

# HARMS Risk Worksheet

Score the AI use your table just performed in Lab 3.

## H Humans affected

Who touches or is touched by this AI use? (users, bystanders, the people in the data)

---



---

## A Adverse outcome

What is the worst plausible outcome — not the worst imaginable?

---



---

## R Risk defenses in place

What stands between today and that outcome? What did we lower without raising something else?

---



---

## M Mitigation & owner

What one defense do we add, and whose name is on it?

---



---

## S Signal & receipt

What evidence would prove the defense worked — and where do we keep it?

---



---

**Highest risk** (one sentence)

---

**Mitigation** (one sentence)

---

Note for facilitators: field labels align to Ken's HARMS framework; swap in canonical wording from the foundation site if it differs.