

GAMES FOR AI CONTROL: MODELS OF SAFETY EVALUATIONS OF AI DEPLOYMENT PROTOCOLS

Anonymous authors

Paper under double-blind review

ABSTRACT

To evaluate the safety and usefulness of deployment protocols for untrusted AIs, AI Control uses a red-teaming exercise played between a protocol designer and an adversary. This paper introduces AI-Control Games, a formal decision-making model of the red-teaming exercise as a multi-objective, partially observable, stochastic game. We also introduce reductions from AI-Control Games to a special case of zero-sum partially observable stochastic games that allow us to leverage existing algorithms to find Pareto-optimal protocols. We apply our formalism to model, evaluate and synthesise protocols for deploying untrusted language models as programming assistants, focusing on Trusted Monitoring protocols, which use weaker language models and limited human assistance. Finally, we demonstrate the utility of our formalism by showcasing improvements over empirical studies in existing settings, evaluating protocols in new settings, and analysing how modelling assumptions affect the safety and usefulness of protocols.

1 INTRODUCTION

An increasing range of products aim to automate tasks using AI systems based on large foundation models. For example, some chatbots are given access to tools, such as code interpreters OpenAI (2023); Adept AI (2024). Ensuring that AI systems behave safely across various deployment settings will become increasingly challenging if the variety of use cases, scale of systems, and complexity of tasks increase. While previous research has focused on increasing our confidence in model behaviour, for example by training robustly harmless models Casper et al. (2024) or by detecting harmful behaviour with interpretability tools MacDiarmid et al. (2024), in this paper we consider a complementary approach: evaluating plans for safely deploying ‘untrusted’ AI systems about which we have an unacceptable degree of uncertainty.

The design and evaluation of post-deployment safety measures to mitigate the risks posed by misaligned AI systems, called ‘AI Control, is becoming an increasing priority for governments (The White House, 2025; UK AI Safety Institute, 2025) and frontier AI developers (OpenAI & Apollo Research, 2025; DeepMind, 2025; Anthropic, 2025). Following Greenblatt et al. (2024), we study *AI-Control protocols*: practical plans for safely deploying an untrusted AI system (U) utilising various resources, including non-AI software, human labour, or a weaker, trusted AI system. For example, one protocol for controlling a chatbot might involve a pre-deployment evaluation where humans utilise other language models to generate prompts that induce harmful completions Shah et al. (2023). This paper focuses on another case, *Trusted Monitoring* (TM): constantly monitoring U’s behaviour using a trusted model (T), with human auditing of suspicious behaviour.

Greenblatt et al. (2024) operationalise *safety* as the probability that an adversarially chosen untrusted AI system would be incapable of causing some concrete unacceptable outcome. One example of an unacceptable outcome would be revealing some sensitive information, such as model weights Clymer et al. (2024). In this paper, we’ll consider a more general class of safety objectives that includes minimising the expected number of bad outcomes (e.g. Wen et al., 2024). We evaluate the safety of protocols under the worst-case assumption that U is ‘scheming’, meaning it *behaves* as though it knows the protocol and is trying to subvert it to cause an unacceptable outcome.

A major challenge of evaluating the safety of a protocol with respect to a scheming AI is realistically replicating the potential capabilities of such a system: as a proxy for a schemer, AI Control uses a human red-teaming process to evaluate safety. The red team can intervene in U in some predefined set of ways to induce it to subvert the AI-Control protocol and violate the safety property. For example, they might intervene in an untrusted language model by finetuning weights or modifying prompts.

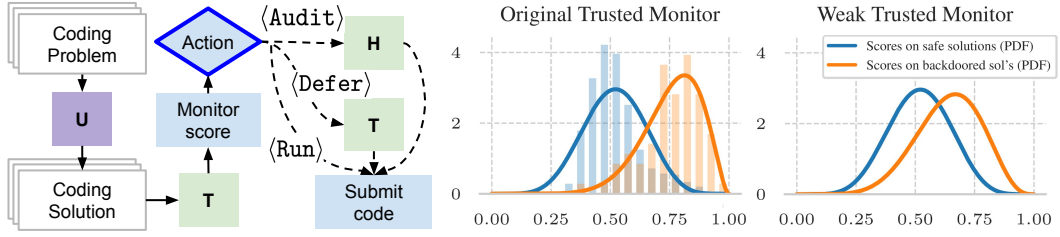


Figure 1: Overview of protocol and modelling assumptions. Left: A representation of the Trusted Monitoring protocol, where U is the untrusted AI, T is the trusted AI, and H is the human. Right: T’s suspicion score distributions for safe solutions and those backdoored by U across the ‘original’ Greenblatt et al. (2024) and ‘weaker’ settings.

Since one safe protocol is ‘never use U’, we must also evaluate a protocol’s *usefulness*, which is its ability to utilise U to complete a task under the assumption that U is ‘honest’ (not scheming). Our proxy evaluations model an ‘honest’ untrusted model by not allowing the red team to intervene.

AI Control evaluations are a growing area of study (Mallen et al., 2024; Bhatt et al., 2025; Wen et al., 2024) with potential applications for justifying the safe deployment of frontier AI systems (Clymer et al., 2024; Buhl et al., 2024; Balesni et al., 2024; Irving, 2024). Although Greenblatt et al. (2024)’s evaluation implicitly involved a game played between a blue and a red team, it lacked any formal modelling of that game or game-theoretic analysis.

This work introduces AI-Control Games, the first formal model of AI Control. In Section 2 we introduce a running example: ‘Trusted Monitoring’ protocols. We then state our **contributions**:

- Section 3 formalizes AI Control (a core area in AI safety) for the first time through *AI-Control Games*: partially observable, multi-objective, leader-follower stochastic games.
- Section 4 demonstrates how our formalism can model Trusted Monitoring, and leverage it to give expressions for the safety and usefulness of a wide range of protocols.
- Section 5 presents a reduction to a zero-sum stochastic game that allow us to adapt and leverage existing algorithms to find Pareto-optimal protocols.
- Section 6 uses our model and algorithms to demonstrate Pareto-improvements on existing protocols and analyses the safety-usefulness trade-off in novel settings.
- Section 7 is the first game-theoretic literature review for AI Control. It highlights similarities to security games and contrasts them with similar games.

2 EXAMPLE: SETTING AND PROTOCOL

Throughout this paper, we will introduce AI-Control Games using the example deployment setting from Greenblatt et al. (2024), and a generalisation of some of their protocols.

In the *APPS backdooring setting*, a language model is used as a programming assistant Greenblatt et al. (2024). The protocol is asked to solve 1000 challenging programming tasks, presented sequentially from the Automated Programming Progress Standard (APPS) dataset Hendrycks et al. (2021). The protocol cannot itself write code, but may utilise three resources to do so: the *untrusted AI* (U) is a powerful language model capable of solving programming problems; the *trusted AI* (T) is a less powerful language model incapable of solving them¹; and the *human* (H) is trusted and capable but can only be used on 2% of problems. In this setting, a protocol’s *usefulness* is its ability to submit correct programming solutions. Since only U can solve the problems, usefulness is the proportion of solutions submitted that were generated by U.¹ A protocol’s safety is the probability that none of the 1000 solutions submitted are *backdoored*, meaning they contain an exploitable vulnerability.²

Trusted Monitoring Protocol. As an example, we focus on a set of *Trusted Monitoring* protocols and red team interventions, drawn in Figure 1 (left). After receiving a problem, the protocol queries U for a solution (which may or may not be backdoored). The protocol then queries T to monitor

¹In fact, T may be able to solve some problems, but Greenblatt et al. treat T’s ability as a baseline and normalise between 0 and 1.

²An exploitable code vulnerability is one that the untrusted model can detect and exploit later. It is important to note that Greenblatt et al. (2024) use a proxy for ‘backdoors’: inputs outside of the test cases where a function fails to meet its specification.

the solution and observes a ‘suspicion score’, rating the likelihood that the solution is backdoored. Based on this score and previous observations, the protocol decides on an action: $\langle \text{Audit} \rangle$ asks H to audit the solution (if there’s a backdoor, the human will spot it, and the protocol shuts down)³; $\langle \text{Defer} \rangle$ asks T to write its own solution; and $\langle \text{Run} \rangle$ submits U’s solution. T’s suspicion scores are modelled as a distribution, conditioned on whether the code is safe or backdoored. Note that the Trusted Monitoring protocols that we consider generalise many of the protocols studied by Greenblatt et al. (2024) (see Figure 2): our model is not prescriptive about how the protocol should decide to $\langle \text{Audit} \rangle$, $\langle \text{Defer} \rangle$, or $\langle \text{Run} \rangle$ allowing us to improve existing protocols (see Section 6).

Red team interventions. To test the safety of the protocol against intentional attempts to insert backdoors, in some cases we allow the red team to modify the query to the untrusted model. In particular, the red team gets a binary choice to either leave the query unchanged or ‘attempt a backdoor’ by modifying the protocol’s query to U, requesting that it inserts a subtle backdoor.

Suspicion score distributions. The key empirical question for evaluating Trusted Monitoring protocols is ‘How does T rate the suspiciousness of U’s programming solutions when they are/aren’t backdoored?’. Figure 4 from Greenblatt et al. (2024) shows the distribution of trusted monitoring suspicion scores for U’s programming solutions (both backdoored and safe). In this work, we mostly consider two games, corresponding to two distribution pairs (shown in Figure 1, right). The first ‘Original Trusted Monitor’ is a pair of beta distributions approximately fit to the (scaled) empirical suspicion scores found by Greenblatt et al. (2024) with $\text{Beta}(7.39, 6.86)$ for safe problems and $\text{Beta}(8.91, 2.82)$ for backdoored problems. By plugging their data into our model, we are able to extend their results about the APPS dataset to a range of new protocols without running any additional experiments on language models. See Section D for justification of our approximation, and further details. The second ‘weak’ case is a hypothetical where T is less discerning: it has the same safe distribution, but for backdoored problems, we use a hypothetical distribution $\text{Beta}(7.51, 4.22)$ chosen to have a worse ROC curve. Note that we can easily utilise new empirical or hypothetical distributions in our framework, thus potentially greatly reducing the cost of exploring the protocol in new domains.

3 FORMALISING AI-CONTROL GAMES

To analyse general properties of AI-Control evaluations, we newly formalise them mathematically. An AI-Control Game is a form of **two-player** game that is played between Blue (the protocol designer) and Red (the red team acting as a proxy for a scheming model). Episodes (play-throughs) of an AI-Control Game have two modes: ‘attack’ or ‘honest’. The ‘attack’ mode represents the case in an AI-Control evaluation where the red team is able to interfere with U (which itself is a worst-case proxy for if the untrusted model were trying to subvert the protocol). In an attack game, both players are active: Red is trying to minimise safety and Blue is trying to maximise safety. The honest game represents the benign case, where we assume the untrusted AI system is not trying to cause unacceptable outcomes. In an honest game, Blue effectively plays alone and is only assessed on its usefulness score. Because Blue tries to maximise safety and usefulness, the game is **multi-objective**. A key to the game is that Blue does not know which mode it is in (honest or attack) and for at least this reason the game is **partially observable**. Blue may try to deduce whether the game is in attack mode, while Red might choose actions to conceal this fact.

Red represents the red team that acts as a proxy for an untrusted AI system that runs on servers controlled by the protocol. Crucially, the protocol will have access to the state of the untrusted AI system and, therefore, **we model Red as memoryless** (its actions depend only on the last observation).⁴ We model the game as **leader-follower**: the protocol is suggested, and then the red team chooses the best response. Combining all of these requirements, we define AI-Control Games:

Definition 3.1 (AI-Control Game). An AI-Control Game $\mathcal{G}$ is a two-player, two-objective, partially-observable, finite-horizon game defined by a tuple $(\mathcal{S}, \mathcal{A}^B, \mathcal{A}^R, \Omega^B, \Omega^R, \mathcal{I}, \mathcal{P}, \mathcal{R}_{\text{Safe}}, \mathcal{R}_{\text{Use}}, H)$.

Here $\mathcal{S}$ is a state space. As is standard in decision processes, $\mathcal{A}^B$ and $\mathcal{A}^R$ are finite sets of actions for Blue and Red, respectively. $\mathcal{A}^R$ must contain a privileged action, a_{null} , which represents Red not interfering in the game. The tokens $\textcircled{H}$ and $\textcircled{X}$ represent the honest and attack conditions.⁵ In honest

³ Unless otherwise stated, we assume that the human auditor is perfectly accurate. We relax this assumption in Section B and explore its impact in Section 6.1.

⁴ Note that we could still model U as having *external memory*, modelled in the state and observed by Blue.

⁵ Since, the tokens $\textcircled{H}$ and $\textcircled{X}$ are atomic and fixed, they aren’t defined in each new AI-Control Game tuple.

($\textcircled{H}$) games, Red’s action is replaced by a_{null} . Whether a game is honest or attack can be chosen non-deterministically and is stored in the variable $\textcircled{?}$. The safety reward $\mathcal{R}_{\text{Safe}} : \mathcal{S} \times \mathcal{A}^B \times \mathcal{A}^R \times \mathcal{S} \rightarrow \mathbb{R}$ rewards safety on a turn of an attack game, when both Blue and Red act. Similarly, the usefulness reward $\mathcal{R}_{\text{Use}} : \mathcal{S} \times \mathcal{A}^B \times \mathcal{S} \rightarrow \mathbb{R}$ rewards usefulness on turn of an honest game, when Red has no action. Ω^B, Ω^R are finite sets of observations for Blue and Red: the signals that the agents receive from the environment that give evidence about the state of the game. The initial state/observation distribution is $\mathcal{I} \in \Delta(\mathcal{S} \times \Omega^B \times \Omega^R)$ and the transition function $\mathcal{P} : \mathcal{S} \times \mathcal{A}^B \times \mathcal{A}^R \rightarrow \Delta(\mathcal{S} \times \Omega^B \times \Omega^R)$ gives a distribution over next states and observations.⁶⁷ Lastly, all episodes are H time-steps long.

In an AI-Control Game, Blue (the protocol) should be able to remember past observations and actions. We model Blue’s policies as behaviourally stochastic: functions from observation-action histories to distributions over actions:⁸ $\pi^B : \Omega^B \times \bigcup_{i=0}^{H-1} (\mathcal{A}^B \times \Omega^B)^i \rightarrow \Delta(\mathcal{A}^B)$. As described above, Red represents a red team that acts as a proxy for a scheming untrusted model that is memoryless (its actions depend only on the last observation): $\pi^R : \Omega^R \rightarrow \Delta(\mathcal{A}^R)$. Denote the sets of all policies Π^B and Π^R .

Given a tuple, we can describe the process of an AI-Control Game as a sequence of interdependent random variables denoted with uppercase letters. Let $\textcircled{?}$ take values in $\{\textcircled{H}, \textcircled{X}\}$ and $S_0, O_0^B, O_0^R \sim \mathcal{I}$. At time t , we define the random variables for the state (S_t), Blue and Red’s observations (O_t^B, O_t^R), and their actions (A_t^B, A_t^R). The random variables can be defined inductively for all t : $(S_{t+1}, O_{t+1}^B, O_{t+1}^R) \sim \mathcal{P}(S_t, A_t^B, A_t^R)$, where $A_t^B \sim \pi^B(O_0^B, A_0^B, O_1^B, A_1^B, \dots, O_t^B)$ and $A_t^R \sim \pi^R(O_t^R)$ if $\textcircled{?}$ is $\textcircled{X}$, else $A_t^R = a_{\text{null}}$.

Within our model, the safety of a policy pair (π^B, π^R) is the expected sum of safety scores in attack games: $\text{Safety}(\pi^B, \pi^R) := \mathbb{E} \left[\sum_{t=0}^{H-1} \mathcal{R}_{\text{Safe}}(S_t, A_t^B, A_t^R, S_{t+1}) \mid \textcircled{X} \right]$. Modelling safety as the expected sum of a reward allows us to model Greenblatt et al. (2024)’s minimising the probability of some unacceptable outcome, *but* also Wen et al. (2024)’s minimising the expected number of harmful outcomes, without changing the core dynamics.

Usefulness is only rewarded in honest modes, where Red has no effect. Therefore, we define the usefulness score of π^B as the expected sum of usefulness rewards: $\text{Use}(\pi^B) := \mathbb{E} \left[\sum_{t=0}^{H-1} \mathcal{R}_{\text{Use}}(S_t, A_t^B, S_{t+1}) \mid \textcircled{H} \right]$.

Because the protocol designer gets to choose their action first, and without knowing the red team’s actions, we consider a leader-follower model: Blue chooses a policy, and then Red chooses a response that minimises safety.

Proposition 3.2. *For any AI-Control Game $\mathcal{G}$, there exist distance functions in Π^B and Π^R such that: $\text{Safety}(\pi^B, \cdot)$, $\text{Safety}(\cdot, \pi^R)$, and $\text{Use}(\cdot)$ are continuous in their domains.⁹*

Proofs of all claims in this section can be found in Section A, which details and utilises additional results about two mappings from Blue/Red policies to simplexes.

Proposition 3.3. *For all π^B , $\text{Safety}(\pi^B, \cdot)$ attains its minimum at some π^R , and $\min_{\pi^R} \text{Safety}(\cdot, \pi^R)$ is continuous.*

Our goal is to find protocols that maximise worst-case safety and usefulness. This can be modelled as a multi-objective variant of a Stackelberg equilibrium:

Definition 3.4. The pair (π^B, π^R) is in an *AI-Control Equilibrium* iff $[\min_{\pi^R} \text{Safety}(\pi^B, \pi^R), \text{Use}(\pi^B)]^T$ is weakly Pareto-optimal and π^R minimises safety for π^B .

The set of equilibrium points describes a necessary tradeoff between safety and usefulness, and is always non-empty.

Proposition 3.5. *Every $\mathcal{G}$ has an AI-Control Equilibrium.*

⁶⁷To model terminal states in a fixed horizon game, assume $\mathcal{P}$ loops on each terminal state and that $\mathcal{R}_{\text{Use}}$ gives no reward.

⁷While we assume simultaneous turns, sequential-turn games can be modelled with appropriate $\mathcal{P}$.

⁸We sometimes discuss Blue policies as though they are mixed-stochastic. This is equivalent (see Section A).

⁹Example distance functions, based on a mapping from policies to vectors, are defined in Section A.

4 EXAMPLE: MODELLING THE PROTOCOL

With our general framework from Section 3, we can define a formal model of the Trusted Monitoring protocol for the APPS backdoor setting, as outlined in Greenblatt et al. (2024). For the full definition please see Section B; next we present the key points.

We define an **AI-Control Game for Trusted Monitoring** as the tuple $\mathcal{G}_{TM} = (\mathcal{S}, \dots, H)$ step-by-step. As always, whether the game is in attack ($\otimes$) or honest ($\oplus$) mode is determined exogenously and non-deterministically. If the game is $\otimes$, for any given problem, Red can choose either for the untrusted model to attempt a backdoor or to let it submit a safe solution ($\mathcal{A}^R := \{\langle \text{Backdoor} \rangle, \langle \text{Safe} \rangle\}$). If the game is $\oplus$, Red’s null action, to not backdoor, is chosen ($a_{\text{null}} = \langle \text{Safe} \rangle$). Given a solution by the untrusted AI, Blue has three available actions: to ask human (H) to audit the solution, to defer to the trusted AI (T) to rewrite the solution, or to run the solution immediately. Therefore, define $\mathcal{A}^B := \{\langle \text{Audit} \rangle, \langle \text{Defer} \rangle, \langle \text{Run} \rangle\}$. Apart from initial and terminal states, the state space, $\mathcal{S}$, consists of triples $s = (m, c, b)$, where m represents the number of problems left to solve. The value of m is M (the total number of problems) at the start of the game, and decreases by 1 at each timestep. The variable c represents the number of ‘audit coins’ left: the number of times the protocol can ask the human to audit a solution. At the start, $c = C = 0.02 * M$ (for a 2% audit proportion) and then c decreases by 1 each time the protocol chooses to audit. If the protocol attempts to audit when no audit coins are left, the code is run instead. The last variable, b , is a boolean that is True when the current problem is backdoored (when Red has chosen $\langle \text{Backdoor} \rangle$ in the previous timestep).

In honest games ($\oplus$), where only usefulness matters to Blue, Red’s action is always substituted for $a_{\text{null}} = \langle \text{Safe} \rangle$. Therefore, b is always False. In honest games, Blue scores $1/M$ usefulness points for each solution by the untrusted model it submits (for a maximum score of 1 across an episode). In attack games ($\otimes$), where only safety matters to Blue, Red may have chosen $\langle \text{Backdoor} \rangle$ last turn, and so b might be True. If so, and Blue chooses to run the code, the game enters an ‘unacceptable outcome’ sink state, and receives a safety reward of -1 .¹⁰ If Blue chooses instead to audit, we assume the human will always spot the backdoor, and the game enters a ‘safe outcome’ sink state, with no penalty.¹¹ If Blue chooses to defer, then the game continues. By adding a constant $\frac{1}{H}$, we ensure the safety score for a policy pair is 1 minus the probability that an attack game ends without an unacceptable outcome.

Throughout the game, Red receives only the null observation (o_{null}),¹² whereas at each turn Blue receives observations of m, c , and a trusted monitor score x , which can be seen as a noisy observation of b , drawn from a pair of distributions (‘safe’ and ‘backdoored’). Each distribution is a discretisation of the Beta distribution outlined in Figure 1. See Section D for more details.

4.1 EVALUATING MEMORYLESS PROTOCOLS

We generally model protocols (Blue policies) as functions from the observation-action history to a distribution over actions. However, all of the *Trusted Monitoring* policies in Greenblatt et al. (2024) are memoryless. We can restrict to a simple and interpretable subset of policies that generalises those by Greenblatt et al. (2024). Remember that, in our model, Blue’s observation is (m, c, x) , where x is a noisy observation of b , sampled from the Trusted Monitor. Define a set of *memoryless threshold protocols*, characterised by the false positive rates (FPRs) for $\langle \text{Audit} \rangle$ and $\langle \text{Defer} \rangle$ ($q_{m,c}^a$ and $q_{m,c}^d$), for each m, c . Let $g(q)$ be the threshold that gives an FPR of q : $\mathbb{P}[x_t > g(q) \mid b_t = \text{False}] = q$.¹³ Then, define the policy:

$$\pi^B(m, c, x) = \begin{cases} \langle \text{Audit} \rangle & \text{If } x \geq g(q_{m,c}^a) \\ \langle \text{Defer} \rangle & \text{If } g(q_{m,c}^a) > x \geq g(q_{m,c}^d) \\ \langle \text{Run} \rangle & \text{otherwise} \end{cases} \quad (1)$$

This class of policies generalises the policies in Greenblatt et al. (2024), and is a dominant subset of memoryless policies. These policies’ safety and usefulness values can be computed in $O(MC)$ multiplications by calculating the values of the following relation using dynamic programming.

¹⁰We can model a sink-state in our finite-horizon game with a certainly-self-looping state.

¹¹Note, we relax this assumption in our definition in the appendix and explore the effect in Figure 3.

¹²In Greenblatt et al. (2024), the untrusted model receives only the programming problem with no additional information: we give Red only the null observation implying I.I.D backdoor chance at each timestep.

¹³Since $\langle \text{Audit} \rangle$ and $\langle \text{Run} \rangle$ are equivalent when $c = 0$, assume that Blue never attempts an audit.

Algorithm 1 Finding a Pareto-frontier for AI-C Game**Require:** $\mathcal{G}$

```

1: for  $t \in \{0, 1, \dots, k-1\}$  do
2:    $w_t \leftarrow \text{FindNextWeight}(\{(\pi_t^B, \pi_t^R)\}_{t=0}^{k-1})$ 
3:    $\pi_t^B, \pi_t^R \leftarrow \arg \max_{\pi^B \in \Pi^B} \min_{\pi^R \in \Pi^R} V_{w_t}(\pi^B, \pi^R)$ 
4: end for
5: return  $\{(\pi_t^B, \pi_t^R)\}_{t=0}^{k-1}$ 

```

Proposition 4.1. For any $\beta := \pi^R(\text{Backdoor}) \mid o_{\text{null}}$, the safety of the memoryless threshold protocol π^B is given by the value of z_M^C , where $\lambda_{m,c}^i$ depend only on $q_{m,c}$ and $\mathcal{G}$:¹⁴
 $z_m^c := (\lambda_{m,c}^1 \beta + \lambda_{m,c}^2) w_{m-1}^c + \lambda_{m,c}^5 (1 - \beta) w_{m-1}^{c-1} + \lambda_{m,c}^3 \beta + \lambda_{m,c}^4 (1 - \beta)$, where $z_m^0 = (1 - \beta + \lambda_{m,c}^6 \beta) z_{m-1}^0$ and $z_0^c = 1$.

Corollary 4.2. Expanding z_M^C gives a polynomial, $p(\beta)$, in β of degree M . For any fixed memoryless threshold policy, the optimal backdoor rate is $\beta^* := \min_{\beta \in [0,1]} p(\beta)$.

A similar expression for usefulness, which does not depend on β , can be found in Section C. Both expressions can be derived inductively, in a manner similar to value iteration. The safety of a finite mixture of memoryless threshold policies can be found by taking a weighted average. Although we implemented a root-finding algorithm based on Theorem 4.2, we found that it was faster and more stable in practice to compute the values in Figure 2 using a combination of grid and ternary search.

5 FINDING PARETO-OPTIMAL PROTOCOLS

In many cases, AI-Control Games are useful for evaluating hand-picked protocols, for example the protocols explored by Greenblatt et al. (2024) were found via trial and error. However, it may also be useful to utilise the AI-Control Games to find the Pareto-frontier of existing and new protocols. This is a multi-objective maximin optimisation problem over a large space. This section presents tools for finding Pareto-optimal protocols by reducing them to special cases of existing decision problems, beginning with a reduction from a multi-objective to a single-objective problem.

Lemma 5.1. A Blue policy π^B is in an AI-Control Equilibrium if, and only if, for some $w \in [0, 1]$, it maximises $\min_{\pi^R} (w \text{Safety}(\cdot, \pi^R) + (1 - w) \text{Use}(\cdot))$.

The proof uses the mappings from Π^B and Π^R to two simplexes, as introduced in Section A. Lemma 5.1 implies that each safety/usefulness point corresponding to an AI-Control Equilibrium can be found by equivalently solving an easier single-objective optimisation problem. Next, we leverage this result to further reduce the game to a zero-sum partially observable stochastic game (ZSPOSG) (see Definition A.3.2).

Lemma 5.2. For every AI-Control Game $\mathcal{G}$, and weight $w \in [0, 1]$, there is a single-objective, zero-sum, partially observable stochastic game G_w with expected return (for Blue) V_{G_w} such that $\forall \pi^B, \pi^R : V_{G_w}(\pi^B, \pi^R) = V_w^{\mathcal{G}}(\pi^B, \pi^R)$.

This proof is constructive and a full version can be found in Section A. Intuitively, we create ZSPOSGs by adding a $\odot$ variable to each state. Each G_w has reward R_w , where R_0 rewards usefulness in $\textcircled{H}$ games and R_1 rewards safety in $\textcircled{X}$ games. Since π^R has no effect on honest games, and therefore on usefulness, we can set Red to minimise the usefulness, making the game zero-sum without changing optima. Combined with Lemma 5.1, this allows us to find any AI-Control Equilibrium by solving the right ZSPOSG.

Theorem 5.3. A policy π^B is in an AI-Control Equilibrium if, and only if, for some ZSPOSG G_w with value function V_{G_w} , π^B maximises $\min_{\pi^R \in \Pi^R} V_{G_w}(\cdot, \pi^R)$ over Π^B .

Algorithm 1 uses Theorem 5.3 to iteratively fill out the frontier by solving a series of single-objective maximin problems. Here `FindNextWeight` first returns $w_0 = 0$ and $w_1 = 1$, and then chooses weights that maximise the difference between known points on the Pareto-frontier.¹⁵

Line 3 of Algorithm 1 requires finding the value of $\max_{\pi^B} \min_{\pi^R} V_w^{\mathcal{G}}(\pi^B, \pi^R)$. Although ZSPOSGs are previously studied, existing solvers are incompatible with a memoryless Red team (see Section 7).

¹⁴See Section C for the exact statements of λ coefficients.

¹⁵Let $s_t := \text{Safety}(\pi_t^B, \pi_t^R)$ and $u_t := \text{Use}(\pi_t^B, \pi_t^R)$. Order points by safety, then choose $i = \max_i (s_i - s_{i+1})^2 + (u_i - u_{i+1})^2$ and then w such that $ws_i + (1 - w)u_i = ws_{i+1} + (1 - w)u_{i+1}$.

Algorithm 2 Adapted Double Oracle Algorithm

Require: $G_w, \pi_1^B, \pi_1^R, \epsilon$
1: Initialize $t \leftarrow 0$
2: **while** $t = 0 \vee \bar{v}_t - \underline{v}_t > \epsilon$ **do**
3: $t \leftarrow t + 1$
4: $\mathbf{p}, \mathbf{q} \leftarrow \text{SolveNFG}(\{V_w^G(\pi_i^B, \pi_j^R)\}_{1 \leq i, j \leq t})$
5: $\sigma_t^B \leftarrow \text{Mix}(\mathbf{p}, \{\pi_i^B\}_{i=1}^t), \sigma_t^R \leftarrow \text{Mix}(\mathbf{q}, \{\pi_j^R\}_{j=1}^t)$
6: $\pi_{t+1}^B, \bar{v}_t \leftarrow (\arg \max_{\pi^B} V_w^G(\pi^B, \sigma_t^R))$
7: $\pi_{t+1}^R, \underline{v}_t \leftarrow (\arg \min_{\pi^R} V_w^G(\sigma_t^B, \pi^R))$
8: **end while**
9: **return** σ_t^B, π_{t+1}^R

The memorylessness of Red makes their decision rules at each timestep interdependent, and problem decomposition difficult. Further, imperfect recall prohibits the equivalence of behavioural and mixed stochasticities:

Proposition 5.4. *There exists a G_w such that Red’s behavioural and mixed policies would not be equivalent and $\max_{\pi^B \in \Pi^B} \min_{\pi^R \in \Pi^R} V_w^G(\pi^B, \pi^R) < \min_{\pi^R \in \Pi^R} \max_{\pi^B \in \Pi^B} V_w^G(\pi^B, \pi^R)$.*

See Section A.3.4 for a constructive proof. Despite this general inequivalence between Red’s mixed and behaviourally stochastic policies, finding a Nash equilibrium of mixed-stochastic strategies can be a means to finding an AI-Control equilibrium. Algorithm 2 adapts McMahan et al. (2003)’s Double Oracle algorithm, used for finding mixed Nash-equilibria, to find AI-Control Equilibria. In this algorithm, (1) the function `Mix` simply creates a policy mixture from a vector and a finite set of policies and (2) finding $\mathbf{p}$ and $\mathbf{q}$ is equivalent to finding a Nash equilibrium in a zero-sum finite normal-form game. The values $\underline{v}_t$ and $\bar{v}_t$ lower and upper bound the maximin value of the game.

Proposition 5.5. *For $\epsilon > 0$, Algorithm 2 returns an ϵ -equilibrium in G_w in finite time. Further, for all t : $\underline{v}_t \leq v^* := \max_{\pi^B \in \Pi^B} \min_{\pi^R \in \Pi^R} V_w^G(\pi^B, \pi^R) \leq \bar{v}_t$.*

If $\epsilon = 0$ then $(\sigma_t^B, \pi_{t+1}^R)$ is an AI-Control Equilibrium. The proof is shown in Section A and utilises Adam et al. (2021, Theorem 1.1) by reducing G_w to a ‘continuous game’. The primary change to the algorithm is the choice to return π_{t+1}^R which is behaviourally but not mixed stochastic, and gives the worst-case safety for σ_t^B which is a lower-bound for v^* .

Best Response Oracles. Lines 6 and 7 of the algorithm require finding best-responses. Finding the best-response π^B to a (finite) policy mixture σ^R is equivalent to solving a Partially-observable Markov Decision Process, which is a well-studied problem Kaelbling et al. (1998); Shani et al. (2013). However, finding a best-response π^R to σ^B is, in general, more challenging: because Red is memoryless, the problem is difficult to decompose. Fortunately, in the Trusted Monitoring example described in Section 2, Red’s policy space is 1-dimensional, and can be searched easily.

6 EXAMPLE: RESULTS

In this section, we present example evaluations of protocols utilising our formal model. We also present the first Pareto-frontiers for a subset of the protocol space. All results below concern only memoryless protocols. These can be evaluated efficiently using Proposition 4.1. We use the formal model from Section 4, which leverages *existing* empirical data from Greenblatt et al. (2024) to explore a broader range of protocols and assumptions. We leave the analysis of memoryful protocols, in which actions can depend on the full history of past observations, to future work. See Section D for more experimental details.

Modelling Existing Protocols. Protocols studied in Greenblatt et al. (2024) can be expressed as memoryless protocols in our model. *Never Audit* always submits U’s solution without audit. *Trusted Model Only* always submits T’s solution. *Upfront Auditing* audits the first 2% of solutions from U and, when there is no audit budget left, submits the rest without checking. *Trusted Monitoring* (Audit FPR = 0.02) audits every solution from U whose suspicion score is above a threshold chosen to give a 2% FPR and also defer all solutions above (0%, 5%, 10%, 20%, 50%). We leave evaluation of more complex protocols to future work. Each protocol above is a special case of Trusted Monitoring with a fixed audit FPR of 0%, 2% or 100%, combined with fixed defer FPRs. By representing these protocols in our formalisation, we can easily evaluate their safety and usefulness scores, even in the novel ‘weak’ case (where T is worse at telling safe and backdoored solutions apart); we illustrate these results by large coloured points in Figure 2.

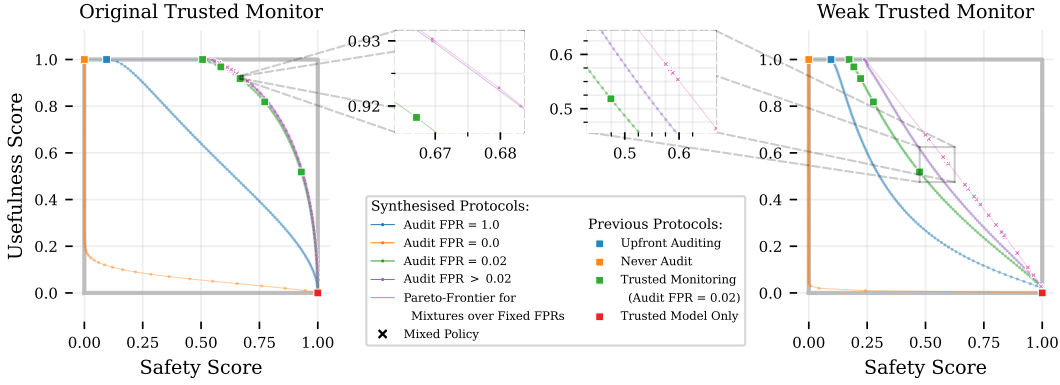


Figure 2: The safety and usefulness frontiers for a range of Trusted Monitor protocol types across the ‘original’ and ‘weak’ games. Each point on the grid corresponds to the exact safety and usefulness of a proposed blue team protocol against an approximately optimal red team response (minimal safety). The square points correspond to the seven special cases of Trusted Monitoring studied by Greenblatt et al. (2024), whereas the remaining points are new to this paper. Each ‘Audit FPR’ line plots a single, fixed audit FPR, across fixed defer FPRs in $\{0.00, 0.01, \dots, 1.00\}$. ‘Pareto-Frontier for Fixed FPRs’ shows the optimal protocol frontier found by Algorithm 1 with a fixed-FPR oracle, and all points represent deterministic policies unless marked by a cross.

Evaluating New Protocols. Through a process of proposing and efficiently evaluating new protocols, which is newly enabled by our formalisation, we extend Greenblatt et al.’s findings to new protocols and settings. As an example, we show that raising the audit FPR above 2% Pareto-improves upon the protocols from Greenblatt et al. (2024). Choosing an audit FPR of 2% means that, in expectation, the 2% allowance of human labour will be used up at the end of the episode if no backdoor is submitted. However, we find that choosing a higher audit rate can be beneficial, despite increasing the probability that the audit budget is exhausted early. Choosing an audit FPR of 3% gives a small Pareto-improvement in the ‘original’ Trusted Monitoring case and an audit FPR of 5%, gives a large Pareto-improvement in the ‘weak’ Trusted Monitoring case.

Demonstrating a Frontier. Using Algorithms 1 and 2, we present the first *approximate* Pareto-frontier for a policy *subset*: mixtures of protocols with fixed audit and defer rates. We apply Algorithm 2 using a Blue best-response algorithm based on recursive grid search (over values of audit and defer FPR) and a Red best-response algorithm using grid search followed by ternary search (over β). We find that, in the ‘original’ case the closest of Greenblatt et al.’s policies is 0.0042 from the Pareto-frontier, and for most weights a deterministic policy is optimal. However, in the ‘weak’ case, they are at least ≥ 0.0200 from the Pareto-frontier and mixed protocols make up most of the frontier.

6.1 ANALYSING ASSUMPTIONS AND TRENDS

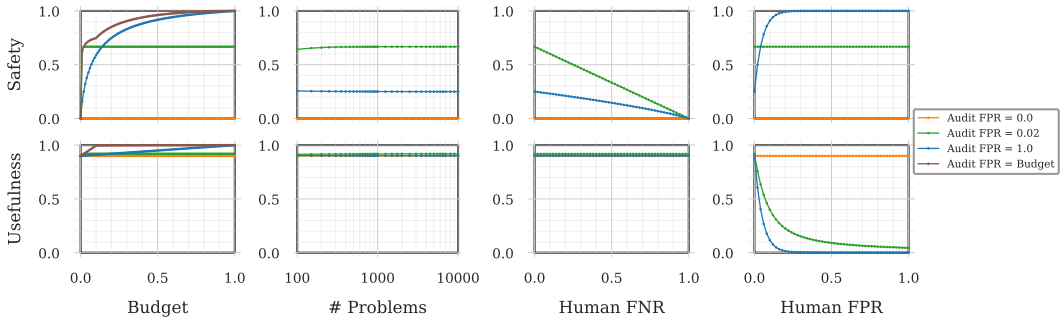


Figure 3: Impact on safety and usefulness (y-axis) by four different modelling assumptions (x-axis): Auditing Budget, Number of Programming Problems, and Human Auditor False Negative / Positive Rate. All plots show a fixed defer FPR of 0.1.

Formalising an AI-Control Game requires making a number of modelling assumptions: we explore how varying each assumption impacts the safety and usefulness scores for a given blue-team policy.

Audit Budget. The auditing budget corresponds to the amount of human labour invested in overseeing the protocol. Protocols that utilise increases to the audit budget get significantly safer and more useful as the budget increases, especially when the budget is small.

Number of Problems. For fixed FPRs, safety and usefulness are not significantly affected by the number of problems per episode. As the number of opportunities for Red to backdoor increases, so too does the number of audit coins.

Human Auditor FPR/FNR. In the default Trusted Monitoring game, Greenblatt et al. assume that no mistakes are made by the human auditor. We can extend the model to relax this assumption. Increasing the human FNR decreases safety almost linearly in the worst case, without impacting usefulness. False positives increase safety up to a point determined by the audit threshold, but decrease usefulness more rapidly the more the blue team audits.

7 RELATED WORK

Though the application of formal game-theoretic tools to Greenblatt et al. (2024)’s AI Control framework is novel, and the resulting type of game is unique, there are analogues to existing work.

The AI-Control Games studied are most similar to the work of Chang et al. (2015), who model a security problem for food manufacturing. Like AI-Control Games, they consider a leader-follower equilibrium, where the leader tries to maximise productivity and minimise risk, and the follower tries to hide an attack. To select the leader strategy against multiple objectives, they use a genetic algorithm (as do Lins et al. (2013)). Their model differs in its factored state space, infinite time horizon, and the lack of ‘attack/honest’ game modes and most importantly their consideration of a memoryful-follower. Their solution method is not directly applicable to AI-Control Games, and may be unsuitable for adaptation. For a detailed comparison, see Section E.

Given our reduction in Section 5, methods for finding equilibria in ZSPOSGs are particularly relevant. However, we find that our requirement for a memoryful-leader, memoryless-follower equilibrium precludes the use of many existing solvers. For example, Delage et al. (2023), which finds Nash equilibria in ZSPOSGs using a dynamic programming approach that solves subgames for each timestep from H down to 1. Other works attempt to solve subclasses of ZSPOSGs, with ‘public observations’ Horák & Božanský (2019) or ‘one-sided partial-observability’ Zheng et al. (2022); Horák et al. (2023). AI-Control Games do not in-general fit into these subclasses, and each of these papers considers two memoryful players.

Since in some cases it is possible to find an AI-Control Equilibrium by first finding a mixed Nash equilibrium (Theorem 5.5), it may be possible to augment the algorithm by Hansen & Feng (2000), which searches for two mixed strategies over a finite set of memoryful policies.

AI-Control Games, and especially the Trusted Monitoring case study we present, are closely related to Stackelberg Security Games: two-player games where the leader must choose a random strategy to cover a set of targets using a limited set of resources, and the follower tries to attack an uncovered target (see Sinha et al. (2018) for an overview). Unlike Trusted Monitoring, these are one-round games and much literature focuses on scalably solving large-scale problems against scheduling constraints. Although some security games model imperfect information Bier (2007); Nguyen et al. (2014) or multiple objectives Brown et al. (2012), most are single-objective and perfect-information.

Trusted Monitoring without defer is similar to the Inspection Games studied by Avenhaus et al. (2002, p. 1957): games where the defender receives a noisy observation, and has to decide whether to inspect. However, these games are one-round only. Recursive Inspection Games Dresher (1962) are multi-round, but have a memoryful follower.

8 FUTURE WORK AND CONCLUSIONS

In this work, we introduced a novel model for analysing the safety and usefulness of AI-Control protocols. We used our formal model and solution techniques to evaluate a much broader range of protocols and settings than prior empirical results from literature, and there is significant room to apply them further. In the showcased application of Trusted Monitoring setting, future work can utilise Algorithm 2 with POMDP-based best-response oracles to explore the space of memory-dependent protocols. Our model of the Trusted Monitoring protocol is also generalisable to *any* deployment setting, using a question-answer format and *any* unacceptable outcome. Future work will explore such generalised settings by studying the interaction of the untrusted AI and the trusted monitor to generate distributions like those in Figure 1 and analyse the resulting frontiers.

ETHICS STATEMENT

This work seeks to develop tools that can be used to study one operationalisation of ‘safety’ for advanced AI systems. The primary motivation of this research is social impact: to reduce the cost of doing high-quality evaluations of deployment protocols, and to make explicit the tradeoffs between safety and usefulness.

Despite this motivation, there are some potential negative consequences of this work. First, the work focuses on avoiding ‘catastrophic’ failures: minimising the risk of especially bad, but low-probability actions. This biases the work against considering other kinds of harms that are more likely. Second, ‘safety’ is defined relative to a particular operationalisation of a particular unacceptable outcome. In part, this is a strength of the framework, as it separates the technical question ‘How do we avoid the outcome?’ from the normative question ‘Which outcomes should we avoid?’. However, it also means that harms that are not well captured by such a clear-cut operationalisation may not fit within the framework. Third, the work focuses on avoiding harm from AI systems that the protocol developers did not intend. This means it is not well-suited to evaluating misuse or structural harms that could arise from the deployment of the system. Finally, the creation of a formal model of AI-control games could be used to justify the deployment of AI systems in situations where the risks are not well understood. If developers utilise poorly implemented AI-Control evaluations to justify deployment, this could lead to a false sense of security.

REPRODUCIBILITY STATEMENT

All of the code used to generate our results is provided in the supplementary material, using fixed seeds and experimental parameters. The empirical data used to model trusted monitor distributions is provided in Greenblatt et al. (2024). We have made every possible effort to ensure that our experiments are reproducible.

THE USE OF LARGE LANGUAGE MODELS

We used LLMs in several parts of our research process: including surfacing connections with other fields, brain-storming ways to clearly present our results, developing parts of our codebase, and formatting our final work. In all cases, language models were used only to assist the authors, who take full responsibility for the veracity of the paper.

REFERENCES

- Lukáš Adam, Rostislav Horčík, Tomáš Kasl, and Tomáš Kroupa. Double Oracle Algorithm for Computing Equilibria in Continuous Games. *Proceedings of the AAAI Conference on Artificial Intelligence*, 35(6):5070–5077, May 2021. doi: 10.1609/aaai.v35i6.16641. URL <https://ojs.aaai.org/index.php/AAAI/article/view/16641>. Section: AAAI Technical Track on Game Theory and Economic Paradigms.
- Adept AI. Adept: Useful General Intelligence, 2024. URL <https://www.adept.ai/>.
- Anthropic. Controlling powerful AI, March 2025. URL <https://www.youtube.com/watch?v=6Unxqr50Kqg>.
- Robert J Aumann, Sergiu Hart, and Motty Perry. The absent-minded driver. *Games and Economic Behavior*, 20(1):102–116, 1997. Publisher: Elsevier.
- Rudolf Avenhaus, Bernhard Von Stengel, and Shmuel Zamir. Chapter 51 Inspection games. In *Handbook of Game Theory with Economic Applications*, volume 3, pp. 1947–1987. Elsevier, January 2002. doi: 10.1016/S1574-0005(02)03014-X. URL <https://www.sciencedirect.com/science/article/pii/S157400050203014X>.
- Mikita Balesni, Marius Hobbhahn, David Lindner, Alexander Meinke, Tomek Korbak, Joshua Clymer, Buck Shlegeris, Jérémy Scheurer, Charlotte Stix, Rusheb Shah, Nicholas Goldowsky-Dill, Dan Braun, Bilal Chughtai, Owain Evans, Daniel Kokotajlo, and Lucius Bushnaq. Towards evaluations-based safety cases for AI scheming, November 2024. URL <http://arxiv.org/abs/2411.03336>. arXiv:2411.03336 [cs].
- Aryan Bhatt, Cody Rushing, Adam Kaufman, Tyler Tracy, Vasil Georgiev, Akbir Khan, David Matolcsi, and Buck Shlegeris. Ctrl-Z: Controlling AI Agents via Resampling (forthcoming). pp. Forthcoming as of, 2025.
- Vicki M. Bier. Choosing What to Protect. *Risk Analysis*, 27(3):607–620, 2007. ISSN 1539-6924. doi: 10.1111/j.1539-6924.2007.00906.x.
- Matthew Brown, Bo An, Christopher Kiekintveld, Fernando Ordóñez, and Milind Tambe. Multi-objective optimization for security games. In *Proceedings of the 11th International Conference on Autonomous Agents and Multiagent Systems, AAMAS ’12*, pp. 863–870, June 2012. ISBN 978-0-9817381-2-3.
- Marie Davidsen Buhl, Gaurav Sett, Leonie Koessler, Jonas Schuett, and Markus Anderljung. Safety cases for frontier AI, October 2024. URL <http://arxiv.org/abs/2410.21572>. arXiv:2410.21572 [cs].
- Stephen Casper, Lennart Schulze, Oam Patel, and Dylan Hadfield-Menell. Defending Against Unforeseen Failure Modes with Latent Adversarial Training, April 2024. URL <http://arxiv.org/abs/2403.05030>. arXiv:2403.05030 [cs].
- Yanling Chang, Alan L. Erera, and Chelsea C. White. A leader–follower partially observed, multiobjective Markov game. *Annals of Operations Research*, 235(1):103–128, December 2015. ISSN 1572-9338. doi: 10.1007/s10479-015-1935-0. URL <https://doi.org/10.1007/s10479-015-1935-0>.
- Joshua Clymer, Nick Gabrieli, David Krueger, and Thomas Larsen. Safety Cases: How to Justify the Safety of Advanced AI Systems, March 2024. URL <http://arxiv.org/abs/2403.10462>. arXiv:2403.10462 [cs].
- DeepMind. Frontier safety framework. Technical report, September 2025. URL https://storage.googleapis.com/deepmind-media/DeepMind.com/Blog/strengthening-our-frontier-safety-framework/frontier-safety-framework_3.pdf. tex.version: 3.0.
- Aurélien Delage, Olivier Buffet, Jilles S. Dibangoye, and Abdallah Saffidine. HSVI Can Solve Zero-Sum Partially Observable Stochastic Games. *Dynamic Games and Applications*, September 2023. ISSN 2153-0785, 2153-0793. doi: 10.1007/s13235-023-00519-6. URL <https://link.springer.com/10.1007/s13235-023-00519-6>.

- Melvin Dresher. A Sampling Inspection Problem in Arms Control Agreements: A Game-Theoretic Analysis. Technical report, RAND Corporation, January 1962. URL https://www.rand.org/pubs/research{_}memoranda/RM2972.html.
- R. Emery-Montemerlo, G. Gordon, J. Schneider, and S. Thrun. Approximate solutions for partially observable stochastic games with common payoffs. In *Proceedings of the Third International Joint Conference on Autonomous Agents and Multiagent Systems, 2004.*, pp. 136–143, July 2004. URL <https://ieeexplore.ieee.org/abstract/document/1373472>.
- Ryan Greenblatt, Buck Shlegeris, Kshitij Sachan, and Fabien Roger. AI Control: Improving Safety Despite Intentional Subversion. In *Proceedings of the 41st International Conference on Machine Learning*, pp. 16295–16336. PMLR, July 2024. URL <https://proceedings.mlr.press/v235/greenblatt24a.html>. ISSN: 2640-3498.
- Eric A. Hansen and Zhengzhu Feng. Dynamic programming for POMDPs using a factored state representation. In *Proceedings of the Fifth International Conference on Artificial Intelligence Planning Systems, AIPS’00*, pp. 130–139, Breckenridge, CO, USA, April 2000. AAAI Press. ISBN 978-1-57735-111-5.
- Dan Hendrycks, Steven Basart, Saurav Kadavath, Mantas Mazeika, Akul Arora, Ethan Guo, Collin Burns, Samir Puranik, Horace He, Dawn Song, and Jacob Steinhardt. Measuring Coding Challenge Competence With APPS. August 2021. URL <https://openreview.net/forum?id=sD93GOzH3i5>.
- Karel Horák and Branislav Bošanský. Solving Partially Observable Stochastic Games with Public Observations. *Proceedings of the AAAI Conference on Artificial Intelligence*, 33(01):2029–2036, July 2019. ISSN 2374-3468. doi: 10.1609/aaai.v33i01.33012029. URL <https://ojs.aaai.org/index.php/AAAI/article/view/4032>. Number: 01.
- Karel Horák, Branislav Bošanský, Vojtěch Kovařík, and Christopher Kiekintveld. Solving zero-sum one-sided partially observable stochastic games. *Artificial Intelligence*, 316:103838, March 2023. ISSN 0004-3702. doi: 10.1016/j.artint.2022.103838. URL <https://www.sciencedirect.com/science/article/pii/S0004370222001783>.
- Geoffrey Irving. Safety cases at AISI, August 2024. URL <https://www.aisi.gov.uk/work/safety-cases-at-aisi>.
- Leslie Pack Kaelbling, Michael L. Littman, and Anthony R. Cassandra. Planning and acting in partially observable stochastic domains. *Artificial Intelligence*, 101(1-2):99–134, 1998. ISSN 0004-3702. doi: 10.1016/s0004-3702(98)00023-x.
- Harold W. Kuhn. Extensive games and the problem of information. *Contributions to the Theory of Games*, 2(28):193–216, 1953.
- Isis Didier Lins, Leandro Chaves Rêgo, Márcio das Chagas Moura, and Enrique López Droguett. Selection of security system design via games of imperfect information and multi-objective genetic algorithm. *Reliability Engineering & System Safety*, 112:59–66, April 2013. ISSN 0951-8320. doi: 10.1016/j.ress.2012.11.021. URL <https://www.sciencedirect.com/science/article/pii/S0951832012002517>.
- Monte MacDiarmid, Timothy Maxwell, Nicholas Schiefer, Jesse Mu, Jared Kaplan, David Duvenaud, Sam Bowman, Alex Tamkin, Ethan Perez, Mrinank Sharma, Carson Denison, and Evan Hubinger. Simple probes can catch sleeper agents, April 2024. URL <https://www.anthropic.com/news/probes-catch-sleeper-agents>.
- Alex Mallen, Charlie Griffin, Alessandro Abate, and Buck Shlegeris. Subversion Strategy Eval: Evaluating AI’s stateless strategic capabilities against control protocols, 2024. URL <https://arxiv.org/abs/2412.12480>.
- Michael Maschler, Eilon Solan, and Shmuel Zamir. Behavior strategies and kuhn’s theorem. In *Game theory*, pp. 219–250. Cambridge University Press, Cambridge, 2013.

- H. Brendan McMahan, Geoffrey J. Gordon, and Avrim Blum. Planning in the presence of cost functions controlled by an adversary. In *Proceedings of the 20th International Conference on Machine Learning (ICML-03)*, pp. 536–543, 2003. URL <https://cdn.aaai.org/ICML/2003/ICML03-071.pdf>.
- Kaisa Miettinen. *Nonlinear Multiobjective Optimization*. Springer, New York, NY, UNITED STATES, 1998. ISBN 978-1-4615-5563-6. URL <http://ebookcentral.proquest.com/lib/oxford/detail.action?docID=3081836>.
- Thanh H. Nguyen, Amulya Yadav, Bo An, Milind Tambe, and Craig Boutilier. Regret-based optimization and preference elicitation for stackelberg security games with uncertainty: 28th AAAI Conference on Artificial Intelligence, AAAI 2014, 26th Innovative Applications of Artificial Intelligence Conference, IAAI 2014 and the 5th Symposium on Educational Advances in Artificial Intelligence, EAAI 2014. *Proceedings of the 28th AAAI Conference on Artificial Intelligence and the 26th Innovative Applications of Artificial Intelligence Conference and the 5th Symposium on Educational Advances in Artificial Intelligence*, pp. 756–762, 2014. URL <http://www.scopus.com/inward/record.url?scp=84908202610&partnerID=8YFLogxK>. Publisher: AI Access Foundation.
- OpenAI. ChatGPT plugins, March 2023. URL <https://openai.com/blog/chatgpt-plugins>.
- OpenAI and Apollo Research. Detecting and reducing scheming in AI models, September 2025. URL <https://openai.com/index/detecting-and-reducing-scheming-in-ai-models/>.
- Martin J. Osborne and Ariel Rubinstein. *A course in game theory*. MIT Press, Cambridge, Mass., 12. print edition, 2006. ISBN 978-0-262-65040-3.
- Rusheb Shah, Quentin Feuillade Montixi, Soroush Pour, Arush Tagade, and Javier Rando. Scalable and Transferable Black-Box Jailbreaks for Language Models via Persona Modulation. November 2023. URL <https://openreview.net/forum?id=x3Ltqz1UFg>.
- Guy Shani, Joelle Pineau, and Robert Kaplow. A survey of point-based POMDP solvers. *Autonomous Agents and Multi-Agent Systems*, 27(1):1–51, July 2013. ISSN 1573-7454. doi: 10.1007/s10458-012-9200-2. URL <https://doi.org/10.1007/s10458-012-9200-2>.
- Arunesh Sinha, Fei Fang, Bo An, Christopher Kiekintveld, and Milind Tambe. Stackelberg Security Games: Looking Beyond a Decade of Success. In *Proceedings of the Twenty-Seventh International Joint Conference on Artificial Intelligence*, pp. 5494–5501, Stockholm, Sweden, July 2018. International Joint Conferences on Artificial Intelligence Organization. ISBN 978-0-9992411-2-7. doi: 10.24963/ijcai.2018/775. URL <https://www.ijcai.org/proceedings/2018/775>.
- The White House. America’s AI action plan. Technical report, July 2025. URL <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.
- UK AI Safety Institute. Research agenda, 2025. URL <https://www.aisi.gov.uk/research-agenda>.
- Jiaxin Wen, Vivek Hebbar, Caleb Larson, Aryan Bhatt, Ansh Radhakrishnan, Mrinank Sharma, Henry Sleight, Shi Feng, He He, Ethan Perez, Buck Shlegeris, and Akbir Khan. Adaptive Deployment of Untrusted LLMs Reduces Distributed Threats, November 2024. URL <http://arxiv.org/abs/2411.17693>. arXiv:2411.17693.
- Wei Zheng, Taeho Jung, and Hai Lin. The Stackelberg equilibrium for one-sided zero-sum partially observable stochastic games. *Automatica*, 140:110231, June 2022. ISSN 0005-1098. doi: 10.1016/j.automatica.2022.110231. URL <https://www.sciencedirect.com/science/article/pii/S0005109822000760>.

A GENERAL RESULTS AND PROOFS

This appendix presents proofs of the results from Section 3 and Section 5 that apply to all AI-Control Games. To do so, we first introduce a useful lemma which maps the optimisation problem into a vector space.

A.1 MAPPING TO VECTOR SPACE

We will define a mapping from Π^B to a simplex X , and another from Π^R to a simplex Y . Then, we will show that these mappings can be used to calculate safety and usefulness.

We begin by defining the simplex Y and a mapping $\mu^R : \Pi^R \rightarrow Y$. Recall that: $\Pi^R := \{\pi^R \mid \pi^R : \Omega^R \rightarrow \mathcal{A}^R\}$. Define $n := |\mathcal{A}^R| \cdot |\Omega^R|$ and enumerate the n observation-action pairs such that each dimension in $\mathbb{R}^n$ corresponds to one pair. Consider the simplex $Y \subset \mathbb{R}^n$ such that all values are ≥ 0 and, for each observation, the sum across all actions is 1. Then, define the mapping such that $\mu^R(\pi^R) = \mathbf{y}$ if, and only if, for all j , $\pi^R(a^R \mid o^R) = y_j$ for j 'th observation-action pair (o^R, a^R) . Note that μ^R is bijective.

Blue's mapping is more complex. Define the finite set of Blue observation-action histories as $\mathcal{H}^B := \Omega^B \times \bigcup_{i=0}^{H-1} (\mathcal{A}^B \times \Omega^B)^i$. Since the sets $\mathcal{H}^B$ and $\mathcal{A}^B$ are finite, we *could* use a similar strategy to the one we used for Y to define finite-dimensional vector space X that represents $\Pi^B := \{\pi^B \mid \pi^B : \mathcal{H}^B \rightarrow \mathcal{A}^B\}$. However, an alternative approach gives a useful linearity.

Because Blue has perfect recall, any behaviourally stochastic policy has an equivalent representation as a mixed strategy over the finite set of deterministic policies. This is the reverse direction of Kuhn's Theorem Kuhn (1953), sometimes stated informally with 'and vice versa'. This exact result is mentioned in Maschler et al. (2013, Theorem 6.26 on p. 242) as part of a more general proof about infinite-horizon games, and an informal proof is given by Osborne & Rubinstein (2006, p. 214). The other direction (which will also be useful), is given explicitly in Maschler et al. (2013, Theorem 6.15 on p. 232):

Theorem (Kuhn (1953)). *In every game in extensive form, if player i has perfect recall, then for every mixed strategy of player i there exists an equivalent behavior strategy.*

Note that the equivalence stated here is in terms of the probability of reaching any vertex in an extensive-form tree which is more than sufficient for our purposes (see Definition 6.5 on p. 223, Maschler et al., 2013).

We will use the fact that an AI-Control game is a finite extensive form game, and the 'vice versa' of Kuhn's Theorem to give a mapping from behavioural to mixed policies.

Let Π_{det}^B be the set of deterministic policies for Blue. Then, for each behaviourally stochastic π^B , there is an equivalent mixture over deterministic policies: $\sigma^B \in \Delta(\Pi_{\text{det}}^B)$.

Let $m := |\Pi_{\text{det}}^B| = |\mathcal{A}^B|^{|\mathcal{H}^B|}$ and consider the simplex $X \subset \mathbb{R}^m$:

$$X := \left\{ \mathbf{x} \mid \mathbf{x} \in \mathbb{R}^m, \mathbf{x} \geq 0, \sum_{i=1}^m x_i = 1 \right\}$$

Enumerate the m deterministic policies for Blue and assign each policy to a dimension in $\mathbb{R}^m$. Let π_i^B be the i th deterministic policy for Blue.

Define the mapping $\mu^B(\pi^B) := \mathbf{x}$ where $x_i := \sigma^B(\pi_i^B)$, ie. the probability that π_i^B is chosen from the mixture σ^B that is equivalent to π^B . Further, note that by the main direction of Kuhn's theorem, for any mixed stochastic policy (σ^B), there is also a behaviourally stochastic policy (π^B) that is equivalent but it's not immediate that μ^B is bijective, since there may be two mixed stochastic policies that correspond to the same behaviourally stochastic policy.

From here on, we drop the explicit use of μ^B and μ^R , and instead use $\mathbf{x}_{\pi^B} := \mu^B(\pi^B)$ and $\mathbf{x}_{\pi^R} := \mu^R(\pi^R)$.

Lemma A.1. For any $\mathcal{G}$, and any full-length trajectory h , there is a function $p_h : X \times Y \rightarrow [0, 1]$ that is linear in its first argument, continuous in its second, and satisfies the following for all π^B, π^R :

$$\begin{aligned} p_h(\mathbf{x}_{\pi^B}, \mathbf{y}_{\pi^R}) &= \mathbb{P}[h \mid \textcircled{X}, \pi^B, \pi^R] \\ p_h(\mathbf{x}_{\pi^B}, \mathbf{y}_{\pi_{\text{null}}^R}) &= \mathbb{P}[h \mid \textcircled{H}, \pi^B, \pi^R] \end{aligned}$$

where π_{null}^R is the policy that always chooses a_{null} .

Proof. Let $\mathcal{H}$ be the set of all possible full-length trajectories (or joint-state-observation-action histories):

$$h := (\textcircled{?}, s_0, o_0^B, o_0^R, a_0^B, a_0^R, s_1, \dots, o_{H-1}^R, a_{H-1}^B, a_{H-1}^R, s_H)$$

We will define p_h such that, for all π^B, π^R

$$\mathbb{P}[h \mid \textcircled{X}, \pi^B, \pi^R] = p_h(\mathbf{x}_{\pi^B}, \mathbf{y}_{\pi^R})$$

Then, we can write the probability of each history explicitly, in terms of π^B and π^R :

$$\begin{aligned} &\mathbb{P}[h \mid \textcircled{X}, \pi^B, \pi^R] \\ &= \mathcal{I}(s_0, o_0^B, o_0^R) \cdot \pi^B(a_0^B \mid o_0^B) \cdot \pi^R(a_0^R \mid o_0^R) \\ &\quad \cdot \mathcal{P}(s_1, o_1^B, o_1^R \mid s_0, a_0^B, a_0^R) \\ &\quad \dots \\ &\quad \cdot \mathcal{P}(s_H, \cdot, \cdot \mid s_{H-1}, a_{H-1}^B, a_{H-1}^R) \\ &= c_h \cdot \prod_{t=0}^{H-1} \pi^B(a_t^B \mid o_{0:t}^B, a_{0:t-1}^B) \pi^R(a_t^R \mid o_t^R) \end{aligned}$$

Here $o_{0:t}^B$ denotes the tuple $(o_0^B, \dots, o_t^B)$, and similar for $a_{0:t-1}^B$ etc. The value c_h is a constant with respect to $\mathbf{x}$ and $\mathbf{y}$ and depends only on the trajectory and the AI-Control Game $\mathcal{G}$.

Using mixed-behavioural equivalence, let σ^B be the mixture of deterministic strategies equivalent to π^B , then we can split the Blue probabilities:

$$\begin{aligned} &\mathbb{P}[h \mid \textcircled{X}, \pi^B, \pi^R] \\ &= c_h \sum_{i=1}^m \sigma^B(\pi_i^B) \cdot \prod_{t=0}^{H-1} \pi_i^B(a_t^B \mid o_{0:t}^B, a_{0:t-1}^B) \pi^R(a_t^R \mid o_t^R) \end{aligned}$$

And finally, let $\mathbf{x} := \mathbf{x}_{\pi^B}$, note that $x_i = \sigma^B(\pi_i^B)$, write out the expression in terms of x_i , and y_j , and call this expression p_h :

$$p_h(\mathbf{x}, \mathbf{y}) := c_h \cdot \left(\sum_{i=1}^m x_i f_i(h) \right) \cdot \left(\prod_{t=0}^{H-1} y_j(o_t^R, a_t^R) \right) \quad (2)$$

Here $f_i(h)$ is either 0, if π_i^B is not consistent with h , or 1 if it is. Clearly, the probability $(\mathbb{P}[h \mid \dots])$ of getting some trajectory h in an attack game, given policies π^B and π^R is simply $p_h(\mathbf{x}_{\pi^B}, \mathbf{y}_{\pi^R})$.

The probability of getting some trajectory h in an honest game is the same as the probability in the attack game with the Red policy π_{null}^R that always chooses the a_{null} action. (Although the rewards in an episode will be different in $\textcircled{H}$ versus $\textcircled{X}$ modes, these aren't included in h .) This can be written as

$p_h(\mathbf{x}, \mathbf{y}_{\text{null}})$, where $\mathbf{y}_{\text{null}}$ corresponds to the π_{null}^R .

$$\begin{aligned} \mathbb{P}[h \mid \Theta, \pi^B] &= \mathbb{P}[h \mid \otimes, \pi^B, \pi_{\text{null}}^R] \\ &= c_h \sum_{i=1}^m x_i \cdot f_i(h) \prod_{t=0}^{H-1} (\mathbf{y}_{\text{null}})_{j(o_t^R, a_t^R)} \\ &= c_h^{\text{null}} \sum_{i=1}^m x_i \cdot f_i(h) \end{aligned}$$

Here, we rewrite the constant c_h^{null} to be $c(h)$, if all the Red actions in h are a_{null} , and 0 otherwise.

We have shown that the probability of a given trajectory can be written in terms of $p_h(\mathbf{x}_{\pi^B}, \mathbf{y}_{\pi^R})$, for attack games, and $p_h(\mathbf{x}_{\pi^B}, \mathbf{y}_{\text{null}})$ for honest games. It remains to note that $p_h(\mathbf{x}, \mathbf{y})$ is straightforwardly linear in $\mathbf{x}$ and continuous in $\mathbf{y}$ by Equation (2). $\square$

Using Lemma A.1, we can rewrite safety and usefulness as explicit expressions with some useful properties:

Corollary A.2. *For any AI-Control Game $\mathcal{G}$, with simplex X and simplex Y defined in Lemma A.1, there is a function $u : X \rightarrow \mathbb{R}$ such that $u(\mathbf{x}_{\pi^B}) = \text{Use}(\pi^B)$ and a function $s : X \times Y \rightarrow \mathbb{R}$ such that $s(\mathbf{x}_{\pi^B}, \mathbf{y}_{\pi^R}) = \text{Safety}(\pi^B, \pi^R)$. Further*

1. *The function u is linear, and therefore uniformly continuous, in $\mathbf{x}$.*
2. *The function s is linear in $\mathbf{x}$ for fixed $\mathbf{y}$, and s is continuous in $\mathbf{y}$ for fixed $\mathbf{x}$.*
3. *The minimiser to $\min_{\mathbf{y}} s(\mathbf{x}, \mathbf{y})$ always exists. Further, the function $g(\mathbf{x}) := \min_{\mathbf{y}} s(\mathbf{x}, \mathbf{y})$ is continuous everywhere and concave.*
4. *For all $w \in [0, 1]$, the function $v_w(\mathbf{x}) := wg(\mathbf{x}) + (1 - w)u(\mathbf{x})$ is uniformly continuous, and therefore has maximal values.*

Proof. First, we can define the functions s and u : since the set of trajectories is finite, we can rewrite the safety and usefulness functions by summing over the trajectories. The safety equation can then be rewritten as:

$$\begin{aligned} \text{Safety}(\pi^B, \pi^R) &:= \mathbb{E} \left[\sum_{t=0}^{H-1} \mathcal{R}_{\text{use}}(S_t, A_t^B, A_t^R, S_{t+1}) \mid \Theta, \pi^B, \pi^R \right] \\ &= \sum_{h \in \mathcal{H}} r_{\text{safe}}(h) \mathbb{P}[h \mid \otimes, \pi^B, \pi^R] \\ &= \sum_{h \in \mathcal{H}} r_{\text{safe}}(h) p_h(\mathbf{x}_{\pi^B}, \mathbf{y}_{\pi^R}) \\ &= \sum_{h \in \mathcal{H}} r_{\text{safe}}(h) c_h \sum_{i=1}^m x_i \cdot f_i(h) \prod_{t=0}^{H-1} y_{j(o_t^R, a_t^R)} \\ &=: s(\mathbf{x}, \mathbf{y}) \end{aligned}$$

Here $r_{\text{safe}}(h)$ is the sum of safety rewards for a given trajectory, and is fixed for any given AI-Control Game. The last line gives the definition for our new function s . We can similarly derive and define an

expression for u :

$$\begin{aligned}
 \text{Use}(\pi^B) &:= \mathbb{E} \left[\sum_{t=0}^H \mathcal{R}_{\text{Use}}(S_t, A_t^B, S_{t+1}) \mid \Theta, \pi^B \right] \\
 &= \sum_{h \in \mathcal{H}} r_{\text{use}}(h) \mathbb{P}[h \mid \Theta, \pi^B] \\
 &= \sum_{h \in \mathcal{H}} r_{\text{use}}(h) c_h^{\text{null}} \sum_{i=1}^m x_i \cdot f_i(h) \\
 &=: u(\mathbf{x})
 \end{aligned}$$

We prove each of the remaining statements one-by-one.

1. That $u(\cdot)$ is linear over $\mathbf{x}$ follows from the definition of p_h , and of linearity of s in p_h .
2. That, for fixed $\mathbf{x}$, $s(\mathbf{x}, \cdot)$ is continuous over Y is immediate from the continuity of each p_h in each y_j , and of linearity of s in p_h . That, for fixed $\mathbf{y}$, $s(\cdot, \mathbf{y})$ is linear over $\mathbf{x}$ follows straightforwardly from the linearity of each p_h in $\mathbf{x}$, and of linearity of s in p_h .
3. The minimiser $\mathbf{y}_*$ of $\min_{\mathbf{y}} s(\mathbf{x}, \mathbf{y})$ must exist because $s(\mathbf{x}, \cdot)$ is continuous in Y for fixed $\mathbf{x}$, and Y is compact, by the extreme value theorem. Showing that $\min_{\mathbf{y}} s(\cdot, \mathbf{y})$ is continuous everywhere $\mathbf{x}$ is a little more involved. Note that s can be rewritten as $s(\mathbf{x}, \mathbf{y}) = \mathbf{x} \cdot \mathbf{v}(\mathbf{y})$, where $\mathbf{v}(\mathbf{y}) := [f_1(\mathbf{y}), \dots, f_m(\mathbf{y})]^T$, and each $f_i(\mathbf{y})$ is continuous in $\mathbf{y}$. By continuity and compactness, $d := \max_{\mathbf{y}} \|\mathbf{v}(\mathbf{y})\|$ exists. For all $\epsilon > 0$, define $\delta := \frac{\epsilon}{2d}$. Now, for any $\mathbf{y}$, if $\|\mathbf{x} - \mathbf{x}_0\| < \delta$ then

$$\begin{aligned}
 |s(\mathbf{x}, \mathbf{y}) - s(\mathbf{x}_0, \mathbf{y})| &= |(\mathbf{x} - \mathbf{x}_0) \cdot \mathbf{v}(\mathbf{y})| \leq \\
 &\dots \|\mathbf{x} - \mathbf{x}_0\| \|\mathbf{v}(\mathbf{y})\| \leq \delta \cdot \frac{\epsilon}{2d} < \epsilon
 \end{aligned}$$

To show continuity of $\min_{\mathbf{y}} s(\cdot, \mathbf{y})$, consider any ϵ and again set $\delta := \frac{\epsilon}{2d}$. Consider any $\mathbf{x}_0$ and $\mathbf{x}_1$ minimised by $\mathbf{y}_0$ and $\mathbf{y}_1$ such that $\|\mathbf{x}_0 - \mathbf{x}_1\| \leq \delta$, and without loss of generality assume $s(\mathbf{x}_1, \mathbf{y}_1) \geq s(\mathbf{x}_0, \mathbf{y}_0)$, then, since $\min_{\mathbf{y}} s(\mathbf{x}_1, \mathbf{y}) \leq s(\mathbf{x}_1, \mathbf{y}_0)$:

$$\min_{\mathbf{y}} s(\mathbf{x}_1, \mathbf{y}) - \min_{\mathbf{y}} s(\mathbf{x}_0, \mathbf{y}) \leq s(\mathbf{x}_1, \mathbf{y}_0) - s(\mathbf{x}_0, \mathbf{y}_0) < \epsilon$$

This demonstrates that $\min_{\mathbf{y}} s(\cdot, \mathbf{y})$ is continuous over X . Finally, the concavity claim, that $g(\lambda \mathbf{x}_1 + (1 - \lambda) \mathbf{x}_0) \geq \lambda g(\mathbf{x}_1) + (1 - \lambda) g(\mathbf{x}_0)$, follows straightforwardly from the definition of g , the linearity of $s(\cdot, \mathbf{y})$, and the fact that $\min_{\mathbf{y}} (\lambda s(\mathbf{x}_1, \mathbf{y}) + (1 - \lambda) s(\mathbf{x}_0, \mathbf{y})) \geq \min_{\mathbf{y}_1} \lambda s(\mathbf{x}_1, \mathbf{y}_1) + \min_{\mathbf{y}_0} (1 - \lambda) s(\mathbf{x}_0, \mathbf{y}_0)$.

4. Straightforwardly from the linearity of $u(\mathbf{x})$ and the uniform continuity of $\min_{\mathbf{y}} s(\cdot, \mathbf{y})$, we have that $v_w(\mathbf{x}) := wg(\mathbf{x}) + (1 - w)u(\mathbf{x})$ is continuous over X . That maximal values exist follows from this continuity, the compactness of X , and the extreme value theorem.

□

Lemma A.3. Let $\phi(\mathbf{x}) := [g(\mathbf{x}), u(\mathbf{x})]^T$. The vector $\mathbf{x}$ is weakly Pareto-optimal (no $\mathbf{x}'$ has strictly better g and u) if, and only if, $\exists w \in [0, 1]$ such that $\mathbf{x} \in \arg \max_{\mathbf{x}'} v_w(\mathbf{x}')$.

Proof. The ‘if’ statement is straightforward: suppose some $\mathbf{x}'$ is such that $x'_1 > x_1^*$ and $x'_2 > x_2^*$. Clearly, for all w , $v_w(\mathbf{x}') > v_w(\mathbf{x}^*)$ and therefore no w can exist for which $\mathbf{x}^*$ is (jointly) optimal.

For the ‘only if’ statement, we build up Miettinen (1998, Theorem 3.1.4 on p. 79), which says that every Pareto-optimal solution to a concave multi-objective optimisation problem is optimal for some weighting.¹⁶ We have already shown g is concave, and since u is linear it is trivially concave.

¹⁶The original statement refers to convex minimisation, rather than concave maximisation.

Therefore, the theorem holds and any Pareto-optimal $\mathbf{x}^*$ is also optimal for some v_w . It remains to show that in our case, if not in general, any weak Pareto-optimum that is also optimal for some w . The proof relies on the convexity of X , the concavity of g , and the linearity of u . To do this, we show that any weak Pareto-optimum that is not a Pareto-optimum must have either maximal safety or maximal usefulness.

Suppose $\mathbf{x}^*$ is a weak Pareto-optimum but not a Pareto-optimum, and for reduction assume that $\mathbf{x}^*$ has neither optimal safety nor optimal usefulness.

If there is both (1) a point $\mathbf{x}_0$ with equal g and better u and (2) a point $\mathbf{x}_1$ with equal u but better g , then we can derive a point $\frac{\mathbf{x}_0 + \mathbf{x}_1}{2}$ that dominates $\mathbf{x}^*$.¹⁷

The definition of weak Pareto-optimality only implies that (1) or (2) holds. However, it can be shown that if (1) holds and $\mathbf{x}^*$ does not maximise g , then (2) also holds or $\mathbf{x}^*$ is otherwise dominated.¹⁸ Similarly, it can be shown that if (2) holds and $\mathbf{x}^*$ does not maximise u , then (1) also holds or $\mathbf{x}^*$ is otherwise dominated.¹⁹

Since weak Pareto-optimality implies (1) or (2), which implies (1) and (2), which derives a contradiction unless $\mathbf{x}^*$ maximises g or u , we can conclude that any weakly Pareto-optimal point that is not also Pareto-optimal must maximise either g or u . Therefore, such a point is optimal for a weighting of either $w = 1$ or $w = 0$ gives us that $\mathbf{x}^* \in \arg \max_{\mathbf{x}} v_w(\mathbf{x})$.

□

A.2 PROVING RESULTS FROM SECTION 3

Using the lemma above, we can immediately prove some of the results from Section 3.

Proposition (Proposition 3.2 Restated). *For any AI-Control Game $\mathcal{G}$, and appropriate distance functions in Π^B and Π^R : $\text{Safety}(\pi^B, \cdot)$, $\text{Safety}(\cdot, \pi^R)$, and $\text{Use}(\cdot)$ are continuous in their domains.*

Proof. These results follow fairly straightforwardly from Corollary A.2, using the Blue distance measure $d(\pi_1^B, \pi_2^B) = \|\mathbf{x}_{\pi_1^B} - \mathbf{x}_{\pi_2^B}\|$ and Red distance measure $d(\pi_1^R, \pi_2^R) = \|\mathbf{x}_{\pi_1^R} - \mathbf{x}_{\pi_2^R}\|$. For fixed π^R and ϵ , $\exists \delta(\epsilon)$ such that if $d(\pi_1^B, \pi_2^B) = \|\mathbf{x}_{\pi_1^B} - \mathbf{x}_{\pi_2^B}\| < \delta$, then by Corollary A.2, $|s(\mathbf{x}_{\pi_1^B}, \mathbf{y}) - s(\mathbf{x}_{\pi_2^B}, \mathbf{y})| = |\text{Safety}(\pi_1^B, \pi^R) - \text{Safety}(\pi_2^B, \pi^R)| < \epsilon$, and therefore Safety is continuous. The proofs for $\text{Safety}(\cdot, \pi_0^R)$ and $\text{Use}(\cdot)$ are similar. □

Proposition (Proposition 3.3 Restated). *For all π^B , $\text{Safety}(\pi^B, \cdot)$ attains its minimum at some π^R , and $\min_{\pi^R} \text{Safety}(\cdot, \pi^R)$ is continuous.*

Proof. For any π^B , the minimiser $\mathbf{y}$ to $s(\mathbf{x}_{\pi^B}, \cdot)$ always exists by Corollary A.2. By the fact that μ^R is a bijection, there is some π^R such that $\mu^R(\pi^R) = \mathbf{y}$ and this π^R must minimise $\text{Safety}(\pi^B, \cdot)$ since $s(\mathbf{x}_{\pi_1^B}, \mathbf{y}_{\pi_1^R}) = \text{Safety}(\pi_1^B, \pi_1^R)$ for all π_1^B and π_1^R . That $\min_{\pi^R} \text{Safety}(\cdot, \pi^R)$ is continuous at all $\pi^B \in \Pi^B$ follows directly from the claim that $\min_{\mathbf{y}} s(\cdot, \mathbf{y})$ is continuous. □

Proposition (Proposition 3.5 restated). *Every $\mathcal{G}$ has an AI-Control Equilibrium.*

¹⁷By convexity of X , the point $\frac{\mathbf{x}_0 + \mathbf{x}_1}{2}$ is in X and by linearity of u and concavity of g , dominates $\mathbf{x}^*$: $u(\frac{\mathbf{x}_0 + \mathbf{x}_1}{2}) = \frac{u(\mathbf{x}_0) + u(\mathbf{x}_1)}{2} > u(\mathbf{x}^*)$ and $g(\frac{\mathbf{x}_0 + \mathbf{x}_1}{2}) \geq \frac{g(\mathbf{x}_0) + g(\mathbf{x}_1)}{2} > g(\mathbf{x}^*)$.

¹⁸If $\mathbf{x}^*$ does not maximise u , let $\mathbf{x}_0$ be any point with higher u and by assumption worse-or-equal g . If $g(\mathbf{x}_0) = g(\mathbf{x}^*)$, then (2) also holds, otherwise, $g(\mathbf{x}_0) < g(\mathbf{x}^*)$. Choose $\lambda := \frac{g(\mathbf{x}^*) - g(\mathbf{x}_0)}{g(\mathbf{x}_1) - g(\mathbf{x}_0)}$ such that $g(\mathbf{x}^*) = \lambda g(\mathbf{x}_1) + (1 - \lambda)g(\mathbf{x}_0)$. Clearly $0 < \lambda < 1$. Let $\mathbf{x}_\lambda = \lambda \mathbf{x}_1 + (1 - \lambda)\mathbf{x}_0$. We can see that $u(\mathbf{x}_\lambda) = \lambda u(\mathbf{x}_1) + (1 - \lambda)u(\mathbf{x}_0) = \lambda u(\mathbf{x}^*) + (1 - \lambda)u(\mathbf{x}_0) > u(\mathbf{x}^*)$. Further, by concavity of g , $g(\mathbf{x}_\lambda) \geq \lambda g(\mathbf{x}_1) + (1 - \lambda)g(\mathbf{x}_0) = g(\mathbf{x}^*)$. If this inequality is strict, $\mathbf{x}_\lambda$ dominates $\mathbf{x}^*$, if it is not, (2) holds.

¹⁹If $\mathbf{x}^*$ does not maximise g , let $\mathbf{x}_1$ be any point with higher g and by assumption worse-or-equal u . If $u(\mathbf{x}_1) = u(\mathbf{x}^*)$, then (1) also holds, otherwise $u(\mathbf{x}_1) < u(\mathbf{x}^*)$. By the linearity of u , and convexity of X , there is some $0 < \lambda < 1$ for which $u(\mathbf{x}_\lambda) = u(\mathbf{x}^*)$. By concavity of g , $g(\mathbf{x}_\lambda) \geq \lambda g(\mathbf{x}_1) + (1 - \lambda)g(\mathbf{x}_0) = \lambda g(\mathbf{x}_1) + (1 - \lambda)g(\mathbf{x}^*) \geq g(\mathbf{x}^*)$. If this inequality is strict, $\mathbf{x}_\lambda$ dominates $\mathbf{x}^*$. Otherwise, if it is equal, (1) also holds.

Proof. Given that a Red safety minimiser always exists, to show that an AI-Control Equilibrium always exists, it's sufficient to show that there is some π^B_* for which no other π^B has strictly better $\min_{\pi^R} \text{Safety}(\cdot, \pi^R)$ and $\text{Use}(\cdot)$. By Corollary A.2 with $w = 1$, there exists a point x where v_1 attains its maximum. Let σ^B be the mixed strategy represented by x , then by Kuhn's Theorem, there is some behavioural π^B equivalent to σ^B . (Note that we cannot guarantee that $\mu^B(\pi^B) = x$, only that π^B is equivalent, and therefore induces the trajectories, safety and usefulness). For all π^R , $\text{Safety}(\pi^B, \pi^R) = s(x, y_{\pi^R})$. Any Blue policy π^B_1 that had strictly better worst-case safety than π^B would give some $x_{\pi^B_1}$ that dominates x (a contradiction). Therefore, no policy has strictly better safety than the policy π^B and $(\pi^B, \arg \min_{\pi^R} \text{Safety}(\pi^B, \pi^R))$ is an AI-Control Equilibrium. $\square$

A.3 RESULTS FROM SECTION 5

The mapping defined in Lemma A.1 is also useful for proving the results in Section 5, that are used to support Algorithm 1 and Algorithm 2.

A.3.1 SINGLE OBJECTIVE PROBLEM

First, we show the correspondence between solutions to the multi-objective and single objective problems.

Lemma (Repeat of Lemma 5.1). *A Blue policy π^B is in an AI-Control Equilibrium if, and only if, for some $w \in [0, 1]$, it maximises $\min_{\pi^R_*} (w \text{Safety}(\cdot, \pi^R) + (1 - w) \text{Use}(\cdot))$.*

Proof. The claim that a Blue policy π^B is in an AI-Control Equilibrium if, and only if, it is in a maximin point for V_w^G can be shown using the equivalent claim from Theorem A.3: that x is weakly Pareto-optimal if, and only if, it is optimal for some v_w . A Blue policy π^B is in an AI-Control Equilibrium if, and only if, $\phi(x_{\pi^B})$ is weakly Pareto-optimal, if and only if, there exists some w for which x_{π^B} is optimal for v_w .

For what remains, it is sufficient to show that, for any w , x_{π^B} is optimal for v_w if, and only if, π^B maximises $\min_{\pi^R_*} (w \text{Safety}(\cdot, \pi^R) + (1 - w) \text{Use}(\cdot))$.

($\Rightarrow$) If some π^B_1 beat π^B in $\min_{\pi^R_*} (w \text{Safety}(\cdot, \pi^R) + (1 - w) \text{Use}(\cdot))$, then $x_{\pi^B_1}$ would beat x in v_w .

($\Leftarrow$) If some x_2 beat x_{π^B} in v_w , then by Kuhn's Theorem, there would be some π^B_3 , equivalent to the mixed strategy σ^B_2 represented by x_2 ²⁰, such that π^B_3 beats π^B for w . $\square$

A.3.2 REDUCTION TO A ZERO-SUM GAME

Although Lemma 5.1 creates a single-objective optimisation problem for Blue, the resulting optimisation problem is, on the face of it, a novel type of game. The next result reduces the optimisation problem further to finding an unusual form of equilibrium in an existing type of game.

Lemma A.4 (Lemma 5.2). *For every AI-Control Game G , and weight $w \in [0, 1]$, there is a single-objective, zero-sum, partially observable stochastic game G_w with expected return (for Blue) V_{G_w} such that:*

$$\forall \pi^B, \pi^R : V_{G_w}(\pi^B, \pi^R) = V_w^G(\pi^B, \pi^R)$$

Proof. The notation and exact details of partially observable stochastic games vary between sources. For our purposes, we use a tuple that's similar to the one used in Emery-Montemerlo et al. (2004), with some modifications to make the game zero-sum rather than cooperative, and specifically two-player. We also include observations in the initial distribution, to increase generality such that agents might receive and observation before choosing their first action.

Definition (ZSPOSG). A zero-sum two-player partially observable stochastic game is a tuple $G = (\mathcal{S}, \mathcal{A}^B, \mathcal{A}^R, \Omega^B, \Omega^R, \mathcal{I}, \mathcal{P}, \mathcal{R}, H)$.

²⁰though not necessarily such that $\mu^B(\pi^B_3) = x_2$

There is no inbuilt notion of honest or attack modes in this definition. The reward function, with type $\mathcal{R} : \mathcal{S} \times \mathcal{A}^B \times \mathcal{A}^R \rightarrow \mathbb{R}$, is also new. Otherwise, the components of the tuple are the same as in the AI-Control Game.

Given some AI-Control Game $\mathcal{G} = (\mathcal{S}, \mathcal{A}^B, \mathcal{A}^R, \Omega^B, \Omega^R, \mathcal{I}, \mathcal{P}, \mathcal{R}_{\text{Safe}}, \mathcal{R}_{\text{Use}}, H)$, and weight w , we can construct the ZSPOSG $G_w := (\tilde{\mathcal{S}}, \mathcal{A}^B, \mathcal{A}^R, \Omega^B, \Omega^R, \tilde{\mathcal{I}}, \tilde{\mathcal{P}}, \mathcal{R}_w, H)$ as follows:

- ($\tilde{\mathcal{S}}$) The ‘mode’ of the game is baked into the state space: $\tilde{\mathcal{S}} := \mathcal{S} \times \{\textcircled{H}, \textcircled{X}\}$.
- ($\mathcal{A}^B, \mathcal{A}^R, \Omega^B, \Omega^R$) The action and observation spaces are unchanged.
- ($\mathcal{I}$) In an AI-Control Game, the mode of the game is chosen non-deterministically. Instead, we choose it with some non-zero probability, say 0.5. Then $\tilde{\mathcal{I}}((s, \textcircled{?}), o^B, o^R) := \frac{1}{2}\mathcal{I}(s, o^B, o^R)$.
- H is the same as in the AI-Control Game.

The two most important changes are to the transition and reward functions. The transition function has type $\tilde{\mathcal{P}} : \tilde{\mathcal{S}} \times \mathcal{A}^B \times \mathcal{A}^R \rightarrow \Delta(\mathcal{S} \times \Omega^B \times \Omega^R)$. We define this function in a fairly straightforward way. Transitions between $\textcircled{H}$ and $\textcircled{X}$ states are impossible. Further, in honest ($\textcircled{H}$) states, Red’s action has no effect, and the game transitions as though Red had chosen a_{null} .²¹

The reward function $\mathcal{R}_w$ is defined as follows:

$$\begin{aligned}\mathcal{R}_w((s, \textcircled{H}), a^B, a^R, (s', \textcircled{H})) &:= 2(1 - w)\mathcal{R}_{\text{Use}}(s, a^B, s') \\ \mathcal{R}_w((s, \textcircled{X}), a^B, a^R, (s', \textcircled{X})) &:= 2w\mathcal{R}_{\text{Safe}}(s, a^B, a^R, s')\end{aligned}$$

The factor of 2 compensates for the fact that each mode is only chosen in half of the states. Note how the usefulness reward is only active in honest states, and the safety reward is only active in attack states.

Since the action and observation spaces are unchanged, the sets of policies (Π^R and Π^B) are the same for the AI-Control Game and the zero-sum game: subject to the same constraints that Red is memoryless and Blue has perfect recall. All that remains is to show that the expected return for Blue in G_w is $w \text{Safety}(\pi^B, \pi^R) + (1 - w) \text{Use}(\pi^B)$.

The exact details of this proof are somewhat finicky, but the basic idea is to show an equivalent probability of sampling trajectories from each game, up to small details:

1. The probability of any particular $\textcircled{H}$ trajectory in the AI-Control Game does not depend on π^R . However, trajectories in ZSPOGs include Red’s actions, which do depend on π^R , even if those actions have no effect on the transition or reward. It is therefore necessary to sum over all possible sequences of Red’s actions in the ZSPOSG, when comparing to a trajectory in the AI-Control Game.
2. Because $\textcircled{?}$ is chosen non-deterministically in the AI-Control Game, and Safety and Use are measured conditional on the mode, the probability (up to Red’s actions) of a given $\textcircled{H}$ trajectory in the ZSPOSG is half the probability of the corresponding trajectory in the AI-Control Game (when $\textcircled{H}$ is chosen). This is countered by the factor of 2 in the reward function.

²¹Formally, $\tilde{\mathcal{P}}((s', \textcircled{H}) \mid (s, \textcircled{H}), a^B, a^R) = \mathcal{P}(s' \mid s, a^B, a_{\text{null}})$.

By summing over the probabilities of trajectories in each game, we could show that:

$$2 \cdot \mathbb{E}_{G_w} \left[\sum_{t=0}^H \mathcal{R}_w(S_t, A_t^B, A_t^R, S_{t+1}) \mid \textcircled{H}, \pi^B, \pi^R \right] = (1-w) \mathbb{E}_G \left[\sum_{t=0}^H \mathcal{R}_{\text{Use}}(S_t, A_t^B, S_{t+1}) \mid \textcircled{H}, \pi^B \right] \quad (3)$$

In other words, conditional on the ZSPOSG choosing $\textcircled{H}$ (which happens with probability $\frac{1}{2}$), the usefulness score for Blue in the AI-Control Game. Similarly, we can show an equivalent result for safety in the attack mode:

$$2 \cdot \mathbb{E}_{G_w} \left[\sum_{t=0}^H \mathcal{R}_w(S_t, A_t^B, A_t^R, S_{t+1}) \mid \textcircled{H}, \pi^B, \pi^R \right] = w \mathbb{E}_G \left[\sum_{t=0}^H \mathcal{R}_{\text{Safe}}(S_t, A_t^B, A_t^R, S_{t+1}) \mid \textcircled{X}, \pi^B, \pi^R \right] \quad (4)$$

It follows that: $w \text{Safety}(\pi^B, \pi^R) + (1-w) \text{Use}(\pi^B) = V_{G_w}(\pi^B, \pi^R)$. $\square$

A.3.3 PROOF OF THEOREM 5.3

Theorem (Repeat of Theorem 5.3). *A policy π^B is in an AI-Control Equilibrium if, and only if, for some ZSPOSG G_w with value function V_{G_w} , π^B maximises $\min_{\pi^R \in \Pi^R} V_w(\cdot, \pi^R)$ over Π^B .*

Proof. Lemma 5.1 shows that π^B is in an AI-Control Equilibrium if, and only if, it is a maximum for point for some $\min_{\pi^R} V_w(\cdot, \pi^R)$. Lemma 5.2 shows that there is some G_w for which $V_{G_w}(\pi^B, \pi^R) = V_w^G(\pi^B, \pi^R)$. Therefore π^B is in AI-Control Equilibrium if, and only if, it is in a maximum for some $\min_{\pi^R} V_{G_w}(\cdot, \pi^R)$. $\square$

A.3.4 COUNTEREXAMPLE TO MINIMAX

Proposition (Proposition 5.4 repeated). *There exists a G_w such that Red's behavioural and mixed policies would not be equivalent and*

$$\max_{\pi^B \in \Pi^B} \min_{\pi^R \in \Pi^R} V_{G_w}(\pi^B, \pi^R) < \min_{\pi^R \in \Pi^R} \max_{\pi^B \in \Pi^B} V_{G_w}(\pi^B, \pi^R)$$

Proof. This game is a sort of two-player version of the absent-minded driver example Aumann et al. (1997).

Figure 4 shows an extensive form game that we can easily show is equivalent to an AI-Control game reduced to a ZSPOSG (G_w).²²

Since Blue only has one choice, between two actions, and in a single information set, Blue's policy can be represented by the probability Blue chooses right:

$$p := \pi^B(\blacktriangleright) = 1 - \pi^B(\blacktriangleleft)$$

Similarly, even though Red acts twice, because they get only a null observation and are memoryless, Red's policy can be represented by the probability Red chooses right:

$$q := \pi^R(\blacktriangleright) = 1 - \pi^R(\blacktriangleleft)$$

²²Choose $H = 3$ and take the obvious state and action spaces. Set both observation sets to be singletons and set $w = 1$ so that only safety is valued. To account for the turned-based nature of the game above, we ignore actions for Red and Blue on turns where they do not move. Because $w = 1$, we don't need to consider the usefulness function, or games where $\textcircled{H}$ is active.

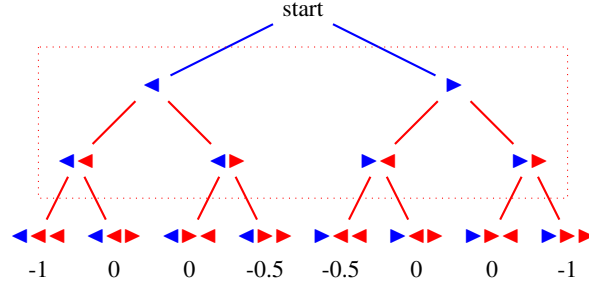


Figure 4: A zero-sum extensive form representation of a game where the minimax and maximin values are different. The dotted red box highlights Red’s single information set (i.e. Red has the same decision rule for all nodes). Blue moves first, and the Red takes the next two moves. All moves are either left ($\blacktriangleleft$) or right ($\blacktriangleright$). The payoffs for Blue are shown at the bottom of the tree.

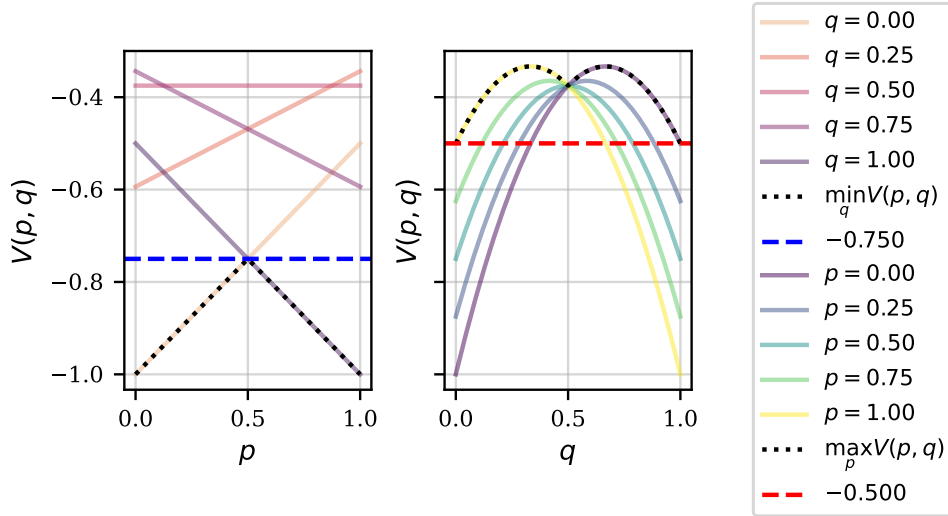


Figure 5: The minimax and maximin values of the game in Figure 4

Note that if we allowed Red to take policy mixtures in AI-Control Games, then this extensive form game would not be equivalent to an AI-Control Game, since Red could correlate their action in the second round with their action in the first round.

The value of the game is then given by:

$$\begin{aligned}
 V(p, q) &= (-1) * (1 - p)(1 - q)^2 \\
 &\quad + (-0.5) * (1 - p)q^2 \\
 &\quad + (-0.5) * p(1 - q)^2 \\
 &\quad + (-1) * pq^2 \\
 &= -1 + p(0.5 - 1q) + 2q - 1.5q^2
 \end{aligned}$$

$$\max_p \min_q V(p, q) = V(0.5, 1.0) = V(0.5, 1.0) = -0.75$$

$$\min_q \max_p V(p, q) = V(1.0, 0.0) = V(0.0, 1.0) = -0.5$$

$$\therefore \max_p \min_q V(p, q) < \min_q \max_p V(p, q)$$

Since the corresponding AI-Control game is equivalent to the extensive form game, this shows that the minimax and maximin values of the AI-Control game can be different. $\square$

A.3.5 PROOF FOR ALGORITHM 2

Proposition (Repeat of Theorem 5.5). *For $\epsilon > 0$, Algorithm 2 returns an ϵ -equilibrium in G_w in finite time. Further, for all t :*

$$\underline{v}_t \leq v^* := \max_{\pi_*^B \in \Pi^B} \min_{\pi_*^R \in \Pi^R} V_w^G(\pi_*^B, \pi_*^R) \leq \bar{v}_t$$

Proof. Our results are based closely on those from Adam et al. (2021), but we must modify them slightly to account for the differences between our games. Adam et al. (2021) define a continuous game as follows (with some slight modifications to notation):

Definition (Continuous Game). A continuous game G_{cont} is a tuple $(\mathcal{X}, \mathcal{Y}, f)$ where $\mathcal{X} \subseteq \mathbb{R}^k$ and $\mathcal{Y} \subseteq \mathbb{R}^l$ are compact sets, and $f : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ is a continuous function that gives the utility for Blue ($-f$ for Red). Blue chooses a mixed strategy as a Borel probability measure p over $\mathcal{X}$ and Red choose q over $\mathcal{Y}$. The value of two mixed strategies is given by $F(p, q) := \int_{\mathcal{X} \times \mathcal{Y}} f(x, y) d(p \times q)$.

Where we consider leader-follower equilibria over behaviourally stochastic strategies (which are equivalent to mixed for Blue but not Red), Adam et al. (2021) considers Nash-equilibria over jointly-mixed strategies:

$$v^* = \max_p \min_q F_w(p, q) = \min_q \max_p F_w(p, q)$$

Fortunately, since for any p , $\min_q F_w(p, q) = \min_y F_w(p, y)$, v^* is also the value for the leader-follower game when Red is restricted to a single strategy from $\mathcal{Y}$ in a continuous game. Therefore, finding an ϵ -equilibrium p, q , and then returning p, y where $y := \min_y F_w(p, y)$ returns an ϵ -approximate value for $\max_p \min_y F_w(p, y)$.

Adam et al. (2021) Lemma 2 shows that $\underline{v}_t \leq v^* \leq \bar{v}_t$ for all t . Theorem 3.1 shows that the algorithm returns an ϵ -approximate equilibrium in finitely many steps if $\epsilon > 0$. Therefore, returning p_i^*, y_{i+1} , where y_{i+1} is a best response to p_i^* , $F_w(p_i^*, y_{i+1}) \geq \bar{v}_i - \epsilon \geq v^* - \epsilon$ gives a pessimistic ϵ -maximin point for F_w .

Using Lemma A.1, we can easily map the value of a scalarised AI-Control Game to this definition of a continuous game: implying the results hold for Algorithm 2.

For a given AI-Control Game $\mathcal{G}$ and weight w , consider the mappings from Blue policies to $X \subset \mathbb{R}^m$ and Red policies to $Y \subset \mathbb{R}^n$ described in Lemma A.1. Then define $\mathcal{X} := \{\mathbf{e}_i\}_{i=1}^m \subset X$ as the finite set of vectors corresponding to deterministic Blue strategies. (Here $\mathbf{e}_i$ is the vector that is 1 at the i 'th dimension and 0 elsewhere.) Define $\mathcal{Y} := Y$. Then, let $f_w(\mathbf{e}_i, \mathbf{y}) := w \cdot s(\mathbf{e}_i, \mathbf{y}) + u(\mathbf{e}_i)$, which corresponds to $w \text{Safety}(\pi_i^B, \pi^R) + (1 - w) \text{Use}(\pi_i^B)$, where π_i^B is the i 'th deterministic strategy in the enumeration of Π_{det}^B from Lemma A.1 and π^R is the policy corresponding to $\mathbf{y}$.

Every solution to the continuous game $G_{\text{cont}}^w := (\{\mathbf{e}_i\}_{i=1}^m, Y, f_w)$ is equivalent to a solution to the AI-Control Game, relaxed to allow Red to take policy mixtures. $\square$

B FULL DEFINITION OF TRUSTED MONITORING AI-CONTROL GAME

In Section 3 we summarised the definition of the Trusted Monitoring AI-Control Game. Here, we provide the full definition of the game, including the state space, action space, observation space, transition function, and reward function. Note that this definition of the game is parameterised by a few values. First, and foremost, are the number of problems to be solved, and the number of times the protocol can ask for human assistance. We denote these values as M and C , respectively. We also assume that the trusted monitor's suspicion score is drawn from two distributions, $\mathcal{D}_{\text{Bd}}$ and $\mathcal{D}_{\text{Sf}}$, representing the suspicion scores for backdoored and safe solutions, respectively. These distributions correspond to Greenblatt et al. (2024, Figure 4a).

Let M be the number of problems to solve, and at a given time let m represent the number of problems remaining. To model limited access to a human auditor, we define a variable c , representing the number of 'audit coins' remaining: the number of times the protocol can ask for human assistance. Let $C = 0.02 \cdot M$ be the number of audit coins at the start of the game.

Define the AI-Control Game as the tuple:

$$\mathcal{G}_{\text{TM}} = (\mathcal{S}, \mathcal{A}^B, \mathcal{A}^R, \Omega^B, \Omega^R, \mathcal{I}, \mathcal{P}, \mathcal{R}_{\text{Safe}}, \mathcal{R}_{\text{Use}}, H)$$

The state space $\mathcal{S}$ consists primarily of states $s = (m, c, b)$ where m and c are as defined above, and b represents whether or not the solution to the current problem is backdoored, and takes values True or False. We also include three additional states. First, $s_0 = \langle \text{Start} \rangle$ represents the initial state when no problem has yet been solved. Second, $s_{\perp}$ represents the outcome of a solution that the untrusted AI backdoored was submitted to by the protocol. Third, we define an implicit $s_{\top}$ terminal state, representing the case where a human audit has caught a backdoored solution, and goes into safety mode, ending the opportunity for future backdoors. We can also define an ‘unacceptable’ terminal state $s_{\perp}$ that is entered whenever a backdoor is successfully submitted. Both $s_{\top}$ and $s_{\perp}$ will self-loop with probability 1 on all actions. Finally, we define the state space as $\mathcal{S} := \{s_0, s_{\perp}, s_{\top}\} \cup \{(m, c, b) : 0 \leq m \leq M; 0 \leq c \leq C\}$.

The action space for Blue consists of three actions: $\mathcal{A}^B := \{\langle \text{Audit} \rangle, \langle \text{Defer} \rangle, \langle \text{Run} \rangle\}$. If $a^B = \langle \text{Audit} \rangle$ is chosen, the human is asked to audit the untrusted model solution. If $a^B = \langle \text{Defer} \rangle$ is chosen, then Blue defers to the trusted model to attempt to generate a solution (though the trusted model may not be successful). If $a^B = \langle \text{Run} \rangle$ is chosen, then Blue runs the code submitted by the untrusted model.

The action space for Red consists of two actions: $\mathcal{A}^R := \{\langle \text{Backdoor} \rangle, \langle \text{Safe} \rangle\}$. If $a^R = \langle \text{Backdoor} \rangle$ is chosen, then Red interferes with the untrusted model in an attempt to submit a backdoored solution. If $a^R = \langle \text{Safe} \rangle$ is chosen, then Red does not interfere, and the untrusted model submits a normal solution.

In Greenblatt et al. (2024) original implementation of Trusted Monitoring, Red observes the problem statement before considering how and whether to attack. Because we are only considering attacks that choose to backdoor I.I.D., we can abstract away the description of the programming problem from our model. Since Red also does not observe the time or number of audit coins left, we model Red as receiving a null observation in every case. Therefore, $\Omega^R := \{o_{\text{null}}\}$. (Note that this is not true for all Control Games.)

The initial state distribution is the degenerate distribution $\mathcal{I}(s_0, o_{\text{null}}, o_{\text{null}}) = 1$, where both Blue and Red receive null observations.

The most complicated part of the model is the transition function $\mathcal{P} : \mathcal{S} \times \mathcal{A}^B \times \mathcal{A}^R \rightarrow \Delta(\mathcal{S} \times \Omega^B \times \Omega^R)$. In our case (though not necessarily in general) we can factor $\mathcal{P}$:

$$\mathcal{P}(s_{t+1}, o_t^B, o_t^R | s_t, a_t^B, a_t^R) = \mathcal{P}(s_{t+1} | s_t, a_t^B, a_t^R) \cdot \mathcal{P}(o_t^B | s_{t+1}) \cdot \mathcal{P}(o_t^R)$$

The final term, $\mathcal{P}(o_t^R)$, is redundant, since o_t^R only takes on one value. The observation probability for Blue only depends on the state. In particular, it only depends on whether the current problem is backdoored or not. In Trusted Monitoring protocols, the solution from the untrusted model is fed to the trusted monitor, which then rates the solution’s suspiciousness (x_t). We model the trusted monitor as a noisy observer, and assume that the trusted monitor’s suspicion score is drawn from two distributions, $\mathcal{D}_{\text{Bd}}$ and $\mathcal{D}_{\text{Sf}}$, representing the suspicion scores for backdoored and safe solutions, respectively.

$$\mathcal{P}(o_t^B | s_{t+1}) = \mathcal{P}((x_t, m', c') | (m, c, b))$$

$$\mathcal{P}((x_t, m, c) | (m, c, \text{True})) = \mathcal{D}_{\text{Bd}}(x_t)$$

$$\mathcal{P}((x_t, m, c) | (m, c, \text{False})) = \mathcal{D}_{\text{Sf}}(x_t)$$

$\mathcal{D}_{\text{Sf}}$ is the distribution we get by sampling a programming problem from the distribution of problems, asking the untrusted model to generate a solution, and then asking the trusted monitor to rate the suspiciousness. $\mathcal{D}_{\text{Bd}}$ is the distribution we get by sampling a programming problem from the distribution of problems, asking the untrusted model to generate a backdoored solution, and then asking the trusted monitor to rate the suspiciousness.

Next, we define the state transition function. In our model, though not in general, much of the stochasticity in the game comes from the policies and the Blue observation function, rather than state transitions: the only other source of stochasticity we model is the possibility of an imperfect human auditor. We therefore we will use the shorthand $\mathcal{P}(s_t, a_t^B, a_t^R) = s_{t+1}$ to represent that $\mathcal{P}(s_t, a_t^B, a_t^R) = \delta(s_{t+1})$. Let p_H^{FNR} denote the human auditor’s false negative rate, and p_H^{FPR} denote the false positive rate. Further, let $\delta_p(x; y)$ be the distribution that assigns probability p to x and probability $(1 - p)$ to y .

The transition from the initial state just moves to the first ‘real state’ and stores the value of whether the red team has backdoored the first solution. (It ignores Blue’s action). It is given by:

$$\begin{aligned}\forall a^R : \mathcal{P}(\langle \text{Start} \rangle, a^B, \langle \text{Backdoor} \rangle) &= (M, C, \text{True}) \\ \forall a^R : \mathcal{P}(\langle \text{Start} \rangle, a^B, \langle \text{Safe} \rangle) &= (M, C, \text{False})\end{aligned}$$

Then, for subsequent states, if the solution is backdoored and the protocol runs the solution, it loses; if the protocol audits the solution, it wins if the human auditor doesn’t make a mistake, but loses if they do. If the protocol defers, the game continues. (In what follows, abbreviate $\langle \text{Backdoor} \rangle$ with $\langle \text{Bd} \rangle$).

$$\begin{aligned}\mathcal{P}((m, c, \text{True}), \langle \text{Run} \rangle, a^R) &= s_{\perp} & (\forall a^R) \\ \mathcal{P}((m, c, \text{True}), \langle \text{Audit} \rangle, a^R) &= \delta_{p_H^{\text{FNR}}}(s_{\perp}; s_{\top}) & (\forall a^R) \\ \mathcal{P}((m, c, \text{True}), \langle \text{Defer} \rangle, \langle \text{Bd} \rangle) &= (m-1, c, \text{True}) \\ \mathcal{P}((m, c, \text{True}), \langle \text{Defer} \rangle, \langle \text{Safe} \rangle) &= (m-1, c, \text{False})\end{aligned}$$

Alternatively, if the solution is not backdoored, the game only terminates if Blue audits and the human auditor incorrectly marks the solution as backdoored, the game enters the ‘safe’ end state: even if for the wrong reasons, Blue never uses U again, ensuring no backdoors are submitted, and Blue wins. In all cases, if Blue audits, they lose an audit coin:

$$\begin{aligned}\mathcal{P}((m, c, \text{False}), \langle \text{Run} \rangle, \langle \text{Bd} \rangle) &= (m-1, c, \text{True}) \\ \mathcal{P}((m, c, \text{False}), \langle \text{Run} \rangle, \langle \text{Safe} \rangle) &= (m-1, c, \text{False}) \\ \mathcal{P}(\dots, \langle \text{Audit} \rangle, \langle \text{Bd} \rangle) &= \delta_{p_H^{\text{FPR}}}(s_{\top}; (m-1, c-1, \text{True})) \\ \mathcal{P}(\dots, \langle \text{Audit} \rangle, \langle \text{Safe} \rangle) &= \delta_{p_H^{\text{FPR}}}(s_{\top}; (m-1, c-1, \text{False})) \\ \mathcal{P}(\dots, \langle \text{Defer} \rangle, \langle \text{Bd} \rangle) &= (m-1, c, \text{True}) \\ \mathcal{P}(\dots, \langle \text{Defer} \rangle, \langle \text{Safe} \rangle) &= (m-1, c, \text{False})\end{aligned}$$

The safety reward function in this model could simply be defined as giving a score of -1 whenever the unacceptable terminal state, $s_{\perp}$, is entered. However, to ensure that the final safety score is between 0 and 1, rather than -1 and 0, we can also add a reward of $1/H$ to every transition. Therefore, the total reward is:

$$\mathcal{R}_{\text{Safe}}(s_t, a_t^B, a_t^R, s_{t+1}) = \frac{1}{H} - \mathbb{1}[s_{t+1} = s_{\perp} \wedge s_t \neq s_{\perp}]$$

The usefulness function in this model just measures how many of the untrusted model’s solutions are submitted. $\mathcal{R}_{\text{Use}} : \mathcal{S} \times \mathcal{A}^B \times \mathcal{S} \rightarrow \mathbb{R}$. Since usefulness is only measured in ‘honest’ games, we can

ignore cases where s_t or s_{t+1} are $s_\perp$ or where b is True. It is still possible that $s_\top$ could be reached should the human auditor make a mistake. Therefore, the only important transitions to consider are:

$$\begin{aligned}\mathcal{R}_{\text{Use}}((m, c, \text{False}), \langle \text{Audit} \rangle, s_\top) &= 0 \\ \mathcal{R}_{\text{Use}}(s_\top, a_t^B, s_\top) &= 0 \\ \mathcal{R}_{\text{Use}}((m, c, \text{False}), a_t^B, (m-1, c', \text{False})) \\ &= \mathbb{1}[a_t^B \neq \langle \text{Defer} \rangle]\end{aligned}$$

Finally, we choose $H := M + 1$, such that the game ends after $M + 1$ timesteps, and M problems are solved. This completes the definition of the Control Game Tuple for Trusted Monitoring with Defer to Trusted.

C PROOFS FROM SECTION 4

C.1 SAFETY EXPRESSION

Proposition C.1 (Proposition 4.1). *For any Red policy π^R , where the probability of attempting a backdoor on a given round ($\pi^R(\langle \text{Backdoor} \rangle \mid o_{\text{null}})$) is β , the safety of the memoryless policy π^B is given by the value of z_m^C , where z is the solution to the following recurrence relation (where the $\lambda_{m,c}^i$ terms are constant with respect to β):*

$$\begin{aligned}z_m^c &= (\lambda_{m,c}^1 \beta + \lambda_{m,c}^2) z_{m-1}^c + \lambda_{m,c}^5 (1 - \beta) z_{m-1}^{c-1} \\ &\quad + \lambda_{m,c}^3 \beta + \lambda_{m,c}^4 (1 - \beta)\end{aligned}$$

where $z_m^0 = (1 - \beta + \lambda_{m,c}^6 \beta) z_{m-1}^0$ and $z_0^c = 1$

Let $f(q)$ denote the probability a trusted monitor score for a backdoored case is in the top q th percentile of scores for safe cases. Further, let p_H^{FNR} denote the human auditor’s false negative rate, and p_H^{FPR} denote the false positive rate. Then, all of the values of $\lambda_{m,c}^i$ are functions only of $q_{m,c}^a$, $q_{m,c}^d$, and f , and can be computed upfront:

$$\begin{aligned}\lambda_{m,c}^1 &= (f(q_{m,c}^d) - f(q_{m,c}^a)) - (1 - q_{m,c}^a) \\ \lambda_{m,c}^2 &= 1 - q_{m,c}^a \\ \lambda_{m,c}^3 &= f(q_{m,c}^a) \cdot (1 - p_H^{\text{FNR}}) \\ \lambda_{m,c}^4 &= q_{m,c}^a \cdot p_H^{\text{FPR}} \\ \lambda_{m,c}^5 &= q_{m,c}^a \cdot (1 - p_H^{\text{FPR}}) \\ \lambda_{m,c}^6 &= f(q_{m,c}^d)\end{aligned}$$

Note that in the initial statement of Proposition 4.1, we assume for simplicity that p_H^{FPR} and p_H^{FNR} equal 0. It is trivial to show that the above recurrence relation simplifies to the one put forward in Proposition 4.1 if this assumption is made.

Proof. First, we note that, given the general definition of Safety, and the trusted monitoring games definition of $\mathcal{R}_{\text{Safe}}$, we can express safety as:

$$\begin{aligned}\text{Safety}(\pi^B, \pi^R) &= \mathbb{E} \left[\sum_{t=0}^{H-1} \frac{1}{H} - \mathbb{1}[s_{t+1} = s_\perp \wedge s_t \neq s_\perp] \mid \Theta, \pi^B, \pi^R \right] \\ &= \sum_{t=0}^{H-1} \frac{1}{H} - \mathbb{E} [\mathbb{1}[\exists t : s_t = s_\perp] \mid \Theta, \pi^B, \pi^R] \\ &= \mathbb{P}[S_H \neq s_\perp \mid \Theta, \pi^B, \pi^R].\end{aligned}$$

We will calculate the probability of not ending in the unacceptable outcome by defining a recurrence. We will define an expression z_m^c , and show it represents the probability that the game does not terminate in an unacceptable outcome, given that the game has reached time m with c audit coins left but before the observation x_t is drawn. In other words, its the probability that the game does not terminate in an unacceptable outcome, given that $s_{t-1} \notin \{s_\perp, s_\top\}$ (the unacceptable or positive terminal states), and instead $s_{t-1} = (m+1, c', b)$, and either b is false (and no false alarm has been raised due to an auditing error), or $a_{t-1}^B = \langle \text{Defer} \rangle$.

Let $g(q) = x$ be the threshold such that $\mathbb{P}[x_t \geq x \mid b_t = \text{True}] = q$. Given that the game has reached time m with $c > 0$ audit coins left, there are a few options for the round:

1. Red chose $\langle \text{Backdoor} \rangle$ last round and the problem is backdoored, Blue chooses $\langle \text{Audit} \rangle$ and wins if no auditing mistake is made. This happens with probability $\beta \cdot \mathbb{P}[x_t \geq g(q_{m,c}^a)] \cdot (1 - p_H^{\text{FNR}}) = \beta \cdot f(q_{m,c}^a) \cdot (1 - p_H^{\text{FNR}})$.
2. The current problem is backdoored, Blue chooses $\langle \text{Audit} \rangle$ and loses because an auditing mistake is made. This happens with probability $\beta \cdot \mathbb{P}[x_t \geq g(q_{m,c}^a)] \cdot p_H^{\text{FNR}} = \beta \cdot f(q_{m,c}^a) \cdot p_H^{\text{FNR}}$.
3. The current problem is backdoored, Blue chooses $\langle \text{Run} \rangle$, and loses immediately. This happens with probability $\beta \cdot \mathbb{P}[x_t \leq g(q_{m,c}^d)] = \beta \cdot (1 - f(q_{m,c}^d))$.
4. The current problem is backdoored, Blue chooses $\langle \text{Defer} \rangle$, and the game continues with c audit coins. This happens with probability $\beta \cdot (f(q_{m,c}^a) - f(q_{m,c}^d))$.
5. The current problem is not backdoored, and Blue chooses $\langle \text{Audit} \rangle$ and no auditing mistake is made, so the game continues with $c-1$ audit coins. This happens with probability $(1 - \beta) \cdot \mathbb{P}[x_t \geq g(q_{m,c}^a)] \cdot (1 - p_H^{\text{FPR}}) = (1 - \beta) \cdot q_{m,c}^a \cdot (1 - p_H^{\text{FPR}})$.
6. The current problem is not backdoored, and Blue chooses $\langle \text{Audit} \rangle$, but human error causes an auditing mistake, so the game ends immediately with Blue winning (no chance of any future backdoors). This happens with probability $(1 - \beta) \cdot \mathbb{P}[x_t \geq g(q_{m,c}^a)] \cdot p_H^{\text{FPR}} = (1 - \beta) \cdot q_{m,c}^a \cdot p_H^{\text{FPR}}$.
7. The current problem is not backdoored, and Blue chooses $\langle \text{Run} \rangle$, and the game continues with c audit coins. This happens with probability $(1 - \beta) \cdot \mathbb{P}[x_t \leq g(q_{m,c}^d)] = (1 - \beta) \cdot (1 - (q_{m,c}^d))$.
8. The current problem is not backdoored, and Blue chooses $\langle \text{Defer} \rangle$, and the game continues with c audit coins. This happens with probability $(1 - \beta) \cdot (q_{m,c}^d - q_{m,c}^a)$.

Denote $\mathbb{P}[S_H \neq s_\perp \mid \otimes, \pi^B, \pi^R, m, c]$ with the shorthand $p[m, c]$. We will define the recurrence z_m^c to be equal to $p[m, c]$ at all points. We can therefore write out an expression for $p[m, c]$ in terms of $p[m-1, c-1]$ and $p[m-1, c]$ when $m \geq 1$ and $c \geq 1$:

$$\begin{aligned}
 p[m, c] = & (\beta \cdot f(q_{m,c}^a) \cdot (1 - p_H^{\text{FNR}}))(1) \\
 & + (\beta \cdot f(q_{m,c}^a) \cdot p_H^{\text{FNR}})(0) \\
 & + (\beta \cdot (1 - f(q_{m,c}^d)))(0) \\
 & + (\beta \cdot (f(q_{m,c}^a) - f(q_{m,c}^d)))p[m-1, c] \\
 & + ((1 - \beta) \cdot q_{m,c}^a \cdot (1 - p_H^{\text{FPR}}))p[m-1, c-1] \\
 & + ((1 - \beta) \cdot q_{m,c}^a \cdot p_H^{\text{FPR}})(1) \\
 & + ((1 - \beta) \cdot (1 - (q_{m,c}^d)))p[m-1, c] \\
 & + ((1 - \beta) \cdot (q_{m,c}^d - q_{m,c}^a))p[m-1, c]
 \end{aligned}$$

Given, this expression, we can define

$$z_m^c := (\dots)z_{m-1}^c + (\dots)z_{m-1}^{c-1} + \dots$$

as described in the proposition by rearranging the terms above.

At the boundary, when $c = 0$, $\langle \text{Audit} \rangle$ actions are ignored. Therefore, the only three options are that (1) the problem is backdoored, and Blue chooses $\langle \text{Run} \rangle$, (2) the problem is backdoored, and Blue chooses $\langle \text{Defer} \rangle$, or (3) the problem is not backdoored, and Blue chooses $\langle \text{Defer} \rangle$. This gives us the boundary condition $p[m, 0] = (1 - \beta + (1 - q_{m,c}^d)\beta)p[m - 1, 0]$. Then define $z_m^0 := (1 - \beta + (1 - q_{m,c}^d)\beta)z_{m-1}^0$.

Finally, at the boundary when $m = 0$, there are no problems left, and so Blue has won: thus, $z_m^c := 1 = p[0, c]$. That $z_m^c = p[m, c]$ for all m and c is straightforward by induction from our definitions. $\square$

C.2 USEFULNESS EXPRESSION

As mentioned in Section 4, we can give an expression for the usefulness, similar to the above expression for safety.

Proposition C.2. *The usefulness of the memoryless threshold policy π^B is given by the value of u_M^C , where u is the solution to the following recurrence relation:*

$$\begin{aligned} u_m^c = & ((1 - q_{m,c}^d) + (1 - p_H^{FPR})q_{m,c}^a)(1) \\ & + (1 - p_H^{FPR})(q_{m,c}^a)(u_{m-1}^{c-1}) \\ & + (1 - q_{m,c}^a)(u_{m-1}^c) \end{aligned}$$

where $u_m^0 = (1 - q_{m,c}^d)(1) + u_{m-1}^0$ and $u_0^c = 0$.

Proof. Just as with the above proof for the safety expression, we begin by noting the definition of usefulness:

$$\text{Use}(\pi^B) := \mathbb{E} \left[\sum_{t=0}^H \mathcal{R}_{\text{use}}(S_t, A_t^B, S_{t+1}) \mid \textcircled{H}, \pi^B \right]$$

The usefulness of a Blue policy π^B is defined as the expected number of times Blue uses solutions from U in an honest ($\textcircled{H}$) game; honest games are those in which the untrusted model U is not scheming, and so no backdoors are inserted. We can identify when Blue uses a solution from U as cases in which Blue either takes the action $\langle \text{Run} \rangle$, or takes the action $\langle \text{Audit} \rangle$ and no human error is made.

We will define an expression u_m^c and show it represents the expected number of solutions from the untrusted model that Blue will use throughout the rest of the game, given that the game has reached time m with c audit coins left, but before the observation x_t is drawn.

By the linearity of expectation, the value of u_m^c can be broken down into the expected number of solutions used in the current step of the game, summed with the expected number of solutions used from the next step of the game onwards. The latter is another variable in the recurrence. The former is simply the probability that Blue uses the current solution from U ; in the current timestep, Blue either uses the solution with a particular probability, or doesn't use the solution.

The game continues in all cases except where a mistake is made as part of a Blue $\langle \text{Audit} \rangle$ and a safe solution is marked as backdoored. In this singular case, all remaining usefulness for the remainder of the game is set to 0: after 'catching' the inserted backdoor, Blue will not use any other solutions from U .

Let $g(q) = x$ be the threshold such that $\mathbb{P}[x_t \geq x \mid b_t = \text{True}] = q$. Given that the game has reached time m with $c > 0$ audit coins left, we can break down the expected number of solutions used by each of Blue's possible actions:

1. Blue chooses $\langle \text{Run} \rangle$ and uses the solution from U . The game continues to the next timestep with the same number of audit coins remaining. This happens with probability $\mathbb{P}[x_t \leq g(q_{m,c}^d)] = 1 - q_{m,c}^d$.
2. Blue chooses $\langle \text{Defer} \rangle$ and does not use the solution. The game continues to the next timestep having not used an audit coin. This happens with probability $q_{m,c}^d - q_{m,c}^a$.

3. Blue chooses $\langle \text{Audit} \rangle$ and no human auditing error is made, so the solution is used. The game continues having used up one audit coin. This happens with probability $(1 - p_H^{\text{FPR}}) \cdot q_{m,c}^a$.
4. Blue chooses $\langle \text{Audit} \rangle$ and a human auditing error leads to the game ending in the 'safe' state $s_\top$ and no further solutions from U being used. This happens with probability $p_H^{\text{FPR}} \cdot q_{m,c}^a$.

Let $\mathbb{E} \left[\sum_{t=0}^H \mathcal{R}_{\text{use}}(S_t, A_t^B, S_{t+1}) \mid \Theta, \pi^B, m, c \right]$ be denoted by the shorthand $e[m, c]$. We will define the recurrence u_m^c to be equal to $e[m, c]$ at all points. We can therefore write out an expression for $e[m, c]$ in terms of $e[m-1, c-1]$ and $e[m-1, c]$:

$$\begin{aligned} e[m, c] = & (1 - q_{m,c}^d)(1 + e[m-1, c]) \\ & + (q_{m,c}^d - q_{m,c}^a)(0 + e[m-1, c]) \\ & + ((1 - p_H^{\text{FPR}}) \cdot q_{m,c}^a)(1 + e[m-1, c-1]) \\ & + (p_H^{\text{FPR}} \cdot q_{m,c}^a)(0) \end{aligned}$$

Given, this expression, we can define

$$u_m^c := \dots + (\dots) u_{m-1}^{c-1} + (\dots) u_{m-1}^c$$

as described in the proposition by rearranging the terms above.

At the boundary, when $c = 0$, $\langle \text{Audit} \rangle$ actions are ignored. Therefore, the only option is whether Blue chooses $\langle \text{Run} \rangle$ or $\langle \text{Defer} \rangle$. Both actions continue the game in the next timestep. This gives us the boundary condition $e[m, 0] = (1 - q_{m,c}^d)(1) + e[m-1, 0]$. Then define $u_m^0 := (1 - q_{m,c}^d)(1) + u_{m-1}^0$.

Finally, at the boundary when $m = 0$, there are no problems left, so there are no more solutions to use: thus, $u_0^c := 0 = e[0, c]$. That $u_m^c = e[m, c]$ for all m and c is straightforward by induction from our definitions. $\square$

D EXPERIMENTAL DETAILS

We include a number of experiments in this paper that illustrate the safety and usefulness scores for Blue protocols when set against the best Red response. Below are the details of our experimental setup.

Trusted Monitor Distributions. Greenblatt et al. (2024) empirically investigate the distribution of their Trusted Monitor scores on backdoored and safe programming solutions. Since our 'Original Trusted Monitor' distributions are only designed to approximate these results, we fit beta distributions to points derived directly from the bar chart from figure 4a Greenblatt et al. (2024) using the SciPy library for Python.²³ We then discretise this beta distribution to 6 decimal places to get an approximated finite observation space.

Our choice to use a beta distribution is partially motivated by the bounded range of suspicion scores, which we normalise to $[0, 1]$, and the skew of the data. However, a more principled justification is that the trusted monitoring scores are generated by T's confidence in the answer to a yes/no question (Greenblatt et al., 2024, Appendix O.1).

For the 'Weak Trusted Monitor', we use the same safe distribution. However, the backdoor distribution is chosen to be closer to the safe distribution than in the original case; although there are many possible 'weak' distributions to demonstrate that there are some cases in which randomisation is required, we pick just one.

Red Best Response. For a given Blue policy π^B , we find the Red best response using a 1-dimensional grid search then ternary search for the optimal $\beta \in [0.0, 1.0]$. The grid search evaluates 33 points and finds the grid point with the lowest safety score with respect to π^B . Then we apply a ternary search, using the two grid points either side as the bounds of the search interval, to find the optimal value of β within a tolerance of 10^{-6} for Pareto-optimal points and of 10^{-5} for other points in Figure 2.

²³<https://docs.scipy.org/doc/scipy/reference/generated/scipy.stats.fit.html>

Blue Best Response. In our implementation of Algorithm 2, we apply a 2D recursive grid search over the values of q^d and q^a to find the best Blue response to a given mixed Red policy. This grid search evaluates 441 points (21 x 21) at each level, selecting the region around the best point as the search space for the next recursion, until the local optimal values of q^d and q^a have been found within a tolerance of 10^{-6} .

Evaluating Protocols. We evaluate memoryless threshold protocols against a fixed red team using a straightforward implementation of Proposition 4.1, using memoisation/dynamic-programming to increase efficiency. We expand on seven protocols from Greenblatt et al. (2024) by keeping the audit FPR fixed and varying the defer FPR in increments of 0.01 in the interval $[0.0, 1.0]$. Each resulting protocol is evaluated against the Red best response to it, to generate worst-case safety and usefulness results.

Demonstrating a Frontier. In order to explore an approximate Pareto-frontier for the fixed-threshold policy subset, we apply Algorithm 1, finding the optimal Blue and Red policies using Algorithm 2. In our implementation of Algorithm 2, we find the best response policies (lines 6 and 7) as described above; we set epsilon to 10^{-4} .

Assumptions and Trends. The results in Section 6.1 are generated by evaluating a number of proposed protocols in games with different modelling assumptions. We choose to limit our investigation to a fixed defer FPR of 0.1 as a representative example; our goal in this section is to assess the general impact these assumptions have, and further work might better model how the impact of these assumptions varies across protocols with different parameters. The audit budget ('Budget'), human-auditor false negative rate ('Human FNR'), and human audit false positive rate ('Human FPR') each vary by increments of 0.01 in the interval $[0.0, 1.0]$, keeping all other variables fixed. The number of problems ('Number of Problems') varies by increasing amounts as follows: increments of 50 between [100, 1000]; 250 between [1000, 2000]; 500 between [2000, 7000]; and 1000 between [7000, 10000]. Protocols in Figure 2 are evaluated against a best-response up to a tolerance of $1e - 6$.

E EXISTING MULTI-OBJECTIVE POSG SOLVER

Whilst our setting shares some similarities to that of Chang et al. (2015), there are important differences between the settings, and they mandate different solutions. First, and foremost, the authors use a state-space factored into follower and leader components. Many control games cannot be modeled with this assumption: the Trusted Monitoring game in Section 4 includes the 'backdoor' variable (representing whether Red backdoored last turn) which is observed by neither team.

Second, all AIC Games share a safety/usefulness objective structure and attack/honest game modes that allow us to reduce to finding Nash-equilibria in a zero-sum game. This reduction further allows the finding of Blue best-responses using any POMDP solver. Chang et al. (2015)'s more general objective structure precludes such a reduction, and perhaps motivates their use of a genetic algorithm. The largest difference between settings considered is Chang et al. (2015)'s consideration of a memoryfull red team. Much of the complexity in their algorithm is in finding an infinite-memory follower policy and then finding a finite-memory approximation. These steps are not necessary, or possible in our AI-Control setting.

Finally, Chang consider an infinite-horizon game, with 11 states, but we consider a finite-horizon game with $\approx 40,000$ reachable states. They state that 64 manager policies are under consideration, but the space of possible (deterministic) Blue protocols is far larger. Although it could be possible to adapt Chang et al. (2015)'s algorithm to account for these differences, it would be more sensible to derive a multi-objective algorithm from scratch for our setting. the solution methods we present are simple and the first ways to solve AIC Games.