

LLMs for Social Scientists

A Working Draft

Yunus Emre Tapan

2026-07-17

Table of contents

LLMs for Social Scientists	3
Preface	4
The arc of the book	4
Who this is for	5
How to read this draft	5
About the author	5
Citing this draft	5
Acknowledgments	5
Roadmap	6
Roadmap	7
Current state	7
How chapters get here	8
References	9
References	10

LLMs for Social Scientists

A Working Draft

Preface

This is a working draft. I am writing this book in the open. Chapters appear here as they become readable, not when they are perfect; each one carries a status badge, and the [roadmap](#) shows what is drafted, what I have taught before, and what is still to come. If something is unclear, wrong, or missing — tell me in the comments at the bottom of any chapter. That is the point of drafting in public.

Social scientists have always studied things we cannot fully see — institutions, norms, beliefs. Now we work alongside a new kind of opaque object: large language models. We use them to code text, summarize interviews, and simulate respondents, usually without knowing what happens between our prompt and their answer.

This book is my attempt to change that, one layer at a time. It grew out of teaching — the BLISS workshops at Northeastern, the AIDE summer bootcamp, and a SICSS-Istanbul lecture on what I call the *anthropology of machines* — and it keeps the spirit of those rooms: everything by hand first, everything connected to a real research task, and honest about what we do and do not know.

The arc of the book

The book moves through four ways of relating to a language model:

1. **Use** — the plumbing. Running open models with Hugging Face, calling frontier models through APIs, and scaling from one prompt to ten thousand documents (*LLMs in Action*).
2. **Steer** — getting better answers without touching a single weight: prompting informed by the evidence rather than folklore, and context engineering — RAG, GraphRAG, and research tools you can put in front of colleagues (*Steering LLMs*).
3. **Read** — mechanistic interpretability: opening the model and observing what happens inside, from the logit lens to sparse autoencoders. This is where social-scientific instincts — coding, interpretation, validity — turn out to matter most (*Reading LLMs*).
4. **Change** — fine-tuning and adapting models for your own research programs (*Changing LLMs*).

Before all that, *Foundations* builds the conceptual floor — neural networks, embeddings, transformers, attention — by hand, on paper, with numbers small enough to trust. And a *Python refresher* gets you set up if you are coming back to code after a while.

Who this is for

Graduate students and researchers in the social sciences who want to use language models seriously: not just prompt them, but understand them well enough to defend a method section. Basic Python helps; no machine learning background is assumed. Every technique in this book is taught against real research tasks — classifying documents, annotating text at scale, building corpora, interrogating what a model represents.

How to read this draft

- **Status badges.** Every chapter opens with its status: `outline` → `drafted` → `taught` → `revised`. Drafted chapters are real but unpolished; taught chapters have survived contact with a classroom.
- **Comments.** Each page has a discussion thread at the bottom (via GitHub). You can also [open an issue](#) or suggest an edit with the links in the sidebar.
- **Code.** Code is shown, not executed, in the rendered book; companion notebooks are linked from chapters as they are released. Code is MIT-licensed; the text is [CC BY-NC 4.0](#).

About the author

Yunus Emre Tapan is a PhD student in Political Science at Northeastern University, studying social networks and political narratives with computational methods. He is the founder of [BLISS](#) — the Boston LLMs Initiative for Social Sciences — and has taught LLM methods at the AIDE summer bootcamp and SICSS-Istanbul.

Citing this draft

Tapan, Y. E. (2026). *LLMs for Social Scientists* (working draft). <https://yemretapan.com/llms4socialscientists>

Acknowledgments

This book owes its existence to the communities I taught it in first: the **BLISS** community at Northeastern, supported by the NULab for Digital Humanities and Computational Social Science; the **AIDE** summer bootcamp; and **SICSS-Istanbul**.

Roadmap

Roadmap

This book is written in public, chapter by chapter. This page is the honest map: what exists, what state it is in, and what is coming. It changes as the book grows.

Status legend

Badge	Meaning
<code>outline</code>	Structure and sources decided; prose not yet written
<code>drafted</code>	Full draft exists; awaiting my revision — read with that in mind
<code>taught</code>	Material has been taught in a workshop or lecture at least once
<code>revised</code>	Revised after teaching and feedback; stable
<code>planned</code>	On the roadmap; tell me in the comments if you want it sooner

Current state

Part	Chapter	Status	Taught at
Python Refresher	Python for Text Analysis	<code>drafted</code>	AIDE 2025
Foundations	Neural Networks by Hand	<code>drafted</code>	AIDE 2025, BLISS
Foundations	Embeddings: Concepts as Vectors	<code>drafted</code>	AIDE 2025, SICSS-Istanbul 2026
Foundations	Transformers & Attention by Hand	<code>drafted</code>	AIDE 2025, BLISS
LLMs in Action	Open Models with Hugging Face	<code>drafted</code>	AIDE 2025
LLMs in Action	Working with APIs	<code>drafted</code>	AIDE 2025
LLMs in Action	Batch Processing at Research Scale	<code>drafted</code>	AIDE 2025
Steering LLMs	Prompting, Informed by the Literature	<code>drafted</code>	—
Steering LLMs	Retrieval-Augmented Generation	<code>drafted</code>	AIDE 2025
Steering LLMs	GraphRAG: Adding Structure	<code>drafted</code>	AIDE 2025

Part	Chapter	Status	Taught at
Steering LLMs	Building Research Tools	drafted	BLISS
Reading LLMs	Why Interpretability? An Anthropology of Machines	drafted	SICSS-Istanbul 2026
Reading LLMs	The Logit Lens	drafted	SICSS-Istanbul 2026
Reading LLMs	Sparse Autoencoders & Features	drafted	SICSS-Istanbul 2026
Reading LLMs	Activation Patching	planned	—
Reading LLMs	Circuits	planned	—
Changing LLMs	Fine-Tuning for Research	planned	—

*(Chapters marked **drafted** are being revised for publication here — they appear in this public version as they are ready.)*

How chapters get here

Each chapter follows the same path: I teach the material somewhere → the teaching materials become a draft → the draft is revised in my own words → it is published here with comments open → your feedback shapes the revision. If a **planned** chapter matters to your research, say so in the comments below — it genuinely affects what I write next.

References

References

The bibliography is being rebuilt as chapters are revised: while this book is a working draft, most chapters cite their sources inline (with links) and close with a curated *Further reading* list. Entries below are collected automatically from formal citations as they are added.