

E

Glossary

These are the terms used in this book, described the way the chapters use them. Where a term has a broader technical meaning, I've kept to what a leader needs in order to follow the decision. The chapter in parentheses is where the term is developed.

Agent. A product description suggesting that software can take actions, such as reading or changing records, sending messages, or triggering another process. Ask the vendor to demonstrate what those actions are and where a person's approval is required. (Chapter 2)

AI governance. The responsibilities and processes an organization uses to decide how AI may be used, support its implementation, and respond to what happens during use. (Chapter 8)

AI use inventory. A maintained record of AI uses in an organization, including purchased products, internally developed tools, capabilities embedded in existing software, and uses staff have begun on their own. Recording a use in the inventory does not approve it. (Chapter 8)

Business associate. Under HIPAA, a party that performs certain functions or services for a covered entity involving protected health information. Qualified reviewers should determine how the rules apply to a particular arrangement. (Chapter 10)

Confabulation. The term NIST's Generative AI Profile uses for generative AI producing inaccurate material in confident, readable language, also commonly called hallucination. (Chapter 2)

Context Cycle. The working sequence this book uses to apply the Context Principle: Observe, Context, Interpret, Decide, Act, and Reflect/Refine, then observe again. A team can return to an earlier step when new information changes its understanding. (Chapter 4)

Context Principle. My framework for decisions: understanding must precede intervention, in proportion to the potential impact of the decision. This book applies it to AI in healthcare. (Chapters 1 and 4)

Covered entity. Under HIPAA, a health plan, a healthcare clearinghouse, or a provider conducting covered electronic transactions. (Chapter 10)

Decision record. A short written account of a decision: what is permitted, the stakes, what the team established, the cost of delay, who is responsible, and when the decision returns for review. (Chapters 1 and 6)

Decision Triangle. Three considerations for judging how much understanding a decision needs: the cost of being wrong, the difficulty of reversing the harm, and the scope of human impact. A serious concern in one area deserves attention even when the other two appear manageable. (Chapter 6)

Drift. Changes in a model's input data, or in the relationships a model relies on, after it was developed or evaluated. Drift can affect performance, but it does not necessarily produce a slow, predictable decline. (Chapter 2)

External validation. Evaluating a prediction model's performance on new participant data, which might come from a later period, another organization, or a different population. Evidence of predictive performance does not, by itself, establish that using the tool improves care. (Chapters 2 and 9)

Fine-tuning. Additional training that adapts an existing model. Giving a tool a document to consult is not, by itself, the same as fine-tuning it. (Chapter 2)

Four Lenses. Four perspectives for developing understanding: Evidence, Context, Expertise, and Lived Experience. They overlap in practice. (Chapter 5)

Generative model. A model that produces content, such as text, images, or audio. (Chapter 2)

Illusion of understanding. The risk that a well-written answer is mistaken for evidence that the underlying problem has been understood. (Chapter 1)

Inference. Using a model to produce an output, as distinct from training it. Using a model does not necessarily retrain it. (Chapter 2)

Intervention. A change we introduce: a tool, a policy, a revised process, or a different assignment of responsibility. An AI feature added to existing software is an intervention even when nobody requested it and no purchase is involved. (Chapter 4)

Large language model (LLM). A type of model used for language tasks, including generating text. (Chapter 2)

Model and product. The model is the underlying system that produces an output. The product is what your organization uses, which may combine a model with instructions, document retrieval, access controls, and connections to other applications. Evidence about the model alone won't tell you how the whole product performs. (Chapters 2 and 9)

Pilot. A label for a limited trial. Calling something a pilot doesn't make its consequences small; the information involved, who sees the output, and what someone might do because of it determine the risk. (Chapters 3 and 13)

Plan-Do-Study-Act (PDSA). An improvement method the Institute for Healthcare Improvement describes for planning a test, carrying it out, studying the results, and using the learning to guide the next test. (Chapters 4 and 13)

Predictive model. A model that estimates an outcome or assigns a score from the information available to it. (Chapter 2)

Premortem. A discussion in which a team imagines that a project has failed, identifies possible reasons, and plans how to address them. (Chapter 6)

Proportional Understanding. Matching our preparation to the possible consequences of a decision. (Chapter 6)

Proxy. A measure used in place of the outcome we actually care about, such as healthcare spending used in place of health need. A proxy can be useful for one decision and unsuitable for another. (Chapters 1 and 11)

Psychological safety. In Amy Edmondson's research, whether team members feel able to take interpersonal risks, such as raising a question or admitting a mistake. (Chapter 12)

Retrieval-augmented generation (RAG). An approach that combines retrieval of source material with generation of an answer. A product might retrieve passages from an organization's policies and make them available to the language model as it responds. (Chapter 2)

Review path. A route through governance matched to the proposed use, such as a delegated review for a well-contained test or deeper specialist review for a use with more consequential effects. (Chapter 8)

Sufficient Understanding Test. Five questions for judging whether understanding is sufficient for a proposed action: lens coverage and the weakest area, proportional depth, the uncertainties that matter, meaningful input from affected people, and an investigated plausible failure. Completing it does not authorize a use. (Chapters 1 and 6)

Training. Adjusting a model's internal parameters using data. (Chapter 2)

Understanding. Making sense of information in the setting where we intend to act, including how the work actually happens and what people are expected to do with the result. (Chapter 4)

Understanding Debt. The costs that can accumulate when changes are repeatedly made without enough understanding for the stakes. Later, someone has to resolve those questions while also doing the work. (Chapter 7)

Whole-organization approach. Following the consequences of a change far enough to understand who is affected and what they need. It scales to the decision; it doesn't require every department to approve every experiment. (Chapter 2)

Workaround. An unofficial practice, such as a separate tracker or a sequence of phone calls, that compensates for something the formal process doesn't provide. Understanding what it does comes before removing it. (Chapter 7)