



University of
Pittsburgh

School of Computing
and Information

Proposal Defense

Doctor of Philosophy in Information Science

"Rapport Before Correction: Peripheral Cue Gatekeeping in Conversational AI Health Misinformation Debunking"

by Rr. Nefriana

Date: July 21, 2026

Time: 10:00 a.m. – Noon

Place: Room 502, Information Sciences Building, 135 N
Bellefield Ave, Pittsburgh PA 15260

Committee:

- Yu-Ru Lin, Professor, Department of Informatics and Networked Systems, School of Computing and Information, University of Pittsburgh
- Daqing He, Professor, Department of Informatics and Networked Systems, School of Computing and Information, University of Pittsburgh
- Angela E.B. Stewart, Assistant Professor, Department of Informatics and Networked Systems, School of Computing and Information, University of Pittsburgh
- Jamie Zelazny, Assistant Professor, Department of Health & Community Systems, School of Nursing, University of Pittsburgh

Abstract:

Health misinformation does not just mislead—it delays treatment, discourages care, and erodes trust in health systems. Existing debunking methods have struggled to change deeply held beliefs, and what prevents people from accepting corrective information—and how this varies across populations—remains poorly understood, making it difficult to design interventions that effectively address those barriers.

The Elaboration Likelihood Model (ELM) offers a framework for understanding these barriers: people process information either through careful argument evaluation or through reliance on surface-level cues such as warmth or empathy, depending on their motivation and ability to engage. I hypothesize that deliberately structuring peripheral cues before correction—a mechanism I call peripheral cue gatekeeping—is a plausible extension of ELM to conversational AI debunking contexts. Specifically, peripheral cue gatekeeping may address processing barriers by increasing motivation and ability to engage in careful argument processing, and by reducing perceived threat and shifting motivation toward fair argument evaluation for those processing defensively.

To test these ideas empirically, this proposed study includes three research components. RC1 examines what barriers prevent people from accepting corrective health information and how processing unfolds sequentially. Demographic comparisons are additionally conducted to explore, in a preliminary and descriptive manner, whether certain groups may face distinct barriers. Together, RC1 findings provide interpretive context for RC2 and RC3 and inform the refinement of RC2's experimental design.



University of
Pittsburgh

School of Computing
and Information

RC2 tests whether rapport building before correction — with peripheral cue gatekeeping as one plausible mechanism — produces greater belief change than correction without rapport building. The centerpiece is a randomized controlled trial using a large language model (LLM)-based chatbot, as a potential tool for health misinformation debunking. RC1 findings will inform whether an additional baseline condition is needed, depending on whether peripheral cue presence or sequencing emerges as the more critical mechanism. RC2 further examines for whom this works best and through what cognitive and, secondarily, relational mechanisms. Finally, RC3 analyzes RC2 conversation transcripts to characterize what peripheral cues and argument content LLMs deploy, and tests whether their presence and placement predict belief change.

This study extends ELM research to conversational AI, hypothesizing how peripheral cue gatekeeping may activate receptiveness to corrective arguments. Empirically, it provides an account of processing barriers to health misinformation correction, an empirical test of rapport-first debunking, and an account of the conversational strategies that predict belief change in LLM-based debunking.