
Research Methodology for Public Policy

Lecture 10 — Qualitative Methods 1

Dr. Colin Kuehl

2026-07-21



Northern Illinois
University



**SCHOOL OF
PUBLIC POLICY**
CHIANG MAI UNIVERSITY

Welcome Back & Agenda

Quick Recap: Lectures so far

- Lectures so far one feature: N is large, and the payoff is a precise *average* effect
- Today we talk about small N and processes in single cases

Today's Agenda

- Why and when qualitative, small- N methods are the right tool
- Case selection logic: most similar, most different, typical, deviant, and more
- Interviews: structure, recruitment, and access
- Surveys
- Participant observation and ethnography
- Triangulation, bias, and practical fieldwork guidance
- A preview of where Lecture 11 picks up (mixed methods)

Why Qualitative, Small- N Methods?

Two Traditions, One Goal

- Quantitative methods (Lectures 1–9): explain variation across many cases, using numbers
- Qualitative methods: understand process, meaning, and mechanism — often within one or a few cases
- Same underlying goal: rigorous, defensible inference about the social and political world
- Different tools for different questions — not a hierarchy

Effects of Causes vs. Causes of Effects

- Quantitative research, especially causal inference, focuses on the **effects of causes**: does treatment X move outcome Y , on average, across many units?
- Small- N qualitative research flips the question: it focuses on the **causes of effects** — how did *this particular outcome* come to be?
- Neither framing is more rigorous than the other — they answer different questions about the same social world
- A good small- N theory still has to do real work: identify the cause and its implications, explain the outcome across the full scope of cases it claims to cover, and specify whether the cause was necessary, sufficient, or merely contributory

What Qualitative Methods Add

Quantitative Strengths

- Generalizability across many cases
- Precise estimates of average effects
- Replicable, transparent procedures

Qualitative Strengths

- Depth: *how* and *why*, not just *whether*
- Uncovering mechanisms and process
- Discovering concepts you didn't know to look for

Telling Stories Numbers Miss

- Small- N research can be **inductive** (building theory from the data) or **deductive** (testing an existing theory against the data)
- Researchers have used small- N methods to surface stories that large- N data sets miss almost by construction
- Examples from the discipline: rural political identity in Wisconsin, the secondary marginalization of the most disenfranchised within a community, the political incorporation of Afro-Caribbean immigrants
- A common thread: each used in-depth, small- N fieldwork to challenge or complicate an existing large- N theory

When to Reach for Qualitative Tools

- Your research question is “how” or “why,” not just “how much”
- You’re at an early stage of theory-building, not theory-testing
- N is small by necessity — rare events, elite access, single cases of major importance
- You need context, meaning, or process that numbers alone can’t capture
- Contrast: Lecture 9’s experiments isolate *whether* a treatment works; small- N work explains *how* it worked, for whom, and why

Designing Qualitative Research

Research Questions Suited to Qualitative Inquiry

- Good qualitative questions are usually open-ended and process-oriented: “How did X happen?” “Why did actors choose Y over Z ?”
- Same standards from Lecture 3 still apply — feasible, interesting, ethical — but “novel” and “testable” look a little different here
- Often more inductive: the theory and the question can evolve as you learn more about the case

The Case Study Logic

- A **case** is a bounded instance of the phenomenon you care about — a country, a policy episode, an organization, a decision
- Case studies trade breadth for depth: one (or a few) cases studied intensively, instead of many cases studied thinly
- This is a deliberate design choice, not a “lesser” version of a large- N study
- The central design problem is not “how do I analyze this case” — it’s **which case(s) should I even pick?**

Case Selection Logic

Purposive, Not Random

- Qualitative research rarely uses random sampling — the goal isn't statistical generalization, it's analytical insight
- **Purposive sampling:** deliberately select cases because of what they can teach you, given your question
- **Snowball sampling:** existing participants refer you to others — useful for hard-to-reach populations
- Case selection is “setup to use a small number of cases in order to go into a deep dive into a specific subject.” (*EMPS Ch. 9*)

A Menu of Case Selection Strategies

EMPS Ch. 9 (drawing on Seawright and Gerring 2008) organizes case selection around the logic you need for your question — we'll walk through four core strategies, then a few extensions.

- Most Similar
- Most Different
- Typical case
- Deviant case

Most Similar Systems Design

- Compare at least two cases that are alike on (almost) every variable except the key independent variable — isolating its effect on the outcome
- “When using most similar cases all independent variables other than the key independent variable or dependent variable would be similar.” (*EMPS Ch. 9*)
- Worked example: two neighborhoods matched on income, religion, and education, differing mainly in racial composition
- The catch: “matching any particular cases by exact characteristics is essentially impossible in the social sciences.” (*EMPS Ch. 9*)

Most Different Systems Design

- Compare cases that differ on nearly everything *except* the key independent variable and the outcome
- Logic: if the same X – Y relationship shows up despite all those differences, it's probably not an artifact of context
- EMPS Ch. 9 is candid about its weaknesses: this “tends to be the weaker route to take in comparing two cases,” though still useful under the right circumstances
- Useful for arguing a finding is *generalizable*, not just a local quirk

Mill's Methods, Underneath It All

- Most-similar and most-different designs trace back to J.S. Mill's “method of difference” and “method of agreement”
- The logic mirrors everything else this course has covered: hold things constant (or vary them) to isolate what's actually driving an outcome
- Easy to state, hard to satisfy in practice — truly “similar” or truly “different” cases are rare, and the researcher must be honest about the compromise

Case Selection Logic, Side by Side

Design	Key Advantages
Most Similar Systems Design	- Easier control over variables
	- Focused comparison allows for identifying key differences
	- Can reveal subtle but crucial differences between cases
	- Enables deeper and more detailed analysis of a few cases
	- Effective for testing hypotheses and theories in controlled environments
Most Different Systems Design	- Useful for identifying universal factors and generalizable explanations
	- Encourages diverse case selection, offering broader insights
	- Can discover robust causal mechanisms that work across different contexts
	- Overcomes regional or cultural biases in research
	- Effective for challenging or confirming existing theories across a wide variety of contexts

- Most Similar buys you control and depth; Most Different buys you breadth and generalizability
- Notice the tradeoff is the same one from Lecture 2: internal validity vs. external

...And the Costs of Each

Design	Key Disadvantages
Most Similar Systems Design	- Overemphasis on similarities, limiting scope
	- Limited generalizability
	- Difficulty controlling for all similarities
	- Risk of omitted variables
	- Complex causality can obscure clear findings
Most Different Systems Design	- Risk of oversimplifying complex phenomena
	- Difficulty identifying causal mechanisms
	- Challenges in selecting cases
	- Problem of equifinality (different causes, same outcome)
	- Limited understanding of context-specific factors

- Most Similar risks omitted variables and limited generalizability
- Most Different risks **equifinality**: different causal paths producing the same outcome, which can obscure the mechanism you're looking for

Typical Case

- A **typical case** is representative of a broader category — it illustrates how things “normally” work
- “The typical case should be well defined by an existing model which allows for the researcher to observe problems within the case rather than relying on any particular comparison.” (*EMPS Ch. 9*)
- Great for confirming *or* disconfirming an existing theory — you don’t need a comparison case if the existing model already tells you what to expect
- The payoff depends entirely on showing the case truly is representative of the category you claim it belongs to

Deviant Case

- A **deviant case** cannot be explained by existing theory — it's the anomaly, not the rule
- Deviant cases function as exploration: they check an established theory for blind spots and can surface previously unidentified explanations
- EMPS Ch. 9 frames the payoff as conceptual: deviant cases let you show a population once assumed to fit a general pattern is “more complex than previously understood”
- You can have one deviant case or several — the point is variation the dominant theory doesn't predict

Other Selection Approaches, and How to Choose

- **Extreme case:** very high or very low on your key variable — maximizes variation on the dimension you care about, which (unlike in regression) is often a *feature*, not a bug, for small- N work
- **Diverse cases:** deliberately pick low/medium/high values to illustrate the full range — good for generating new hypotheses
- **Influential case:** an outlier that may be driving a large- N result — but EMPS Ch. 9 notes this category mostly matters in *large- N* diagnostics, not small- N design
- Choosing among all of these starts with your question: Most Similar / Most Different test *comparisons*; Typical / Deviant test *fit to an existing theory* — and every choice has to be defended in writing, not treated as a footnote

Collecting Qualitative Data: Interviews

Setting Up a Small- N Project

- EMPS Ch. 9 frames small- N design as a chain: research question → goals → conceptual framework → method → validity check
- Field notes are described as “a researcher’s best friend” — demographics, noises, emotions, norms, anything that situates the data
- The chapter is blunt about the time cost: it is “not uncommon for researchers to eventually live in the places they are studying”

Interviews: The Spectrum

- Structured, semi-structured, and unstructured interviews — a spectrum of how much you let the conversation lead
- Semi-structured is the common default in policy research: a guide of core questions, with room to follow up
- EMPS Ch. 9 distinguishes the **structured interview** (a rigid set of questions, the same for every participant) from the **semi-structured interview** (a flexible strategy that allows the researcher to deviate within the confines of the research question)

Building the Interview Guide

- Decide your structure (structured vs. semi-structured) *before* you contact anyone
- The interview guide lays out the core questions and themes — it's your map, even when you depart from it
- Two categories of people you might want to talk to:
 - **Key informants:** experts who can speak *about* the population — academics, community leaders, politicians
 - **Representative sample:** people who lived the experience directly, giving a firsthand account

Key Informants vs. Representative Sample

“What is important to understand about the difference between the representative sample and key informants is that the sample is giving a firsthand account of their experiences, while a key informant is mainly given their observation and experiences of the representative sample from an outside perspective.” (EMPS Ch. 9)

- Both are legitimate, but they answer different questions — don't substitute one for the other without noticing
- A strong design often uses both: informants to understand context, the sample to understand lived experience

Recruitment and Access

- Recruitment is hard, and there's no single trick that works every time
- Factors that shape whether someone agrees to participate: occupation, region, retirement status, vulnerability, and sponsorship from within their own network
- **Snowball sampling**: ask current participants to refer you to others in their network — often the most effective tool for hard-to-reach populations
- Face-to-face contact and prompt follow-up tend to outperform cold emails or surveys

Logistics That Actually Matter

- Bring: two recorders (one is a backup), tissues, your interview guide, a consent form, and a small token of appreciation if you can manage it
- Keep the interview to about an hour — a sign of respect for your participant's time
- Take meticulous notes *during* the interview, not just after
- Always ask permission for a possible follow-up before you leave

Elite and Practitioner Interviews

- Elite interviews — with officials, practitioners, senior staff — raise distinct access and positionality challenges
- Power runs the other direction: the participant may control what gets said, and may be skilled at managing interviewers
- Be alert to what an elite interviewee has an incentive to emphasize, omit, or spin
- Cross-checking elite accounts against documents or other interviews is not optional — it's the standard

Survey Research - Both Qual and Quant

What is Survey Research?

Why Surveys?

A lot of political science is ultimately about **people** — their opinions, preferences, and behaviors.

- How do Thais feel about military rule? What predicts support for welfare spending? Does trade exposure affect protectionist attitudes?
- These are fundamentally questions about **individuals** that require individual-level data

Surveys provide a direct window into public opinion at scale — something no other single method does as efficiently.

Definition

Survey research is a quantitative and qualitative method with two key characteristics: (1) variables of interest rely on respondents' *self-reported* measures, and (2) *sampling* plays a central role in connecting results to a broader population. (Clipperton et al.,

A Brief History: Three Eras (Groves 2011)

Era of Invention

1930–1960

- Survey design basics established
- Face-to-face and mail surveys
- Likert scale invented
- Major polling firms founded (Gallup)

Era of Expansion

1960–1990

- Quantitative social science booms
- Federal funding increases
- CATI (computer-assisted telephone interviewing)
- Telephone directories as sampling frames

1990–Present

- Mobile phones displace landlines
- Caller ID reduces contact rates
- Internet opens new channels
- Online panels and MTurk
- Lower response rates, new challenges

Each era adapted to technological and social change — the core logic of survey research stayed constant; the *tools* evolved.

The Four Stages of Survey Research

Every survey research project — from a Gallup poll to a PhD dissertation — follows the same basic structure:

1. Develop the Survey → 2. Sample → 3. Field the Survey → 4. Analyze Results

Each stage has its own challenges and potential sources of **error**. A mistake at any one stage can compromise everything downstream.

Today we walk through each stage carefully.

Conceptualisation & Operationalisation

Stage 1: Developing the Survey

Before writing a single question, answer these:

- **What** is the purpose of the survey? What are we trying to learn?
- **Whose** opinions are we measuring? Who is our target population?
- **How** can we best ask them?

Two foundational concepts:

Conceptualisation

Clearly defining the key variables of interest. What do we mean by “democracy,” “corruption,” or “support for trade”? Different definitions lead to different conclusions.

Operationalisation

Turning an abstract concept into something measurable. For example, operationalising *democracy* as a binary (0 = non-democracy, 1 = democracy) or as a continuous index

Example: Operationalising Democracy

Concept: Democracy

Operationalisation	Variable type	Captures
Presence of regular elections	Binary (0/1)	Minimal procedural definition
Freedom House score	Ordinal (1–7)	Civil liberties + political rights
V-Dem Liberal Democracy Index	Continuous (0–1)	Checks, balances, participation
Self-reported satisfaction with democracy	Likert (1–5)	Subjective citizen assessment

The *same* concept operationalised differently can produce substantially different results.

Always justify your measurement choice and acknowledge its limits

Writing Good Survey Questions

Three Criteria for Good Questions

After conceptualising and operationalising your variables, you write the questions. Poor question wording is one of the most common — and most avoidable — sources of measurement error.

Three criteria (EMPS Ch. 6):

- **Understandable** — all respondents interpret the question the same way. Avoid jargon; if you must use technical terms, define them. Ask one thing at a time.
- **Clear** — simple, plain language; specific enough that only one type of answer fits. *“Do you approve of how the President is handling the economy?”* is clearer than *“Is the President doing a good job?”*
- **Answerable** — avoid hypothetical or counterfactual questions. Respondents cannot reliably report on things they haven’t experienced or haven’t thought about.

Common Pitfalls to Avoid

Double-barreled questions

Ask only one thing at a time.

Bad: “Do you support free trade and open immigration?”

Better: Ask two separate questions.

Leading questions

Neutral wording captures real opinion;
loaded wording pushes respondents toward
a predetermined answer.

Bad: “How damaging do you think the government’s wasteful spending has been?”

Better: “How do you evaluate the government’s recent spending decisions?”

Social desirability bias

Respondents give answers that seem socially acceptable rather than truthful.

- Church attendance is consistently *over-reported* in surveys (Hadaway et al. 1993)
- Sensitive sexual or health behaviors are routinely *under-reported*

Mitigation strategies:

- Place sensitive questions at the *end* of the survey
- Emphasise anonymity and confidentiality
- Use indirect or list experiments for

Question Order and Wording Effects

Small differences in wording or order can produce large swings in results:

Pew Research Center — Iraq War (EMPS Ch. 6)

Wording A: “Would you favor or oppose taking military action in Iraq to end Saddam Hussein’s rule?” → **68% favor, 25% oppose**

Wording B: Same question + “even if it meant U.S. forces might suffer thousands of casualties?” → **43% favor, 48% oppose**

Order effects work similarly — the question asked *first* frames how respondents interpret everything that follows (Schuman & Presser 1996).

Always pilot your survey. Read every question aloud. Ask: could this be misread? Does the order prime a particular answer?

Sampling

Population vs. Sample

It is usually impossible to survey an entire **population** — so we draw a **sample** and use it to make inferences about the population.

Key concepts

- **Population:** the entire group we want to understand (e.g., all Thai adults; all registered voters in Chiang Mai)
- **Sample:** a subset drawn from the population
- **Sampling frame:** the list of all units that *could* be sampled — ideally matches the population exactly
- **Coverage error:** when the sampling frame does not perfectly match the target population

The 1936 Literary Digest disaster

The Digest mailed 10 million surveys and received 2.4 million responses — and still got the election badly wrong. The sampling frame (telephone directories and car registration lists) systematically excluded poorer Americans, producing a sample that was wildly unrepresentative of the electorate.

Size is no substitute for representativeness.

Probability Sampling Methods

Method	How it works	Best when
Simple Random Sampling (SRS)	Every unit has equal probability of selection	Population is homogeneous; full list available
Stratified Sampling	Divide into homogeneous strata; randomly sample within each	Key subgroups must be represented (e.g., race, region)
Cluster Sampling	Divide into heterogeneous clusters; randomly sample <i>clusters</i>	Full population list unavailable; geographically dispersed

- As a rule of thumb, $n \approx 1,000$ is sufficient for national-level inference about American (or Thai) adults
- Oversampling of minorities is common in studies of racial/ethnic politics (e.g.,

Non-Probability Sampling

Sometimes probability sampling is infeasible. Researchers then turn to **non-probability samples**, most commonly **convenience samples** — whoever is easily accessible.

Common convenience samples

- University subject pools (political science undergraduates)
- MTurk and Prolific online panels
- Volunteers responding to advertisements

These samples tend to be **younger, more educated, more politically aware**, and (in the U.S.) more liberal than the general population.

When is it acceptable?

- **Pilot testing** — checking whether questions work before fielding on a representative sample
- **Experimental studies** — internal validity depends on *random assignment*, not sample representativeness (recall the experiments lecture)
- **Limit claims accordingly** — results from a student sample should not be generalised to all adults

Fielding the Survey

Survey Modes: Trade-offs

Mode	Response Rate	Anonymity	Cost	Limitations
Face-to-face	Highest	Low	High	Interviewer effects; logistically complex
Telephone	Moderate	Moderate	Moderate	Not suitable for long surveys; declining landline coverage
Online	Variable	High	Low	Computer

Analyzing Survey Results

From Sample Statistics to Population Parameters

We want to estimate **population parameters** (e.g., true public support for a policy) from **sample statistics** (e.g., the percentage in our sample who support it).

Because samples involve random variation, any estimate will carry uncertainty:

$$\text{Confidence Interval} = \hat{p} \pm \text{Margin of Error}$$

- **Margin of error** quantifies sampling variability — larger samples produce smaller margins
- Political scientists conventionally use **95% confidence intervals**
- A margin of error of $\pm 3\%$ on a 1,000-person national survey is typical

Margin of error only captures **random sampling error**. It does not correct for bad question wording, non-response bias, or coverage error — all of which can produce results that are precise but systematically wrong.

Common Sources of Error

Error type	What it is	Example
Sampling error	Random variation from drawing a sample	Unavoidable; shrinks with larger n
Coverage error	Sampling frame $\neq$ target population	Online survey misses elderly non-users
Non-response bias	Non-respondents differ from respondents	Undocumented immigrants opt out of immigration survey
Measurement error	Question wording distorts true attitudes	Leading questions; social desirability
Respondent fatigue	Quality deteriorates in long surveys	Straight-lining; random clicking

Advantages, Disadvantages & Applications

Strengths of the Survey Method

- **Direct measurement** of individual-level opinions, preferences, and behaviour — exactly what micro-level theories predict
- **Versatility**: can measure attitudes, values, demographics, behaviour, and even embedded experiments (survey experiments)
- **Scale**: large samples enable subgroup analysis and generalisable inference
- **Longitudinal potential**: repeated surveys (ANES, GSS, WVS) allow trend analysis over decades
- **Cost-effective** relative to depth: online platforms make large- n surveys cheaper than ever

Major data sources used in political science

ANES (American National Election Studies), GSS (General Social Survey), WVS (World Values Survey), Eurobarometer, Afrobarometer, Pew Global Attitudes Survey, YouGov/CCES

Limitations of the Survey Method

Surveys cannot do everything:

- **Cannot establish causation** on their own — observational data, no random assignment
- **Snapshots** — a survey captures opinion at one moment; attitudes may shift
- **Self-report limitations** — respondents may not accurately recall past behaviour, may not have formed opinions on all issues, and may answer to please rather than inform
- **Shallow depth** — surveys reveal *what* people think but not always *why*. Open-ended questions and qualitative follow-up are needed for depth
- **Large-N, by design** — small, hard-to-reach populations may require different methods entirely

The survey method is powerful for description and correlation at scale. It is the *wrong* tool when you need to understand causal mechanisms, trace decision-making processes, ⁴³

Participant Observation and Ethnography

Participant Observation, Defined

- **Participant observation:** the researcher participates in an organization, community, or group activity *as a member*, rather than only watching from outside
- “It involves a researcher immersing themselves within a community... [and requires] that the researcher build a strong bond of trust with the observed community.”
(*EMPS Ch. 9*)
- Originally an anthropological method, now widely used across political science and policy research

The Active/Passive, Overt/Covert Grid

- Two design choices, usually made with an IRB: how *active* is your participation, and how *overt* is your presence?
- Passive + covert: the community may have no idea they are being studied
- Active + overt: the community knows who you are and why you're there — but your presence can itself change what you observe
- Neither option escapes the basic tension: visibility changes behavior; invisibility raises ethical questions

Covert Observation in Practice



- Covert observation buys you access to behavior people might suppress if they knew they were being watched
- It also raises the sharpest ethical and consent questions of any method in this course — IRB review is not a formality here

A Field-Defining Example

- Mary Pattillo studied the Black middle class by moving into the Chicago neighborhood she would write about in *Black Picket Fences*
- She led the local church choir, joined the community's action group, and coached cheerleading at the local park — full immersion, not a site visit
- Her own account ties this back to her research question: the participant observation was “essential to describing the black middle class” in that setting
- This is the upper bound of commitment — most projects won't go this far, but it shows what the method can buy you

What Participant Observation Is *For*

- Best suited to **theory-building**: it surfaces process and mechanism that a survey or interview alone would miss
- EMPS Ch. 9 recommends pairing it with another method (often interviews), so the community can also narrate its own story in its own words
- The core discipline: meticulous field notes, written promptly, with explicit attention to the researcher's own biases as they form

Observing Policy in Public



- Not every setting requires deep immersion — public meetings, hearings, and agency proceedings are observable without covert access
- Still demands the same rigor: systematic notes on who spoke, what was contested, 49

Ethnography

- **Ethnography** is the broader project of studying a group of people and their lived experience, typically through sustained immersion and relationship-building
- Overlaps heavily with participant observation — the distinguishing feature is the emphasis on documenting culture, meaning, and relationships, not just behavior
- A field-defining example: Robert Vargas's *Wounded City* captures fear, frustration, and empowerment among residents of a Chicago neighborhood negotiating turf wars between gangs, police, and the alderman's office — insight that “cannot simply be surveyed”
- Especially valuable for **underrepresented communities**, where survey response rates are often low and trust in outside researchers may be limited — it documents emotions, attitudes, and relationships that are “sometimes impossible to capture in quantitative work” (*EMPS Ch. 9*)

A Brief Word on Focus Groups and Process Tracing

- Two other small- N tools — **focus groups** (group interviews that surface shared narratives) and **process tracing** (tracing the causal mechanism within a single case) — round out the small- N toolkit
- Lecture 11 takes these up in depth, alongside mixed-methods designs that combine qualitative and quantitative evidence
- For today: know they exist and roughly what problem each solves, so the menu of small- N tools feels complete

Practical Guidance: Bias, Access, and Triangulation

Access and Rapport

- Unlike a public data set, qualitative access usually has to be *earned*: through gatekeepers, sponsorship, reputation, or sustained presence
- Plan for it explicitly in your research design — “I’ll just go interview people” is not a method, it’s a hope
- Snowball sampling, key informants, and institutional affiliations are common ways researchers build a path in
- Rapport can be built through how you present yourself — dress, demeanor, shared markers of identity — but it has to be genuine, not performed; what builds trust with a group of blue-collar workers may differ from what works with college students
- Rapport is a means to better data, not an end in itself — don’t let it become a substitute for systematic note-taking

Reflexivity and the Researcher's Own Bias

- Small- N fieldwork puts the researcher's own perspective directly into the data-collection process — there is no way to fully step outside it
- Discipline against this: constant self-questioning about what you are noticing, what you might be missing, and why
- Keep both mental notes and a notepad; write up observations immediately after leaving the field, while memory is still fresh

Documents as a Complementary Source

Primary Sources	vs	Secondary Sources
A piece of evidence created by someone at the time of the event. Examples include: • Letters • Diaries • Government records • Autobiographies • Artifact • Computer software 		Information created by someone who was not present at an event, after an event happened. Examples include: • Newspaper articles • Textbooks • Biographies • Encyclopedias • Dictionaries • Atlases 

- Government records, organizational archives, media coverage, meeting minutes, personal papers
- Primary sources (created at the time, by someone present) and secondary sources

Triangulation

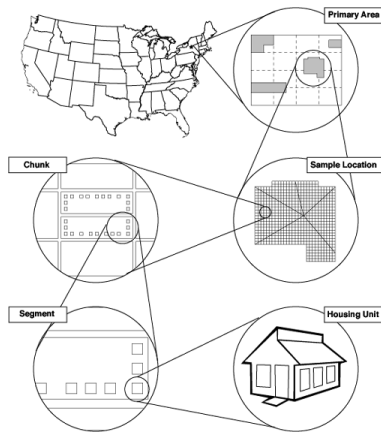
- **Triangulation:** corroborate a finding using multiple data sources or methods, so no single account carries the whole argument
- The disadvantage of small- N work is *not* less data — interviews, field notes, and archives generate plenty — it's a limited pool of *cases*
- “A responsible researcher will have to consider whether their case selection is representative... and how best to ensure that they are getting a diverse set of voices.” (*EMPS Ch. 9*)
- Combining small- N methods with each other, or with quantitative evidence, is the standard way to strengthen the resulting argument — a preview of Lecture 11

Why Large Samples Don't Solve Every Problem



- The 1948 “Dewey Defeats Truman” headline followed a famously wrong prediction, built on a poll sample that systematically missed certain voters
- The lesson generalizes: a large N does not protect you from a *bad* sample, and a small N is not automatically a *weak* one — what matters is whether the cases

Sampling Logics, Contrasted



- Large- N , survey-style sampling (Lecture 6) works through stages — area, segment, household — to approximate a random draw from a population
- Small- N case selection works the opposite direction: not “give every unit a

Key Terms & Looking Ahead

Key Terms from Today

- Purposive / snowball sampling
- Most similar / most different design
- Typical / deviant case
- Key informant
- Participant observation
- Overt / covert observation
- Ethnography
- Triangulation

Before Next Session

Lecture 11: Qualitative Methods 2 — Mixed Methods

- Read: QRM, Ch. 4–6 (sampling and data collection), if not yet finished
- Read: EMPS, Ch. 9 (*Small-N*), remainder — focus groups and process tracing
- Bring a research question you're considering for your final project — we'll workshop what a mixed-methods design for it might look like

Discussion

Take a policy outcome you'd normally study with regression. Which case-selection strategy from today — most similar, most different, typical, or deviant — would teach you the most about it, and why?

Questions?

Thank you — see you in the Practicum this afternoon.

See You in the Practicum

