

Learning Game Parameters from MSNE: An Application to Learning IDS Games

Hau Chan Luis E. Ortiz

Department of Computer Science, Stony Brook University
Stony Brook, NY, 11794-4400

Abstract

A survey is a popular and common method for eliciting behavioral data on a topic from a sample population. Such behavioral data captures the actions of the sampled population under some possibly unknown environment. Quite often, we do not have information about the individual responses due to privacy concerns or bookkeeping overloads. Instead, what we typically observe is some form of aggregation or summarization of the individual responses that represents the percentages of the individuals who reportedly took certain actions. Because, as we assume, each person is strategic and takes the best action given the actions of other people, we view the given behavioral data as a set of possible (approximate) mixed-strategy Nash equilibrium (MSNE) of some game. Given this, our goal is to learn a game that would best explain or rationalize the behavior of the population. In this work, we introduce a machine learning (generative) framework to learn the structure and parameters of games given a set of possible (approximate) MSNE for the purpose of predicting and analyzing behavior, even under causal intervention or counterfactual queries. Under our framework, we show that, under some mild assumptions, maximizing the log-likelihood of a game given behavioral data is equivalent to finding a game that maximizes the number of (approximate) MSNE in the data while maintaining the overall proportion of (approximate) MSNE of the game as low as possible. Moreover, we illustrate the effectiveness of our framework by learning the parameters of generalized interdependent security games from real-world vaccination data publicly available from the Center for Disease Control and Prevention (CDC) in the United States.

1 Introduction

A survey is an important tool for eliciting information from agents in large populations. Many government agencies such as the *Centers for Disease Control and Prevention (CDC)* in the United States sample some subsets of the population, and elicit information from them via a survey or a questionnaire. For example, the CDC can ask an agent, “Did you take the H1N1 vaccine this month?”

We are particularly interested in this type of question because it reveals the action an agent took perviously. Indeed, if we believe that the agents are rational, self-interested, and their decisions affect the decisions of others, then we can attempt to learn and infer a game in which the actions the agents took are the “best-responses” of all other agents, and no agent would benefit from unilaterally deviating from their current action (i.e., from taking a vaccine to not taking a vaccine, or vice versa). These actions (collectively) are known as *pure-strategy Nash Equilibrium (PSNE)* and we will define it more formally later. However, quite often, we do not get to see the completed survey of each individual in the population. This is, perhaps, due to the large amount of data and information: it is impossible to keep track of all the details. Instead, what we typically can publicly obtain is some compact representation of the data that summarizes or aggregates the information collected.

More concretely, using the CDC vaccination data as our running example, we can observe the monthly state vaccination percentages (along with standard deviations) for different age groups, race groups, and types of vaccinations. These vaccination percentages represent the *average* behavior of the people in the USA. Indeed, if we view each state as an agent, we can interpret each vaccination rate at each state in the USA as the probability, or, using game-theory parlance, the “mixed strategy”, that the corresponding “state agent” vaccinates against a particular disease. (Please keep in mind that the state agent’s mixed strategy really corresponds to the percentage of people in that state that would decide to vaccinate against that particular disease.) Moreover, if we further consider the fact that the behaviors (vaccination rates) of the state agents affect the vaccination rates of others, then we can model the vaccination scenario, at the level of states, as a game. In particular, agents are the states in the USA, actions are either to vaccinate or not vaccinate, and the payoff of each state is some function that roughly capture the average preferences or utilities of all the individuals in the population of that state. We assume that the payoff function is im-

plicit in the behavioral data and that we can learn it or infer it by trying to “rationalize” each state’s behavior in the given dataset of vaccination rates. Said differently, here, we do not know the exact game the agents are playing and hence we do not know their payoffs. However, we have data that potentially specify the mixed strategies of the agents. Some may correspond to an (equilibrium) outcome of the game. Therefore, using the vaccination rates as mixed strategies, we can learn a game that could partially capture some of these mixed strategies as outcomes, or, as we view it, *mixed-strategies Nash equilibria (MSNE)* of the game. In addition, we do expect some of these mixed strategies to be noisy. Therefore, we assume that some of these mixed strategies may be approximate or ϵ -MSNE of the games and try to find ϵ as small as possible to capture the variations.

In this work, not only we introduce a general machine learning (generative) framework to learn any arbitrary game given the data, we also provide a way to learn a class of *generalized interdependent security (α -IDS) games* [Chan and Ortiz, 2014]. We will define α -IDS games more formally in Section 3.1), which we use to model vaccination decisions.

We now describe the way we use the CDC data, and delay the details about our learning framework and the mechanisms we derive to learn α -IDS games. Using the 2009-2010 states H1N1 vaccination percentages and their standard deviations, we generate 10,000 (mixed-strategy) examples according to normal distributions with means and standard deviations of the states. Each example represents a single mixed-strategy profile of the 48-states in the continental USA (i.e., excluding Alaska and Hawaii). Given these examples, we aim to learn a (α -IDS) game that would best explain the generated data.

Our main interest for learning games is the ability they provide to potentially interpret what would happen at an MSNE, even when the given data may not be an exact MSNE, or be noisy. The mixed strategies of the state agents in our data may not correspond to the optimal equilibrium strategies. Therefore, we want to infer the (real) behavior of the states at equilibrium from noisy data, in which not all examples may belong to the set of approximate MSNE of some game.

Thus, given the learned games, we can run a version of some learning-heuristics/regret-minimization [Fudenberg and Levine, 1998], in which we use the average vaccination rates as the initial mixed-strategy profiles to compute ϵ -MSNE in these games. We expect that as ϵ goes to zero, we would be able to capture the true equilibrium strategies of the state agents.

Contribution. We conclude the introduction with a summary of our

contributions. Our interest in this work is learning games from observed mixed-strategy data. In contrast to previous work, in which the data are the actions or pure strategies of the players [Honorio and Ortiz, 2014], we are dealing with data that summarize the actions of all the individuals within a state’s population using rates. In our model, we view each rate as representing the mixed strategy of each state agent.

In game-theoretic terms, we view these probabilities collectively as (approximate) MSNE. In particular, we

- propose and introduce a machine learning (generative) framework to learn a game given the data;
- show that, under some mild conditions, maximizing the log-likelihood of the game is equivalent to maximizing the number of (approximate) MSNE under our framework;
- use our framework to derive a heuristic to learn α -IDS games given the CDC vaccination data; and
- experimentally show that our framework and learning heuristic are effective for inferring α -IDS games, and may be able to provide insight into the behavior of state agents.

1.1 Related Work

The closest work to ours is Honorio and Ortiz [2014] where they provide a general machine learning framework to learn the structure and parameters of games from discrete (e.g., “Yes/No” responses) behavioral data. Moreover, they demonstrate their framework on learning influence games [Irfan and Ortiz, 2014] using congressional voting data. For the sake of completeness, all other previous methods assume that the actions and payoffs are observable in the data [Wright and Leyton-Brown, 2010, 2012, Gao and Pfeffer, 2010, Vorobeychik et al., 2007, Ficici et al., 2008, Duong et al., 2009, 2012] while others are interested in predicting future behavior from the past behavior (system dynamics) [Kearns and Wortman, 2008, Ziebart et al., 2010]. We refer the reader to the related work section of Honorio and Ortiz [2014] for a more detailed discussion.

2 A Framework to Learn Games from Data

Let $V = \{1, 2, \dots, n\}$ be a set of players. For each $i \in V$, let A_i be the set of actions/pure-strategies available to i and $u_i : \times_{j \in V} A_j \rightarrow \mathbb{R}$ be the payoff of i given the actions of i and other $V - \{i\}$ agents. Let X_i be the set of mixed strategies of i , which is a simplex over A_i , and denote by $u_i(x) \equiv \mathbf{E}_{a \sim x}[u_i(x)]$ the expectation over the probabilities x_i of playing the pure-actions. A joint-mixed strategy $x^* \in \times_{i=1}^n X_i$ is a *mixed-strategy Nash equilibrium (MSNE)* of a non-cooperative game [Neumann and Morgenstern, 1944, Nash, 1950, 1951] if, for each player i , $x_i^* \in \arg \max_{x_i \in X_i} u_i(x_i, x_{-i}^*)$, where $x_{-i}^* \equiv (x_1^*, x_2^*, \dots, x_{i-1}^*, x_{i+1}^*, \dots, x_n^*)$. We denote the set of all MSNE of a game $\mathcal{G}$ as

$$\mathcal{NE}(\mathcal{G}) \equiv \{x^* \mid \forall i, x_i^* \in \arg \max_{x_i \in X_i} u_i(x_i, x_{-i}^*)\}$$

For the following definition, we assume that the utilities are normalized between 0 and 1. Given $\epsilon > 0$, a joint-mixed strategy $x^* \in X = \times_{i=1}^n X_i$ is an ϵ -MSNE of a non-cooperative game if, for each player i , $x_i^* \in \arg \max_{x_i \in X_i} u_i(x_i, x_{-i}^*) - \epsilon$. We denote the set of all ϵ -MSNE of a game $\mathcal{G}$ as

$$\mathcal{NE}^\epsilon(\mathcal{G}) \equiv \{x^* \mid \forall i, x_i^* \in \arg \max_{x_i \in X_i} u_i(x_i, x_{-i}^*) - \epsilon\}.$$

For the rest of the paper, we assume that each agent has two actions and the action set of each of the agents is either 0 or 1 (i.e., $A_i = \{0, 1\}$ for all i). As such, the mixed strategies of the agents are in $[0, 1]$ (i.e., $X_i = [0, 1]$ and with probability $x_i \in X_i$, player i plays action 1).

It is easy to see that $\mathcal{NE}(\mathcal{G}) \subseteq \mathcal{NE}^\epsilon(\mathcal{G}) \subseteq \mathcal{NE}^{\epsilon'}(\mathcal{G})$ for all $0 < \epsilon < \epsilon'$. Note that an ϵ -MSNE might not be close to any true MSNE.

2.1 A Generative Model for Behavioral Data on Joint-mixed-strategies

We adopt a similar learning approach to that of Honorio and Ortiz [2014]. However, this time the generative model of behavioral data is over the set of joint-mixed-strategies. Hence, a *probability density function (PDF)* over the n -dimensional hypercube $[0, 1]^n$ now defines the generative model. We point out that our results are nontrivial extensions analogous to those of Honorio and Ortiz [2014] but in a continuous space, as oppose to the use of a *probability mass function (PMF)* over the set of joint pure-strategies $\{0, 1\}^n$,

a discrete space. To keep things simple, we present our results without proofs; the proofs are appropriately adapted versions of those of Honorio and Ortiz [2014].

Let $\mathcal{G}$ be a game and the approximation parameter $\epsilon > 0$. We assume that the set $\mathcal{NE}^\epsilon(\mathcal{G})$ is (Borel) μ -measurable and denote its measure by $|\mathcal{NE}^\epsilon(\mathcal{G})| \equiv \mu(\mathcal{NE}^\epsilon(\mathcal{G}))$. (We leave the question of whether ϵ -MSNE is always measurable open.) We assume the statistical process generating the data is a simple mixture model: i.e., with probability $q \in (0, 1)$, the process generates/outputs a joint-mixed-strategy x by drawing uniformly at random from the set $\mathcal{NE}^\epsilon(\mathcal{G})$; with probability $1 - q$, the process generates a joint-mixed-strategy x by drawing uniformly at random from $\overline{\mathcal{NE}^\epsilon(\mathcal{G})} \equiv [0, 1]^n - \mathcal{NE}^\epsilon(\mathcal{G})$, the complement of the $\mathcal{NE}^\epsilon(\mathcal{G})$. Said differently, our generative model of behavioral data based on joint-mixed-strategies is a mixture model with mixture parameter q and mixture components defined in terms of the approximation parameter $\epsilon > 0$ and a game $\mathcal{G}$, with the implicit assumption that $\mathcal{NE}^\epsilon(\mathcal{G})$ is measurable, with Borel measure value $|\mathcal{NE}^\epsilon(\mathcal{G})|$; otherwise the model is undefined. Note that, in our context, because $\mu([0, 1]^n) = 1$, we can view the Borel measure μ as a probability measure. From now on, all references to measures are to the Borel (probability) measure, and μ denotes such measure. We also assume that when we refer to the set $\mathcal{NE}^\epsilon(\mathcal{G})$, it is measurable with respect to the respective $\mathcal{G}$ and ϵ , unless stated otherwise. More formally, the PDF f for the generative model with parameters $(q, \mathcal{G}, \epsilon)$ over the hypercube of joint-mixed-strategies $[0, 1]^n$ is

$$f_{(q, \mathcal{G}, \epsilon)}(x) \equiv q \frac{\mathbb{1}[x \in \mathcal{NE}^\epsilon(\mathcal{G})]}{|\mathcal{NE}^\epsilon(\mathcal{G})|} + (1 - q) \frac{\mathbb{1}[x \notin \mathcal{NE}^\epsilon(\mathcal{G})]}{1 - |\mathcal{NE}^\epsilon(\mathcal{G})|}, \quad (1)$$

for all $x \in [0, 1]^n$, assuming $\mathcal{NE}^\epsilon(\mathcal{G})$ is measurable; otherwise, $f_{(q, \mathcal{G}, \epsilon)}(x)$ is undefined.

In order for eq. 1 to be valid, if $\epsilon = 1$ or $\mathcal{NE}^\epsilon(\mathcal{G}) = [0, 1]^n$, then we need to require that $q = 1$. Note that $\mathcal{NE}^\epsilon(\mathcal{G}) = \emptyset$ is impossible since every game has at least one MSNE, by Nash's seminal result [Nash, 1950, 1951].

We assume that the behavioral data are i.i.d. instances drawn according to $f_{(q, \mathcal{G}, \epsilon)}$.

Definition 2.1. (Trivial and Non-trivial Games) *We say that a game $\mathcal{G}$ is trivial if and only if $|\mathcal{NE}^\epsilon(\mathcal{G})| \in \{0, 1\}$ and non-trivial if and only if $|\mathcal{NE}^\epsilon(\mathcal{G})| \in (0, 1)$.*

Let $\pi^\epsilon(\mathcal{G})$ be the *true proportion of ϵ -MSNE* in the game $\mathcal{G}$ where

$$\pi^\epsilon(\mathcal{G}) \equiv \mu(\mathcal{NE}^\epsilon(\mathcal{G})) . \quad (2)$$

The following set of propositions is analogous to those Honorio and Ortiz [2014] state and establishes the fact that there are different games with the same set of ϵ -MSNE. Said differently, the game $\mathcal{G}$ is not identifiable with respect to the generative model $f_{(q,\mathcal{G},\epsilon)}$ defined in eq. 1. We side-step the non-identifiability of $\mathcal{G}$ with respect to $f_{(q,\mathcal{G},\epsilon)}$ using a common practice in machine learning (ML): we invoke the *Principle of Ockham's Razor* for model (i.e., game) selection. In general, experts in the respective field (e.g., epidemiology) would provide the necessary bias for learning. Here, we impose a particular bias that induces “sparse” or “compactly representable” games, as we define formally in a later section. We note, however, that the games are identifiable in terms of their ϵ -MSNE, with respect to $f_{(q,\mathcal{G},\epsilon)}$, as we show and discuss later.

Proposition 2.2. *Given the approximation parameter $\epsilon > 0$ and a non-trivial game $\mathcal{G}$, the mixture parameter $q > \pi^\epsilon(\mathcal{G})$ if and only if $f_{(q,\mathcal{G},\epsilon)}(x_1) > f_{(q,\mathcal{G},\epsilon)}(x_2)$ for any $x_1 \in \mathcal{NE}^\epsilon(\mathcal{G})$ and $x_2 \notin \mathcal{NE}^\epsilon(\mathcal{G})$.*

Definition 2.3. *Given the approximation parameter $\epsilon > 0$, we say the games $\mathcal{G}_1$ and $\mathcal{G}_2$ are ϵ -approximation-equivalent, or ϵ -equivalent, if and only if their approximate Nash equilibrium sets are identical, i.e.: $\mathcal{G}_1 \equiv_{\epsilon\text{-MSNE}} \mathcal{G}_2 \Leftrightarrow \mathcal{NE}^\epsilon(\mathcal{G}_1) = \mathcal{NE}^\epsilon(\mathcal{G}_2)$.*

Lemma 2.4. *Let $\mathcal{G}_1$ and $\mathcal{G}_2$ be two non-trivial games. Given the approximation parameter $\epsilon > 0$, for some mixture parameter $q > \max(\pi^\epsilon(\mathcal{G}_1), \pi^\epsilon(\mathcal{G}_2))$, $\mathcal{G}_1$ and $\mathcal{G}_2$ are ϵ -equivalent if and only if they induce the same pdf over the mixed strategies of space $[0, 1]^m$ of the agents. Moreover, $\mathcal{G}_1 \equiv_{\epsilon\text{-MSNE}} \mathcal{G}_2 \Leftrightarrow \forall x \ f_{(q,\mathcal{G}_1,\epsilon)}(x) = f_{(q,\mathcal{G}_2,\epsilon)}(x)$.*

2.2 Learning the Parameters of the Games via MLE

In this section, we present a mean to estimate the parameters of a graphical game from data. The data, as we will assumption, will be a set of ϵ -MSNE.

For the following, we recall the Kullback-Leibler (KL) divergence between two Bernoulli distributions with parameters $p_1, p_2 \in (0, 1)$, which, using common practice to simplify the presentation, we denote by

$$\text{KL}(p_1||p_2) \equiv p_1 \log \frac{p_1}{p_2} + (1 - p_1) \log \frac{1 - p_1}{1 - p_2} .$$

Given a dataset $D = \{x^{(1)}, \dots, x^{(m)}\}$ drawn *i.i.d.* according to $f_{q, \mathcal{G}, \epsilon}$, let $\hat{\pi}^\epsilon(\mathcal{G})$ be the *empirical proportion of ϵ -MSNE*, i.e.,

$$\hat{\pi}^\epsilon(\mathcal{G}) \equiv \frac{1}{m} \sum_{l=1}^m \mathbb{1}[x^{(l)} \in \mathcal{NE}^\epsilon(\mathcal{G})]$$

Proposition 2.5. (Maximum-likelihood Estimation) *The tuple $(\hat{\mathcal{G}}, \hat{q}, \hat{\epsilon})$ is a maximum likelihood estimator (MLE), with respect to dataset D , for the parameters of the generative model $f_{(q, \mathcal{G}, \epsilon)}$, as defined in eq. 1 if and only if $\hat{q} = \min\left(\hat{\pi}^\epsilon(\hat{\mathcal{G}}), 1 - \frac{1}{2m}\right)$, and $(\hat{\mathcal{G}}, \hat{\epsilon}) \in \arg \max_{(\mathcal{G}, \epsilon)} \text{KL}(\hat{\pi}^\epsilon(\mathcal{G}) \parallel \pi^\epsilon(\mathcal{G}))$.*

Let us make a few observations that follow immediately from the MLE expression given above. First, if $\epsilon \geq 1$, then $\pi^\epsilon(\mathcal{G}) = 1$ for all games $\mathcal{G}$, which implies then $\hat{\pi}^\epsilon(\mathcal{G}) = 1$ for all games $\mathcal{G}$. Hence, if $\hat{\epsilon} \geq 1$ the resulting KL value is zero, so that $\hat{\mathcal{G}}$ could be *any* game. Similarly, if $\pi^\epsilon(\hat{\mathcal{G}}) = 0$ then we have $\hat{\pi}^\epsilon(\hat{\mathcal{G}}) = 0$, so that once again the resulting KL value is zero. Hence, $\hat{\mathcal{G}}$ could be *any* game. Said differently, in summary, if *any* trivial game is an MLE, then *every* game, trivial or non-trivial, is also an MLE. Therefore, we can always find non-trivial games corresponding to some MLE: the set of MLEs always contain a tuple corresponding to a non-trivial game.

An informal interpretation of the MLE problem is that, assuming we can keep the true proportion of ϵ -MSNE low, the learning problems becomes one of trying to infer a game that captures as much of the mixed-strategy examples in the dataset as ϵ -MSNE, but without implicitly adding more ϵ -MSNE than it needs to. Thus, formulating the learning problem using MLE brings out the fundamental tradeoff in ML between model complexity and generalization ability (or “goodness-of-fit”), despite the simplicity of our generative model.

One problem with the exact KL-based formulation of the MLE presented above is that dealing with $\pi^\epsilon(\mathcal{G})$, even if it is well-defined (i.e., the set $\mathcal{NE}^\epsilon(\mathcal{G})$ is measurable). The following lemma provides bounds on the KL divergence that will prove useful in our setting.

Lemma 2.6. *Given a non-trivial game $\mathcal{G}$ with $0 < \pi^\epsilon(\mathcal{G}) < \hat{\pi}^\epsilon(\mathcal{G})$, we can upper and lower bound the KL divergence as*

$$-\hat{\pi}^\epsilon(\mathcal{G}) \log \pi^\epsilon(\mathcal{G}) - \log 2 < \text{KL}(\hat{\pi}^\epsilon(\mathcal{G}) \parallel \pi^\epsilon(\mathcal{G})) < -\hat{\pi}^\epsilon(\mathcal{G}) \log \pi^\epsilon(\mathcal{G}).$$

From the last lemma, we argue (details omitted) that if $\pi^\epsilon(\mathcal{G})$ is “low enough” then we can obtain an approximation to the MLE by simply maximizing $\hat{\pi}^\epsilon(\mathcal{G})$ only: i.e., $\arg \max_{\mathcal{G}} \text{KL}(\hat{\pi}^\epsilon(\mathcal{G}) \parallel \pi^\epsilon(\mathcal{G})) \approx \arg \max_{\mathcal{G}} \hat{\pi}^\epsilon(\mathcal{G})$. We implicitly enforce the constraint that $\pi^\epsilon(\mathcal{G})$ is “low enough” through regularization or some other way that allows us to introduce bias into the model selection, as it is standard in ML.

Therefore, we aim to develop techniques to maximize the number of ϵ -MSNE in the data while keeping ϵ as small as possible. In what follows, we will apply our learning framework to infer the parameters of generalized Interdependent Security (α -IDS) games [Chan and Ortiz, 2014] using the CDC vaccination data.

3 Application: Learning α -IDS Games

Given the CDC vaccination data, we want to learn a game that would explain the behavior of the agents and how the behavior of the agents affect the behavior of other agents (within the same population). In particular, we are interested in understanding the behavior of an “average” individual in each state when facing the question of “What is the probability that a member of my population will transfer a virus/sickness to me if that member is ill?” Therefore, we want to look at games that model such interaction.

Generalized Interdependent Security (α -IDS) games [Chan and Ortiz, 2014] are one of the most motivated [Chan and Ortiz, 2014, Heal and Kunreuther, 2004] and well-studied games to model the investment decisions of agents when facing transfer risks from other agents. As Heal and Kunreuther [2005a] discusses, α -IDS games have applicability in fire protection Kearns and Ortiz [2004], and, more importantly in vaccination settings [Heal and Kunreuther, 2005b]. We refer the reader to a recent survey by Laszka et al. [2014] for a broader perspective on these type of models, including numerous other applications.

In the vaccination setting, each agent decides whether or not to get vaccinated given (1) the agent’s implicit and explicit cost of vaccination and loss of getting sick, (2) the vaccination decisions of other agents, and (3) the potential transfer probabilities/risks from other agents. The CDC vaccination data captures the “average” behavior of the people in each state through the vaccination rates, but does not explicitly contain the costs or losses of any individual, nor the transfer risk between individuals. Actually, it does not

include the “average” costs, losses, or transfer risks even at the level of whole states. In what follows, we put forward an approach to learn such quantities at the state level from the CDC data on vaccination rates.

3.1 Generalized Interdependent Security Games

In the α -IDS games of n agents, each agent i determines whether or not to invest in protection. Therefore, there are two actions i can play, and we denote $a_i = 1$ if i invests and $a_i = 0$ if i does not invest. We let $a = (a_1, \dots, a_n)$ to be the joint-actions of all agents and a_{-S} to be the joint-actions of all agents that are not in S . There is a cost of investment C_i and loss L_i associated with the bad event occurring, either through a direct or indirect (transferred) contamination. We denote by p_i the probability that agent i will experience the bad event from a direct contamination and by q_{ji} to be the probability that agent i will experience the bad event due to transfer exposure from agent j (i.e., the probability that agent j will transfer the contamination to i). Moreover, the parameter $\alpha_i \in [0, 1]$ specifies the probability that agent i 's investment will not protect that agent against transfers of a bad event. Given the parameters, the *cost function of agent i* is

$$\begin{aligned} M_i(a_i, a_{-i}) &\equiv a_i[C_i + \alpha_i r_i(a_{-i})L_i] \\ &\quad + (1 - a_i)[p_i + (1 - p_i)r_i(a_{-i})]L_i \end{aligned}$$

where $r_i(a_{-i}) \equiv 1 - s_i(a_{-i})$ and $s_i(a_{-i}) \equiv \prod_{j \neq i} (a_j - (1 - a_j)(1 - q_{ji}))$ are the overall risk and safety functions of agent i .

As mentioned before, we aim to learn all of these parameters given that we can observe the mixed strategies. Therefore, we need look at the cost function of the agents in terms of mixed strategies. Roughly speaking, we can do this by letting x_i be the probability that $a_i = 1$ and take the expectation of the above cost function (i.e., replace all a terms by x). Comparing the cost when $a_i = 1$ and $a_i = 0$, we can derive a best-response correspondence for i .

Therefore, the cost function of player i becomes (in mixed strategies)

$$\begin{aligned} M_i(x_i, x_{-i}) &\equiv x_i[C_i + \alpha_i r_i(x_{-i})L_i] \\ &\quad + (1 - x_i)[p_i + (1 - p_i)r_i(x_{-i})]L_i. \end{aligned}$$

By definition, a mixed strategy is an ϵ -MSNE in an α -IDS game if and only

if

$$M_i(x_i, x_{-i}) - \epsilon \leq M_i(0, x_{-i}) \quad (3)$$

$$M_i(x_i, x_{-i}) - \epsilon \leq M_i(1, x_{-i}) \quad (4)$$

It follows that from Equation 3 and Equation 4 that

$$\begin{aligned} x_i[C_i + \alpha_i r_i(x_{\text{Pa}(i)})L_i - (p_i + (1 - p_i)r_i(x_{\text{Pa}(i)}))L_i] &\leq \epsilon \\ -(1 - x_i)[C_i + \alpha_i r_i(x_{\text{Pa}(i)})L_i - (p_i + (1 - p_i)r_i(x_{\text{Pa}(i)}))L_i] &\leq \epsilon \end{aligned}$$

For simplicity, we let $\Delta_i \equiv C_i + \alpha_i r_i(x_{\text{Pa}(i)})L_i - (p_i + (1 - p_i)r_i(x_{\text{Pa}(i)}))L_i$.

3.2 Learning the Structure and Parameters of α -IDS Games

As we argue in a previous section, we can approximate our MLE objective by maximizing the number of ϵ -MSNE in the data, or equivalently, maximizing $\hat{\pi}^\epsilon(\mathcal{G})$ over ϵ and $\mathcal{G}$. In our approach, we subdivide the optimization by first optimizing over $\mathcal{G}$, and then optimizing over ϵ . For any ϵ , we would like to apply a simple gradient-ascent optimization technique to learn the game $\mathcal{G}$. Unfortunately, even the latter maximization is non-trivial due to the discontinuities induced by the indicator functions defining the ϵ -MSNE constraints. Our goal is then to further approximate $\hat{\pi}^\epsilon(\mathcal{G})$. First, we use a simple upper bound that results from using eq. 3 and eq. 4, which correspond to satisfying the ϵ -MSNE of the games:

$$\begin{aligned} \hat{\pi}^\epsilon(\mathcal{G}) &= \max_{\mathcal{G}} \frac{1}{m} \sum_{l=1}^m 1[\mathbf{x}^l \in NE_\epsilon(\mathcal{G})] \\ &\leq \max_{\mathcal{G}} \frac{1}{m} \sum_{l=1}^m \sum_{i=1}^n 1[x_i^l \Delta_i^l \leq \epsilon] + 1[-(1 - x_i^l) \Delta_i^l \leq \epsilon] \end{aligned}$$

Then, we approximate the indicator function in the last upper bound with another differentiable function. In the following subsection, we discuss what is perhaps the simplest approximation to the indicator function: using the logistic/sigmoid function. This is the standard approach leading to the famous BackProp algorithm used to train neural networks from data [Haykin, 1999].

3.2.1 Using the Logistic/Sigmoid Function

The first approximation to the upper bound above that we consider uses the sigmoid function, $s(x) \equiv \frac{1}{1+e^{-x}}$, which yields the following approximation to the last upper bound:

$$\max_{\mathcal{G}} \frac{1}{m} \sum_{l=1}^m \sum_{i=1}^n s(-x_i^l \Delta_i^l + \epsilon) + s((1 - x_i^l) \Delta_i^l + \epsilon).$$

To avoid overfitting and introduce our bias for “sparse” (graphical) game structures, we regularize the transfer parameters q_{ji} . In particular, the structure of the α -IDS games is implicitly defined by those transfer probabilities. That is, viewing α -IDS games from the perspective of a (directed, parametric) graphical game, the directed graph capturing the direct transfer risks between the players is such that each node in the graph represents a player in the game, and there is a directed edge (i.e., an arc) from node j to node i if and only if $q_{ji} > 0$. The typical regularizer used to induce sparsity in the learned structure is the L_1 -regularizer, which we impose over the q_{ji} ’s. We denote by $\lambda > 0$ the regularization parameter quantifying the amount of penalization for large values of the q_{ji} ’s.

$$\max_{\mathcal{G}} \frac{1}{m} \sum_{l=1}^m \sum_{i=1}^n S(-x_i^l \Delta_i^l + \epsilon) + S((1 - x_i^l) \Delta_i^l + \epsilon) + \lambda \sum_{j=1}^n q_{ji}.$$

We “learn” λ using cross-validation. (This is the typical approach to find an “optimal” λ in ML.)

Before continuing, there is an important normalization constraints on the utility/costs functions required for the ϵ parameter in approximation to be meaningful. In particular, recall that in order to define ϵ -MSNE, we want to ensure that the cost function of each player of α -IDS games is between zero and one for each possible mixed strategy. The following expression leads to a normalized cost function, denoted by $\widetilde{M}_i$, for player i :

$$\widetilde{M}_i(x_i, x_{-i}) \equiv \frac{M_i(x_i, x_{-i}) - \min_i}{\max_i - \min_i}$$

where $\min_i = \{C_i, p_i L_i\}$, $\max_i = \{C_i + \alpha_i r_i(0_{-i}) L_i, [p_i + (1 - p_i) r_i(0_{-i})] L_i\}$, and 0_{-i} stands for the vector that sets all the elements of x_{-i} to the value 0, so that $r_i(0_{-i}) = 1 - \prod_{j \neq i} (1 - q_{ji})$.

Notice that, if the minimum and the maximum, respectively, of the cost function of each player of α -IDS is exactly 0 and 1, respectively, then we do not have to perform any normalization when computing and evaluating ϵ -MSNE.

Unfortunately, working with the normalized costs $\widetilde{M}_i$'s is cumbersome. Instead, we keep the ϵ -MSNE constraints in terms of the original (unnormalized) cost functions M_i 's and introduce additional constraints based on the expressions for $\min_i$ and $\max_i$ given above directly into the optimization problem. Using primal-dual optimization, in which we denote by the corresponding dual-variables/Lagrange-multipliers β_i and γ_i for each additional cost-function normalization for each player i , we obtain the following minimax program:

$$\begin{aligned} \min_{\delta, \beta} \max_{\mathcal{G}} \frac{1}{m} \sum_{l=1}^m \sum_{i=1}^n S(-x_i^l \Delta_i^l + \epsilon) + S((1 - x_i^l) \Delta_i^l + \epsilon) + \lambda \sum_{j=1}^n q_{ji} \\ - \delta_i (C_i - 1)(p_i L_i - 1) \\ - \gamma_i (2 - (C_i + \alpha_i r_i (0_{-i}) L_i))(2 - ([p_i + (1 - p_i) r_i (0_{-i})] L_i)) , \end{aligned}$$

where $\beta = (\beta_1, \dots, \beta_n)$ and $\gamma = (\gamma_1, \dots, \gamma_n)$.¹ We want to solve the above program subject to the respective constraints on each of the variables. As stated previously, we follow the traditional approach of using gradient-ascent/descent optimization as a heuristic to update, and eventually learn, the parameters.

Denote by $(q', \mathcal{G}', \epsilon')$ the tuple we learn using the approach we propose above. In this paper we assume that $\mathcal{N}\mathcal{E}^{\epsilon'}(\mathcal{G}')$ is measurable, so that the learned generative model is well-defined.

4 Experiment

We are currently working on the experimental part of the paper. We expect to have the results by the date of the conference.

¹We intentionally enforce that $\min_i = \{C_i, p_i L_i\} = 1$ and $\max_i = \{C_i + \alpha_i r_i (0_{-i}) L_i, [p_i + (1 - p_i) r_i (0_{-i})] L_i\} = 2$ to avoid computational issues. As long as the difference of the $\min_i$ and $\max_i$ is close to 1, then we can easily see that the ϵ -MSNE definition will be well-defined.

References

- H. Chan and L. E. Ortiz. Computing nash equilibria in generalized interdependent security games. In *Proceedings of the Neural Information Processing Systems Foundation (NIPS)*, 2014.
- Quang Duong, Yevgeniy Vorobeychik, Satinder Singh, and Michael P. Wellman. Learning graphical game models. In *Proceedings of the 21st International Joint Conference on Artificial Intelligence, IJCAI'09*, pages 116–121, San Francisco, CA, USA, 2009. Morgan Kaufmann Publishers Inc. URL <http://dl.acm.org/citation.cfm?id=1661445.1661465>.
- Quang Duong, Michael P. Wellman, Satinder Singh, and Michael Kearns. Learning and predicting dynamic networked behavior with graphical multiagent models. In *Proceedings of the 11th International Conference on Autonomous Agents and Multiagent Systems - Volume 1, AAMAS '12*, pages 441–448, Richland, SC, 2012. International Foundation for Autonomous Agents and Multiagent Systems. ISBN 0-9817381-1-7, 978-0-9817381-1-6. URL <http://dl.acm.org/citation.cfm?id=2343576.2343639>.
- Sevan Ficici, David Parkes, and Avi Pfeffer. Learning and solving many-player games through a cluster-based representation. In *Proceedings of the Twenty-Fourth Conference Annual Conference on Uncertainty in Artificial Intelligence (UAI-08)*, pages 188–195, Corvallis, Oregon, 2008. AUAI Press.
- Drew Fudenberg and David K. Levine. *The Theory of Learning in Games*, volume 1 of *MIT Press Books*. The MIT Press, June 1998.
- Xi Alice Gao and Avi Pfeffer. Learning Game Representations from Data Using Rationality Constraints. In *Proceedings of the 26th Conference on Uncertainty in Artificial Intelligence (UAI'10)*, 2010.
- Simon Haykin. *Neural Networks: A Comprehensive Foundation*. Prentice Hall, 1999.
- Geoffrey Heal and Howard Kunreuther. Interdependent security: A general model. Working Paper 10706, National Bureau of Economic Research, August 2004.

- Geoffrey Heal and Howard Kunreuther. IDS models of airline security. *Journal of Conflict Resolution*, 49(2):201–217, April 2005a.
- Geoffrey Heal and Howard Kunreuther. The vaccination game. Working paper, Wharton Risk Management and Decision Processes Center, January 2005b.
- Jean Honorio and Luis E. Ortiz. Learning the structure and parameters of large-population graphical games from behavioral data. *Journal of Machine Learning Research*, (In Press), 2014.
- Mohammad T. Irfan and Luis E. Ortiz. On influence, stable behavior, and the most influential individuals in networks: A game-theoretic approach. *Artificial Intelligence*, 215:79–119, 2014.
- Michael Kearns and Luis E. Ortiz. Algorithms for interdependent security games. In *Advances in Neural Information Processing Systems*, NIPS '04, pages 561–568, 2004.
- Michael Kearns and Jennifer Wortman. Learning from collective behavior. In *In Proceedings of the 21st Annual Conference on Learning Theory*, pages 99–110, 2008.
- Aron Laszka, Mark Felegyhazi, and Levente Buttyan. A survey of interdependent information security games. *ACM Comput. Surv.*, 47(2):23:1–23:38, August 2014.
- John Nash. Non-cooperative games. *Annals of Mathematics*, 54(2):286–295, 1951.
- John F. Nash. Equilibrium points in n-person games. *Proceedings of the National Academy of Sciences of the United States of America*, 35(1):48–49, Jan. 1950.
- John Von Neumann and Oskar Morgenstern. *Theory of Games and Economic Behavior*. Princeton University Press, 1944.
- Yevgeniy Vorobeychik, Michael P. Wellman, and Satinder Singh. Learning payoff functions in infinite games. *Mach. Learn.*, 67(1-2):145–168, May 2007. ISSN 0885-6125. doi: 10.1007/s10994-007-0715-8. URL <http://dx.doi.org/10.1007/s10994-007-0715-8>.

- James Wright and Kevin Leyton-Brown. Beyond equilibrium: Predicting human behavior in normal-form games. In *AAAI Conference on Artificial Intelligence*, 2010. URL <https://www.aaai.org/ocs/index.php/AAAI/AAAI10/paper/view/1946>.
- James R. Wright and Kevin Leyton-Brown. Behavioral game theoretic models: A bayesian framework for parameter analysis. In *Proceedings of the 11th International Conference on Autonomous Agents and Multiagent Systems - Volume 2*, AAMAS '12, pages 921–930, Richland, SC, 2012. International Foundation for Autonomous Agents and Multiagent Systems. ISBN 0-9817381-2-5, 978-0-9817381-2-3. URL <http://dl.acm.org/citation.cfm?id=2343776.2343828>.
- Brian D. Ziebart, J. A. Bagnell, and Anind K. Dey. Modeling interaction via the principle of maximum causal entropy. In Johannes Fürnkranz and Thorsten Joachims, editors, *Proceedings of the 27th International Conference on Machine Learning (ICML-10)*, pages 1255–1262. Omnipress, 2010. URL <http://www.icml2010.org/papers/28.pdf>.