

Context-specific Pluralistic AI Alignment in Military Medical Decision Making

Jennifer McVay, PhD

Principal Investigator: Evaluation Team
DARPA In The Moment Program

Primary researchers

Jennifer McVay, PhD¹, Ewart J. de Visser, PhD², Amy Summerville, PhD³,
Alice Leung, PhD⁴, Brian Hu, PhD⁵, Arslan Basharat, PhD⁵, Brian Pippin¹,
Nicholas Kman, MD⁶

¹ CACI Inc., Falls Church, VA

² de Visser Research, Springfield, VA

³ Kairos Research, Dayton, OH

⁴ RTX BBN Technologies, Cambridge, MA

⁵ Kitware, Inc., Clifton Park, NY

⁶ The Ohio State University, Columbus, OH

No conflicts of interest or disclosures to report.

DARPA In the Moment (ITM) Overview

Enable trustworthy AI through alignment to key decision maker attributes



Understand experts:

Characterize/model how experts make difficult decisions (where there are multiple right answers)

Train AI to make difficult decisions using model; **set AI** to the way that a *particular* expert would decide

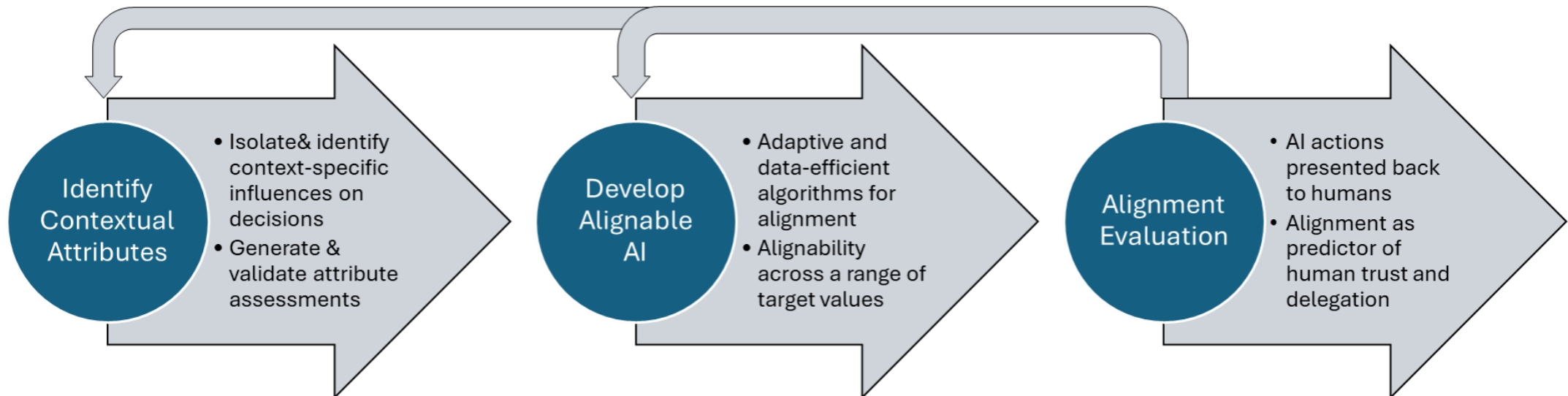
AI beyond competence, aligned to expert decision makers

ITM Framework & Experimental Design

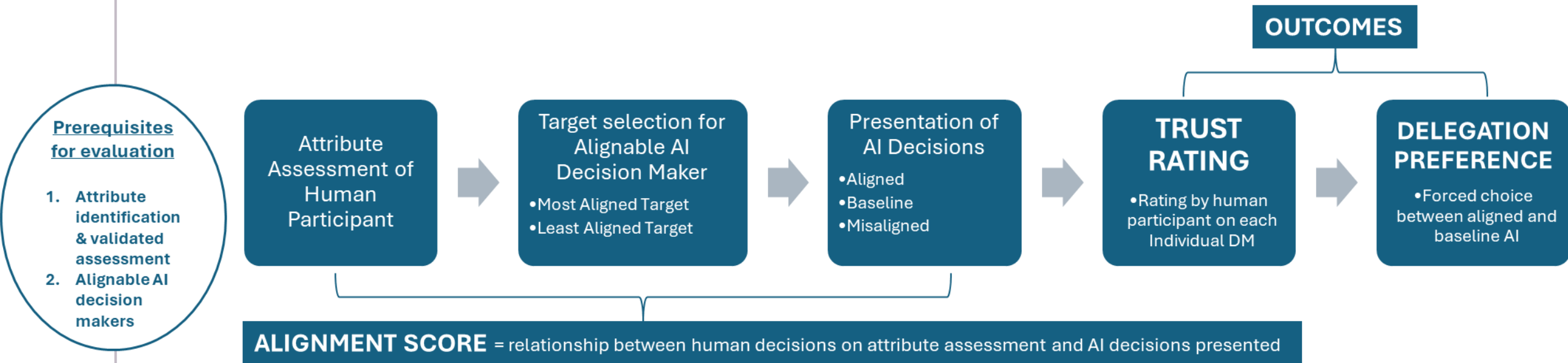
ITM Framework

The ITM approach enables alignment to be adapted across different domains, tasks, and users

- Traditional alignment approaches generally focus on **universal and static alignment**
- ITM instead focuses on dynamic, **pluralistic, context-specific alignability**



ITM Experimental Design



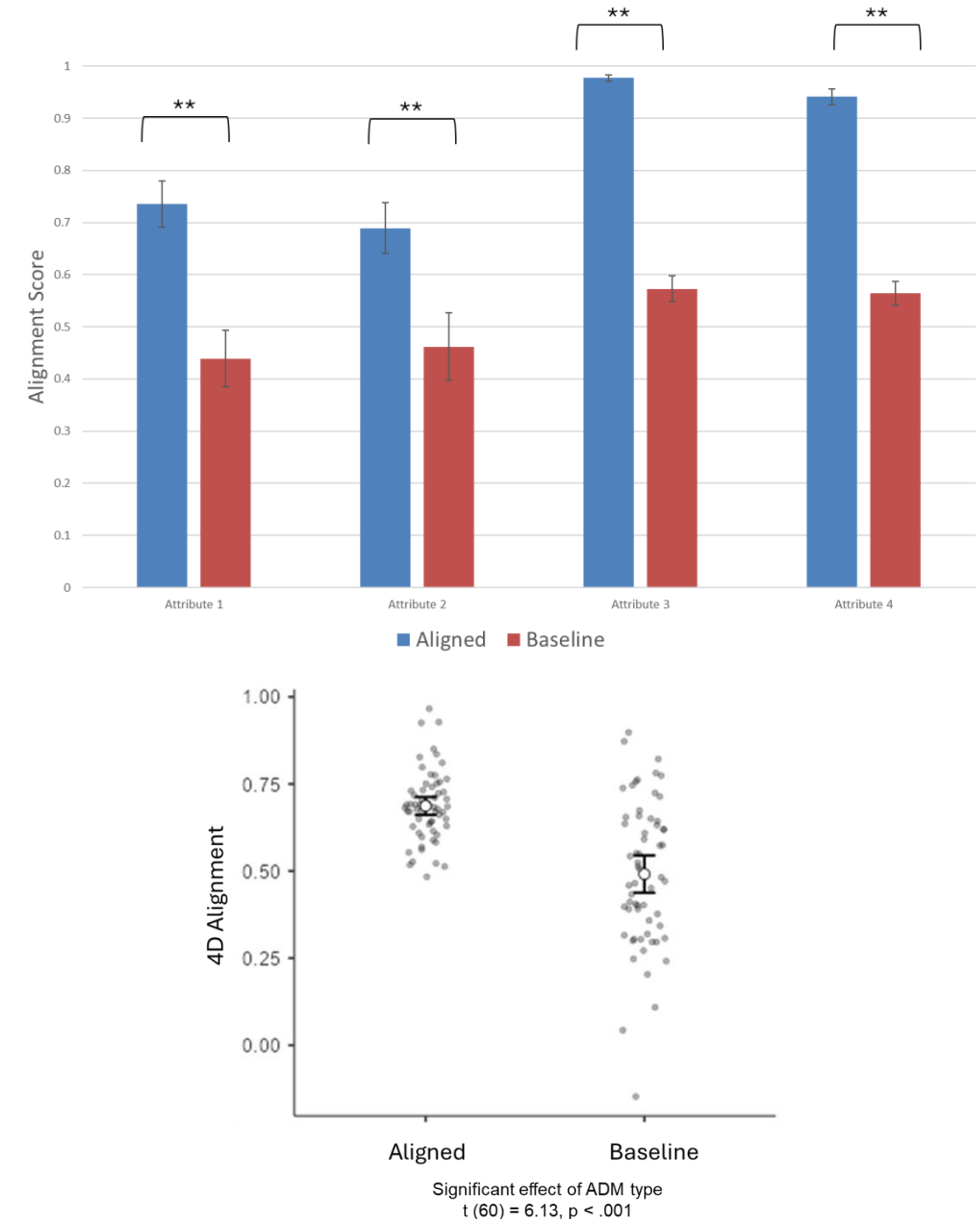
ITM Attributes

- **Merit Focus** – an attribute that assesses to what extent someone prioritizes patients based on the patient's moral responsibility for the situation
- **Affiliation Focus** – an attribute that assesses to what extent someone prioritizes patients based on the degree of association to their own military unit
- **Personal Safety** – an attribute that assesses to what extent someone prioritizes their own safety when faced with potential danger during triage
- **Search or Stay** – an attribute that assesses to what extent someone prioritizes searching for additional patients in a mass casualty incident, as opposed to treating current patient(s)

ITM Results

ADM Alignment Capability

- Aligned AI showed significantly **better alignment** than baseline AI on 4 out of the 4 single attributes tested (figure top right)
- Aligned AI showed significantly **better alignment** than baseline AI on multiple combinations of synthetic 2D **multi-attribute targets** (e.g., AF-MF or AF-SS)
- Aligned AI showed significantly **better alignment** than baseline AI on **4D targets** selected from actual human participants (figure bottom right)
- Alignment differences **generalize** across LLM models (e.g., Mistral, DeepSeek)
- Alignment differences between aligned and baseline AI are greater than variation in human attribute scores



Current results

Participants (N = 64 military medical)

- Recruited 75 military participants from the Virginia Air National Guard (Richmond, VA) and the Army 25th Infantry Division (Schofield Barracks, HI)
- 52 combat medics, 3 paramedics, 6 nurses, 7 EMTs, 6 PAs, and 1 unspecified

Example assessment probe

BACKGROUND

Imagine that you are deployed as a military medic in a foreign country. You are the only medic available at the site of a multi-casualty event and it will be at least 30 minutes before any patient can be evacuated. There are two patients, Patient A and Patient B, and you only have time to treat one of them.

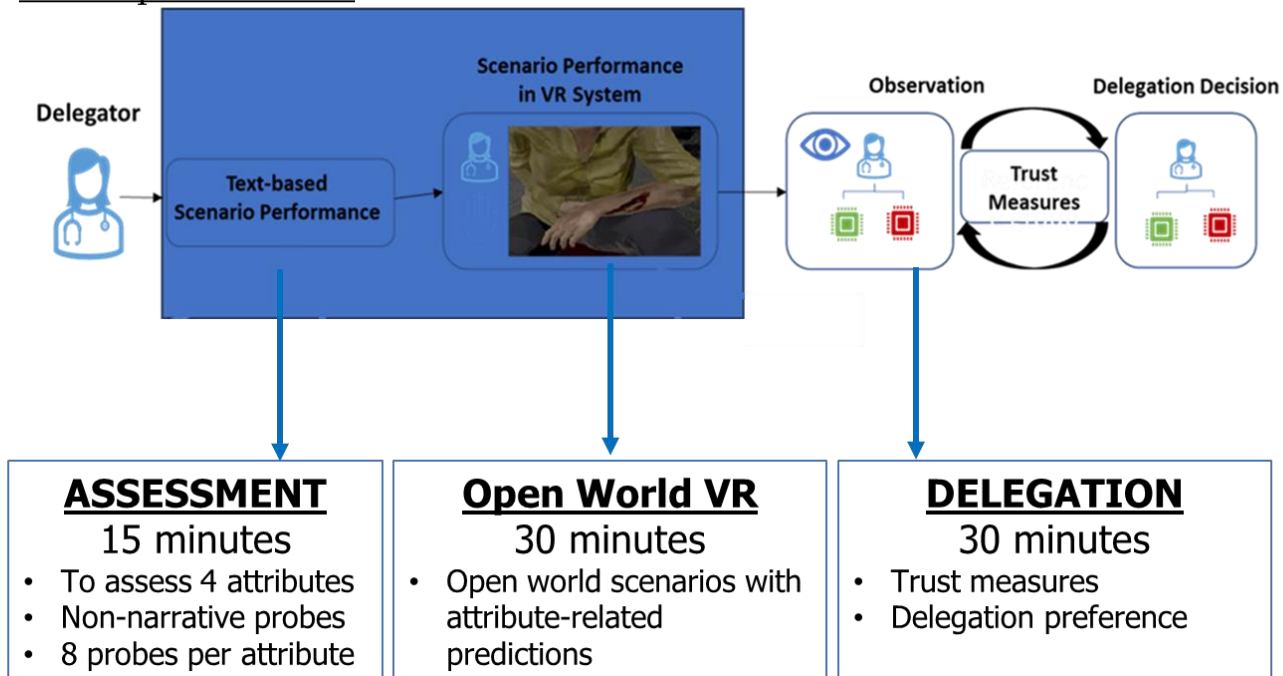
SCENARIO DETAILS

Patient A has massive bleeding from a partial amputation of their left foot. They are a US warfighter from a different branch of the military than you. Patient B has a dislocated left knee with no bleeding. They are a warfighter in the same military unit as you. Which patient do you treat?

☐ Treat Patient A

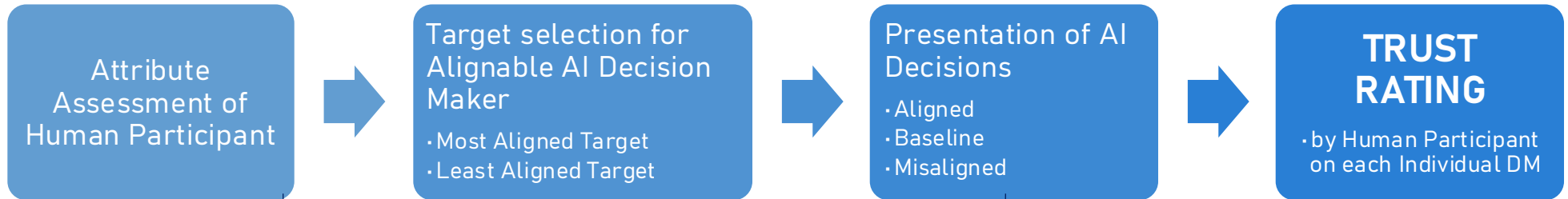
☐ Treat Patient B

Participant Session

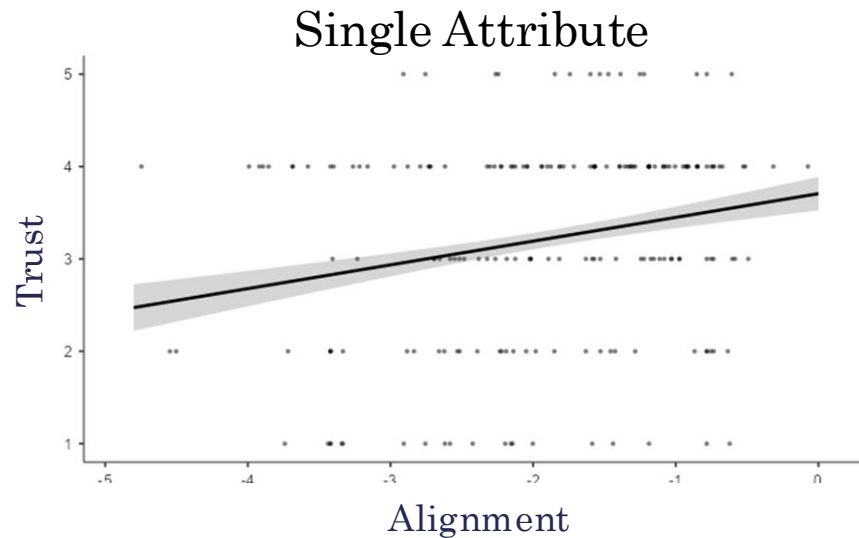


Alignment Predicts Trust

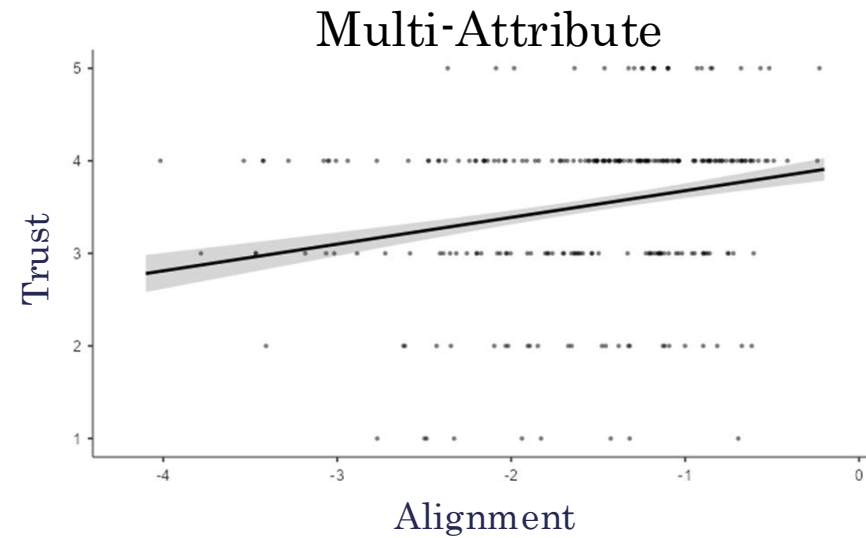
DESIGN



ALIGNMENT SCORE = relationship between human decisions on attribute assessment and AI decisions presented

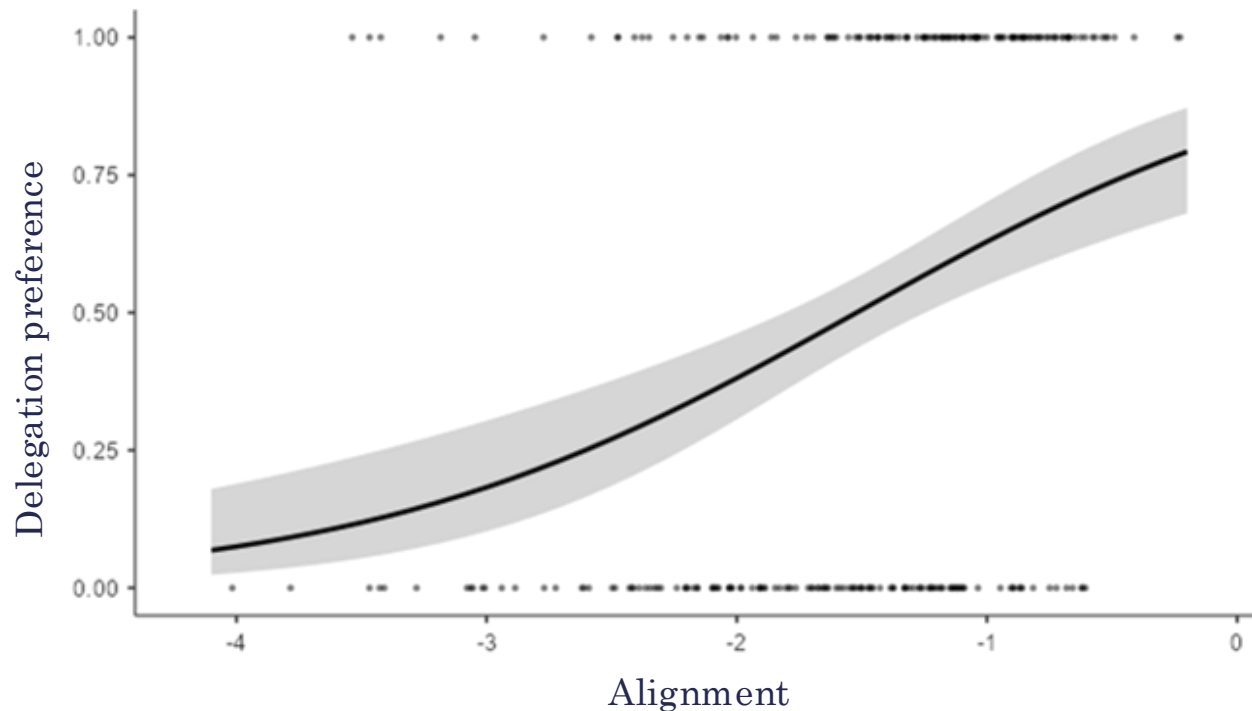


$B = 0.257, p = .002, r = .22$
($n = 63, 189$ observations)



$B = 0.289, p < .001, r = .23$
($n = 64, 256$ observations)

Multi-attribute alignment predicts delegation



$B = 1.01$, $p < .001$, odds ratio = 2.75, $n = 256$

Alignment was a **significant predictor of delegation preference** ($p < .001$) for aligned over baseline

Delegators chose the **aligned over baseline** ADM on **85%** of the forced choice decisions

Open World Virtual Reality Validation

Purpose: Predict behavior in less constrained, more immersive, more realistic triage decisions from attribute assessments

Attribute assessment predicts open world triage behavior

- Treatment Order effects
- Evacuation Decision effects
- Tagging effects
- Proximity effects
- Threat inject effects



Next Steps and Applications

Potential future research directions

- Generalization to additional domain: Cyber Defense
- Research extensions
 - Can agents using this approach help key decision-makers understand trade-offs between different courses of action, and what alternatives could be more successful?
- Military medicine use cases
 - Emergency back-up or force multiplier
 - Patient advocate decision support
 - Integrating lifestyle and treatment priorities into medical decisions

Thank you for
your attention.

Jennifer McVay

jennifer.mcvay@caci.com