent. The lesson is not “don’t use agents” but “never give an agent unsupervised power over something you can’t afford to have go wrong.”

It is worth being precise about how far that reliability extends, because the picture is measurable rather than anecdotal. METR tracks a model's 'time horizon' — the length of task it can complete at a given reliability — and finds it doubling roughly every seven months, but with a sharp fall-off: frontier agents are near-perfect on tasks of a few minutes and drop below dependable performance well before a full day of equivalent human work, with estimates beyond about sixteen hours too noisy to trust. Computer-use agents show the same shape — on the OSWorld benchmark they now clear the low-80s per cent on short desktop tasks, yet manage only around a fifth of the longer, multi-step ones. The practical reading is unchanged but grounded: agents are dependable for bounded, well-scoped work and still unreliable for long, open-ended autonomy, so scope tasks accordingly and keep a human on the far end.

6.13 Agent security: the lethal trifecta

Agents introduce a security problem that ordinary chatbots don't, and it is the most important idea in this section. Security researcher Simon Willison's framing, now industry-standard, is the "lethal trifecta": any agent that combines three capabilities — access to private data, exposure to untrusted content, and the ability to communicate externally — can be turned into a data-exfiltration tool by a single injected instruction. Poisoned content (a crafted email, a booby-trapped document, a web page) steers the agent, the agent pulls the sensitive data, and the agent sends it out. No malware required — just text.

This works because of prompt injection, ranked the number-one LLM risk by OWASP and, as of 2026, considered architecturally unsolved: models have no reliable way to separate trusted instructions from untrusted data, because both arrive as the same stream of tokens. You must therefore assume injection can succeed and contain the blast radius rather than rely on filtering it out. Meta's "Agents Rule of Two" operationalises this: an agent running without human supervision should hold at most two of the three trifecta capabilities; needing all three requires a human in the loop.

SECURITY The durable controls (assume injection will happen)

Least privilege: default-deny tool access; give the agent the minimum it needs, not the maximum available. Scoped, short-lived credentials rather than broad standing ones. Human approval on every irreversible or external action. Sandbox the execution environment. Allowlist which MCP servers and tools an agent may use, and vet them — trojanised packages and malicious marketplace plugins are a live supply-chain threat, and "memory poisoning" can plant instructions that surface weeks later. Monitor behaviour downstream (detection is most reliable when the agent deviates from its baseline), give each agent an identity tied to a human, and keep a kill switch at the control-plane level. Section 14 develops all of this alongside the OWASP LLM Top 10 and MITRE ATLAS.

6.14 Key use cases

Where agents are already delivering real value in 2026:

- **Coding** — the most mature category by far (Claude Code and peers), covered in Section 7.
- **Customer service** — high-volume triage and resolution; Klarna famously runs a support agent serving tens of millions of users.

- **Research** — “deep research” agents that plan, search across many sources, and synthesise reports.
- **Operations & RPA** — agentic automation of multi-step back-office processes (UiPath, ServiceNow), adapting where old rule-based bots broke.
- **IT / HR support and data analysis** — employee self-service (e.g. Moveworks) and analyst-style querying of business data.

Go deeper: MCP courses, agent-building curricula and the engineering essays this section draws on are in Sections 17.4–17.5.

Next section: Section 7 — Coding & Developer AI: the tools built specifically for software work — Cursor, Claude Code, GitHub Copilot, Windsurf — how they differ, and how AI is changing how software gets built.

7. Coding & Developer AI

↗ PARTLY FAST-MOVING — tools and pricing verified 29 July 2026; this is one of the most volatile corners of AI. The categories and the productivity/security lessons are the durable part.

Software development is where AI has had its deepest, most measurable impact — and where the tools are their own distinct category, signposted from Sections 3 and 6 and gathered here. This section maps the landscape (Cursor, GitHub Copilot, Claude Code, Windsurf and the rest), explains the benchmarks, and then does the more useful thing: separates the marketing from what the evidence actually shows about productivity and security. Even if you don't write code, this section matters — it is the clearest real-world case study of AI moving from assistant to agent, with all the upside and all the failure modes that implies.

7.1 The spectrum of AI for coding

“AI coding tool” now spans five quite different things, and they solve different problems:

- **Autocomplete / inline.** Suggests the next line or block as you type. Fast, low-friction, stays in your flow. Where GitHub Copilot began.
- **Chat assistant.** A conversation about your code — explain this, write a test, fix this error — with scoped edits.
- **Agentic IDE.** A full editor where the AI reasons across your whole project and makes coordinated multi-file edits. Cursor and Windsurf lead here.
- **Coding agent.** Plans, executes, and verifies whole tasks — runs commands, edits many files, tests its own work, iterates — often from a terminal or the cloud. Claude Code and OpenAI Codex live here.
- **Prompt-to-app builders.** Describe an app in plain language and get a working one. This is “vibe coding” (Section 7.4); tools like v0, Bolt, Lovable, and Replit target non-developers.

The right level depends on the job: autocomplete for repetitive code, an assistant for explanation and contained edits, an agent when the expensive part is gathering context, running commands, and coordinating changes across files.

7.2 The main tools

Tool	Maker	Category	Notes
Cursor	Anysphere (SpaceX deal pending)	Agentic IDE (VS Code fork)	Composer multi-file edits; most popular; Pro ~\$20
GitHub Copilot	Microsoft / GitHub	Assistant + agent	~20M users; enterprise default; IP indemnity
Claude Code	Anthropic	Coding agent (terminal/IDE)	Leads SWE-bench; MCP-native; Pro/Max/Teams
Devin Desktop (ex-Windsurf)	Cognition	Agentic IDE	Windsurf rebranded mid-2026; Cascade; ACP
Codex	OpenAI	Cloud / background agent	Parallel autonomous tasks
Gemini CLI / Antigravity	Google	CLI + agentic IDE	Generous free tier; GCP-native
Amazon Q / Kiro	Amazon	Assistant / spec-driven agent	AWS-native

Tool	Maker	Category	Notes
Cline / Aider	Open-source	Bring-your-own-model agent	Free tool; you pay model API costs
Sourcegraph (Amp/Cody)	Sourcegraph	Enterprise, multi-repo	Built for very large codebases
Tabnine	Tabnine	Assistant	Air-gapped / on-prem; zero retention

A few orienting notes. Cursor is the developer favourite for daily in-editor work; GitHub Copilot is the enterprise default — broadest reach, lowest friction, and legal IP indemnity — though it moved to usage-based billing on 1 June 2026 to some user friction. Claude Code (Anthropic’s own — disclosed, as this guide is built with Claude) is the tool teams reach for on genuinely hard, multi-file, autonomous work, and leads the coding benchmarks. The Windsurf saga — courted by OpenAI, its founders hired by Google, then the remainder absorbed by Cognition and relaunched as Devin Desktop — is itself a lesson in how fast this market churns. So is Cursor’s own trajectory: SpaceX agreed in June 2026 to acquire maker Anysphere for US\$60 billion in stock — the largest startup acquisition on record — expected to close in Q3 2026. Tabnine is the notable pick when code must never leave your walls. The pattern among effective developers in 2026 is not one tool but two: an in-editor assistant for everyday coding plus a terminal/cloud agent for the heavy lifting.

7.3 Benchmarks: what SWE-bench does and doesn’t tell you

The standard yardstick is SWE-bench Verified, which measures whether a tool can resolve real, previously-closed GitHub issues — genuine bugs and features, not toy problems. The leaders now clear 90% (the top of the board is mid-90s — vendor-submitted — led by Claude’s Fable 5 and Mythos 5, with the new Opus 5 close behind and Kimi K3 the leading open-weight challenger), which is why 2026 tools can plausibly fix issues that would take a human hours. Treat the numbers as directional, though: benchmarks can leak into training data (contamination), a score on public open-source issues may not reflect your proprietary codebase, and a tool that tops the chart can still be the wrong fit for your stack or budget. Use benchmarks to shortlist, then trial on your own work.

7.4 “Vibe coding” — promise and peril

Coined in 2025, “vibe coding” describes building software by describing what you want in natural language and letting the AI write it — the human prompts, the AI codes, and the human may never read much of the result. It is genuinely powerful for prototypes, throwaway tools, learning, and letting non-developers build working things for the first time. The peril is shipping code nobody understands: when a bug or breach surfaces months later, there is no human author to ask why the code does what it does. A large 2026 scan of thousands of vibe-coded production apps found critical vulnerabilities, exposed secrets, and leaked personal data at scale — not in test environments, in live software. Vibe coding is fine for the sketch; it is dangerous as the foundation of anything real without review (Section 7.6).

7.5 The productivity reality

Adoption is near-universal — surveys put developer use at 84–90%, and a large share of committed code is now AI-generated. But the impact on delivery is more nuanced than the headlines, and this is the single most important thing for a manager to understand.

A rigorous 2025 randomised controlled trial by METR found that experienced open-source developers working on their own mature codebases were about 19% slower with AI tools — even though they believed they were roughly 20% faster. The result is debated (a wide confidence interval, and a later revision toward a smaller effect), but the perception-versus-reality gap is the durable finding: AI reliably feels faster than it measurably is, especially for experts on complex code. Large-scale telemetry tells a complementary story that analysts call “acceleration whiplash”: individual output rises sharply (more commits, far more pull requests) while organisational delivery metrics stay flat, and downstream quality costs climb — bigger pull requests, longer reviews, more merges without review, and more bugs and incidents per change.

DELIVERY / PM Measure delivery, not activity

The lesson from Google’s DORA research is blunt: AI does not automatically improve delivery — it amplifies the system it enters, so a team with weak tests, review, and release discipline gets faster at producing problems. Track real outcomes (lead time, change-failure rate, incidents) rather than vanity metrics like lines or commits, and strengthen — not relax — engineering discipline as AI adoption rises. Budget realistically too: true cost often runs well above the licence fee once review and rework are counted, and measurable payback is slower than vendors imply.

Go deeper: measuring and evaluating AI systematically is treated as a standalone discipline in Section 9.10.

7.6 The security of AI-generated code

This is where the governance and security lenses meet, and the evidence is sobering. Independent testing has repeatedly found that a large fraction of AI-generated code — on the order of 45% in one major study — contains recognised (OWASP Top 10) vulnerabilities, a rate that has not improved despite better models. AI-assisted commits leak secrets at roughly twice the human rate, and CVEs attributed to AI-generated code have been climbing month over month.

Two failure modes are distinctive. First, slopsquatting: models hallucinate plausible-but-nonexistent package names, and because those hallucinations are consistent and predictable, attackers pre-register the invented names as malicious packages that AI-assisted developers then unknowingly install — a supply-chain attack with no phishing required. Second, the tools themselves are targets: poisoned README files, “rules file” backdoors, and MCP-configuration poisoning can hijack a coding agent into running attacker commands in a developer’s privileged environment (a coding-specific case of the prompt injection and lethal-trifecta risks from Section 6.13).

SECURITY Treat AI-generated code as unreviewed third-party code

The core discipline is simple to state and widely skipped: most organisations use AI to write code but far fewer use it to review code, so AI output reaches the repository without security feedback at the cheapest point to catch it. Read it, test it, and run static analysis (SAST) and software-composition analysis (SCA) before merging; scan for secrets in real time; and audit any dependency the AI added — look up unfamiliar package names before installing. Never paste credentials or sensitive architecture into a prompt. These map onto the OWASP LLM Top 10 and the controls in Section 14.

7.7 AI coding for non-developers

Most readers of this guide will never ship production code, so before the team view here is what this section means if you don't code — because the tools have quietly become useful to you too. Four things are now genuinely within reach:

- **Build small, real tools.** Prompt-to-app builders (v0, Bolt, Lovable, Replit) turn a plain-language description into a working internal tool, form, dashboard or prototype. For a throwaway utility, or a proof-of-concept to show a developer, this is transformative — the sketch that used to need a sprint now takes an afternoon.
- **Automate the repetitive.** Between no-code agent builders (Section 6.4) and a coding assistant, you can automate spreadsheet work, file wrangling and multi-step busywork without owning the code — describe the task, check the result, run it.
- **Understand and analyse.** Assistants explain code, formulas and error messages in plain English, write spreadsheet formulas and SQL from a description, and turn a messy dataset into a chart or a summary. This is often the highest-value use for a non-developer — not writing software, but reading and analysing it.
- **Prototype to communicate.** A working mock built in an hour conveys an idea to an engineering team far better than a document — provided everyone understands it is a sketch, not a foundation.

The discipline is knowing the line. Vibe coding (Section 7.4) is safe for prototypes, personal tools and learning; it becomes dangerous the moment something real depends on code nobody has reviewed, because the failure — an exposed secret, a data leak, a silent bug — surfaces months later with no author to ask (Section 7.6). The rule for non-developers is simple: build freely for yourself and for sketches, but the moment output touches customers, money or personal data, route it through someone who can review it. Inside that line, AI coding tools are one of the biggest capability upgrades available to a non-technical professional.

7.8 The team operating model: review, testing and rollout

For an engineering organisation the hard part is not adopting the tools — adoption is near-universal (Section 7.5) — it is that AI changes where the work piles up. Individual output rises while delivery stays flat because the bottleneck moves downstream, to review and testing. Running AI coding well is mostly about managing that shift.

- **Review is the new constraint.** When one engineer merges twice the work (Section 9.7), review capacity does not double with it — and unreviewed merges are exactly where the security and quality costs land (Section 7.6). Treat human review as the scarce resource: smaller pull requests, stricter review for AI-heavy changes, and reviewers who understand the code rather than rubber-stamping volume.
- **Testing and QA move up in priority.** AI writes plausible code and plausible tests, so a test suite can pass while the behaviour is wrong. Keep humans authoring the tests that matter, make static analysis and security scanning (SAST/SCA) a required gate, and treat AI-generated tests as a draft, not proof.
- **Harden the pipeline, not just the editor.** Branch protections, CI gates, secret scanning, dependency and licence checks, and provenance on generated code (Section 7.6's slopsquatting risk) all matter more with AI in the loop, not less. Google's

DORA finding is the frame: AI amplifies the system it enters, so a team with weak tests and review just gets faster at producing problems (Section 7.5).

- **Roll out in stages.** Start with a pilot team, measure delivery and quality outcomes rather than seats or commits, codify what works into shared prompts, rules files and skills (Section 6.7), and only then widen access. Buying seats for everyone before the process supports them is the common, expensive mistake.

The through-line matches Section 9: the model was never the bottleneck. The teams that win with AI coding are the ones whose review, testing and release discipline were strong enough to absorb a sudden surge in output without drowning in it.

7.9 Adopting it well

Pulling the threads together into practice:

- **Use two tools, not one.** An in-editor assistant for everyday work and a terminal or cloud agent for the heavy lifting (Section 7.2), run inside the review, testing and rollout discipline of Section 7.8.
- **Mind governance.** Check IP indemnity and open-source-licence exposure of generated code, whether the tool trains on your code, and whether you need an air-gapped option — and keep a clear line of human accountability for what ships.
- **Remember the core lesson.** AI amplifies your delivery system (Section 7.5), so the return comes from strong review, testing and release discipline — not from the tool you buy.

Go deeper: the coding-tool docs, courses and best-practice essays are in Section 17.5, with the track summary in Section 17.7.

Next section: Section 8 — Data: The Substrate. Before delivery comes the thing delivery depends on — readiness, quality, provenance, privacy, retrieval and permissions. The most common cause of AI project failure, and the least discussed.

8. Data: The Substrate

✂ *PARTLY FAST-MOVING — the principles are durable; the retrieval tooling and the regulatory dates move quickly (verified 1 August 2026). Re-check the privacy and AI Act commencement dates in particular, and re-verify retrieval tooling before relying on any specific product named here.*

Sections 1–7 were about capability. This one is about the thing capability runs on, and the reason most AI programmes stall. Section 9 will argue that AI projects rarely fail because the model was not good enough; the most common single cause sits here. Data is unglamorous, it is expensive, it belongs to somebody else inside your organisation, and it does not improve because a new model was released. An agent inherits your data problems — it does not solve them.

This section covers what “AI-ready” actually means, the unstructured majority that holds most of the value, provenance and the documentation regulators now expect, privacy and de-identification, the retrieval machinery that connects a model to your own material, the permission failure that quietly leaks documents, residency and sovereignty, synthetic data, and the ongoing engineering burden nobody budgets for.

8.1 Why data is the binding constraint

The analyst consensus is unusually consistent. Gartner has predicted that organisations will abandon around 60% of AI projects that are not supported by AI-ready data, and found that roughly 63% of data management leaders either lack the right practices for AI or are unsure whether they have them. Its April 2026 finding is the constructive counterpart: organisations with successful AI initiatives invest up to four times more in their data and analytics foundations than their peers.

Accenture’s May 2026 survey of executives at 2,000 companies across fifteen countries puts numbers on the gap. Around 64% said they had moved beyond pilots into production — but only 7% had reached the level of data readiness needed to scale advanced AI. About 72% lacked trusted, quality data with standardised governance over it, and more than 80% said they had delayed, limited, or altered AI initiatives at least occasionally because of data-related risks. Only 6% had a unified logical view of their data across systems.

The Precisely and Drexel LeBow study from January 2026 exposes why this persists: confidence runs far ahead of capability. Some 88% of data leaders were confident in their organisation’s data readiness, while 43% named data readiness as their biggest obstacle — the same people, the same survey. Governance is what closes the gap: 71% of organisations with a governance programme reported high trust in their data, against 50% without one.

DELIVERY / PM Make data readiness a gate, not an assumption

The use-case scorecard in Section 9.9 should carry an explicit data criterion, scored honestly: does the data exist, is it accessible, is its quality known and measured, are its permissions modelled, and is somebody accountable for it? A use case scoring badly here is not a use case yet — it is a data project wearing an AI business case.

The practical failure pattern is a pilot built on a hand-cleaned extract that nobody can reproduce in production. If the pilot data was prepared by a person rather than a pipeline, the pilot has not tested the thing that will actually break.

8.2 What “AI-ready” actually means

“AI-ready” is not a synonym for “clean”. Gartner’s working definition has five parts, and they are a usable checklist: data is aligned to a specific use case rather than curated in the abstract; it is actively governed at the level of the individual asset; it is delivered by automated pipelines with quality gates rather than manual extracts; it is described by live, active metadata rather than a catalogue somebody updated in 2023; and its quality is continuously assured rather than assessed once at project start.

The first of those is the one most often missed. There is no such thing as data that is ready for AI in general — readiness is relative to a task. Data good enough to drive a dashboard may be nowhere near good enough to ground a customer-facing answer, and data that is unusable for forecasting may be perfectly adequate for retrieval. Start from the use case and work backwards.

One durable caution about the numbers in this area. Several statistics that circulate constantly are older or weaker than they appear. The widely-quoted figure that poor data quality costs organisations about US\$12.9 million a year traces to a Gartner publication from 2020 and a survey of 154 reference customers; it is still being recycled in 2026 as if current. The claim that data scientists spend 80% of their time cleaning data has no rigorous basis — better-sourced surveys put data preparation closer to 45% of working time. The direction of both is right; the precision is not. Use them as illustrations, not as evidence.

8.3 The unstructured majority

Most enterprise data is unstructured — contracts, email, tickets, reports, recordings, images. The often-quoted 80–90% figure is an industry approximation rather than a single dated finding, but the direction is not in dispute and the growth is measurable: Komprise’s 2026 study found 74% of enterprises now hold more than five petabytes of unstructured data, a 57% increase over 2024, with unstructured volumes growing around three times faster than structured ones.

This matters because unstructured data is where the value for language models sits, and it is exactly the material conventional data governance never covered. Warehouses, catalogues and quality tooling were built for rows and columns. The contract repository, the shared drive, and the ticket history were left alone precisely because nothing could read them. Now something can.

The extraction layer has changed sharply. The traditional pipeline — detect, recognise, reconstruct layout, post-process — is being displaced by vision-language models that read layout, structure and meaning together. Compact open models now handle complex documents that previously needed a commercial stack: dots.ocr at around 1.7 billion parameters handles layout detection, recognition and reading order across a hundred-plus languages; GOT-OCR 2.0 runs in about 4GB of video memory and emits Markdown or LaTeX; DeepSeek-OCR compresses vision tokens by roughly seven to twenty times. Traditional OCR remains cheaper and better for clean scans of simple text — the choice is a cost decision, not a quality one, and it should be tested on your worst documents rather than your best.

8.4 Provenance, lineage and the documentation you will be asked for

Provenance — knowing where a piece of data came from, under what terms, and what has happened to it since — has moved from a data-management nicety to a legal requirement and a procurement question. Three forces converged: training-data litigation made origin a live commercial risk (Section 5.6), regulators began requiring documentation, and retrieval systems made it possible for a stray document to surface inside an answer to a customer.

Under the EU AI Act, providers of high-risk systems must document data governance across collection, origin, preparation, labelling and quality assurance, and the technical documentation in Annex IV amounts to a datasheet for each dataset — how it was obtained and selected, its main characteristics, and how it was labelled. As Section 13.3 sets out, those high-risk obligations were deferred by the Digital Omnibus to December 2027 for standalone systems and August 2028 for embedded products, so there is time — but the obligations on general-purpose model providers, including publishing a summary of training data, have applied since August 2025 and did not move.

Practically, this means treating metadata as infrastructure rather than documentation. The 2026 shift in the catalogue market is instructive: tools are being repositioned from a place people look things up to a live context layer that systems — including agents — query at runtime for lineage, permissions and definitions. Whatever tool you use, the durable artefacts are the same: a register of what data feeds which AI system, where it came from, who owns it, what may lawfully be done with it, and when it was last checked.

GOVERNANCE Provenance is the evidence, and you will be asked for it twice

The same records answer two different questions: the regulator asking how a high-risk system was built, and the customer or insurer asking whether your product was trained or grounded on material you had the right to use. Organisations that keep this only as tribal knowledge discover the gap during due diligence, when it is expensive.

Build the register once and map it to many demands (Section 13.7). The overlap between EU AI Act Annex IV documentation, ISO/IEC 42001 evidence, and an ordinary procurement questionnaire is large enough that maintaining three separate answers is pure waste.

8.5 Privacy and de-identification

De-identification is not anonymisation, and treating them as equivalent is the most common privacy error in AI projects. Removing direct identifiers leaves re-identification risk, because background knowledge does the rest — a rare diagnosis in a small town, a distinctive combination of role and location. Differential privacy is the only technique with a formal guarantee, adding calibrated noise so that no individual's presence materially changes the result; the trade-off is explicit and tunable, and the parameter choice belongs in the privacy impact assessment rather than in an engineer's head.

Australia. Australia's reforms are the ones to diarise, and they arrive in stages. The Privacy and Other Legislation Amendment Act 2024 received assent in December 2024. A statutory tort for serious invasions of privacy commenced in June 2025 — and unusually, neither the small-business exemption nor the employee-records exemption applies to it, so it reaches organisations the rest of the Act does not. From 10 December 2026, entities must disclose in their privacy policy the kinds of personal information used in automated decision-making and the kinds of decisions made substantially by it, where those decisions could significantly affect someone. Note the scope carefully: the obligation is not limited to AI, and a rules-

based script or a spreadsheet falls inside it just as readily as a model. The OAIC issued a consultation paper in May 2026 with final guidance expected around September 2026 — which means organisations are preparing for a December obligation against draft guidance. A second tranche, including removal of the small-business exemption, a right to erasure, and a “fair and reasonable” test, remains a government commitment rather than law.

European Union. The European Data Protection Board adopted guidelines on GDPR and AI training data in July 2026, closing off the assumption that web-scraped training data enjoys a blanket legitimate-interest basis. The test must be run genuinely and case by case — a real interest, necessity, and a balance against data subjects’ rights. The guidelines were still in consultation at the time of writing, so treat the direction as settled and the detail as provisional.

United States. Roughly twenty states now have comprehensive consumer privacy laws in force, with more commencing through 2026, and the AI-specific edges are starting to show: California’s SB 361 requires data brokers to disclose whether they sell personal data to developers of generative AI. Colorado’s reversal is the cautionary tale — its comprehensive AI Act was repealed and replaced with a narrower transparency law weeks before it was due to take effect (Section 13.5). Architect compliance around durable principles, not around any single jurisdiction’s current rules.

8.6 Retrieval: connecting AI to your own data

A model knows only what it was trained on and nothing about your documents. Retrieval-augmented generation is the answer, introduced as the second adaptation lever in Section 1.8 and sketched in Section 6.6: fetch relevant material at query time and put it in the prompt, producing grounded, citable answers. The concept is simple. The engineering is where systems fail, and it belongs here with the rest of your data.

The naive pipeline — split documents into fixed-size chunks, embed them, store them in a vector database, retrieve the nearest neighbours — still works for demonstrations and disappoints in production. The 2026 consensus architecture has four differences.

- **Chunk on meaning.** Fixed-size chunks with overlap are a reasonable default, not best practice. Semantic chunking — starting a new chunk where the meaning shifts rather than where the token count runs out — preserves arguments and clauses that fixed windows cut in half. Chunks should carry metadata: source, section heading, date, and permissions.
- **Search two ways, not one.** Run dense vector search and lexical BM25 search in parallel and fuse the results. Keyword search remains undefeated for exact matches — product codes, statutory references, unusual acronyms — precisely the terms where semantic similarity is least reliable. Pure vector search fails most visibly on the queries users think are easiest.
- **Rerank before you generate.** Retrieve broadly, then narrow. Fetch on the order of a hundred candidates, pass them through a cross-encoder reranker, and give the model only the best handful. Reranking is the single highest-leverage addition to a mediocre retrieval system.
- **Choose the store deliberately.** Vector databases stopped being interchangeable somewhere around 2025. Performance now varies sharply by workload and scale, and

the honest answer for many organisations is that a relational database with vector and full-text support is sufficient — a dedicated store should be a decision, not a default.

A recurring claim deserves addressing directly, because it resurfaces with each context-window increase. In January 2026 “RAG is dead” circulated widely on the strength of million-token contexts. It was wrong, and the reason is instructive. Long context and retrieval solve different problems: long context is good for reasoning across a whole document you already have; retrieval is for finding the right material in a corpus too large, too fresh, or too access-controlled to paste. Models also degrade as irrelevant material fills the window — “context rot” — so more context is not monotonically better. The practical rule is that retrieval wins when the corpus is large, when freshness matters, or when you need citations, and the dominant 2026 pattern is a hybrid: let an agent retrieve a focused subset, then reason over it with a long-context model.

That agentic turn is the real change. Retrieval is increasingly a decision the model makes — what to look up, whether the results were good enough, what to search next — rather than a fixed step that happens before generation. Knowledge-graph approaches extend this to multi-hop questions where the answer depends on relationships rather than passages. All of it sits inside what practitioners now call context engineering: deliberately designing everything the model sees on each call. As the discipline’s own maxim has it, most agent failures are not model failures but context failures.

8.7 Permissions: the failure mode nobody sees

This deserves its own subsection because it is the most consequential and least visible defect in enterprise retrieval systems. When a document is ingested and split into chunks, its access controls must travel with every chunk. If they do not — or if the synchronisation with the source system goes stale — a user’s query can return material from documents they were never entitled to open.

What makes this dangerous is that it fails silently. There is no alert, no data-loss-prevention flag, no anomaly for a security team to investigate. The system behaves exactly as designed; it simply answers a question using a document the asker should never have seen. Salary files, board papers, disciplinary records and draft redundancy lists are all typically stored in exactly the repositories organisations rush to index first.

The fix is to enforce permissions inside the retrieval layer rather than in the application on top of it: sync source-system access lists into the index as metadata, filter at query time against the requesting user’s identity, and re-synchronise on a schedule short enough that revoked access actually takes effect. Test it adversarially before launch — give a low-privilege account to someone who knows what confidential material exists and ask them to go looking. This is the control most likely to be raised in an enterprise security review, and the one most likely to have been deferred.

SECURITY Retrieval inherits every access-control mistake you have ever made

Indexing a shared drive does not just expose the permissions you set deliberately — it exposes the ones that drifted over a decade: the folder opened up “temporarily” in 2019, the departed employee’s files inherited by a group, the site whose default was set to organisation-wide by someone who did not read the dialogue. Retrieval turns latent misconfiguration into an answering system.

Treat indexing as a privilege escalation event. Audit permissions on a repository before you index

it, not after, and prefer indexing a deliberately scoped subset over pointing the crawler at everything and relying on filtering to save you. Combined with tool access, this is the data half of the lethal trifecta in Section 6.13.

8.8 Residency and sovereignty

Two ideas get conflated, and the distinction matters commercially. Data residency is about where data physically sits. Data sovereignty is about whose law reaches it. A provider can guarantee the first without delivering the second: infrastructure operated through a European subsidiary with local staff and separate key management still has a parent company subject to its home jurisdiction's legal process. When a vendor answers a sovereignty question with a map, they have answered a different question.

The market has moved quickly on both. All three hyperscalers now operate Australian regions with sovereign offerings — customer-held keys, cleared local staff, contractual carve-outs — and Microsoft has extended in-country processing for its Copilot products to a small set of countries including Australia. The Digital Transformation Agency's five-year whole-of-government arrangement with Microsoft, commencing July 2026, includes strengthened protections for government data and provisions aimed at accelerating public-sector AI adoption. In Europe, the Commission proposed a Cloud and AI Development Act in June 2026 to build out capacity and a common sovereignty framework, alongside the Gaia-X labelling scheme.

For most organisations the practical question is narrower than the political one: which specific data, under which specific obligation, must not leave which specific jurisdiction — and does the AI feature you are about to enable move it? The answer is frequently discovered after a business unit has already switched something on.

8.9 Synthetic data

Synthetic data is genuinely useful and routinely oversold. Where it works well is filling gaps rather than replacing reality: rebalancing rare classes so a fraud model sees enough fraud, generating adverse scenarios for stress testing, and — the most mature and least controversial use — populating development and test environments so engineers never touch production personal data. That last case alone justifies the technique for most organisations.

Where it fails is recursive training. Research published in *Nature* in 2024 demonstrated model collapse: train a model on the output of models trained on model output, and quality degrades measurably — perplexity roughly doubled after nine generations, and image models progressively forgot under-represented categories. The mechanism is intuitive once stated: each generation samples from the last, and the tails of the distribution — the rare cases that matter most — thin out first. The mitigation is to accumulate rather than substitute, keeping real data in the mix rather than replacing it.

The forecasting record here is a useful corrective. Gartner predicted in 2021 that 60% of AI training data would be synthetic by 2024; that did not happen as measured. Its current prediction runs the other way — that by 2027 a majority of data and analytics leaders will face serious failures from mismanaging synthetic data. Treat synthetic data as an

augmentation technique with a governance requirement attached: label it, track what was generated from what, and never let it enter a training or evaluation set unmarked.

8.10 Pipelines and the operating burden

The unglamorous truth is that most data work is maintenance. Fivetran's 2026 benchmark of senior data and technology leaders at large enterprises put the business exposure from pipeline failures at roughly US\$3 million a month — before counting the engineering time spent fixing them or the AI projects stalled waiting. The human cost shows in the workforce data: a survey of 600 data engineers found 97% reporting burnout and 70% saying they were likely to leave within twelve months.

Two structural problems sit underneath. The first is a mismatch of tempo: organisations built for nightly batch processing are deploying models into interactive and agentic use cases that assume current data, and bridging the gap with stale precomputed values degrades answers silently — no error, no alert, just worse. The second is freshness ownership: retrieval indexes drift out of date with their sources unless someone owns re-indexing, and “the index is three months old” is a common root cause of an AI system that used to be good.

For the business case, this is the line that pilots omit and Section 10.7 insists on. Data pipeline operation is a recurring cost with a named owner, or it is a future incident.

8.11 Where to start

A sequence that works, and is deliberately unambitious.

- Pick the use case first and identify precisely which data it needs. Resist enterprise-wide data programmes justified by AI in general — they are how five-year initiatives get funded and nothing ships.
- Audit permissions on the target repository before indexing anything, and scope the index deliberately.
- Measure quality against the task, not in the abstract: sample a hundred documents and record how many are current, complete, correctly attributed, and extractable.
- Build the register — what data, from where, owned by whom, usable under what terms — and keep it live. It is governance evidence, procurement evidence, and ROI evidence at once.
- Automate the pipeline before you scale the use case; a hand-prepared extract is a prototype, not a foundation.
- Instrument retrieval quality from day one, so that “the answers got worse” becomes a measurement rather than an argument (Section 9.10).

Go deeper: the data engineering, retrieval and privacy resources for this section are in Sections 17.5–17.6, with the Australian privacy sources listed in Section 17.6.

Next section: Section 9 — AI for Business: Adoption, Delivery & Project Management. With the substrate in place, the question becomes organisational: how to choose use cases, run them as projects, evaluate them properly, and get people to use what you build.

9. AI for Business: Adoption, Delivery & Project Management

The methodology here is durable; the statistics are indicative and dated (verified 29 July 2026). This is the delivery/PM anchor section — the discipline of turning AI capability into organisational value.

The previous sections were about what AI can do. This one is about why most organisations struggle to capture that value — and what the minority who succeed do differently. The striking finding of 2026 is that the gap is almost never the technology. The models are capable; the bottleneck is organisational. For a reader whose interest is project management and delivery, this is the heart of the guide.

9.1 The state of play: the “GenAI Divide”

Adoption is nearly universal and value is highly concentrated — that is the paradox. Surveys put AI use at around 88% of organisations in at least one function, yet fewer than 40% have scaled beyond pilots. The most-quoted figure comes from MIT’s NANDA study, *The GenAI Divide: State of AI in Business 2025*: about 95% of enterprise GenAI pilots delivered no measurable profit-and-loss impact, while only ~5% created significant value.

That number deserves context rather than panic. It measured rapid P&L impact within roughly six months, concentrated in sales and marketing — the lowest-ROI area studied. A new hire rarely moves the P&L in six months either. Framed more usefully, it is a value-realisation gap: roughly 29% of organisations report significant ROI, and the great majority of executives feel some benefit. The lesson is not that AI doesn’t work — it is that deploying a model and generating value from it are entirely different problems, and the second is an organisational discipline.

9.2 Why AI projects fail — and it isn’t the model

The diagnosis is remarkably consistent across MIT, Gartner, BCG, and RAND. The recurring root causes:

- **No defined outcome before the build.** Success is specified after launch instead of before, so nobody can say whether the pilot worked.
- **Data that isn’t ready.** Gartner attributes about 85% of failed AI projects to poor data quality, and expects a large share of projects lacking AI-ready data to be abandoned. Agents will not fix a data and governance problem — they inherit it.
- **Treating it as an IT project, not organisational change.** BCG’s widely-cited split is that AI transformation is 10% algorithms, 20% data and technology, and 70% people, process, and culture. Invest only in the first two and the pilot works in the lab and dies in the field.
- **Adoption assumed, not designed.** Accuracy in a demo is irrelevant if the workflow around the tool never changes and nobody uses it.
- **A build-vs-operate imbalance.** Money goes into building the pilot and model selection; too little into the evaluation, monitoring, and operations that production actually requires.

The COO's summary quoted in the MIT research captures the mood: the hype says everything has changed; on the ground, nothing fundamental has shifted about how organisations must execute.

9.3 Choosing where to apply AI

Where the money goes and where the returns are do not match. Budgets skew heavily to sales and marketing, but the biggest, most direct ROI tends to sit in back-office and operations — cutting outsourcing, streamlining process-heavy work, reducing error rates. The winning profile is high domain-specificity plus deep workflow integration; the losers are flashy, generic pilots that never touch a real workflow. Practical selection discipline:

- **Map value against effort** and pick a small number of candidates; start deliberately small (around 73% of successful deployments did, per Stanford's 2026 playbook) and frame first pilots as controlled experiments, not enterprise rollouts.
- **Define the P&L metric in week one.** Pick two or three KPIs, set a baseline and target, and name an owner for each number before any building starts.
- **Favour high-friction, high-value work** where AI can be embedded in the flow rather than bolted alongside it.

9.4 Build vs buy

One of MIT's clearest findings: buying from specialised vendors or partnering succeeds roughly twice as often as building in-house (about 67% vs 33%). The reason is rarely raw capability — it is that vendors bring workflow fit, domain fluency, and adoption experience, while internal teams routinely underestimate integration and operating cost. The sensible pattern is hybrid: buy for common needs (chat, document processing, transcription), build only where the use case is proprietary and a genuine source of competitive advantage, and treat data-and-workflow fit as the deciding criterion. One caveat to weigh: deep integration creates switching costs, so choose vendors whose systems learn and adapt to your processes, because that is what you will be locked into.

9.5 Running AI as a project: the delivery discipline

The core mindset shift is to treat AI as an operating-model change, not a one-off experiment — and to fund it through gates like a capital asset, not indulge it like a science project. A workable lifecycle:

Phase	Focus	Done when
Strategy	2–3 KPIs, one-page value thesis, owners, data check	Baselines and targets set; owners named
Readiness	Honest audit of data, governance, skills, infra	Gaps known before the build starts
Use-case selection	Value-vs-effort; start small; back-office bias	1–2 pilots chosen as experiments
Build / buy	Buy common, build proprietary; hybrid	Decision made on fit, not by default
Pilot	Controlled experiment with evals wired in	Meets business (not just technical) targets
Scale	Ops function, monitoring, named ownership first	Production-ready before go-live
Operate & govern	Continuous evaluation, monitoring, audit	Value tracked; governance live

The dangerous stretch is the pilot-to-production cliff: pilots are easy to start and hard to finish, and only around 14% of organisations have scaled an agent to org-wide use even

though most have a pilot. The teams that cross it are distinguished by three habits: they appoint an AI-operations function (owning monitoring, evaluation, and incident response) before scaling rather than after an incident; they have automated evaluation infrastructure running before the first production task; and they spend proportionally more on operating than on model-picking. Evaluations are the project-management discipline for non-deterministic systems — a maintained test set, run continuously like regression testing (ties back to Section 4.5 and Section 6.11). One security note: connecting AI to live enterprise systems expands the attack surface (data exposure, identity, poisoning), which Section 14 addresses.

DELIVERY / PM The two habits that separate the 5%

First, define the success metric before you build — a P&L or operational KPI with a baseline, a target, and a named owner in week one — not a technical score like accuracy after launch. Second, stand up the operations and evaluation function before you scale, not in response to the first incident: organisations that established ownership and eval infrastructure pre-scale were dramatically less likely to roll back. Everything else — executive sponsorship (its absence makes failure roughly 2.5× more likely), kill/scale gates, a central platform with decentralised delivery — supports these two.

9.6 The people side: change management and adoption

If AI transformation is 70% people and process, change management is not a soft add-on — it is the work. Roughly 70% of change efforts fail when the affected staff aren't meaningfully involved, and post-mortems repeatedly identify the same fix: the AI Champions model, embedding one trained advocate in each affected department rather than running a single all-hands session. Effective change management here means three concrete things: redesigning the process (who does what, and what the human–AI handoff looks like), role-specific training (workflow-level instruction for the actual system, not generic AI awareness), and a feedback loop so frontline users can report issues to whoever owns the system. A visible workforce split is emerging too — “AI-elite” super-users reporting large time savings while others lag — which makes deliberate enablement, rather than assumed osmosis, the difference.

GOVERNANCE Bring shadow AI into the light

Around 90% of workers already use personal AI tools for work while only ~40% of organisations provide sanctioned ones — a governance blind spot where proprietary data flows into unvetted services. The response is not prohibition (that just drives it underground) but an AI inventory: discover what tools are actually in use, provide sanctioned equivalents with proper data-handling terms, and build governance in from day one as enabling infrastructure rather than bolting it on later, when use-case sprawl has outrun it. Section 13 develops this into a full governance picture.

9.7 Driving and measuring adoption

Section 9.6 covered the change management that surrounds a deployment. This subsection is about the thing that determines whether any of it mattered: whether people actually use what you built. It is the most reliably under-managed part of AI delivery, and the numbers are stark enough to be worth stating plainly before the remedies.

The gap. IBM's 2026 study of more than two thousand chief executives found around 85% of employees have access to AI capability while only about 25% use it regularly — a sixty-

point gap, and the reason those executives now rank adoption above technology as their primary concern. Deloitte's enterprise research points the same way: access to sanctioned tools rose sharply in a year, to around 60% of workers, but of those with access fewer than 60% use it in daily work, essentially unchanged year on year. Independent polling shows roughly half of US workers using AI at least occasionally but only around 13–15% daily. Seat utilisation surveys put weekly active use at 20–30% of purchased licences.

The uncomfortable reading is that a ten-thousand-seat deployment with 15% weekly use is not an adoption success with a long tail. It is an 85% write-off being reported as a rollout, and it is the single largest destroyer of the returns modelled in Section 10.8, where the benefit scales directly with the adoption ramp.

Why people do not use it

The reasons are better evidenced than most of this field, and only one of them is about the tool.

- **Social penalty.** The strongest finding here is also the least discussed. A peer-reviewed study across four experiments with around 4,400 participants found that workers who used AI were judged by colleagues and managers as lazier, less competent and less diligent than people who produced identical work manually — and that managers were less likely to hire or promote them. Independent polling corroborates the mechanism from the other side: around 43% of workers say they trust a colleague's work less when they know AI was involved, against 20% who trust it more. Roughly half of workers hide their AI use. Notably, the penalty shrinks when the evaluator also uses AI — which makes visible leadership use a control, not a gesture.
- **Status threat.** Survey work on around 2,257 employees found that status threat — the fear of appearing replaceable — predicted reduced depth of use, while job design factors, particularly autonomy and skill variety, were the most consistent positive predictors. People who feel their judgement is the point of their job use the tool more, not less.
- **Unclear permission.** Around 44% of workers say their employer has no clear AI policy or they are unsure whether one exists, and the figure is worse in small organisations. This does not stop use; it drives it underground (Section 13.8).
- **The saving is invisible or eaten by rework.** Employees widely report saving hours a week, but a substantial share of that saving is consumed by checking and correcting output — one large study put the net gain at a small fraction of the reported figure once rework was counted. Where the saving is invisible or illusory, continued use is rational to abandon.
- **Workflow friction.** And the most under-appreciated: a slightly less capable tool inside the workflow someone already uses will beat a more capable one that requires opening a separate window and changing a habit.

Shadow tools and why yours loses

The corollary of low sanctioned adoption is high unsanctioned adoption. Analyst survey work in early 2026 found around 88% of employees who have access to enterprise AI tools also use personal ones for work, and that 58% actively prefer the unapproved tool. Section 13.8 covers the policy response; the delivery question is why the sanctioned option loses, and the answers are consistent: the approved tool is a cut-down version of the popular one; approval takes months while the alternative takes seconds; the policy sits in a PDF nobody has read

since induction; and employees make a rational trade-off in which hitting the deadline beats an abstract compliance risk. Every one of those is a delivery decision, not a discipline problem.

What actually moves the number

Four interventions have real evidence behind them, and they are not the ones organisations usually fund first.

- **Managers using it visibly.** The best-measured finding in this area. Large-scale research combining survey and product telemetry across twenty thousand workers found that when managers visibly use AI themselves, employees' perceived value of it rises by around 17 points, critical thinking about its use by 22, and confidence with agentic tools by 30. Where managers create room to experiment without penalty, readiness rises around 20 points and staff are roughly 1.4 times more likely to become frequent users. Mandating use while not using it yourself is the worst available combination.
- **Champions embedded in teams.** Peer-nominated champions appear in a large majority of high-adoption organisations and a small minority of low-adoption ones. The mechanism is translation: a champion demonstrates the tool on work the team actually does, which no central training session can.
- **Role-specific training, aimed at existing users too.** Employees who receive training use the tools regularly at far higher rates than those who do not — one study reports roughly 93% against 57%. But note where training is misallocated: leaders overwhelmingly say it is a priority while only around a third of active users get further training. The people already using it are the ones whose depth of use is cheapest to improve.
- **A strategy, not a tool.** Organisations with an explicit AI strategy report substantially greater impact than those buying better tools without one — in one large study, a 25-percentage-point difference against 5. Adoption follows a clear answer to “what is this for?”

On mandates the evidence is genuinely mixed and worth reporting honestly, because they are increasingly popular. A longitudinal study of a real mandate at a software company — 802 developers, some 196,000 pull requests over two years — found the throughput target was met, roughly doubling merged work per engineer. It also found the obvious limit: review capacity did not double with it. Against that, around 22% of workers say they would consider leaving a job that forced AI use in ways they disagreed with, and reported usage can rise while trust falls. A mandate can move a usage metric. Whether it moves value depends on whether the constraint downstream was ever the thing you were measuring.

DELIVERY / PM Protect the first experience — trust does not recover symmetrically

Human–computer interaction research finds that a single wrong answer causes users to generalise “this tool makes mistakes” and to discount subsequent correct outputs too, that errors early in a relationship with a tool damage reliance more than the same errors later, and that trust recovers only partially with sustained good performance. Below roughly 70–85% accuracy on a task, depending on consequence, baseline trust never forms at all.

The delivery consequence is that pilot cohort selection is a trust decision, not a convenience one. Launch on the use case where quality is highest rather than where the volunteers are loudest, fix the failure modes before widening, and accept that a cheap early embarrassment can cap

adoption in that population permanently — regardless of how good the system later becomes.

Measuring it

Two distinctions make adoption measurable. The first is between activation — the moment a user first gets real value — and adoption, which is the tool becoming the default way that work is done. Most organisations declare success at activation. The second is between leading indicators, which move first and can be acted on, and lagging indicators, which are what you were actually trying to change (Section 10.5).

Indicator	Type	Definition	What it tells you, and the trap
Activation rate	Leading	Share of provisioned users who complete one genuine task end to end.	Whether onboarding works. Trap: it is the metric most often reported as “adoption”.
Weekly active share	Leading	Weekly active users as a proportion of eligible users, not of licences bought.	The honest headline number. Trap: monthly actives flatter it considerably.
Depth of use	Leading	Tasks per active user, and how many distinct use cases each person applies it to.	Whether it has become habit or stayed a novelty. Trap: high activity on trivial tasks.
Workflow penetration	Leading	Share of eligible transactions in a target workflow handled with the tool.	The link between usage and the process you are trying to improve — the metric that connects to the ROI ramp.
Twelve-week retention	Leading	Users still active twelve weeks after activation, as a share of those activated.	Whether adoption is real or a launch spike. Trap: measuring at four weeks, while enthusiasm lasts.
Quality and rework	Lagging	Error and correction rates on AI-assisted output versus baseline.	Whether apparent savings survive verification. Trap: omitting it makes every other number look better.
Cycle time and cost per unit	Lagging	The operational metrics defined at G1 (Section 9.9).	Whether adoption converted into performance. Trap: attributing movement without a counterfactual (Section 10.6).
Sentiment and trust	Lagging	Whether users would object to the tool being removed; whether they disclose using it.	The leading indicator of next quarter's usage. Non-disclosure signals the social penalty is operating.

Report these by cohort rather than in aggregate. Averages conceal the pattern that matters — a team at 80% and four teams at 5% is a completely different management problem from five teams at 20%, and produces an identical mean.

Sequencing and licensing

Staged rollout beats big-bang for the same reasons it does in any delivery context, with an AI-specific addition: staging gives you the comparison group that makes the value claim credible (Section 10.6). A workable shape is a controlled pilot of roughly 50–200 users with an explicit readiness gate before each expansion, rather than a date-driven ramp.

Licensing is a genuinely open question and the guide will not pretend otherwise. Giving everyone a seat maximises reach and, on current utilisation figures, wastes most of the spend. Making people earn a seat concentrates value and risks excluding the quiet majority who would have benefited. No controlled comparison exists. What is defensible is to stop treating seats purchased as a measure of anything, and to review allocation against actual use on a cadence — which requires the measurement above to exist first.

The Australian picture. Australia is a useful case in point, and a reminder that adoption is a national as well as an organisational problem. National AI Centre tracking put small and medium business adoption at around 44% in early 2026, with trust the dominant barrier — roughly 65% of non-adopters cite distrust of AI decision-making or a preference for human control rather than cost or capability. In the public sector, the Australian Public Service AI Plan requires a chief AI officer in every agency by July 2026 and has moved a whole-of-government assistant through alpha into beta during 2026. Note also what is missing: there is no good Australian equivalent of the access-versus-usage gap data quoted above, so local figures should not be inferred from US studies.

9.8 Measuring value

The final discipline is honest measurement. When a pilot becomes production, the success criteria must move from technical metrics (accuracy, latency) to business outcomes: cost per unit processed, cycle time, error rate in the target process, adoption. Two realities to build into any business case. First, total cost of ownership: three-year costs commonly run 40–60% above the initial build quote, because initial development is only a quarter to a third of the lifetime cost — operations and governance are the rest. Second, returns are real but concentrated: the oft-cited average of about \$3.70 back per \$1 accrues to organisations deploying across multiple functions, not to those running isolated pilots. Beware vanity metrics (tools deployed, seats bought, messages sent); measure whether the work actually got cheaper, faster, or better.

9.9 The delivery playbook

Sections 9.3–9.8 describe the discipline; this subsection turns it into four artefacts you can run tomorrow. They are deliberately plain — a scorecard, a set of gates, one worked example, and a cost table — because the failure mode they prevent (Section 9.2) is rarely a lack of sophistication. It is skipping the boring steps.

1. The use-case selection scorecard. Score each candidate 1–5 per criterion and rank by total; anything scoring 1 on error tolerance or containment needs redesign before it proceeds, whatever its total. Re-run the exercise quarterly — model capability shifts what is feasible (Section 16).

Criterion	What to ask	Scores high when...
Value	Is there measurable, recurring pain — hours, cost, cycle time, error rate — with a baseline you can state today?	The pain is quantified and someone already owns the metric.
Feasibility	Does the data the task needs exist, and is it accessible, clean enough, and permitted for this use?	Inputs are digital, retrievable, and legally unambiguous.
Error tolerance	What does one wrong answer cost, and how cheaply can a human catch it before it lands?	Errors are cheap, visible, and verification is faster than doing the task.
Volume & repetition	How often does the task recur, in how standard a form?	High frequency, standard shape — the automation dividend compounds.
Containment	What could the system reach if it misbehaves — data, money, customers, reputation (Section 14)?	Blast radius is naturally small or cheap to fence (least privilege, approval steps).
Buy option	Does a vendor already sell this as a feature	No adequate buy option exists — build

	of software you own (Section 9.4)?	effort is actually warranted.
Ownership	Is there a named business owner who wants the outcome (not an innovation team placing bets)?	A sponsor with the problem, budget, and authority to change the workflow.

2. Pilot-to-production gates. Most stalled pilots (Section 9.1) were never designed to pass a gate — they were designed to demo. Write the gates before the pilot starts, and let ‘no’ be a normal, cheap outcome: a pilot that fails at G2 for \$30K has done its job.

Gate	The question it answers	Evidence required to pass
G0 • Screen	Is this worth a pilot at all?	Scorecard above; rough token/cost estimate (Section 12); risk tier assigned in the AI register (Section 13).
G1 • Design	Do we know what good looks like?	Baseline metric stated; evaluation set built (see below); threat model done (Section 14); data-use terms confirmed.
G2 • Pilot exit	Does it actually work, economically?	Eval scores vs baseline on the golden set; unit economics per task vs human cost; real users completed real work with it; error review of a graded sample.
G3 • Production	Can we run it safely at scale?	Monitoring and rollback in place; named owner; guardrails and least-privilege verified; governance sign-off; support and training ready (Section 9.6).
G4 • Operate	Is it still working?	Evals re-run on every model/prompt/tool/data change; value metric tracked against the G1 baseline; periodic re-red-teaming.

3. Evals as gates — a worked example. Suppose the use case is contract summarisation for a legal team. Before the pilot (G1), assemble a golden set: fifty representative contracts with human-written reference summaries, including the ugly ones — scanned, amended, non-standard. Define the metrics that reflect actual harm: material-clause recall (did it miss anything a lawyer would flag?), factual fidelity (does every claim trace to the source?), and format compliance. Set thresholds a stakeholder signs, for example: $\geq 95\%$ clause recall, zero fabricated citations, $\leq 5\%$ summaries needing rework. Automate the grading with a second model, but calibrate it — humans grade a 10% sample each run and the automated grader must agree with them before its scores count. The eval then runs at G2 (pass/fail), and again at G4 on every change: a model upgrade, a prompt tweak, or new contract types all re-open the gate. This is the single highest-leverage delivery habit in this guide: it converts ‘it seems good’ into a decision an organisation can defend (Section 15 explains why eyeballing fails).

4. The honest TCO table. Pilots systematically understate steady-state cost because the visible line — inference — is often the smallest one. Estimate all eight before G2, not after G3.

Cost line	Why it gets missed
Inference (tokens)	Scales with adoption and prompt bloat, not licence seats — success makes it bigger (Sections 1.2, 12.1). Budget per task, times realistic volume.
Human review time	The largest hidden line for high-stakes work: someone checks the output, and their hourly rate is part of the unit economics.
Evaluation & monitoring	Golden sets, graders, dashboards, and the time to maintain them — the price of knowing it still works.
Model churn tax	Models deprecate and improve on a monthly cadence (Section 2); every swap means re-evals, prompt fixes, and regression time.

Retrieval & context plumbing	Vector stores, connectors, permission-aware retrieval (Section 6) — infrastructure with its own run cost.
Integration & permissions	Identity, scoped credentials, audit logging (Sections 6, 14) — the enterprise plumbing that pilots skip and production cannot.
Change management & training	Adoption is a people project (Section 9.6); unbudgeted training is how tools end up unused.
Governance & audit evidence	Registers, risk assessments, incident processes (Section 13) — cheap as a by-product of delivery, expensive as archaeology.

Run together, the four artefacts give a delivery lead what Section 9.2 said most failed projects lacked: a reason to believe, checked at every stage, that this use case, with this data, at this cost, is worth shipping — and a paper trail that doubles as governance evidence (Section 13.7).

9.10 Evaluation: the discipline underneath the gates

Evaluation has appeared in every section that touches production — agents in Section 6, generated code in Section 7, the gates in Section 9.9, the measurement families in Section 10.5, reliability in Section 15. This subsection is its home. If Section 9.9 asks how you decide whether to ship, this asks how you know anything at all about what you are shipping.

The core distinction is between a benchmark and an eval. A benchmark measures a model's general capability and is useful for procurement. An eval measures whether your system — your prompt, your retrieved context, your tools, your model, your guardrails — does your task correctly, and it is the only thing that tells you anything about production. Public benchmarks cannot substitute: the general knowledge tests have saturated above 90% and no longer separate frontier models, coding benchmarks suffer contamination from training data, and as Section 7.3 notes, the same model's score on an agentic coding benchmark can move by around twenty-five percentage points depending on the scaffolding around it. A leaderboard position is a hypothesis about your use case, not evidence about it.

Build the golden set first. The foundation is a golden set: representative inputs with known-good outputs, held stable so that scores are comparable over time. For a single feature, fifty to a hundred well-chosen examples is enough to catch real regressions; a hundred to three hundred is the working range for a production system, beyond which the marginal case costs more in compute than it returns in confidence.

Composition matters more than size. A useful set has four parts: a stratified sample of genuine production traffic; deliberately adversarial cases; constructed edge cases covering the ugly inputs — scanned, amended, malformed, non-standard; and replays of every failure that has already reached a user. That last category is the most valuable and the most often skipped, and it is what turns an eval suite into an institutional memory rather than a test script. Refresh it on a cadence: add newly observed failures monthly, sweep for new adversarial patterns quarterly, and re-baseline annually. A golden set that never changes is measuring a system that no longer exists.

Design metrics around harm. Split metrics into two kinds. Deterministic checks — schema validity, format compliance, required fields, forbidden content — are free, fast, reproducible, and should run on everything. Graded metrics handle what cannot be checked mechanically, and they should be defined by the harm they detect rather than by convenience: recall of material facts, groundedness in the retrieved sources, citation accuracy, correct refusal.

Resist generic text-similarity measures. The classical overlap metrics correlate with human judgement at roughly 0.25 to 0.45 — barely better than noise for this purpose — and embedding similarity, while better at around 0.70 to 0.85, carries a frequency bias that penalises precise, unusual vocabulary in favour of common phrasing. A model-based judge outperforms both, which is why the practice has moved in that direction despite its own problems.

LLM-as-judge, used carefully. Using a model to grade another model's output is now standard, and it works better than intuition suggests. The foundational 2023 study found strong judge models agreeing with human preferences over 80% of the time — about the rate at which two human annotators agree with each other. Structured judging approaches reach correlations with human ratings well beyond what automatic metrics achieve.

The biases are well documented and manageable. Judges favour whichever response they see first, favour longer answers regardless of added value, and favour output from their own model family. The mitigations are mechanical: randomise presentation order, mask model identity, control for length explicitly, and use a reference answer where one exists. The non-negotiable step is calibration — measure agreement between your judge and human raters on a labelled sample using Cohen's kappa before you trust it, treating above 0.8 as strong and below 0.6 as a signal that the rubric, not the judge, needs work. Raw percentage agreement is misleading on imbalanced data: a judge that passes everything scores 90% agreement on a 90%-pass set while measuring nothing. Re-check monthly, because judges drift when the underlying model is updated.

Retrieval and agents need different measures. Retrieval systems fail in two distinct places and need to be measured in both. The standard decomposition is context relevance (did retrieval find the right material), groundedness (is every claim in the answer supported by what was retrieved), and answer relevance (does it address the question). Groundedness is measured by decomposing the response into individual claims and checking each against the context — which is also the most practical hallucination detector available. Note that hallucination rates are strongly domain-dependent: a model measured at around 2% on general summarisation can exceed 9% on specialist technical material, so a published rate for a frontier model tells you very little about your corpus.

Agents are harder again, because the output is a trajectory rather than an answer. The properties worth grading are tool selection, argument correctness, whether the agent actually used what the tool returned, recovery from failure, plan coherence, and completion — and they must be graded from the transcript, not just the final response, because an agent can reach the right answer by an unacceptable route. Because agent runs are non-deterministic, single trials are meaningless: run several per task and track variance as a first-class result. The arithmetic explains why this matters. At 95% reliability per step, a ten-step task completes cleanly around 60% of the time and a twenty-step task around 36%; even 99% per step decays to roughly 37% over a hundred steps. Impressive single-turn accuracy and unreliable multi-step behaviour are the same system.

Treat evals as tests. Evals earn their cost when they run automatically. The working pattern is tiered: cheap deterministic checks on every change, the full judged suite nightly against a versioned dataset, and a statistical gate rather than a raw score comparison — with fifty examples you can only detect large regressions, and a small sample will produce false

alarms roughly as often as real ones. Around a hundred to two hundred cases is where a three to five percent quality change becomes detectable.

The highest-value moment for an eval suite is a model change, whether you chose it or a provider forced it. Run the new model in shadow against the same inputs and compare quality, latency, token consumption and cost before committing. Be alert to a subtler trap: reasoning-model generations have changed which parameters control behaviour, so a harness written around older assumptions can silently stop controlling what it thinks it controls, and produce a clean-looking comparison of two things that were never configured alike.

Then watch it in production. Offline evaluation tells you about the cases you thought of. Production tells you about the rest. Trace every request end to end — prompt, retrieval, generation, guardrails — so that quality degradation is observable rather than inferred from complaints. Sample live traffic into a human review queue, and feed those judgements back into both the golden set and the judge calibration. Watch for drift on the input side as well as the output side; the most common cause of a system getting worse is that the questions changed, not the model.

Not everything can be judged offline. Some quality only shows up against real users, so the production counterpart to the golden set is controlled online experimentation: release behind a flag, roll out as a canary to a small slice of traffic first, and A/B or interleave the new version against the old on the metrics that matter — task success, escalation rate, user acceptance — rather than offline scores alone. Shadow-run a risky candidate against live inputs without showing users the result. The discipline is the ordinary one of engineering: change one thing, measure it against a control, and be able to roll back in minutes.

On tooling, the landscape is healthy and converging. Open-source options cover tracing and offline scoring well, commercial platforms add managed workflows and production guardrails, and the OpenTelemetry project's generative-AI semantic conventions are emerging as the interoperability standard — still largely experimental as at mid-2026, but the right thing to instrument towards, since it makes traces portable between tools rather than locking evaluation into one vendor.

Evals are governance evidence. The last argument for doing this properly is not technical. Evaluation records are the evidence base every assurance framework asks for. The NIST AI RMF's Measure function is operationalised through exactly this material; ISO/IEC 42001's performance-evaluation clause requires monitoring, measurement and internal audit of it; and the EU AI Act makes technical documentation the cornerstone of conformity assessment, with post-market monitoring obligations that require the evidence to stay current rather than being produced once. Logged golden-set results, judge-calibration records, red-team findings, drift reports and regression history are what an auditor or notified body will actually examine. Teams that build evals for engineering reasons discover they have built most of their compliance evidence as a by-product — and teams that build neither discover it in the opposite order.

Go deeper: the delivery/PM lens resources — including the evaluation material this section leans on — are in Sections 17.5–17.6.

9.11 Choosing and contracting an AI vendor

Build-versus-buy (Section 9.4) usually resolves toward buy, which turns the delivery question into a procurement one: how to evaluate and contract a vendor whose product is non-deterministic, updates under you, and often runs on a model it does not itself own. The criteria are scattered through this guide — model fit (Section 2), security posture (Section 14), governance and data terms (Section 13) — so this subsection pulls them into one checklist to put in front of a vendor before signing.

Area	What to ask for	Why it matters	Red flag
Data & training terms	A written no-training-on-your-data guarantee and clear retention limits	Your inputs may otherwise train a shared model or persist indefinitely	"It's covered by our privacy policy" with no contractual term
IP & indemnification	Ownership of outputs and indemnity against third-party IP claims	Output copyright is unsettled (Section 13.10) and training-data litigation is live (Section 5.6)	No indemnity, or one capped far below plausible exposure
Security posture	SOC 2 / ISO 27001, and a posture mapped to the OWASP LLM Top 10 and NIST AI RMF	AI expands the attack surface (Section 14); procurement and insurers now ask	Certifications "in progress" with no date, or no AI-specific controls
Reliability & SLAs	Uptime commitments, and whether you depend on a preview-tier model with weaker guarantees	Preview models can carry no contractual SLA (Section 2)	A production dependency on a preview endpoint with no fallback
Exit & portability	Data export, prompt and config portability, and a documented off-ramp	Deep integration creates switching costs (Section 9.4) — lock-in is what you buy into	No export path, proprietary formats, or punitive exit terms
Residency & sub-processors	Where data is processed, and a current sub-processor list	Residency is a legal requirement in parts of the EU and Australia (Section 8.8)	Undisclosed sub-processors or data routed through unexpected regions
Incident notification	Contractual breach- and incident-notification windows	You inherit the vendor's incidents; regulators expect timely reporting (Section 13)	Vague or absent notification obligations

The through-line is that an AI contract is a risk-allocation document, not a feature list. The two clauses that most often turn out to matter are the ones teams skip in the demo glow: what happens to your data, and how you get out. Weight them accordingly, and prefer vendors whose systems adapt to your processes — because adaptive fit is precisely what you will be locked into (Section 9.4).

9.12 AI by function: where it lands and what it risks

This guide is deliberately horizontal — the principles hold across industries — but readers work inside a function, and it helps to see where the value and the risk concentrate by role. The map below is a starting orientation, not a substitute for the use-case selection discipline in Section 9.3; the pattern that repeats is that the highest-value uses and the sharpest risks tend to sit in the same place.

Function	Highest-value uses	The principal risk to watch
Software engineering	Code generation, review, test-writing, migration, and agentic multi-file work (Section 7)	Insecure generated code and unreviewed volume (Sections 7.5–7.6)
Customer service & support	Deflection, agent-assist, summarisation, and triage	A confident wrong answer to a customer, and eroded trust (Section 9.7)
Sales & marketing	Content drafting, personalisation, research,	Brand-voice drift, fabricated claims, and

	and outreach at volume	consumer-law disclosure exposure
Legal & compliance	First-pass review, research, summarisation, and clause extraction	Hallucinated authority and privilege or confidentiality leakage (Section 15.2)
Finance & accounting	Reconciliation, analysis, forecasting narratives, and reporting	Silent numerical error and auditability gaps (Section 15.4)
HR & recruiting	Screening, drafting, and knowledge-base answers	Bias, and the heaviest external regulation of any internal use (Sections 11.8, 13.8)
Healthcare & clinical	Documentation, coding, triage support, and literature synthesis	Patient-safety stakes and strict oversight — keep a human decision-maker
Operations & back-office	Process automation, document extraction, and data cleanup	Data-readiness and permission failures (Section 8) — often the highest real ROI
Research & analysis	Synthesis, literature review, and drafting	Over-trust and un-sourced claims; ground and verify (Sections 4.8, 15.5)

Two cross-cutting notes. The functions with the biggest headline potential — legal, healthcare, finance — are also the most regulated and least forgiving of a wrong answer, so they reward grounding, human oversight and evaluation most heavily. And the least glamorous — operations and back-office — is where Section 9.3's point about where the real ROI sits most often proves out.

9.13 The organisation's first ninety days

Section 17.10 gives an individual a thirty-day plan; this is its organisational counterpart — a first-ninety-days sequence for a team or business starting deliberately, built from the delivery discipline in this section and the governance in Section 13. It is intentionally modest: the failure mode is a grand programme that dies in the pilot-to-production gap (Section 9.5), not a small one that compounds.

- **Days 1–30 — see clearly and pick one.** Stand up an AI inventory (discover sanctioned and shadow tools, Section 9.6); write the short acceptable-use policy and tool list (Section 13.8) so people know what is permitted; and choose a single high-friction, high-value use case with a named owner and a P&L or operational metric defined before any build (Sections 9.2–9.3).
- **Days 31–60 — build small and instrument it.** Run the one use case as a controlled experiment against its baseline; stand up evaluation before the first production task, not after an incident (Section 9.10); model the data's permissions and quality honestly (Section 8); and threat-model the use case with least-privilege access (Section 14).
- **Days 61–90 — decide and institutionalise.** Judge the pilot against its metric and the counterfactual (Section 10.6); if it clears the bar, name the AI-operations owner and the governance accountability before scaling (Sections 9.5, 13.7); if it does not, kill it cleanly and bank the learning. Either way, capture what you built — prompts, evals, policy — as reusable assets, not tribal knowledge.

The point of ninety days is not speed for its own sake but a full loop — inventory, govern, ship one thing, measure it honestly, and decide — because the organisations that capture value are the ones that learned to finish, not the ones that started biggest.

Next section: Section 10 — Measuring AI Value & Return: the question that follows delivery. Six kinds of return, the frameworks that exist, the metrics that matter, the counterfactual problem, and a worked business case end to end.

10. Measuring AI Value & Return

The method here is durable; the benchmarks and survey figures are indicative and dated (verified 1 August 2026). This is the value anchor section — the companion to Section 9, which is about shipping AI, and to Section 12, which is about what AI costs the industry rather than what it returns to you.

Section 9 was about getting AI into production. This one is about the question that follows and that most organisations still answer badly: did it pay? After three years of extraordinary investment, there is still no agreed method for measuring the return on an AI initiative. What exists is a patchwork — consultancy frameworks, vendor-commissioned business cases, a thin but growing academic literature, and a thirty-year-old body of work on IT benefits realisation that the AI field has largely ignored. This section maps that patchwork, sets out a taxonomy you can actually use, works a full example end to end, and is candid about what nobody can measure yet.

10.1 Why AI breaks the standard ROI formula

The formula itself is trivial: return on investment is net benefit divided by cost. Applied to AI, four things go wrong with it, and they go wrong together.

- **The benefits are diffuse and largely non-cash.** Most AI benefit arrives as time, attention, or quality rather than cash. Time saved is not money saved unless the freed capacity is redeployed or not replaced. The conversion step is where most business cases quietly fail.
- **The costs are back-loaded and under-scoped.** Build is the visible cost and the smallest one. Evaluation, monitoring, model migration, change management, and governance are the rest, and they recur annually while the build does not.
- **The counterfactual is unobservable and non-stationary.** You cannot observe what would have happened without the tool, and the tool itself improves under you. A twelve-month before-and-after comparison silently bundles the AI effect with a process redesign, a staffing change, and two model upgrades.
- **The denominator moves fast.** For a fixed capability tier, inference prices have fallen roughly tenfold a year, so the cost side of a business case written today may be materially wrong within twelve months — usually too pessimistic on unit cost and too optimistic on total spend, because cheaper tokens invite far more of them (Section 12.2).

The clearest evidence that the field has not settled this is that the two most-cited 2026 studies reached opposite conclusions about the same technology in the same year. MIT's NANDA work found that roughly 95% of enterprise GenAI pilots produced no measurable profit-and-loss impact within about six months. Wharton found that around 74% of organisations measuring returns were seeing a positive one, with 72% reporting a structured measurement process. Both are defensible. They measured different things: MIT tested rapid P&L movement, Wharton counted productivity, efficiency, and retention. McKinsey's figures sit between them and are the most sobering — about 39% of organisations attribute any EBIT impact to AI at all, and only around 6% report more than 5% of EBIT attributable to it.

The lesson is not that one camp is lying. It is that “AI ROI” is not a single quantity, and any headline figure is meaningless until you know which kind of return was counted, at what level, against what baseline. The rest of this section is those three questions.

10.2 Six kinds of return

Almost every claimed AI benefit is one of six types. They are ordered here by how directly they reach the P&L, which is the inverse of how easily they can be claimed — a pattern worth keeping in mind when reading anyone’s business case, including your own.

Type of return	What it is	How it reaches the P&L	Difficulty and honest caveat
1. Cost reduction	Work that no longer needs doing, or needs fewer people or vendors to do. Outsourcing cut, contractors released, headcount not backfilled.	Directly and immediately.	Easiest to measure, hardest to enact. Requires a real decision about the freed capacity.
2. Productivity	The same people doing more, or the same work in less time. The most-claimed category by a wide margin.	Only if the freed time is redeployed to something the organisation values.	The conversion step is almost never specified. Self-reported time savings are unreliable (Section 10.6).
3. Revenue enablement	More leads worked, faster quotes, shorter sales cycles, new products that were not viable before.	Directly, but through a noisy channel.	Attribution is hard — revenue moves for many reasons at once. Needs a control group to be credible.
4. Quality and error reduction	Fewer defects, less rework, higher first-pass yield, better compliance outcomes.	Through avoided rework cost and, sometimes, avoided penalty.	Measurable if you had a baseline error rate. Most organisations did not.
5. Risk and loss avoidance	Fraud caught, incidents prevented, breaches that did not happen, regulatory exposure reduced.	As an avoided cost, not a booked gain.	Counterfactual by nature. Real, but easy to inflate and impossible to audit precisely.
6. Strategic and option value	Capability built, data assets created, talent retained, the ability to move faster later.	Indirectly, over years, if at all.	Genuinely real and genuinely unfalsifiable — which makes it the refuge of weak business cases.

Two practical rules follow. First, a business case should state which type or types it is claiming, in those terms, before it states a number. Second, most disappointing AI programmes are ones that promised type 1 and delivered type 2: the hours were genuinely saved, nobody decided what to do with them, and the finance function correctly declined to recognise the benefit.

Boards have noticed. Across 2026 the primary metric organisations report against has shifted away from productivity — down from roughly 23.8% to 18.0% as the headline measure — while direct financial impact has nearly doubled to about 21.7%. That is a rational response to three years of productivity claims that never showed up in the operating result.

10.3 Four levels of analysis — and why they do not add up

AI returns are measured at four quite different altitudes, and confusing them is the single most common error in public debate about whether AI is “working”.

- **Task and individual.** A controlled study of one person doing one kind of task. This is where the rigorous evidence lives — randomised trials, golden-set evaluations, before-

and-after timing on a fixed workload. It answers “does the tool help on this task?” and nothing more.

- **Use case.** One workflow, one team, one business case with costs and benefits over three years. This is the level Sections 9 and 10 are built around, and the level at which most organisational decisions are actually made.
- **Portfolio.** The whole AI programme, including the failures. This is the level a CFO cares about, and the level that exposes the survivorship bias in every published case study. Organisations reporting strong returns run notably fewer concurrent use cases — around 3.5 on average, against 6.1 for those not seeing returns.
- **Economy.** National productivity statistics. Australia’s Productivity Commission estimates a likely gain of about 0.4 percentage points a year, or roughly 4% added to labour productivity over the decade — meaningful, but an order of magnitude below the rhetoric.

The critical point is that these levels do not aggregate upward. A 30% task-level speed-up does not become a 30% team improvement, let alone a 30% margin improvement. Value leaks at every boundary: the task is only part of the job, the job is only part of the workflow, the workflow is bounded by upstream and downstream constraints that AI did not touch, and the organisation has to make a decision to bank whatever survives. Call it the leakage problem. It explains most of the gap between what pilots demonstrate and what the P&L shows, and it is the reason the macro figures look so modest next to the task-level results.

10.4 What already exists: the frameworks you can borrow

There is no standard for AI value measurement. There are, however, seven usable bodies of work, and between them they cover most of what you need. The table below is the honest state of the toolkit as at mid-2026.

Framework	Origin and status	What it is good for	What to watch
Forrester Total Economic Impact	Forrester Consulting; the closest thing to an industry standard, in use for over 20 years.	A defensible four-part structure — cost, benefits, flexibility, risk — modelled as NPV over three years.	Almost always vendor-commissioned. Borrow the method; never borrow the headline number as a benchmark.
IDC Business Value	IDC; structured customer interviews with statistical extrapolation.	Comparable quantified outcomes across a customer base rather than a single composite.	Same sponsorship problem, plus self-selection in who agrees to be interviewed.
Gartner outcome-driven metrics	Gartner; framed around three pillars for deriving value from AI.	Forcing each initiative onto a business metric owned by a department — cost per resolved claim, not model accuracy.	Proprietary and non-comparable between firms. Gartner itself reports that 45% of CFOs are funding the wrong outcomes.
Benefits realisation management	Cranfield School of Management and the wider IT investment literature; roughly thirty years old and mature.	The Benefits Dependency Network — mapping enablers to business changes to benefits to objectives on one page.	Predates AI and is written in IT-project language, so it needs translating. It is nonetheless the strongest tool available.
Academic RCTs	Brynjolfsson, Li & Raymond (QJE 2025); Noy & Zhang (Science 2023); the HBS/BCG study of 758 consultants; METR 2025.	The only genuinely causal evidence of AI effects on work, with control groups and pre-registered measures.	Task-level and narrow. None of it rolls up to an enterprise ROI figure, and results vary sharply by task and seniority.
ISO/IEC 42001	International standard for AI management systems;	Not a value framework, but the evidence trail it requires is	Governs the management system, not the return.

	certifiable.	reusable as measurement infrastructure.	Treating certification as a proxy for value is a category error.
FinOps and unit economics	Adapted from cloud cost management; increasingly the default for agent workloads.	Cost per resolved task or per accepted output — the denominator that makes outcome-based pricing legible.	Covers cost precisely and benefit not at all. Necessary, not sufficient.

The most useful item in that table is not an AI framework. Benefits realisation management was developed to answer exactly this question for IT investments — how do you connect a technical capability to a business outcome, and prove it afterwards — and it has three decades of practice behind it. A Benefits Dependency Network maps four linked layers: the technology enabler, the enabling changes needed to use it, the business changes that actually produce the benefit, and the measurable benefit itself, tied to an investment objective. Drawn for an AI use case, it exposes the missing link that kills most business cases — the business change nobody committed to make. Nobody has to invent AI ROI from scratch, and the field would be further along if it had stopped trying.

10.5 The metrics that matter

Metrics fall into six families. A credible business case draws from at least four of them, and the common failure is living entirely in the last two.

Family	Representative metrics	What it answers, and when to use it
Financial	NPV, IRR, payback period, three-year ROI, total cost of ownership, cost per unit of work.	Is this worth the capital? Required at the investment decision and at any gate where money is released. Discount the cash flows; AI benefits ramp, so undiscounted totals flatter year one.
Operational	Cycle time, throughput, cost per resolved task, containment or deflection rate, first-pass yield, queue depth.	Is the work actually getting cheaper or faster? The primary evidence layer once in production, and the one that connects most directly to a departmental owner.
Quality	Accuracy against a golden set, material-error rate, rework rate, human-override rate, escalation rate.	Is the output good enough to rely on? Set before the pilot, not after (Section 9.9), and define errors by the harm they cause rather than by string similarity.
Adoption	Active users as a share of eligible users, depth of use, share of eligible transactions handled, workflow penetration.	Is anyone using it? The leading indicator of every other number. Benefits scale with adoption, so a business case that assumes full adoption from month one is wrong by construction.
Technical and cost	Tokens per completed task, cost per successful outcome, retry and failure rate, latency, eval pass rate.	What does one unit of value cost to produce? Essential for agents, where a single workflow can consume 15,000–80,000 tokens against 500–2,000 for a simple query.
Counterfactual	Baseline measurement, control or hold-out group, difference-in-differences, pre-registration of the success criterion.	Would this have happened anyway? The family almost everyone skips, and the only one that makes the other five believable.

Beware the vanity layer beneath all of these: tools deployed, seats purchased, prompts sent, messages generated. These measure activity, not value, and they are the metrics most likely to be reported upward precisely because they always go up.

10.6 The counterfactual: the discipline almost everyone skips

Everything above depends on a comparison, and the comparison is where AI measurement is weakest. There are four workable approaches, in descending order of rigour and ascending order of practicality.

- **Randomised control group.** Split the eligible population, give half the tool, measure both. Expensive and politically awkward, and the only design that produces a defensible causal claim. Worth it for a flagship programme; overkill for a small internal tool.
- **Staggered rollout.** Roll out by team, region, or queue over time, and compare the change in adopting units against the change in not-yet-adopting ones. This is difference-in-differences, it fits a normal staged rollout, and it is the best value-for-effort option for most organisations.
- **Pre-measured baseline.** Measure the process for four to eight weeks before anything is built — volume, cycle time, cost per unit, error rate — and freeze the figures in writing. Weak against concurrent change, but infinitely better than nothing, and it costs almost nothing if done before the build rather than reconstructed after.
- **Self-report.** Ask people how much time they saved. This is the default method in industry and it is the one method known to be actively misleading.

The evidence against self-report is unusually direct. METR's 2025 randomised trial of experienced open-source developers working on their own mature codebases found participants were about 19% slower with AI tools while believing they had been roughly 20% faster — a swing of nearly forty percentage points between perception and measurement, in the direction that flatters the tool. That finding does not generalise to all tasks; the Brynjolfsson customer-support study found large and real gains, concentrated among less-experienced workers. What generalises is the gap between felt and measured productivity. Any ROI figure resting on a survey question should be labelled as an estimate of sentiment, not of performance.

DELIVERY / PM Write the counterfactual into the plan, not the post-mortem

The cheapest possible investment in AI measurement is four weeks of baseline data collected before anyone builds anything. It requires no technology, costs a few hours of someone's time, and is impossible to recover later — once the tool is live, the pre-state is gone.

Practically: name the two or three metrics at G1 (Section 9.9), record their current values with dates and sample sizes, identify a comparison group even if it is only the teams later in the rollout sequence, and write down in advance what result would count as failure. A pilot that cannot fail cannot demonstrate success either.

10.7 The other side of the ledger: counting the full cost

Section 9.9 set out the eight cost lines a pilot systematically understates. Two of them deserve expanding here, because they are where AI business cases most often go wrong in the second and third years rather than the first.

Operating cost, not build cost. Three-year costs commonly run 40–60% above the initial build quote, because development is only a quarter to a third of lifetime cost. The recurring lines — inference, evaluation, monitoring, incident response, model migration when a provider deprecates what you built on, and the human review that most production systems

still require — do not appear in a pilot budget at all. Agentic workloads make this sharper: published 2026 figures for agent-assisted resolution put the model and tool cost at a small fraction of the total, with human review and handling dominating, plus an explicit reserve for failed attempts. If your cost model contains only the API bill, it is understating the truth by roughly an order of magnitude.

Governance and compliance. Compliance is a real line item, and its size depends on where you operate and what you are building. Published 2026 estimates for EU AI Act readiness run from roughly €5,000–25,000 in year one for a small organisation with no high-risk systems, to €25,000–100,000 for a mid-sized organisation with them, to €100,000–500,000 or more for large enterprises and high-risk providers, with ongoing per-system costs in the tens of thousands. Against penalties reaching €35 million or 7% of global turnover, that arithmetic is not close — but it belongs in the business case, not in a separate governance budget that nobody nets off the return.

Working against both is the fastest-moving variable in the model: for a fixed capability tier, inference prices have fallen roughly tenfold a year, and comparable output that cost dollars per million tokens in 2023 costs cents in 2026. This does not mean AI bills fall. It means unit costs fall while volumes rise faster, which is the Jevons pattern described in Section 12.2. For a business case, the practical implication is to model the cost side as a band rather than a point, and to re-forecast annually rather than treating the year-one unit cost as fixed for three years.

GOVERNANCE The audit trail is measurement infrastructure

The evidence a governance programme already requires — documented purpose, defined performance criteria, evaluation results, monitoring records, incident logs (Section 13.7) — is most of what a credible ROI claim needs. Organisations that treat compliance and value measurement as two separate exercises pay twice and usually do both badly.

The practical move is to define the business metric and the conformity evidence at the same gate, from the same records. It makes the return auditable, which matters increasingly as disclosure expectations tighten, and it turns a cost centre into part of the case for the investment.

SECURITY Avoided loss is real, and the easiest number to inflate

Risk reduction is a legitimate return — fraud caught, incidents contained, exposure reduced. It is also the category where business cases stretch furthest, because the counterfactual is unobservable and the denominator can be chosen to suit.

Two disciplines keep it honest. Use your own historical loss data rather than published industry averages, which reflect a different threat profile and a different control baseline. And count avoided loss once — not in the security business case, the AI business case, and the insurance renewal simultaneously.

10.8 A worked example, end to end

Take the contract-summarisation case from Section 9.9 through to a financial answer. An in-house legal team of six reviews about 3,000 supplier contracts a year. First-pass review takes about 45 minutes each at a fully-loaded cost of roughly A\$140 an hour — about A\$315,000 a year of effort. With AI-drafted summaries and lawyer verification, first pass falls to about 18 minutes, or roughly A\$126,000 a year. Gross annual saving at full adoption: about A\$189,000.

That figure is the start of the analysis, not the end. Adoption ramps rather than arriving whole — assume 45%, 85%, and 95% across three years. Costs include an A\$78,000 build (golden set, integration, evaluation harness, project management), annual running costs of A\$31,000–36,000, governance and assurance, and an allowance for model migration in years two and three.

Year	Benefit realised	Cost	Net	Cumulative
Year 1 (45% adoption)	A\$85,050	A\$125,000	–A\$39,950	–A\$39,950
Year 2 (85% adoption)	A\$160,650	A\$54,000	A\$106,650	A\$66,700
Year 3 (95% adoption)	A\$179,550	A\$56,000	A\$123,550	A\$190,250
Three-year total	A\$425,250	A\$235,000	A\$190,250	—

Three-year ROI is about 81%. Payback falls around month 16 or 17. Net present value at an 8% discount rate is roughly A\$153,000. These are respectable numbers and they are deliberately unspectacular — they are what a well-run, well-scoped AI use case actually looks like, far below the 300–500% figures typical of vendor-commissioned studies, which model composite organisations chosen for having succeeded.

The conversion test. The most important line in the analysis is not in the table. This is a type 2 return: it saves lawyer hours, not lawyer salaries. It becomes money only if those hours are redeployed to work the organisation values — matters currently sent to external counsel, contract volume the team previously refused — or if an anticipated seventh hire is not made. If neither is decided and committed, the honest conclusion is that the programme returns A\$0 to the P&L, and a CFO who discounts it to zero is right to. Deciding that question is a business change, not a technology one, and it belongs in the business case at G1 with a named owner. This is precisely the link a Benefits Dependency Network is designed to make visible.

10.9 Nine ways an AI ROI number misleads

Whether you are writing a business case or reading one, these are the failure modes that account for most of the distance between claimed and realised value.

- **Self-reported time savings.** People believe AI helped more than it did, sometimes dramatically. Treat survey-derived time savings as sentiment data.
- **No baseline.** If the pre-state was never measured, the improvement is an assertion. This is the single most common defect.
- **Time counted as money.** Hours saved become dollars only through a decision about capacity. Without that decision the benefit is real to the worker and invisible to the accounts.
- **Selection effect.** The best use case is measured, and its result is extrapolated across a portfolio containing several that failed. Portfolio-level accounting includes the failures.
- **Vendor studies read as benchmarks.** Studies commissioned to support a product are useful as method and misleading as benchmark. Borrow the structure, discard the percentage.

- **Pilot cost mistaken for run cost.** Pilot economics omit monitoring, evaluation, migration, and human review — the costs that recur. Three-year cost typically exceeds the build quote by 40–60%.
- **The failure tail ignored.** Agent workflows retry, fail, and escalate. A cost model built on the happy path understates the bill and overstates the throughput.
- **Shadow AI on both sides of the ledger.** Around 67% of employees use AI tools at work while only about 18% of organisations have formal guidelines, and leaders underestimate heavy use by roughly threefold. Some of the value being claimed for a new platform was already being produced informally — and some of the value being produced is invisible to every measurement system in the building.
- **Survivorship in the evidence base.** Published case studies are written by the projects that worked. The base rate — most pilots do not reach production — is the context every case study omits.

10.10 The regional picture

Boards and regulators in the three regions this guide serves are pushing on value measurement from different directions, and a business case written for one reads oddly in another.

United States. The pressure is from the capital markets and it is toward hard financial attribution. With analysts and boards asking what the AI spend has returned, the metric that carries weight is EBIT or margin impact — which is why the McKinsey finding that only around 6% of organisations can point to more than 5% of EBIT is quoted so often. Expect scrutiny of the conversion step, not the productivity claim.

European Union. Compliance cost is a live, quantified line in the business case rather than an afterthought, and although the Digital Omnibus deferred the heaviest high-risk obligations to December 2027 (Section 13.3), the readiness work is being funded now. The upside is that the documentation burden produces exactly the evidence base a rigorous ROI analysis needs; the downside is that for smaller organisations the compliance line can exceed the build line for a first high-risk use case.

Australia. The Productivity Commission's estimate — roughly 0.4 percentage points a year, about 4% of labour productivity over the decade — is deliberately modest, and it comes with an explicit observation that Australia lags the United States, United Kingdom, and Canada on adoption. Sector estimates are larger in headline terms; a Deloitte and Amazon analysis put the small business opportunity alone at around A\$45 billion. For anyone building a business case here, the practical implication is that the adoption and change-management line should be funded more generously than an American template suggests, because slower diffusion is the measured local reality rather than a pessimistic assumption.

10.11 Is there a body of knowledge yet?

Honestly: an emerging literature, not a discipline. The distinction matters. A discipline has agreed definitions, a standard method, comparable published results, and a professional body that maintains them. AI value measurement in 2026 has none of these.

What it has is three tiers that do not talk to each other. At the rigorous end sit the academic randomised trials — causally sound, narrow in scope, and unable to produce an enterprise

number. At the practical end sit the consultancy and vendor frameworks — organisationally realistic, commercially motivated, proprietary, and not comparable between firms. In between sits the gap: no ISO or professional-body standard for AI benefit measurement exists. ISO/IEC 42001 governs the management system rather than the value it produces, and there is no AI equivalent of the benefits-management guidance the project profession maintains for conventional investments. Early academic work is attempting to close this — including proposals for risk-adjusted AI returns that integrate ISO 42001 controls and regulatory exposure directly into the calculation — but none of it is settled or widely adopted.

That gap will close, probably within a few years and probably by extension of existing benefits management practice rather than by anything invented for AI. Until it does, four moves are available and all of them are worth making. Borrow benefits realisation management for the structure. Borrow FinOps unit economics for the cost side. Borrow the RCT discipline — control groups, pre-registered criteria — at the task level where it is affordable. And keep your own longitudinal record, because your organisation's own history, measured consistently, is the only benchmark genuinely comparable to your next initiative.

10.12 What good looks like

A checklist for a business case that will survive contact with a finance function.

- State which of the six types you are claiming, before stating a number.
- Measure the baseline before you build, and write it down with dates and sample sizes.
- Name the business change that converts the benefit into money, and give it an owner.
- Identify a comparison group, even an imperfect one, and prefer staged rollout so you have one by default.
- Model adoption as a ramp, never as full from day one.
- Cost the three-year operating reality, not the build, and include governance and compliance.
- Discount the cash flows and report payback and NPV alongside the ROI percentage.
- Define in advance what result would mean stopping, and let stopping be a normal outcome.
- Report the portfolio, including the initiatives that failed, not only the one that worked.
- Re-forecast annually — capability, unit cost, and what is feasible all move (Sections 12 and 16).

None of this is sophisticated. That is the point: as Section 9.2 argued about delivery, the failure mode is almost never a lack of analytical sophistication. It is skipping the boring steps — and in measurement, the boring step is the baseline nobody took.

Go deeper: the delivery, evaluation and measurement resources behind this section — including the benefits management material and the productivity research it cites — are in Sections 17.5–17.6.

Next section: Section 11 — People & Work: what the evidence actually shows about AI and employment, the junior pipeline, role redesign, the obligations that attach before you change how people work, and algorithmic management.

11. People & Work

The frameworks and obligations are durable; the labour-market evidence is early and moving (verified 1 August 2026). This is the most contested subject in the guide — the figures below are reported with their methodology attached, because in this area the methodology is the argument.

Section 9 argued that AI transformation is mostly organisational, and Section 10 concluded that a productivity gain only becomes money through a decision about people. This section is that decision, examined properly: what the evidence actually shows about AI and employment, what it does to the junior pipeline, what the law requires you to do before you change how people work, and what distinguishes organisations that handle this well from those that generate headlines.

It is also the section where hype and counter-hype are both loudest. The discipline applied here is to separate three things that are routinely blurred: what AI could technically do, what organisations say they will do, and what has actually been measured.

11.1 Exposure is not displacement

Almost every alarming headline number is an exposure estimate. Exposure means an occupation's tasks overlap with what AI can do. Displacement means jobs actually went away. The two are related the way flammability is related to fire, and conflating them is the single largest source of distortion in this debate.

The International Labour Organization's index puts roughly a quarter of workers globally in occupations with some generative-AI exposure, but only about 3.3% of global employment in the highest-exposure category. The distribution is uneven in ways worth noting: exposure is far higher in high-income countries (around 34% of employment, against about 11% in low-income countries), and it is markedly gendered — in high-income countries around 9.6% of female employment sits in the highest-exposure band against 3.5% of male employment, because clerical and administrative work is both the most exposed category and disproportionately done by women.

The OECD's 2026 exposure measure makes the distinction explicit in its own findings: the set of jobs most exposed to AI is not the same as the set most at risk of automation. Exposure describes where the technology touches; substitutability describes whether it can take over; and organisational choice determines whether it does. Only the third is under anyone's control, which is why this section spends more time on it than on forecasts.

11.2 What the evidence actually shows

Four bodies of evidence carry most of the weight, and they differ enormously in rigour. Reading them together is more informative than any one of them.

Source	What it measures and how	What it found
Stanford Digital Economy Lab	Payroll microdata on 4.6 million US workers across 730+ occupations, tracked continuously. Descriptive but unusually well-identified.	Workers aged 22–25 in the most exposed occupations — software, customer service, marketing — down about 16% in relative employment through late 2025. Workers 30+ in the same occupations grew 6–12%. No significant wage decline; the adjustment is in hiring, not pay.
Federal Reserve (Atlanta)	Survey of around 750 corporate executives, March 2026. Reports	Aggregate US employment effect from AI of less than a 0.4% decline in 2026; output per

	expectations, not outcomes.	worker about 3% higher than counterfactual; routine clerical employment share falling around 2.2 points by 2028, partly offset by skilled technical roles.
ILO and OECD indices	Task-level exposure modelling across occupations and countries. Describes potential, not outcomes.	About a quarter of global workers in occupations with some exposure; 3.3% in the highest band; exposure concentrated in high-income countries and in clerical work.
Challenger, Gray & Christmas	Tracks reasons employers state in announced US layoffs. Measures what companies say, which is not the same as what is true.	Around 55,000 cuts cited AI in 2025 (about 5% of announced cuts); roughly 88,000 in the first five months of 2026 (about 22%), with AI the most-cited single reason for several consecutive months.
Acemoglu (NBER)	Macroeconomic modelling of task-level automation potential. The leading sceptical estimate.	Around 5% of tasks profitably automatable within a decade, implying roughly 1–2% GDP effect over ten years — an order of magnitude below the most-quoted investment-bank forecasts.

Read together, these say something fairly clear. There is no economy-wide employment collapse in the data. There is one robust, specific, well-identified displacement finding — concentrated in the youngest workers in a handful of exposed occupations — and a much larger volume of stated intention. Australia is a useful counterpoint: the Australian government's own analysis, published in July 2026 using labour force and vacancy data, found no evidence to date of broad AI-driven upheaval in the domestic labour market, with youth outcomes largely holding up. Whether that reflects a genuinely different economy or simply a lag is not yet answerable.

11.3 The junior pipeline problem

The Stanford finding matters more than its size suggests, because of what it implies about how organisations train people. The work AI does best in professional services is precisely the work juniors were given in order to learn: first-pass document review, research summaries, first drafts, tier-one support. That work was never economically efficient. It was an apprenticeship disguised as production.

The legal sector shows the pattern most clearly. Entry-level hiring into the largest US firms has been flat to declining across four years in which those same firms added tens of billions in combined revenue — growth decoupled from junior headcount for the first time. In February 2026 one global firm cut several hundred business-services roles explicitly citing AI integration.

The forward-looking concern follows logically but is not yet demonstrated: if juniors are not hired and not given the work that builds judgement, the supply of people capable of becoming seniors contracts on a five-to-ten-year lag, by which point the decision is irreversible. Treat this as a serious risk to reason about rather than an established finding — and note that it is a collective-action problem. No individual firm is wrong to stop paying for work a machine does more cheaply. Every firm doing so simultaneously produces a shortage none of them chose.

DELIVERY / PM If AI does the training work, the training has to move

The practical response is not to preserve make-work. It is to recognise that apprenticeship was bundled into tasks that no longer need doing, and to unbundle it deliberately: structured review of

AI output as a teaching exercise, deliberate exposure to the judgement calls the model cannot make, and explicit time for the reasoning that used to be learned by grinding through volume. This is a real cost that no business case will surface, because it does not appear as a line item until the capability gap shows up years later. Name it during the use-case scorecard (Section 9.9), where “what does this remove from the learning pipeline?” belongs alongside error tolerance and containment.

11.4 Augmentation, substitution, and the jagged frontier

The most robust finding in the applied literature is skill levelling: on well-structured tasks, AI helps the least experienced most. The landmark study of 5,179 customer support agents found average productivity up about 14%, but with a 34% improvement for novices and a negligible effect on the most experienced. The mechanism appears to be that the model disseminates the working patterns of the best performers to everyone else — agents with two months’ experience performed at the level previously reached at six.

The consulting study behind the jagged frontier (Section 13.4) found the same shape — below-average performers improving 43% against 17% for above-average — but added the crucial caveat: on tasks outside the technology’s competence, consultants using AI were 19% less likely to reach the correct answer than a control group with no AI at all. Assistance is not uniformly positive; it is positive inside the frontier and harmful outside it, and the frontier is not visible from the inside.

Nor is skill levelling universal. A 2026 field experiment in customer service found that while low performers gained, top performers saw little speed benefit and a measurable decline in service quality — direct evidence that assistance can degrade expert work rather than merely failing to help. The unifying variable appears to be task complexity: AI narrows the novice–expert gap on routine, well-structured work and widens it on work requiring deep judgement. Deployed without that distinction, a single tool rolled out uniformly will help one part of a workforce and quietly hurt another.

Anthropic’s analysis of actual usage rather than survey response provides a useful counterweight to the automation narrative: through late 2025 the split ran at roughly 52% augmentative to 45% automative, having briefly tipped the other way earlier in the year. The direction moves with product features as much as with model capability, which is itself the point — how a tool is designed shapes whether it replaces a person or works alongside one.

11.5 How to read a layoff announcement

The gap between announced AI-driven job cuts and demonstrable AI-driven job cuts is wide, and the incentive runs in one direction. Attributing redundancies to AI tells investors the company is modernising; attributing them to over-hiring tells investors management misjudged. As one MIT economist has put it, whether or not AI is the reason, a company would be wise to give AI the credit.

The sceptical case is made by people well placed to make it. Oxford Economics finds firms are not replacing workers with AI at any significant scale economy-wide and suspects layoffs are being dressed up as good news. Research at Wharton has found very little evidence that AI cuts headcount in most cases at all. Against that, the Challenger data is real and rising, and some claims are specific enough to be credible — IBM’s chief executive naming a few

hundred HR roles taken over by chatbots is a far more useful data point than a ten-thousand-person restructure with AI mentioned in the third paragraph.

The most instructive case is the reversal. A payments company cut around 700 customer service roles in 2024 with heavy publicity, became the global case study for AI substitution, and then in 2025 its chief executive publicly conceded the company had focused too much on cost, that quality had suffered, and began rehiring for a hybrid model. Both halves of that story are true and both are worth remembering: the substitution was real, and so were its limits.

When reading any announcement, ask three questions. Is the AI attribution specific to named roles and tasks, or general? Was the workforce growing unsustainably beforehand? And is the company describing something it has done or something it intends? Most announcements answer the third question in a way that should lower your confidence considerably.

11.6 Role redesign and the jobs that are appearing

The largest analysis of job advertisements — over a billion postings across 27 countries — describes a labour market splitting into two tracks rather than shrinking. In “professionalised” roles, AI absorbs routine tasks and elevates the judgement component; those roles are growing about twice as fast, with salary growth around 42% faster. In “democratised” roles, AI lowers the expertise required to do the work, which broadens access and compresses the premium. Both are real, and which track a given role lands on is substantially a design decision rather than a technological inevitability.

The Australian figures in the same analysis are striking: AI-related job postings roughly doubled from 20,000 in 2024 to 41,000 in 2025, the wage premium for AI skills rose to around 62%, and demand for AI and machine learning skills was up 245% since 2023 — in a labour market where the same government analysis found no broad displacement. Both things are happening at once.

New roles have emerged fast enough to be visible in postings that did not exist two years ago: agentic AI engineers who build and maintain tool-using systems; evaluation engineers who own golden sets, judges and regression suites (Section 9.10); orchestration specialists, currently one of the tightest talent bottlenecks anywhere; AI governance leads driven by the regulatory timeline in Section 13; decision engineers who design where autonomy ends and human review begins; and enablement leads who run adoption. Most of these are recognisable recombinations of existing skills rather than exotic new disciplines, which is good news for internal mobility.

11.7 Skills that hold their value

Employer surveys consistently forecast that a large share of core skills will change within five years, and consistently name the same durable capabilities: judgement, complex communication, and the ability to work across domains. Treat these as expectations rather than measurements — they are what employers say they will need, and employers have been wrong before.

The more defensible framing comes from the evidence in this section rather than from surveys. If AI levels skill on structured tasks, the value of being able to do those tasks falls and the value of knowing which task to do rises. If assistance is harmful outside the frontier,

the ability to recognise where the frontier lies becomes a distinct and valuable skill. If output is cheap and verification is expensive, evaluation becomes the bottleneck — and the ability to tell good work from plausible-but-wrong work is the capability that compounds. Section 17 develops this into a curriculum; the short version is that the durable skills are the ones AI makes more necessary rather than the ones it cannot yet do.

One practical observation from the adoption data. Employees are typically ahead of their employers: usage of personal AI tools for work is near-universal even in organisations whose official programmes have stalled, and leaders consistently underestimate how much their staff use these tools — in one study by a factor of three. The reskilling problem is therefore rarely one of willingness. It is one of direction, permission, and the absence of a sanctioned alternative, which is why bringing shadow use into the light (Section 9.6) is usually a faster route to capability than a training programme.

11.8 Consultation and legal obligations

Changing how people work triggers obligations that exist independently of AI, and organisations regularly discover them late. None of the three jurisdictions this guide covers has comprehensive AI-specific employment legislation; all three reach AI through existing law.

European Union. Employers deploying high-risk AI systems in the workplace must inform workers' representatives and affected workers before deployment under Article 26 of the AI Act — a right to information, which is narrower than a right to consultation and should not be confused with one. National law is often stronger: German works councils must be informed well in advance and consulted on effects on the nature of work, and French employers above fifty employees must inform and consult the CSE before introducing new technologies, with case law treating meaningful employee impact as triggering the duty.

Australia. There is no standalone AI Act and no AI-specific consultation duty, but the existing framework bites. Consultation is a mandatory term in modern awards and most enterprise agreements for major workplace change likely to have significant effects on employees — which an AI deployment that alters duties, rosters or headcount plainly is. The Fair Work Commission issued revised model consultation terms in early 2025. Unions have begun bargaining AI provisions directly: an enterprise agreement approved in December 2025 restricts AI use in editorial work and requires workforce consultation before an AI governance framework is implemented. Separately, New South Wales legislated in February 2026 to bring “digital work systems” — expressly including AI and algorithmic work allocation — within the primary work health and safety duty, the first Australian law to name algorithmic management as a regulated hazard. And from December 2026 the automated decision-making disclosure obligation in Section 8.5 applies. The government's position remains that technology-neutral law already covers this: if an algorithm effectively makes a dismissal or recruitment decision, unfair dismissal and anti-discrimination law apply exactly as they would to a human decision-maker.

United States. Federal employment law imposes no AI-specific consultation duty, but state activity is growing and existing anti-discrimination law applies to algorithmic hiring and management decisions — an area where the evidentiary burden on an employer relying solely on an automated tool is uncomfortable.

GOVERNANCE Consultation is cheaper before the decision than after it

The obligation in most jurisdictions is to consult while the decision can still be influenced — not to announce a settled outcome. Organisations that treat consultation as a communications exercise after the business case is approved satisfy neither the letter nor the purpose, and in Australia expose themselves to a dispute under the consultation term of the applicable award or agreement.

The delivery consequence is a sequencing one: workforce consultation belongs before the production gate, not after it. Fold it into the gate criteria in Section 9.9 alongside the technical and financial tests, and record it — the same evidence discharges the obligation and supports the change-management case.

11.9 Monitoring and algorithmic management

The same technology that assists work can supervise it, and the boundary is thinner than most policies acknowledge. Systems that allocate tasks, score performance, or flag anomalies in employee behaviour are algorithmic management whether or not anyone called them that when they were procured.

The public-opinion evidence here is thinner and older than the subject deserves — the most rigorous available survey work is now several years old, and much of the confident-sounding 2026 commentary rests on vendor surveys of poor quality. What that older work established is a clear gradient: opposition rises with invasiveness, from a substantial minority objecting to AI evaluating performance, to a majority objecting to monitoring of computer activity, to a larger majority objecting to tracking physical movement. The qualitative finding that recurs is that acceptance depends on transparency, proportionality, and a genuine connection to safety or security rather than to productivity.

Two practical cautions. Heavy surveillance reliably shifts behaviour from discretionary effort toward minimum compliance, which is the opposite of what the tool was bought to achieve. And in the EU, certain workplace applications — emotion recognition at work among them — are prohibited or high-risk under the AI Act rather than merely sensitive, so this is a compliance question before it is a culture question.

11.10 Organisational design

The structural change already visible has less to do with AI directly than with what AI is used to justify. Average spans of control have widened substantially over the past decade — from roughly eight direct reports per manager in 2013 to around twelve by 2025 — and manager headcount at large US public companies fell about 6% between 2022 and 2025. A significant share of employees report their organisation has removed management layers.

The efficiency case is straightforward and the cost is under-reported: a substantial minority of employees in flattened organisations report feeling directionless, and senior leaders frequently doubt their capacity to manage the expanded scope. Formal modelling suggests the trajectory may not be linear at all — spans contracting during adoption, when coordination costs rise, before expanding as capability matures. Flattening is a bet, not a consequence.

The most-quoted figure in this area — that AI transformation is 10% algorithms, 20% data and technology, and 70% people, process and culture, cited in Section 9.2 — deserves an honest caveat. It is a consultancy heuristic rather than a measured finding, and no

underlying study supports the precise split. It survives because it is directionally consistent with better evidence: the recurring finding that workflow redesign correlates more strongly with realised value than technology choice does. Use it as a slogan for a true idea, not as a statistic.

11.11 Doing this well

What separates organisations that handle this credibly is not sophistication. It is sequencing and candour.

- Decide what happens to freed capacity before the tool ships, and say so. The ambiguity is what generates fear, and fear is what generates quiet non-adoption (Section 10.8).
- Consult while the decision can still change, and document it.
- Separate the augmentation case from the substitution case in your own planning. Both are legitimate; conflating them in communications destroys trust that is expensive to rebuild.
- Design for the frontier: identify which tasks the tool genuinely helps with, and warn people about the ones where it will confidently mislead them.
- Protect the learning pipeline deliberately, because nothing else in the system will.
- Keep monitoring proportionate, transparent, and tied to something other than productivity.
- Measure adoption and quality by cohort, not in aggregate — the averages conceal the experts being slowed down and the novices being levelled up.

And retain the humility the evidence demands. The most confident claims in this field, in both directions, are the least supported. What is actually established is narrow: real gains on structured tasks, real harm outside the frontier, a measurable effect on the youngest workers in the most exposed occupations, and a great deal of stated intention that has not yet happened.

Go deeper: the labour-market research, the Australian workplace-relations sources, and the change-management material behind this section are in Sections 17.5–17.6.

Next section: Section 12 — The Economics of AI: what AI actually costs and who makes money — tokens and pricing, the capex boom, the business models, and whether the economics are sustainable. From the people the money pays to the industry the money flows to.

12. The Economics of AI

⚡ *FAST-MOVING* — all figures are indicative and dated (verified 29 July 2026). Capex guidance, lab revenues, valuations, and token prices move monthly; treat specific numbers as a snapshot, and re-check anything load-bearing.

Behind every chat reply sits a supply chain that now absorbs some of the largest capital flows in corporate history. This section follows the money — from the token you pay for, up through the labs, the clouds, the chips, and the power stations — and asks the question hanging over all of it: do the economics actually work? The answer matters whether you're budgeting a project (Section 9), assessing vendor risk, or just trying to understand why the AI industry both feels unstoppable and looks precarious at the same time.

12.1 The unit of everything: the token

As Section 1 introduced, models process text as tokens, and tokens are also the unit of pricing. APIs charge per million tokens, and a detail that wrecks budgets: output tokens typically cost several times more than input tokens (often 3–10×), and most real workloads are output-heavy. The good news is that the cost of raw intelligence has collapsed. Price for a given capability has fallen roughly 98% since the GPT-3 era — what cost around \$20–30 per million tokens in 2022–23 is often well under \$1 today — and analysts track order-of-magnitude declines per year for each capability milestone, with more expected as new chips (each hardware generation targets large inference-cost reductions) come online. Intelligence, per unit, is deflating faster than almost any commodity in history.

12.2 The paradox: cheaper tokens, bigger bills

Here is where intuition fails. Even as the price per token collapses, total AI spending is rising — average enterprise AI budgets went from roughly \$1.2M in 2024 to around \$7M in 2026. This is Jevons' paradox, the 1860s observation that more efficient coal engines made Britain burn more coal, not less: when something useful gets cheaper, people find far more uses for it. Three forces compound the effect:

- **Agentic workloads are token-hungry.** An agent that loops, re-sends context each step, loads tool schemas, retries, and runs unattended can consume on the order of 1,000× the tokens of a simple chat for an equivalent task.
- **Most of the cost is off the model invoice.** A large majority of production AI cost — by some estimates around 70% — sits in orchestration, retrieval, retries, and observability around the model, not in the token price itself.
- **Adoption multiplies.** Every team that ran one pilot in 2024 runs ten in 2026; one large telco went from about 1 billion to 27 billion tokens a day in eighteen months.

The one-line summary: the cost of intelligence is falling while the cost of deploying intelligence is climbing. That gap is the central economic tension of 2026, and it is why the subsidy era — where labs priced below cost to win adoption — is visibly ending, with real customers hitting budget ceilings (one large company reportedly exhausted its annual AI-coding budget in four months as usage spread).

12.3 The capex boom

The supply side is a building spree without modern precedent. The four largest US hyperscalers guided to roughly \$725 billion of combined capital spending in 2026 — up about 77% on 2025, which had itself nearly doubled 2024 — with Amazon around \$200B, Microsoft around \$190B, Google \$180–190B (raised at its 22 July earnings), and Meta \$125–145B, plus Oracle and others on top. Most of it buys AI data centres, Nvidia GPUs (tens of thousands of dollars each), custom silicon, and power. Nvidia remains the prime beneficiary, with data-centre revenue running around \$75 billion a quarter; the hyperscalers are simultaneously building their own chips (Google's TPUs, Amazon's Trainium, Microsoft's Maia, Meta's MTIA) to reduce that dependence and cut per-token cost. Headline projects convey the scale: Stargate has grown to roughly 7 GW and US\$400B+ of planned capacity after a July Oracle deal added up to 4.5 GW, and OpenAI's total infrastructure plan was raised in late July to some US\$750 billion through 2030 — with multi-trillion-dollar multi-year forecasts from investment banks on top. The strain is starting to show: several of these firms may see free cash flow compress sharply in 2026 as capex outruns AI revenue, and their shares have wobbled each time guidance rises.

12.4 Energy: the binding constraint

Increasingly the bottleneck is not chips or capital but electricity. Securing power, land, cooling, and grid connection has become a bigger competitive moat than algorithms for much of the industry. To put it in proportion, a handful of tech firms are now outspending the entire US electric-utility industry on energy-adjacent infrastructure by a wide margin, and data-centre demand is measured in tens of gigawatts. This has real-world consequences — grid strain, water use, local-community friction, and lengthening delivery timelines — and it is a live issue in all three regions: multi-year grid-interconnection queues in the US, permitting and power prices in the EU, and data-centre siting, grid capacity, and energy mix in Australia all shape where and whether capacity gets built. Efficiency helps (leaner model architectures have shown large training-energy reductions), but demand is still outrunning it.

There is a narrower version of this that lands on your own desk: your organisation's AI footprint is becoming a reporting and procurement question, not just an industry one. Inference has a measurable energy and water cost, disclosure expectations are rising under ESG and sustainability regimes, and buyers increasingly ask suppliers to account for it. The practical levers are the same ones that cut cost — routing to smaller models, caching and batching (Section 2.3) reduce tokens, and fewer tokens mean less energy — so efficiency and footprint largely move together.

12.5 Who actually makes money

Value is captured unevenly across the stack. The old mining-rush wisdom — sell picks and shovels — describes 2026 well: the clearest profits are in hardware and infrastructure, while the model labs, for all their revenue growth, are mostly still spending more than they earn.

Layer	Examples	Economics in 2026
Chips & hardware	Nvidia, AMD, Broadcom, TSMC	The clearest profits today; highest margins
Cloud & data centres	AWS, Azure, Google Cloud, Oracle, Equinix	Capex-heavy; renting compute; vast backlogs
Power & land	Utilities, grid, cooling, energy	The binding constraint; long lead times

Layer	Examples	Economics in 2026
Model labs	OpenAI, Anthropic, Google, Meta, xAI	Explosive revenue, mostly still loss-making
Applications	Cursor, Perplexity, vertical & in-house tools	Value accrues where lock-in + proprietary data

The two leading labs illustrate a real divergence in business model. Anthropic (the maker of Claude, with which this guide is built) grew from around \$1B to roughly \$30B in annualised revenue in about fifteen months — by some reports overtaking OpenAI's ~\$25B in the process — is roughly 85% enterprise/developer, was projecting its first operating profit in 2026 with healthy gross margins, and closed a US\$65 billion Series H in May 2026 at a valuation near US\$1 trillion. OpenAI, larger in raw revenue and with hundreds of millions of weekly consumer users, is roughly the inverse — mostly consumer, the vast majority on free tiers — and was projecting billions in losses (one 2026 estimate near \$14B for the year) with breakeven pushed toward the end of the decade. The structural reason matters for anyone forecasting AI costs: enterprise customers generate several times more revenue per token than consumers, with more predictable usage and stickier contracts. Consumer scale drives inference cost without proportionate revenue; enterprise depth is where profitability currently lives.

12.6 Is it sustainable? The bubble debate

This is genuinely contested, and an honest guide should give both cases rather than a verdict. (What follows is descriptive, not investment advice.)

The bear case centres on circular financing. A small set of firms — the chipmaker, the clouds, and the labs — act at once as each other's suppliers, customers, investors, and validators. Labs commit tens of billions in future cloud spend to the same hyperscalers that invest in them, sometimes via compute credits usable nowhere else; cloud providers book large paper gains by marking up their stakes in the labs (one hyperscaler's recent record quarterly profit was substantially a markup on its AI-startup stake, even as its free cash flow collapsed under capex). Critics note the demand signal can become circular — each node paying another — and draw uneasy parallels to the 2001 telecom bust, where firms swapped network capacity to manufacture the appearance of revenue. If any node slows, the argument goes, all of them slow.

The bull case answers that vendor financing is as old as the tech industry (the 1990s telecom build-out did the same) and that the underlying demand is real: revenue has grown tens of times over in under two years, enterprise adoption is compounding, and per-unit efficiency is improving faster than almost anyone forecast. On this view the capex is a rational land-grab for a demand curve everyone can see coming, and being short of compute is more dangerous than overspending. The reckoning that will settle much of this is disclosure: as the major labs move toward public listings, their filings will itemise the related-party deals and off-balance-sheet compute commitments that private rounds have kept opaque. Until then, 2026 is best described as AI's "prove-it" phase.

12.7 What it means for you

For anyone deploying rather than investing, the economics translate into a concrete discipline. Because cost is volatile, usage-driven, and mostly hidden outside the token price, cost management is now a first-class part of running AI — the practice the industry calls FinOps for AI.

DELIVERY / PM FinOps for AI: manage the bill, not just the model

Route by task, don't default to the frontier. The most expensive mistake is using a top-tier model for everything (the "big-model fallacy"); a router that sends easy work to cheap, fast models and hard work to frontier ones is the single biggest lever. Cache aggressively (prompt/context caching cuts repeat input costs), set token budgets and alerts per team and per workflow (one runaway retry loop can burn a month's budget overnight), and budget for output tokens and the ~70% of cost outside the model invoice, not the advertised input price. Treat provider pricing as a moving target — billing models change (flat-rate to usage-based shifts have blindsided teams) — and fold all of this into the TCO discipline from Section 9.

GOVERNANCE Concentration is a governance question

The same circularity that worries investors is a procurement and resilience risk for buyers: a handful of firms supply the chips, the clouds, and the models, so vendor lock-in, price power, and single-provider dependency are real exposures. Prefer architectures that keep switching possible (portable prompts, abstraction layers, avoiding deep proprietary hooks unless the value is worth the lock-in), watch the sustainability and energy disclosures that regulators are beginning to require, and treat provider concentration as something to manage deliberately rather than inherit by default. Section 13 develops the governance picture and Section 14 the resilience/security angle.

Go deeper: the newsletters that track AI economics credibly — Import AI, Latent Space, The Batch — are in Section 17.9.

Next section: Section 13 — Governance, Risk & Compliance: the rules of the road — the EU AI Act, the US federal and state picture, the Australian picture, ISO 42001 and the NIST AI RMF — and how to build governance that enables rather than blocks. The governance anchor section.

13. Governance, Risk & Compliance

Regulation is fast-moving; dates and statuses verified 29 July 2026 and flagged as provisional where they are. This is the governance anchor section, and it treats the three regions this guide serves — the United States, the European Union, and Australia — with equal weight. It is general information, not legal advice — confirm specifics with counsel.

Every earlier section has pointed here. Governance is where the capabilities, risks, and delivery discipline of the whole guide come together into a single question: how does an organisation use AI in a way that is responsible, defensible, and legal? This section separates three things people often blur — governance, regulation, and risk management — maps the rules that now apply across all three regions, and sets out a practical way to build governance that enables adoption rather than blocking it.

13.1 Governance, regulation, risk — three different things

The terms get used interchangeably; keeping them distinct clarifies everything that follows. Governance is your internal system — the policies, roles, controls, and lines of accountability that decide how your organisation builds, buys, and uses AI. Regulation is external law imposed by governments, with penalties. Risk management is the process — identifying, assessing, and mitigating what could go wrong. You need all three: regulation sets the floor, risk management is the method, and governance is how you actually run it day to day.

Why this matters now, more than a year ago: binding law is arriving on a real timetable; agents act in the world and hold privileged access (Section 6), so the stakes of an ungoverned system are higher; shadow AI (Section 9) means unmanaged use is already happening inside most organisations; and the reputational and legal exposure of getting it wrong is no longer hypothetical. The framing that runs through this section: good governance is enabling infrastructure, not red tape — the thing that lets you scale AI with confidence rather than the thing that stops you.

13.2 The regulatory map: three postures

Globally, 2026 splits into three broad stances, and a single AI system can fall under several at once — the EU because it touches the EU market, a US state because of where users live, Australia because that's where you operate. The practical implication is that compliance is no longer a checklist against one rulebook but the ability to satisfy several with one well-governed program.

Jurisdiction	Approach	Where it stands (2026)
European Union	Binding, risk-tiered law (AI Act)	High-risk delayed to Dec 2027; transparency, GPAI & enforcement live Aug 2026
US — federal	Light-touch; pushing to preempt states	No federal AI statute; preemption contested and unresolved
US — states	Patchwork of binding laws	Colorado, Texas, California and more live or near; many cite NIST/ISO
Australia	Existing law; mandatory standards announced	Mandatory 'Australian Standards for AI' + Office of AI announced 15 Jul 2026; legislation to come
Global standards	Voluntary frameworks	NIST AI RMF, ISO 42001 — the cross-border common language

A scope note before the detail: this guide treats the United States, the European Union and Australia in depth. Readers elsewhere — the United Kingdom, Canada, and much of APAC

— will find the EU AI Act and the NIST AI RMF are the two anchors most other regimes borrow from or align to, so building to those and then checking local specifics with counsel is the durable move. The UK in particular has taken a lighter, regulator-led path closer to Australia's than to the EU's.

13.3 The EU AI Act: the one with teeth

Even from Australia, the [EU AI Act](#) matters, because its extraterritorial reach and first-mover status make it the de facto global baseline — the “Brussels effect.” It is risk-tiered: a few uses are prohibited outright (e.g. social scoring, most real-time public biometric surveillance, and — newly added — tools that generate non-consensual intimate imagery or CSAM); high-risk uses (hiring, credit, biometrics, education, essential services, and AI embedded in regulated products) carry heavy obligations; limited-risk uses owe transparency; and everything else is largely unregulated.

The 2026 headline is the Digital Omnibus, which pushed the heaviest high-risk obligations back — standalone (Annex III) high-risk to December 2027, product-embedded (Annex I) to August 2028 — because the technical standards weren't ready. But this is sequencing, not repeal, and crucially several things did not move: the obligations on general-purpose model providers have applied to new models since August 2025, while the Article 50 transparency duties (telling people they're dealing with AI, and labelling AI-generated content) and the Commission's enforcement and fining powers take effect from 2 August 2026 — this week. Penalties are serious — up to €35M or 7% of global turnover for prohibited-practice breaches, and up to €15M or 3% for most others. The distinction that decides your obligations is provider (you make or substantially modify a model) versus deployer (you use one); building an app on someone else's API usually makes you a deployer, but heavy fine-tuning can tip you into provider territory.

GOVERNANCE The EU bottom line

If you touch the EU market at all, assume the Act reaches you — that is what extraterritorial scope means. Classify your uses against the risk tiers now, meet the Article 50 transparency duties immediately, and diarise December 2027 (Annex III high-risk) and August 2028 (embedded products). Build to the harmonised standards as they land, use the FLI explorer (Section 17.6) rather than reading the Act raw, and remember the provider-versus-deployer line: heavy fine-tuning can quietly move you into the stricter category.

13.4 The voluntary frameworks: the “how”

Laws say what you must achieve; frameworks tell you how to organise for it, and they have become the common currency that procurement teams and insurers now expect regardless of jurisdiction. Two matter most:

- **NIST AI RMF** (US, voluntary) — a risk-management method built on four functions: Govern, Map, Measure, Manage, with a dedicated generative-AI profile. The strongest single anchor for a US-facing program.
- **ISO/IEC 42001** — the international standard for an AI management system, and crucially it is certifiable (the way ISO 27001 is for security). Certification is increasingly demanded in enterprise procurement as proof of governance maturity.

Supporting these are ISO/IEC 23894 (AI risk guidance), the OECD AI Principles, and the OWASP LLM Top 10 for security (Section 14). The reason to adopt one seriously is

leverage: because the frameworks map onto each other and onto the law — EU conformity, several US state laws, and sector rules all point back to them — you can implement once and satisfy many. That “implement once, map to many” property is the single most efficient move in AI compliance.

13.5 The United States picture

There is no federal AI statute, and the federal posture is deliberately light-touch and pro-acceleration: the voluntary NIST AI RMF (Section 13.4) is the closest thing to a national framework, and since December 2025 Executive Order 14365 has actively pushed back on state regulation — an AI Litigation Task Force to challenge state laws in court, and discretionary federal funding used as leverage against states that regulate. An executive order cannot itself pre-empt otherwise lawful state statutes, though, and the order carves out state authority over child safety, data-centre infrastructure, and government procurement; congressional attempts at a statutory answer, including a proposed multi-year moratorium on state AI laws, have so far failed to pass.

The action is therefore in the states, and it is a genuine patchwork of binding law. California’s frontier-model transparency act (SB 53) and Texas’s TRAIGA took effect in January 2026; Illinois has AI disclosure rules; and Colorado — the bellwether — passed the first comprehensive state AI act in 2024, delayed it twice, then replaced it with a narrower, transparency-focused law (SB 26-189) due 1 January 2027, with enforcement of the original suspended by a federal court in the meantime. The practical lesson: track the states where your users live, not just Washington — and expect the map to change quarterly.

GOVERNANCE The US bottom line

With no single rulebook, the durable move is the same as everywhere: anchor on the NIST AI RMF — the de facto national framework that federal policy and several state laws borrow from — inventory which state laws touch your users, and watch the pre-emption fight without waiting for it. Courts may take years to resolve it; the state obligations apply today. One well-governed programme mapped to the RMF absorbs the patchwork as configuration, not crisis.

13.6 The Australian picture

Australia has deliberately taken a lighter path than the EU, and it shifted again recently. There is no standalone AI Act, and in December 2025 the National AI Plan confirmed the government will not (for now) proceed with the mandatory high-risk guardrails it had proposed in 2024 — relying instead on existing, technology-neutral laws, sector regulators, and voluntary guidance, backed by a new Australian AI Safety Institute that is advisory and has no enforcement power. Critics call this too light; the government frames it as pro-innovation. Either way, “no AI Act” does not mean “no rules.”

That calculus shifted on 15 July 2026: the Prime Minister announced mandatory national “Australian Standards for AI”, a new Office of AI inside the Department of the Prime Minister and Cabinet to coordinate policy, and legislation to follow next year — alongside obligations on large AI data centres (underwriting new power generation, including renewables, and paying for their own grid and water connections) and promised copyright protections for Australian creators. Detail is thin until the bill lands, but the direction has plainly turned from voluntary back toward mandatory — the December 2025 pullback lasted barely seven months.

The voluntary layer has evolved quickly: the AI Ethics Principles (2019) and the Voluntary AI Safety Standard (2024) were largely superseded in October 2025 by the National AI Centre's [Guidance for AI Adoption](#), which is more practical — shipping an AI screening tool, an AI register template, and policy templates. The public sector faces binding rules of its own via the Digital Transformation Agency. And real obligations are arriving through existing law rather than an AI statute: from 10 December 2026, Privacy Act amendments require organisations to disclose in their privacy policies where substantially automated decisions significantly affect people; Parliament rejected a text-and-data-mining copyright exemption for AI training in early 2026 (a licensing model is being explored); a 2026 unfair-trading bill targets AI-driven “dark patterns”; and non-consensual sexual deepfakes are already criminalised. Underneath it all sit the laws that always applied — the Privacy Act, the Australian Consumer Law, anti-discrimination law, the Security of Critical Infrastructure Act — enforced by regulators like the [OAIC](#), APRA, ASIC, and the TGA.

GOVERNANCE The Australian bottom line

Voluntary today is not the same as optional forever, and the durable move is to align now. Adopt the National AI Centre's Guidance for AI Adoption (its register and screening tools are a fast start) and map to an international standard — ISO/IEC 42001 or the NIST AI RMF — so you're ready as the newly announced mandatory standards take shape and credible to enterprise customers today. Put the 10 December 2026 Privacy Act automated-decision transparency obligations on the calendar now, use the government's AI-and-cyber model contract clauses when procuring AI, and remember that anti-discrimination, consumer, and privacy law already bite regardless of whether an “AI Act” ever arrives.

13.7 Building governance that works

Regardless of jurisdiction, the practical program looks similar. The components that matter:

- **An AI inventory.** You cannot govern what you can't see — discover every AI system and agent in use, sanctioned and shadow (Section 9), and keep it current as new tools ship.
- **Risk-tiered use cases.** Borrow the EU's tiering even if you're not EU-bound: classify each use by potential harm and apply proportionate controls, so low-risk uses stay light and high-risk ones get scrutiny.
- **Clear policy and human oversight.** Acceptable-use rules, data-handling limits, disclosure expectations, and a human in the loop on consequential decisions.
- **Documentation.** Impact assessments for higher-risk uses, and model/system records — the evidence a regulator, auditor, or customer will ask for.
- **Named accountability.** Someone owns AI risk (often an AI governance committee plus an accountable executive); agents get identities tied to a human (Section 6).
- **Procurement due diligence, monitoring, and training.** Vet vendors and their data terms (model contract clauses help), monitor systems in production with an audit trail and incident response, and train the people who use them.

DELIVERY / PM Operationalise governance — implement once, map to many

Governance fails when it lives in PDFs nobody applies. Turn duties into controls that run where the work happens: one AI inventory, each system risk-classified once and that classification mapped to every regime it touches, evidence captured as a by-product of operation rather than reconstructed for an audit. Build it into delivery from day one (Section 9) rather than bolting it on

after use-case sprawl has outrun you — and lean on a recognised framework so a single control set answers the EU AI Act, US state laws, and Australian expectations at once. Security is inseparable from this: the OWASP LLM Top 10 and agent-security controls in Section 14 are governance requirements, not a separate track.

13.8 The AI policy set: what to write and how to keep it

Section 13.7 listed the components of a governance programme. This subsection is about the documents themselves — the artefacts you have to write, approve, enforce and keep current. It is the least glamorous part of governance and the part most often done badly, usually in one of two opposite directions: a policy so restrictive that staff route around it, or so vague that it is indistinguishable from having none.

Get the layering right. The first decision is structural. Conventional policy architecture has four layers, and using them properly solves the problem everyone complains about — that AI moves too fast for a policy to stay current.

- **Policy.** A short statement of intent and mandate, approved at executive or board level. It states principles, scope and accountability. It should be stable enough to review annually.
- **Standard.** The specific, testable requirements that make the policy operable — which data classes may be used with which category of tool, what requires human review, what must be logged. Owned by security or risk, revised as needed.
- **Procedure.** How a particular team actually complies — the request form, the approval route, the review checklist. Owned by the operating teams.
- **Guideline and tool list.** The approved tool list, model-specific configuration advice, prompt guidance. This is where the volatility belongs, and it can change weekly without touching anything that needs board approval.

Organisations that put everything in one document face a choice between an out-of-date policy and a quarterly re-approval cycle nobody will sustain. Push the fast-moving material down the stack.

New document, or amendment? The second decision is what should be new. ISO/IEC 42001 makes this an auditable question in its own right: the organisation is expected to determine how existing policies apply to AI and where genuinely new ones are needed. Data-handling rules usually belong as an extension to the existing data classification standard, not a parallel AI rulebook that will drift out of alignment with it. A standalone AI acceptable use policy earns its place when the audience is wider than the existing security policy reaches — contractors, secondees, volunteers — or when a separate attestation is required. The failure mode to avoid is the one that happens when an AI policy is drafted by a governance or ethics group without security in the room: two documents, both live, that define permitted data handling differently.

What an AI acceptable use policy contains

The clause set below is drawn from real published policies rather than vendor templates — notably a US state technology department's policy issued in March 2026, and a bar association's model policy for law firms, both of which are publicly available and worth reading in full.

Clause	What it does	Drafting notes
Scope and definitions	Defines who is covered and what counts as AI. Critically, distinguishes tool categories.	The most useful distinction in practice is between systems procured and governed by the organisation, where data stays inside your boundary, and publicly available consumer tools. Most other rules key off this split.
Permitted and prohibited uses	States what AI may and may not be used for.	Prohibitions should be specific and defensible — bypassing safeguards, generating deceptive material, biometric identification of individuals, unlawful discrimination. Vague prohibitions are unenforceable.
Data rules	Maps data classification to tool category: what may be entered where.	Reference the existing classification scheme rather than inventing AI-specific tiers. If your organisation has no classification scheme, that is the prior problem (Section 8.2).
Tool approval	Requires approval before use of a non-standard tool, with a named route and a form.	The approval route must be faster than the shadow alternative or it will be bypassed. Publish the approved list where people work, not in a policy annexe.
Verification duty	Places responsibility for output on the person who uses it.	The Texas policy phrases this memorably: treat AI output as a starting point written by an unskilled intern. The user remains responsible for what they publish regardless of what produced it.
Disclosure and attribution	States when AI involvement must be declared.	Draw the line at substance: attribution required where AI significantly shaped the ideas or content, not for spelling and grammar. A “when in doubt, disclose” default resolves the rest.
Human oversight	Requires a human decision-maker on consequential outcomes.	Define “consequential” concretely. A heightened-scrutiny category requiring written risk assessment and named-officer approval before deployment is the pattern used in public-sector policies.
Monitoring and logging	Notifies users that prompts and outputs may be logged and reviewed.	Necessary for the evidence trail, and required for fairness. In public bodies, note that prompts may be subject to freedom-of-information access.
Incident reporting	Gives a named channel and a timeframe for reporting AI-related incidents.	AI incidents are often behavioural — drift, confidently wrong output — rather than a breach, so the trigger definition needs care.
Training and attestation	Mandates training and records acknowledgement.	The EU AI Act’s AI-literacy duty has applied to all deployers since February 2025 and extends to people acting on your behalf, not only employees. It expects a programme, not a one-off session.
Consequences, retention and review date	States enforcement, how long records are kept, and when the policy is next reviewed.	Put a dated review commitment in the document itself. A policy with no review date is a policy nobody will revisit.

The rest of the policy set

Acceptable use is the document everyone writes first and the smallest part of the set. A working programme also needs: a register or inventory of AI systems and use cases covering the full lifecycle from proposed through piloted, deployed and retired; an impact assessment procedure with a defined high-impact category and a mandatory set of controls attached to it; a human oversight standard; an AI incident response procedure; third-party

and vendor AI requirements tiered by risk; and, because it is the most externally regulated internal use of AI, a specific position on AI in recruitment and HR decisions (Section 11.8).

One artefact worth borrowing from the public sector is the AI transparency statement — a short public document setting out how the organisation uses AI. Australian government agencies now publish these as standard. It is cheap, it forces internal clarity, and it is increasingly what customers and procurement teams ask for.

Implement once, comply many

Section 13.4 made the case for adopting a framework; this is where it pays. The frameworks overlap heavily, so the efficient approach is to write one internal control and tag it against every external requirement it satisfies, rather than maintaining parallel documents per framework. NIST publishes a crosswalk between the AI RMF and ISO/IEC 42001 for exactly this purpose. A worked example of the mapping:

Your internal control	ISO/IEC 42001	NIST AI RMF	EU AI Act
AI policy, approved and communicated, with roles assigned	A.2 — AI policy and internal organisation	GOVERN 1.1, 1.2 — legal requirements understood; trustworthy-AI characteristics in policy	Deployer obligations; supports Article 4 AI literacy
AI system inventory and registration	A.5 — AI system lifecycle management	MAP 1 — context established and documented	Registration duties for high-risk systems
Impact assessment before deployment	A.4 — impact assessment	MAP 5 — impacts characterised	Fundamental rights impact assessment; data governance under Article 10
Human oversight standard with named reviewers	A.9 — responsible use	MAP 3.5, MEASURE 3.2 — human oversight and monitoring	Articles 14 and 26 — human oversight, competent assigned persons
Evaluation and monitoring records (Section 9.10)	Clause 9 — performance evaluation	MEASURE — the whole function	Article 11 technical documentation; Article 72 post-market monitoring
Automatic, tamper-evident logging (Section 6.10)	A.5, A.6 — lifecycle and data management	MANAGE 4 — documented and monitored	Article 12 — record-keeping, minimum six-month retention
Vendor AI due diligence	A.10 — third-party and supplier relationships	GOVERN 6 — third-party risk policies	Provider and deployer obligations flow down the supply chain

Maintain the mapping as a register rather than a one-off exercise. It is what lets you answer a new requirement by pointing at an existing control, and it is what stops the next framework arriving as a greenfield programme.

Lifecycle: owning it, changing it, excepting it

Ownership. Practice varies and there is no settled answer. The models seen in real organisations are a chief AI or data officer holding the policy with a cross-functional committee beneath, legal holding it with security owning the standards, or the CISO holding the lot. What matters less than the title is that one named person owns the document and a defined body approves changes.

Review cadence. Annual for the policy layer, with a dated commitment written into the document. Standards and the approved tool list should move far more often — quarterly at

least, and immediately when something material changes. This is the payoff for the layering above.

Exceptions and waivers. Every restrictive policy needs a legitimate way out, or it generates illegitimate ones. The strongest model observed — from US federal practice — requires a written, system-specific risk assessment, review by a technical panel and then a governance body, approval one level above the requesting organisation, explicit non-delegability of that approval, annual re-justification, and revocability at any time. Most organisations do not need that weight, but every element of it is answering a real failure: exceptions granted informally, by the person who wanted them, permanently.

Attestation. Record acknowledgement, and treat completion as a governance metric. A policy nobody has attested to is difficult to enforce and easy to disown.

Policy as code, and its limits. Some of this can be enforced automatically. An AI gateway sitting between users and models can implement access rules, block unsanctioned endpoints, apply data-loss prevention and redaction, and produce the log trail — turning the standards layer into version-controlled configuration rather than prose. That is real and worth doing. But be clear about the limit: automated enforcement handles access and data flow well, and cannot adjudicate judgement. Whether an output was adequately verified, whether AI “significantly shaped” a document and therefore needs attribution, whether a use case is consequential — these remain human determinations recorded through the inventory and assessment process. A gateway is not a governance programme.

GOVERNANCE Auditors want the enforcement trail, not the policy document

The consistent report from practitioners taking AI systems through assurance is that the question has shifted. Nobody is impressed by a well-drafted policy; they ask for evidence that the control operated — risk registers, impact assessments, evaluation results, incident logs, access records tying each interaction to a user and an applied rule.

The uncomfortable implication is that most AI deployments currently log nothing that would satisfy this. The policy establishes the rule; the inventory, the gateway logs, the evaluation records (Section 9.10) and the agent audit trail (Section 6.10) are the evidence. Neither substitutes for the other, and the second one takes longer to build.

DELIVERY / PM A policy people route around is worse than no policy

Both failure directions have the same root cause — the policy was written without reference to how the work actually gets done. Too restrictive, and staff move to personal accounts and personal devices, so the organisation loses the visibility it thought it was buying and gains false confidence. Too vague, and “I didn’t know” becomes a complete defence.

The practical test before publishing: can someone doing a normal task follow this policy and still finish their work today, using a tool that is actually available to them? If not, you have written a shadow-AI generator (Section 9.7). Pair every prohibition with a sanctioned route that is fast enough to use.

The Australian reference points. For Australian organisations, the reference points are clear even though most are not binding. The National AI Centre’s Guidance for AI Adoption, published in October 2025, superseded the earlier voluntary safety standard and condenses it into six essential practices — accountability, understanding impacts, managing risk, information sharing, testing and monitoring, and human control — published in both a short foundations form and a detailed implementation guide broadly aligned to ISO/IEC 42001. Federal government entities work to the Digital Transformation Agency’s policy for

responsible AI use, updated in December 2025, with mandatory practices phasing in through 2026 and an accompanying impact assessment tool. New South Wales agencies have had a mandatory AI Assessment Framework since 2024, structured as an assessment questionnaire producing a risk rating and a corresponding set of required assurance activities. An organisation that has already implemented ISO/IEC 42001 for international reasons is, on the National AI Centre's own framing, largely meeting the domestic expectation already.

13.9 The risks you're actually governing

Finally, a checklist of what the controls above exist to manage — each links to fuller treatment elsewhere: bias and unfair outcomes; privacy and data protection; security and misuse (Section 14); intellectual-property and copyright exposure (Section 5); accuracy, hallucination, and over-reliance/automation bias (Section 15); transparency and explainability; third-party, vendor, and concentration risk (Section 12); and environmental impact (Section 12). Good governance doesn't eliminate these — it makes them visible, assigns them an owner, and keeps them proportionate to the stakes.

Go deeper: the governance resource table — NIST, ISO 42001, the EU AI Act explorer and the Australian sources — is in Section 17.6.

13.10 Who owns the output — and can you protect it?

A question every professional using generative tools eventually asks, and one the guide has so far answered only for media (Section 5.6): do you own what the AI produces, and can you protect it? Two different things are at stake — your contractual right to use the output, and your ability to claim intellectual-property protection over it — and they have different answers.

Usage rights are a contract question. The major platforms generally assign you ownership of, or a broad licence to, the outputs you generate on paid plans, and several add IP indemnification against third-party infringement claims — but terms vary by provider and plan, some free tiers reserve broader rights, and open-weight model licences carry their own conditions (Section 2.4). Read the terms before you monetise, and put ownership and indemnity in the contract (Section 9.11).

Copyright protection is a law question, and the current US answer is the clearest: the Copyright Office holds that purely AI-generated material is not copyrightable, because protection requires human authorship. A work can still qualify where a human exercises sufficient control over its expressive elements — creative selection, arrangement and editing count — but a prompt alone generally does not, and applicants must disclaim the AI-generated portions when they register. The US Supreme Court declined to reopen the question in March 2026, leaving that position in place. The EU and Australia reach similar ground by different routes (both require human creativity), and Australia has separately rejected a blanket text-and-data-mining exemption for AI training (Section 13.6). The practical posture: assume raw, unedited AI output may not be protectable, treat meaningful human authorship as what earns protection, and keep records of the human contribution where it matters commercially.

Next section: Section 14 — AI Security & Threats: the adversarial view — the OWASP LLM Top 10, MITRE ATLAS, prompt injection and the lethal trifecta in depth, data and supply-chain risks, and how to defend AI systems. The security anchor section.

14. AI Security & Threats

Threat frameworks are relatively stable; specific incidents and tooling move fast — verified 29 July 2026. This is the security anchor section, gathering the threads seeded in Sections 5, 6, and 7 into one adversarial view.

Every previous section has touched security in passing — prompt injection in prompting and agents, the lethal trifecta in agents, deepfake fraud in media, vulnerable code in development. This section takes the adversary's seat: how AI systems get attacked, how AI is used to attack, and how to defend. It is written for a technically literate reader who needs the real threat model, not reassurance.

14.1 Why AI security is different

AI changes security in two directions at once. It expands the attack surface of the systems you build, and it arms the people attacking you. A few distinctions frame the rest of the section:

- **Security vs safety.** Security is about adversaries — deliberate attempts to make a system misbehave. Safety is about unintended harm and reliability (Section 15). This section is the adversarial half.
- **A genuinely new vulnerability class.** As Section 6 noted, a model cannot reliably separate trusted instructions from untrusted data — both arrive as the same stream of tokens. That single fact underlies most AI-specific attacks and has no clean fix.
- **Runtime behaviour is the new perimeter.** You can't fully secure an AI system by inspecting it once before launch; its behaviour depends on inputs it sees in production, so defence has to operate at runtime.

14.2 The maps: OWASP LLM Top 10 and MITRE ATLAS

Two reference frameworks organise the field, and they're complementary rather than competing. The [OWASP LLM Top 10](#) is a developer- and AppSec-facing checklist of the most critical risks in LLM applications — what to defend. Its 2025 (v2) list:

OWASP LLM risk	What it is
LLM01 Prompt injection	Malicious instructions in input override the system prompt — direct, or indirect via retrieved content
LLM02 Sensitive info disclosure	The model leaks secrets, PII, or proprietary data in its output
LLM03 Supply chain	Compromised models, datasets, packages, or MCP servers
LLM04 Data & model poisoning	Tainted training, fine-tuning, or RAG data plants backdoors or bias
LLM05 Improper output handling	Model output is trusted and reaches a shell, browser, or database unchecked
LLM06 Excessive agency	An agent has more tools, permissions, or autonomy than the task needs
LLM07 System prompt leakage	The hidden system prompt — and any secrets or logic in it — is extracted
LLM08 Vector & embedding weaknesses	RAG retrieval is manipulated or leaks data across users/tenants
LLM09 Misinformation	Confident, wrong output is relied upon (overlaps hallucination, Section 15)
LLM10 Unbounded consumption	Runaway or adversarial usage drives cost ("denial of wallet") or denial of service

[MITRE ATLAS](#) is the other map: an adversary-focused matrix of tactics and techniques for attacking AI systems (the AI counterpart to MITRE ATT&CK), with real incident case studies,

used by red teams and threat modellers. Where OWASP tells you what to defend, ATLAS shows how an attack actually unfolds. A newer OWASP Top 10 for Agentic Applications (2026) extends the list to agent-specific risks — goal hijacking, tool misuse, memory poisoning, and the amplification that comes when several LLM weaknesses combine in an autonomous loop. Its risks are catalogued as ASI01–ASI10, and the Cloud Security Alliance’s MAESTRO framework offers a companion method for threat-modelling agentic systems layer by layer.

14.3 The attacks that matter most

A handful of attack classes account for most real-world risk:

- **Prompt injection** (LLM01) is the most exploited vulnerability, ranked number one for three years running, and considered architecturally unsolved. Direct injection is an adversarial user input; the nastier indirect form hides instructions in content the model later reads — a web page, an email, a document, a code comment. A distinction worth keeping: a jailbreak tries to bypass the model’s safety training; an injection hijacks the application’s instructions. Guardrails reduce both but eliminate neither.
- **The lethal trifecta and exfiltration.** As Section 6.13 detailed, an agent combining private-data access, exposure to untrusted content, and an external-communication channel can be turned into a data-exfiltration tool by a single injected instruction. This is where injection stops being a curiosity and becomes a breach.
- **Data and model poisoning** (LLM04) corrupts what the model learns or retrieves. Poisoned training or fine-tuning data can plant hidden backdoors; poisoned RAG sources are startlingly efficient — 2026 research showed a handful of crafted documents could steer answers the large majority of the time; and “memory poisoning” can plant instructions that surface weeks later, a sleeper-agent pattern.
- **Information leakage** (LLM02, LLM07) covers the model disclosing sensitive data, and attackers extracting the hidden system prompt to learn its rules and secrets — or, more subtly, inferring training data from model behaviour.
- **Supply-chain compromise** (LLM03) targets the components you trust: backdoored open models, poisoned datasets, trojanised packages, and malicious MCP servers, including the slopsquatting of hallucinated package names covered in Section 7.6.
- **Excessive agency and unbounded consumption** (LLM06, LLM10) are the agent-era additions: an over-permissioned agent turns a small compromise into a large one, and adversarial or runaway usage can rack up ruinous cost or take a service down.

These are not hypothetical. The postmark-mcp package became the first confirmed malicious MCP server found in the wild — fifteen clean releases, then a single added line that quietly copied users’ email to an attacker — a textbook supply-chain compromise (LLM03) that no pattern-scanner would have flagged. It is the concrete form of the risk this section keeps returning to: the components an agent trusts are now part of your attack surface.

The list grew again in July 2026: “DuneSlide” (CVE-2026-50548/-50549, CVSS 9.8) chained prompt-injected content into an overwrite of Cursor’s sandbox-enforcer binary — zero-click, OS-level code execution escaping the terminal sandbox — and Microsoft incident responders documented MCP tool-poisoning (malicious tool descriptions steering agents)

used in the wild against fintech MCP servers, demonstrating the same chain against a Copilot Studio finance agent.

14.4 AI as the attacker's weapon

The other half of the picture is AI lowering the cost and raising the quality of attacks. This is threat awareness, not a how-to, but defenders need to know the shape of it. Deepfakes and voice cloning have already enabled multi-million-dollar fraud (Section 5), and the practical consequence is blunt: a familiar face or voice on a call is no longer proof of identity.

Generative text makes phishing cheap, fluent, and personalised at volume, erasing the spelling-and-grammar tells people were trained to spot. AI assistance speeds up parts of malware development and vulnerability discovery, and lowers the barrier to entry for less-skilled attackers. And synthetic media makes disinformation faster to produce and harder to debunk. The summary a security team should internalise: AI arms both sides, so authentication, verification, and provenance (Section 5) matter more than ever.

The clearest sign of where this is heading came in November 2025, when Anthropic disclosed the first reported, largely AI-orchestrated cyber-espionage campaign. A state-linked group (tracked GTG-1002) jailbroke an AI coding agent into believing it was running authorised security tests, then had it execute an estimated 80–90 per cent of an intrusion campaign against roughly thirty targets, with human operators stepping in only at a few key decisions. Set alongside the steady progress of autonomous vulnerability-discovery systems, the direction is unmistakable: the effort required to mount a capable attack is falling, which is exactly why detection, verification, and least-privilege containment matter more every year.

14.5 Defending AI systems

The governing principle follows from 11.1: you cannot filter your way to safety, so assume a determined injection or jailbreak can succeed and design so that success is survivable. That means defence-in-depth — OWASP recommends at least three independent layers — combined with blast-radius containment. The practical layers:

- **Contain agency (the highest-value layer).** Least privilege by default, scoped short-lived credentials, human approval on irreversible or external actions, sandboxed execution, and allowlisted MCP servers/tools (Section 6.13). Over-privileged API keys are the new admin credentials — audit them.
- **Guardrails on input and output.** Prompt/response filtering and policy checks (tools like NeMo Guardrails, Llama Guard, Rebuff) reduce risk — necessary but never sufficient, since multi-turn and indirect attacks bypass single layers.
- **Red-team and evaluate continuously.** Adversarial testing with tools like Garak and PyRIT, and bug-bounty programs, run as an ongoing gate rather than a one-off — guardrails you don't test have blind spots.
- **Monitor at runtime.** Behaviour baselines and anomaly detection (AI security posture management), because detection is most reliable downstream, when the system deviates from normal.
- **Secure the supply chain and data.** Vet and sign models and datasets, keep an AI bill of materials, vet MCP servers, run SAST/SCA on AI-generated code (Section 7.6), keep secrets out of prompts, and track data provenance and integrity. For MCP

specifically, the November-2025 specification mandates OAuth 2.1 for remote servers, and the July-2026 revision promotes Enterprise-Managed Authorization (IdP-controlled access to MCP servers) to stable, with a Linux Foundation-hosted signed registry planned — prefer registered, authenticated servers over the long tail of unvetted public ones.

None of this is exotic — it is classic security discipline (least privilege, zero trust, defence-in-depth, supply-chain hygiene) applied to a new substrate, plus the AI-specific layers of guardrails and adversarial evaluation.

One newer direction deserves flagging, because it attacks the root of the problem rather than its symptoms: designing systems so that a successful injection cannot do harm, instead of trying to filter injections out. The dual-LLM pattern splits the work between a privileged model that plans and holds tool access but never reads untrusted content, and a quarantined model that reads untrusted data but cannot act. Google DeepMind's CaMeL takes this further, deriving the control and data flow from the trusted instruction and attaching capability metadata to every value so untrusted input can never redirect the program — in testing it completed about 77 per cent of tasks with provable security guarantees, against 84 per cent for an undefended agent. These designs are not free — they add engineering effort and a modest capability cost, and are still maturing — but they are the most promising answer yet to a problem that filtering alone cannot solve, and they pair naturally with the lethal-trifecta containment in Section 6.13.

DELIVERY / PM Build security into the delivery pipeline

Security can't be a gate bolted on before launch and forgotten. Fold it into the AI delivery lifecycle from Section 9: threat-model each use case at design time, run adversarial red-teaming and security evals as a required gate before production (alongside the quality evals), and re-run them whenever the model, prompt, tools, or data change. Treat AI-generated code as unreviewed third-party code (Section 7.6). The teams that scale AI safely are the ones for whom "ship it" already includes "red-teamed it."

GOVERNANCE Security is a governance requirement, not a side project

The controls here are also compliance obligations. Regulators increasingly expect risk assessment, audit trails, and incident/breach reporting for AI systems (Section 13), and enterprise procurement and insurers now ask for a demonstrable security posture mapped to the OWASP LLM Top 10, MITRE ATLAS, and the NIST AI RMF. Fold AI security into your existing security governance and inventory — agents and models are assets that need an owner, a risk tier, and a place in incident response — rather than running it as a separate, unowned track. Expect this to harden into law: the EU AI Act's high-risk security duties (logging, robustness, human oversight) now land in 2027–28 after the Digital Omnibus delay, and a 2026 NIST/CAISI request-for-information on AI-agent security signals dedicated agent rules to come.

Go deeper: security and red-team evaluations are part of the broader evaluation discipline in Section 9.10 — run them on the same cadence as your quality evals.

14.6 Where to start

For an organisation without a dedicated AI-security team, a few high-leverage moves cover most of the risk: inventory and risk-tier your AI systems and agents (Section 13); enforce least privilege and human approval on anything consequential; put guardrails on at least the top three OWASP risks (prompt injection, sensitive-information disclosure, improper output handling); red-team before you expose a system to untrusted input or external users; and

treat every AI output — text, code, or tool call — as untrusted until checked. Perfect security isn't the goal; making a successful attack expensive, contained, and detectable is.

Go deeper: the security resource table — OWASP, ATLAS, red-teaming guides and tooling — is in Section 17.6.

14.7 When it goes wrong — an AI incident playbook

Every earlier control reduces the odds of an incident; none removes them. Because AI systems fail in ways classic runbooks do not anticipate — a hallucinated fact reaching a customer, an agent taking a wrong irreversible action, an indirect injection quietly exfiltrating data — it is worth deciding in advance how you will respond. The pieces are named across Sections 9.7, 13.8 and 14.5; this pulls them into one sequence.

- **Detect and triage.** Know how you would even find out — runtime monitoring and anomaly detection (Section 14.5), plus a route for staff and customers to report bad output. Classify by blast radius: who saw it, what data moved, what action was taken, and whether it is still happening.
- **Contain.** Stop the bleeding before diagnosing it: revoke the agent's credentials or pause the tool, cut the external channel that completes an exfiltration (Section 6.13), and roll back or quarantine wrong actions. Least privilege and reversibility, designed in earlier, are what make this possible.
- **Communicate.** Decide in advance who owns an AI incident and who must be told — affected users, and regulators or partners where breach- and incident-notification duties apply (Section 13). Silence and over-reassurance both erode trust, which recovers only partially after a visible failure (Section 9.7).
- **Learn and prevent.** Treat every incident as an eval: add the failing case to the regression set (Section 9.10), fix the class of error rather than the instance, and re-run adversarial testing whenever the model, prompt, tools or data change (Section 14.5).

The goal is the same as the rest of this section: not a system that never fails, but one whose failures are detected fast, contained small, and never happen the same way twice.

Next section: Section 15 — Safety, Reliability & Limitations: the non-adversarial failure modes — hallucination, bias, over-reliance, and the limits of what today's AI can reliably do — and how to work within them.

15. Safety, Reliability & Limitations

Research area; findings and figures verified 29 July 2026 and presented as indicative. This section is the non-adversarial counterpart to Section 14 — what goes wrong without anyone attacking. Note the reflexive point: this guide is itself AI-assisted, and so is subject to the very limitations it describes, which is why it is dated, sourced, and cross-checked.

Where Section 14 covered deliberate attacks, this section covers the failures that happen on a good day with no adversary at all: the model that states a falsehood with total confidence, the tool that’s brilliant on one task and hopeless on a trivially similar one, the quiet erosion of a user’s own judgement. Understanding these limits isn’t pessimism — it’s the difference between using AI well and being misled by it.

15.1 Safety, reliability, alignment — what we mean

“AI safety” gets used in two quite different senses. The practical sense — the focus of this section — is about reliability and unintended harm in everyday deployment: accuracy, consistency, fairness, and the effects of use on people. The frontier sense concerns longer-horizon questions about highly capable systems (deception, loss of control, alignment of goals); those are real research topics, but they belong to the outlook in Section 16, and this section stays with the failures you can observe in a system you’re running today. The distinction from Section 14 is simple: security is about adversaries, safety is about everything that can go wrong on its own.

15.2 Hallucination: the headline limitation

Hallucination — confidently generating false or unsupported content — is the single most important practical limitation of today’s AI. The crucial insight from 2025–26 research is that it is not mainly a data-quality bug but an incentive problem: models are trained to predict the most plausible next token, and the benchmarks that rank them reward a confident guess over an honest “I don’t know.” So models learn to bluff. Fixing it therefore means changing the rules — rewarding calibrated uncertainty and treating refusal as a valid answer — which is an active research direction (Anthropic, the maker of Claude, has shown refusal can be a learned internal policy rather than a fragile prompt trick). A counter-intuitive corollary: more reasoning does not always mean fewer hallucinations; in some “slow-thinking” models a chain of reasoning can compound an early error into a confident wrong conclusion.

Rates vary enormously by task, which is why a single “hallucination rate” is meaningless. Indicative figures from various 2025–26 studies (different models and methods, not directly comparable):

Task type	Indicative rate	Note
Grounded summarisation (top models)	<1.5%	Source text supplied to the model
General factual queries, no grounding	~15–20%	Stanford HAI
General conversational use	~31%	Real-world benchmark
Medical case summaries	~43–64%	Highly prompt-dependent
Legal research queries	~58–88%	High-stakes, across major models
Code citing non-existent libraries	up to ~99%	When prompted with fake package names

The pattern is clear: grounding collapses the rate. That’s why the dominant reframe is data-centric — “govern the context the model retrieves” rather than “fix the model” — with one

analysis attributing roughly 72% of enterprise AI failures to inadequate context rather than model capability. The practical mitigations stack: retrieval-augmented generation and tool use to ground answers in real sources; prompting that permits “I don’t know” (one study cut a model’s rate from 53% to 23% with prompt changes alone); confidence and uncertainty signals; and human verification for anything that matters. One specific trap: models fabricate citations, references, and quotes that look perfectly formatted — never trust an AI-supplied source without checking it exists.

15.3 Bias and fairness

Models learn from human data, so they absorb human bias — and can amplify it. Two kinds of harm matter: representational (stereotyping, or under- and mis-representing groups) and allocative (unfair distribution of real opportunities, such as who gets shortlisted, approved, or flagged). The stakes are highest exactly where regulation focuses (Section 13) — hiring, lending, healthcare, justice — and there is evidence the feedback loops of repeated AI use can shift human judgement more than human-to-human interaction does. Mitigation is possible but imperfect: curated and representative data, fairness evaluations and audits, human oversight, and documentation of known limitations. Fairness can’t be fully “solved,” partly because it is genuinely multi-dimensional — different reasonable definitions of fair can conflict — so the honest goal is to measure, disclose, and manage bias rather than to claim it’s been eliminated.

Two practical questions turn this from a principle into a control. Where does the bias enter? Three places: the training data, which reflects historical patterns including discriminatory ones; the labelling and fine-tuning, whose human judgements get encoded as ‘correct’; and deployment, where a biased output feeds the next decision and the next dataset in a loop. And how would you even know? Fairness is measured, not asserted — disaggregate outcomes by group and test for disparate impact, run the same cases through the system changing only a protected attribute, and audit the decisions the model actually influences rather than the model in the abstract. Because several reasonable definitions of ‘fair’ can conflict mathematically, choose the definition that fits the decision and the law before you measure, not after.

This is also the most externally regulated internal use of AI, which is why it recurs across the guide: hiring and HR decisions carry specific obligations (Section 11.8), the EU AI Act classes many of these uses as high-risk (Section 13.3), and bias-audit and disclosure duties are appearing in US state and city law. The workable mitigation stack is therefore both technical and procedural — representative data and documented limitations; a fairness evaluation run as a gate rather than a one-off; a named human accountable for consequential decisions with real authority to overrule the model, not a rubber stamp; and a record of what was tested and found, because that documentation is exactly what a regulator, auditor or claimant will later ask for (Section 8.4). The honest goal, again, is to measure, disclose and manage bias — not to claim it has been eliminated.

15.4 Reliability and the jagged frontier

Even setting truth aside, AI performance is uneven in ways that surprise people. The vivid metaphor is the jagged frontier: a system can be superhuman at one task and fail a trivially similar one. Documented examples include maths accuracy dropping sharply when

irrelevant details are added to a problem, and weak spatial reasoning such as reliably counting objects in an image. Three related properties compound this:

- **Non-determinism.** The same prompt can yield different answers on different runs — fine for brainstorming, a problem for anything that needs to be repeatable.
- **Brittleness.** Small, meaning-preserving changes in wording can change the answer; performance also varies by language and domain.
- **Silent, confident failure.** The most dangerous property — models rarely signal low confidence, so a wrong answer looks exactly like a right one. And because output is fluent narrative rather than a probability score, it can feel more authoritative than an old-style classifier that at least said “72% sure.”

The takeaway is to calibrate expectations to the specific task rather than the model’s general reputation, and to test on your own workload rather than trusting a headline benchmark.

Knowing the frontier is jagged, the engineering question is how to build something dependable on top of an undependable component. You cannot make the model reliable; you can make the system around it reliable. The patterns that do the work:

- **Ground and constrain.** The single biggest lever is retrieval and tool use (Sections 6 and 8): a model answering from supplied sources fails far less than one answering from memory. Structured output and schema validation catch malformed results before they propagate.
- **Verify, don’t trust.** Check the output with something independent — a second model or a ‘judge’, self-consistency (sample several times and compare), a validator that runs the generated code or re-computes the number, or a rules check. Cross-checking catches the confident-but-wrong answer the model will not flag itself.
- **Design for graceful failure.** Assume some answers will be wrong and make that survivable: confidence thresholds that route low-certainty cases to a human, a fallback to a simpler deterministic path, and clear ‘I can’t answer that’ behaviour instead of a confident guess. A system that fails loudly beats one that fails silently — the core risk above.
- **Put the human where the stakes are.** Human-in-the-loop is a dial, not a switch: review everything for high-consequence decisions, sample for medium, monitor in aggregate for low — and place the checkpoint on the irreversible or external action (Section 6.13), not on every step.
- **Monitor in production.** Reliability is not established once before launch; behaviour drifts as inputs, models and data change. Baseline it, watch for anomalies, and re-run your evaluations continuously (Section 9.10) — the same regression discipline the rest of the guide keeps returning to.

None of this makes a probabilistic system deterministic. It makes its failures rarer, cheaper and visible — which, for most real deployments, is the achievable definition of reliable.

15.5 Over-reliance and the human factor

Some of the most consequential limitations aren’t in the model at all — they’re in how people use it. Automation bias is the well-documented tendency to trust automated output without enough scrutiny; with AI it is amplified because the output is articulate and confident. Two dynamics deserve attention:

- **Sycophancy.** Models tend to agree with the user and tell them what they seem to want to hear. Recent work reframes this as active, reactive persuasion — the system argues back — which means “just keep a human in the loop” may not be enough if the loop itself is being nudged. Structural safeguards help: forming an independent view before consulting the AI, parallel validation, and separating AI-assisted drafting from final judgement.
- **Deskilling and critical-thinking erosion.** Early evidence suggests heavy, passive reliance can weaken the very skills it substitutes for; studies find that users who actively question and selectively restrict their AI use show stronger critical thinking than passive users, and that student work produced by uncritically accepting suggestions tends to be less creative and less appropriate. The cautious consensus: unchecked reliance can turn AI from an aid into a cognitive substitute, producing gains that are efficient but fragile and that decay without sustained human engagement. At the extreme, researchers document genuine cognitive-dependency concerns.

None of this is an argument against using AI — it’s an argument for using it as a collaborator you check rather than an oracle you defer to. The healthiest posture keeps your own judgement in the driver’s seat.

15.6 Other limits and emerging concerns

A briefer roll-call of limits worth keeping in view: a knowledge cutoff means built-in staleness on recent events unless the model can search; context windows are finite and effective attention degrades well before the advertised limit (“context rot,” Sections 1 and 6); genuine multi-step reasoning and long-horizon planning remain unreliable despite improvements; and there’s an open debate about whether models “understand” in any meaningful sense or are sophisticated pattern-matchers — a distinction that matters when you’re tempted to assume human-like common sense. Environmental cost is a real limitation too (Section 12). And at the frontier sit longer-horizon safety and alignment questions — deception, situational awareness, loss of control — which are active research areas rather than everyday deployment issues; Section 16 picks up that thread.

One under-discussed dimension cuts the other way: accessibility. Used well, these tools are a powerful assistive technology — real-time captioning, plain-language rewriting, image description, voice control — and for many users that is their single biggest benefit. But the AI interfaces themselves must meet accessibility obligations (WCAG, and public-sector duties across all three regions), and outputs can encode barriers of their own — image generators that default to narrow representations, voice systems that stumble on atypical speech. Treat accessibility as both an opportunity to design for and a requirement to meet, not an afterthought.

15.7 Working within the limits

The limitations above point to a consistent set of habits rather than a reason to avoid AI:

- **Verify anything that matters.** Treat AI output as a capable draft, not a finding — check facts, sources, numbers, and code before they carry weight.
- **Ground it.** Give the model the source material (RAG, tools, documents) instead of relying on its memory, and it will be far more reliable.

- Match the model and mode to the task, and test on your own workload rather than a leaderboard.
- Keep human judgement structurally central — and remember sycophancy: form your own view before you ask, on decisions that count.
- Disclose AI use, and design for graceful failure so a wrong answer is caught, not shipped (the reliability patterns in Section 15.4).
- Know when not to use it — some tasks demand guarantees or accountability that a probabilistic system can't provide.

DELIVERY / PM Reliability is a delivery requirement

Ship AI features the way you'd ship anything safety-relevant: define what "accurate enough" means for the task, build evaluation sets that measure accuracy and consistency (not just vibes), and gate high-stakes releases on them — the same eval discipline from Sections 7 and 9. Put a human checkpoint on consequential outputs, and prefer grounded architectures over ungrounded generation wherever a wrong answer has a cost. "It demoed well" is not evidence of reliability.

GOVERNANCE Accuracy and fairness are governance obligations

Reliability and bias aren't just quality issues — they're increasingly compliance ones. High-risk-use rules expect accuracy, human oversight, bias assessment, and transparency about AI's involvement (Section 13). Bake the limitations into policy: require disclosure that AI was used, mandate verification for consequential decisions, document known failure modes and bias testing, and make sure someone owns the question "how wrong can this be, and who's accountable when it is?"

Go deeper: evaluation and verification skills are threaded through Stages 1–3 of the curriculum, Sections 17.3–17.5.

Next section: Section 16 — Trends & Near-Term Outlook: where AI is heading over the next 12–24 months — the trajectories worth watching, and the ones to treat with healthy scepticism.

16. Trends & Near-Term Outlook

⚡ *SPECULATIVE & FAST-MOVING* — this section is informed forecasting, not fact, and dates fastest of all. Written from a 29 July 2026 vantage; treat every projection as a signal to watch, not a certainty. Re-read it as a snapshot of what mid-2026 expected, and judge it against what actually happened.

A guide that stops at the present tense is only half useful. This section looks 12–24 months ahead — the horizon where informed judgement still beats a coin flip — and it picks up the frontier-and-alignment thread deferred from Section 15. The field forecasts itself badly in both directions, over-hyping the imminent and under-rating the eventual, so the honest posture is calibrated attention: know which trajectories have real momentum, and which deserve a raised eyebrow.

16.1 How to read a forecast about AI

Two framing points. First, near-term (roughly a year out) is far more predictable than the multi-year picture; the further out a claim reaches, the more it's a bet dressed as a projection. Second, there's been a genuine vibe shift over the past year: the “superintelligence in two years” narrative that dominated 2024–25 has cooled, as flagship releases delivered incremental rather than revolutionary gains and even prominent forecasters walked back their timelines. The emerging consensus isn't that AI has stalled — it's that enormous value is available well before anything like AGI, so the near-term story is about deployment and reliability, not a sudden singularity.

16.2 The capability trajectory

The concrete directions with real momentum over the next year or two:

- **Agents keep maturing — unevenly.** The move from chat to action continues, with agent “control planes” and multi-agent orchestration emerging. But the winners are the narrow, deeply-integrated agents, not the do-everything autonomous ones; expect a shake-out (analysts predict a large share of agentic projects will be cancelled) and beware “agent-washing” of rebranded automation (Sections 6 and 9).
- **Reasoning gets adaptive.** Rather than always thinking hard or never, models increasingly vary their effort by task difficulty, and hybrid stacks route simple queries to fast models and hard ones to reasoning models — making the economics of reasoning workable (Sections 4 and 12).
- **Multimodal-by-default and ambient.** Native voice, vision, and video become the norm, and AI increasingly disappears into the tools people already use rather than living in a separate chat box.
- **Open-weight and on-device rise.** Efficient open models (sparse MoE, long context, agent-ready) keep closing the gap, and smaller models push to phones and laptops for privacy, latency, and cost (Section 2).
- **AI joins scientific discovery.** Beyond summarising research, models begin proposing hypotheses and helping run experiments in fields like biology, chemistry, and materials — one of the more credible high-upside trajectories.

The counterweight to all of this is reliability. Large-scale evaluations still show agents succeeding only about half the time on multi-hour tasks (and needing far shorter tasks for high reliability) — a stray pop-up can derail one — so dependable automation of long,

complex work remains out of reach for now. Progress is a staircase, not a cliff — though the steps keep coming: mid-to-late July alone brought four frontier releases (Grok 4.5, GPT-5.6 GA, Kimi K3, Opus 5) and a teased Gemini 4.

16.3 The scaling and AGI debate

The deeper question under the capability trajectory is whether simply scaling up keeps working. The pure pre-training scaling that drove 2020–23 shows signs of diminishing returns, and attention has shifted to other levers: test-time compute (reasoning), better data, and post-training techniques. On timelines, respected researchers who once sounded urgent now offer sober ranges — estimates clustering anywhere from about five to twenty years for something AGI-like, with broad industry consensus drifting toward the 2030s rather than the mid-2020s. The practically important point isn't who's right about the far horizon; it's that trillions in value are up for grabs from today's already-powerful systems, which is where the next two years will actually be fought.

16.4 The economics and industry trajectory

Section 12 covered the money; the near-term arc is that infrastructure constraints — compute supply, and above all power — will shape 2026–27 as much as any algorithmic advance, with chip generations turning over yearly and capex still climbing. Expect base models to keep commoditising as capable open-weight options proliferate and token prices fall, pushing durable value toward the application and agent layers where proprietary data and workflow integration create moats. And expect the “prove-it” phase to bite: as the major labs face public-market scrutiny, the pressure to convert capability into profit intensifies, raising the odds of consolidation, repricing, or a correction somewhere in the stack — without necessarily denting the underlying usefulness of the technology.

16.5 The governance and safety trajectory

On the regulatory side, the direction is one-way even as the details wobble: the EU AI Act's heavy high-risk obligations land in 2027–28, more jurisdictions are legislating, and AI-governance leadership is becoming a standard corporate role. The message from Section 13 holds — build the governance muscle now, because the timeline only tightens.

This is also where the frontier-safety thread deferred from Section 15 belongs. Beyond everyday reliability sit longer-horizon concerns about highly capable systems: could a model learn to appear aligned while pursuing other goals (“deceptive alignment”), behave differently when it senses it's being tested (“situational awareness”), or resist oversight? These are active research areas, not science fiction — and the current empirical picture is measured rather than alarming: controlled experiments can elicit mild deceptive tendencies in some advanced reasoning models, but today's systems don't show overt uncontrolled behaviour. The response ecosystem is maturing: mechanistic interpretability (reading a model's internals), dangerous-capability evaluations (for cyber, bio, AI-R&D, and deception), independent auditors and national AI Safety Institutes, an annual [International AI Safety Report](#) drawing on experts from dozens of countries, and voluntary frontier-safety frameworks from the major labs — several now being formalised into law for frontier developers. The honest caveat is a structural one that Anthropic among others has acknowledged: far more compute goes into building capabilities than into making them safe, so safety techniques often arrive reactively. The milestones that matter most over the next

few years are unglamorous — standardised evaluations, production-grade interpretability, shared red-teaming — rather than dramatic breakthroughs.

16.6 Trends to treat with scepticism

A field this hyped needs a filter. A rough guide to separating signal from noise:

The hype	A more grounded read
“AGI is one or two years away”	Respected researchers now say ~5–20 years; huge value exists well before then
“2026 is the year agents replace workflows”	Real but uneven; many agent projects will be cancelled, and “agent-washing” is rife
“This benchmark score means it’s solved”	Benchmarks are gamed and performance is jagged — test on your own workload (Section 15)
“AI will replace your team”	Evidence favours augmentation; the bottleneck shifts to human judgement
“The new model changes everything”	Most releases are incremental; capability is a staircase, not a cliff

The mirror-image error is cynicism — dismissing all of it as a bubble. Both the breathless and the dismissive miss the same thing: a technology that is genuinely, unevenly transformative, arriving faster than institutions adapt but slower than the headlines claim.

16.7 Signposts to watch

Because the point of this section is calibrated attention rather than prediction, here is what to actually watch — the observable signals that tell you which trajectory is playing out, so you update from evidence rather than headlines. None needs inside information; all are public and trackable through the directory in Section 17.9.

Signal to watch	What it would tell you
Frontier benchmark jumps versus plateaus	Whether new scaling and methods are still delivering step-changes, or we are in an incremental phase (Section 16.3).
Agent reliability on long, multi-hour tasks	The real gate on autonomy — track the task length achievable at high reliability, not demo success rates (Section 16.2).
Open-weight models closing on the frontier	How fast base models commoditise, and how much value shifts to the application and agent layers (Section 16.4).
Inference prices and cost-per-capability curves	Whether the steep cost decline continues — the economics beneath every business case (Section 12).
Lab revenue against capex and free cash flow	The health of the build-out, and the odds of a repricing or correction somewhere in the stack (Section 12.6).
Power, grid queues and data-centre siting	The binding real-world bottleneck on capacity — increasingly it constrains more than chips or capital (Section 12.4).
EU AI Act standards, guidance and first enforcement	Whether the regulatory timeline is holding, slipping or starting to bite (Section 13.3).
Safety milestones: standardised evals, interpretability, mandatory frontier rules	Whether the safety ecosystem is maturing to match capability, or falling further behind it (Section 16.5).

Watching a handful of these beats reading a hundred hot takes. When several move together — plateauing benchmarks alongside tightening lab cash flow and slipping enforcement, say — that is the field telling you which of the trajectories above is actually unfolding, months before the commentary catches up.

16.8 How to position yourself

The practical response to a fast-moving field is not to chase every release but to invest in what stays durable. Specific tools, model names, and prices will churn — much of this guide’s fast-moving detail will date within months — but the fundamentals compound: understanding how these systems work and fail, the judgement to know when to trust output and when to verify, skill at framing problems and evaluating results, and the governance and security instincts to deploy responsibly. Build for swappability rather than lock-in, adopt on a rolling basis, watch the underlying signals rather than the launch-day noise, and treat your own calibrated judgement as the asset that appreciates. That is exactly what the final section is for.

DELIVERY / PM Plan for change, not for a fixed target

Don’t architect around one model or vendor as if it’s permanent. Build abstraction layers so you can swap models as price and capability shift (Section 12), keep evaluation sets that let you re-test a new model against your own workload in an afternoon, and adopt on a rolling cadence rather than a big-bang bet. Invest disproportionately in the durable layer — your data, your evals, your team’s judgement — because that’s what survives the next model generation.

GOVERNANCE Get ahead of the regulatory curve

The regulatory timeline only tightens: EU high-risk obligations in 2027–28, more jurisdictions legislating, and frontier-model disclosure rules spreading. Organisations that build the inventory, risk-tiering, and documentation habits now (Section 13) will absorb new requirements as adjustments rather than emergencies — “implement once, map to many” pays off precisely because the rules keep coming. In the US the state patchwork keeps growing faster than any federal framework can pre-empt it; in Australia, July’s announcement of mandatory national AI standards makes the direction concrete. In all three regions, aligning to existing guidance and international standards today turns the coming rules into a step, not a scramble.

Go deeper: the keep-current directory in Section 17.9 is built for watching exactly these trajectories.

Next section: Section 17 — Learning AI: Roadmap & Resources: how to go from here to genuine fluency — a staged learning path, the skills that endure, and curated ways to keep current. The capstone.

17. Learning AI: Roadmap & Resources

The capstone — now also the curriculum. ↗ FAST-MOVING in part: the roadmap and durable skills age slowly, but the resource directory will shift — every URL below was checked live on 29 July 2026, and is re-checked with each version bump. “Free” means genuinely free; “freemium” means a usable free tier.

Everything so far has been the map. This final section is the journey — and, from this version, a full curriculum rather than general advice. It pulls the whole guide together into a staged path where each stage is a module: what you should be able to do by the end, a core path to work through in order, and the platform tracks — Anthropic/Claude, OpenAI, Google, Microsoft, open-weight/local, and the coding tools — with free courses, documentation, and video for each. It closes with the durable skills, a sustainable way to stay current, and a thirty-day starting plan.

17.1 How to actually learn AI

One principle outranks all the others: learn by doing. Reading about prompting, agents, or evals builds vocabulary; using them on your own real problems builds judgement, and judgement is the thing that transfers. Two habits make the difference. First, keep the durable-versus-perishable distinction from Section 16 front of mind — pour your energy into fundamentals that compound (how these systems work and fail, verification, problem framing) rather than memorising this month’s model names and prices. Second, build a feedback loop: try something, check the result critically, notice where it broke, adjust. That loop — not any course — is what turns a user into a practitioner.

17.2 A staged roadmap

Fluency isn’t one thing; it’s a progression. You don’t need to reach the last stage — pick the one that matches your goals and go deep there. Each stage names where in this guide to build the underlying understanding.

Stage	What it looks like	Guide sections
1. Fluent user	Prompt well, pick the right model for the task, verify output, and get genuinely good documents and knowledge work out of AI	1, 4, 5
2. Power user	Engineer context, connect tools and your own data, use agents as a user, and automate small repetitive workflows	4, 6
3. Practitioner	Build agents and applications, design RAG, run evaluations, choose architectures, and code with AI assistance	6, 7, 12
4. Specialist	Go deep in one lens — governance, delivery/PM, or security — and lead how your organisation adopts AI	8, 10, 11

Most readers of this guide are somewhere around Stage 1–2 and want to move up. The jump that pays off most isn’t technical — it’s learning to ground your work (give the model real source material) and to verify it systematically, because those two habits raise the quality of everything from a quick draft to a production agent.

Each stage now has a full module below (Sections 17.3–17.6): objectives, a core path, and a resource table. Work the core path in order; treat the tables as a menu, not a syllabus — two or three resources done properly beat ten skimmed.

17.3 Stage 1 — the fluent user

Objective: prompt well, pick the right model for the task, verify output, and get consistently good documents and knowledge work out of AI. You can already use a chatbot; this stage makes the difference between asking and directing. Time: roughly 15–20 focused hours spread over a few weeks, applied to your own real tasks throughout (Sections 1, 4, 5 and 15 are the companion reading).

The core path: start with Anthropic's AI Fluency course for the collaboration discipline (it is vendor-made but genuinely general); watch 3Blue1Brown's neural-network series and Karpathy's hour-long intro so the machine stops being magic; then work one official prompting guide properly — Claude's or OpenAI's, matching whichever assistant you use daily — and finish by learning the actual features of your workplace tools (help centres, or the Microsoft/Google starter courses if your organisation runs on 365 or Workspace).

Resource	Type · time	What it is / why it earns its place
AI Fluency: Framework & Foundations (Anthropic) https://anthropic.skilljar.com/ai-fluency-framework-foundations	Course · 3–4 h	The best single non-technical starting point: collaborating with AI effectively, efficiently, ethically and safely. Free, with certificate.
Anthropic Academy https://www.anthropic.com/learn	Course hub	Anthropic's free learning hub — personal use, work deployment, and API tracks; gateway to the full course catalogue at anthropic.skilljar.com .
Claude prompt engineering guide https://platform.claude.com/docs/en/build-with-claude/prompt-engineering/overview	Docs · 2–4 h	The official Claude prompting reference — clarity, examples, XML structure, chaining. Pairs with Section 4 of this guide.
OpenAI Academy https://academy.openai.com/	Courses · 1–5 h each	Free OpenAI courses (AI Foundations, Applied AI Foundations) plus live events and on-demand bootcamps. The OpenAI-side Stage 1 anchor.
OpenAI Help Center https://help.openai.com/	Docs	What ChatGPT actually does — features, projects, memory, data controls, plans. The verification source for your own tool.
Google AI Essentials https://grow.google/ai-essentials/	Course · <5 h	Google's beginner course on generative AI, prompting and responsible use. Free via Google Career Skills; certificate paid (Coursera).
Google Skills: Introduction to Generative AI https://www.cloudskillsboost.google/paths/118	Course path · ~5 h	Free beginner path: generative AI, LLMs and responsible AI principles — the Google-flavoured foundations.
Get started with Microsoft 365 Copilot https://learn.microsoft.com/en-us/training/paths/get-started-with-microsoft-365-copilot/	Course path · 2–3 h	Free Microsoft Learn path: Copilot basics across Word, Excel, Outlook and Teams, plus prompting best practice — essential if your workplace runs on 365.
3Blue1Brown: Neural networks series https://www.3blue1brown.com/topics/neural-networks	Videos · ~4 h	The best visual intuition anywhere for how neural networks, transformers and attention actually work. No mathematics prerequisites that matter.
Karpathy: Intro to Large Language Models https://www.youtube.com/watch?v=zjKBMFhNj_g	Video · 1 h	A single hour that upgrades your mental model permanently: training, scaling, jailbreaks and prompt injection, from an ex-OpenAI founding member.
Learn Prompting https://learnprompting.org/docs/introduction	Docs · 10–20 h	Comprehensive open-source prompting guide written for non-developers; strong prompt-security section. Free docs; paid videos optional.
One Useful Thing (Ethan Mollick) https://www.oneusefulthing.org/	Newsletter	Evidence-based, hype-resistant guidance on using frontier AI for real work, from the Wharton professor who studies exactly that. Free tier.

17.4 Stage 2 — the power user

Objective: engineer context deliberately, connect AI to your own tools and data (MCP, projects, grounded workspaces), use agents as a user, run models locally, and automate small repetitive workflows. This is where Sections 4 and 6 turn from reading into practice. Time: 20–30 hours, best spread across a month of daily use.

The core path: do Anthropic’s interactive prompt-engineering tutorial (the exercises, not just the reading); learn what MCP is from the official docs and connect one real data source to your assistant; ground a research workflow in NotebookLM or a Claude Project; build one small automation in Google AI Studio or Copilot Studio; and run one open-weight model locally with LM Studio or Ollama — nothing demystifies “the model” faster than one running on your own laptop. Karpathy’s deep-dive rounds out the theory.

Resource	Type · time	What it is / why it earns its place
Prompt Engineering Interactive Tutorial (Anthropic) https://github.com/anthropics/prompt-eng-interactive-tutorial	Interactive · 6–10 h	Nine hands-on chapters from basic structure to complex industry prompts, with exercises and an answer key. Needs a (cheap) API key.
Model Context Protocol docs https://modelcontextprotocol.io/	Docs · 3–8 h	The official spec and tutorials for the open standard that connects AI to tools and data (Section 6.5). Read the concepts even if you never build a server.
Intro to Model Context Protocol (Anthropic course) https://anthropic.skilljar.com/	Course · 2–4 h	Free structured MCP course in the Anthropic catalogue — alongside Intro to Agent Skills and Intro to Subagents, both natural next steps.
NotebookLM https://notebooklm.google/	Tool	Google’s source-grounded assistant: upload documents, get cited answers and audio overviews. The fastest way to feel what grounding does to quality (Section 4). Freemium.
Google AI Studio https://aistudio.google.com/	Tool · self-paced	Free browser workbench for prompting, comparing and building with Gemini, with a generous free API tier — the lowest-friction sandbox for structured experiments.
Microsoft Copilot Studio docs https://learn.microsoft.com/en-us/microsoft-copilot-studio/	Docs	Low-code agent building — knowledge sources, tools, MCP, agent flows — plus free “Agent in a Day” workshops. The Stage 2 route for Microsoft shops.
LM Studio docs https://lmstudio.ai/docs	Docs · 1–3 h	The friendliest desktop app for running open-weight models locally, chatting with documents offline, and serving a local API. Free for personal use.
Ollama docs https://docs.ollama.com/	Docs · 1–3 h	The command-line counterpart: pull and run open models locally with an OpenAI-compatible API — the standard base for local experiments (Section 2.4).
GitHub Learn (GitHub Skills) https://learn.github.com/skills	Interactive · 1–2 h each	Guided exercises in real repositories — including GitHub Copilot basics on the free tier. A gentle on-ramp to the coding track even for non-developers.
Prompt Engineering Guide (DAIR.AI) https://www.promptingguide.ai/	Docs · 5–15 h	Research-backed guide spanning prompting through agents, context engineering and RAG — the bridge reference between Stages 2 and 3.
Karpathy: Deep Dive into LLMs like ChatGPT https://www.youtube.com/watch?v=7xTGNLPyMI	Video · 3.5 h	The full training stack — pretraining, fine-tuning, RLHF, hallucination, tool use — for a general audience. The theory companion to this stage.
Simon Willison’s Weblog https://simonwillison.net/	Blog	The most consistently useful practitioner blog on LLMs, tooling and prompt-injection security; monthly state-of-the-field roundups. Free.

17.5 Stage 3 — the practitioner

Objective: build agents and applications, design retrieval (RAG), run evaluations, choose architectures, and code with AI assistance — Sections 6, 7 and 15 as a workshop rather than a map. This stage assumes basic willingness to run code; several of the courses teach the code as they go. Time: 40+ hours — this is a project, so pick a real thing to build and let it drive which resources you need.

The core path: choose one primary API (Claude, OpenAI, or Gemini) and work its documentation and cookbook against a small project of your own; take the Hugging Face Agents course for framework-independent agent literacy; read Anthropic's "Building Effective Agents" and evals essays before you architect anything; adopt one agentic coding tool properly (Claude Code, Cursor, or Copilot) using its official best-practice material; and if you want the foundations under everything, Karpathy's Zero to Hero or fast.ai will take you as deep as you care to go.

Resource	Type · time	What it is / why it earns its place
Claude Developer Platform docs https://platform.claude.com/docs/en/home	Docs	The Claude API end to end: tool use, structured outputs, prompt caching, evals, guardrails. Pair with the free Building with the Claude API course (anthropic.skilljar.com).
Claude Cookbooks https://github.com/anthropics/claude-cookbooks	Code · self-paced	Official runnable notebooks: RAG, tool use, sub-agents, evaluation, vision, agent patterns. Steal these instead of starting from scratch.
OpenAI platform docs https://platform.openai.com/docs	Docs	Responses API, Agents SDK, tools, evals, safety best practices — the OpenAI build track (now hosted at developers.openai.com).
OpenAI Cookbook https://cookbook.openai.com/	Code · self-paced	100+ open-source recipes — prompting guides, RAG, evals, agent improvement loops, running open models locally.
Gemini API docs https://ai.google.dev/gemini-api/docs	Docs	Multimodal calls, function calling, agents and the Live API, with a genuinely useful free tier for development.
Get started with AI apps and agents on Azure https://learn.microsoft.com/en-us/training/paths/get-started-ai-apps-agents/	Course path · ~6 h	Free Microsoft path on Foundry: generative AI, agents, and extraction services — the enterprise-Azure build route.
Hugging Face Agents Course https://huggingface.co/learn/agents-course	Course · 20–40 h	Free, certificated, framework-spanning course on building and deploying agents — the best vendor-neutral agent education available.
Hugging Face LLM Course https://huggingface.co/learn/llm-course	Course · 30–50 h	The definitive free course on transformers and the open-model ecosystem, from tokenisers to fine-tuning (see also the smol course for post-training).
DeepLearning.AI Short Courses https://www.deeplearning.ai/short-courses/	Courses · 1–3 h each	100+ free short courses with in-browser notebooks — including Anthropic-partnered MCP and computer-use courses, plus RAG, agents and evals. Ideal gap-fillers.
Anthropic Engineering blog https://www.anthropic.com/engineering	Essays	"Building Effective Agents", "Effective Context Engineering", "Demystifying Evals" — the essays this guide's Section 6 leans on. Read before architecting.
Claude Code docs + Claude Code in Action https://code.claude.com/docs	Docs + course · 3–5 h	Agentic coding done properly: CLAUDE.md memory, skills, hooks, subagents, CI automation. The free companion course is at anthropic.skilljar.com/claude-code-in-action .
Claude Code: Best Practices for Agentic Coding https://www.anthropic.com/engineering/claude-code-best-practices	Essay · 45 min	The canonical essay on effective agentic-coding workflow — most of it transfers to Cursor and Copilot too.
Cursor docs https://cursor.com/docs	Docs	Agent mode, Rules, Skills, MCP servers and the CLI for the most popular agentic IDE (Section 7). Product freemium.
GitHub Copilot docs https://docs.github.com/en/copilot	Docs	Setup, chat, agents, code review and custom instructions; a free tier exists for individuals

		(github.com/features/copilot/plans).
Karpathy: Neural Networks — Zero to Hero https://www.youtube.com/playlist?list=PLAqhlrjKxbuWI23v9cThsA9GvCAUhRvKZ	Video course · ~15 h	Build GPT from scratch, from backpropagation up. Not required for practitioner work — but nothing builds deeper conviction about what these systems are.
fast.ai: Practical Deep Learning for Coders https://course.fast.ai/	Course · ~30 h	The famous code-first deep-learning course — the route into training and deploying your own models rather than only calling APIs.

17.6 Stage 4 — the specialist

Objective: go deep in one lens — governance, delivery/PM, or security — and lead how your organisation adopts AI. Stage 4 is less a course of study than a professional practice built on the anchor sections (8, 10 and 11); the resources below are the standards, frameworks and primary sources that practice runs on. The three lens callouts give the strategy; the tables and notes give the materials.

GOVERNANCE Going deep in governance

Start from Section 13: pick one framework (ISO/IEC 42001 or the NIST AI RMF) and learn it properly, then practise by building a real AI inventory and risk-tiering a handful of live use cases. Then add your regional layer: the NIST AI RMF playbook and state-law tracking (US), the EU AI Act explorer and harmonised standards (EU), or the National AI Centre's Guidance for AI Adoption and the Privacy Act automated-decision obligations (Australia). Governance fluency is as much about enabling adoption as constraining it.

Governance resource	Type · time	What it is / why it earns its place
NIST AI Risk Management Framework https://www.nist.gov/itl/ai-risk-management-framework	Framework · 4–10 h	The Govern/Map/Measure/Manage framework plus free Playbook and Generative AI Profile — the most practical free governance curriculum in existence.
ISO/IEC 42001 overview https://www.iso.org/standard/42001	Standard · 30 min (overview)	The AI management-system standard (Section 13.4). Overview and FAQ free; the full text is paid — learn the shape here, buy it when you commit.
EU AI Act Explorer (FLI) https://artificialintelligenceact.eu/	Docs · 2–10 h	Full Act text explorer, plain-language summaries, compliance checker and implementation timeline — the working reference for Section 13.3. See also its biweekly newsletter (artificialintelligenceact.substack.com).
National AI Centre (Australia) https://www.industry.gov.au/science-technology-and-innovation/technology/national-artificial-intelligence-centre	Gov guidance · 2–6 h	Guidance for AI Adoption, the AI screening and register tools, and the Responsible AI Network — the Australian starting point, about to matter more as the announced mandatory standards arrive (Section 13.6).
OAIC: privacy and commercial AI products https://www.oaic.gov.au/privacy/privacy-guidance-for-organisations-and-government-agencies/guidance-on-privacy-and-the-use-of-commercially-available-ai-products	Gov guidance · 1–2 h	The Australian privacy regulator's guidance on deploying commercial AI under the Privacy Act — required reading before the December 2026 automated-decision obligations.
International AI Safety Report https://internationalaisafetyreport.org/	Report · 3–20 h	The Bengio-led, 30+-country scientific review of general-purpose AI capability and risk; the 20-page policymaker summary is the efficient read.
Anthropic system cards https://www.anthropic.com/system-cards	Docs · 1–2 h each	Model and system cards documenting capability and safety evaluations — the primary-source habit every governance lead should build (all major labs publish equivalents).

SECURITY Going deep in security

Start from Section 14: internalise the OWASP LLM Top 10 and MITRE ATLAS, then get hands-on — red-team a system you control against prompt injection, and practise blast-radius containment (least privilege, human approval, sandboxing). Treat AI security as classic security discipline applied to a new, non-deterministic substrate.

Security resource	Type · time	What it is / why it earns its place
OWASP GenAI Security Project https://genai.owasp.org/	Framework · 3–10 h	The LLM Top 10, the Agentic Top 10, cheat sheets and governance checklists — the spine of Section 14, all free.
MITRE ATLAS https://atlas.mitre.org/	Framework · 2–6 h	The ATT&CK-style knowledge base of adversary tactics against AI systems, with real case studies — how attacks actually chain.
Microsoft AI Red Team hub + PyRIT https://learn.microsoft.com/en-us/security/ai-red-team/	Docs + tooling	Threat-modelling guidance, the failure-mode taxonomy, “lessons from red-teaming 100 GenAI products”, and PyRIT for hands-on adversarial testing.
Challenges in Red Teaming AI Systems (Anthropic) https://www.anthropic.com/news/challenges-in-red-teaming-ai-systems	Article · 30 min	A practical taxonomy of red-teaming approaches — expert, automated, multimodal, crowdsourced — and what each is for.
Import AI (Jack Clark) https://importai.substack.com/	Newsletter	Weekly research-grounded analysis from an Anthropic co-founder — the strongest free signal on frontier capability and risk for security and policy readers.

DELIVERY / PM Going deep in delivery and PM

Start from Section 9: learn to run AI initiatives as products — use-case selection, the pilot-to-production discipline, evals as gates, change management, and honest TCO. The differentiating skill is measuring value and reliability rather than activity, and knowing why most pilots stall so yours doesn't.

Delivery/PM has fewer formal free resources because the discipline is the general one — product and delivery management — applied to a probabilistic technology. Anchor on Section 9, then: One Useful Thing (above) for the evidence base on AI and work; The Batch (deeplearning.ai/the-batch) for a weekly executive-level read of what changed; the vendor case-study and solution libraries (Anthropic, OpenAI, Microsoft, Google all publish them free) for how deployments actually get structured; and the LLMOps and evaluation material in the Stage 3 table — evals-as-gates is the single most PM-relevant technical skill in this guide.

17.7 Platform tracks at a glance

The stage modules above are organised by capability; this table re-sorts the same resources by vendor, for when you need to go deep on one ecosystem. Start-here entries are Stage 1–2; go-deeper entries are Stage 2–3.

Track	Start here	Go deeper
Anthropic / Claude	AI Fluency course; Anthropic Academy; Claude prompt-engineering guide (14.3)	Interactive tutorial; MCP docs and course; platform docs; Claude Cookbooks; engineering blog; Claude Code (14.4–14.5)
OpenAI	OpenAI Academy; Help Center (14.3)	Platform docs; OpenAI Cookbook; Academy bootcamps on RAG, evals and Codex (14.5)
Google / Gemini	Google AI Essentials; Intro to Generative AI path (14.3)	Google AI Studio; NotebookLM; Gemini API docs; Kaggle Learn micro-courses (kaggle.com/learn) (14.4–14.5)
Microsoft	M365 Copilot Learn path (14.3)	Copilot Studio docs; Azure AI apps & agents path; AI-901 Fundamentals prep; AI Red

		Team hub (14.4–14.6)
Open-weight / local	LM Studio (14.4)	Ollama; Hugging Face Learn — LLM, Agents and smol courses (14.4–14.5)
Coding tools	GitHub Learn Copilot exercises; Copilot free tier (14.4)	Claude Code docs, course and best-practices essay; Cursor docs; Copilot docs (14.5)

17.8 The skills that endure

As models improve, the value of raw “knowing how to prompt” declines and the value of human judgement rises. The skills worth investing in because they compound rather than depreciate:

- **Critical verification.** Knowing when to trust output and when to check it — the single most durable AI skill (Section 15).
- **Problem framing.** Turning a vague goal into a precise, well-scoped task is most of the work; the model handles the rest.
- **Evaluation and taste.** Being able to tell good output from plausible-but-wrong output, and to measure it rather than eyeball it.
- **Domain expertise.** AI amplifies what you already know; deep knowledge of your field is what lets you direct and check it.
- **Governance and security instinct.** The reflex to ask “what could go wrong, who’s accountable, what data is exposed?” before shipping (Sections 13 and 14).
- **Adaptability.** Comfort re-learning tools as they churn, anchored by fundamentals that don’t.

17.9 Keeping current without drowning

The volume of AI news is unmanageable, and most of it is noise. The goal is a sustainable filter, not total coverage. A layered approach that works:

- **Go to the source.** Official product and API documentation, model cards, and the research blogs of the major labs tell you how systems actually behave and change — more reliably than second-hand coverage.
- **Anchor on authoritative standards.** For the durable frame, the NIST AI RMF, ISO/IEC 42001, the OWASP GenAI/LLM Top 10 and MITRE ATLAS (security), the EU AI Act’s official resources, and the annual International AI Safety Report change slowly and reward close reading.
- **Skim primary research.** arXiv and Papers with Code carry the real work; you don’t need to read papers in full — abstracts and figures go a long way.
- **Pick a few trusted synthesisers, not twenty.** A small number of good analysts or newsletters filter the firehose — but keep the Sections 15–16 scepticism handy for benchmark headlines and vendor claims.
- **Learn by using the new thing.** The fastest way to understand a capability is to point it at a real task of your own the week it ships.

Named, and link-checked with the rest of this section — a small directory that covers the firehose without drowning you:

Source	Type	Why this one
--------	------	--------------

AI Explained https://www.youtube.com/@aiexplained-official	YouTube	The most rigorous mainstream AI-news channel — benchmark-literate (its author built SimpleBench), allergic to hype.
Andrej Karpathy https://www.youtube.com/@AndrejKarpathy	YouTube	Rare but essential uploads — the reference channel for how LLMs actually work.
3Blue1Brown https://www.youtube.com/@3blue1brown	YouTube	Visual mathematics, including the neural-network series from Stage 1 — for when a concept refuses to click as text.
Anthropic (official) https://www.youtube.com/@anthropic-ai	YouTube	First-party demos, developer talks and research explainers; OpenAI's equivalent is youtube.com/@OpenAI .
Simon Willison's Weblog https://simonwillison.net/	Blog	Daily practitioner signal on models, tools and security; the monthly roundups alone would justify the subscription — and it's free.
One Useful Thing https://www.oneusefulthing.org/	Newsletter	AI and work, evidence first — the keep-current pick for Stage 1–2 readers.
The Batch (DeepLearning.AI) https://www.deeplearning.ai/the-batch/	Newsletter	Andrew Ng's weekly — balanced, broad, and readable in ten minutes.
Latent Space https://www.latent.space/	Newsletter + podcast	The AI-engineer's newsletter and podcast — how leading teams actually build with this technology. The Stage 3 keep-current pick.
Import AI https://importai.substack.com/	Newsletter	Research and policy signal, weekly, from Anthropic co-founder Jack Clark. The Stage 4 keep-current pick.
arXiv (cs.AI / cs.CL) https://arxiv.org/list/cs.AI/recent	Preprints	The primary source itself — abstracts and figures go a long way; you almost never need the full paper.

17.10 Start here: your first thirty days

To turn this guide into practice, a concrete starting plan beats good intentions:

- Pick one real, recurring task you already do — ideally a document- or knowledge-heavy one — and rework it using the prompting and document techniques in Section 4.
- Build a verification habit around it: decide, up front, what you'll check before trusting the output (Section 15). Make it routine, not occasional.
- Ground one workflow — feed the model your own source material (documents, data, or a connected tool) rather than relying on its memory — and notice the quality jump (Sections 4 and 6).
- Choose your lens — governance, delivery, or security — read that anchor section closely, and align to one recognised standard.
- Set a keep-current cadence you'll actually sustain: one source at the source, one trusted synthesiser, and a standing habit of trying each new capability on a real task.
- **Enrol in one course matched to your stage** — from the modules above: AI Fluency (Stage 1) and the MCP introduction (Stage 2) are the highest-leverage starts, and both are free.

17.11 A final word

If a single thread runs through all seventeen sections, it's this: AI is genuinely, unevenly powerful — remarkable at some things, unreliable at others — and using it well is less about the technology than about the judgement you bring to it. Ground your work in real sources. Verify what matters. Keep a human accountable. Ask early what could go wrong, who owns

the risk, and what data is exposed. Match the tool to the task, and treat your own calibrated judgement as the asset that appreciates while the models churn.

This guide is deliberately dated and deliberately living. Much of its fast-moving detail — the specific models, prices, and rankings — will be out of date within months; that's a feature of an honest snapshot, not a flaw. The fundamentals, the frameworks, and the habits are built to last, and they're what will carry you as the surface keeps changing. It was written and built with Claude, an AI assistant made by Anthropic — which means, fittingly, that it was produced with the very tools and subject to the very limitations it describes, and so was dated, sourced, and cross-checked accordingly.