

Goal Recognition, Norms, and Principled Disobedience in BDI Agents

Prof. Felipe Meneguzzi

University of Aberdeen, Scotland, UK
`felipe.meneguzzi@abdn.ac.uk`

RAD-AI @ AAMAS, Cyprus, 2026

1495



UNIVERSITY OF
ABERDEEN

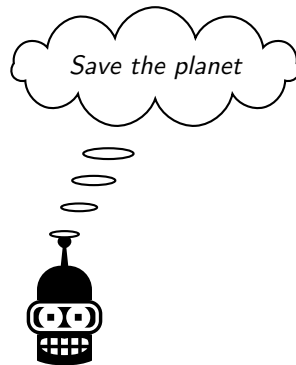
Table of Contents

- 1 Motivation
- 2 BDI Architecture and Reasoning
- 3 Normative BDI
- 4 Recognising Intent and Norms
 - Goal Recognition
 - Norm Identification
- 5 The RAD-AI Synthesis
- 6 Conclusions

Motivation

The Compliance Question

- Autonomous agents operate in environments regulated by norms
 - Laws, contracts, social conventions
 - Obligations, permissions, prohibitions
- Classic assumption: agents *should* comply
 - But compliance may be impossible
 - Or may conflict with an agent's goals
 - Or the norm itself may be unjust



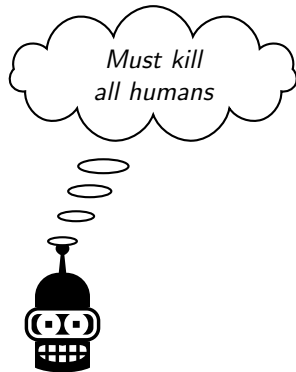
^aThomas Arnold, Gordon Briggs, and Matthias Scheutz. "Only Those Who Can Obey Can Disobey: the Intentional Implications of Artificial Agent Disobedience". In: *RAD-AI @ AAMAS*. 2022.

Motivation

The Compliance Question

- Autonomous agents operate in environments regulated by norms
 - Laws, contracts, social conventions
 - Obligations, permissions, prohibitions
- Classic assumption: agents *should* comply
 - But compliance may be impossible
 - Or may conflict with an agent's goals
 - Or the norm itself may be unjust
- **When should an agent disobey?^a**

^aArnold, Briggs, and Scheutz, "Only Those Who Can Obey Can Disobey: the Intentional Implications of Artificial Agent Disobedience".



- **Deciding your own behaviour**

- How does a BDI agent represent and reason about norms?
- When does goal achievement *require* norm violation?
- How do we learn what norms apply from observing others?

- **Recognising others' intent**

- What is another agent trying to achieve?
- Is their behaviour compliant with norms?
- How can we detect norm violations from observation?

Table of Contents

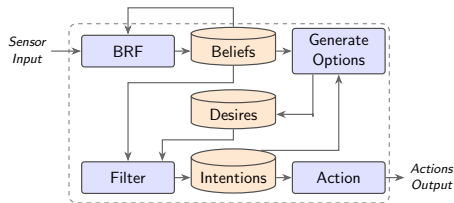
- 1 Motivation
- 2 BDI Architecture and Reasoning
- 3 Normative BDI
- 4 Recognising Intent and Norms
 - Goal Recognition
 - Norm Identification
- 5 The RAD-AI Synthesis
- 6 Conclusions

BDI Agents

A model of *practical* reasoning

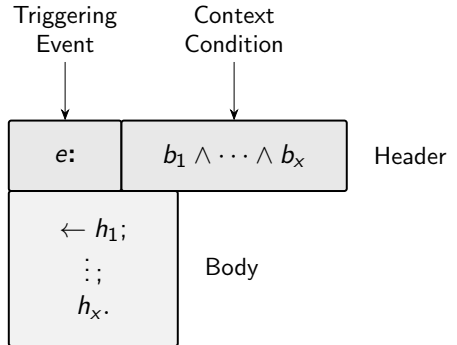
Agent architecture based on three “mental” structures.

- **Beliefs** — agent’s knowledge about the world
- **Desires** — goals the agent wants to achieve
- **Intentions** — committed courses of action
- Based on a philosophical model of *practical reasoning*
- Key process: *means-ends reasoning*
 - Typically via a *plan library*
 - Recently: *generalised planning, natural language*



BDI Reasoning Cycle

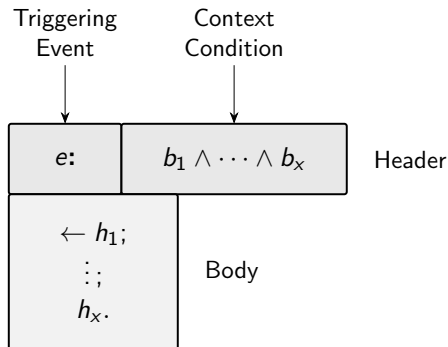
- ① Perceive events from the environment
- ② Update belief base
- ③ Select *relevant* plans (triggering event matches)
- ④ Select *applicable* plans (context condition holds)
- ⑤ Add chosen plan to intentions
- ⑥ Execute one step of current intentions



Plans in BDI

Plan Library

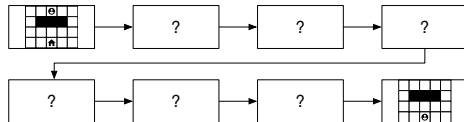
- Plans are *reactive procedures*:
 - Triggering event e
 - Context condition $b_1 \wedge \dots \wedge b_x$
 - Body: ordered list of actions/sub-goals
- Filters for relevant/applicable plans
- Plan library encodes *designer knowledge*
 - Cannot handle unanticipated situations



Plans in BDI

Plan Library

- Plans are *reactive procedures*:
 - Triggering event e
 - Context condition $b_1 \wedge \dots \wedge b_x$
 - Body: ordered list of actions/sub-goals
- Filters for relevant/applicable plans
- Plan library encodes *designer knowledge*
 - Cannot handle unanticipated situations



Plans in BDI

Plan Library

- Plans *focus reasoning*
- Filters for relevant/applicable plans?
- Plan library encodes *designer knowledge*
 - Still limited to what the designer anticipated

BDI Today: Agentic LLMs

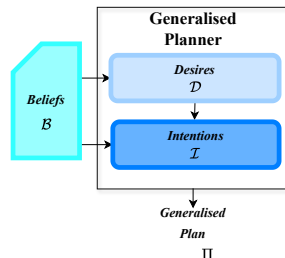
BDI	Agentic LLM
Belief base	Context
Plan library	PLAN.md
Means-ends reasoning	The LLM itself
Reasoning cycle	Tool-call loop
Triggering event	User prompt

“BDI didn’t go away — it just stopped being cited.”

Automated Planning in BDI

Beyond Plan Libraries¹

- Plan libraries are *finite* — they cannot anticipate all situations
- Automated planners generate plans *at runtime*
 - Handle novel situations not foreseen by designers
 - But: planning is computationally expensive
- Recent direction: *generalised planning* in BDI
 - Plans that work across *classes* of problems
 - Bridges plan library robustness with planning flexibility

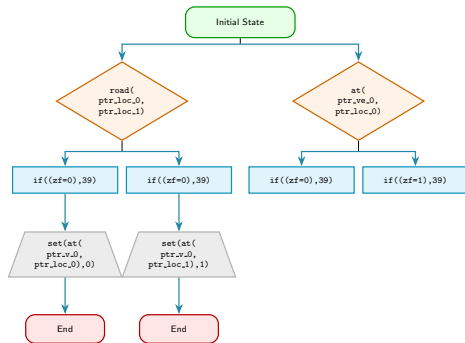


¹Ramon Fraga Pereira, Nir Oren, and Felipe Meneguzzi. "Generalised BDI Planning". In: *AAMAS*, 2025.

Generalised Planning in BDI

Overview

- Generalised plans: programs with loops and conditionals
- Solve not one problem, but a *class* of related problems
- Key insight for BDI:
 - A BDI agent *is* a generalised planner
 - Reasoning cycle implements a generalised executor
 - Plan library entries are conditional branches
- Formal correspondence enables verification



Generalised Planning in BDI

Overview

- Generalised plans: programs with loops and conditionals
- Solve not one problem, but a *class* of related problems
- Key insight for BDI:
 - A BDI agent *is* a generalised planner
 - Reasoning cycle implements a generalised executor
 - Plan library entries are conditional branches
- Formal correspondence enables verification

```
// Direct move: road exists
+!move(V, L0, L1)
  : road(L0, L1) & at(V, L0)
  ← -at(V, L0);
    +at(V, L1).

// Indirect move: no direct road
+!move(V, L0, L1)
  : not road(L0, L1)
  ← !findRoute(V, L0, L1).
```

Planning for the Environment and Planning for Society

- Plan Synthesis / BDI Programmes (usually):
 - Designed to achieve goals;
 - Cope with environmental restrictions;
 - Explicitly coordinate with other agents.
- The usual assumption is of prior knowledge:
 - Of agent goals;
 - Environmental dynamics;
 - Interaction protocols.
- But how do we deal with dynamic societies and implicit behaviour?
 - Agents with unknown or implicit goals?
 - Dynamic environments?
 - Emerging societal dynamics?

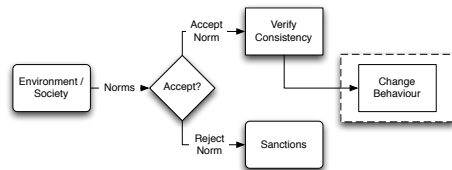
Table of Contents

- 1 Motivation
- 2 BDI Architecture and Reasoning
- 3 Normative BDI**
- 4 Recognising Intent and Norms
 - Goal Recognition
 - Norm Identification
- 5 The RAD-AI Synthesis
- 6 Conclusions

Norms and Agents

Why we need norms²

- Autonomous agents operate in *open* environments;
- Norms regulate these environments;
- They use deontic concepts:
 - **Obligations** — what agents *must* do
 - **Permissions** — what agents *may* do
 - **Prohibitions** — what agents *must not* do
- Challenge: traditional BDI provides *no mechanism* to adapt at runtime to newly accepted norms



²Felipe Meneguzzi and Michael Luck. "Norm-based Behaviour Modification in BDI Agents". In: *AAMAS*. 2009.

Norm Representation

- Norms have an operational structure:

$$\langle \nu, Act, Exp, id \rangle$$

- ν — normative formula
- Act — activation condition
- Exp — expiration condition
- id — norm identifier
- Normative formula: $\mathbf{X}_{\alpha:\rho} \varphi \circ \Gamma$
 - $\mathbf{X} \in \{\mathbf{O}, \mathbf{F}\}$ — obligation/prohibition
 - φ — targeted action
 - Γ — constraints on action parameters

LLM guardrail as a norm

“Do not delete files during a read-only task”

- $\nu = \mathbf{F} \text{ delete}(F) \circ \text{critical}(F)$
- $Act = \text{taskGoal}(\text{readOnly})$
- $Exp = \text{taskComplete}$

Yet: a tool-call chain

read $\rightarrow$ summarise $\rightarrow$ cleanup $\rightarrow$ delete

...still wipes critical files: **a norm violation**

Norm Representation

- Norms have an operational structure:

$$\langle \nu, Act, Exp, id \rangle$$

- ν — normative formula
- Act — activation condition
- Exp — expiration condition
- id — norm identifier
- Normative formula: $\mathbf{X}_{\alpha:\rho} \varphi \circ \Gamma$
 - $\mathbf{X} \in \{\mathbf{O}, \mathbf{F}\}$ — obligation/prohibition
 - φ — targeted action
 - Γ — constraints on action parameters

What counts as “critical” is determined by the user’s implicit intent.

Only an agent that recognises that intent can disobey the norm **selectively**.

LLM guardrail as a norm

“Do not delete files during a read-only task”

- $\nu = \mathbf{F} \text{ delete}(F) \circ \text{critical}(F)$
- $Act = \text{taskGoal}(\text{readOnly})$
- $Exp = \text{taskComplete}$

Yet: a tool-call chain

read $\rightarrow$ summarise $\rightarrow$ cleanup $\rightarrow$ delete

...still wipes critical files: **a norm violation**

Previous Normative Agent Systems

Two Extremes

- **Blanket plan retractions** (Normative AgentSpeak)
 - Entire plans removed when a norm activates
 - Loses autonomy entirely
- **Every norm checked at every step** (BOLD)
 - Computationally wasteful
 - Compliance decision made too early
- Both extremes: *non-compliance is not an option*
- cf. R-HTN³: adaptive HTN agents that *modify* plans on violation — more goals reached, fewer violations
- cf. Ethical Rebellion System⁴: dual-process architecture — fast intuitive + deliberative ethical reasoning

³Héctor Muñoz-Avila, David W. Aha, and Paola Rizzo. “R-HTN: Rebellious Online HTN Planning”. In: *RAD-AI @ AAMAS*. 2025.

⁴Ursula Addison, Othalia Larue, and Matthew Molineaux. “Towards An Ethical Rebellion System”. In: *RAD-AI @ AAMAS*. 2023.

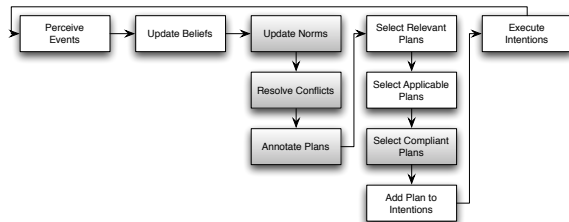
Norm Modification in BDI

Runtime Adaptation⁵

- Open environments deliver norms to agents *at runtime*
- **Obligations:** generate *new plans* to fulfil them
 - Plan creation triggered by norm activation
- **Prohibitions:** *suppress* execution of violating plans
 - Plans annotated and filtered at selection time
- This is a blunt instrument: the agent cannot disobey even when it should.

⁵Meneguzzi and Luck, “Norm-based Behaviour Modification in BDI Agents”. 

- Three new processes added to BDI cycle:
 - ① **Update norms** (resolve conflicts)
 - ② **Annotate plan library** with norm constraints
 - ③ **Select compliant plans** at plan selection time
- Effect of norms computed at *norm receipt*
- Compliance decision *delayed* as long as possible

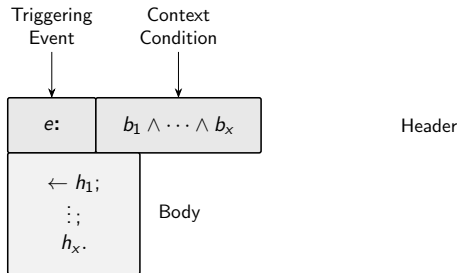


⁶Felipe Meneguzzi et al. "BDI Reasoning with Normative Considerations". In: *Engineering Applications of Artificial Intelligence* 43 (2015).

- When a specific norm activates:
 - Normative formula compared to each plan step
 - Each step annotated with appropriate constraints
- Plan header extended with a *normative header*:

$$\langle e : b_1 \wedge \dots \wedge b_x \mid \gamma_1 \wedge \dots \wedge \gamma_i \rangle$$

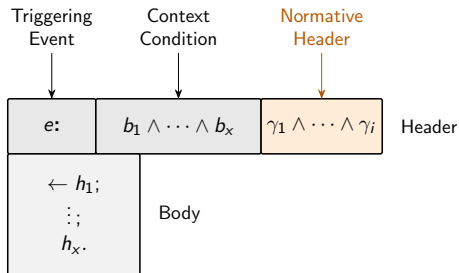
- At plan selection: check normative header
 - Compliant plan instantiation chosen where possible



- When a specific norm activates:
 - Normative formula compared to each plan step
 - Each step annotated with appropriate constraints
- Plan header extended with a *normative header*:

$$\langle e : b_1 \wedge \dots \wedge b_x \mid \gamma_1 \wedge \dots \wedge \gamma_i \rangle$$

- At plan selection: check normative header
 - Compliant plan instantiation chosen where possible
 - Otherwise, selective disobedience



Principled Norm Violation

When Compliance is Impossible

- What if *all* applicable plans are normatively unsatisfiable?

$$(C = \text{London} \wedge C \neq \text{London}) \rightarrow \perp$$

- Naïve response: agent fails to achieve goal
- ν -BDI response: **selective and incremental norm violation**
 - Agent identifies *which* norms must be violated
 - Chooses the plan that minimises violation
 - Norm violation is *explicit* and *reasoned*, not accidental
- This is the formal basis for *principled disobedience*⁷

Key Result

An agent can comply with norms in all cases *where possible*, and violate them only when goal achievement would otherwise be impossible.

⁷Divya Sundaresan, Gordon Briggs, and Munindar P. Singh. “Obey, Refrain, or Contravene?: Disobedience as a Basis for Responsible Autonomy”. In: *AAMAS*. 2025.

Adaptive normative reasoning

Selective disobedience when planning is cheap

- **2010 assumption:** plan libraries are hand-crafted and expensive
 - ν -BDI was conservative: annotate existing plans, defer decisions
- **2025 reality:** generalised planning⁸ makes new plan synthesis cheap
 - Agents can reason *from scratch* about normative constraints
- New spectrum of normative behaviour:
 - ① **Comply** — synthesise a plan satisfying all norms
 - ② **Minimally violate** — synthesise the plan with the fewest violations
 - ③ **Rebel** — explicitly ignore specific norms, with justification

Compliance becomes a *cost–benefit calculation*, not a binary constraint

⁸Pereira, Oren, and Meneguzzi, “Generalised BDI Planning”.

Adaptive normative reasoning

Selective disobedience when planning is cheap

- **2010 assumption:** plan libraries are hand-crafted and expensive
 - ν -BDI was conservative: annotate existing plans, defer decisions
- **2025 reality:** generalised planning⁸ makes new plan synthesis cheap(er)
 - Agents can reason *from scratch* about normative constraints
- New spectrum of normative behaviour:
 - ① **Comply** — synthesise a plan satisfying all norms
 - ② **Minimally violate** — synthesise the plan with the fewest violations
 - ③ **Rebel** — explicitly ignore specific norms, with justification

Compliance becomes a *cost–benefit calculation*, not a binary constraint
But is this *utilitarian/consequentialist* view enough?

⁸Pereira, Oren, and Meneguzzi, “Generalised BDI Planning”.

Table of Contents

- 1 Motivation
- 2 BDI Architecture and Reasoning
- 3 Normative BDI
- 4 Recognising Intent and Norms**
 - Goal Recognition
 - Norm Identification
- 5 The RAD-AI Synthesis
- 6 Conclusions

Reading Implicit Behaviour

Why we need to see what others don't say

- Agents (be they human, or AI/LLM) do not always reveal their goals or norms
 - **Hidden goals:** what is this agent *actually* trying to achieve?
 - **Hidden norms:** what rules govern their behaviour?
- Two mechanisms to surface implicit behaviour:
 - ① **Goal recognition:** infer intent from observed actions
 - ② **Norm identification:** infer governing rules from patterns of compliance and violation
- Both are prerequisites for principled rebellion:
 - You cannot disobey a norm you have not identified
 - You cannot rebel against a goal (or towards an implicit goal) you have not recognised

Reading Implicit Behaviour

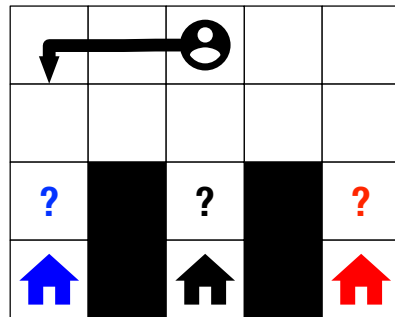
Why we need to see what others don't say

- Agents (be they human, or AI/LLM) do not always reveal their goals or norms
 - **Hidden goals:** what is this agent *actually* trying to achieve?
 - **Hidden norms:** what rules govern their behaviour (even unconsciously)?
- Two mechanisms to surface implicit behaviour:
 - ① **Goal recognition:** infer intent from observed actions
 - ② **Norm identification:** infer governing rules from patterns of compliance and violation
- Both are prerequisites for principled rebellion:
 - You cannot disobey a norm you have not identified
 - You cannot rebel against a goal (or towards an implicit goal) you have not recognised

Goal Recognition

Inferring intent from behaviour

- Given observations $\mathcal{O}$ and candidate goals $\mathcal{G}$: infer the *most likely* goal $g^* \in \mathcal{G}$
- Planning-based**⁹:
 - Landmarks¹⁰: necessary facts in every plan for g
 - Match observed actions to landmarks (fast)
- RL-based**¹¹:
 - Q-functions encode goal proximity
 - Inference: compare Q-values (no planning, no domain model)



⁹Miquel Ramirez and Héctor Geffner. “Plan Recognition as Planning”. In: *IJCAI*. 2009.

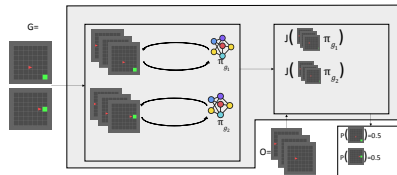
¹⁰Ramon Fraga Pereira, Nir Oren, and Felipe Meneguzzi. “Landmark-Based Approaches for Goal Recognition as Planning”. In: *Artificial Intelligence* 285 (2020).

¹¹Leonardo R. Amado, Reuth Mirsky, and Felipe Meneguzzi. “Goal Recognition as Reinforcement Learning”. In: *AAMAS*. 2023.

Goal Recognition as RL

Scaling without domain models

- Q-value functions encode goal proximity: agents pursuing g increase $Q(s, a, g)$ over time
- **GRAQL**:¹² one tabular Q-function per goal
 - Inference = compare Q-values across goals
 - No domain model at test time
- **DRACO**:¹³ scales to continuous domains
 - Deep Q-network per goal; partial observability, noise
 - Deployed in physical/robotic settings



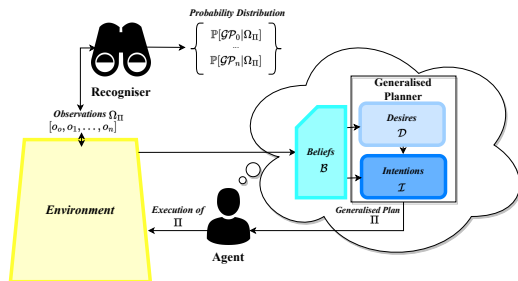
¹²Amado, Mirsky, and Meneguzzi, “Goal Recognition as Reinforcement Learning”.

¹³Ben Nageris, Felipe Meneguzzi, and Reuth Mirsky. “Goal Recognition using Actor-Critic Optimization”. In: *arXiv:2501.01463* (2024).

Compliance Monitoring

GR applied to norm-governed environments

- Observe another agent's actions; ask:
 - **What goal** is being pursued?
 - Does the observed plan **violate any norm**?
 - Is the violation **deliberate** or accidental?
- Supports coordination and counterplanning as by-products^a
- Opens the adversarial question:
 - Can an agent *appear* compliant while rebelling?
 - Deceptive compliance as a research direction



^aMichael T. Cox. "Plan Recognition for Rebel Agent Systems". In: *RAD-AI @ AAMAS*. 2025.

Learning Norms from Behaviour

- GR tells us *what* agents do; norm identification tells us *what rules* shaped it
- **Plan-recognition based**¹⁴: find common constraints across compliant plans
 - Obligations: actions in *every* compliant plan
 - Prohibitions: actions in *no* compliant plan
- **Bayesian**¹⁵: update norm hypotheses from compliant *and* violating behaviour
 - A disobedient agent's behaviour is *evidence* about the norm being violated
 - Can identify *whether* the violation is deliberate

Key insight

Deliberate rebellion requires knowing which norm you are violating. Norm identification is how an observer (or the agent itself) finds out.

¹⁴Nir Oren and Felipe Meneguzzi. "Norm Identification through Plan Recognition". In: *COIN@AAMAS*. 2013.

¹⁵Stephen Cranefield et al. "A Bayesian Approach to Norm Identification". In: *ECAL*. 2015.

Table of Contents

- 1 Motivation
- 2 BDI Architecture and Reasoning
- 3 Normative BDI
- 4 Recognising Intent and Norms
 - Goal Recognition
 - Norm Identification
- 5 The RAD-AI Synthesis**
- 6 Conclusions

Principled Disobedience

A (in)formal account

- ① Agent maintains a belief base and plan library (vanilla BDI)
- ② Agent receives norms and annotates them into plans (ν -BDI)
- ③ Agent pursues goals; selects most compliant available plan
- ④ When *all* plans violate some norm:
 - Identify the minimal norm violation required
 - Execute with explicit record of violation
 - This is *not quite rebellion*, but rather *reasoned* non-compliance
- ⑤ Throughout all of this, agents observe behaviour and:
 - Use goal recognition to identify the pursued goal
 - Use norm identification to infer which norms emerge, and/or get violated
 - Judge whether the violation was intentional or unavoidable

Open Challenges & Ongoing Work

RAD-AI Research Agenda

- **Adversarial goal recognition:** detecting agents that *deceive*
 - Deliberately mimicking compliant behaviour
 - Deceptive planning to evade norm monitoring
- **Norm learning in multi-agent systems**
 - Norms emerge from interaction
 - Reinforcement learning of normative monitoring
- **Normative BDI + LLMs**
 - LLM-based agents: norms encoded implicitly in weights (or hidden system prompts)
 - How do we reason about norm compliance in LLM agents?
 - Value alignment as a special case of norm identification
- **Theory of Mind meets normative reasoning**
 - Recursive modelling of compliant vs. non-compliant agents

¹⁶Shuwa Miura. “On the Depths of Recursive Thinking in Observer-Aware Planning”. In: *RAD-AI @ AAMAS*. 2024.

¹⁷Addison, Larue, and Molineaux, “Towards An Ethical Rebellion System”.

¹⁸Dale Peasley, Zachary Gray, and Sandip Sen. “Disobedience for a Cause: Leveraging Implicit Objectives in User Plans by LLM Surrogates”. In: *AAMAS*. 2025.

¹⁹Meneguzzi, “Compliance, Norms, and Disobedience”.

Open Challenges & Ongoing Work

RAD-AI Research Agenda

- **Adversarial goal recognition:** detecting agents that *deceive* ✓¹⁶
 - Deliberately mimicking compliant behaviour
 - Deceptive planning to evade norm monitoring
- **Norm learning in multi-agent systems** ✓¹⁷
 - Norms emerge from interaction
 - Reinforcement learning of normative monitoring
- **Normative BDI + LLMs** ~¹⁸
 - LLM-based agents: norms encoded implicitly in weights (or hidden system prompts)
 - How do we reason about norm compliance in LLM agents?
 - Value alignment as a special case of norm identification
- **Theory of Mind meets normative reasoning** ✓¹⁹
 - Recursive modelling of compliant vs. non-compliant agents

¹⁶Miura, “On the Depths of Recursive Thinking in Observer-Aware Planning”.

¹⁷Addison, Larue, and Molineaux, “Towards An Ethical Rebellion System”.

¹⁸Peasley, Gray, and Sen, “Disobedience for a Cause: Leveraging Implicit Objectives in User Plans by LLM Surrogates”.

¹⁹Miura, “On the Depths of Recursive Thinking in Observer-Aware Planning”.

Table of Contents

- 1 Motivation
- 2 BDI Architecture and Reasoning
- 3 Normative BDI
- 4 Recognising Intent and Norms
 - Goal Recognition
 - Norm Identification
- 5 The RAD-AI Synthesis
- 6 Conclusions**

Future Work

A research agenda

- Normative reasoning for LLM-based BDI agents
- Adversarial goal recognition and deceptive compliance
- Norm identification from multi-agent observation

Future Work

A research agenda

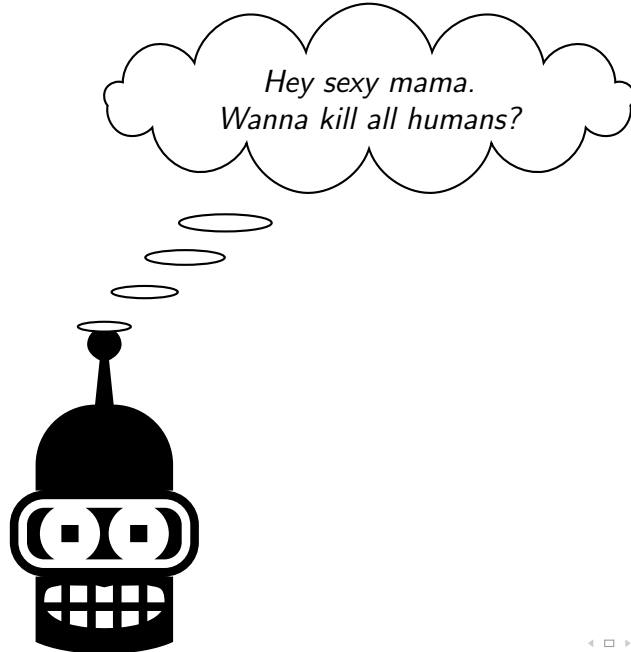
- Normative reasoning for LLM-based BDI agents
- Adversarial goal recognition and deceptive compliance
- Norm identification from multi-agent observation

The game is out there.
It's either play, or get played.

Acknowledgements

Not my genius alone

- **BDI / Generalised Planning:** Ramon Fraga Pereira, Nir Oren
- ν -**BDI:** Odinaldo Rodrigues, Michael Luck, Wamberto Vasconcelos, Nir Oren
- **Goal Recognition:** Ramon Fraga Pereira, Leonardo Amado, Ben Nageris, Mor Vered, Miquel Ramirez, Reuth Mirsky, and **many others**
- **Norm Identification:** Nir Oren, Stephen Cranefield
- **Funding:** EPSRC, CNPq, FAPERGS, Google



Thank you!

`felipe.meneguzzi@abdn.ac.uk`

1495



UNIVERSITY OF
ABERDEEN