

DISSC

1

Toward More Reproducible Research: Why and How

September 30, 2025

Yale DISSC-Library-ISPS ORR series

Limor Peer

Anthony Lollo

Maurice Dalton

Yale Data-Intensive Social Science Center



Open & Reproducible Research

Welcome & introductions

- Introduction:
 - About the ORR series
 - Reproducibility: What and why
- Reproducibility: How



[Anthony Lollo](#)



[Maurice Dalton](#)

Upcoming ORR events

DISSC EVENT

Tips from a Data Archive: Preparing and Working with Replication Packages

TUE OCT 21, 2025

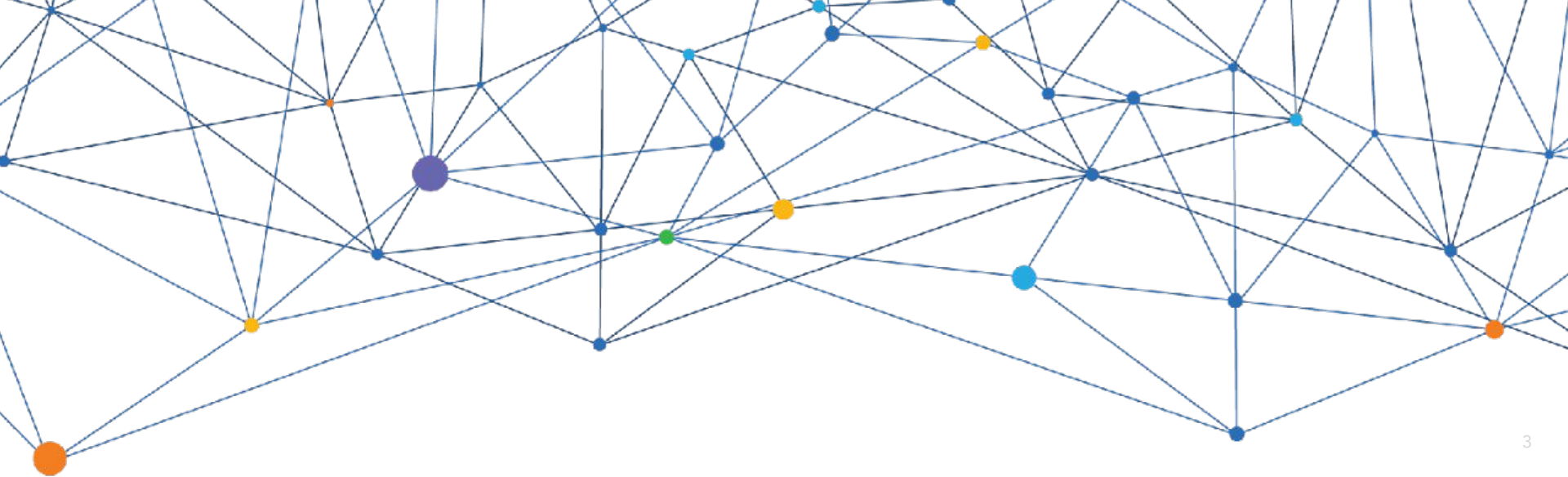
DISSC EVENT

Basics of Research Data Management

TUE NOV 18, 2025

ONLINE

Yale *Data-Intensive Social Science Center*
Open & Reproducible Research

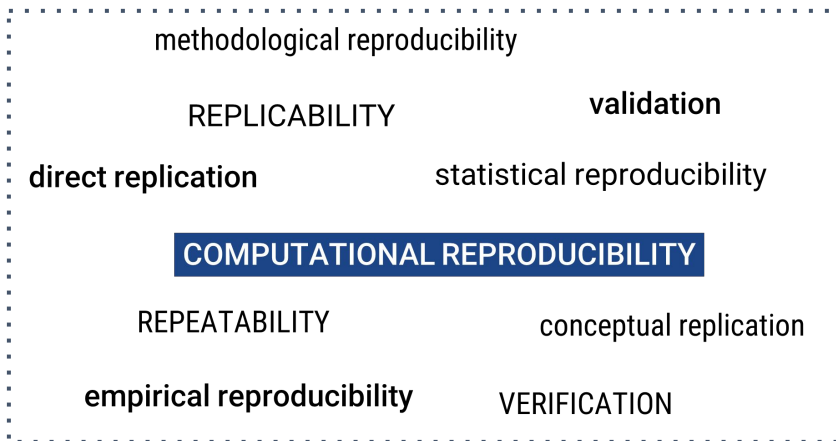


On reproducibility (Limor Peer)

Reproducibility: Definition

Reproducible: A result is reproducible when the *same* analysis steps performed on the *same* dataset consistently produces the *same* answer.

This may be confusing...



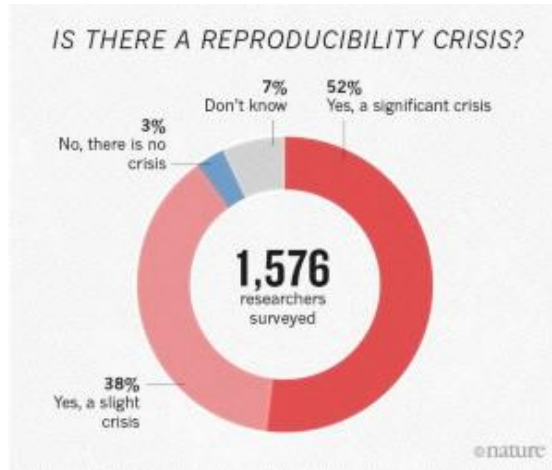
We follow the [The Turing Way](#)

		Data	
		Same	Different
Analysis	Same	Reproducible	Replicable
	Different	Robust	Generalisable

4

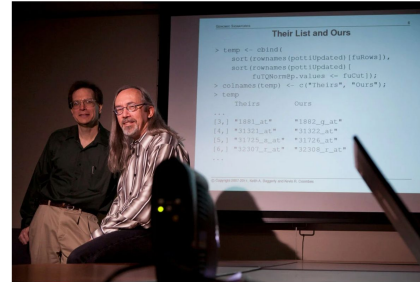
Why reproducibility?

1. Ethical and credible science



Baker, M. (2016). 1,500 scientists lift the lid on reproducibility. Nature, 533(7604), 452–454. <https://doi.org/10.1038/533452a>

How Bright Promise in Cancer Testing Fell Apart



Keith Baggerly, left, and Kevin Coombes, statisticians at M. D. Anderson Cancer Center, found flaws in research on tumors. Michael Stravato for The New York Times

By Gina Kolata

July 7, 2011

When Juliet Jacobson was terrified, but her treatment didn't work, she entered a rare club. Doctors would not fix it.

THE CONVERSATION

Arts | Culture | Economy | Business | Education | Environment | Energy | Ethics | Religion | Health | Medicine | Politics | Society | Science | Technology

THE REINHART-ROGOFF ERROR – or how not to Excel at economics

April 22, 2015 4:40pm EDT

Data and computer code should be made publicly available at an early stage – or else.

Authors: Jonathan Borwein (Liverpool John Moores University), David H. Bailey (P.O. Box 240000, University of California, Davis)

Last week we learned a famous 2010 academic paper, relied on by political big hitters to bolster arguments for austerity cuts, contained significant errors; and that those errors came down to misuse of an Excel spreadsheet.

Sadly, these are not the first mistakes of this size and nature when handling data. So what on Earth went wrong, and can we fix it?

Trusting data provided by Francesca Gino, a former advisee and co-author, made me complicit.

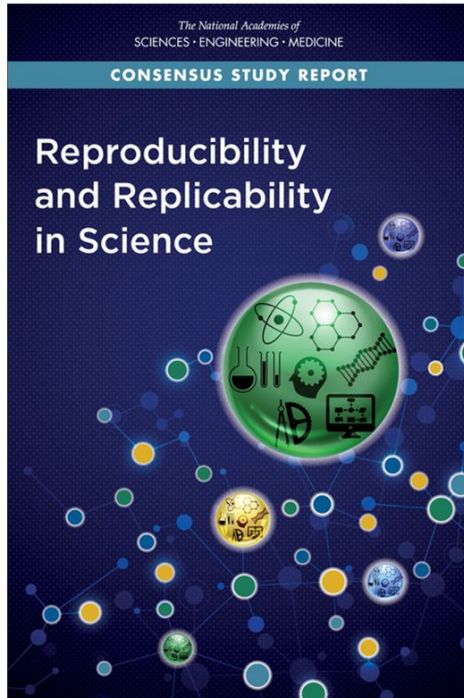


September 8, 2025

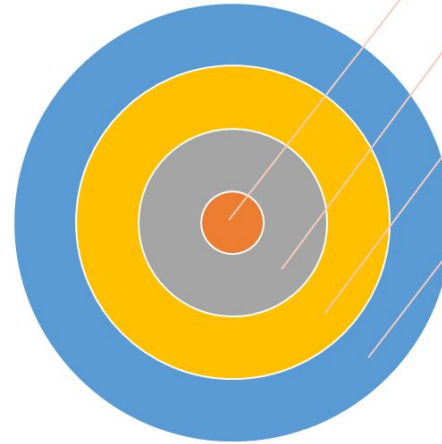
July 15, 2021, I received an email informing me and several co-authors that our well-cited 2012 research paper was fraudulent. The paper claimed to show that you could increase people's honesty by simply having them sign a promise to tell the truth before filling out a form, rather than after, as is traditionally done. Leif Nelson, Joe Simmons, and Uri Simonsohn planned to go to the fraud on Data Colada, their well-respected academic blog. They offered us an to respond.

Open & Reproducible Research

Science is self-correcting



Questions about the reliability of research results that have been raised across all scientific disciplines...



1. Are the data and analysis laid out with sufficient transparency and clarity that the results can be checked?

2. If checked, do the data and analysis offered in support of the result in fact support that result?

3. If the data and analysis are shown to support the original result, can the result reported be found again in the specific study context investigated?

4. Can the result reported, or the inference drawn, be found again in a broader set of study contexts?

National Academies of Sciences, Engineering, and Medicine. (2019). *Reproducibility and replicability in science*. National Academies Press. <https://doi.org/10.17226/25303>

Why reproducibility?

2. Scientific norm

- Makes research accessible and inclusive
- Allows us to accumulate reliable claims about the world
- Helps avoid duplication of effort and waste of resources
- Helps identify new areas of research

Based on Marwick <https://doi.org/10.17605/OSF.IO/MJ4K3>



Image credit:
<https://www.acm.org/publications/policies/artifact-review-badging>

Why reproducibility?

3. Good practice

- Improve your productivity
- Verify your own results
- Enables others to extend your work
- Verify/disprove other's results
- Survive the tech evolution
- Enable community maintenance & support

Based on Ivie & Thain <https://doi.org/10.1145/3186266>



Image credit:
<https://book.the-turing-way.org/reproducible-research/overview/overview-benefit>



Why reproducibility?

4. Increasingly required by funders and journals

What is Gold Standard Science?

As detailed in [Executive Order 14303](#), Gold Standard Science refers to science conducted in a manner that is:

- Reproducible.
- Transparent.
- Communicative of error and uncertainty.
- Collaborative and interdisciplinary.
- Skeptical of its findings and assumptions.
- Structured for falsifiability of hypotheses.
- Subject to unbiased peer review.
- Accepting of negative results as positive outcomes.
- Without conflicts of interest.

9

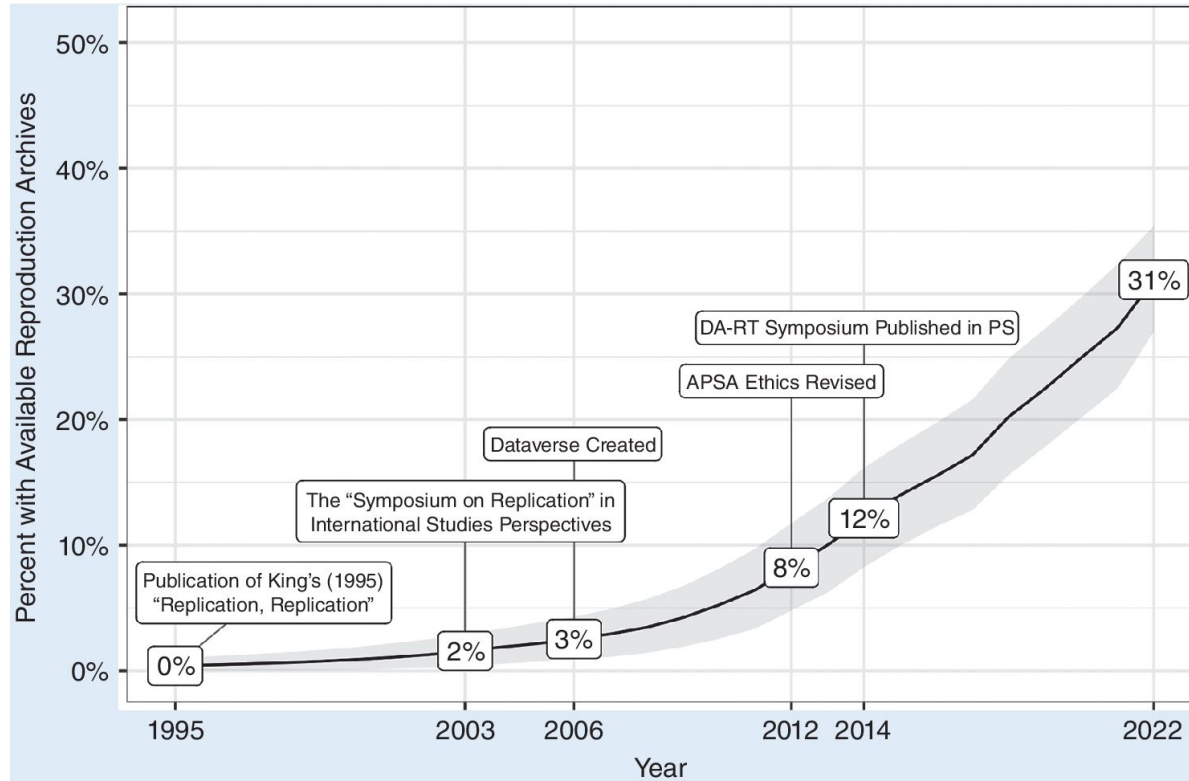
<https://www.nsf.gov/policies/gold-standard-science>

August 22, 2025

Yale *Data-Intensive Social Science Center*

Open & Reproducible Research

Journal data and code (or “replication”) policies



10

Percentage of Quantitative Articles with Reproduction Archives, 1995–2022 (published in political science journals that supply reproduction archives)

Yale Data-Intensive Social Science Center

Open & Reproducible Research



Journal replication policies

Political Science journals guidelines

APSR (<https://apsanet.org/publications/journals/american-political-science-review/guidelines-for-reproducibility/>)

JOP <https://www.journals.uchicago.edu/journals/jop/data-replication>

AJPS <https://ajps.org/guidelines-for-accepted-articles/>

PA

<https://www.cambridge.org/core/services/aop-file-manager/file/678f5f7189ee789cdb3a1df2/replication-guidelines-2025-v1.pdf>

Economic journals guidelines

AEA Data Editor (for all AEA journals)

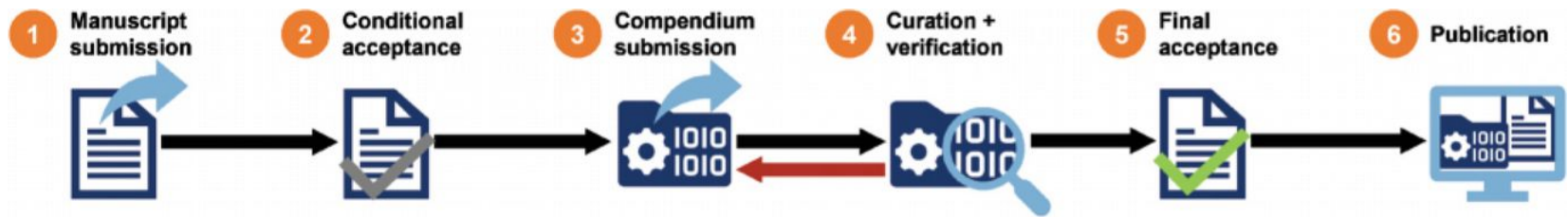
<https://aeadataeditor.github.io/aea-de-guidance/preparing-for-data-deposit.html>

See more...

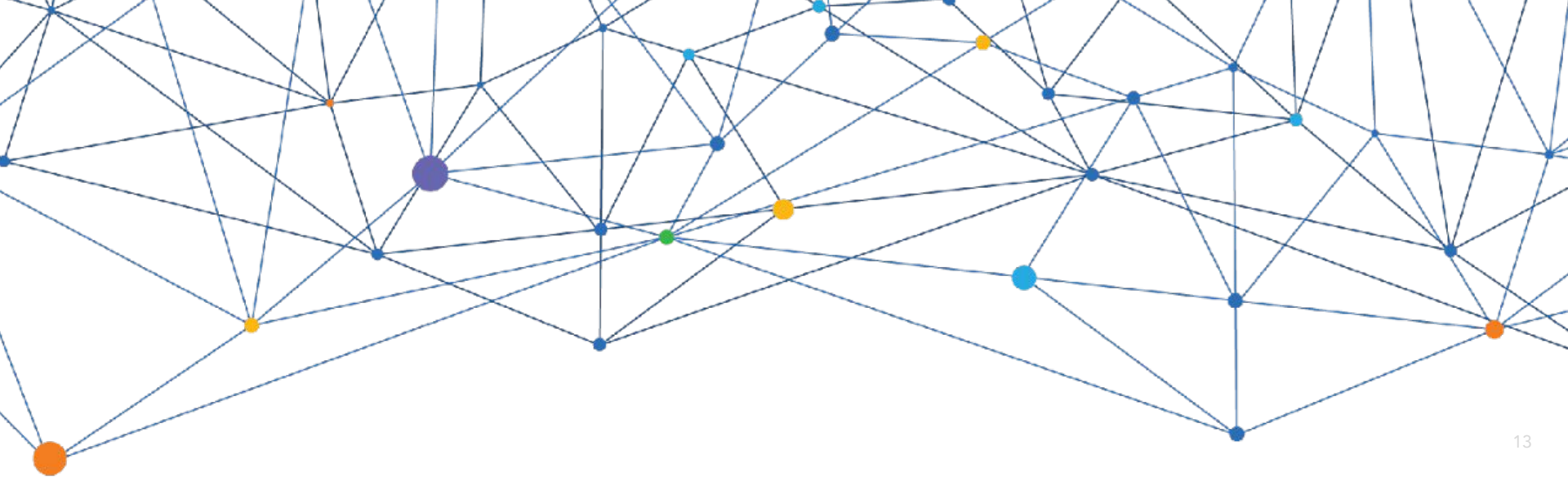
<https://gking.harvard.edu/pages/data-sharing-and-replication> (scroll down to journal policies) Open & Reproducible Research

Yale *Data-Intensive Social Science Center*

Reproducible research publication workflow



12



Example 1 (Anthony Lollo)

The Private Provision of Public Services: Evidence from Random Assignment in Medicaid

14

Danil Agafiev Macambira, Michael Geruso, Anthony Lollo, Chima D. Ndumele & Jacob Wallace

Conditionally accepted at American Economic Review
Data and code deposit accepted

Project Overview

Research Question

What happens when a private firm (*generally for-profit*) provisions the delivery of Medicaid services?

Empirical Strategy

Random assignment of ~100,000 Medicaid beneficiaries across Medicaid plans in Louisiana, 2012-2015.

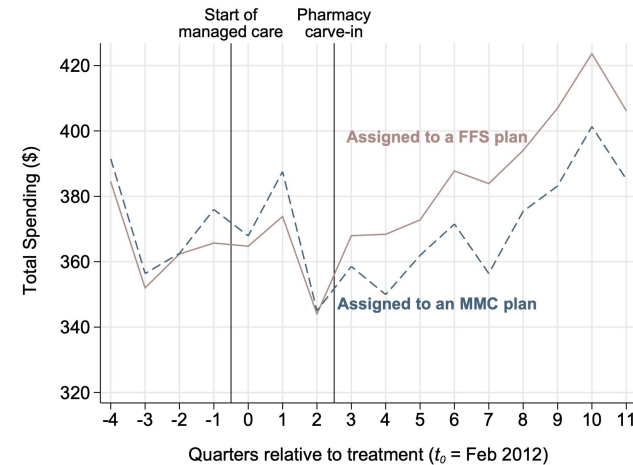
Data

Individual-level Medicaid claims for ~1.5M beneficiaries 2010-2018 (~300M claims). HIPAA Limited data.

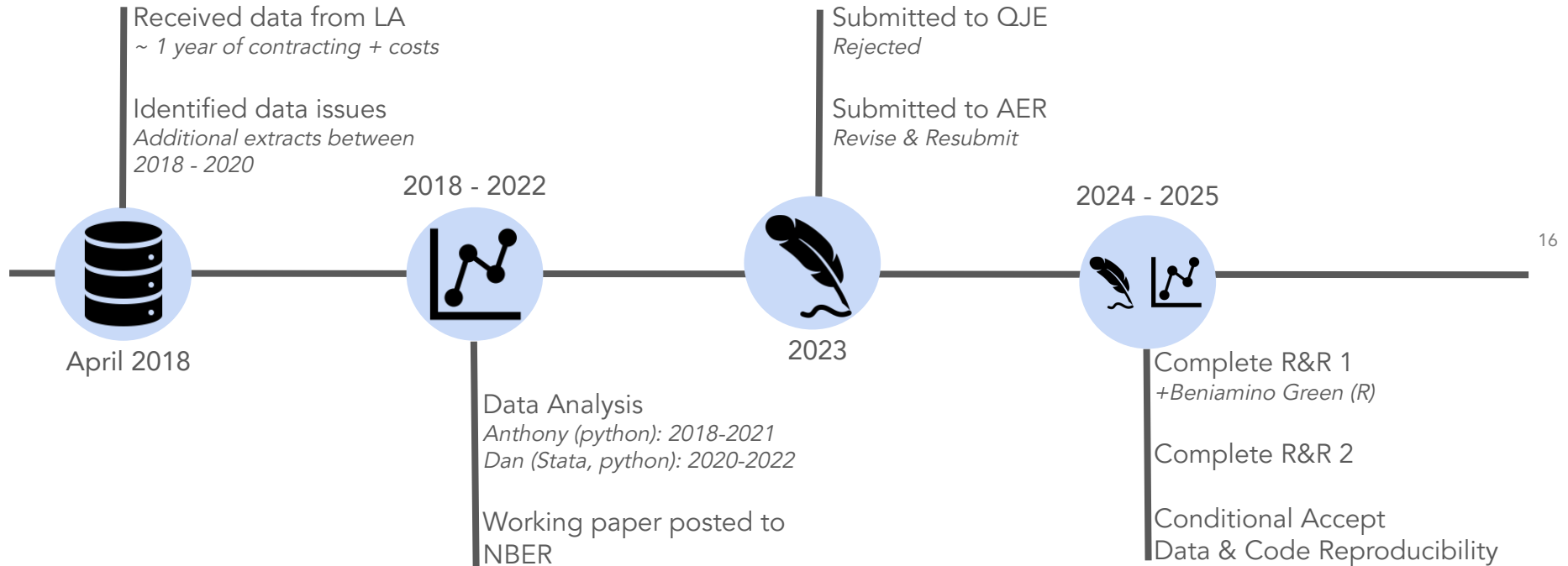
Overall findings

Relative to a government plan, private plans:

- Decrease health care spending, primarily through brand-to-generic substitutions
- Potentially harm quality



Project Timeline: 7 years from data to conditional acceptance



Journal and Reproducibility Considerations

- Broad, general interest topic paired with rich, individual level-data and unprecedented randomization across public and private Medicaid plans.
→ Aiming for a “Top 5” economics journal
- Aware of journal reproducibility requirement
→ Knew all main results needed to be fixed and fully reproducible when posting the Working Paper to NBER. From then on if something changed there needed to be a justification.
Especially true during the two revisions

17

Summary of Challenges for Journal Reproducibility

1. Analyses and revisions took place over 7 years
 - multiple individual contributors
 - multiple programming languages (python, R, Stata)
2. HIPAA-limited data which cannot be shared.
 - For others to obtain the data it would likely require > 1 year of contracting and cost > \$30K.
 - Given government data systems, it's possible the same raw data no longer exists.
3. Accumulation of technical debt after several resubmissions and revisions.
4. A lot of code: 9 figures, 6 tables, 18 appendix figures, 17 appendix tables. > 2 weeks to run fully from start to finish.
5. Spent time standardizing data and building modular analytic libraries to re-use data across projects.
 - Since this step is between the raw data and analytic starting points for this project it all needed to be included in the replication.

18

Data and Code Deposit

Ideal

- A self-contained collection of everything required to reproduce all analyses without any manual steps.

Reality for our project

- HIPAA limited data can't be shared and is too burdensome for others to obtain
 - No tables and figures in the paper can be reproduced by our data and code deposit

19

In these circumstances the Data and Code Deposit:

- Carefully documents how others could obtain the data
 - Who to contact, describes format and cadence of files, provides data dictionaries and crosswalks, explains manual steps in transferring data
- Provides all code and structure required to reproduce all figures and tables *once data is obtained*
- Highlights any known or anticipated replication issues
 - E.g., requested address information will be *most recent*, will not match April 2018 request
- Commits to preserving data for 5+ years and reasonably helping others in replication attempts/data acquisition

Data and Code Deposit

Name
> analysis
> data
> figures
> Healthcare_Data
> ResourceData
> standardization_code
README.md
Files.xlsx
README.pdf
packages-as-installed.txt
analysis.py
create_cols.py
denials_mechanisms.py
main_AA.py
main_DD.py
post_pipeline.py
supplemental_DD.py
wts_mechanisms.py
setup_stata.do
env.yaml
pipeline.yaml

- Folder structure
- All code to perform analyses
 - .py files
 - .yaml files (pipeline build system)
 - .do files (also within folders)
- README
- Environment setup
 - packages-as-installed.txt
 - setup_stat.do
- Metadata entered into deposit archive

Geographic Coverage ?

Louisiana

Time Period(s) ?

1/1/2010 – 12/31/2016 (2010-2016)

Collection Date(s) ?

2018 – 2020

Universe ?

Medicaid beneficiaries in Louisiana between 2010 and 2016

Data Type(s) ?

administrative records data

20

Preparing the Research for Replication

From the start:

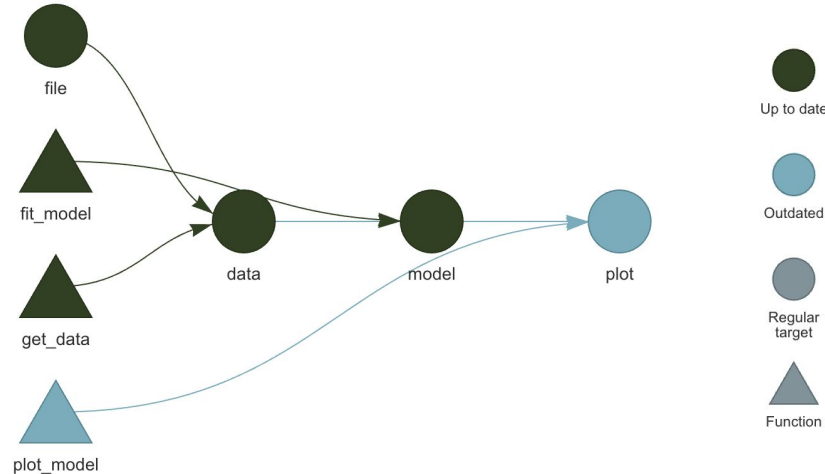
- Prioritized organization and documentation throughout the research project
 - Carved out specific time to refactor code and make codebase more modular
 - Important to strike the right balance between organization and research progress
- Avoided any manual modifications throughout (file moving/renaming. Table formatting!)
- Made sure randomness was controlled (set seeds)

21

Preparing the Research for Replication

Along the way

- Adopted open source tools to manage analysis pipelines (ploomber/luigi – python; targets – R)
 - Ensures that all upstream changes propagate to downstream analyses
 - Encodes knowledge about workflow organization that gets lost when analysts change or over time



22

Preparing the Research for Replication

At the end

- Rewrote all R code in python to completely remove any dependence
- Rewrote most of the stata code in python
 - Minimized I/O and data handoffs
 - Ensured figure/table formatting was consistent
- Created a comprehensive README
- Performed automated and manual checks to ensure no PII/PHI was uploaded

23

```
# Recreate folder structure
import os
inputpath = ...
outputpath = ...

for dirpath, dirnames, filenames in os.walk(inputpath):
    structure = os.path.join(outputpath, dirpath[len(inputpath):])
    if not os.path.isdir(structure):
        os.mkdir(structure)
```

Data and Code Deposit Timeline

When we submitted the first R&R we created most of the replication package

- Editor indicated we should start getting this ready
- Presumed the majority of “large” edits were complete
- Possibility of an additional R&R meant it wasn’t worthwhile to finalize every detail

Once conditionally accepted after 2nd R&R, we finalized the replication package

- Given 1 month to submit the replication

After 2 months we received feedback from the data editor

- Revised and resubmitted in 2 weeks → deposit accepted the following week

24

Feedback from Data Editor

Required revisions were minimal because the replication package was already extensive and well documented.

1. Include data citations and precise access modality
2. Attest that flagged “potential PII” was not PII (it was not)
3. Update data attestation statements to conform to their exact wording
4. Provide data dictionaries and codebooks
5. Update environmental setup
 - a. Stata – indicate package versions
 - b. Python - rename requirements to packages-as-installed and remove OS specific packages

25

Learnings from the Data and Code Deposit

- Keep track of data citations/access modality, incorporate them into the manuscript and replication package
 - Easy to lose track of dates accessed, exact websites visited, and any manual steps
 - It's a lot more work to try to do this 5+ years after the fact
- Stick with descriptive naming, not ordinal naming conventions – revisions will inevitably break the order
 - Explicitly track where all figures and tables are coming from

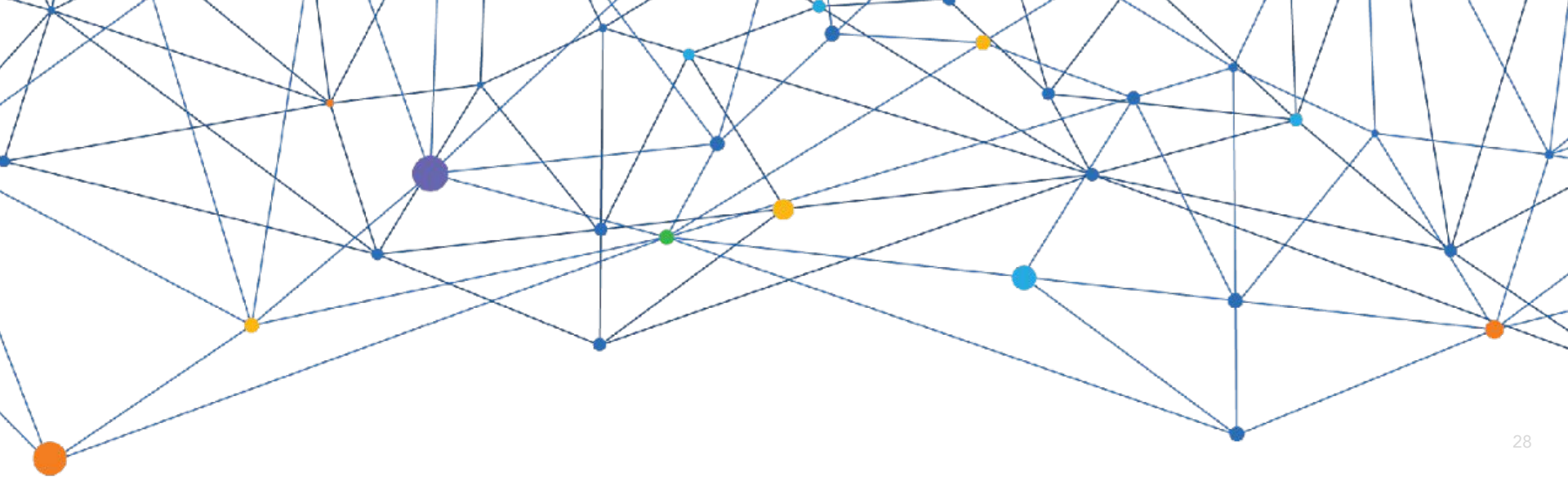
26

Exhibit	Ploumber Pipeline Task	Method File	Method	Line	Stata .do file(s) called by `subprocess` within python method	Notes
Figure 1	create_figure_1	analysis.py	create_fig_1_first_stage	1883	/figures/figure_1_first_stage/code/timeseries_plots_AA.do	
Figure 2	create_app_figure_adnl_ev	analysis.py	create_appendix_fig_adnl_event_studies	4731	None	Figure focus changed in revision, not reflected in method name Outputs individual panels, faceted in LaTeX
Figure 3	create_figure_3	analysis.py	create_fig_3_rx_qt_use	1955	/figures/figure_3_rx_qt_use_eventstudies/code/eventstudies.do	Outputs individual figures, faceted in LaTeX
Figure 4	create_figure_4	analysis.py	create_fig_4_in_asgn_plan	2048	/figures/figure_4_in_asgn_plan/code/timeseries_plots_AA.do	Revision changed to in assigned "model", not reflected in method name. See line 2207 for `outcome='in_asgn_model`
Figure 5	create_figure_5	analysis.py	create_fig_5_DD_timeseries	2091	/figures/figure_5_DD_timeseries/code/timeseries_plots_AA.do /figures/figure_5_DD_timeseries/code/timeseries_plots_DD.do	Outputs individual figures, faceted in LaTeX
Figure 6	create_figure_6	analysis.py	create_fig_6_util_mgt_denials	2134	/figures/figure_6_util_mgt_denials/code/timeseries_plots_AA.do	Outputs individual figures, faceted in LaTeX
Figure 7	create_figure_8	analysis.py	create_fig_8_within_class_substitutions	2281	/figures/figure_8_within_class_substitutions/code/dose_response.do	Figure ordering changed in revision, not reflected in method names
Figure 8	create_app_figure_10	analysis.py	create_appendix_fig_10_MMC_v_FFS_de	4053	/figures/app_figure_10_MMC_v_FFS_denials/code/dose_response.do	Figure ordering changed in revision, not reflected in method names

Learnings from the Data and Code Deposit (cont)

- Minimize the number of programming languages – ideally 1
 - Datatypes and import/output operations can be a real headache and cause really hard to debug problems.
 - Using Stata via subprocess has character length limits. `"/Anthony"` versus `"/Ben"`
- Do not hardcode paths – minimize this to a single spot and make everything else relative. Make sure paths are OS agnostic [e.g., `os.path.join('x', 'y', 'z')`]
- Be as specific and detailed as possible.
 - Assume the project will be handed off to someone else, or that you won't revisit the project for several months.

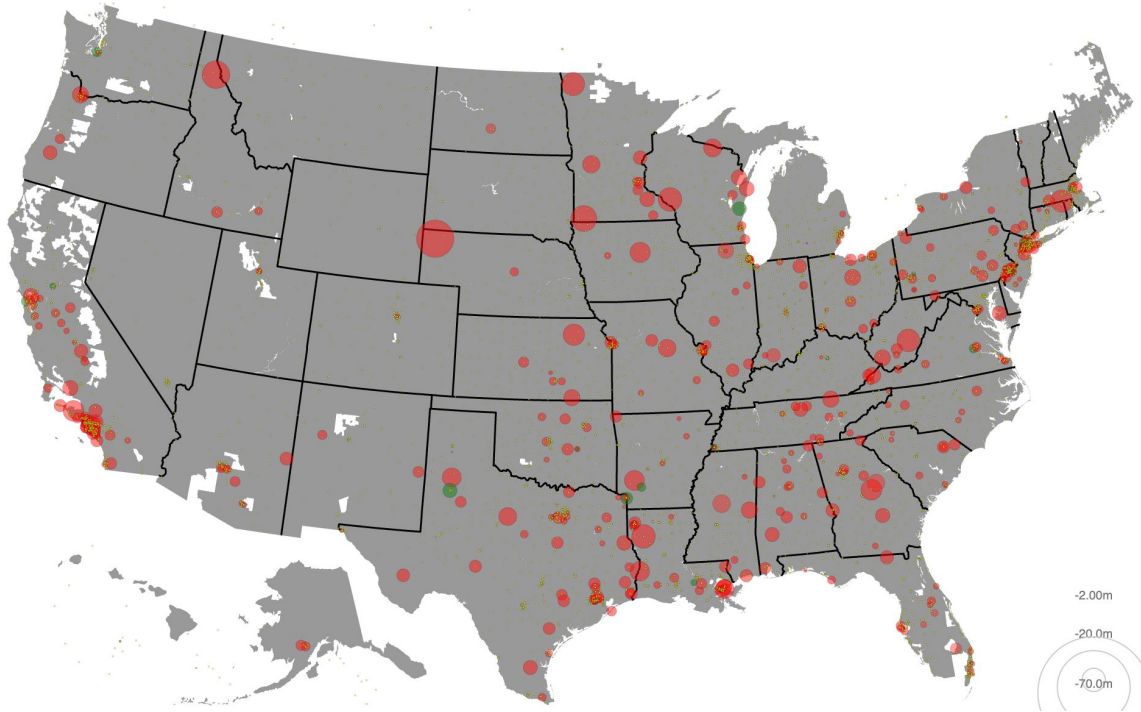
27



Example 2 (Maurice Dalton)

The study: Are Hospital Quality Indicators causal?

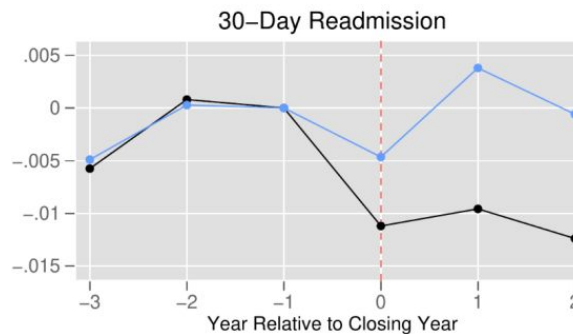
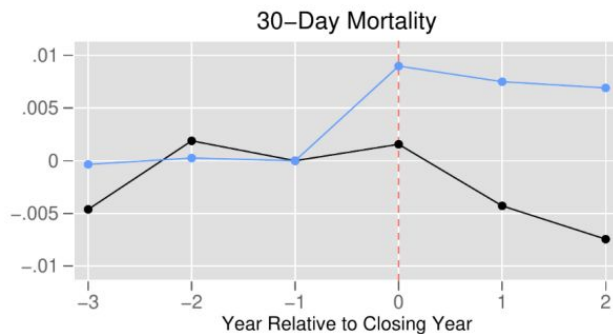
- Using over 20 years of Medicare Claims covering over 30 million patients



What did we find?

- Hospitals quality indicators overstate the causal impact of hospitals on outcomes
- Mortality and readmission rates by less than 10%
- Inpatient cost and length of stay by closer to 40%.
- Hospital closures reduce patient mortality by shifting patients to higher quality hospitals
- The effect varies widely depending on the relative quality of the closing hospital.
- See working paper [here](#)

30



Research takes a long time

- Any guesses when we started this project?
- Currently conditionally accepted, September 16th, with a deadline of submitting replication package by October 16th

Validating Hospital Quality Indicators and the Causal Effect of Hospital Closures

By AMITABH CHANDRA, MAURICE DALTON, AND DOUGLAS O. STAIGER*

We evaluate the validity of commonly used hospital quality indicators using hospital closures that reallocate large numbers of patients to hospitals of different quality. Using over 20 years of Medicare claims for over 30 million patients admitted with five common diagnoses, we find that hospital quality indicators overstate the causal impact of hospitals on mortality and readmission rates by less than 10% but overstate hospital impacts on inpatient cost and length of stay by closer to 40%. On average, hospital closures reduce patient mortality by shifting patients to higher quality hospitals, but the effect varies widely depending on the relative quality of the closing hospital. Simulations suggest that narrow networks limiting admissions to hospitals in the lowest quartile of mortality would reduce mortality by 1.4 percentage points among affected patients.

Some best practices that translate to high reproducibility

- Version control
- Build system
- Create packages of any methods for others to use
- Tables and figures should be generated by the code

Version control

- Pull requests allow you to pool changes into meaningful ideas then easily look back at the changes
- Requires access to git front end like github or gitlab

-
- * **August 2 2021 (compiled the 3rd)** see MR #4
 - bug: in the spillover regressions checks
 - spillover/reallocation check, add two simpler regressions which are easier to code up and explain
 - some measure of distance of CAH to other hospitals

33

```

444 445      * this measure zipm1_spill_wt`l' is then summed over all non closed hospital within a
      zipcode
445      -      gen hclose_flag=hclose_tdiyr==diag_yearpre
446      -      gen not_hclose_flag=hclose_tdiyr!=diag_yearpre
446      +      gen hclose_flag=hclose_tdiyr==baseyear
447      +      gen not_hclose_flag=hclose_tdiyr!=baseyear
447 448      bysort `benegeo' baseyear: egen zip_hclose_tm1 = max(hclose_flag)

```

Build system

- *Goal* is to replicate results with single command
- Simplest could be a single script
- While fully featured build systems often add complexity, e.g. explicitly needing to define outputs from each step, they manage which part of the pipeline needs to be rerun when changes are made
- Examples
 - Makefile/Snakemake good for simple projects but complexity increases quickly
 - Data Version Control (DVC) is a nice tool that versions your code and data and come with a build system. I have been implementing this more and more in my projects.

34

Example: Make file that built this presentation

```
M Makefile
1 # Makefile for building the ORR presentation
2 # Created: September 25, 2025
3
4 # Variables
5 PRESENTATION_FILE = orr-presentation/orr-presentation.qmd
6 OUTPUT_DIR = orr-presentation/docs
7
8 # Default target
9 all: build
10
11 # Build the presentation
12 build:
13     quarto render $(PRESENTATION_FILE)
14
15 # Clean the output directory
16 clean:
17     rm -rf $(OUTPUT_DIR)/*.html
18
19 # Rebuild everything from scratch
20 rebuild: clean build
21
22 # Preview the presentation
23 preview:
24     quarto preview $(PRESENTATION_FILE)
25
26 # Help command
27 help:
28     @echo "Available targets:"
29     @echo "  make       : Build the presentation"
30     @echo "  make build  : Same as above"
31     @echo "  make clean  : Remove generated HTML files"
32     @echo "  make rebuild : Clean and build again"
33     @echo "  make preview : Start a live preview server"
34
35 .PHONY: all build clean rebuild preview help
36
```

35

Managing tables and figures through the lifecycle

- Tables and figures should be generated as part of your build process
- We employed a large **org-mode** file which linked tables and figures to titles and explanations
- In that doc, we tagged table or figure with manuscript numbers when they moved to manuscript, .e.g. T1-Summary of closures
- Labeling scheme that avoids current table and figures numbers will save headaches as you rewrite the paper

36

Packages make things more reproducible

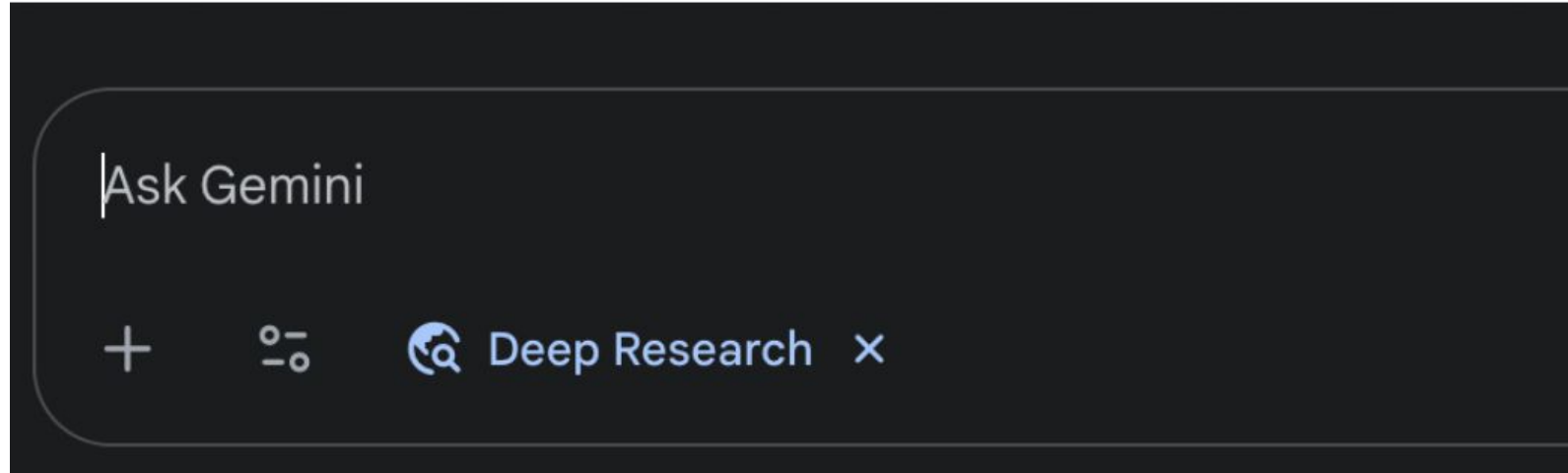
- We used an empirical Bayes user written program from a previous publication studying teacher quality as a starting point
- Some modifications were needed to make it specific to hospitals
- Some improvements were made to the codebase and documentation
- Currently available on [github](#) with a doc site found [here](#). Note we developed package after the fact so still in testing phase as part of replication package

37

ORR process

- AEA has a well documented website
- I decided to make the best of their documentation by employing AI

ORR and AI tools



39

Help me create a replication package as defined by <https://aeadataeditor.github.io/aea-de-guidance/>. Help create a document that an LLM can use to go through and create a replication package.

Gemini created a plan

The Document for the LLM Agent

Here is a step-by-step guide that you can use as a knowledge base for your AWS Bedrock agent. This document will guide the LLM in creating a compliant replication package.

Step 1: Create the README File

The README file is the most critical part of the replication package. It should contain the following sections:

1. Data Availability and Provenance

- * Data Availability Statement: Clearly state where the data can be accessed.
- * Data Provenance: For each dataset used, provide a detailed description of its source. If the data was collected by the author, explain the collection process.
- * License for Data: Specify the license under which the data is released (e.g., Creative Commons Attribution 4.0).

2. Dataset List

- * Provide a list of all datasets included in the package, with a brief description of each.

3. Computational Requirements

- * Software Requirements: List all software needed to reproduce the results, including version numbers (e.g., Stata 17, Python 3.9).
- * Memory, Runtime, and Storage: Estimate the computational resources required.

4. Description of Programs/Code

- * License for Code: Specify the license for the code (e.g., MIT License, BSD).
- * Code to Include: Detail what code is included and what it does. This should encompass everything from data cleaning and analysis to generating tables and figures.

5. Instructions to Replicators

- * Provide a clear, step-by-step guide on how to use the code to reproduce the results.

6. List of Tables and Programs

- * Create a mapping between the tables and figures in the paper and the code files that generate them.

Example 2 (Maurice Dalton)

Step 2: Organize the Data and Code

A well-structured replication package is easier to understand and use. Follow this directory structure:

/ (root)

README.txt (or README.pdf)

/data/

/raw/ (original, unprocessed data)

/analysis/ (data ready for analysis)

/code/ (all scripts and programs)

/results/ (output from the code, such as tables and figures)

Important Considerations:

No ZIP files within the package: The entire package will be zipped for submission, but there should be no zip files inside of it.

Confidential Data: If confidential data is used, create two separate packages: a public package with all non-confidential data and code, and a private package with the confidential data. The public README should explain how to access the confidential data.

Step 3: Choose a License

The AEA recommends a dual-license approach:

For Code: A permissive license like the MIT License or Modified BSD License.

For Data and other materials: The Creative Commons Attribution 4.0 International Public License.

A LICENSE.txt file should be included in the root directory of the package.

Step 4: Prepare the Metadata

When submitting the replication package to a repository (like openICPSR), you will need to provide the following metadata:

Title: "Data and Code for: [Title of article]"

Authors: The names of the individuals who prepared the data and code.

Summary: A brief description of the project and the data.

Yale *Data-Intensive Social Science Center*
Open & Reproducible Research

Cursor + plan + AWS Bedrock

- Use this document to create a replication package using Cursor hooked into AWS Bedrock for data privacy using an older version of claude.
- Pass the plan and code base and let it reorganize
- Ask it to generate to create synthetic data

Currently

- I have structured much of the project in a reproducible way
- The work left is to go through and make sure that we are including only necessary files
- Created a package for our empirical Bayes methods after the fact, but want to test that it replicates our results
- Make sure that tables and figures, especially during the review process are easy to follow

42

DISSC

Please give us your feedback!

DISSC Programming Evaluation
Form



Questions?

Thank You!

<https://dissc.yale.edu/>

Yale *Data-Intensive Social Science Center*

Open & Reproducible Research





Extra slides



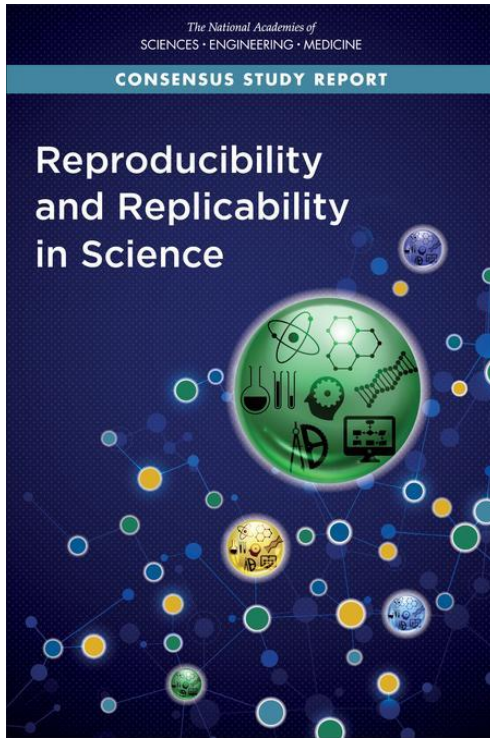
Data Availability Statements

Data availability statements provide information about where data may be found and under what conditions they may be accessed.

Data from the Socioeconomic High-resolution Rural Urban Geographic Dataset on India, Version 1.0 (Asher and Novosad, 2019) is used in this paper. The full dataset and documentation can be downloaded from <https://doi.org/10.7910/DVN/DPESAK>.

45

The following examples are broadly applicable: [Taylor & Francis](#) ; [SpringerNature](#) ; [PLOS](#) ; [Wiley](#). See also [Social Science Data Editors](#).



RECOMMENDATION 4-1:

To help ensure the reproducibility of computational results, **researchers should convey clear, specific, and complete information about any computational methods and data products that support their published results in order to enable other researchers to repeat the analysis**, unless such information is restricted by nonpublic data policies. That information should include the data, study methods, and computational environment...

46

Transparency and Openness Promotion Guidelines (TOP) for journals

Research Practices			
Practice	Level 1: Disclosed	Level 2: Shared and Cited	Level 3: Certified
Study Registration	Researchers stated whether or not a study was registered—and, if so, where and when it was registered.	Researchers registered the study and cited the registration.	A party independent from the researchers certified that the study was registered at an appropriate time and the registration was complete per best-practice for the study design.
Study Protocol	Researchers stated whether or not the study protocol is available—and, if so, where and when it was shared.	Researchers publicly shared the study protocol and cited the protocol.	A party independent from the researchers certified that the study protocol was shared at an appropriate time and the study protocol was complete per best-practice for the study design.
Analysis Plan	Researchers stated whether or not the analysis plan is available—and, if so, where and when it was shared.	Researchers publicly shared the analysis plan and cited the analysis plan.	A party independent from the researchers certified that the analysis plan was shared at an appropriate time and the analysis plan was complete per best-practice for the study design.
Reporting Transparency	Researchers stated whether or not they used a reporting guideline—and, if so, which guideline.	Researchers publicly shared the completed reporting guideline checklist and cited the reporting guideline.	A party independent from the researchers certified that the researchers adhered to the appropriate reporting guideline for the study design.
Materials Transparency	Researchers stated whether or not materials are available—and, if so, where.	Researchers cited materials deposited in a trusted repository by themselves or others.	A party independent from the researchers certified that materials were deposited and documented per best-practice for the type of materials.
Data Transparency	Researchers stated whether or not data are available—and, if so, where.	Researchers cited data deposited in a trusted repository by themselves or others.	A party independent from the researchers certified that data were deposited with metadata per best-practice for the type of data.
Analytic Code Transparency	Researchers stated whether or not analytic code is available—and, if so, where.	Researchers cited analytic code deposited in a trusted repository by themselves or others.	A party independent from the researchers certified that analytic code was deposited and documented per relevant best-practice.

Over 5,000 signatories

47

Yale Data-Intensive Social Science Center

Open & Reproducible Research

Transparency and Openness Promotion Guidelines (TOP) for journals

Verification Practices	
Practice	Definition
Results Transparency	A party independent from the researchers verified that results have not been reported selectively based on the nature of the findings. To verify, the independent party can check that the study registration, protocol, and analysis plan match the final report—and the final report acknowledges any deviations.
Computational Reproducibility	A party independent from the researchers verified that reported results reproduce using the same data and following the same computational procedures. To verify, the independent party can check that they obtain the same results using data and code deposited in a trusted repository.
Verification Studies	
Study Type	Definition
Replication	A study that aims to provide diagnostic evidence about claims from a prior study by repeating the original study procedures in a new sample.
Registered Report	A registered study in which a study protocol and analysis plan are peer reviewed, and the study is accepted in-principle by a publication outlet, before the research is undertaken.
Multiverse	A study in which a single research team examines the research question of interest across different, reasonable choices for processing and analyzing the same data.
Many Analysts	A study in which independent analysis teams conduct plausible alternative analyses of a research question on the same dataset.

Over 5,000 signatories

48

What do researchers think about these policies?

Survey of members of 4
Economics associations

Many comments mentioned that these [data and code sharing] policies **enhance the credibility of economic research** and the credibility of economics as a discipline overall. Some respondents also mentioned that the requirement to put together a replication package caused them to **catch inadvertent mistakes in their own work or caused them to adopt better ways of conducting their empirical research**.

49

Some respondents mentioned that these policies **won't prevent ill-intentioned researchers** from committing misconduct, though others pointed out that these policies make it harder to do so and make it easier to uncover. Some respondents also mentioned that these policies **don't ensure the code is correct** or corresponds to the methods described in the paper (data editors do not check code for correctness; they only check whether it reproduces the results in the paper).

However, these policies enable others to uncover such mistakes.