

Computational & AI-Assisted Methods for Social Sciences

A Research Design Primer

Chapter 7

Computational Analysis and Epistemic Purposes

Ji Ma

cssprimer.org



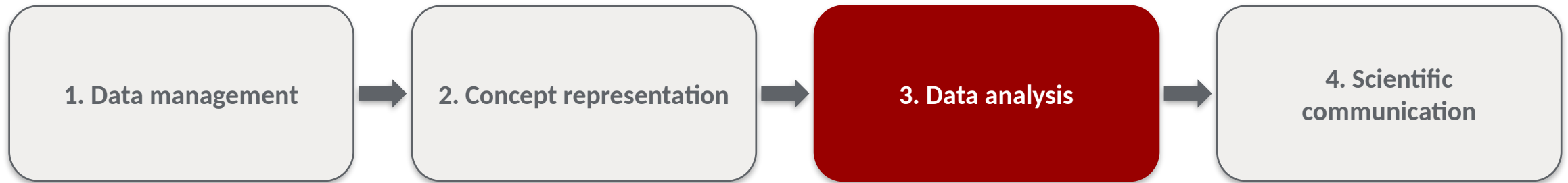
Learning objectives

By the end of this chapter, you can:

- Explain the three computational capacities and where each appears across methods
- Recognize the three challenges those capacities create for inference
- Distinguish distribution from causality approaches to modeling social data
- Match methods and evaluation standards to the four epistemic purposes
- Evaluate the validity claims of AI-assisted methods and synthetic agents

Stage 3 of the pipeline: **data analysis**

Chapter 7



Chapters 7 and 8 develop the third stage: what analysis can claim, and how to practice it.

§7.1

Computational capacities

Iterate, model non-linearity, process high dimensions — what writing code buys you.

Three capacities, one definition

§7.1

- Recall the test from Chapter 1: inference that requires writing and executing code
- What the code buys: iteration, non-linear prediction, high-dimensional processing
- Capacities, not tools — each cuts across many methods

Iteration as optimization and search

§7.1.1

- Gradient descent: guess, compute error, follow the gradient downhill, repeat
- Like seasoning a dish: taste, adjust, taste again
- Dijkstra's algorithm searches million-node networks the same way
- Empirical feedback replaces hand-tuning — but the benchmark matters

Iteration as simulation: rule-based agents

§7.1.1

- Agents follow simple rules across discrete time steps
- Epidemic waves and residential segregation emerge from repeated interaction
- Rerun under different assumptions to compare interventions
- Useful where real experiments are costly, slow, or unethical

Autonomous agents: LLMs as simulated actors

§7.1.1

- Generative agents remember, reflect, plan — and once organized a party from one prompt (Park et al. 2023)
- 1,052 interview-grounded agents beat demographic and persona baselines at predicting their humans' survey answers (Park et al. 2024)
- Realism comes from interview grounding — and inherits its omissions

Iteration as statistical updating

§7.1.1

- Bayes: yesterday's posterior becomes today's prior; evidence accumulates
- Fits how social evidence arrives — poll after poll, wave after wave
- Google Flu Trends overestimated flu once search behavior and algorithms shifted

A result can stabilize because the evidence is genuinely informative — or because the same bias enters the process again and again.

Chapter 7, on convergence and emergence

Social life is not a straight line

§7.1.2

- Stress and performance: an inverted U, not a slope
- Thresholds: a protest gains momentum only past critical mass
- Diminishing returns: each additional study hour helps less
- Interactions: the effect of age depends on education

How a neuron bends a line

§7.1.2

- A unit: weighted sum of inputs, then a nonlinear activation
- $\text{ReLU}(\text{income} - 50)$: zero below \$50k, rising above — a learned threshold
- Stacked layers combine such units into interactions
- Universal approximation guarantees existence, not discovery or generalization

High dimensions: smart summarizers

§7.1.3

- Corpora, panels, and surveys can carry more variables than observations
- PCA and factor analysis have long compressed variables into indices
- Topic models and embeddings do the same for text
- A difference of degree, data type, and workflow — not of kind

AI is not a fourth capacity

§7.1.4

- LLMs run on the same three capacities, at scale
- Role one: coding and labeling — your codebook, a million documents
- Role two: designing agents grounded in interviews or personas
- 2026 survey: 81% of quantitative social scientists tried chatbots; one in five uses coding agents

§7.2

Challenges of computational methods

The same power raises three problems: opacity, meaning, and timing.

Epistemic opacity: the uninspectable path

§7.2.1

- Too many steps, executed too fast, for any human to audit
- No transparent shortcut from inputs to outputs; stochastic paths never seen
- A child-welfare risk score no caseworker can explain to a family or judge
- Rudin: in high stakes, prefer interpretable models of comparable performance

From semantics to application

§7.2.2

- A sentiment score has clear statistical meaning, ambiguous social meaning
- A detected 'community' is dense ties — not shared identity, norms, or action
- Mapping outputs back onto concepts takes theory, context, and judgment

Temporal dynamics: models run in time

§7.2.3

- Two clocks: the model runs in time, and represents a process in time
- Synchronous versus one-by-one updating can change segregation outcomes
- A weather model that takes a million years to run forecasts nothing
- Update order, time steps, and stopping rules are part of the model

§7.3

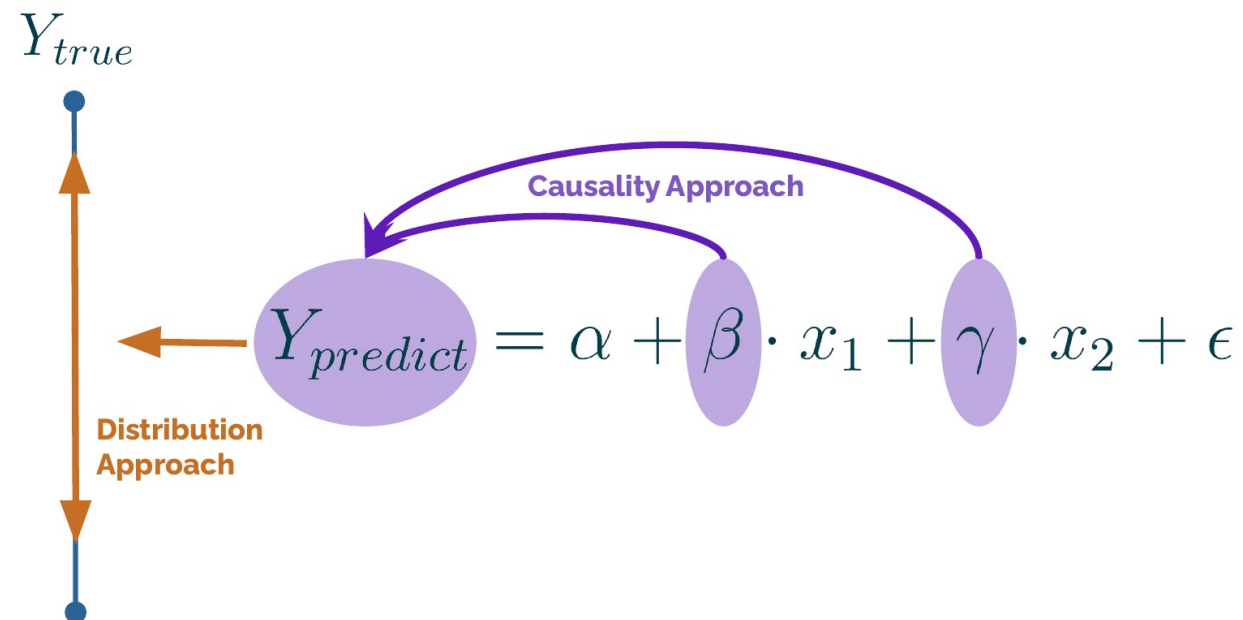
Epistemic purposes

Distribution versus causality — then four distinct research goals, not a ladder.

Two questions you can ask of a model

§7.3.1

Benchmarks vs. Behavioral Evaluations



The distribution approach asks whether predictions match observed outcomes; the causality approach asks what the coefficients mean under intervention.

Structural or agnostic: what can a model recover?

§7.3.1

- Structural: a true data-generating process exists; recover it faithfully
- Agnostic: for social text it rarely does — optimize interpretability instead
- When fit and interpretability diverge, the choice is a modeling commitment

The Hofman 2×2: four epistemic purposes

§7.3.2

Purpose	Core question	Evaluated by
Descriptive	What does the social world look like?	Fidelity and usefulness of the map
Explanatory	What changes under intervention?	Identification of the causal effect
Predictive	How well do we forecast new cases?	Accuracy on held-out data
Integrative	Can we predict when the world itself changes?	Performance out of distribution

Descriptive modeling: mood rhythms in 509 million tweets

§7.3.2

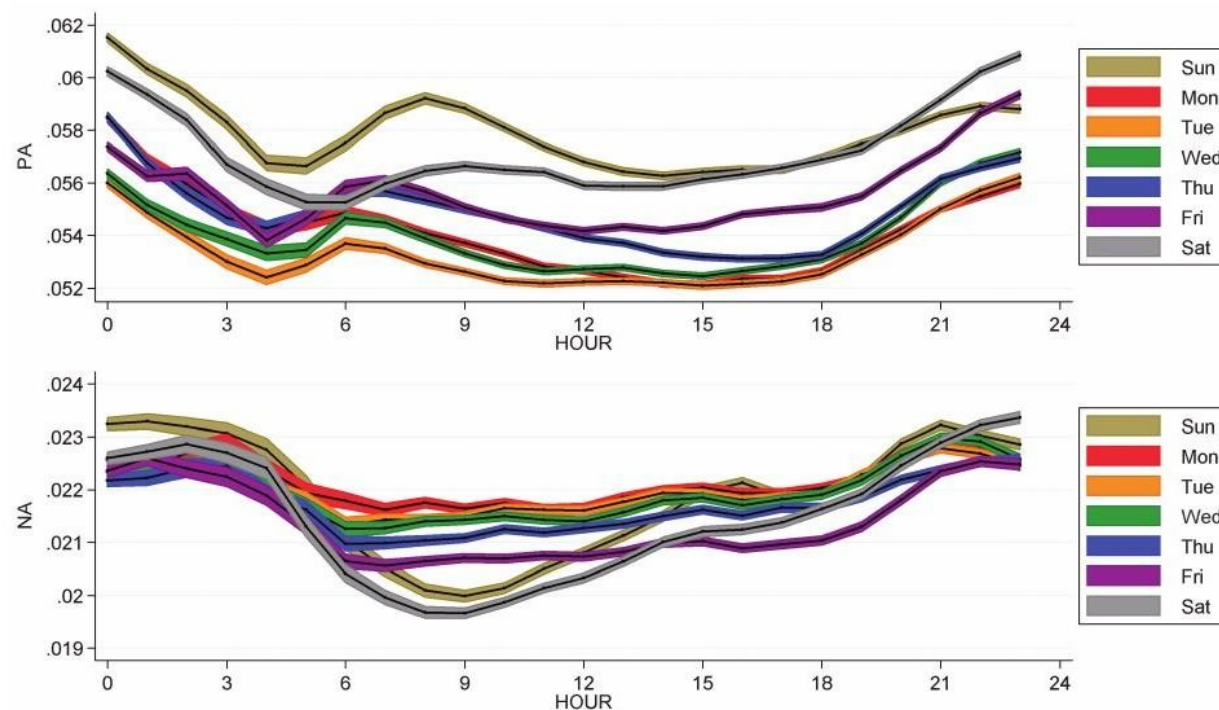


Fig. 1. Hourly changes in individual affect broken down by day of the week (top, PA; bottom, NA). Each series shows mean affect (black lines) and 95% confidence interval (colored regions).

Hourly positive affect (top) and negative affect (bottom) by weekday; weekends run higher and the morning peak arrives later.

Source: Golder & Macy (2011), Fig. 1, Science. © American Association for the Advancement of Science. Included for instructional use.

Descriptive modeling: mapping a bilingual knowledge space

Cited English articles and English and Chinese research topics in one multilingual semantic space; distance means thematic similarity, node size means volume.



Source: Ma (2024), Fig. 5, Nonprofit and Voluntary Sector Quarterly. © SAGE Publications. Included for instructional use.

Takeaway: Description can map abstract intellectual spaces, not only concrete networks.

Explanatory modeling: the opposing-views bot

§7.3.2

- Paid frequent partisan Twitter users to follow an opposing-side bot for a month
- Republicans became more conservative; Democrats shifted little, not significantly
- The value is a causal estimate of one specific intervention

Predictive modeling: proxies that work

§7.3.2

- 50 million Street View images; 22 million vehicles detected and classified
- Neighborhood car patterns estimated income, race, education, and voting
- Early-warning systems rank cases for scarce interventions
- Correlations that forecast well can vanish when policy or behavior shifts

Integrative modeling: predicting under change

§7.3.2

- Specify the change; capture the causal structure; test out of distribution
- Causal ML — double/debiased learning, causal forests — serves the middle requirement
- Moral Machine: prediction critiques the explanation; experiments test the revision
- Predictive failure is informative: it locates the incomplete causal account

Interpretability is not causality

§7.3.2

- Explanations help a decision-maker trust a score or choose how to intervene
- Mug-shot morphing yields hypotheses: 'well-groomed,' 'heavy-faced'
- That is hypothesis generation, not hypothesis testing
- Interpretability is logically independent of causal and predictive properties

§7.4

The AI validity challenge

When the model's own outputs become part of the evidence.

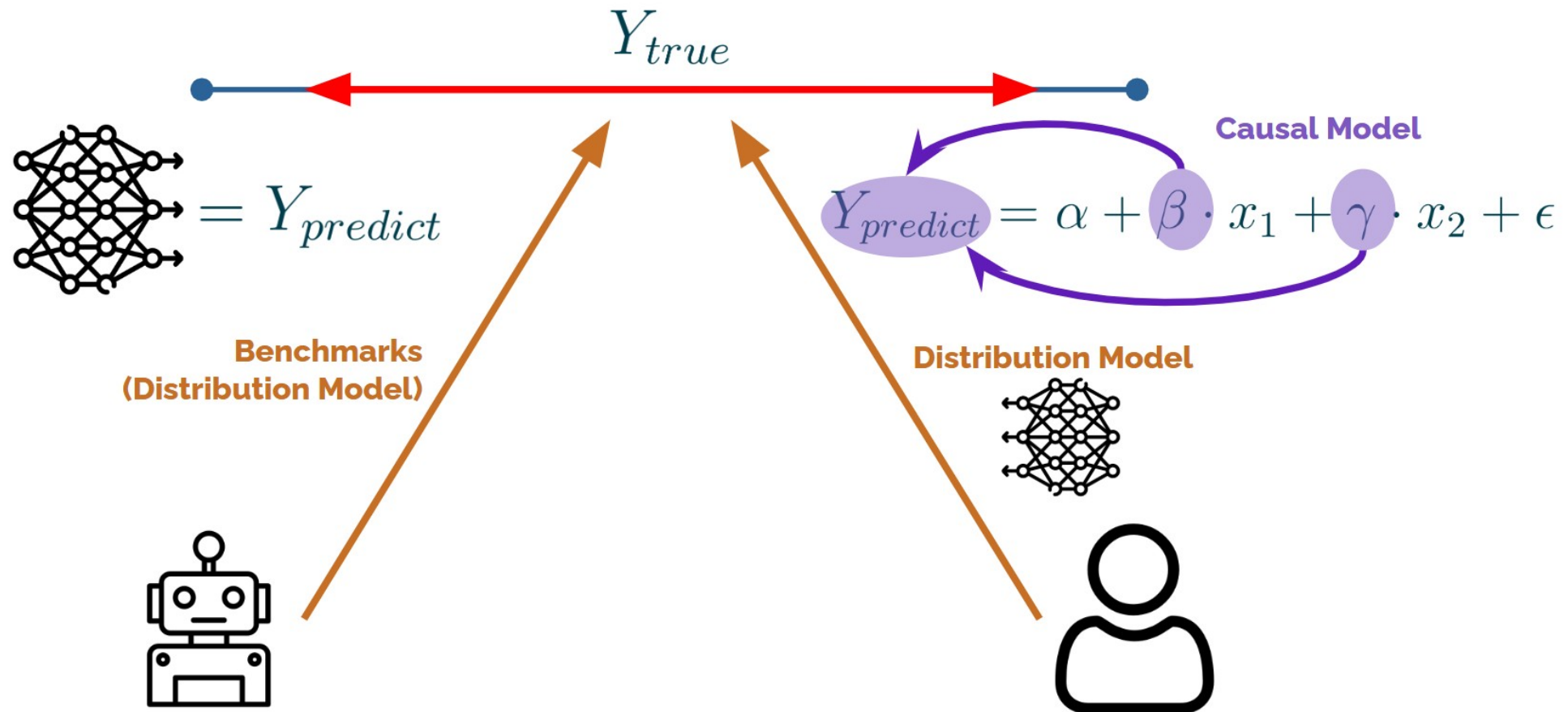
Two AI roles that fail differently

§7.4

- Coding real data: the evidence is still about humans; coding quality is at issue
- Cross-prompt consistency can reflect shared training bias, not accuracy
- Generating data: model behavior is offered as evidence about human behavior
- The second claim is stronger — and fails in documented ways

Two validation logics for AI-assisted studies

§7.4



Benchmark evaluation asks whether outputs match targets on a task; causal-model evaluation asks whether behaviorally relevant relationships hold up under change.

Two validity claims, and neither implies the other

§7.4

- External sociological validity: does model behavior correspond to human subjects?
- Internal representational validity: are constructs meaningfully encoded and manipulable inside the model?
- A model can mimic aggregates on brittle internal shortcuts
- Clean internal encoding does not make a faithful model of humans

Synthetic respondents: the variance collapse

§7.4

- ANES-based personas answering feeling thermometers matched human averages
- But: too little variation, coefficients differ, some reverse sign
- Results shift with small prompt changes and over time
- Picture two curves: same mean, far narrower synthetic spread

Mixed-subjects designs keep humans the gold standard

§7.4

- LLM predictions treated as informative but imperfect proxies
- Prediction-powered inference: efficiency gains without assuming predictions are correct
- Moral Machine reanalysis: silicon sampling alone was biased and overconfident
- 100,000 LLM predictions + 10,000 humans → effective sample $\approx 11,275$

Practice: audit a published computational study

§7.5

- Identify the study's primary epistemic purpose
- Name the computational capacities its method depends on
- Trace how it handles opacity, the semantics gap, and temporal dynamics
- Judge whether its validation matches its purpose
- State what you would change in a replication or extension

You will need

- The CSS Empirical Studies Database, filtered to computational analysis (cssprimer.org/ch07)
- One published study using network analysis, machine learning, embeddings, or another method from this chapter

Deliverable

A one-page analysis answering the five questions, saved to your project repository.

➤ [Studies filtered to computational analysis](https://cssprimer.org/ch07/) · cssprimer.org/ch07/

- Which epistemic purpose do your own analyses usually serve — and does your goal shift as a project moves along?
- Where do you expect the semantics-to-application gap to bite hardest in your work, and how would you bridge it?
- How would epistemic opacity show up in a model you might use, and which remedies would you trust?
- What would you need to make a causal claim from a descriptive or predictive result you admire?

Key concepts from Chapter 7

Iteration: Repeated updating — optimization, simulation, statistical — that learns from empirical feedback at scale

Non-linear modeling: Capturing thresholds, diminishing returns, and interactions via activation functions and layers

High-dimensional processing: PCA, embeddings, topic models compressing many variables into tractable representations

Epistemic opacity: Not being able to know how a model reached its output; limits contestability

Semantics-to-application gap: Linking computational outputs to social concepts requires theory and judgment

Temporal dynamics: Update order, time steps, and stopping rules are part of the model

Descriptive modeling: What does the social world look like? Mapping structure without causal claims

Explanatory modeling: What changes under intervention? Identifying and estimating causal effects

Predictive modeling: Forecasting new observations; useful predictors need not be causal

Integrative modeling: Predicting under change; causal structure that survives distributional shifts

AI validity: External sociological vs. internal representational; neither implies the other

Materials on the companion site

cssprimer.org

[Studies filtered to computational analysis](https://cssprimer.org/ch07/) cssprimer.org/ch07/

The Practice assignment's pre-filtered slice of the studies browser

[Worked notebooks](https://cssprimer.org/notebooks) cssprimer.org/notebooks

The network-analysis notebook with sample data, for students whose corpus is not ready

[Computational-analysis tools](https://cssprimer.org/tools/analysis) cssprimer.org/tools/analysis

The maintained tool list for networks, prediction, and simulation

Where to go next

- Read: Chapter 8 — every concept today has a hands-on exercise on your own data
- Do: the practice audit of one published study from the database
- Browse: the Chapter 7 hub for the studies database and worked examples