

Development and Validation of the Artificial Intelligence Ethical Concerns Inventory - Healthcare (AIECIH)

John Robert Bautisa¹, Rhea Alex², Kenneth R. Fleischmann¹

¹School of Information, The University of Texas at Austin

²School of Human Ecology, College of Natural Sciences, The University of Texas at Austin

Introduction

- The WHO emphasizes the need to interrogate artificial intelligence (AI) ethical issues as a step towards the development of ethical AI-based health technologies.
- One of the challenges in the field of AI ethics is operationalizing AI ethical concerns.

Research Goal

To develop a validated survey tool (AIECIH) that can be used by researchers and policymakers to quantify ethical concerns in the use of AI in healthcare.

Methods

- AIECIH v1 is comprised of 28 items that reflect 10 AI-related ethical concerns.
- Draft items from Martinho, Kroesen, and Chorus (2021).
- 7 experts (4 female) in AI, ethics, and/or medicine were asked to validate AIECIH v1 (via Qualtrics).
- Item (I-CVI) and scale (S-CVI) level content validity index and modified kappa index (k) were used to evaluate AIECIH v1's content validity.

Results

Table 1. Item level results

Item	Agreement	I-CVI	k	Evaluation
Accountability 1	6	0.86	0.86	Excellent
Accountability 2	7	1.00	1.00	Excellent
Accountability 3	5	0.71	0.71	Good
Certification of AI Products 1	6	0.86	0.86	Excellent
Certification of AI Products 2	7	1.00	1.00	Excellent
AI Education 1	6	0.86	0.86	Excellent
AI Education 2	6	0.86	0.86	Excellent
Ethical Design 1	7	1.00	1.00	Excellent
Ethical Design 2	5	0.71	0.71	Good
Explainability 1	6	0.86	0.86	Excellent
Explainability 2	6	0.86	0.86	Excellent
Explainability 3	4	0.57	0.55	Fair
Explainability 4	4	0.57	0.55	Fair
Fairness 1	5	0.71	0.71	Good
Fairness 2	6	0.86	0.86	Excellent
Fairness 3	6	0.86	0.86	Excellent
Fairness 4	6	0.86	0.86	Excellent
Future of Employment 1	7	1.00	1.00	Excellent
Future of Employment 2	4	0.57	0.55	Fair
Future of Employment 3	2	0.29	0.17	Poor
Privacy 1	7	1.00	1.00	Excellent
Privacy 2	4	0.57	0.55	Fair
Privacy 3	6	0.86	0.86	Excellent
Professional autonomy 1	5	0.71	0.71	Good
Professional autonomy 2	5	0.71	0.71	Good
Safety 1	4	0.57	0.55	Fair
Safety 2	3	0.43	0.38	Poor
Safety 3	5	0.71	0.71	Good

Table 2. Summary of evaluation

Evaluation	k range	# of items	%
Excellent	> .78	15	54%
Good	.60 - .77	6	21%
Fair	.40 - .59	5	18%
Poor	< .40	2	7%

S-CVI/Ave (all items) = 0.77

S-CVI/Ave (removed 2 poor items) = 0.80

Table 3. Summary of revisions

Concerns	AIECIH v1: 28 items	AIECIH v2: 27 items	Revisions
1. Privacy	3	3	Edited for clarity
2. Fairness	4	4	Edited for clarity
3. Accountability	3	3	Edited for clarity
4. Safety	3	2	Edited for clarity 1 item removed
5. Professional autonomy	2	2	Edited for clarity
6. Explainability	4	4	Edited for clarity
7. Future of employment	3	3	Edited for clarity 1 item removed 1 item split into two
8. AI Education	2	2	Edited for clarity
9. Certification of AI products	2	2	Edited for clarity
10. Ethical Design	2	2	Edited for clarity

Scan QR code for
AIECIH v2 items



Conclusion

- AIECIH v1 achieved acceptable content validity; results were used to create v2.
- AIECIH v2 is composed of 27 items reflecting 10 AI-related ethical concerns.
- On going data collection to check AIECIH v2's construct validity.

Acknowledgments

- Good Systems Faculty Fellowship
- UT iSchool Boyvey Fellowship