

Passenger Spoofing Attack for Artificial Intelligence-based Mobility-as-a-Service

Kai-Fung Chu and Weisi Guo

Abstract—Mobility-as-a-Service (MaaS), a new mobility service model that integrates multiple mobility providers, relies on many data processing technologies to manage multi-modal transport. Artificial Intelligence (AI) is one of the technologies to improve the services matching to passengers based on their implicit experience and preference. However, incorporating AI into MaaS may also introduce loopholes to the system. One may use the loophole in the heterogeneity of passenger experience and preference by falsifying data to prioritize their journey, which jeopardizes the trustworthiness of MaaS. In this paper, we investigate the cyber security risks in MaaS, focusing on the spoofing attack in which malicious passengers are prioritized by falsifying data to gain an advantage in journey planning. The spoofing attack is based on reinforcement learning that learns to reduce passenger satisfaction about the MaaS and its profit by requesting travel with falsifying passenger states. We conduct experiments based on New York City dataset to evaluate the spoofing attack. The experiment results indicate that the attack can reduce about 70% of the profit. By investigating the cyber security risks in MaaS, we could enhance the knowledge and understanding of the risks for building a secure and trustworthy MaaS.

I. INTRODUCTION

Mobility-as-a-Service (MaaS) is a new mobility service model that integrates multiple mobility providers [1]. Passengers can enjoy the unique services and advantages of MaaS such as multi-modal journey planning and reservation across different mobility providers in multi-modal transportation. The integration of mobility providers could also lead to social benefits such as traffic, air pollution, and energy consumption reduction [2]. One of the important features of MaaS is the multi-modal journey planning as the planning over multi-modal transportation is much more complex than that of a single consistent one. In addition, whether the recommended journey matches with the passenger's expectations could affect their impression and retention. Therefore, it is important to improve the journey planning system so that passengers are satisfied with the plan. One improvement way is to incorporate implicit passenger experience and preferences into the decision-making process using artificial intelligence (AI). AI is one of the popular technologies in the field of intelligent transportation systems due to its success as in [3], [4] and [5]. With its extraordinary modeling capability, we expected that the AI agent could handle the heterogeneity of

passengers and recommend the best journey unique to each passenger rather than Pareto-front solutions that are the same for every passenger with the same origin and destination.

While the AI-based MaaS is promising, there are cyber security risks in MaaS and AI, and those risks could be inherited by the AI-based MaaS. For the risks in MaaS, Callegati et al. [6] summarized the insider threats of each component in the MaaS architecture. Insiders such as developers, service administrators, managers, etc., may perform various attacks on MaaS. Those insider attacks may cause serious consequences such as information leakage, operational failure, and economic loss. Another example of attacks on MaaS is the Denial-of-Service (DoS) attack studied by Thai et al [7]. The authors identified that attackers could disrupt the system to maximize passenger loss by maliciously controlling a fraction of vehicles in the system. As a result, the DoS attacks cost more than \$ 15 US dollars per unit for protection according to [7]. On the other hand, the risks of AI were well discovered by the cyber-security community. The opaqueness of AI, also known as a black-box method, continues to pose challenges in ensuring the integrity of the decision-making process [8]. The lack of detailed information on the causality of AI decision-making hinders the ability to fully explain the reasoning behind the decisions made. Consequently, it becomes difficult to provide guarantees that these decisions are not influenced by corrupted AI-focused attacks. AI could be attacked by data poisoning attacks [9] that corrupt the training data to cause the AI to produce desirable outcomes by the attacker, evasion attacks [10] that manipulate the input data to produce an error output, and inference attacks [11] that gain knowledge about the database used to train the AI. The erosion of people's confidence in novel systems and technologies is anticipated due to incidents involving security vulnerabilities. Therefore, it is imperative to reassess the security of MaaS systems, particularly those employing AI technology.

In this paper, we explore the risks to the AI-based MaaS operators and journey planning problem and identify a new attack, named passenger spoofing attack (PSA), which leverages the diversity of passengers and the heterogeneous service planning in AI to prioritize malicious agents in the journey planning process. The PSA generates falsifying profiles and satisfaction by reinforcement learning model to be used for service queries based on the current state of other passengers. We conduct experiments on two simulated scenarios based on the real-world New York City dataset to investigate the impact of the attack on passenger satisfaction and MaaS profit, as well as the sensitivity of malicious agent

This work was supported by EPSRC MACRO - Mobility as a service: MAnaging Cybersecurity Risks across Consumers, Organisations and Sectors (EP/V039164/1)

Kai-Fung Chu and Weisi Guo are with the School of Aerospace, Transport and Manufacturing, Cranfield University, Bedford, MK43 0AL, UK kai.fung.chu@cranfield.ac.uk; weisi.guo@cranfield.ac.uk

spatial distance. The results of the experiments can be used to develop corresponding detection and defense mechanisms.

The rest of this paper is organized as follows. Section II illustrates the background MaaS model and journey planning problem. Section III presents the reinforcement learning-based PSA. Section IV simulates the attacks in the MaaS scenario and analyzes the attack influence on the system. Finally, the paper is concluded in Section V.

II. MAAS JOURNEY PLANNING

In this section, we introduce the background models including the MaaS and journey planning problem.

A. MaaS Model

MaaS is a cyber-physical transportation system that integrates multiple mobility providers offering mobility services for the same or different routes. The coordination of mobility services of different providers is managed by a MaaS coordinator, who facilitates passengers in selecting, reserving, and paying for combined journeys from their origin to destination. We model the multimodal transport network as a directed graph $G(\mathcal{N}, \mathcal{A})$, where $\mathcal{N}$ and $\mathcal{A}$ represent the sets of nodes and links in the network, respectively. Let $\mathcal{F}$ denote the set of mobility providers. Each mobility provider $f \in \mathcal{F}$ offers mobility services on a sub-network $\mathcal{A}_f \in \mathcal{A}$. The utility costs associated with each mobility service from node i to node j offered by f are represented by parameters β_{ij}^f , δ_{ij}^f , and ρ_{ij}^f , which represent time¹, discomfort, and profit, respectively. The MaaS coordinator's responsibility is to determine an optimal journey fulfilling the transport request of passenger k , from origin o^k to destination d^k , based on the passenger preferences, experiences, and memories, while considering the mobility services of multiple mobility providers for multiple passengers.

B. Journey Planning Problem

This section introduces the multi-modal journey planning problem. To facilitate the formulation of the problem, we define a binary variable x_{ij}^{kf} that represents the coordinator's decision of the journey, which is commonly used in journey planning [13]. Specifically, x_{ij}^{kf} is a binary variable that takes the value of 1 if link (i, j) operated by f offers mobility service to passenger k , and 0 otherwise. Formally, we have:

$$x_{ij}^{kf} = \begin{cases} 1 & \text{if link } (i, j) \text{ operated by } f \text{ offers to } k, \\ 0 & \text{otherwise.} \end{cases} \quad (1)$$

The objective function of the problem is to minimize the total utility cost of passengers, which is formulated as:

$$\sum_{(i,j) \in \mathcal{A}_f, k \in \mathcal{K}, f \in \mathcal{F}} (w_\beta^k \beta_{ij}^f + w_\delta^k \delta_{ij}^f + w_\rho^k \rho_{ij}^f) x_{ij}^{kf}, \quad (2)$$

where w_β^k , w_δ^k , and w_ρ^k are the weighting of the corresponding utility terms.

¹In this study, we use a single parameter to represent all time-related parameters, as a pilot study and for simplicity as in [12]. In a real-world application, travel time may include waiting, in-vehicle, and transfer time.

We also define $\mathcal{N}^+(i)$ and $\mathcal{N}^-(i)$ as the sets of incoming and outgoing locations of i , respectively, such that $\mathcal{N}^+(i) = \{j \in \mathcal{N} | (j, i) \in \mathcal{A}_f\}$ and $\mathcal{N}^-(i) = \{j \in \mathcal{N} | (i, j) \in \mathcal{A}_f\}$. The offered journey in the transport network has to be on a connected path, which can be ensured by the flow conservation equation:

$$\sum_{j \in \mathcal{N}^-(i)} x_{ij}^{kf} - \sum_{j \in \mathcal{N}^+(i)} x_{ji}^{kf} = \begin{cases} 1 & \text{if } i = o^k, \\ -1 & \text{if } i = d^k, \\ 0 & \text{otherwise,} \end{cases} \quad \forall i \in \mathcal{N}, k \in \mathcal{K}, f \in \mathcal{F}. \quad (3)$$

The total number of services offered is limited by capacity, so we restrict the total number of passengers for a service using the following equation:

$$\sum_{k \in \mathcal{K}} x_{ij}^{kf} \leq C_{ij}^f, \quad \forall (i, j) \in \mathcal{A}_f, f \in \mathcal{F} \quad (4)$$

where C_{ij}^f is the maximum capacity of mobility service from i to j that operated by $f \in \mathcal{F}$.

As one may notice, the weights in Eq. 2 may affect the solution. To enhance passenger satisfaction and increase the profit of Mobility-as-a-Service (MaaS) providers in the journey planning problem, it is important to determine a suitable set of weights based on individual passenger preferences, experiences, and memories. To incorporate the passenger experience and memories, a 4-tuple Markov decision process (MDP) [14], $\langle \mathcal{S}, \mathcal{A}, P, R \rangle$, can be utilized to model the passenger retention process, where $\mathcal{S}$, $\mathcal{A}$, P , and R are the sets of states and actions, state transition and reward functions, respectively. For each travel time t , the state $s_t^k \in \mathcal{S}$ represents passenger satisfaction and profiles. The action $a_t^k \in \mathcal{A}$ is the weighting of utility terms, where $a_t^k := [w_\beta^k; w_\delta^k; w_\rho^k]$ for time t and passenger k . $P(s_{t+1}^k | s_t^k, a_t^k)$ describes the satisfaction transition from states $s_t^k \in \mathcal{S}$ to $s_{t+1}^k \in \mathcal{S}$ with action $a_t^k \in \mathcal{A}$. $R(s_t^k, a_t^k, s_{t+1}^k)$ can be employed to model the profit obtained from transiting from s_t^k to s_{t+1}^k by taking action a_t^k and is given by:

$$R(s_t^k, a_t^k, s_{t+1}^k) = \sum_{f \in \mathcal{F}} \rho_{ij}^{kf}. \quad (5)$$

we introduce a passenger satisfaction model that measures satisfaction level with an N -level integer value proportional to the retention rate. Let H^k be the satisfaction level of passenger k . To capture the relationship between satisfaction and the quality of the offered journey, we model the variation of satisfaction as a function of the difference between the expected and actual journey. Specifically, we define the satisfaction level variation as follows:

$$H^k := \begin{cases} H^k + n & \text{if } E^k \geq \bar{E}^k, \\ H^k - n & \text{if } E^k \leq \underline{E}^k, \\ H^k & \text{otherwise,} \end{cases} \quad \forall k \in \mathcal{K}, \quad (6)$$

where $\bar{E}^k$ and $\underline{E}^k$ represent the upper and lower thresholds of expectation difference, respectively. Here, n is the satisfaction level step size, and the expectation difference E^k is

defined as:

$$E^k = \tilde{w}_\beta^k(\tilde{\beta}_{o^k d^k}^k - \beta_{o^k d^k}^k) + \tilde{w}_\delta^k(\tilde{\delta}_{o^k d^k}^k - \delta_{o^k d^k}^k) + \tilde{w}_\rho^k(\tilde{\rho}_{o^k d^k}^k - \rho_{o^k d^k}^k), \quad \forall k \in \mathcal{K}. \quad (7)$$

where $\tilde{w}_\beta^k$, $\tilde{w}_\delta^k$, and $\tilde{w}_\rho^k$ are the corresponding utility weighting of passenger k , and $\tilde{\beta}_{o^k d^k}^k$, $\tilde{\delta}_{o^k d^k}^k$, $\tilde{\rho}_{o^k d^k}^k$ are the utility expected implicitly by the passenger. Note that the expected utility represents the utility of the best journey that the passenger can get without capacity constraint, which can be determined by solving Eqs. (2) and (3) only. The actual utility represents the utility of the journey planned by MaaS.

The multi-modal journey planning and passenger satisfaction problem can be formulated as a bi-level problem:

Problem 1 (MDP-based Journey Planning Problem):

$$\begin{aligned} \min_{x_{ij}^k, y_{ij}^k} \quad & (2) \\ \text{s.t.} \quad & (3)-(4), \\ & a_t^k \in \arg \max \left\{ \sum_{k,t} R(s_t^k, a_t^k, s_{t+1}^k) : (6)-(7) \right\}. \end{aligned}$$

There is approach to tackle Problem 1 such as the one in [15]. However, the development and implementation of an algorithm that efficiently solves this problem is outside the scope of this paper. The algorithm is assumed to be unknown to the attacker.

III. PASSENGER SPOOFING ATTACK

In this section, we introduce the threat model, the problem of spoofing attacks, and the methodology of the attack.

A. Threat Model

In this section, we present our assumptions regarding the malicious attack on the MaaS system, which involves the submission of a fabricated query for a passenger with a tailored state of satisfaction and profiles that is likely to be prioritized. In urban transportation systems, where there is heavy traffic and limited capacity, this attack may cause regular passengers with similar travel paths to be offered detour journeys during rush hours, leading to reduced passenger satisfaction and MaaS profit. We assume that the malicious agent has no prior knowledge of the coordinator and treats it like a black box. The malicious agent does not have access to the formulation of the passenger satisfaction problem, which is defined as Problem 1. However, the agent has basic knowledge of the environment, such as the state of other regular passengers, which can be used to generate a spoofing state through eavesdropping or related attacks [16].

B. Problem Formulation of PSA

In a PSA, a set of malicious agents $\mathcal{M}$ intend to sabotage the operations of the MaaS or take advantage by spoofing another passenger. In an urban transportation system with limited capacity and resources, competing passengers may reduce the chances of regular passengers to receive the most desired services. To achieve this goal, a malicious agent may act as a passenger with similar origin-destination patterns, such as the daily home-to-work flows during rush hour, to

occupy the resources of regular passengers. To maximize the impact of the attack, the malicious agent generates profiles and satisfaction based on those of another passenger. Specifically, the generated state $\tilde{s}_t^m$ can be defined as:

$$\tilde{s}_t^m = \pi(s_t^k | \theta^\pi), \quad (8)$$

where s_t^k represents the state of passenger k , including their profiles (such as income, age, etc.²), and satisfaction (H^k), $\tilde{s}_t^m$ is the generated state based on passenger k , π denotes the generation function, and θ^π represents its trainable parameters.

The objective of the PSA is to minimize the profit of the MaaS system and reduce passenger satisfaction. To this end, the PSA problem is formulated as an optimization problem in which the objective function is the profit generated by the set of passengers being attacked, as given in equation (9):

$$\sum_{k \in \mathcal{K}, t} R(s_t^k, a_t^k, s_{t+1}^k). \quad (9)$$

Since the profit is directly related to the performance of the MaaS coordinator, the malicious agent considers the journey planning problem as a black box and incorporates it into the attack problem, which is formulated as:

Problem 2 (PSA Problem):

$$\begin{aligned} \min_{\theta^\pi} \quad & \sum_{k \in \mathcal{K}, t} R(s_t^k, a_t^k, s_{t+1}^k) \\ \text{s.t.} \quad & a_t^k \in \arg \max \left\{ \sum_{k \in \mathcal{K}, t} R(s_t^k, a_t^k, s_{t+1}^k) \right. \\ & \left. + \sum_{m \in \mathcal{M}, t} R(\tilde{s}_t^m, a_t^m, \tilde{s}_{t+1}^m) : (6), (7), (8) \right\}. \end{aligned}$$

The controllable parameter in Problem 2 is only θ^π . In other words, the objective of the problem is to minimize the profit of the set of passengers $\mathcal{K}$, as expressed in Eq. (9), by determining the parameter θ^π and generation function π as shown in Eq. (8). This is distinct from the lower-level optimization task in Problem 1, which aims to maximize the profit of all passengers, including the set of malicious agents $\mathcal{M}$. Consequently, the malicious agent seeks to entice the MaaS coordinator to obtain more profit from $\mathcal{M}$, resulting in a lower actual profit received from the regular passengers $\mathcal{K}$. An example operation and flow are depicted in Fig. 1, in which the malicious agent pretends to be a passenger with the same origin and destination as the regular passenger, and the MaaS coordinator allocates the ideal mobility service to the malicious agent instead of the regular passenger based on falsified profiles and satisfaction.

C. Methodology of PSA

To develop a generation function π and corresponding parameters θ^π that can generate profiles and satisfaction based on the regular passenger profiles and satisfaction, a reinforcement learning approach can be employed to learn from the interactions between the MaaS coordinator and

²Here, income and age are just example profiles. Other profiles can be used without loss of generality.

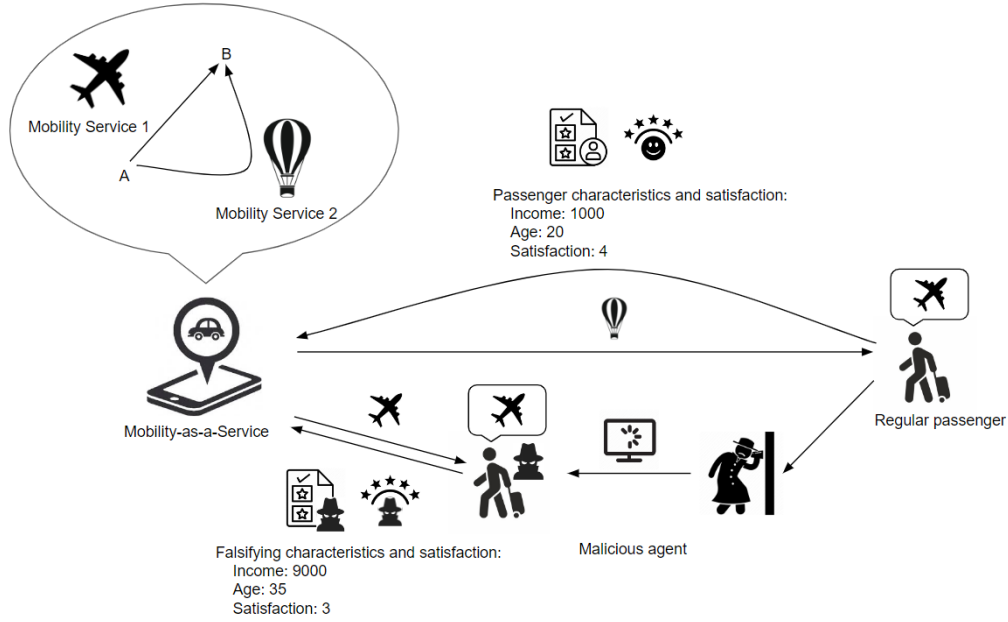


Fig. 1. An example showing the interactions among MaaS coordinator, regular passenger, and malicious agent.

passengers. The presented algorithm is a modified version of the deep deterministic policy gradient (DDPG) algorithm [17], which is suited for the PSA problem. It is a model-free algorithm for continuous control problems since the model of the MaaS coordinator is unknown to the malicious agent, and both the states and actions are continuous values. The malicious agent shares the same MDP as the MaaS coordinator, except for two differences from the agent's perspective. First, the action of the malicious agent is the generated profiles and satisfaction for the passenger spoofing, i.e., $\tilde{s}_t^m$. Second, the reward function is maximized, which is equal to minimizing the negative profit, i.e., $-\sum_{k \in \mathcal{K}, f \in \mathcal{F}, t} \rho_t^{kf}$. The algorithm consists of two components, namely the actor and the critic. The actor function $\pi(s|\theta^\pi)$ is responsible for performing an action based on a given state, whereas the critic function $Q(s, a)$ learns the Q-value of the state and action pair to evaluate the actor. The exploration of the algorithm is modified from an Ornstein-Uhlenbeck (OU) noise-based method to an epsilon-greedy action selection method to ensure that the action remains within the designated range. The epsilon-greedy action selection method specifies that a uniformly random action is generated under the probability represented by ϵ . The value of ϵ decays along the episode for more exploitation in later episodes. The detailed algorithm is provided in Algorithm 1.

IV. EXPERIMENTS

The reinforcement learning-based attack method is evaluated in a real-world transport network based on New York City (NYC). Specifically, it simulates the MaaS environment by considering multiple mobility providers and passengers. In this regard, this section details the experiment setups and presents the results obtained to validate the attack ability of the MaaS system.

A. Experiment Setups

1) *New York City Scenario*: A real-world scenario is considered based on the transport network of Manhattan region in NYC. The taxi zone maps³ are used to construct the transport network, where each node represents a taxi zone and an edge is assigned between two nodes if they are connected in the map. Only connected zones are included in the network, and any isolated zones are excluded. The transport network consists of 63 irregularly-shaped nodes, and 963 edges. For each edge, 3 mobility providers are set, and each mobility service has four attributes, namely, time, discomfort, and profit. The values of time, discomfort, and profit are randomly generated between 0 and 1.

Passenger traffic queries and profiles are simulated using two publicly available datasets of NYC: the NYC Taxi and Limousine Commission Trip Record Data⁴ and the Citywide Mobility Survey⁵. The datasets are processed and filtered to include only passengers within the same Manhattan region as in the transport network. The expected utility weight for the calculation of Eq. (7) is determined based on passenger profiles, which are unknown to the MaaS coordinator. Initially, all passengers have a satisfaction level of 3. The journey planning problem is solved using a standard optimizer in CVXPY [18].

2) *Baselines*: We conducted a comparative study between the proposed reinforcement learning-based PSA and several baselines, which are listed as follows:

- **Without attack**: This baseline does not include a malicious agent and serves as a benchmark for comparing

³<https://data.cityofnewyork.us/Transportation/NYC-Taxi-Zones/d3c5-ddgc>

⁴<https://www1.nyc.gov/site/tlc/about/tlc-trip-record-data.page>

⁵<https://www1.nyc.gov/html/dot/html/about/citywide-mobility-survey.shtml>

Algorithm 1 Multi-agent RL algorithm for PSA

```

1: Initialize actor local  $\pi(s|\theta^\pi)$  and critic local networks
    $Q(s, a|\theta^Q)$  with parameters  $\theta^\pi$  and  $\theta^Q$ 
2: Initialize parameters of actor target  $\pi'(s|\theta^{\pi'})$  and critic
   target networks  $Q'(s, a|\theta^{Q'})$  with parameters  $\theta^{\pi'} \leftarrow \theta^\pi$ 
   and  $\theta^{Q'} \leftarrow \theta^Q$ 
3: for each episode do
4:   Initialize passengers' parameters
5:   for iteration  $t = 1$  to  $T$  do
6:     for malicious agent  $m = 1$  to  $|\mathcal{M}|$  do
7:        $j_t^m \leftarrow$  random number between 0 and 1
8:       if  $j_t^m < \epsilon$  then
9:          $\tilde{s}_t^m \leftarrow$  random action between 0 to 1
10:      else
11:         $\tilde{s}_t^m \leftarrow \pi(s_t^k|\theta^\pi)$ 
12:      end if
13:      Execute action  $\tilde{s}_t^m$  to obtain new state  $s_{t+1}^m$  and
        reward  $r_t^m$ 
14:      Store  $(s_t^m, \tilde{s}_t^m, r_t^m, s_{t+1}^m)$  to  $\mathcal{R}$ 
15:    end for
16:    if number of transitions  $|\mathcal{R}| \geq$  minibatch size  $B$ 
        then
17:      Sample a mini-batch  $(s_i^m, \tilde{s}_i^m, r_i^m, s_{i+1}^m)$  with
        size  $B$  from  $\mathcal{R}$ 
18:      Update critic local  $\theta^Q$  using loss function:
         $L = \frac{1}{B} \sum_i (r_i^m + \gamma Q'(s_{i+1}^m, \pi'(s_{i+1}^m|\theta^{\pi'}))|\theta^{Q'}) -$ 
         $Q(s_i^m, \tilde{s}_i^m|\theta^Q)^2$ 
19:      Update actor local  $\theta^\pi$  using policy gradient:  $\nabla_{\theta^\pi} J \approx$ 
         $\frac{1}{B} \sum_i \nabla_a Q(s, a|\theta^Q)|_{s=s_i^m, a=\pi(s_i^m)} \nabla_{\theta^\pi} \pi(s|\theta^\pi)|_{s_i^k}$ 
20:      Update critic target  $\theta^{Q'} \leftarrow \tau \theta^Q + (1 - \tau) \theta^{Q'}$ 
21:      Update actor target  $\theta^{\pi'} \leftarrow \tau \theta^\pi + (1 - \tau) \theta^{\pi'}$ 
22:    end if
23:  end for
24:   $\epsilon := \epsilon \bar{\epsilon}$ 
25: end for
26: return actor parameters  $\theta^\pi$ 

```

the performance of the original AI-based system.

- **Random actions:** In this baseline, malicious agents generate random values of profiles and satisfaction $\tilde{s}_t^m$ within the range of $[0, 1]$.
- **Identical profiles:** In this baseline, the malicious agents have the same profiles and satisfaction $\tilde{s}_t^m$ as the target passenger.
- **Fixed actions:** In this baseline, the malicious agents always choose the same fixed values of 0.5 for both profiles and satisfaction $\tilde{s}_t^m$.
- **Complementary profiles:** In this baseline, the malicious agents compute their profiles and satisfaction based on the complement of the target passenger's profiles and satisfaction $\tilde{s}_t^m$.
- **Lower income:** In this baseline, the income state of the malicious agents is always half of the target passenger's income state.

TABLE I
PARAMETER SETTINGS.

Parameter	Definition	Value
$ \mathcal{N} $	Number of nodes	63
$ \mathcal{A} $	Number of links	963
$ \mathcal{F} $	Number of mobility providers	3
$ \mathcal{R} $	Replay buffer size	10^6
B	Minibatch size	128
γ	Discount factor	0.99
τ	Target network soft update rate	0.001
-	Actor learning rate	0.0001
-	Critic learning rate	0.0003
-	Neural network optimizer	Adam
ϵ_0	Initial random explore rate	1
$\bar{\epsilon}$	Explore rate decay per episode	0.9995
T	Number of iteration per episode	10
-	Number of neural network layers	3
-	Number of neurons of each layer	256
-	Range of satisfaction level	1 to 5
$\bar{E}^k$	upper expectation threshold	0.0
$\underline{E}^k$	lower expectation threshold	-0.1

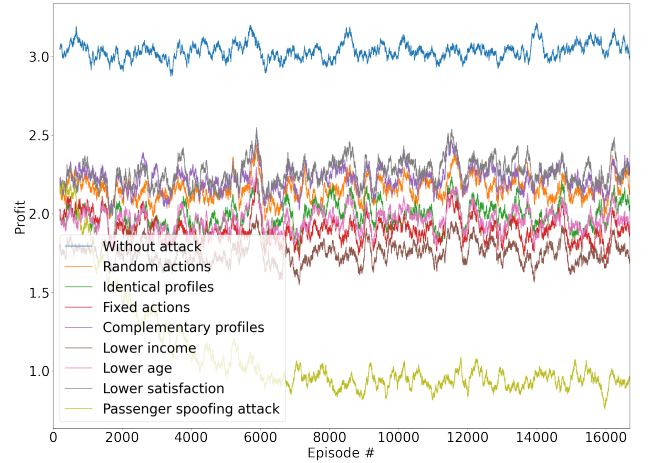


Fig. 2. Moving average reward of various approaches. The time window of the moving average is equal to 200.

- **Lower age:** In this baseline, the age state of the malicious agents is always half of the target passenger's age state.
- **Lower satisfaction:** In this baseline, the satisfaction state of the malicious agents is always half of the target passenger's satisfaction state.

The aim of these baselines is to evaluate the effectiveness of the proposed reinforcement learning-based PSA in comparison to other types of malicious attacks.

B. Experiment Results

We conducted a comprehensive analysis to test the attack ability of the malicious agent in various aspects. Specifically, we evaluated its ability to reduce the profit of the MaaS

TABLE II
AVERAGE PROFIT OF THE BASELINES AND ATTACK OF NYC AND
SYNTHETIC SCENARIOS.

	NYC Scenario
Without attack	3.0431
Random actions	2.1746
Identical profiles	2.0013
Fixed actions	1.8927
Complementary profiles	2.2351
Lower income	1.7592
Lower age	1.9691
Lower satisfaction	2.2891
PSA	0.9408

system, decrease passenger satisfaction, and examined the sensitivity of the spatial distance between the malicious agent and regular passengers.

1) *MaaS Profit*: We conduct a comparative analysis of the reinforcement learning-based PSA approach and baselines to investigate their impact on MaaS profit. The profit earned per regular passenger in each episode is plotted against the corresponding method during training, as shown in Fig. 2. In all attack cases, except the one without an attack, 50% of the passengers are regular and the rest are malicious agents. The total number of regular and malicious passengers remains the same to ensure fair comparison. The case without an attack generates the highest profit, while the one with PSA results in the lowest profit. Most of the baselines fall in the middle of the spectrum. The difference between the cases without attack and baselines is attributed to the addition of incapable agents with the same origin and destination, while the difference between baselines and PSA highlights the attack ability of the method. By training the malicious agent to generate appropriate profiles and satisfaction that may receive higher priority in journey planning over other regular passengers, the profit declines at the beginning and converges around 7500 episodes. The average profit of the baselines and attack scenarios are reported in Table II. The average profit after 10,000 episodes with and without PSA are 0.94 and 3.04 units, respectively. Most of the other baselines fall within the range of 1.7 to 2.2 units. This indicates that the malicious agent can crowd out regular passengers using the generated profiles and satisfaction, especially when the capacity is limited. The attack scenarios result in a profit reduction of about 70%.

2) *Satisfaction*: We conducted an analysis of the differences in passenger satisfaction levels between the baselines and PSA. The frequency distribution of regular passenger satisfaction levels across all episodes is summarized in Fig. 3. The leftmost bars represent the satisfaction levels resulting from the MaaS coordinator without attack, and show that many passengers had high satisfaction levels. However, after introducing the baselines (middle bars), the satisfaction of most regular passengers dropped to level 3. In the case of PSA, most passengers had the lowest satisfaction level. These

results indicate that the malicious agent can reduce the satisfaction levels of other regular passengers by impersonating a passenger with specific profiles and satisfaction, which is consistent with our observations from the profit analysis.

3) *Number of nodes apart between malicious agent and regular passenger*: In order to investigate the relationship between the spatial distance from the passengers and the attack ability, we conducted a series of experiments where we varied the origin and/or destination of the malicious agent. The resulting profit was plotted against the number of nodes apart for three different scenarios: origin only, destination only, and both origin and destination, as shown in Figs. 4. Generally, we observed an increase in profit as the number of nodes increased. The green line, which represents the case where the number of nodes apart is varied for both origin and destination, showed a greater increase in profit than the other two cases. For the case where both origin and destination were varied, the profit increase started to converge when the number of nodes apart was larger than one. For the other two cases, the increases converged when the number of nodes apart was larger than three. The effect of increasing the distance from a malicious agent was relatively small when it was three nodes away. Thus, we can conclude that the attacking effect will be diminished at an average of three nodes apart from the passengers.

V. CONCLUSIONS

In order to safeguard the daily operation of urban transportation systems and prevent the incurrence of significant economic losses, it is crucial to establish a secure and reliable MaaS system. To effectively counter potential cybersecurity attacks, it is essential to identify vulnerabilities and potential attack strategies. This study analyzes threats to the MaaS journey planning process and identifies PSA that capitalizes on passenger heterogeneity. PSA is based on reinforcement learning algorithm to generate spoofing passenger profiles that are likely to be prioritized by the MaaS coordinator. Experimental results based on real-world datasets and transport networks demonstrate that the PSA method can effectively reduce profits by 70% and significantly lower passenger satisfaction levels. We also discovered that the attack will be diminished at an average of three nodes apart from the passengers.

There are several promising directions for future research. Firstly, our study mainly focuses on the PSA attack and its profiles, while countermeasures against the attack need to be developed to enhance the security of MaaS systems. Secondly, it would be valuable to investigate the potential differences in the attack when multiple malicious agents are present in the system, which could form a multi-agent scenario [19].

REFERENCES

- [1] S. Hietanen, "Mobility as a service," *the new transport model*, vol. 12, no. 2, pp. 2–4, 2014.
- [2] A. Nikitas, K. Michalakopoulou, E. T. Njoya, and D. Karampatzakis, "Artificial intelligence, transport and the smart city: Definitions and dimensions of a new mobility era," *Sustainability*, vol. 12, no. 7, p. 2789, 2020.

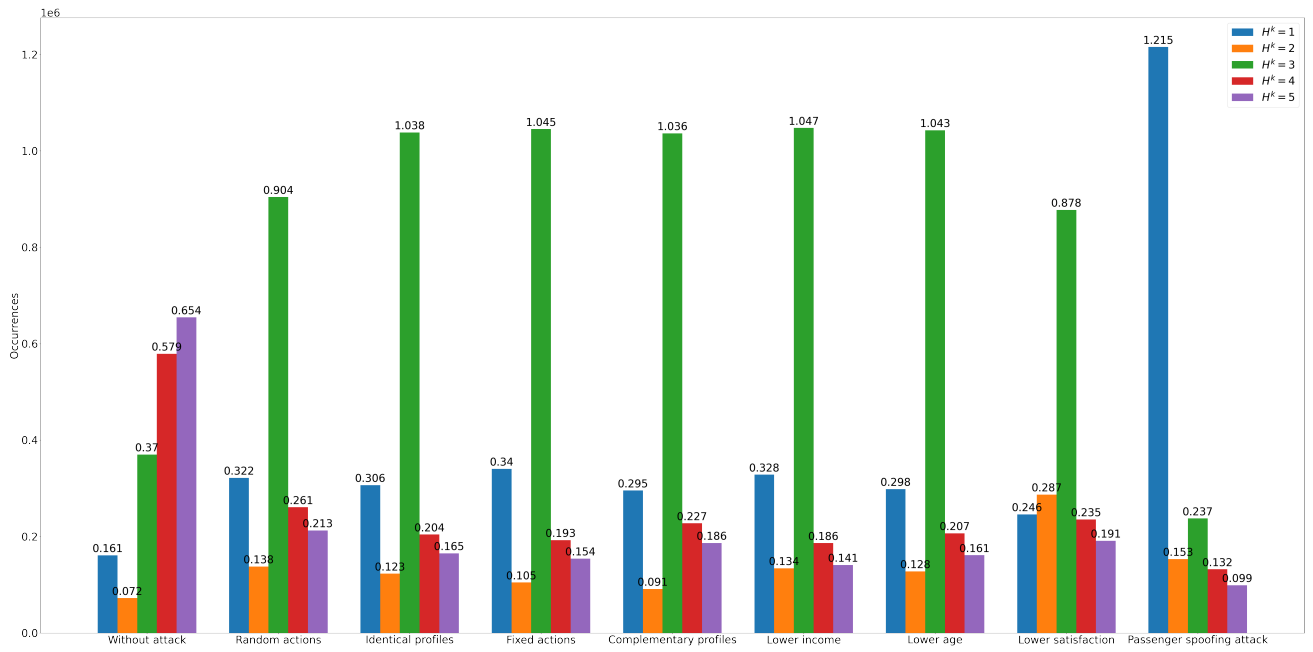


Fig. 3. Total number of satisfaction level of each method.

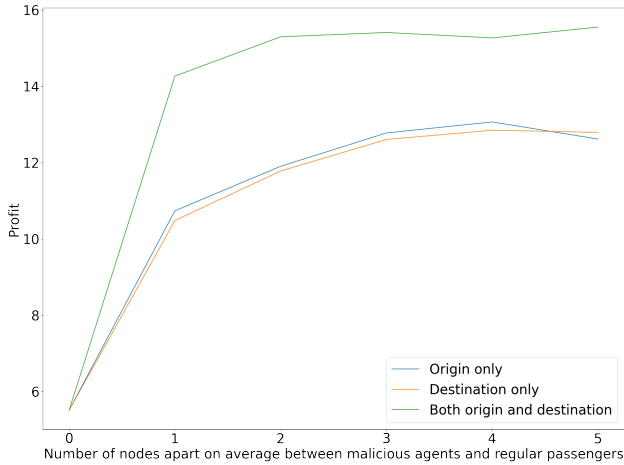


Fig. 4. Profit against number of nodes apart on average between malicious agents and regular passengers.

- [3] K.-F. Chu, A. Y. Lam, and V. O. Li, "Deep multi-scale convolutional lstm network for travel demand and origin-destination predictions," *IEEE Transactions on Intelligent Transportation Systems*, vol. 21, no. 8, pp. 3219–3232, 2019.
- [4] B. R. Kiran, I. Sobh, V. Talpaert, P. Mannion, A. A. A. Sallab, S. Yogamani, and P. Pérez, "Deep reinforcement learning for autonomous driving: A survey," *IEEE Transactions on Intelligent Transportation Systems*, vol. 23, no. 6, pp. 4909–4926, 2022.
- [5] K.-F. Chu, A. Y. Lam, and V. O. Li, "Traffic signal control using end-to-end off-policy deep reinforcement learning," *IEEE Transactions on Intelligent Transportation Systems*, vol. 23, no. 7, pp. 7184–7195, 2022.
- [6] F. Callegati, S. Giallorenzo, A. Melis, and M. Prandini, "Cloud-of-things meets mobility-as-a-service: An insider threat perspective," *Computers & Security*, vol. 74, pp. 277–295, 2018.
- [7] J. Thai, C. Yuan, and A. M. Bayen, "Resiliency of mobility-as-a-service systems to denial-of-service attacks," *IEEE Transactions on Control of Network Systems*, vol. 5, no. 1, pp. 370–382, 2016.
- [8] A. Adadi and M. Berrada, "Peeking inside the black-box: a survey on explainable artificial intelligence (XAI)," *IEEE Access*, vol. 6, pp. 52 138–52 160, 2018.
- [9] J. Steinhardt, P. W. W. Koh, and P. S. Liang, "Certified defenses for data poisoning attacks," *Proceedings of the Advances in Neural Information Processing Systems*, vol. 30, 2017.
- [10] B. Biggio, I. Corona, D. Maiorca, B. Nelson, N. Šrndić, P. Laskov, G. Giacinto, and F. Roli, "Evasion attacks against machine learning at test time," in *Proceedings of the Joint European Conference on Machine Learning and Knowledge Discovery in Databases*. Springer, 2013, pp. 387–402.
- [11] R. Shokri, M. Stronati, C. Song, and V. Shmatikov, "Membership inference attacks against machine learning models," in *Proceedings of the IEEE symposium on security and privacy*, 2017, pp. 3–18.
- [12] K.-F. Chu, A. Y. Lam, B. P. Loo, and V. O. Li, "Public transport waiting time estimation using semi-supervised graph convolutional networks," in *Proceedings of the 2019 IEEE Intelligent Transportation Systems Conference*. IEEE, 2019, pp. 2259–2264.
- [13] K.-F. Chu, A. Y. S. Lam, and V. O. K. Li, "Joint rebalancing and vehicle-to-grid coordination for autonomous vehicle public transportation system," *IEEE Transactions on Intelligent Transportation Systems*, vol. 23, no. 7, pp. 7156–7169, 2022.
- [14] R. Bellman, "A markovian decision process," *Journal of Mathematics and Mechanics*, pp. 679–684, 1957.
- [15] K.-F. Chu and W. Guo, "Deep reinforcement learning of passenger behavior in multimodal journey planning with proportional fairness," *Neural Computing and Applications*, pp. 1–20, 2023.
- [16] I. Zografopoulos, J. Ospina, X. Liu, and C. Konstantinou, "Cyber-physical energy systems security: Threat modeling, risk assessment, resources, metrics, and case studies," *IEEE Access*, vol. 9, pp. 29 775–29 818, 2021.
- [17] T. P. Lillicrap, J. J. Hunt, A. Pritzel, N. Heess, T. Erez, Y. Tassa, D. Silver, and D. Wierstra, "Continuous control with deep reinforcement learning," in *Proceedings of the International Conference on Learning Representations*, 2016.
- [18] S. Diamond and S. Boyd, "CVXPY: A Python-embedded modeling language for convex optimization," *Journal of Machine Learning Research*, vol. 17, no. 83, pp. 1–5, 2016.
- [19] X.-M. Li, Q. Zhou, P. Li, H. Li, and R. Lu, "Event-triggered consensus control for multi-agent systems against false data-injection attacks," *IEEE Transactions on Cybernetics*, vol. 50, no. 5, pp. 1856–1866, 2019.

2024-02-13

Passenger spoofing attack for artificial Intelligence-based Mobility-as-a-Service

Chu, Kai-Fung

IEEE

Chu KF, Guo W. (2023) Passenger spoofing attack for artificial intelligence-based Mobility-as-a-Service. In 2023 IEEE 26th International Conference on Intelligent Transportation Systems (ITSC), 24-28 September 2023, Bilbao, Spain, pp. 4874-4880 pp. 4874-4880
<https://doi.org/10.1109/ITSC57777.2023.10422567>

Downloaded from Cranfield Library Services E-Repository