

How knowledge is stored ↔ *what knowing is*



Knowledge-graph working group · June 2026

Ways of storing knowledge can sit on a continuum

The continuum asks one question of every mechanism:
does the store keep patterns, or named relations?

*after Bourdieu · Schön · others
could also be tacit/explicit
could also be latent/patent*

IMPLICIT

associative · sub-symbolic

*like knowing a friend's voice
instantly, but nothing you
could write down*

or having 'a feel for the game'

EXPLICIT

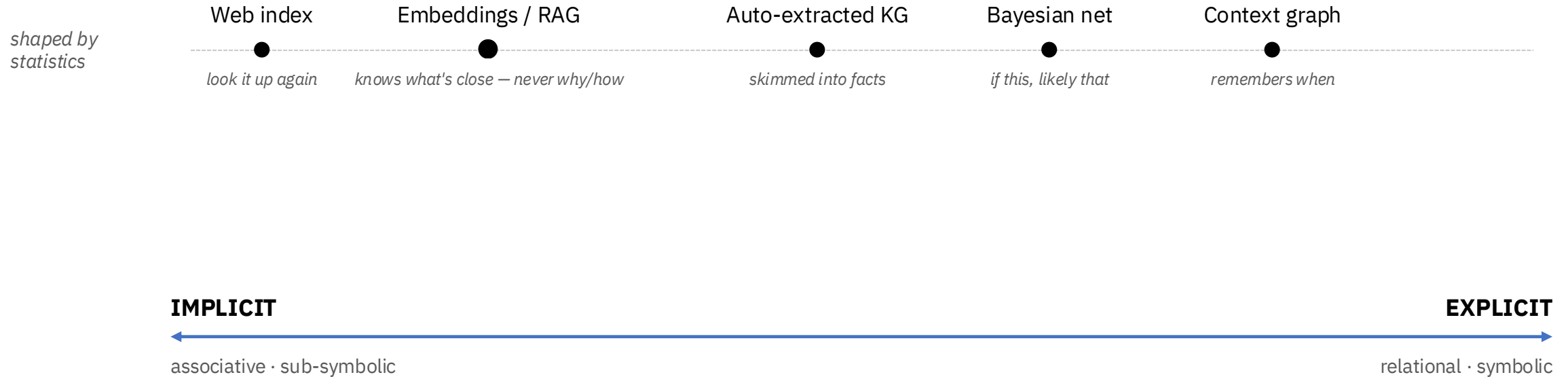
relational · symbolic

*like a recipe...every ingredient
and step written out*

Machines store knowledge across the continuum

after Simon · Ashby · Floridi

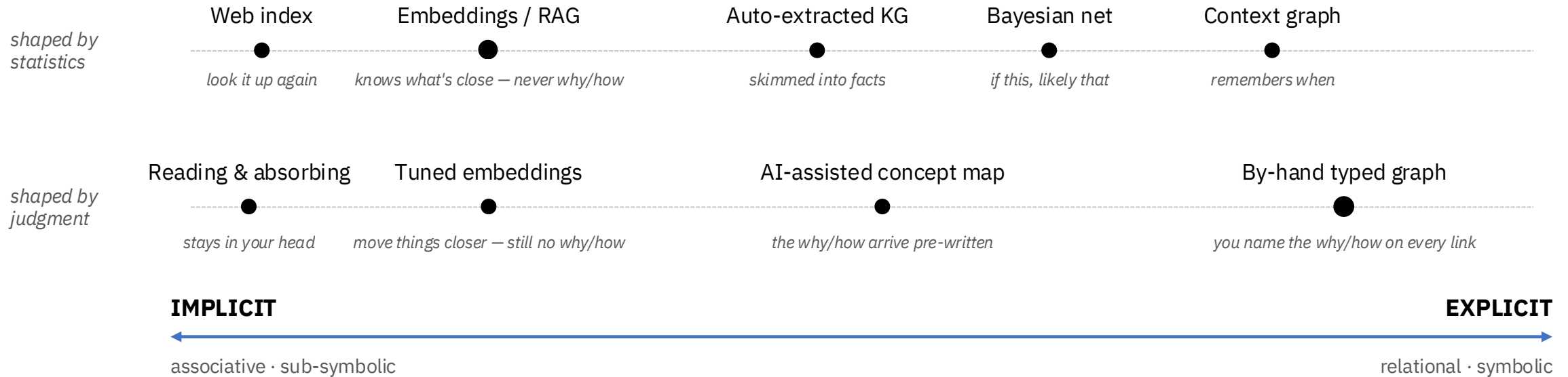
From an index that keeps almost nothing about meaning, to graphs that keep named relations.



People hold knowledge along the same continuum

after Novak & Gowin · Dubberly

From reading that stays in your head,
to a graph where every link is named by hand.



Every mechanism is also a claim about what knowing is

after Bourdieu

When you choose how to store knowledge,
you also choose what *counts as knowing*

Store only similarity, and knowing can only be a feel — you can sense that two things go together, but not say how.

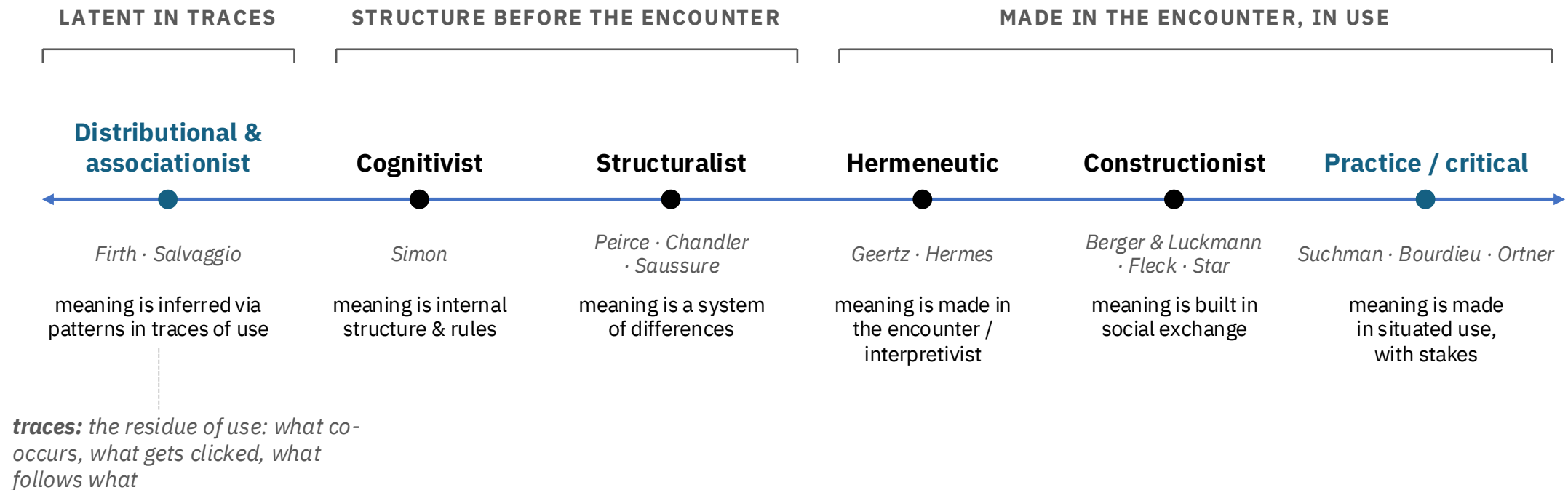
Store named links, and knowing means saying how — and what can be said can be shared, taught, and challenged.

Positions on knowing

entirely attackable — by design

Who, or what, makes the meaning?

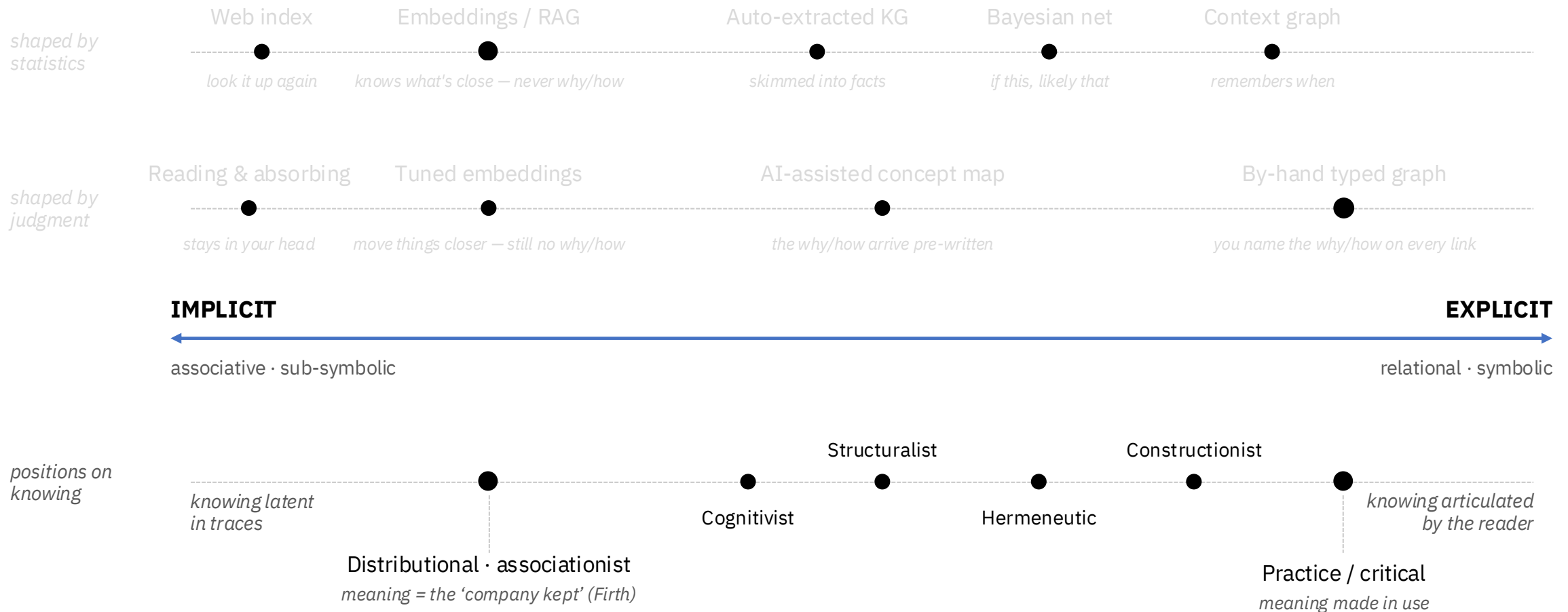
Six positions scholars have defended (placements blur)



The positions land on the same continuum as the tools

*not a complete map —
boundaries for discussion*

Everything on this map orders by how much of the meaning is named
So you can read it vertically!

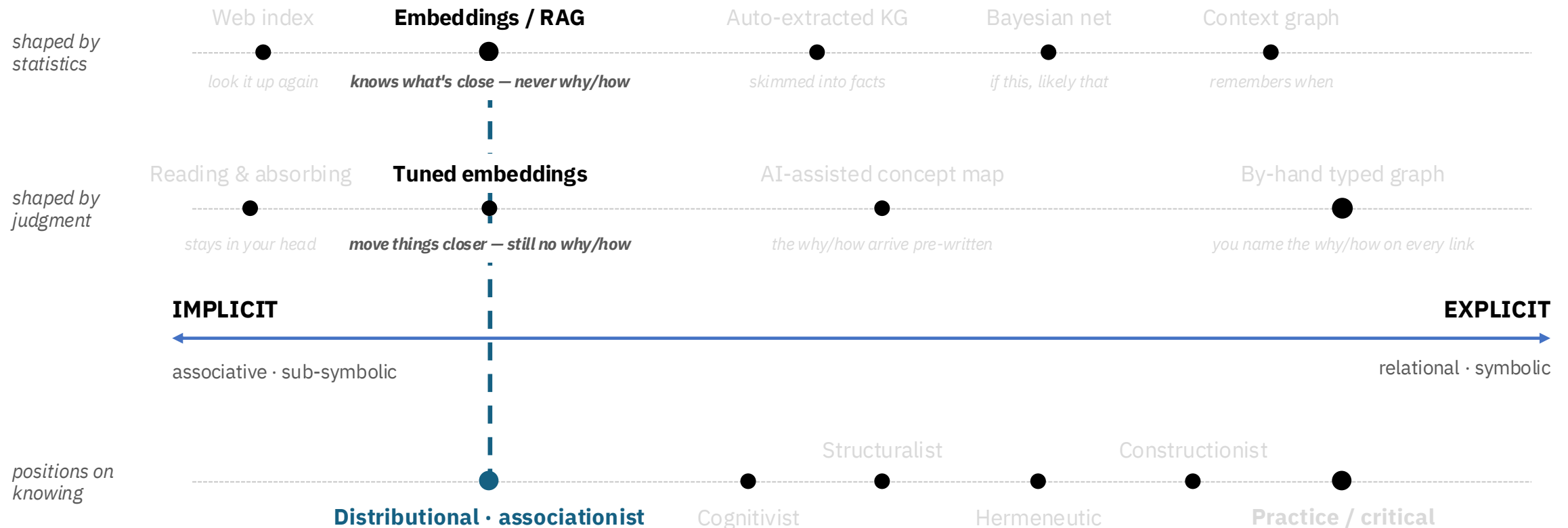


Example move: drop a line straight down

*alignments are claims, not
measurements*

Draw a line from any mechanism down to the position it enacts.

Example: build with embeddings and you've signed on to Firth, wittingly or unwittingly.



**The line is about what the *tool makes*
— not what's in the person's head (what testing can reveal)**

after Winner · McLuhan · Latour

It *shapes* what you practice, what you attend to,
what you can save, what gets rewarded.

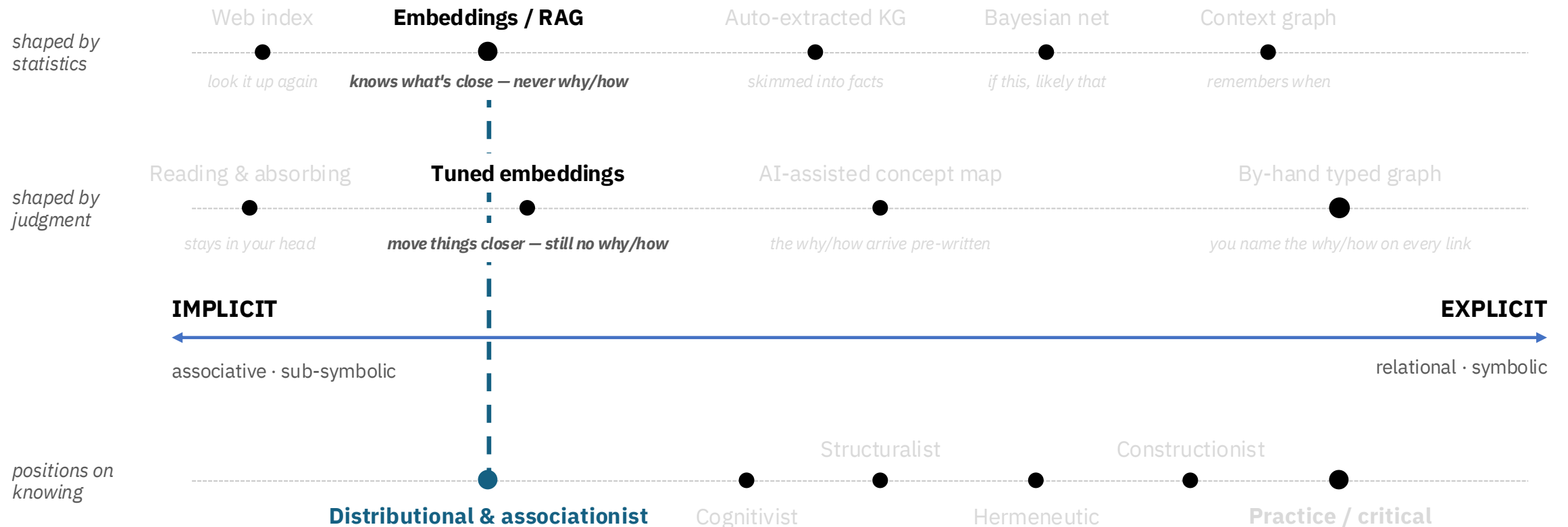
A mechanism enacts a position; it doesn't entail one.

**For teaching, that's the point:
what the artifact can hold is what the practice can grade.**

Embeddings land on distributional, and tuning doesn't move them

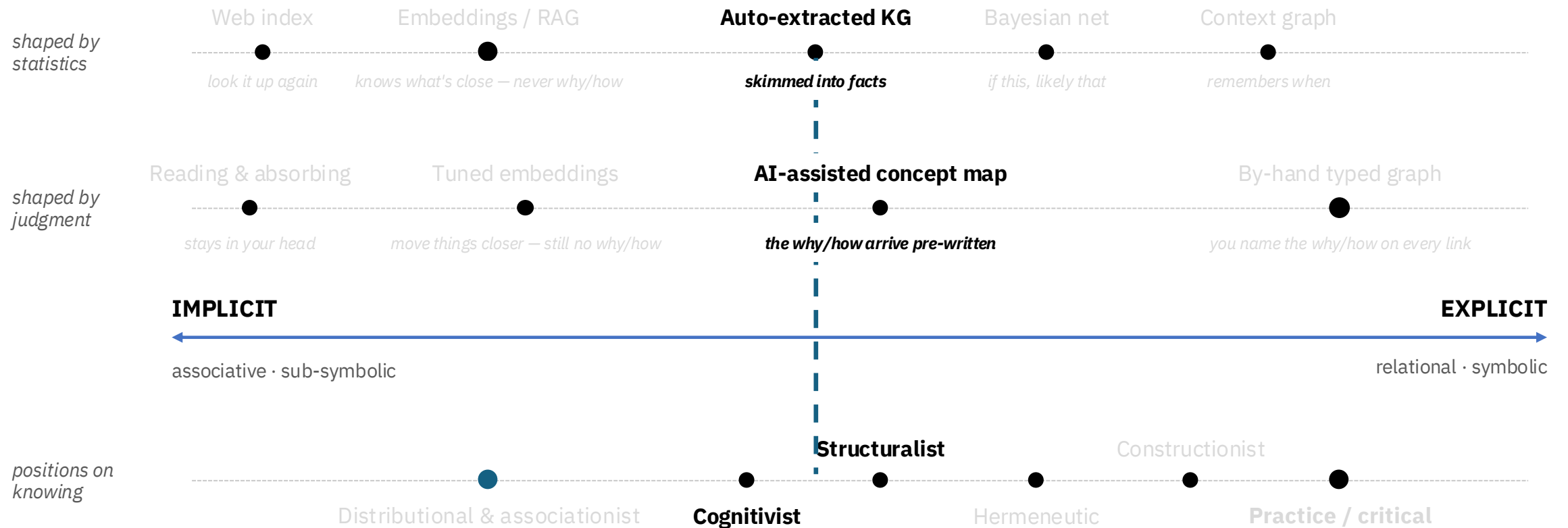
Gain: reach and speed, what's close, across a whole corpus at once.

Loss: never the why — so what was understood stays invisible.



Auto-extracted graph lands mid-band: named links, but *the system* named them

Explicit isn't enough; ask who did the articulating?



The by-hand typed graph lands on practice/critical: every node and edge carries its why/how

Naming the relation is the act of understanding — so the building is the learning (Novak & Gowin)



The map plots artifacts, not readers

Every dot is placed by what a practice leaves behind...the part someone else can inspect.

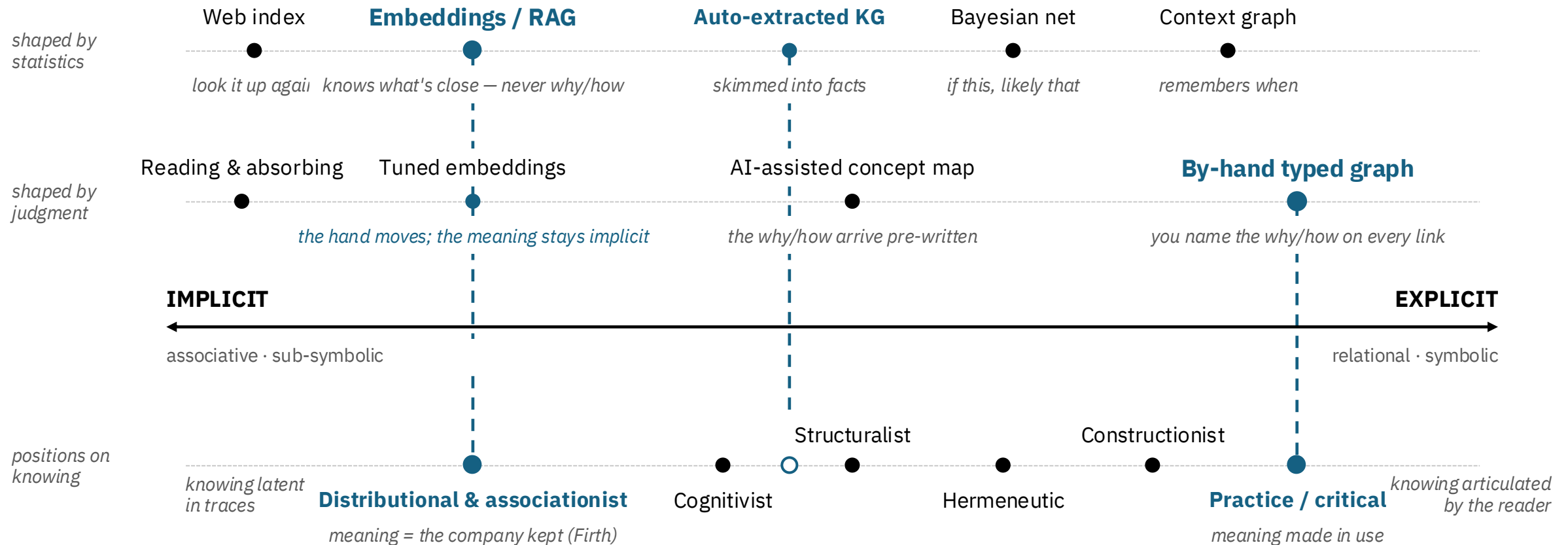
Reading-and-absorbing may be the richest knowing in the room. But it leaves nothing to inspect as an artifact, it is as silent as an embedding

Only on the right does a shallow artifact fail visibly.

An index reads the same whether or not anyone understood. A typed graph cannot hide — every missing why/how shows.

The payoff is range — not one lane, the whole map

Locate a framework on the map, deploy it with dexterity, defend the choice in arguments that last — a repertoire of practice



Not “embeddings bad, graphs good.”

Locate a framework, deploy it with dexterity, defend the choice.

Naming the relation is the act of understanding

Building the graph is the learning, not a report of it. (Novak & Gowin)

And because a shallow graph fails visibly, the work is gradable without an exam — you can see where understanding broke down.