

Integrating Learning with Game Theory for Societal Challenges

Fei Fang
Carnegie Mellon University
feifang@cmu.edu

Security Challenges

Lack of Security Resources



Ansbach attack

A suicide bomb injured at least 12 in Germany's Ansbach, near Nuremberg, on July 24. This is the fourth violent incident in Germany in a week.



Source: Reuters

J. Wu, 25/07/2016

REUTERS

Sustainability Challenges

Lack of Ranger/Conservation Resources



100 years ago

≈ 60,000 tigers

Today

≈ 3,200 tigers

Mobility Challenges

Inefficiencies in Matching/Dispatching



Societal Challenges

Security & Safety



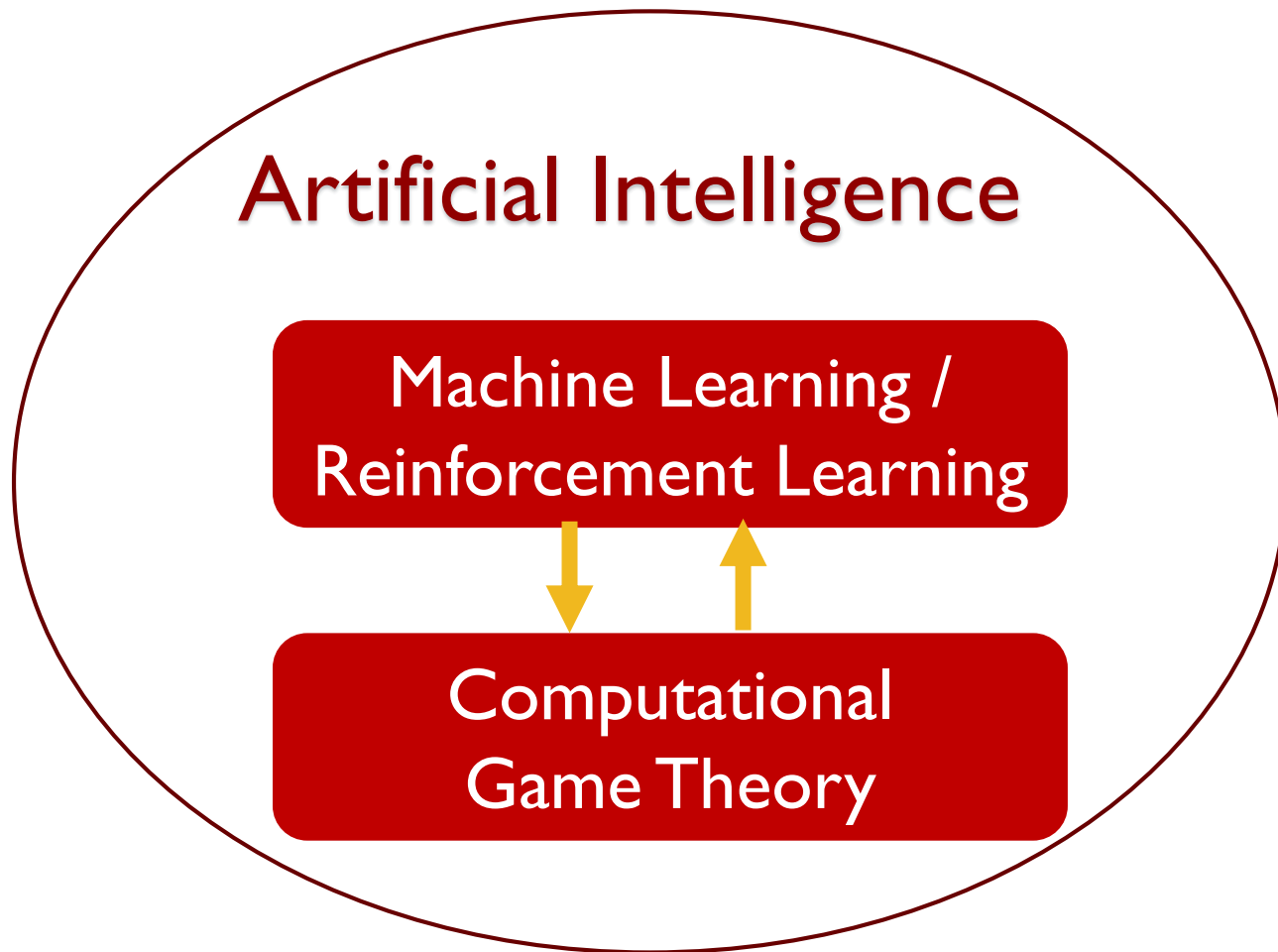
Environmental Sustainability



Mobility



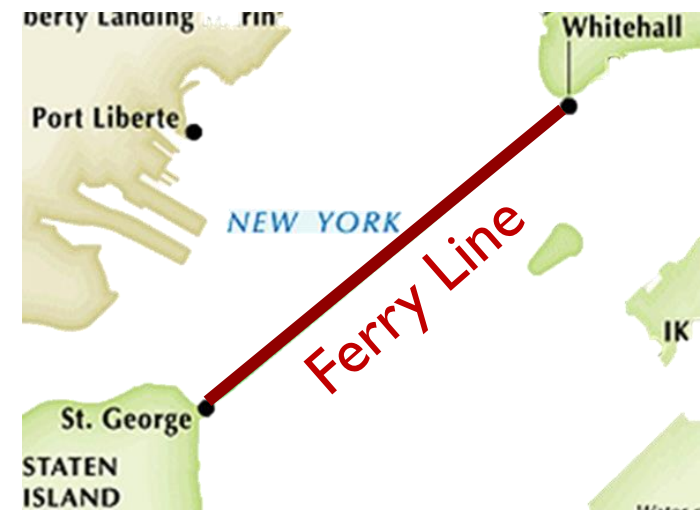
Integrate Learning with Game Theory



Outline

- ▶ Security games
- ▶ Integrating learning and game theory
 - ▶ Data-based game theoretic reasoning
 - ▶ Learning-powered equilibrium computation
 - ▶ End-to-end learning in games
- ▶ Summary

Protecting Staten Island Ferry



Protecting Staten Island Ferry



Protecting Staten Island Ferry



Previous USCG Approach



Protecting Staten Island Ferry



Problem



Game Model and Linear Programming-based Solution

- ▶ Stackelberg game: Leader – Defender, Follower – Attacker
- ▶ Attacker's payoff: $u_i(t)$ if not protected, 0 otherwise
- ▶ Zero-sum → Strong Stackelberg Equilibrium=Nash Equilibrium=Minimax (Minimize Attacker's Maximum Expected Utility)

$$\min_{p_r, v} v$$

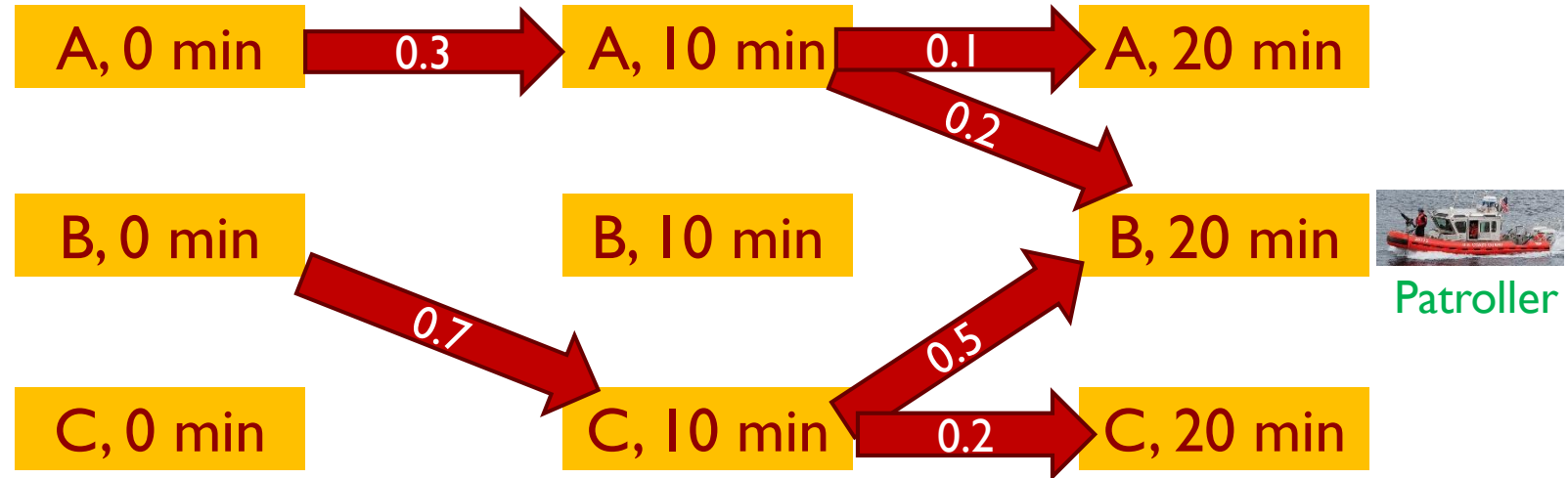
$$\text{s.t. } v \geq \mathbb{E}[U^{att}(i, t)] = u_i(t) \times \mathbb{P}[unprotected(i, t)], \forall i, t$$

Adversary

		10:00:00 AM Target 1	10:00:01 AM Target 1	...	10:30:00 AM Target 3	...
Defender	p_r					
	30%	Purple Route	-5, 5	-4, 4	0, 0	
	40%	Orange Route				
	20%	Blue Route				
						

Advanced Solution

► Flow-based Representation + Critical Time Points

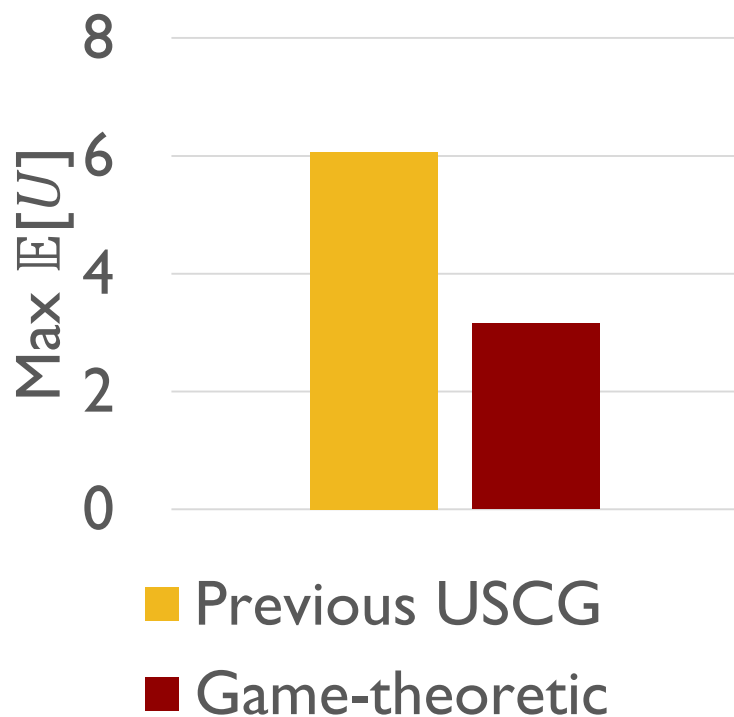


$$\begin{aligned}
 & \min_v \quad f \\
 \text{s.t.} \quad & \sum_{e \in (i,t) \rightarrow} f(e) = \sum_{e \in \rightarrow (i,t)} f(e) \\
 & \sum_{e \in (*,t) \rightarrow} f(e) = 1 \\
 & v \geq \mathbb{E}[U^{att}(i,t)], \forall i, \forall t \in \{t^*\}
 \end{aligned}$$

f → Prob. flow over feasible edges
 } → f is a unit flow
 → Best response

Evaluation: Simulation & Real-World Feedback

Reduce potential risk by 50%



- ▶ Deployed by US Coast Guard
- ▶ USCG evaluation
 - ▶ Point defense to zone defense
 - ▶ Increased randomness
- ▶ Professional mariners:
 - ▶ Apparent increase in Coast Guard patrols

Outline

- ▶ Security games
- ▶ Integrating learning and game theory
 - ▶ Data-based game theoretic reasoning
 - ▶ Learning-powered equilibrium computation
 - ▶ End-to-end learning in games
- ▶ Summary

Repeated and Frequent Poaching Activities

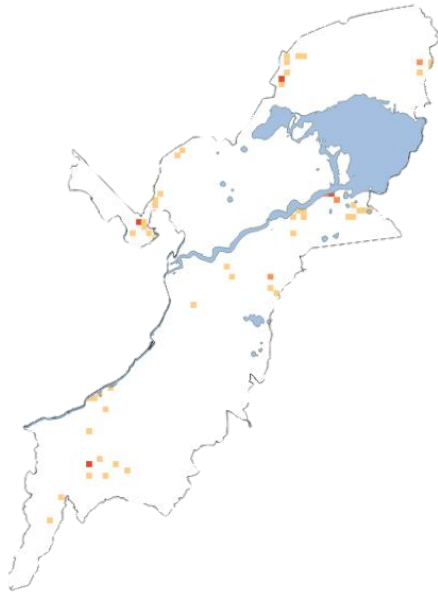


Data SIO, NOAA, U.S. Navy, NGA, GEBCO
Image Landsat
Image IBCAO

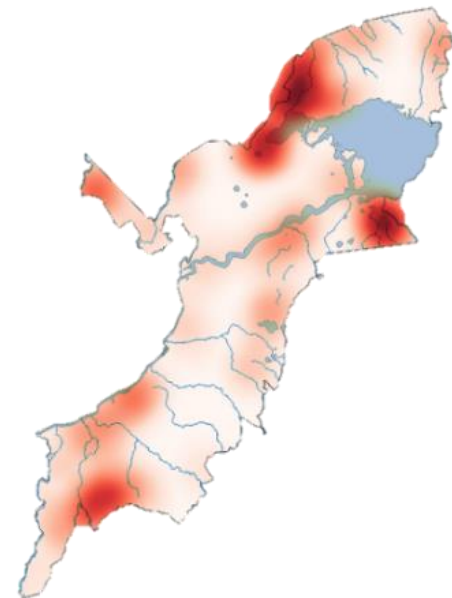
Google earth

Learn Poacher Behavior Model from Real-World Data

12yr of poaching data
from Queen Elizabeth
National Park in Uganda

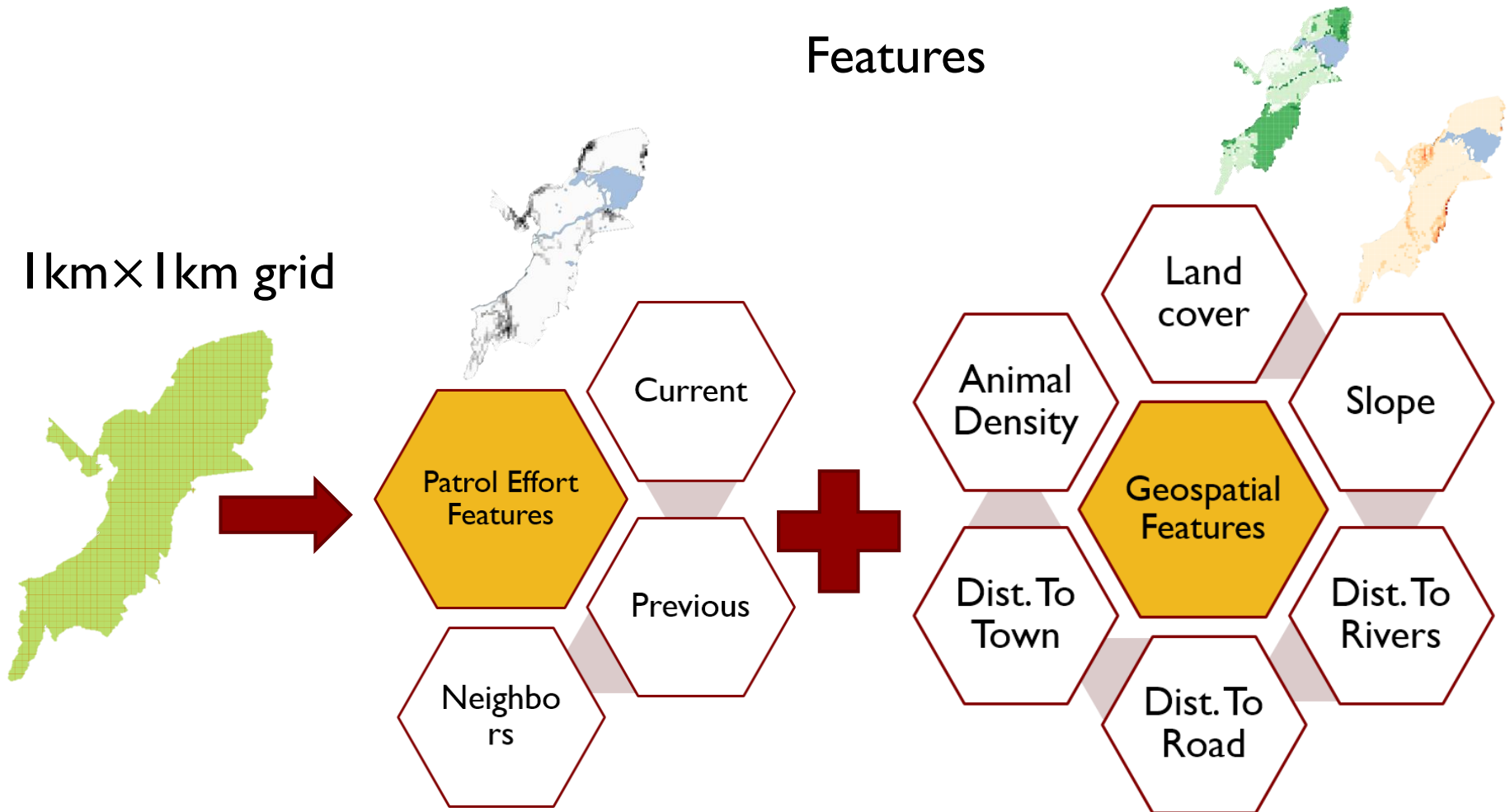


Prob. of snaring or
Prob. of detection



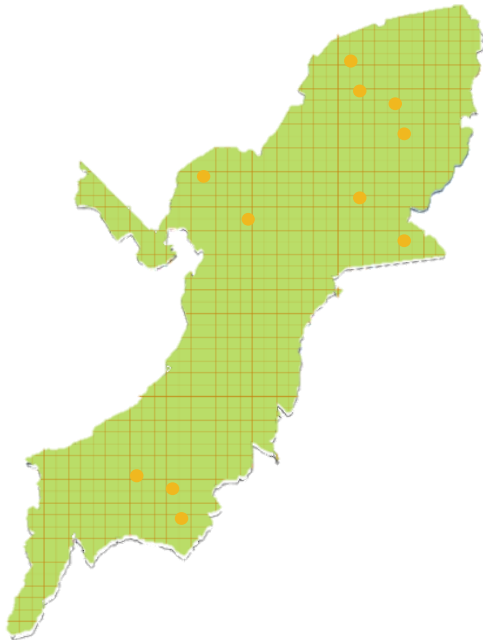
With Uganda Wildlife Authority and Wildlife Conservation Society

Learn Poacher Behavior Model from Real-World Data



Learn Poacher Behavior Model from Real-World Data

Labels

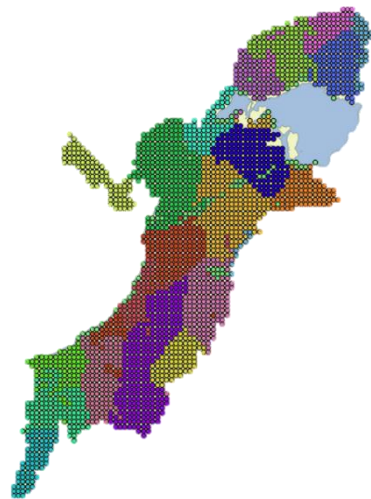
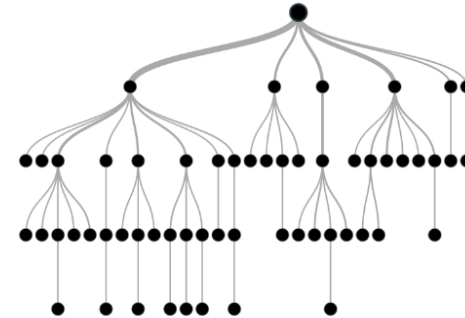
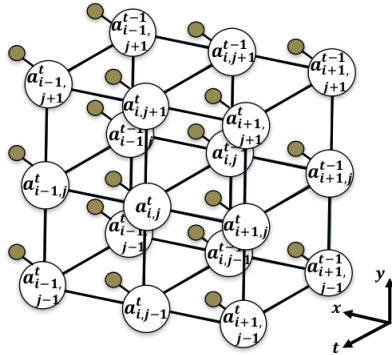


+: Has Poaching

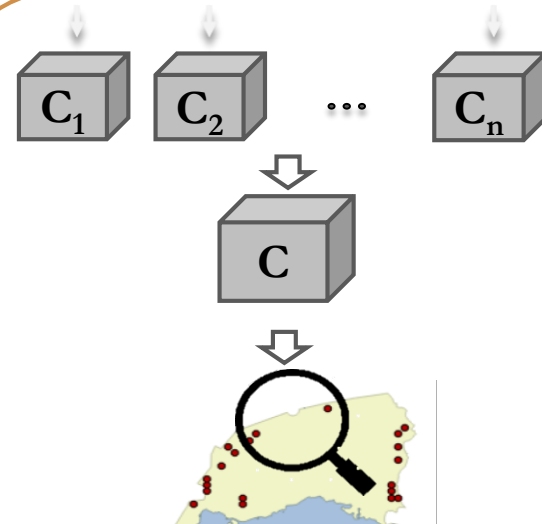
–: No Poaching



Learn Poacher Behavior Model from Real-World Data



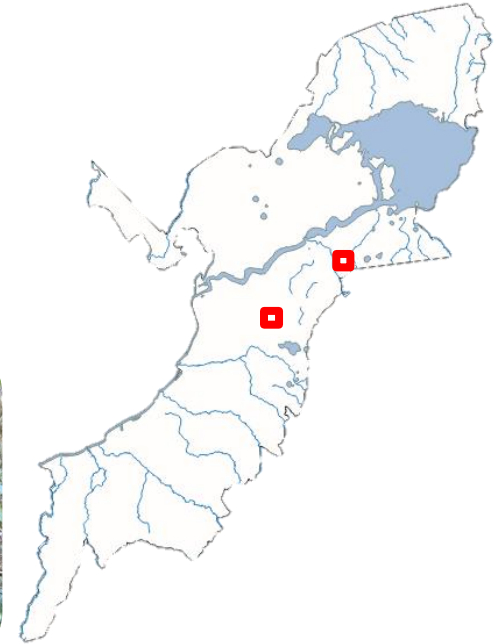
Markov Random Fields



Bagging of Decision Trees

I Month Field Test

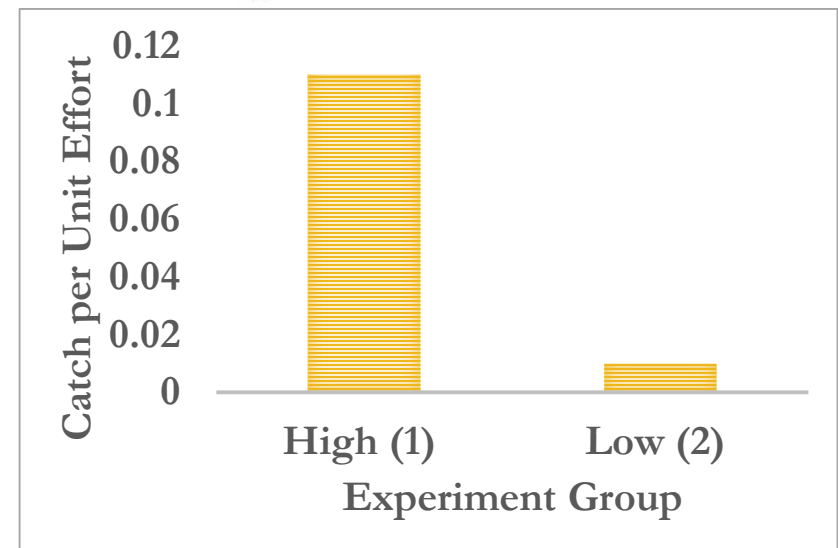
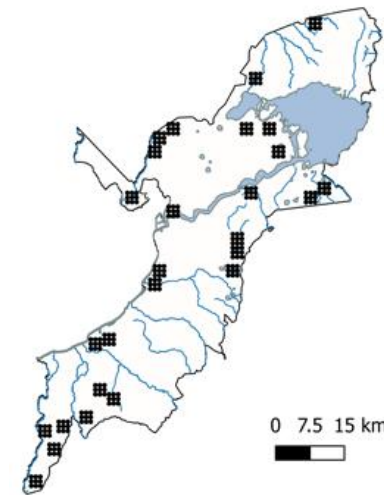
- ▶ Two 9-sq. km areas
 - ▶ Infrequent patrols
 - ▶ Predicted hotspot
- ▶ Findings
 - ▶ 19 litter, ashes, etc.
 - ▶ 1 poached elephant
 - ▶ 1 active snare
 - ▶ 10 antelope snares
 - ▶ 1 roll of elephant snares
- ▶ Snaring hit rates
 - ▶ **Outperform 91% of historical months**



Historical Base Hit Rate	Our Hit Rate
Average: 0.73	3

8 Month Field Test

- ▶ 27 areas, 9-sq km each
- ▶ 2 experiment groups
 - ▶ High: 5 areas
 - ▶ Low: 22 areas
- ▶ 452 km patrolled in total
- ▶ Catch Per Unit Effort (CPUE)
 - ▶ Unit Effort = km walked
 - ▶ Historical CPUE: 0.03
- ▶ Can differentiate H/L threat areas

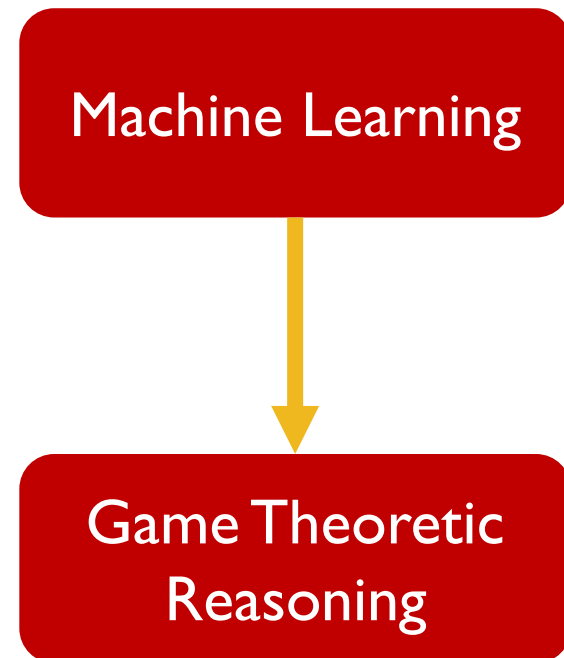


Further Enhancement

- ▶ Augment dataset based on experts' knowledge
- ▶ Field Test in China
 - ▶ Two-day field test in October 2017: 22 snares
 - ▶ 34 patrols from November 2017 to February 2018: 7 snares



Machine Learning + Game Theoretic Reasoning



Game Theoretic Reasoning Based on Learned Model

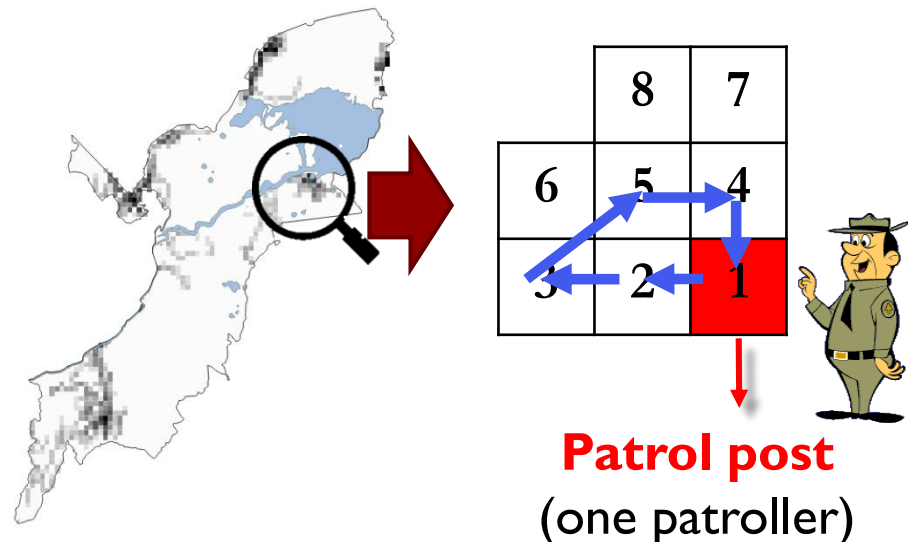
For each cell i : Machine Learning Model

x_i : Current patrol effort at i



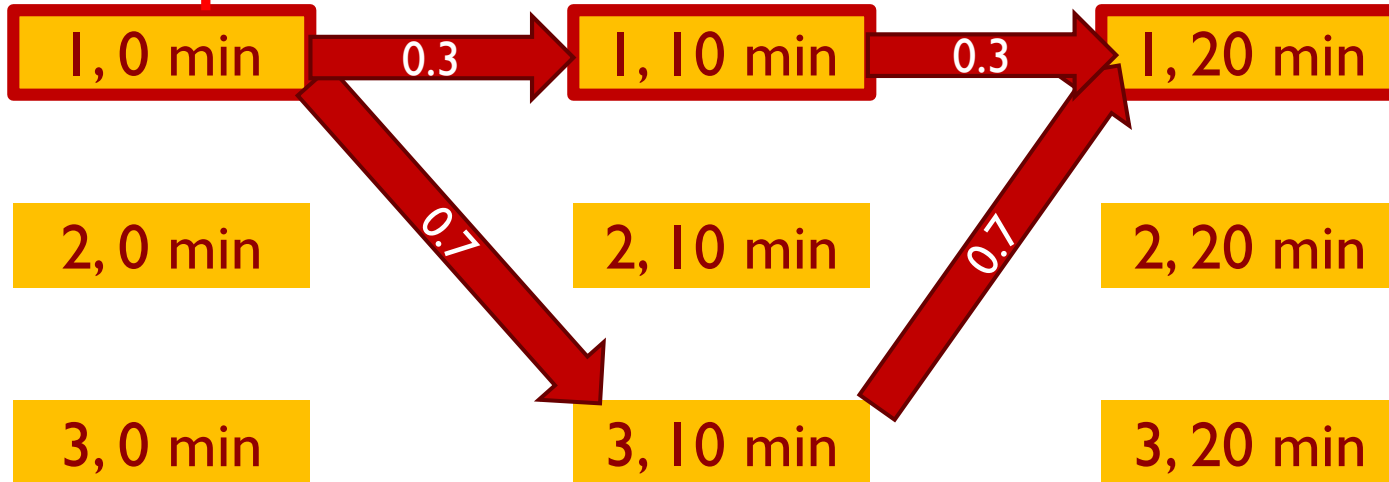
y_i : Prob. of snaring or detection at i in current period

- ▶ Maximize $\sum_i y_i$
- ▶ Scheduling constraint



Game Theoretic Reasoning Based on Learned Model

Patrol post



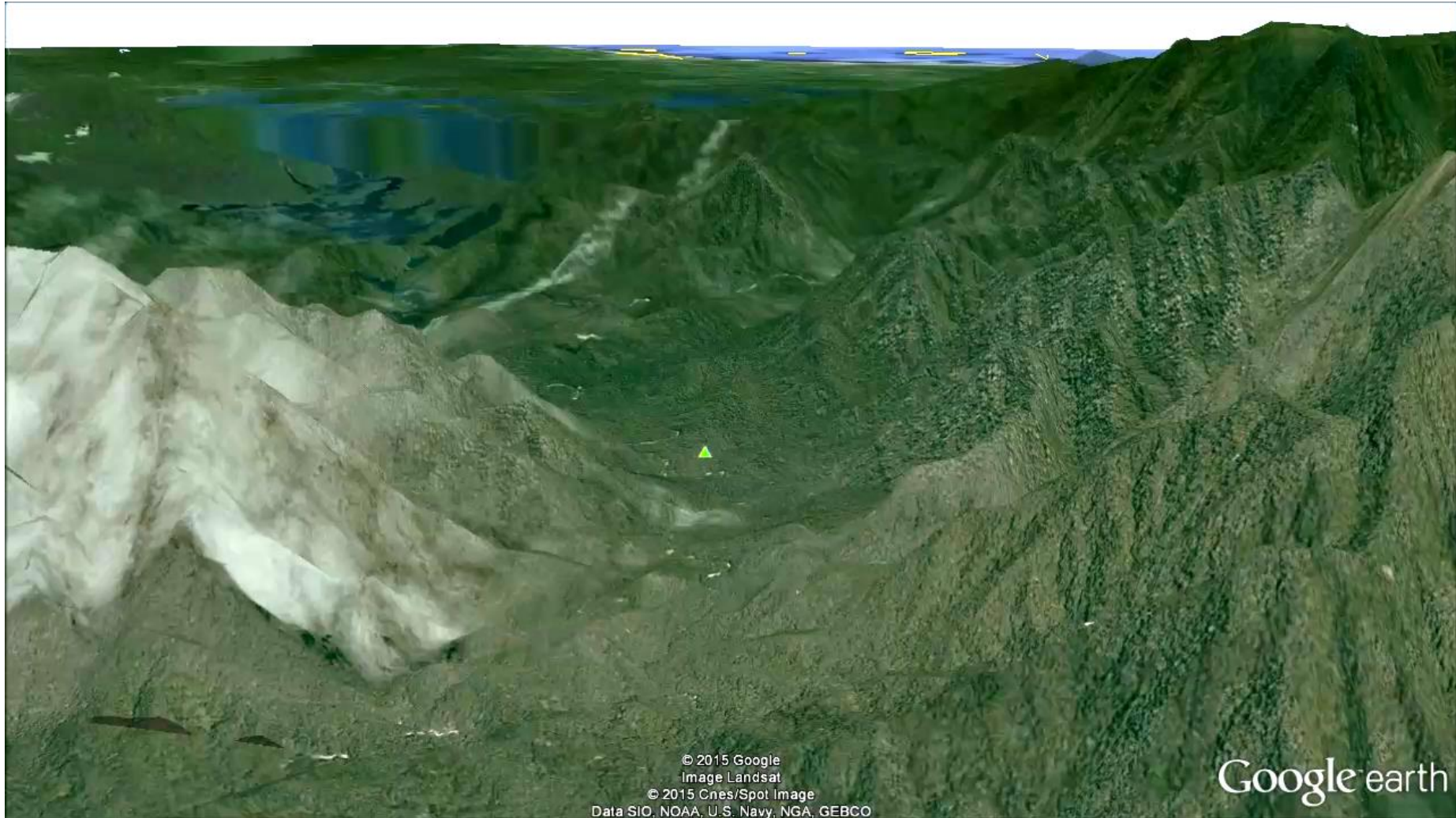
Ranger

$$\begin{aligned} & \max_{x_i} \sum_i g_i(x_i) \\ \text{s.t. } & x_i = \sum_t \sum_{e \in (i,t)} f(e) \\ & f \text{ is a unit flow} \end{aligned}$$

Complex Terrain



Model Terrain to Get Virtual Street Map



30

Deploying PAWS: Field Optimization of the Protection Assistant for Wildlife Security

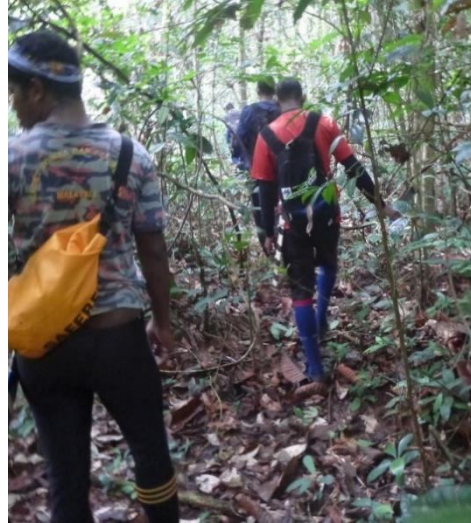
Fei Fang, Thanh H. Nguyen, Rob Pickles, Wai Y. Lam, Gopalasamy R. Clements, Bo An, Amandeep Singh, Milind Tambe, Andrew Lemieux

In IAAI-16: The Twenty-Eighth Annual Conference on Innovative Applications of Artificial Intelligence, February 2016

8/13/2019

Field Test in Malaysia

- ▶ In collaboration with Panthera, Rimba
- ▶ Regular deployment since July 2015 (Malaysia)



Real-World Deployment

Animal Footprint



Tree Mark



Camping Sign



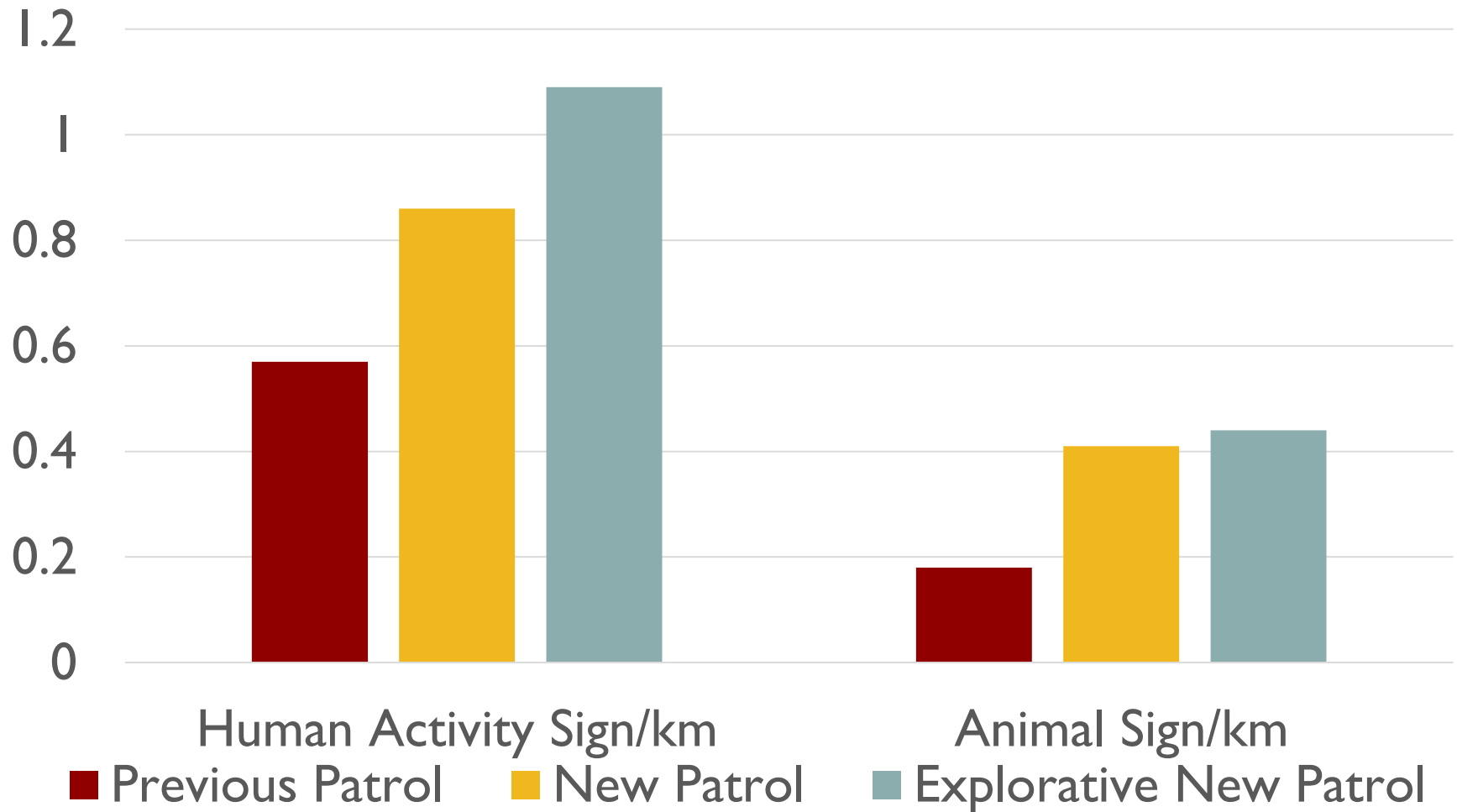
Tiger Sign



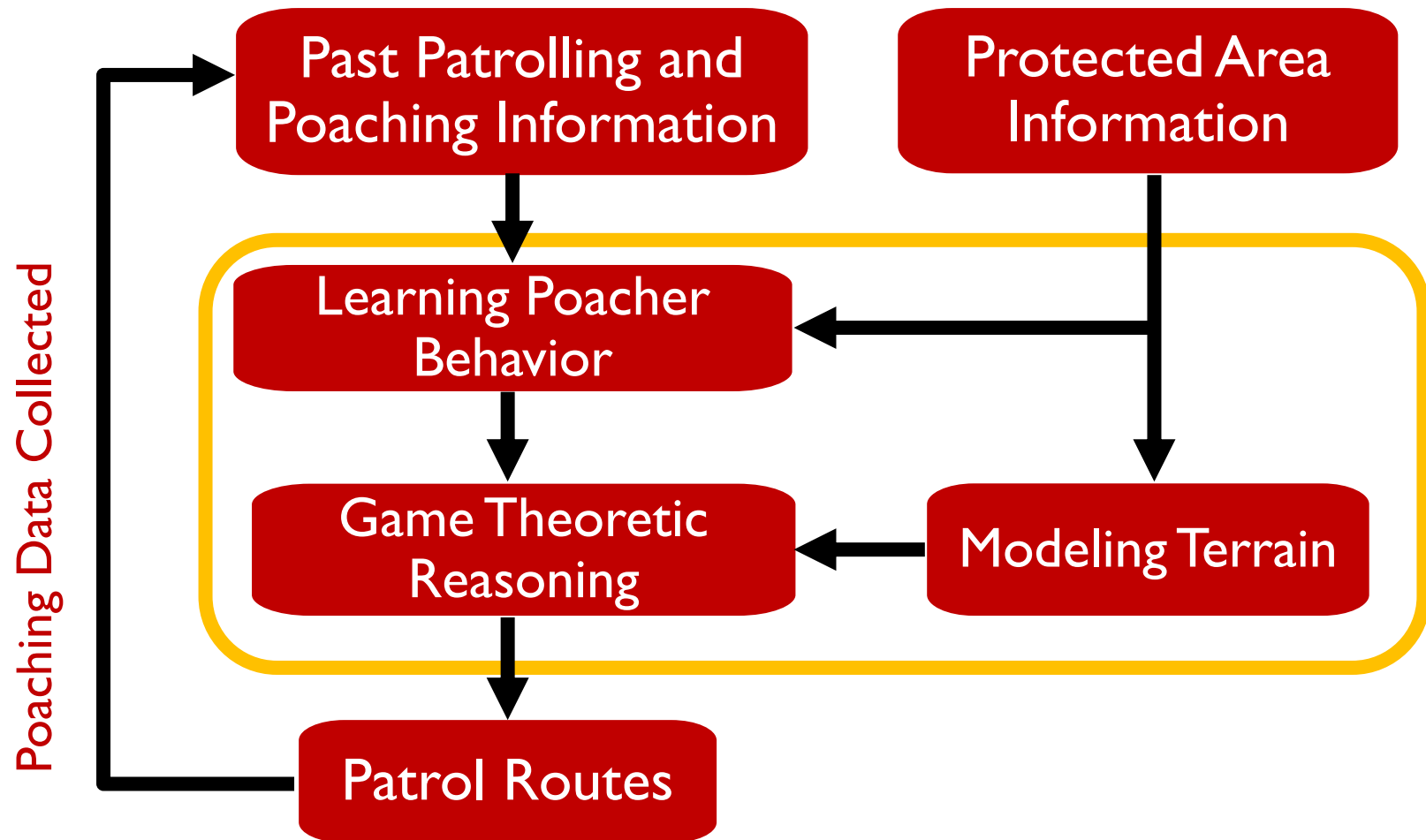
Lighter



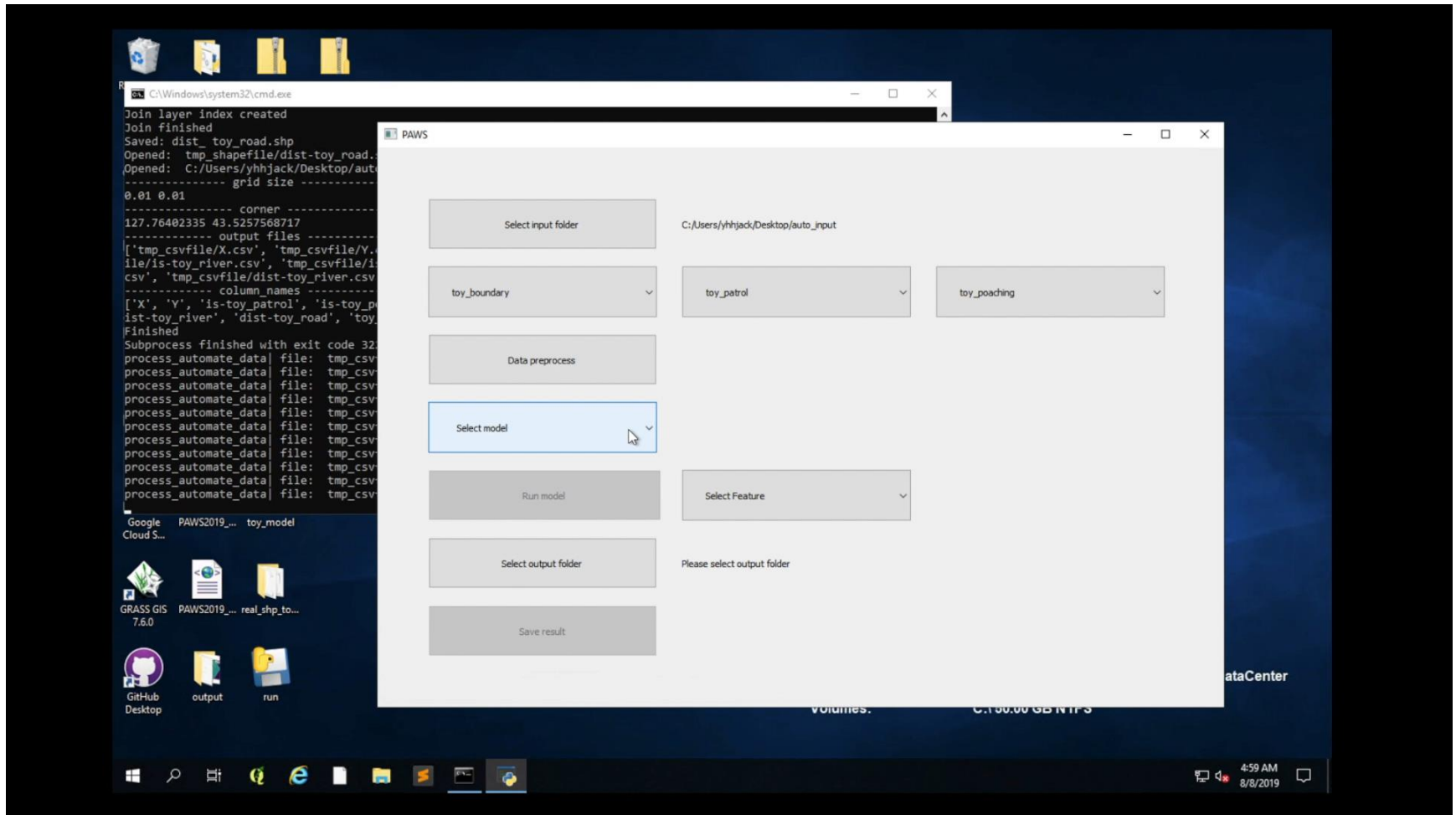
Real-World Deployment



PAWS: Protection Assistant for Wildlife Security



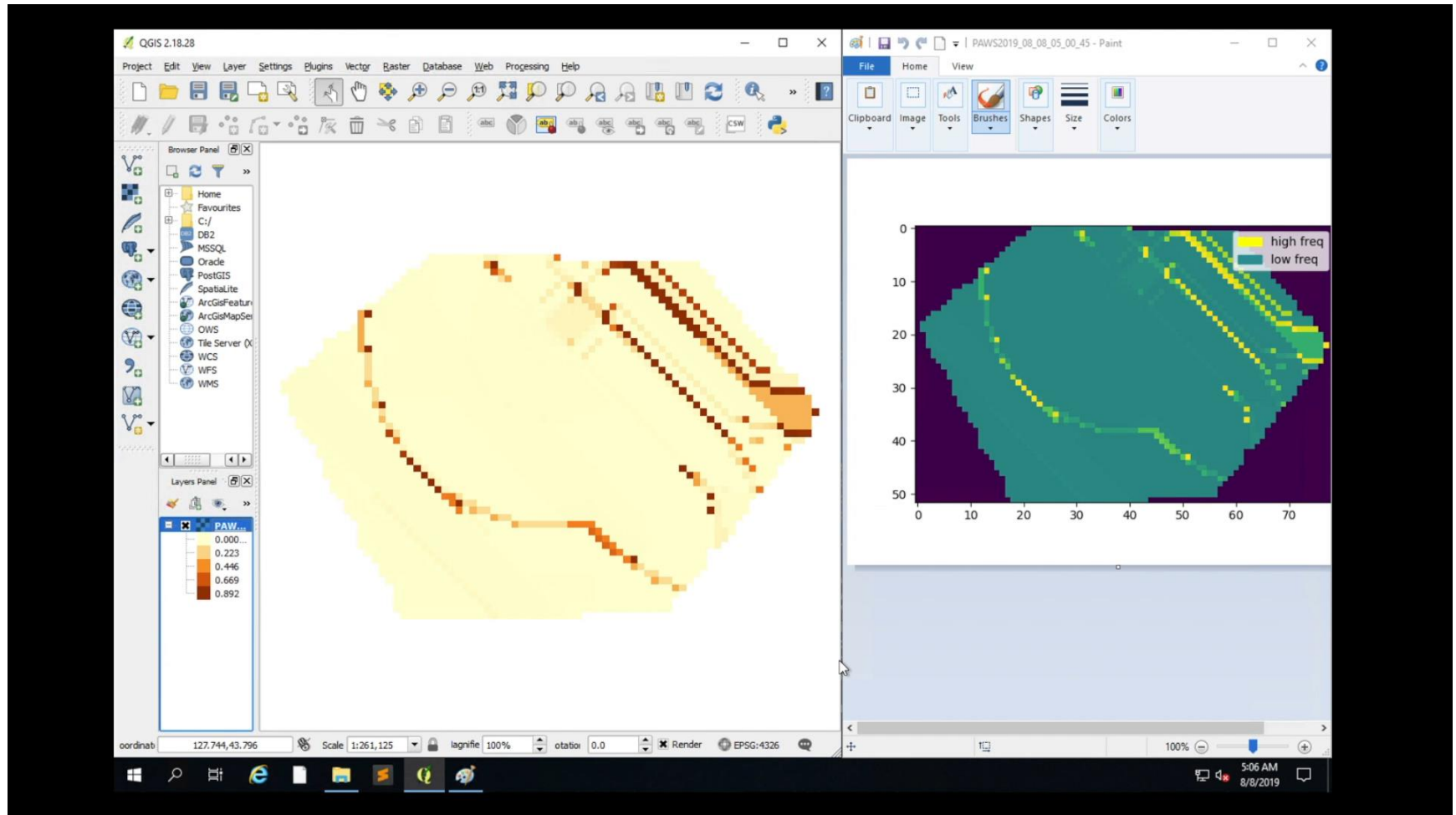
Usable Software Tools



<https://github.com/AlandSocialGoodLab/PAWS>

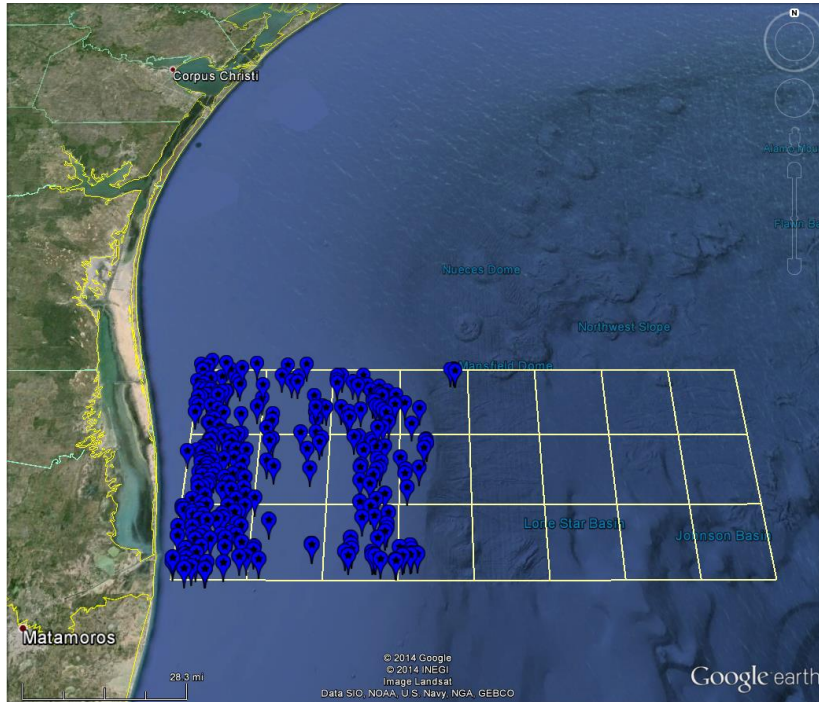


Usable Software Tools

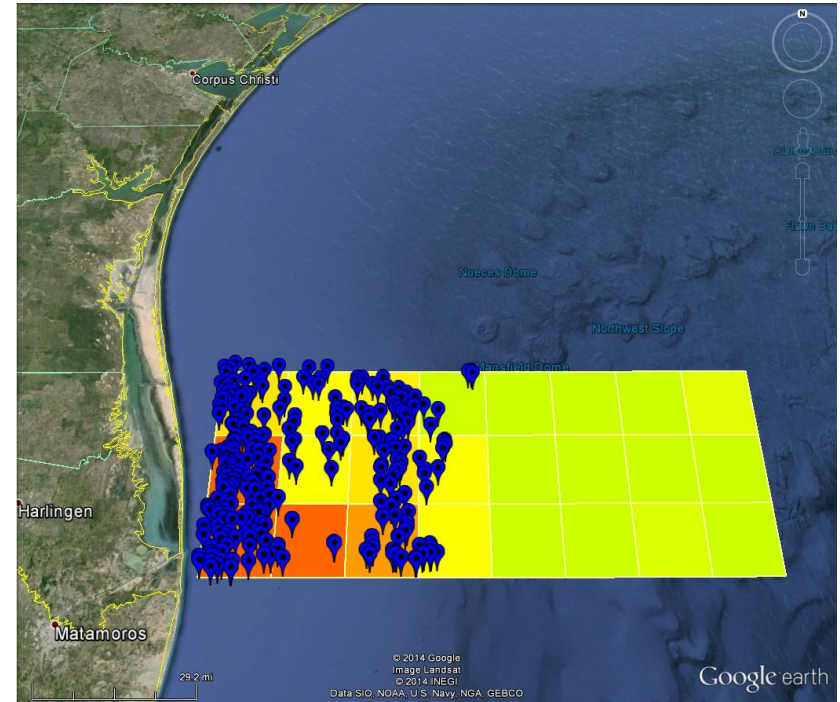


<https://github.com/AlandSocialGoodLab/PAWS>

Other Domains: Reduce Overfishing



Fish Data



Optimal Patrol Strategy

Outline

- ▶ Security games
- ▶ Integrating learning and game theory
 - ▶ Data-based game theoretic reasoning
 - ▶ Learning-powered equilibrium computation
 - ▶ End-to-end learning in games
- ▶ Summary

Patrol with Real-Time Information

- ▶ Sequential interaction
 - ▶ Players make flexible decisions instead of sticking to a plan
 - ▶ Players may leave traces as they take actions
- ▶ Example domain: Wildlife protection



Footprints



Lighters

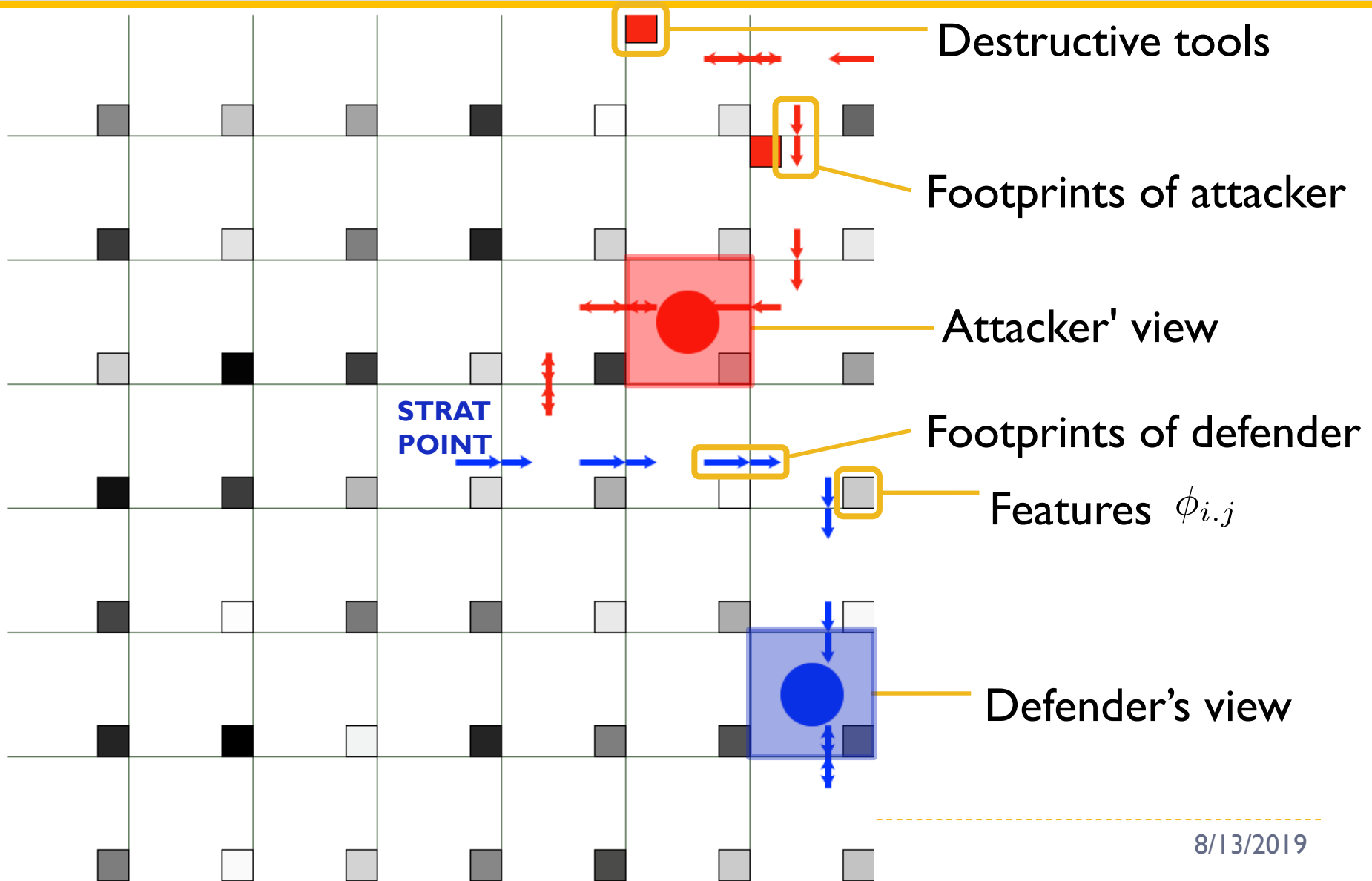


Poacher camp

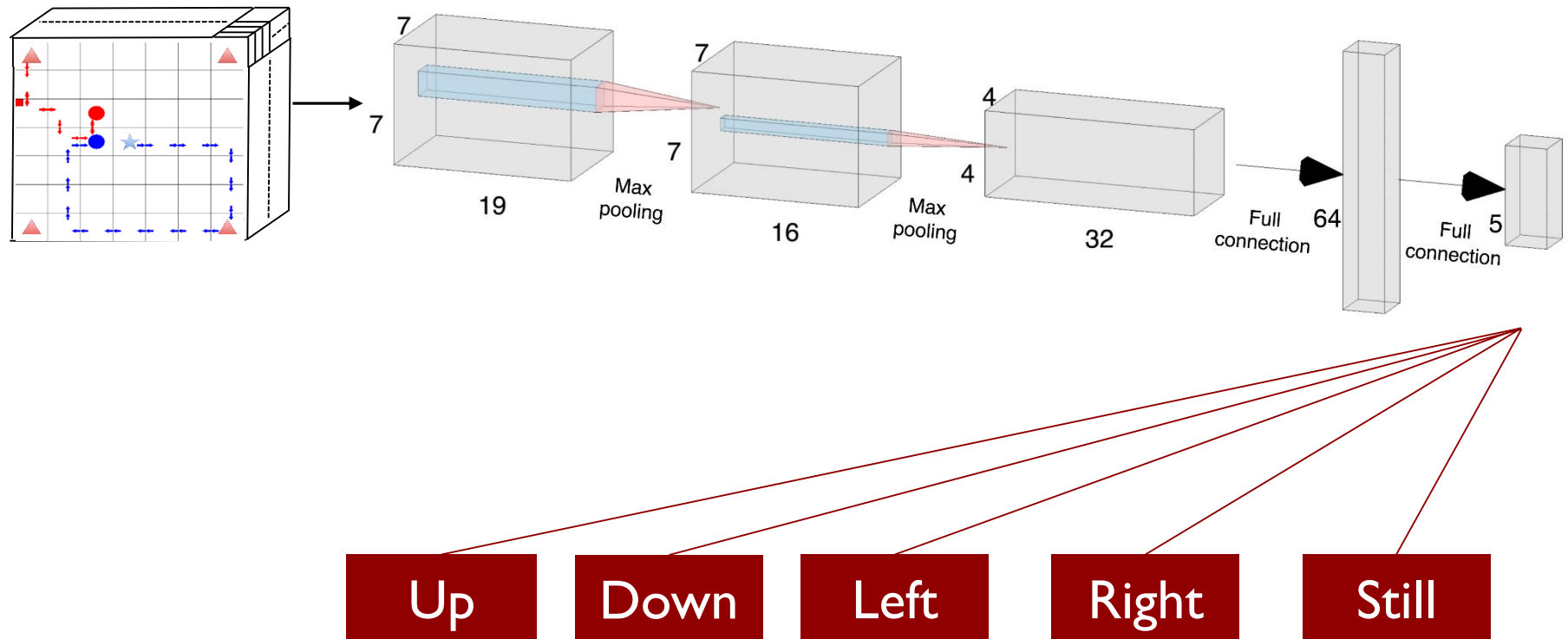


Tree marking

Patrol with Real-Time Information

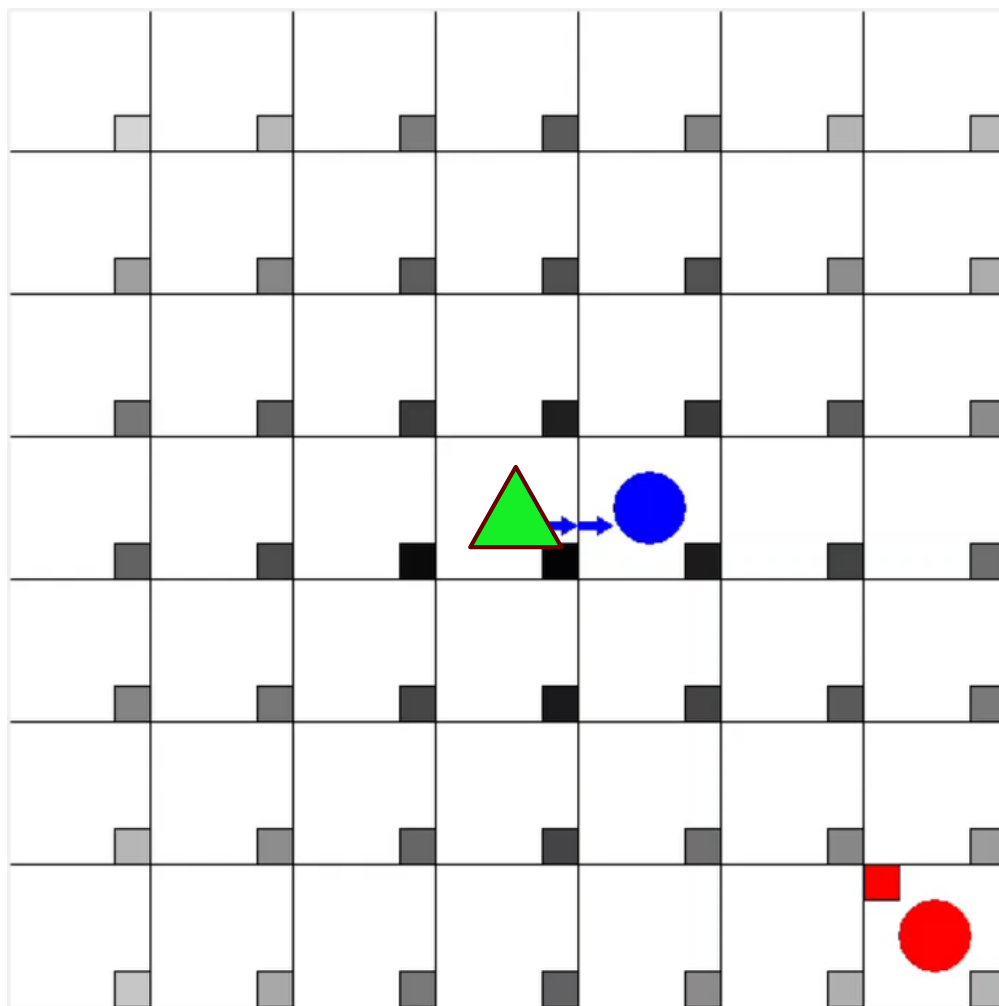


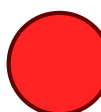
DQN Defender Trained Against Heuristic Attacker



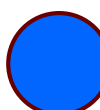
► Q Network: Game state $\rightarrow$ Q-value

DQN Defender Trained Against Heuristic Attacker



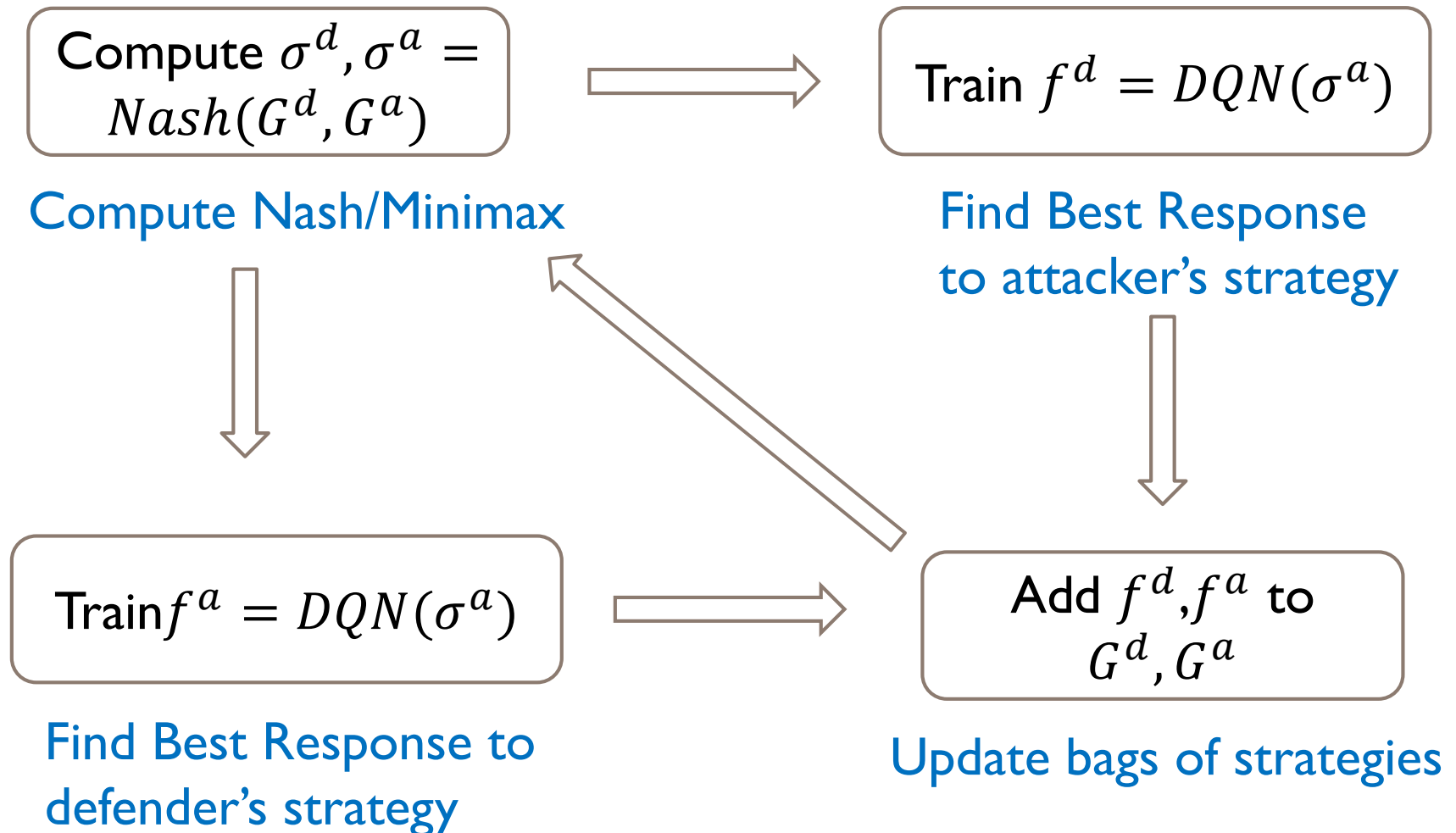
Attacker 

Snares 

Defender 

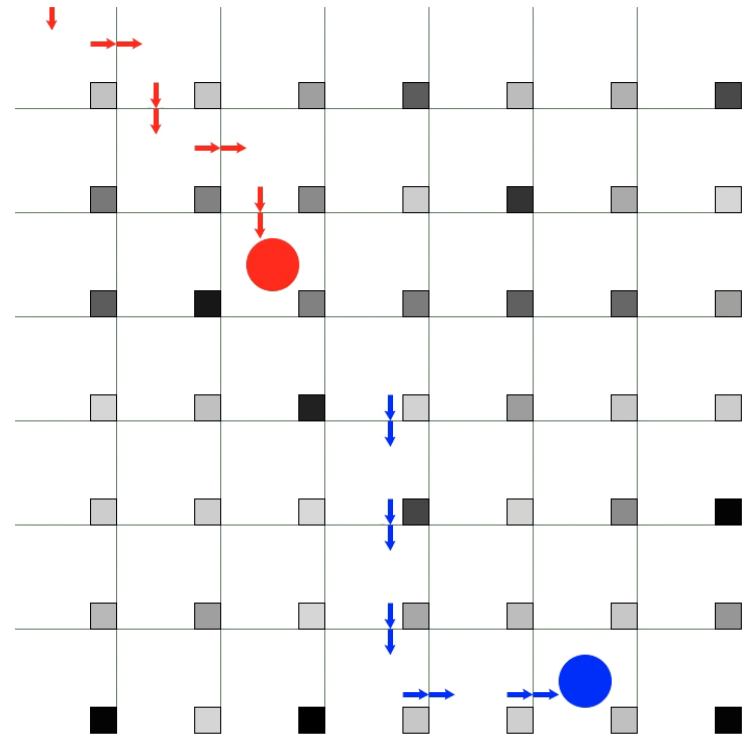
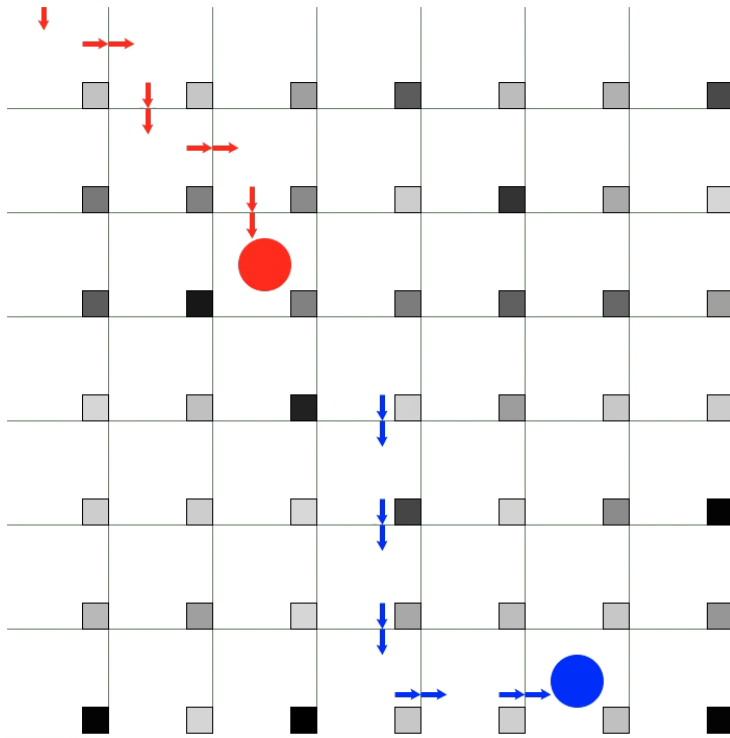
Patrol Post 

Compute Equilibrium: DQN + Double Oracle



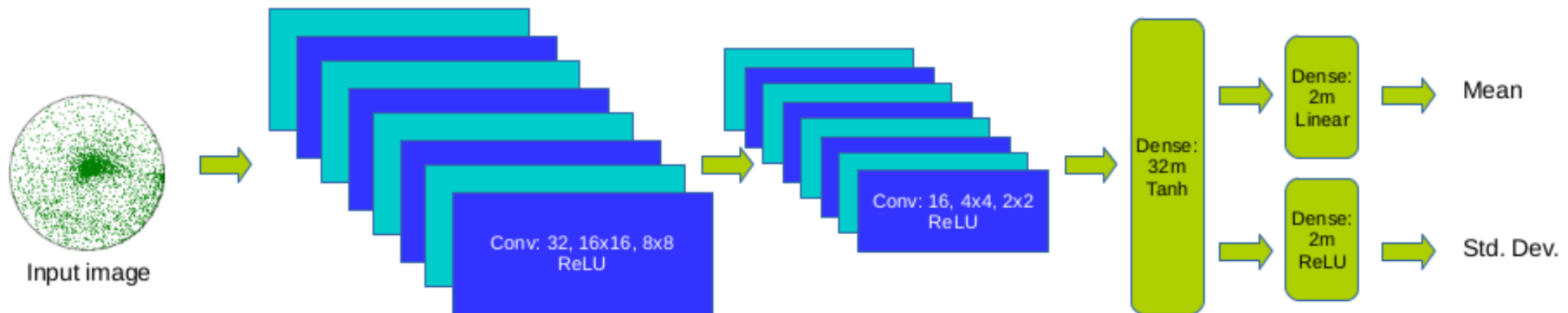
Enhancements

- ▶ Use local modes for efficient and parallized training
- ▶ Start with domain-specific heuristic strategies

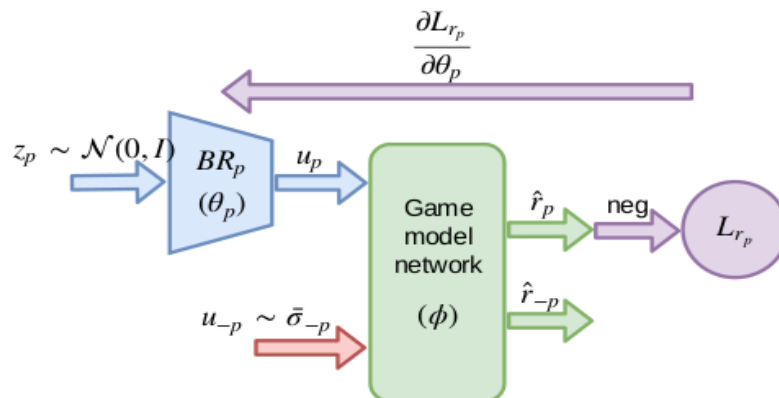


Other Domains: Patrol in Continuous Area

OptGradFP: CNN + Fictitious Play

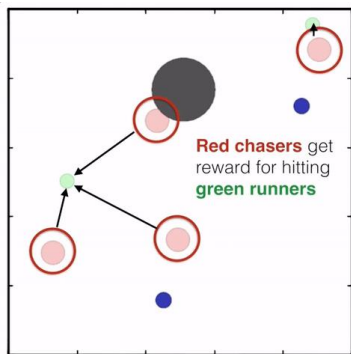


DeepFP: Generative network + Fictitious Play

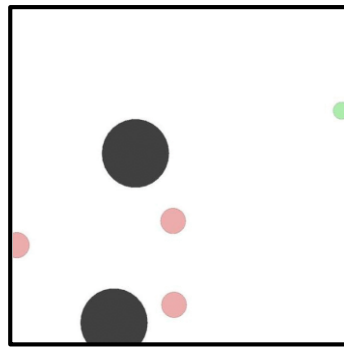


Other Domains: Multiple Agents

Multiple Chasers vs Runners

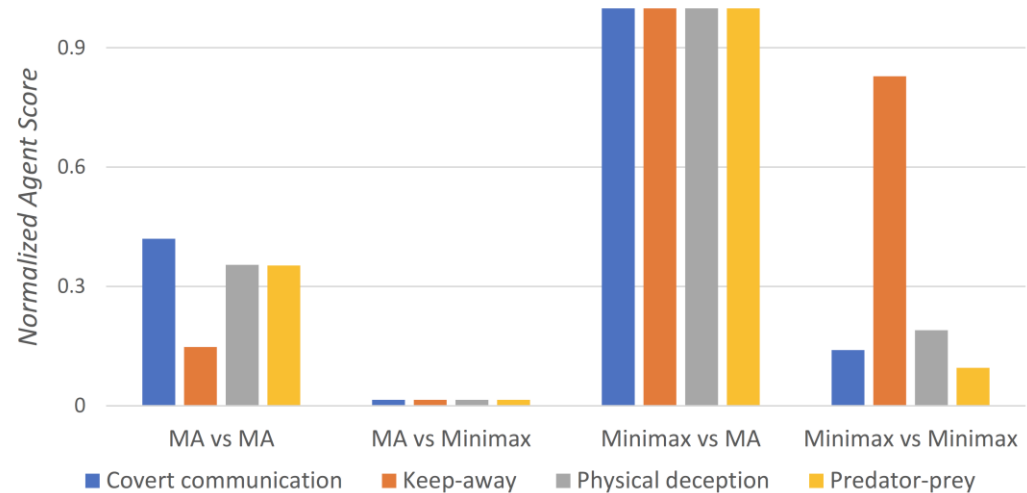


MADDPG



MADDPG (red) vs
M3DDPG (green)

Comparison between MADDPG (MA) and M3DDPG (Minimax)



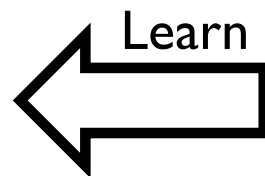
Outline

- ▶ Security games
- ▶ Integrating learning and game theory
 - ▶ Data-based game theoretic reasoning
 - ▶ Learning-powered equilibrium computation
 - ▶ End-to-end learning in games
- ▶ Summary

What game are we/they playing?

- ▶ Common criticism: game parameters are fully known
- ▶ How to learn parameters of **2-player zero sum games** from opponents' or players' actions?

			
	0	b_1	$-b_2$
	$-b_1$	0	b_3
	b_2	$-b_3$	0



i.i.d samples from
equilibrium strategies

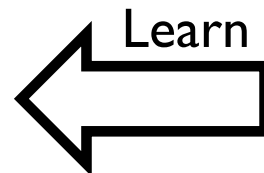
$$a^{(1)} = (\text{rock}, \text{paper})$$

$$a^{(2)} = (\text{rock}, \text{scissors})$$

$$a^{(3)} = (\text{paper}, \text{scissors})$$

Differentiable Learning

			
	0	b_1	$-b_2$
	$-b_1$	0	b_3
	b_2	$-b_3$	0



i.i.d samples from equilibrium strategies

$$a^{(1)} = (\text{rock}, \text{paper})$$

$$a^{(2)} = (\text{rock}, \text{scissors})$$

$$a^{(3)} = (\text{paper}, \text{rock})$$

- ▶ Guess the value of b_i
- ▶ Compute equilibrium of guessed game **Forward pass**
- ▶ Check if the computed equilibrium consistent with data
- ▶ Adjust the value of b_i to increase consistency (update b)
- ▶ Repeat until satisfied

$$b_i := b_i - \frac{\partial L}{\partial b_i} \quad \text{Backward pass}$$

Learning of normal form games

- ▶ QRE = solution of min-max convex-concave problem

$$\begin{aligned} & \min_u \max_v u^T P v - \sum_i v_i \log v_i + \sum_i u_i \log u_i \\ & \text{s.t.} \\ & \quad 1^T u = 1 \\ & \quad 1^T v = 1 \end{aligned}$$

- ▶ KKT conditions:

$$\begin{aligned} P v + \log(u) + 1 + \mu 1 &= 0 \\ P^T u - \log(v) - 1 + \nu 1 &= 0 \\ 1^T u &= 1, \quad 1^T v = 1 \end{aligned}$$

Learning of normal form games

► Forward pass: Apply Newton's Method

$$\begin{bmatrix} \text{diag}(\frac{1}{u}) & P & 1 & 0 \\ P^T & -\text{diag}(\frac{1}{v}) & 0 & 1 \\ 1^T & 0 & 0 & 0 \\ 0 & 1^T & 0 & 0 \end{bmatrix} \begin{bmatrix} \Delta u \\ \Delta v \\ \Delta \mu \\ \Delta \nu \end{bmatrix} = - \begin{bmatrix} Pv + \log(u) + 1 + \mu 1 \\ P^T u - \log(v) - 1 + v 1 \\ 1^T u - 1 \\ 1^T v - 1 \end{bmatrix}$$

► Backward pass: Implicit function theorem

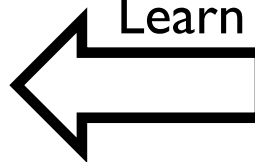
$$\nabla_P L = y_u v^T + u y_v^T,$$

where

$$\begin{bmatrix} y_u \\ y_v \\ y_\mu \\ y_\nu \end{bmatrix} = \begin{bmatrix} \text{diag}(\frac{1}{u}) & P & 1 & 0 \\ P^T & -\text{diag}(\frac{1}{v}) & 0 & 1 \\ 1^T & 0 & 0 & 0 \\ 0 & 1^T & 0 & 0 \end{bmatrix}^{-1} \begin{bmatrix} -\nabla_u L \\ -\nabla_v L \\ 0 \\ 0 \end{bmatrix}$$

Learning in the presence of features

			
	0	$-b_1(x)$	$b_2(x)$
	$b_1(x)$	0	$-b_3(x)$
	$-b_2(x)$	$b_3(x)$	0



i.i.d samples from
equilibrium strategies

$$a^{(1)} = (\text{rock}, \text{paper})$$

$$a^{(2)} = (\text{rock}, \text{scissors})$$

$$a^{(3)} = (\text{paper}, \text{scissors})$$

...

Context

$$x^{(1)} = [0.1, 0.5]$$

$$x^{(2)} = [0.3, 0.7]$$

...

End-to-end learning

Algorithm 1: Learning parameters Φ using SGD

Input: training data $\{(x^{(i)}, a^{(i)})\}$, learning rate η , Φ_{init}

for ep in $\{0, \dots, ep_{\text{max}}\}$ **do**

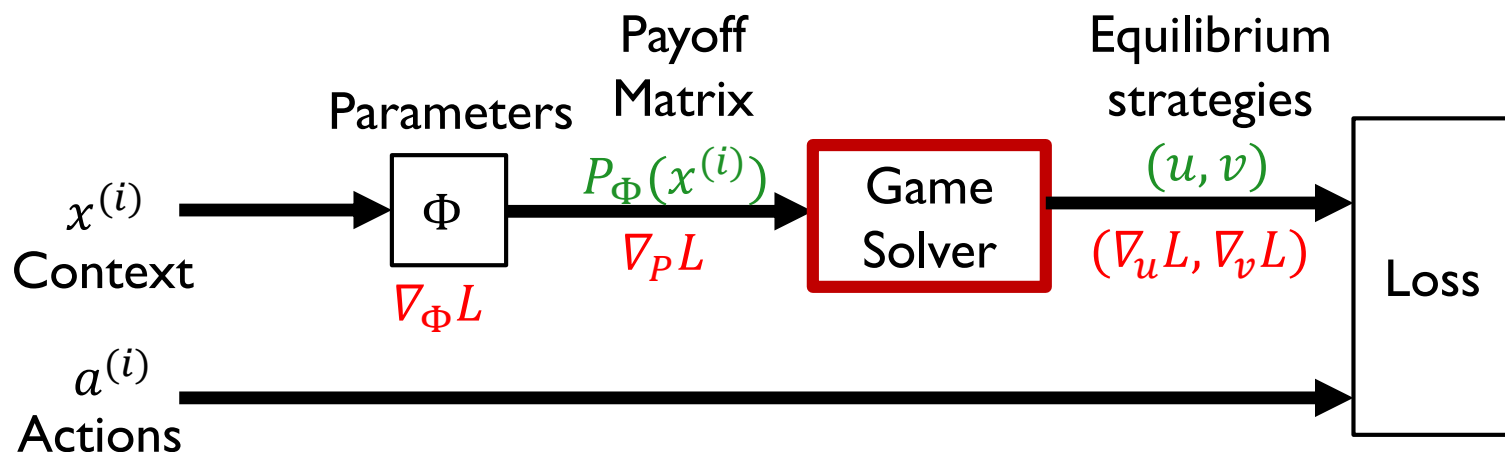
 Sample $(x^{(i)}, a^{(i)})$ from training data;

Forward pass: Compute $P_{\Phi}(x^{(i)})$, QRE (u, v) and loss $L(a^{(i)}, u, v)$;

Backward pass: Compute gradients $\nabla_u L, \nabla_v L, \nabla_P L, \nabla_{\Phi} L$;

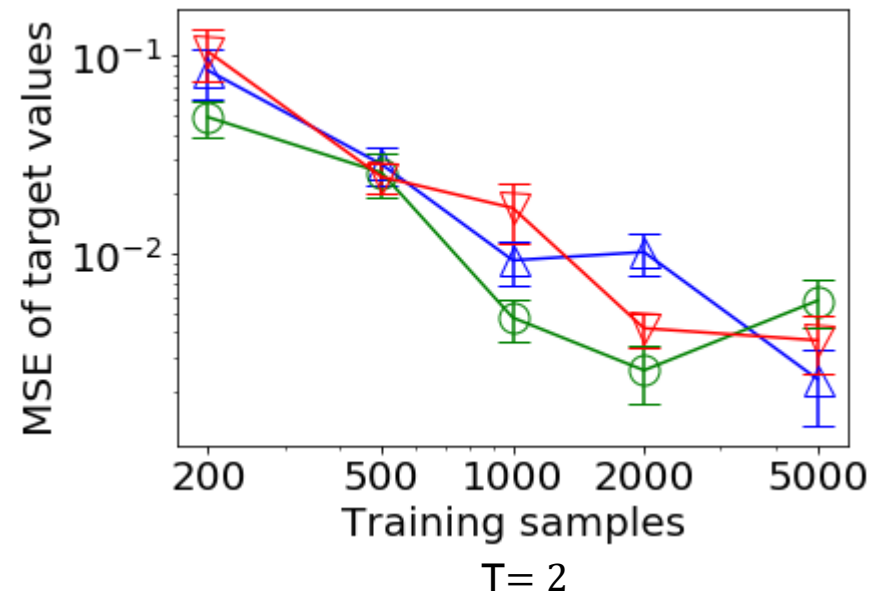
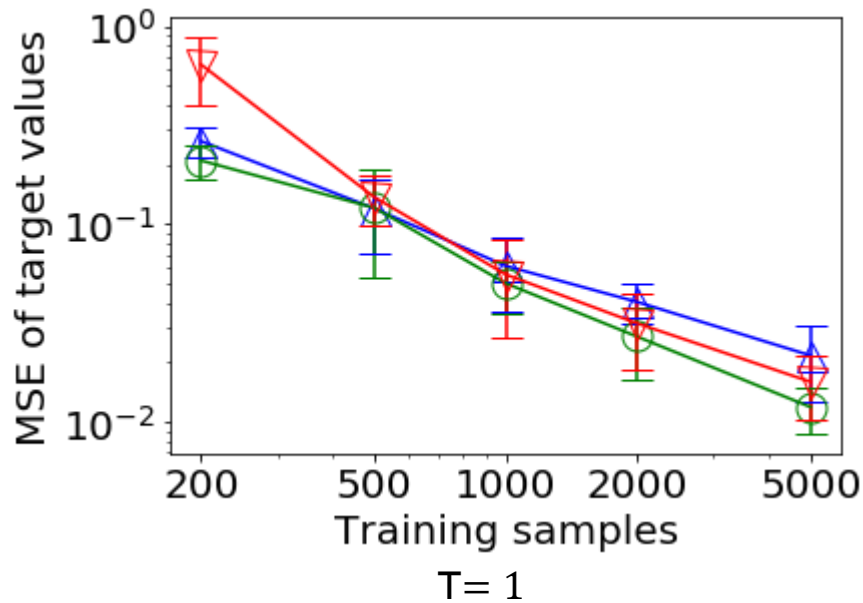
 Update parameters: $\Phi \leftarrow \Phi - \eta \nabla_{\Phi} L$;

end



Resource Allocation Security Game

► $n = 2, r = 5$



Extend to extensive form Games

- Equilibrium is expressed as solution using *dilated entropy regularization*

$$\min_u \max_v u^T P v - \sum_i \sum_a v_a \log\left(\frac{v_a}{v_{p_i}}\right) + \sum_i \sum_a u_a \log\left(\frac{u_a}{u_{p_i}}\right)$$
$$Eu = e, Fv = f$$

$$\nabla_P L = y_u v^T + u y_v^T,$$

where

$$\begin{bmatrix} y_u \\ y_v \\ y_\mu \\ y_\nu \end{bmatrix} = \begin{bmatrix} -\Xi(u) & P & E^T & 0 \\ P^T & \Xi(v) & 0 & F^T \\ E & 0 & 0 & 0 \\ 0 & F & 0 & 0 \end{bmatrix}^{-1} \begin{bmatrix} -\nabla_u L \\ -\nabla_v L \\ 0 \\ 0 \end{bmatrix}$$

Improve Scalability using FOM

- Both the problem in the forward pass and backward pass can be converted into the following format

$$\min_{Ex=x_0} \max_{Fy=y_0} x^T P y + \boxed{\mathcal{E}(x) - \mathcal{F}(y)}$$

strictly convex functions

- This problem can be solved using various first-order iterative methods

Input: $x^{(0)}, y^{(0)}, P, \tau, \sigma$

for i **in** $\{0, \dots\}$ **do**

$\tilde{y} = y^{(i)};$

$x^{(i+1)} = \boxed{\text{BR}_x(x^{(i)}, \tilde{y}; P, \tau)};$

$\tilde{x} = 2x^{(i+1)} - x^{(i)};$

$y^{(i+1)} = \boxed{\text{BR}_y(y^{(i)}, \tilde{x}; P, \sigma)};$

end

BR is smoothed best response

$$\text{BR}_x(\bar{x}, \tilde{y}) = \arg \min_{Ex=x_0} x^T P \tilde{y} + \mathcal{E}(x) + \frac{1}{\tau} D_x(x, \bar{x})$$

$$\text{BR}_y(\bar{y}, \tilde{x}) = \arg \min_{Fy=y_0} -\tilde{x}^T P y + \mathcal{F}(y) + \frac{1}{\sigma} D_y(y, \bar{y}).$$



Takeaway I: AI Has Great Potential for Social Good

Artificial
Intelligence

Machine Learning /
Reinforcement Learning

Computational
Game Theory

Societal Challenges

Security & Safety



Environmental Sustainability

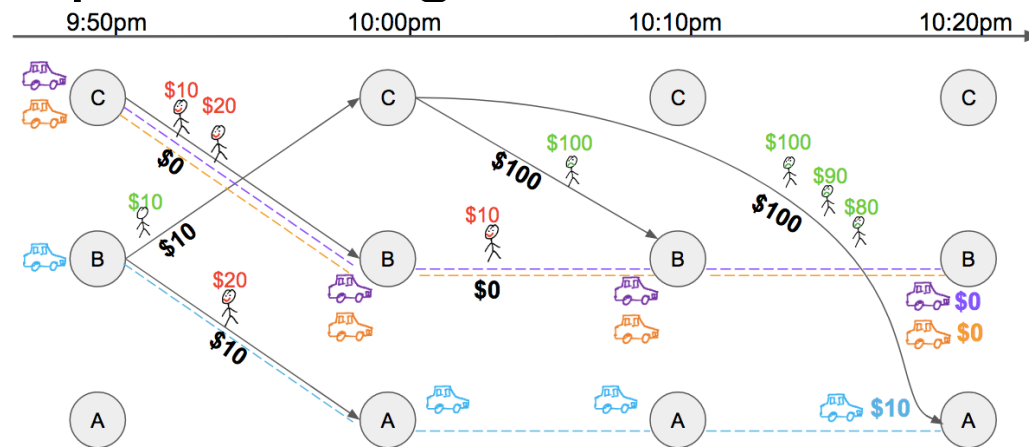


Mobility



Design Dispatching and Pricing Scheme in Ridesharing

► Spatial-Temporal Pricing

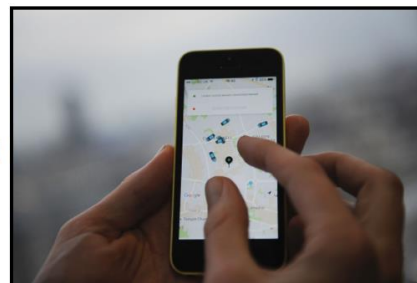


► Two-Stage Dispatching

Schedule a Ride

Tue, Aug 7
6:20 AM - 6:30 AM

Day	Time	Price
Sun Aug 5	4:10	10
Mon Aug 6	5:15	15
Today	6:20 AM	20
Wed Aug 8	7:25 PM	25
Thu Aug 9	8:30	30
Fri Aug 10	9:35	35



**Stage 1:
Scheduled
Requests**

**Stage 2:
Real-Time
Requests**

**Successful
Dispatching**

Takeaway 2: Ways to Integrate Learning and Game Theory

- ▶ Data-based game theoretic reasoning
- ▶ Learning-powered equilibrium computation
- ▶ End-to-end learning in games
- ▶ More?!

Acknowledgment

► Advisors, students and all co-authors!



► Collaborators and partners



► Funding support

