



Community Debate: Is AI helping or hurting our journey to reliable evidence?

Patrick Ryan, PhD

Vice President, Observational Health Data Analytics, Janssen
Research and Development

Assistant Professor, Adjunct, Department of Biomedical
Informatics, Columbia University Medical Center



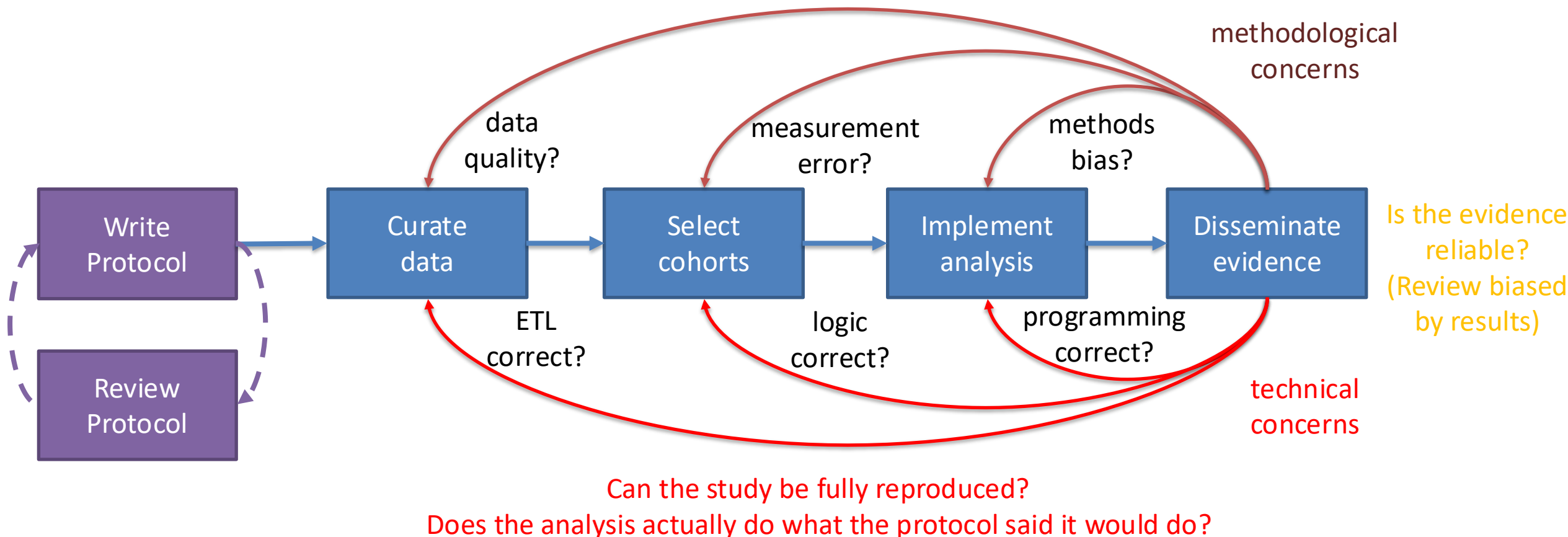
OHDSI's mission

To improve health by empowering a community to collaboratively generate the evidence that promotes better health decisions and better care



Current status quo in observational research makes it challenging to build trust in evidence

Does the study provide an unbiased effect estimate?
Are the findings generalizable to the population of interest?





An interactive journey from data to insights

COHORT_ID	PERSON_ID	AGE	RACE	PAIN_LEVEL	LDL_I	AE
0	115743	61	W	3.84	86	0
0	149648	85	O	5.69	123	0
0	134288	76	W	2.42	184	0
0	180972	44	W	7.24	107	0
0	174582	87	O	9.41	138	0
0	125791	45	O	7.97	91	1
0	193479	75	B	6.77	126	0
1	132963	61	W	0.35	176	0
0	104099	81	B	7.61	104	0
0	117940	83	W	5.99	108	0
0	181071	48	O	3.55	69	0



Join by Web PollEv.com/patrickryan800 Join by Text Send [patrickryan800](https://t.me/patrickryan800) to 22333

What do you think is the current impact of AI on OHDSI's journey to generating and disseminating reliable evidence?

- (A) Predominantly Harmful: AI primarily hinders the journey to reliable evidence (e.g., 0-20% Help / 80-100% Harm) 0%
- (B) More Harmful than Helpful: AI tends to create more obstacles than solutions for reliable evidence (e.g., 21-40% Help / 60-79% Harm) 0%
- (C) Equally Balanced: AI's helpful and harmful impacts on reliable evidence are roughly equal (e.g., 41-59% Help / 41-59% Harm). 0%
- (D) More Helpful than Harmful: AI generally facilitates the journey to reliable evidence more than it impedes it (e.g., 60-79% Help / 21-40% Harm). 0%
- (E) Predominantly Helpful: AI overwhelmingly advances the journey to reliable evidence (e.g., 80-100% Help / 0-20% Harm). 0%



Some observations about the role of AI

- If humans cannot consistently perform a task, it is not reasonable to expect AI will ***consistently*** perform the same task
- If humans cannot define a reproducible process for performing a task, it is not reasonable to expect AI will follow a ***reproducible*** process.
- If humans cannot verify on how well a task has been performed, it is not reasonable to expect AI to be ***verifiable***



An interactive journey from data to insights

COHORT_ID	PERSON_ID	AGE	RACE	PAIN_LEVEL	LDL_I	AE
0	115743	61	W	3.84	86	0
0	149648	85	O	5.69	123	0
0	134288	76	W	2.42	184	0
0	180972	44	W	7.24	107	0
0	174582	87	O	9.41	138	0
0	125791	45	O	7.97	91	1
0	193479	75	B	6.77	126	0
1	132963	61	W	0.35	176	0
0	104099	81	B	7.61	104	0
0	117940	83	W	5.99	108	0
0	181071	48	O	3.55	69	0

Data: COHORT_ID PERSON_ID AGE RACE PAIN_LEVEL LDL_I AE

Insight: OHDSI PEOPLE COCREATE REAL VALID LEARNING