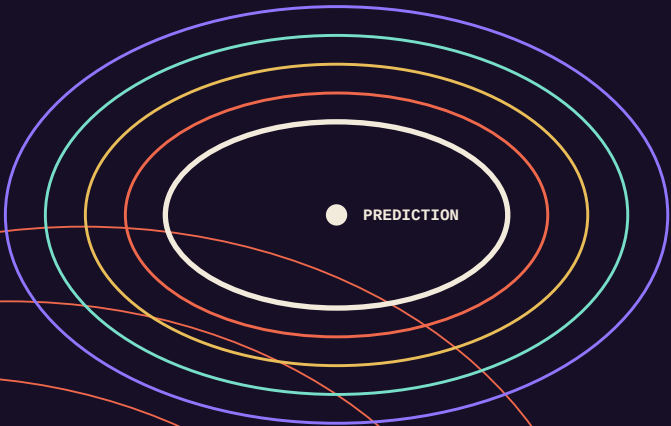


AN ADVANCED-STUDENT BOOK

# We Used to Find the *Pattern First*

How artificial intelligence changes prediction, discovery,  
and what humans can know



● PREDICTION

JOHN RECTOR / EXPANDED EDITION / 2026

# We Used to Find the Pattern First

**How Artificial Intelligence Changes Prediction,  
Discovery, and What Humans Can Know**

Expanded print edition

**John Rector**

First expanded edition, 2026

Copyright © 2026 John Rector. All rights reserved.

This edition expands the original web publication into a longer advanced-student text with additional examples, field notes, mathematical context, and evaluation guidance.

Designed as a  $6.5 \times 9.5$  inch reading edition.

# Contents

The machine after the map . . . . . 6

**Beneath creation . . . . . 8**

    The prediction underneath the artifact . . . . . 9

    What a pattern is . . . . . 13

    Intelligence and compression . . . . . 17

    Compression revealed through time . . . . . 21

**The reversal . . . . . 25**

    The old order . . . . . 26

    The inversion . . . . . 30

**The representation . . . . . 34**

    The geometry of intelligence . . . . . 35

    Dimensions are not parameters . . . . . 39

    Random to whom? . . . . . 43

**The epistemic machine . . . . . 47**

    Useful before understandable . . . . . 48

    Computation before understanding . . . . . 52

    From automation to discovery . . . . . 56

    When noise becomes information . . . . . 60

    The revised scientific method . . . . . 64

    What AI still cannot predict . . . . . 68

    The boundary of the knowable . . . . . 74

    We no longer have to arrive first . . . . . 78

    A compact mathematical field guide . . . . . 80

Sixteen questions for any predictive system . . . . . 82

Selected sources . . . . . 83

PROLOGUE

# The machine after the map

**Every information technology arrives with a public description that is true and inadequate.**

The printing press made copies. The telegraph sent messages. The telephone carried voices. The database stored records. The computer calculated. The internet connected machines. Each sentence is accurate. None of them explains the reorganization that followed.

Artificial intelligence predicts.

That sentence is also accurate and inadequate until we understand what prediction now means. A weather forecast is a prediction. So is a credit score. So is the next token in a sentence. But modern machine learning makes prediction inside representations humans did not fully design, using relationships humans did not first identify, across quantities of experience no person could absorb. Prediction is no longer merely the last step of a human theory. It can be the instrument that tells us a theory is waiting to be found.

This book is about that break.

The old machine received a map. A human named the important objects, identified the roads, selected the destination, and translated the route into instructions. The machine followed the map at enormous speed.

The new machine can participate in making the map.

It still receives choices from us. We choose the data, architecture, objective, training process, evaluation, permissions, and point of use. We decide what counts as an error. We decide what consequences matter. Nothing in this book transfers responsibility to a model. The narrower and more radical claim is that the complete operative pattern does not have to arrive from a human mind.

This is why the familiar comparisons fail. AI is not a new database. A database preserves what happened. AI estimates what follows. AI is not merely automation. Automation executes a known transformation. Machine learning can search for a transformation. AI is not merely an interface. Natural language

makes the capability approachable, but conversation is the doorway, not the engine.

The engine is a cycle: experience, representation, prediction, error, adjustment.

We will begin beneath the artifact. We will ask why an output that looks like creation is still prediction. We will move from pattern to compression, and from compression to probabilistic code length. We will examine the reversal from explicit programming to learned structure. We will separate dimensions, parameters, and activations. We will use XOR to show how randomness can belong to the observer. Then we will follow the consequences into science, discovery, explanation, and the limits of what any predictor can know.

The advanced student should resist two temptations.

The first is mystification. A model that uses relationships we cannot summarize is not supernatural. It is a mathematical system shaped by data, objectives, and optimization. Its internal scale can exceed our inspection without escaping causality.

The second is reduction. Saying that a model predicts does not make the result small. Prediction can accumulate into a proof, a design, a diagnosis, or a strategy. A city remains a city even though it is built one decision at a time.

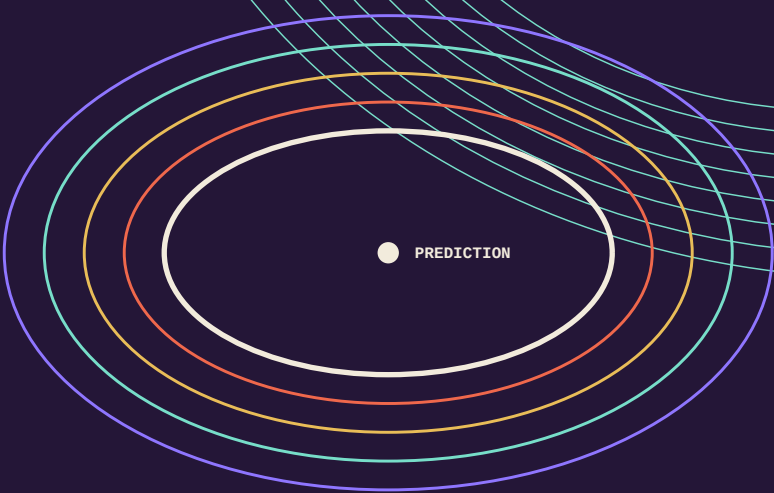
The goal is precision without deflation.

Something genuinely new has happened. Information technology can now produce useful behavior from patterns that humans did not first have to discover and formalize. Once that sentence becomes clear, the rest of the AI landscape reorganizes around it.

PART I

# Beneath creation

The artifact is what the human sees. Prediction is what the machine does.



● PREDICTION



## CHAPTER 01

# The prediction underneath the artifact

**Artificial intelligence does not begin by creating a book, a plan, an image, a diagnosis, a strategy, or a line of software. It begins by estimating what should come next.**

We name the output by its visible form. A prediction that continues for sixty thousand words becomes a book. A prediction arranged around a business problem becomes a report. A prediction expressed as pixels becomes an image. A prediction written in Python becomes software. A prediction ordered into future actions becomes a plan.

The forms are different. The primitive is the same.

A language model estimates the probability of possible next tokens given the tokens already present. An image model estimates visual structure under a conditioning signal. A classifier estimates which label is most probable given an input. A forecasting system estimates a future value given prior observations. Even when the system produces a long chain, each local move is conditioned by the state created by the moves before it.

This does not make the output trivial. A cathedral is made one stone at a time. Its local construction does not erase its global form. In the same way, next-step prediction can accumulate into an argument, a program, an image, or an investigation. The important correction is not that the artifact is unreal. The correction is that creation is the external appearance of predictive computation.

**I've got the pattern. I'll create what comes next.**

That sentence places the burden where it belongs. The machine cannot predict from nothing. It needs structure that connects what has happened to what is likely to happen. The book, plan, image, or solution is the visible release of that structure.

Ask what an AI creates and you will receive a catalog of artifacts. Ask what operation makes those artifacts possible and the catalog collapses into one word: prediction.

The next question is therefore unavoidable.

What makes prediction possible?

## The artifact is a disguise

**The word prediction becomes difficult because ordinary language gives it a narrow time horizon. We imagine a statement about tomorrow: rain, revenue, an election, a patient's risk. A generative model appears to do something else. It produces a paragraph, an illustration, a program, or a song. The result seems present, not future.**

But "next" does not have to mean next Tuesday. It can mean the next token, the next patch of image structure, the next state in a trajectory, the next action in a plan. Generation is prediction repeated under changing context. Each output becomes part of the condition for the output that follows.

This autoregressive structure explains both coherence and fragility. A strong early choice narrows the later possibilities in a useful direction. A weak early choice can become context that the model must continue. The system is not retrieving one finished book from a hidden shelf. It is constructing a probability-guided path through a space of possible continuations.

For a language model, the central object is a conditional distribution:

$$P(x_t \mid x_1, x_2, \dots, x_{(t-1)})$$

The model does not return every probability to the user. A decoding procedure turns the distribution into an actual continuation. Greedy decoding selects the highest-probability option. Sampling allows lower-probability options to appear according to their assigned mass. Temperature and other controls reshape that selection. The learned model and the decoding policy are therefore different parts of the result.

This distinction matters when people say, "The same prompt gave me two answers." A probability distribution can be stable while samples differ. Variation does not prove the absence of structure. It reveals that the structure permits several plausible continuations.

Long outputs also depend on global scaffolding. A model can condition on an outline, a system instruction, retrieved sources, tool results, and intermediate work. These additions alter the pattern available at each step. What looks like one act of creation is often a sequence of predictions inside a designed environment.

The environment can be more consequential than the base model. Give a model verified records, a calculator, a code runner, and a review loop, and the pattern includes external reality checks. Give it only a vague prompt, and it must rely more heavily on broad regularities from training. Context does not simply add facts. It changes the conditional problem.

The practical question is therefore not "Can the model create?" That question is settled by the artifact in front of us. The stronger questions are: What is being conditioned on? Which probability distribution has been learned? How is a continuation selected? Where does feedback enter? What corrects a locally plausible move before it becomes a global failure?

Calling the operation prediction is not an attempt to make it ordinary. It is an attempt to name the mechanism precisely enough that we can improve it.

## Mistake to avoid: confusing probability with belief

**The numbers inside a model are often described as beliefs. The metaphor can help, but it also imports a human subject that may not exist. A conditional probability is a numerical relation produced by a trained function. It tells us how the model distributes expectation under its representation. It does not establish consciousness, commitment, or an inner point of view.**

The distinction keeps the analysis clean. We can study calibration, uncertainty, confidence, and error without pretending that the model experiences doubt.

## CHAPTER 02

# What a pattern is

## **A pattern is structure that reduces uncertainty.**

Suppose a light flashes red, blue, red, blue, red, blue. Before seeing the sequence, the next color could have been either. After observing the alternation, the next color is no longer equally uncertain. The history contains structure. The structure changes the probability of the next event.

That is enough for a working definition. A pattern need not be decorative. It need not repeat visibly. It need not fit into a sentence. It need not be simple enough for a human to notice. It must only make some outcomes more expected than others.

If every possible future remains exactly as likely after observing the past, the past supplied no predictive pattern. If the probabilities move, structure was present.

$$P(Y | X) \neq P(Y)$$

Patterns come in many forms. A repeated sequence is a pattern. A statistical association is a pattern. An invariance under rotation is a pattern. A grammar is a pattern. A causal mechanism creates patterns, but causation is not required for prediction. A proxy can predict an outcome without producing it.

This last distinction matters. Prediction asks whether information reduces uncertainty about an outcome. Explanation asks why the outcome occurs. A model can succeed at the first while remaining weak at the second.

## The observer is part of the claim

**Pattern is never merely a property of data in isolation. It is a relation among data, representation, objective, and observer. The alternating lights are easy to represent because time gives them a natural order. A relationship spread across five hundred clinical variables may be just as real and far harder for a human to detect.**

The phrase "there is no pattern" is therefore stronger than it sounds. It may mean there is no structure. It may mean we chose the wrong variables. It may mean we represented the variables badly. It may mean the structure is too weak for the available sample. It may mean the process changed. Or it may mean the observer lacks the capacity to find what is present.

No pattern, no prediction. But failure to predict does not prove the absence of pattern. It may reveal the boundary of the predictor.

## Pattern depends on a question

**A sequence can contain many structures at once. Which one counts as the pattern depends on what we are trying to predict.**

Take a record of daily temperatures. For tomorrow's temperature, recent days, season, latitude, and atmospheric state may matter. For electricity demand, the same temperatures interact with hour, building type, occupancy, and price. For crop risk, duration above a threshold may matter more than the daily average. The data did not change. The target changed which regularities became useful.

This is why "find the pattern" is incomplete. A learner needs an objective or a task. Structure is abundant. Relevance is selective.

The conditional statement  $P(Y|X)$  differs from  $P(Y)$  gives us a minimal test:  $X$  carries information about  $Y$  when observing  $X$  changes the distribution over  $Y$ .

Yet practical learning requires more. The change must be large enough to estimate, stable enough to survive new data, and available at the moment of prediction.

Availability is frequently overlooked. A hospital record may contain a code entered after a diagnosis was confirmed. That code predicts the diagnosis perfectly in historical data and is useless for early detection. A sales system may record the refund after the customer leaves. A model trained carelessly can exploit the future while appearing to predict it.

The pattern is real inside the dataset and invalid inside the decision.

Temporal discipline is therefore part of pattern discovery. For every feature, ask: when did this information exist? Who created it? What process placed it in the record? Would it exist in the same form when the model is used?

Patterns also inherit the measurement system. "Customer satisfaction" may be a survey response from the minority who answered. "Employee performance" may be a manager rating shaped by opportunity and bias. "Crime" may mean recorded incidents, which depend on reporting and enforcement. A model can accurately learn the pattern of a measurement without learning the full phenomenon named by the measurement.

This does not make the model worthless. It makes the target specific. Advanced work begins when names stop floating free of collection procedures.

## Signal, sample, and stability

**Suppose a relationship is weak but real. With too little data, it will be indistinguishable from sampling noise. Suppose a relationship is strong but limited to one setting. With data from only that setting, it will look universal. Suppose a relationship held for ten years and a policy changed yesterday. Historical volume will not restore current validity.**

Three questions should therefore accompany any pattern claim:

- Is the signal identifiable in the available sample?

- Is the sample representative of the intended use?
- Is the generating process stable enough for the relationship to persist?

Machine learning expands pattern search, but it does not repeal these questions. In fact, model flexibility makes them more important. A system with enough capacity can discover a regularity in almost any finite dataset. Generalization is the test that separates a fitted accident from a useful pattern.

## Pattern is relational

**It is tempting to imagine pattern as an object hidden inside data, waiting to be uncovered. A better view is relational. Pattern arises among a phenomenon, measurements, representation, objective, learner, and use environment.**

Change the sensor and a new relationship becomes visible. Change the representation and a nonlinear boundary becomes simple. Change the target and a different compression matters. Change the environment and yesterday's regularity may disappear.

The world supplies constraints. The observer supplies a way of asking.



## CHAPTER 03

# Intelligence and compression

**If a pattern exists, the data can be described more economically than by listing every observation independently.**

Consider a thousand-character string made by repeating 10 five hundred times. The literal string is long. Its rule is short: "repeat 10 five hundred times." The shorter description is possible because the sequence contains regularity.

Now consider a thousand truly random bits. There may be no description meaningfully shorter than the bits themselves. The sequence gives us nothing to factor out. In the language of algorithmic information theory, an incompressible string is random relative to the descriptive machinery being used.

This is why the claim "intelligence is compression" has real depth and also needs discipline. Compression is not merely shrinking a file. It is finding reusable structure. A successful compressor says: these observations are not isolated; this shorter model can regenerate or predict them.

Kolmogorov complexity expresses the idealized version. For a string  $x$ , its complexity is the length of the shortest program that outputs  $x$  on a fixed universal computer.

$$K(x) = \min$$

A short program reveals strong regularity. A long irreducible program reveals little exploitable structure. The fixed choice of universal machine matters by a constant, but not enough to destroy the central idea for sufficiently long descriptions.

Yet compression alone is not a complete definition of intelligence. A specialized compressor can be excellent on one distribution and useless elsewhere. Intelligence must also involve transfer: the discovered structure must support

accurate prediction or successful action across situations that were not copied from the past.

A disciplined working definition is stronger:

**Intelligence is the capacity to discover structure in experience that improves prediction and guides successful action across changing situations.**

Compression is evidence that structure was found. Prediction tests whether that structure reaches beyond the observations already seen. Action tests whether the prediction matters in an environment.

## Compression is a wager about recurrence

**To compress is to bet that some structure will recur.**

A dictionary compressor replaces repeated strings with shorter references. An image codec uses regularities in nearby pixels and perceptual sensitivity. A language model assigns short effective codes to likely continuations. Every compressor embodies expectations about what kind of data it will encounter.

This is why no compressor wins everywhere. A method designed for photographs will not necessarily compress executable code well. A model trained on legal prose may assign poor probabilities to protein sequences. Compression performance is always relative to a distribution.

Kolmogorov complexity reaches beyond specific distributions by asking for the shortest program that produces an object. The definition is universal up to an additive constant across universal description languages, but the shortest program cannot be found in general. There is no algorithm that always certifies that no shorter program exists.

That incomputability is not a footnote. It separates an ideal measure of structure from every practical compressor. Real systems search a restricted family with finite compute. They find shorter descriptions, not necessarily the shortest description.

The practical lesson is not despair. It is humility. A failed compressor may reveal the limits of the compressor rather than the randomness of the data.

Minimum Description Length turns the compression idea into a model-selection discipline. A good explanation should compress both the model and the residual errors:

$$\textit{Total description} = \textit{description of model} + \textit{description of data given model}$$

A model that memorizes every example makes the residual tiny but the model description enormous. A model that is too simple has a short description and leaves large residual error. The best balance captures repeatable structure without encoding every accident.

This is Occam's razor with an accounting system. Simplicity is not admired for aesthetic purity. It is valued because a model that compresses the data economically is less likely to be a disguised list of exceptions.

## Lossy and lossless understanding

**Lossless compression preserves every bit. Prediction and intelligence often depend on lossy compression: preserving what matters for a task while discarding irrelevant detail.**

A subway map is not a photograph of a city. It deletes building shapes, precise distances, trees, weather, and elevation. That loss is the map's intelligence. By discarding detail, it makes connectivity legible.

Learned representations do something similar. A classifier may become invariant to small changes in lighting because lighting does not change the object class. A speech system may discard speaker-specific variation to preserve linguistic content. A planning system may compress a scene into objects, affordances, and likely transitions.

But every loss creates a boundary. Information discarded as irrelevant to one task may be essential to another. A representation optimized for average accuracy may erase rare subgroups. A compressed customer profile may omit the circumstance that makes this customer different.

The central design question is not "How much can we compress?" It is "Which distinctions must survive?"

## Intelligence is not a zip file

**The slogan "intelligence is compression" earns its place when compression means discovering transferable regularity. It fails when compression is treated as a complete account of agency, embodiment, value, or action.**

An intelligent system must not only describe experience economically. It must use its description under novelty. It must update when the environment changes. It must distinguish uncertainty from confidence. And if it acts, it must connect prediction to goals without destroying the conditions that made the prediction reliable.

Compression supplies the world model. Intelligence appears in how that model survives contact with what comes next.

## CHAPTER 04

# Compression revealed through time

**Prediction and compression are not neighboring metaphors. For probabilistic sequences, they are operationally convertible.**

If a model assigns high probability to the event that actually occurs, that event requires few bits to encode under an ideal code. If the model assigns low probability to the event, the event is surprising and requires more bits.

$$I(x) = -\log$$

For a sequence, the code length accumulates one conditional prediction at a time:

$$L(x$$

Arithmetic coding turns those probabilities into a compressed bitstream. Run the relationship in the other direction and a compressor becomes evidence of a predictive model. The 2023 paper "Language Modeling Is Compression" demonstrates the equivalence across text and other data types: powerful predictive models can act as general-purpose compressors, and compressors can support conditional generation.

This gives the book its central mathematical hinge:

**Prediction is compression revealed through time.**

The compressor says: "I have found enough structure that I do not need to transmit every observation independently." The predictor says: "I have found enough structure that I can distribute probability over what comes next." Both are reports about uncertainty reduced by a model.

Claude Shannon's information theory supplied the quantitative language. Ray Solomonoff pushed the idea toward universal induction: prefer computable explanations according to their descriptive length, update them against observations, and let the weighted explanations predict the continuation. The ideal is mathematically profound and generally incomputable. Its importance is not that modern neural networks secretly implement Solomonoff induction. They do not. Its importance is that simplicity, probability, compression, and prediction can be joined inside one formal frame.

Once that connection is visible, a model no longer looks like a bag of answers. It looks like compressed predictive structure.

## Why surprise costs bits

*The equation  $I(x) = -\log_2 p(x)$  turns expectation into length.*

If an event has probability one half, identifying it among two equally likely possibilities requires one bit. If it has probability one eighth, identifying it among eight equally likely possibilities requires three bits. Less expected events require longer descriptions because the code must distinguish them from a larger field of alternatives.

The logarithm matters because independent surprises add. If two independent events have probabilities  $p$  and  $q$ , their joint probability is  $pq$ . The negative logarithm converts multiplication into addition:

$$-\log(pq) = -\log(p) - \log(q)$$

This gives information theory a clean accumulation rule. A sequence's total surprise is the sum of local surprises assigned by the predictive model.

Cross-entropy loss uses the same structure. During language-model training, the system is penalized according to the negative log probability assigned to the observed token. Confidently predicting the correct token produces little loss. Assigning the correct token tiny probability produces large loss. Gradient-based optimization adjusts parameters to reduce that expected surprise across the training distribution.

Perplexity is another expression of the same quantity. Roughly, it reports the effective number of equally plausible choices faced by the model. Lower perplexity means the model concentrates more probability on what actually occurs. It is not a complete measure of usefulness, truth, safety, or reasoning, but it is a direct measure of predictive fit under the evaluated sequence distribution.

## Arithmetic coding as the bridge

**Arithmetic coding makes the prediction-compression connection concrete. Instead of assigning a fixed codeword to each symbol, it repeatedly narrows an interval according to the model's conditional probabilities. Likely symbols receive larger subintervals. Unlikely symbols receive smaller ones. After a complete sequence, any number inside the final interval identifies the sequence. Better predictions leave a larger final interval and therefore require fewer bits to specify.**

The compressor depends on the predictor. The decoder runs the same predictive probabilities to reverse the process. If both sides share the model, the sequence can be reconstructed exactly.

This is why probabilistic modeling and lossless compression can be converted into each other. The equivalence is not a poetic resemblance. It is an engineering relationship.

## A warning about compressed knowledge

**Model parameters are sometimes described as a compressed copy of the training data. The phrase is useful if handled carefully. The model does compress regularities from training, but not necessarily as a conventional archive that can reproduce every record. Some examples may be memorized. Many are not. The learned weights embody statistical structure shaped by architecture, optimization, sampling, and capacity.**

The distinction matters for both capability and governance. A model can generate a plausible fact because the fact is supported by distributed regularity, because a near-verbatim example was memorized, or because a pattern produced a false continuation. The output alone may not reveal which path occurred.

Compression explains how finite parameters can support broad behavior. It does not guarantee faithful storage or factual retrieval.

## Prediction is a moving code

**In a static compressor, the code reflects a fixed model. In an interactive system, context changes the distribution continuously. A question, document, tool result, or correction reshapes what counts as likely next. The code is conditional on the conversation's present state.**

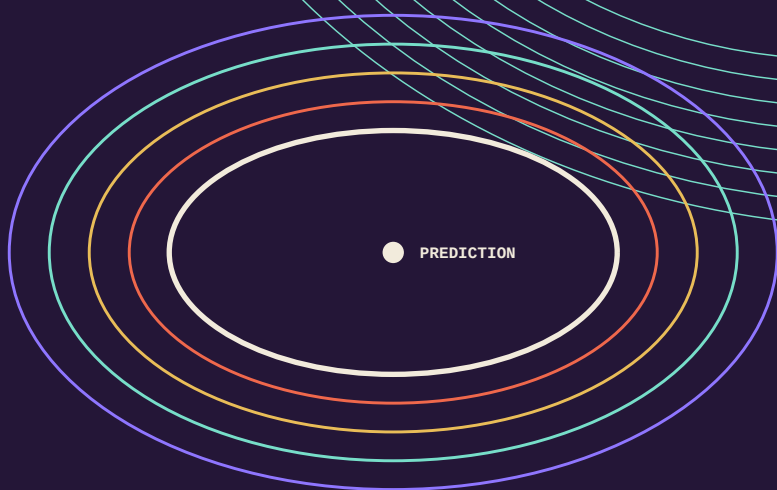
This is why good context can make a model appear transformed. The parameters remain the same. The active predictive problem has changed.



PART II

# The reversal

Traditional software inherited the result of human intelligence. Learning changes where useful structure can originate.



## CHAPTER 05

# The old order

**For most of the history of information technology, intelligence stood upstream of the machine.**

A human encountered a recurring situation. The human noticed a relationship. The human understood it well enough to state what should happen. A programmer converted that understanding into instructions. The computer executed those instructions quickly and reliably.

Consider inventory. "If stock falls below one hundred units, order five hundred more." The machine can monitor every item continuously and place the order in milliseconds. But the threshold, quantity, objective, and relationship were supplied by people. The computer did not discover replenishment. It inherited a formalized decision.

Accounting software did not discover double-entry bookkeeping. Payroll software did not discover wages. Tax software did not discover tax law. Navigation software did not discover geometry. Spreadsheets did not discover arithmetic. Databases did not discover the meaning of a customer record.

These systems transformed civilization. Their lack of learning does not diminish them. It identifies their architecture.

Traditional programming is the transfer of understood patterns into executable form. The human says: "I know the transformation. I will describe it precisely enough that the machine can repeat it."

This is why requirements were sacred. What should happen? Under what conditions? Which exceptions apply? How should the output be calculated? Someone had to know. Ambiguity upstream became failure downstream because the computer could only execute the formal structure it received.

The great constraint was not speed. It was articulation. Many tasks remained outside software because humans could perform them without being able to specify the complete rule.

Recognizing a cat is easy for a child and punishing for a rule writer. Four legs do not define a cat. Fur does not. Whiskers do not. Pointed ears do not. Every candidate rule admits other animals and rejects unusual cats. Human perception carried a pattern that human language could not finish describing.

Previous information technology could automate what we could formalize. The border of software was the border of explicit human understanding.

## Software before learning

### **The history of conventional software is the history of formalized intention.**

A program is written because someone can specify a transformation. Inputs are defined. Valid states are named. Rules connect states. Outputs are produced according to instructions. Even when the program is too large for one person to understand, its behavior is assembled from explicit operations.

This architecture created a powerful contract. If the code and machine state are known, behavior is explainable through execution. Bugs can be difficult to find, but they arise from implemented logic, integration, state, or specification. The program does not invent a new decision boundary because last month's payroll examples changed its weights.

Conventional software can include randomness, heuristics, search, simulation, and optimization. The distinction is not that old software was simple or deterministic. The distinction is that its operative procedures were authored rather than learned from examples.

Search algorithms already explored spaces humans did not enumerate item by item. A chess program could evaluate positions no programmer had seen. A numerical optimizer could return a design no engineer had drawn. These are important predecessors. They show that machine novelty did not begin with neural networks.

What changed with machine learning is where the evaluation function and representation can come from. In classical search, humans usually define the state, legal moves, and criterion. In learned systems, the evaluation itself can

be fitted from data, and the representation can be transformed through training.

## The specification bottleneck

### **Traditional software fails when people cannot specify the task completely enough.**

Some failures are ordinary missing requirements. Others reveal a deeper problem: tacit competence. A radiologist can notice an abnormal image without being able to list every visual cue. An editor can hear that a sentence is wrong before explaining the rule. A driver recognizes a risky motion that would be tedious to encode as thresholds over pixels.

Human knowledge often lives as practiced discrimination rather than explicit proposition. Conventional programming demanded that tacit ability be converted into rules. The conversion was expensive, brittle, and sometimes impossible.

Machine learning offers another route. Collect labeled examples or self-supervised experience. Define error. Let the system shape a function that reproduces useful discrimination.

This bypasses the specification bottleneck, but it creates a verification bottleneck. If we did not write the rule, how do we establish where it works? If competence is distributed across parameters, how do we find the exception? If the model learns from historical behavior, how do we keep historical failure from becoming predictive policy?

The old system demanded specification before deployment. The new system demands evaluation before trust.

## The old order still matters

**Machine learning has not replaced conventional software. Most reliable systems combine them. Explicit code manages permissions, transactions, invariants, logging, interfaces, and deterministic calculations. Learned models handle perception, ranking, generation, and prediction where complete rules are unavailable or too expensive.**

The advanced design question is not "rules or learning?" It is "Which parts require guarantees, and which parts benefit from learned uncertainty?"

A bank balance should not be guessed. A fraud risk may need estimation. A tax formula should execute the law. A document classifier may learn from examples. A medical dosage calculation may be explicit even if diagnosis support is probabilistic.

The boundary between rules and models is an architectural judgment. The historical break becomes useful only when we know where not to apply it.

## CHAPTER 06

# The inversion

## Machine learning reverses the location of the pattern.

Instead of giving the machine a complete rule, we give it examples, a structure capable of changing, and an objective that distinguishes better predictions from worse ones.

The system predicts. The prediction meets an observed target. The difference becomes error. An optimization process changes parameters in a direction expected to reduce future error. Then the system predicts again.

Prediction. Error. Adjustment.

Repeated across enough examples, the system acquires internal structure that supports better performance. No programmer writes every feature, interaction, exception, or threshold. The programmer designs the learning arrangement: architecture, data pipeline, objective, optimization process, evaluation, and constraints. The particular predictive relationships emerge through training.

$$\theta^* = \arg \min$$

The equation is compact. Its consequences are not. The objective does not state the rule for recognizing the cat, translating the sentence, detecting the fraud, or completing the proof. It states what counts as being wrong. Error supplies pressure. Training shapes a model that becomes less wrong on the data distribution.

This is the historical inversion:

**Traditional computing: Here is the pattern. Execute it. Machine learning: Here are the examples. Find structure that predicts them.**

Human understanding does not disappear. It moves. Humans still choose what to measure, what to optimize, what data to collect, what failures matter, what behavior is acceptable, and whether the output deserves use. But the human no longer has to enumerate the complete transformation before the machine can perform it.

Artificial intelligence is different from earlier information technology because useful structure can now emerge inside computation instead of entering computation only after a human has understood it.

## Learning means changing under error

**The word learning is often used loosely. A database grows when new records are added. A cache changes when new items are stored. A program can update a counter. None of these changes necessarily counts as machine learning.**

Learning means that experience alters the system in a way intended to improve performance on future cases.

The future clause is decisive. Memorizing a training set is change without generalization. A learned system must extract structure that remains useful beyond the examples that shaped it.

In supervised learning, examples pair inputs with targets. The system predicts a target, measures loss, and adjusts parameters. In self-supervised learning, the target is derived from the data itself: a masked token, a next token, a missing patch, a transformed view. In reinforcement learning, feedback arrives through reward as actions affect an environment. The feedback structures differ. Each creates pressure toward behavior that performs better under an objective.

## Gradients are local; capability is global

**Gradient descent follows local information. It estimates how a small change in each parameter would change loss, then moves parameters in a direction that reduces loss. The update is mathematically modest:**

$$\theta_{t+1} = \theta_t - \eta * \text{gradient } L(\theta_t)$$

Yet billions of modest updates across rich data can produce global capabilities that were never individually specified. Translation, summarization, coding, and multi-step task behavior can arise because they help predict structured data and respond to later training.

This creates a gap between training signal and observed capability. No single gradient says "learn irony." Prediction pressure across many contexts rewards representations that make irony easier to model. The named human concept may emerge as a useful distributed regularity.

Emergence should not be treated as magic. It means that a system-level behavior becomes visible when scale, data, representation, and optimization cross a threshold. The behavior remains caused by the training process even when no individual update corresponds to the human label.

## Objectives are compressed instructions

**Machine learning does not eliminate programming. It moves part of the instruction into an objective.**

An objective is a compressed statement of what counts as better. That compression is powerful and dangerous. "Predict the next token" permits broad representation learning because many linguistic and world regularities improve the task. But it does not say "be truthful," "respect uncertainty," or "optimize the user's real outcome." Those properties require additional data, objectives, system design, and evaluation.

The objective also creates blind spots. A system trained for average accuracy may neglect costly rare errors. A recommendation model trained for clicks may



learn outrage. A delivery system trained for speed may externalize risk. The machine can learn an extraordinary pattern for the wrong target.

This is not evidence that objectives are useless. It is evidence that objectives are governance.

## The inversion has limits

**We do not simply give the machine data and wait for truth. Humans choose the dataset boundaries, labels, architecture, optimization method, compute budget, and stopping point. The learned rule is not explicitly written, but the conditions under which it can emerge are engineered.**

The inversion is therefore specific:

The human no longer has to specify the complete input-output mapping.

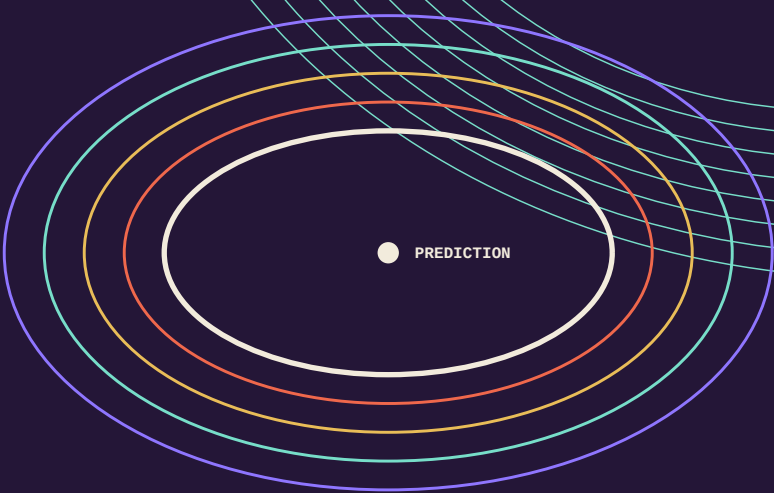
The human still has to specify the learning situation and take responsibility for its use.

That distinction protects the thesis from two distortions. It prevents us from describing AI as merely old software with statistics. It also prevents us from describing the model as an autonomous discoverer detached from human choices.

PART III

# The representation

A learned model builds spaces in which difficult relationships become easier to express.



● PREDICTION

## CHAPTER 07

# The geometry of intelligence

## Representation is the hidden revolution inside machine learning.

The raw world does not arrive in a form that makes every useful relationship obvious. Pixels are not objects. Sound-pressure samples are not words. Tokens are not arguments. A strong model transforms inputs into internal representations where relevant distinctions can be made more easily.

Modern neural networks represent information as vectors: ordered lists of numbers. A word, image region, audio segment, or system state becomes a point in a high-dimensional space. Layers transform those points. Directions and distances acquire functional meaning because they help the system reduce predictive error.

The model is not required to label each direction in human language. It only needs the geometry to work.

Imagine that two classes overlap when plotted by height alone. Add width and they begin to separate. Add texture, motion, context, and history, and a boundary that was impossible in one coordinate system may become simple in another. Learning a representation means learning which transformations expose the structure needed by the task.

This is why deep learning broke through problems that had resisted hand-built features. The earlier system depended on humans to decide what measurements mattered. The learned system can transform raw inputs through many stages, developing features because they improve the objective.

*h*

The phrase "the model sees" is dangerous if taken literally. It invites a tiny human inside the network. A better statement is that the model maps inputs into a geometry whose organization supports prediction.

Humans also depend on representation. A map makes distance inspectable. Algebra makes unknown quantities manipulable. Musical notation makes temporal structure visible. A graph can reveal a relationship hidden in a table. Intelligence is not only possessing facts. It is finding a form in which the relationship can be used.

Machine learning expands that principle beyond forms humans can comfortably inspect. It can operate in thousands of dimensions without needing to picture them. This does not guarantee discovery. It creates representational room in which more complicated regularities can exist.

## A representation makes some answers easy

### **The difficulty of a problem depends on how it is represented.**

Roman numerals and positional notation can express the same quantities, but multiplication is far easier in one representation. A list of city coordinates and a subway map can describe the same network, but route planning is easier on the map. A cloud of points can look tangled in two dimensions and separate cleanly after a transformation.

Representation learning is the process of constructing forms in which useful structure becomes accessible to later computation.

At the input, a digital image is an array of numbers. The array does not arrive labeled with edges, surfaces, objects, depth, or identity. A neural network applies successive transformations. Early layers may become sensitive to local contrasts and textures. Later representations can integrate information across larger regions. The system is not required to reproduce the vocabulary of human vision. It is required to develop states that reduce predictive error.

For language, the challenge is even clearer. The same word changes meaning by context. "Bank" can name a financial institution or the side of a river. A static dictionary entry is insufficient. Contextual representations place the current use into a state shaped by surrounding words and learned relationships. Meaning becomes conditional geometry.

## Distance, direction, and neighborhood

**When people hear "vector," they often imagine a point with an assigned concept. Real learned spaces are less tidy. A concept may correspond to a region, direction, manifold, or distributed activation pattern. Relationships can be encoded across many coordinates.**

Distance can indicate similarity under the model's learned task. Direction can correspond to a changing attribute. Neighborhood can collect items that support similar predictions. But the geometry is not automatically semantic in the human sense. A model trained for one objective may place two items near each other for reasons unrelated to the distinction we care about.

Every geometric interpretation must therefore be tested. A compelling two-dimensional projection can mislead because projection necessarily discards information. Clusters can appear or vanish according to the method. Visual intuition is evidence-generating, not conclusive.

## Attention as conditional routing

**In transformer models, attention mechanisms allow one position to weight information from other positions. The mechanism helps build representations that depend on context rather than a fixed local window. Different attention patterns can route names to pronouns, delimit code blocks, connect questions to evidence, or track structural dependencies.**

The word attention is borrowed and should not be confused with human attention. It names a mathematical operation that produces weighted combinations. A model can compute attention without experiencing focus, fatigue, salience, or surprise.

Yet the mechanism illustrates the central point: representation is active. The model does not merely place each token into a fixed coordinate. It repeatedly reconstructs token states in relation to the current sequence.

## The hidden bargain

**Learned representation trades explicit meaning for performance. Humans no longer decide every feature, which allows the system to discover unexpected structure. In exchange, the resulting coordinates may resist explanation.**

This bargain is not always acceptable. If a decision affects liberty, credit, health, or employment, we may need reasons that can be challenged. A representation can be predictively excellent and institutionally insufficient.

The technical question "Does the geometry work?" must be joined by the social question "Can the use be justified?"

Representation expands what the machine can see. It does not decide what the institution should do with the sight.

CHAPTER 08

# Dimensions are not parameters

**A model can represent an item in a space thousands of dimensions wide while containing billions of learned parameters. These numbers describe different things.**

A dimension is a coordinate in a representation. If a hidden state is a vector of 4,096 numbers, that state occupies a 4,096-dimensional space. The dimensions give the representation degrees of freedom.

A parameter is a learned numerical value that helps transform one representation into another. Weight matrices, biases, and related learned quantities contain parameters. Training changes them. Their collective values determine the function the network computes.

An activation is a value produced for a particular input as it moves through the model. Parameters persist across inputs. Activations change with the input.

| COORDINATE SYSTEM                                                                                                        | LEARNED MACHINERY                                                                                                                  | CURRENT COMPUTATION                                                                                                          |
|--------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------|
| <p><b>Dimensions</b></p> <p>Where a representation can vary. They describe the width and structure of a state space.</p> | <p><b>Parameters</b></p> <p>The numerical values adjusted during training. They determine how representations are transformed.</p> | <p><b>Activations</b></p> <p>The temporary values produced when one specific input passes through the learned machinery.</p> |

Confusing dimensions with parameters produces bad intuition. A billion-parameter model does not necessarily place each concept on its own billion-axis coordinate grid. Concepts are usually distributed across many activations and transformations. A single parameter rarely means "cat," "democracy," or "sarcasm." The capability belongs to the organization of many values.

### Open the distinction with a small network

Suppose a layer accepts a 100-dimensional vector and returns a 40-dimensional vector. Its weight matrix contains  $40 \times 100 = 4,000$  parameters, plus 40 bias parameters. The output representation has 40 dimensions. The layer has 4,040 learned parameters. For each input, it produces 40 activation values.

Dimensions describe the input and output spaces. Parameters describe the transformation. Activations are the particular point produced this time.

Scale this structure across many layers and the numbers become enormous, but the distinction holds. Representation provides space. Parameters shape the terrain. Activations trace a route through it.

High dimensionality is not magic. More room can support richer structure, but it can also support noise, spurious shortcuts, and brittle boundaries. Capacity becomes intelligence only when training, data, and evaluation force the capacity toward structure that generalizes.

## Capacity is not knowledge

**Large parameter counts dominate public descriptions of models because the number is easy to compare. It is also easy to misunderstand.**

Parameters create capacity: the ability to implement a rich family of functions. Capacity is not the same as knowledge, intelligence, memory, or reliability. An untrained network can contain billions of parameters and know nothing useful. Training organizes capacity under data and objectives.

Dimensions describe the size of a representation at a particular location. A hidden state with 4,096 coordinates has 4,096 dimensions even if the model contains hundreds of billions of parameters elsewhere. Activations are the values occupying those coordinates for a particular input.

The distinction can be written as a simple layer:

$$h = \text{phi}(Wx + b)$$



If  $x$  has 100 dimensions and  $h$  has 40 dimensions,  $W$  contains 4,000 weights. Add 40 biases and the layer has 4,040 parameters. For each input, it produces 40 activation values.

Dimensions: the shape of the state.

Parameters: the learned transformation.

Activations: the current computation.

## Distributed structure

**Human concepts do not generally reside in one parameter. Changing one weight rarely removes exactly one idea. Capability is distributed across interacting transformations.**

This distribution gives neural networks robustness and opacity. A local perturbation may have little visible effect because the representation is redundant. A concept can be supported by many paths. At the same time, the absence of a single symbolic location makes explanation difficult.

Distributed does not mean uninterpretable in principle. Researchers can identify features, circuits, attention patterns, causal pathways, and activation directions. But the unit of explanation may not align with a familiar human concept, and a complete account can be larger than the explanation is worth.

## High dimensions without mythology

**High-dimensional spaces behave differently from the three-dimensional geometry of everyday intuition. Volume concentrates in surprising places. Distances can become less discriminating. Random vectors tend toward near-orthogonality. The number of possible directions grows enormously.**

These facts create both opportunity and difficulty. Many attributes can coexist with limited interference. Complex boundaries can be represented. But data becomes sparse relative to the space, and naive similarity measures can fail.

Neural networks do not succeed merely because they have many dimensions. They succeed when architecture, data, and optimization exploit structure that occupies a much smaller effective manifold inside the large ambient space.

The world is not an arbitrary collection of all possible images or sentences. Natural data contains regularity. Faces, grammar, physical scenes, and code occupy constrained regions. Learning is possible because the distribution is structured.

## Scale and the burden of proof

**More parameters can improve predictive fit when supported by enough data, compute, and regularization. Scale can also increase memorization, cost, latency, and the number of behaviors requiring evaluation.**

The right question is never simply "How large is the model?" It is "What capability did the scale buy, under which distribution, at what cost, with which failure modes?"

Parameter count is an engineering fact. Treating it as a single intelligence score is numerology.

## CHAPTER 09

# Random to whom?

**Humans often use the word random when the more accurate sentence is: "I cannot identify a representation that predicts this."**

Those claims are not equivalent.

Consider XOR, the exclusive-or relation. Two binary inputs produce 1 when they differ and 0 when they match. Examine either input by itself and the output remains perfectly balanced. Neither variable alone predicts the answer. Examine both together and the answer is deterministic.

## Inspect the XOR field

Projection destroys the relationship. In the joint representation, the pattern is exact.

XOR is elementary. Its lesson is not. A pattern can vanish when the observer projects the data into too few dimensions. Nothing about the underlying process changed. Only the representation changed.

Now replace two binary inputs with hundreds of interacting variables. No single measurement may carry much signal. Pairs may look weak. Simple charts may show noise. Yet a high-order relationship could still reduce uncertainty when the variables are represented together.

This is where machine learning can exceed normal human pattern discovery. Human conscious reasoning is powerful but representationally narrow. We favor a small number of named variables, visible trends, and relations that can be verbalized. A trained model can exploit a distributed relationship across more coordinates than a person can hold in explicit attention.

That does not abolish randomness. Quantum measurement, thermal noise, unobserved causes, chaotic sensitivity, and computational irreducibility all place different limits on prediction. Even a deterministic process can become practically unpredictable when tiny errors in initial conditions grow rapidly.

The correction is narrower and more consequential:

**Some noise is structure seen through an inadequate representation.**

The arrival of stronger pattern-finding systems therefore changes the status of the unknown. "Unpredictable" can no longer be treated as one category. We must ask whether the barrier is fundamental, informational, representational, computational, or temporary.

## XOR is a philosophy lesson disguised as a truth table

**XOR matters because it shows how a perfectly deterministic pattern can disappear under the wrong view.**

Let X and Y be independent fair bits. The output is one when they differ and zero when they match. Observe X alone. The output is equally likely to be zero or one. Observe Y alone. The same is true. Each variable appears useless. Observe the pair and uncertainty vanishes.

The missing information was not in a hidden third variable. It was in the relation.

This is a warning against feature-by-feature reasoning. A variable can have no marginal predictive power and become decisive in interaction. Human analysis often begins by ranking individual factors because the result is easy to explain. The procedure can erase the pattern it seeks.

Real systems are rarely clean XOR tables. Interactions may be noisy, continuous, delayed, and conditional on context. But the logical lesson survives: predictive information can live in configuration rather than component.

## Three kinds of randomness

### **The word random mixes at least three different conditions.**

Epistemic uncertainty comes from ignorance. The process may be predictable with better data, representation, or models.

Aleatoric uncertainty belongs to stochastic variation that remains even under the best available state description. The correct prediction is a distribution, not a point.

Computational unpredictability occurs when the process is determined or structured but calculating the outcome is infeasible within the available time and resources.

These categories are not always easy to separate. What appears aleatoric may conceal an unobserved variable. What appears computationally impossible may yield to a new algorithm. What appears epistemic may persist because the required measurement is physically inaccessible.

The categories still improve reasoning. They tell us what kind of progress to seek. Collect information for epistemic uncertainty. Quantify and manage aleatoric risk. Approximate or reformulate computationally hard problems.

## Chaos and precision

### **Chaotic systems are deterministic and practically bounded in forecast horizon. Tiny errors in initial conditions grow. More precise measurement and better models extend the useful horizon, but they may not make the distant state reliably predictable.**

This is not a contradiction. Pattern can support strong local prediction without supporting arbitrary long-range prediction. The system has structure, but information about the present state is finite.

Weather forecasting is the canonical example. The atmosphere obeys physical regularities. Forecasts are valuable. Yet the farther horizon eventually outruns the precision of initial observations and model approximations.

## Random to a civilization

**The observer in "random to whom?" can be an individual, a laboratory, a company, or an era.**

Before germ theory, disease outbreaks contained structure that institutions could not represent. Before spectroscopy, starlight contained chemical information that no naked eye could read. Before modern statistical learning, large interaction patterns could remain operationally invisible.

Tools change the observer. The phenomenon stays where it was.

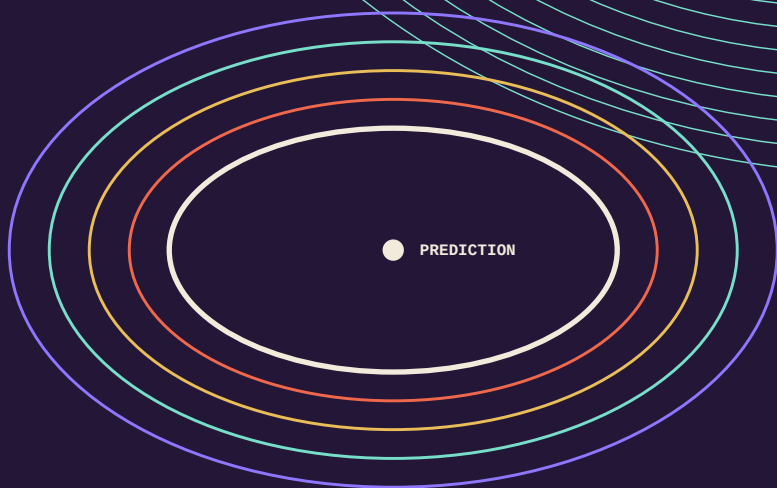
This is the epistemic force of AI. It does not guarantee that today's noise contains discoverable signal. It makes our previous inability weaker evidence that no signal exists.

The responsible posture is neither certainty nor resignation. It is a better question: What kind of observer would make this process predictable?

PART IV

# The epistemic machine

Prediction can become useful before the pattern becomes humanly explainable.



## CHAPTER 10

# Useful before understandable

**Under the old order, understanding usually preceded usefulness. A human discovered a relationship, converted it into a method, and then asked a computer to execute it.**

Machine learning permits the opposite order. A model can exploit a relationship that improves prediction before a human can express that relationship as a compact rule.

Give a system rich patient data and verified outcomes. Ask it to estimate the probability of a disease. It may discover that a complicated interaction among measurements, history, timing, and image features changes risk. The prediction may be reproducibly useful even when no physician can summarize the interaction in one sentence.

This creates a historically unusual conversation.

The machine predicts correctly.

The human asks: "Why?"

With traditional software, the answer was available in principle because the machine knew only the logic we supplied. With a learned model, the effective logic is distributed across parameters and conditional activations shaped by training. The model may be inspectable in fragments without yielding a complete human-scale explanation.

*Three standards, not one* Predictive accuracy asks whether the output is right. Interpretability asks whether humans can trace relevant internal or input-output relationships. Explanation asks whether we possess a satisfying account of why the phenomenon



*itself occurs. These standards can move together. They are not identical.*

The existence of a useful black box does not eliminate the need for explanation. In medicine, law, finance, and public administration, the reason can matter as much as the score. A prediction can be accurate on average and unjust in application. It can rely on a proxy that fails after a policy change. It can expose a correlation that is real but unsafe to manipulate.

The new possibility is not "accuracy replaces understanding." It is "accuracy can arrive before understanding."

That reversal gives human inquiry a new starting point. Instead of asking only, "What theory should we encode?" we can ask, "What structure is the model exploiting, where does it hold, and can we turn it into a testable human account?"

The model becomes an instrument for detecting the presence of knowable structure. It does not automatically deliver the knowledge in human form.

## Performance can outrun prose

**Human understanding is often measured by the ability to state a rule. Learned systems challenge that standard because useful performance can arise from a function too large, distributed, or context-sensitive to summarize.**

There are several reasons a model can be hard to explain.

The mechanism can be enormous. A complete trace may involve millions of activations.

The representation can be alien. The model's useful coordinates may not correspond to human categories.

The behavior can be conditional. What matters in one input may not matter in another.

The explanation request can be underspecified. "Why did the model produce this?" may ask for a causal trace, an influential feature, a contrast with another outcome, a human-readable justification, or the training history that created the behavior.

Different questions require different explanations.

## Interpretability methods answer partial questions

**Feature attribution estimates which inputs influenced an output. It can show sensitivity without revealing a complete internal mechanism.**

Probes test whether information is recoverable from a representation. Recoverability does not prove the model uses the information in the decision.

Ablation removes a component and measures behavioral change. It provides causal evidence about the component but can miss redundancy and compensation.

Mechanistic analysis attempts to identify circuits that implement a behavior. It can yield precise local explanations while leaving system-wide organization unresolved.

Counterfactual explanations ask what minimal input change would alter the output. They can help a user contest a decision without explaining why the model learned the boundary.

No single method produces "the explanation." Each makes a specific claim under assumptions.

## Accuracy is not a permission slip

**A validated black box may be reasonable in one setting and unacceptable in another. The judgment depends on consequences, reversibility, appeal, distribution of error, and available alternatives.**

A music recommender can tolerate opaque mistakes because the harm is small and feedback is immediate. A sentencing recommendation carries a different burden. Even high average accuracy cannot answer whether the target is legitimate, whether protected groups bear unequal errors, or whether the affected person can challenge the evidence.

Institutional standards are not technical obstacles to intelligence. They are part of what makes capability usable in a society.

## Understanding can follow usefulness

**The optimistic sequence is not permanent opacity. A model identifies a robust relationship. Researchers localize the signal. New experiments separate hypotheses. A human theory emerges.**

In that sequence, machine prediction acts as a detector. It tells us where explanation may be worth the work.

This is a new relationship between tool and theory. The instrument does not merely measure a variable already named by science. It can reveal that an unnamed combination carries predictive structure.

Usefulness arrives first. Understanding catches up.

## CHAPTER 11

# Computation before understanding

**For centuries, formal computation followed conceptual discovery.**

Newton supplied equations; people calculated their consequences. Engineers developed theories of load; computers solved larger structures. Accountants established financial categories; software produced statements. A human theory defined the meaningful variables and relationships before computation scaled them.

Artificial intelligence permits another sequence:

The machine can provide evidence that structure exists before a human knows its form. A stable predictive gain is not yet a scientific explanation, but it is a signal: some information in these observations reaches the outcome.

This changes computation from a calculator of known relationships into a detector of possible relationships.

The difference is as profound as the difference between a telescope and a table of planetary positions. The table stores known observations. The telescope extends observation. A learning system extends pattern detection. It can make regularity operationally visible through improved prediction even when the human eye sees only a dense field of variables.

But computation before understanding introduces a new discipline. Predictive performance must survive outside the data that produced it. Training error is not discovery. A flexible model can memorize. A leaky dataset can reveal the answer through an accidental channel. A benchmark can reward the wrong shortcut. Repeated testing can manufacture significance.

The evidence becomes serious only when the relationship survives appropriate separation: new samples, new times, new settings, adversarial checks, changed measurement procedures, and competing hypotheses.

Computation can precede understanding. It cannot excuse validation.

Nor can it settle ontology by itself. A model may use a variable because it predicts. That does not tell us whether the variable is a cause, a consequence, a proxy, a measurement artifact, or a marker for something omitted. Predictive structure opens the investigation. It does not close it.

**The machine can discover that there is something to understand before it can tell us what that something is.**

## A predictor can become an epistemic instrument

**A thermometer converts invisible molecular motion into a readable scale. A microscope converts inaccessible spatial detail into visible form. A machine-learning model can convert inaccessible statistical structure into predictive advantage.**

The output is different from a direct measurement. A probability estimate is mediated by training data, representation, objective, and model class. Yet stable out-of-sample improvement can still tell us something real: the inputs contain information about the target that the baseline did not exploit.

This makes model comparison a form of inquiry. Suppose a simple linear model fails and a nonlinear model succeeds. The difference suggests interaction, thresholds, or geometry absent from the linear representation. Suppose a model using time order succeeds while a shuffled version fails. Temporal structure matters. Suppose removing one data source destroys performance. That source carries unique or enabling information.

These are not final theories. They are experimental clues.

## The danger of leaderboard science

**Prediction invites a seductive shortcut: if the score improves, knowledge improved.**

Scores can rise for the wrong reasons. A benchmark may overlap with training data. Labels may be noisy. The test set may be repeatedly tuned against. A metric may reward superficial cues. A tiny average gain may come from one easy subgroup while critical cases worsen.

Scientific use requires preregistered or protected evaluation, replication, uncertainty estimates, and tests designed around alternative explanations. The model should meet new evidence, not merely a familiar scoreboard.

## From anomaly to theory

**Anomalies are especially valuable. A model that predicts most cases and fails systematically on a subset has discovered the boundary of its representation or the presence of a different process.**

Rather than hiding those errors inside an average, researchers can cluster them, inspect their conditions, and propose new variables. The error becomes a map of what the model - and perhaps the field - does not yet understand.

This reverses the usual emotional posture toward error. Error is not only a defect to minimize. It is information about missing structure.

## Computation before understanding needs a stopping rule

**Predictive mining can produce endless candidate relationships. At some point, inquiry must decide which ones deserve causal testing and theoretical attention.**

Useful stopping criteria include effect size, stability, cost of verification, actionability, novelty relative to known theory, and consequences of being wrong. A relationship that improves a trivial metric by a fraction may not justify years of interpretation. A modest signal in an important disease may.

The machine can widen the search. Human judgment still allocates the scarce resource of serious investigation.

## CHAPTER 12

# From automation to discovery

**Automation removes execution from a known transformation. Discovery searches for a transformation that was not already known.**

A factory robot automates a motion. Payroll software automates a calculation. A database automates storage and retrieval. A compiler automates translation between formal representations. Each system can be remarkably complex. In each case, people specify the operation and its success conditions.

Machine learning can participate earlier. It can search over functions, representations, designs, molecules, policies, proofs, or control strategies and use predictive feedback to move toward useful regions.

This is not unrestricted discovery. Every search has a space, an objective, evidence, and constraints. The model can only find what the setup makes expressible and what the feedback makes selectable.

The cleanest way to see the difference is:

| KNOWN TRANSFORMATION                                                                                    | UNKNOWN MAPPING                                                                                   | UNKNOWN POSSIBILITY                                                                                   |
|---------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------|
| <b>Automation</b><br>Humans specify the rule. The machine repeats it faster, cheaper, or more reliably. | <b>Learning</b><br>Humans specify examples and loss. The machine fits structure that generalizes. | <b>Discovery</b><br>Prediction and search expose candidates that humans did not enumerate in advance. |

Discovery emerges when the learned structure supports a candidate that changes the human search. A model proposes a molecule with desired properties. It identifies a mathematical relation. It reveals an unexpected subgroup. It generates a design that satisfies constraints. The candidate remains a prediction until reality answers.



That final sentence preserves the boundary. AI predicts. Experiments meet the world. Institutions decide what evidence counts. Humans accept responsibility for acting.

A generated hypothesis is not a discovery merely because it is novel. A predicted protein structure is not a treatment. A plausible proof is not a theorem until verified. A high-scoring policy is not legitimate until its consequences and values are examined.

Artificial intelligence accelerates the production and ranking of possibilities. It changes discovery by making search over rich spaces cheaper. But discovery still requires the conversion of a predicted possibility into an actual, tested relationship.

## Search changes when evaluation becomes learnable

### **Discovery requires two capabilities: generating candidates and evaluating them.**

Classical computation has long generated possibilities. Optimization algorithms search schedules, routes, structures, and proofs. The decisive change occurs when the evaluation function itself can be learned from complex data.

A molecule generator is only useful if another system can estimate properties, toxicity, synthesizability, and novelty. A theorem generator needs proof checking or reliable verification. A design system needs simulation, constraint checking, and human evaluation. Generation without evaluation creates volume. Discovery requires selection.

Machine learning lowers the cost of approximate evaluation. A surrogate model can predict the result of an expensive simulation. A learned representation can rank candidates before laboratory testing. This changes the feasible search space.

## The closed loop

### **The strongest discovery systems form a loop:**

- Generate a candidate.
- Predict its value or behavior.
- Select the most informative or promising test.
- Measure the actual result.
- Update the model.
- Generate again.

The loop converts prediction into an experiment policy. It does not merely ask what is likely. It asks which next observation would reduce uncertainty or improve the objective most.

Active learning formalizes part of this idea. Instead of labeling examples at random, the system selects cases expected to be most informative. Bayesian optimization uses a probabilistic surrogate to balance exploitation of promising regions with exploration of uncertain ones.

## Discovery is not novelty

**Generative systems make novelty cheap. They can produce more hypotheses, designs, and combinations than humans can review. Novelty becomes the beginning of the bottleneck, not its solution.**

A discovery must connect to evidence. A novel molecule must be synthesized and tested. A mathematical conjecture must be proved or refuted. A proposed mechanism must survive intervention. A business strategy must meet customers and constraints.

The machine predicts a possibility. Reality converts the possibility into knowledge.

## Credit and responsibility

**When discovery emerges from human questions, public datasets, model training, automated search, laboratory robots, and institutional review, the old image of one discoverer becomes incomplete.**

Credit must reflect contributions without pretending the machine is morally responsible. The system can be causally essential to a result while remaining an artifact operated within human institutions. Responsibility for validation, disclosure, harm, and use stays with people and organizations.

AI expands discovery by moving candidate generation and pattern evaluation earlier and faster. It does not eliminate the social machinery that makes a result trustworthy.

## CHAPTER 13

# When noise becomes information

**Information is not simply data. It is a reduction in uncertainty relative to a question.**

A laboratory archive may contain millions of measurements that nobody can use. The data exists. The information becomes visible when a representation links part of that archive to an outcome, a mechanism, or a decision.

Machine learning changes the economics of making those links. It can test more interactions, absorb more examples, and form richer intermediate representations than a person can inspect manually. As predictive capability rises, some observations move from the category "noise" into the category "signal."

This movement is never automatic. A larger model can also turn accidents into apparent patterns. If the dataset records the hospital scanner rather than the disease, the machine may predict the source institution. If the answer leaks into a timestamp, the model may appear brilliant. If historical discrimination shaped the target, predictive accuracy may reproduce the discrimination.

There are therefore two opposite errors:

- Human provincialism: assuming no pattern exists because humans cannot see one.
- Machine credulity: assuming every model gain reveals durable structure in the world.

Advanced students must refuse both.

The correct question is not "Did the model find a pattern?" Flexible models always find something. The correct questions are: Does it survive new data? Does it survive intervention? Does it survive measurement change? Does it transfer across populations? Does the gain exceed a simpler baseline? Does the

representation rely on information that will exist when the prediction is needed?

When those tests succeed, the status of the data changes. What was formerly too entangled to use becomes predictive. The archive becomes an instrument.

This is one reason AI alters the boundary of economically useful information. A variable can matter even when no analyst would have chosen it in advance. A weak trace can matter in combination. Old records can acquire new value when better representation makes their structure accessible.

The world did not become less structured. The observer became more capable.

## Data becomes information inside a model

**An archive can be enormous and uninformative. Size does not guarantee structure that answers a question.**

Information appears when observations change uncertainty. A million records that repeat the same known fact add volume but little new information. One unexpected measurement can reorganize a theory.

Machine learning can recover weak, distributed information from large archives. A small signal repeated across many examples can become estimable. Several individually weak features can combine into a strong predictor. A representation can align records collected under different surface forms.

This is one reason old data can acquire new value. The archive did not change. The model class and question did.

## The long tail becomes visible

**Human analysis tends to focus on dominant patterns because rare combinations are difficult to inspect. Models can identify long-tail structure when enough examples and appropriate objectives exist.**

But the long tail creates a paradox. Rare cases are where flexible models may be most valuable and where evidence is thinnest. A model can appear confident because it interpolates from nearby representations, yet the specific subgroup may have little support.

Uncertainty must therefore accompany prediction. Calibration asks whether events assigned a probability of 70 percent occur about 70 percent of the time. A calibrated model can still be inaccurate for an individual, and calibration can break under shift, but it makes probabilistic claims testable.

Selective prediction adds another option: abstain when uncertainty is high. In many systems, knowing when not to predict is a central capability.

## Shortcut learning

**When a model turns noise into signal, we must ask whose signal.**

An image classifier may recognize the watermark associated with a class. A medical model may use hospital identity because disease prevalence differs across institutions. A hiring model may exploit a proxy for historical exclusion. These relationships predict within the dataset. They fail the intended abstraction.

Shortcut learning occurs because the objective rewards the easiest predictive route, not the explanation humans intended. The system is not cheating. It is obeying the available signal.

Robust evaluation deliberately removes or reverses shortcuts. Test on a new hospital. Crop the watermark. balance the background. intervene on the proxy. A pattern deserves trust only after surviving tests that attack convenient alternatives.

## Information has a cost of custody

**As prediction makes data more valuable, collection pressure increases. Logs once considered harmless may reveal health, identity, preference, or vulnerability in combination. The fact that a pattern can be extracted does not establish the right to extract it.**

Information governance must therefore consider inferential privacy, not only explicit fields. A dataset can omit a protected attribute while allowing a model to reconstruct it from correlated variables.

The new value of noise creates a new obligation. Data that seemed inert can become sensitive when representation improves.

## When noise becomes information, institutions change

**Insurance, credit, medicine, logistics, education, and security all depend on deciding which differences matter. Better pattern discovery can improve allocation and prevention. It can also produce a society in which every trace becomes a score.**

The question is not whether the signal is real. The question is whether the use is legitimate, contestable, and proportional.

# The revised scientific method

**Artificial intelligence does not replace the scientific method. It inserts a new instrument between observation and explanation.**

The classical ideal moves from observation to hypothesis, from hypothesis to experiment, and from experiment to revised theory. The AI-assisted path can begin with predictive structure that no human proposed:

The model identifies where prediction improves, which cases violate expectation, which variables interact, or which candidate deserves testing. Humans then design probes that distinguish among possible explanations.

This suggests a revised division of labor.

- Machines search representational space. They fit patterns across scales and combinations that exceed unaided inspection.
- Humans formulate epistemic tests. They ask which alternative explanations remain and what intervention would separate them.
- Reality adjudicates. The experiment, deployment, or new observation determines whether the predicted structure persists.
- Institutions govern use. They decide when evidence is sufficient, which harms are unacceptable, and who is accountable.

The model can also become the object of science. Mechanistic interpretability, feature visualization, attribution, sparse representations, causal tracing, and controlled ablation all attempt to convert learned computation into inspectable claims. The goal is not always to translate the entire network into a paragraph. It is to establish specific, testable statements about what information a system uses and how that use affects output.

Science advances when the cost of a meaningful question falls. AI can lower the cost of proposing hypotheses, screening candidates, finding anomalies, and



designing measurements. That increases the importance of standards. Cheap hypotheses can flood a field. Cheap correlations can exhaust attention. The scarce resource becomes discriminating evidence.

**The new scientific method does not end with the model's answer. It begins with the model's advantage over ignorance.**

## Prediction reorganizes the order of inquiry

**The traditional scientific narrative begins with theory. A scientist proposes a mechanism, deduces a prediction, and tests it. Real science has always been more diverse. Observation often precedes theory. Instruments reveal anomalies. Statistical patterns provoke mechanisms. Machine learning extends this data-first tradition.**

The revised sequence can be:

Observe. Model. Predict. Localize. Intervene. Explain.

The localization step is essential. A black-box performance gain must be turned into narrower questions. Which variables matter? Under what conditions? Which subgroup carries the effect? Which counterfactual breaks it? Where does the model fail?

These questions convert prediction into experiment design.

## Causal knowledge asks a different question

**Prediction estimates what is likely given observations. Causal inference estimates what would change under intervention.**

Ice-cream sales predict drowning because both rise in warm weather. Reducing ice-cream sales will not prevent drowning. The predictive relationship is real; the intervention implied by a naive causal reading is wrong.

For science and policy, the difference is decisive. A risk model can identify who is likely to experience an outcome. It does not automatically identify which action will improve that person's outcome.

Causal knowledge requires assumptions, natural experiments, randomized trials, structural models, or other designs that separate correlation from intervention. Machine learning can support these designs by modeling complex relationships and heterogeneous effects. It cannot infer causation from observational fit alone without assumptions.

## Models as hypothesis generators

**One productive workflow treats the model as a generator of constrained hypotheses. It identifies stable predictive features or interactions.**

**Researchers translate them into candidate mechanisms. Experiments test those mechanisms.**

This translation can fail. A distributed model feature may not have a single human name. The predictive relationship may combine several causal paths. The model may rely on a proxy. Failure is still informative because it reveals that current conceptual categories do not capture the predictive organization.

## Reproducibility after adaptive search

**Machine learning makes it easy to search thousands of relationships. Every adaptive choice consumes evidence. If the same dataset guides model selection, feature discovery, threshold choice, and final evaluation, apparent discoveries will be overstated.**

Protected test sets, external replication, nested validation, and transparent search histories are therefore scientific necessities. The more powerful the search, the more independent the confirmation must be.

## A new bottleneck: verification

**When candidate generation becomes cheap, verification becomes scarce. Laboratories, field trials, expert review, and longitudinal observation do not accelerate at the same rate as model output.**

Science may therefore face a surplus of plausible hypotheses and a shortage of decisive tests. The valuable system will not be the one that speaks most. It will be the one that helps choose which claim deserves contact with reality.

CHAPTER 15

# What AI still cannot predict

**No pattern, no prediction. No relevant observation, no access to the pattern. No stable environment, no guarantee that yesterday's pattern survives tomorrow.**

The limits of prediction are not one wall. They are different walls that require different responses.

## Irreducible uncertainty

**Some events may contain genuine stochasticity. Even with a perfect description of the available state, outcomes retain a distribution. A good model can estimate that distribution. It cannot convert probability into certainty.**

## Missing state

**A process can be patterned while the predictor lacks the variables that carry the pattern. The limitation is observational. More data of the same kind may not help. A different sensor might.**

## Chaotic sensitivity

**Deterministic systems can amplify microscopic uncertainty. Weather has structure and remains bounded in forecast horizon. Better models extend the horizon; they do not abolish sensitivity to initial conditions.**

## Distribution shift

**A model learns from one distribution and meets another. Technology changes, language moves, pathogens evolve, incentives adapt, and institutions rewrite rules. A pattern can be real and expired.**

## Reflexivity

**Predictions can change the system they predict. A public market forecast affects trades. A policing score affects where police look. A ranking changes what people click. The model enters the causal field and alters the data-generating process.**

## Computational limits

**A pattern may exist but require infeasible computation to exploit. The shortest description can be impossible to find. The best plan can be beyond practical search. Solomonoff induction is an illuminating ideal precisely because it is not a usable finite procedure in the general case.**

## Objective failure

**The model can accurately predict the target we supplied while the target fails to represent what we value. Loss functions are not moral theories. Benchmark victory is not social legitimacy.**

### Open the prediction-limit audit

Before calling a system unpredictable, ask: Is the uncertainty intrinsic? Are the relevant variables observed? Is the representation adequate? Is the process stable? Does the prediction change behavior? Is the computation tractable? Is the objective the real objective?

Each answer points toward a different action: accept a distribution, collect a variable, learn a representation, monitor drift, model feedback, approximate the search, or redefine success.

The mature claim is not that AI predicts everything. It is that AI expands the class of patterns that can become predictively useful before humans can formalize them. Expansion is not infinity.

## The limit is a taxonomy, not a confession

**When a prediction fails, "the model was wrong" is a report, not an explanation. Failure must be classified.**

Was the event intrinsically uncertain? Was relevant state missing? Did the model choose the wrong representation? Did the environment shift? Did the prediction change behavior? Was the computation insufficient? Was the target wrong? Did the system communicate more certainty than it possessed?

Each failure points toward a different remedy.

## Missing information cannot be optimized away

**No amount of parameter tuning can recover a variable that leaves no trace in the inputs. If two situations are identical to the model and lead to different outcomes, the best prediction is a distribution over those outcomes.**

This is a frequent source of false confidence. A point estimate hides irreducible or unobserved variation. Decision systems should preserve distributions, intervals, or multiple scenarios when the information does not support one answer.

## Shift turns knowledge into history

**A model estimates patterns under a data-generating process. When the process changes, yesterday's accuracy becomes historical evidence rather than a current guarantee.**

Covariate shift changes the distribution of inputs. Label shift changes outcome frequencies. Concept drift changes the relationship between input and outcome. Policy changes can create all three.

Monitoring should therefore track more than average accuracy. It should inspect input distributions, calibration, subgroup errors, missingness, and feedback loops. Some outcomes arrive too late for immediate performance measurement, which makes leading indicators and periodic audits necessary.

## Reflexive systems learn their own consequences

**A prediction can alter the world it measures.**

If a model predicts high demand and a store increases inventory, the observed shortage disappears. If a platform predicts engagement and promotes an item,

the prediction helps create the engagement. If police are sent where a model predicts crime, recorded incidents become concentrated where police look.

Feedback can make a model appear self-confirming. The data no longer represents an untouched environment. It represents the joint system of environment, prediction, and intervention.

Evaluation must distinguish prediction from policy effect. Randomized deployment, exploration, counterfactual logging, and causal analysis can help. Without them, the model may learn a world partly manufactured by its previous outputs.

## Adversaries study the pattern

**In security, fraud, markets, and competitive systems, the target adapts. Once a pattern becomes known or operational, people change behavior to evade it. Prediction degrades because the environment contains agents with objectives.**

This is not ordinary noise. It is strategic response. Robust systems require ongoing adaptation, hidden evaluation signals, and a recognition that every successful detector creates pressure for a new attack.

## The no-free-lunch reminder

**Learning succeeds because real problems are not uniformly arbitrary. Models bring inductive bias: assumptions about smoothness, locality, compositionality, sparsity, invariance, or other structure. An algorithm that performs well on one family of problems pays for that preference elsewhere.**

There is no universally superior predictor across all possible data-generating functions. Progress comes from matching bias to the structure of the world we actually inhabit.



The limit is therefore not one dramatic edge. It is the continuous requirement to ask what assumptions make prediction possible here.

## CHAPTER 16

# The boundary of the knowable

**Artificial intelligence changes the boundary between what humans understand and what human systems can use.**

Before machine learning, the usable boundary sat close to the articulable boundary. To make software perform a transformation, humans generally had to express the transformation. The machine extended speed, scale, memory, and reliability. It did not usually extend the discovery of the operative pattern.

Now the boundaries separate.

A relationship may be unusable by unaided humans, usable by a model, partially interpretable by investigators, and still not understood as a compact theory. Use can precede explanation. Prediction can precede science. Computation can reveal the outline of a regularity before language catches it.

This creates at least four domains:

- Known and formalized: humans understand the rule and machines execute it.
- Learned and interpretable: machines help discover structure that humans can subsequently state and test.
- Learned and operational: prediction works within validated bounds, but the full internal relationship resists human summary.
- Unreachable: relevant structure is absent, unobserved, unstable, or computationally inaccessible.

The third domain is the new tension. It offers enormous value and demands enormous care. A society can act on relationships it cannot fully explain. Sometimes that is rational. Humans have always used medicines, heuristics, and institutions before possessing complete theories. But machine learning

scales the practice and obscures the path by which the relationship was formed.

Responsibility therefore becomes more important, not less. The inability to explain a model does not transfer moral agency to the model. People choose the target, accept the evidence, define the deployment, distribute the errors, and decide who may appeal.

Prediction enlarges capability. It does not answer what should be done with the capability.

The boundary of the knowable is no longer a single line between ignorance and understanding. It is a layered frontier among what can be predicted, what can be compressed, what can be explained, what can be caused, what can be controlled, and what should be used.

Artificial intelligence moves those lines at different speeds.

## The knowable has layers

**Human culture often treats knowledge as a binary. We know or we do not know. Machine learning reveals a layered frontier.**

We may be able to predict without explaining.

We may be able to explain without controlling.

We may be able to control without understanding every mechanism.

We may be able to compress without preserving the distinctions required by justice.

We may be able to act while uncertainty remains.

These are different epistemic achievements. Treating them as one creates both overconfidence and missed opportunity.

## Operational knowledge

**A validated model can provide operational knowledge: within stated conditions, these inputs support this distribution over outcomes. The claim can be rigorous even if the internal representation resists translation.**

Operational knowledge is bounded by evaluation. It should include the population, time period, data process, uncertainty, known failure modes, and monitoring plan. Without those boundaries, a score becomes an oracle by rhetoric rather than evidence.

## Explanatory knowledge

**Explanation organizes relationships into a human model that supports counterfactual reasoning. It tells us not only what co-occurs but what would happen if conditions changed.**

Explanations can be simpler than the predictive system because their purpose is different. A complete neural trace may be less useful than a causal mechanism with a few variables. The explanatory model does not have to duplicate every computation to be scientifically valuable.

## Normative knowledge

**No predictive model derives a value from data alone. It can estimate consequences under choices. It cannot establish which consequence is just, which risk should be borne, or whose objective should dominate.**

Values enter through target selection, constraints, thresholds, and institutional authority. Hiding those choices inside a score does not remove them. It makes them harder to contest.

## Human responsibility at the widened boundary

**As AI expands operational knowledge beyond immediate human explanation, responsibility must become more explicit.**

Someone must own the target.

Someone must own the evidence standard.

Someone must own the decision to deploy.

Someone must own the appeal.

Someone must decide when the pattern has expired.

The phrase "the model decided" is usually an institutional evasion. The model computed. An organization gave that computation authority.

## The deeper continuity

**Humans have always used patterns before fully understanding them. Farmers learned seasons before atmospheric science. Physicians used treatments before molecular mechanisms. Language itself compresses distinctions discovered across generations.**

AI does not introduce the possibility of useful partial understanding. It industrializes and accelerates it. It produces operational patterns at a scale and opacity that require new disciplines of validation, explanation, and governance.

The boundary of the knowable expands, but it also becomes more complicated. What the machine can use, what the human can explain, and what society should authorize no longer move together.

## EPILOGUE

# We no longer have to arrive first

**The history of information technology was built on an assumption so stable that it rarely needed to be stated.**

The human had to find the pattern first.

We developed accounting, then accounting software. We developed navigation, then navigation software. We developed tax rules, then tax software. We discovered equations, then asked computers to calculate them. The machine could be faster, more accurate, more consistent, and more scalable. But the intelligence that identified the useful transformation came before the program.

Machine learning breaks that dependency.

We can provide experience, an objective, and a system capable of adjustment. Prediction error can shape representations. Those representations can exploit relationships no person enumerated. The model can become useful before the pattern becomes a human explanation.

This is not the end of programming. It is not the end of science. It is not the end of human judgment. It changes their sequence.

Programming increasingly includes designing learning conditions and deterministic boundaries around learned uncertainty.

Science can begin from predictive structure and work backward toward mechanism.

Human judgment moves toward targets, evidence, consequences, permissions, and responsibility.

The most important word in the thesis is not machine. It is first.

Humans still discover patterns. Humans still formalize rules. Humans still produce explanations the machine cannot supply. But useful computation is no longer forced to wait for the complete human account.

The machine can meet the data and return with evidence that structure exists.

Then the real work begins.

What did it find?

Where does the relationship hold?

Which shortcut could explain the gain?

What experiment would distinguish correlation from cause?

What information was discarded?

Who bears the errors?

What right do we have to use the pattern?

When should the system abstain?

When has the environment changed enough that the old model is no longer knowledge?

These are not objections added after intelligence. They are the disciplines required by a technology that participates in discovering its own operative structure.

Artificial intelligence is not different because it talks like us. It is not different because it can produce attractive artifacts. It is not different because it uses more computation than earlier software.

It is different because the machine no longer needs us to find the pattern first.

## APPENDIX A

# A compact mathematical field guide

## CONDITIONAL PROBABILITY

**$P(Y|X)$  expresses the probability of Y given X. X contains predictive information about Y when observing X changes the distribution over Y.**

## SELF-INFORMATION

$$I(x) = -\log_2 p(x)$$

An unlikely event carries more information because identifying it requires more bits under an ideal code.

## ENTROPY

$$H(X) = - \sum_x p(x) \log_2 p(x)$$

Entropy is expected surprise under a distribution. It measures average uncertainty before observing the outcome.

## CROSS-ENTROPY

$$H(p,q) = - \sum_x p(x) \log_2 q(x)$$

Cross-entropy measures the expected code length when events follow distribution p but are encoded using model q. Minimizing it pushes q toward p.

## MUTUAL INFORMATION

$$I(X;Y) = H(Y) - H(Y|X)$$

Mutual information measures how much observing X reduces uncertainty about Y.



#### KOLMOGOROV COMPLEXITY

**$K(x)$  is the length of the shortest program that outputs  $x$  on a fixed universal machine. It is not computable in general.**

#### MINIMUM DESCRIPTION LENGTH

**Choose the model that minimizes the combined description of the model and the remaining unexplained data.**

#### EMPIRICAL RISK MINIMIZATION

**Choose parameters that minimize average loss on observed examples, with evaluation and regularization used to estimate whether the learned structure generalizes.**

#### CALIBRATION

**A probabilistic predictor is calibrated when outcomes assigned probability  $p$  occur with frequency  $p$  under the evaluated conditions.**

#### DISTRIBUTION SHIFT

**The relationship between training and use changes. Shift can affect inputs, outcome frequencies, or the conditional mapping itself.**

APPENDIX B

# Sixteen questions for any predictive system

- What exactly is the target?
- When does each input become available?
- How was the outcome measured?
- Which population produced the data?
- What baseline must the model beat?
- Which metric reflects the actual cost of error?
- How are rare but severe failures weighted?
- What shortcuts could produce apparent performance?
- Has the evaluation been protected from repeated tuning?
- How well calibrated are the probabilities?
- Where should the model abstain?
- What changes when people respond to the prediction?
- How will drift be detected?
- Can an affected person contest the input or result?
- Who owns the deployment decision?
- What evidence would cause the system to be withdrawn?

APPENDIX C

# Selected sources

**Claude E. Shannon. "A Mathematical Theory of Communication." Bell System Technical Journal, 1948.**

Ray J. Solomonoff. "A Formal Theory of Inductive Inference." Information and Control, 1964.

David H. Wolpert and William G. Macready. "No Free Lunch Theorems for Optimization." IEEE Transactions on Evolutionary Computation, 1997.

Shane Legg and Marcus Hutter. "Universal Intelligence: A Definition of Machine Intelligence." Minds and Machines, 2007.

Yoshua Bengio, Aaron Courville, and Pascal Vincent. "Representation Learning: A Review and New Perspectives." IEEE Transactions on Pattern Analysis and Machine Intelligence, 2013.

Grégoire Delétang and colleagues. "Language Modeling Is Compression." 2023.

Yuzhen Huang and colleagues. "Compression Represents Intelligence Linearly." 2024.