



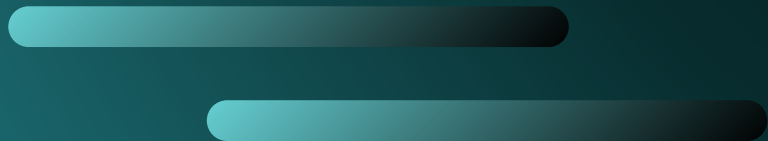
ANALÝZA BEZPEČNOSTNÝCH ÚDAJŮ Z HONEYPOTŮ POMOCÍ VELKÝCH JAZYKOVÝCH MODELOV

Vedúci práce: doc. RNDr. JUDr. Pavol Sokol, PhD. et PhD.

Autor: Sophia Petra Krišáková

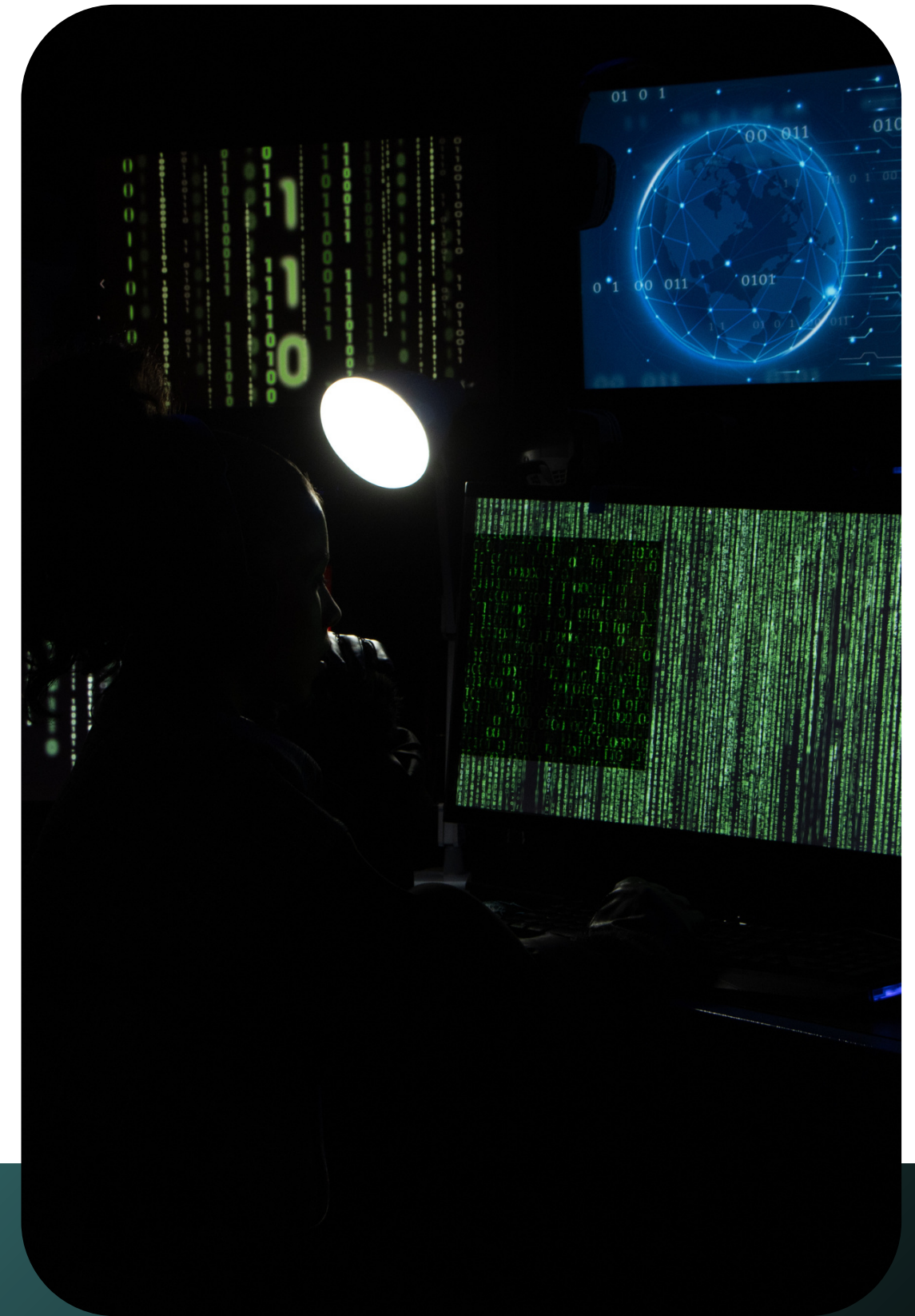
Ústav informatiky, Prírodovedecká fakulta, UPJŠ

21.11.2024



Motivácia

- LLM + Cybersecurity



Honeypoty a dáta z nich

- lákadlo pre útočníkov, simulovaný systém
- zhromažďovanie informácií o metódach útoku
- identifikácia zraniteľností
- maskovaný ako zraniteľný cieľ



- <https://haas.nic.cz/stats/>



	time	time_closed	ip	country	login_successful	commands
0	2024-09-01T00:00:07.720868+00:00	2024-09-01T00:00:08.726674+00:00	112.72.235.19	KR	False	[]
1	2024-09-01T00:00:07.785101+00:00	2024-09-01T00:00:08.789687+00:00	112.72.235.19	KR	False	[]
2	2024-09-01T00:00:07.808650+00:00	2024-09-01T00:00:08.813711+00:00	112.72.235.19	KR	False	[]
3	2024-09-01T00:00:07.865543+00:00	2024-09-01T00:00:08.870439+00:00	112.72.235.19	KR	False	[]
4	2024-09-01T00:00:07.899126+00:00	2024-09-01T00:00:08.903224+00:00	112.72.235.19	KR	False	[]
...	...	...	...	...	...	...
22026	2024-09-01T23:59:42.472663+00:00	2024-09-01T23:59:42.981977+00:00	122.46.16.230	KR	True	[ps grep '[Mm]iner']
22027	2024-09-01T23:59:44.649059+00:00	2024-09-01T23:59:45.102192+00:00	122.46.16.230	KR	True	[ps -ef grep '[Mm]iner']
22028	2024-09-01T23:59:46.752907+00:00	2024-09-01T23:59:47.258398+00:00	122.46.16.230	KR	True	[ls -la /dev/ttyGSM* /dev/ttyUSB-mod* /var/spo...
22029	2024-09-01T23:59:49.092738+00:00	2024-09-01T23:59:49.586087+00:00	122.46.16.230	KR	True	[echo Hi cat -n]
22030	2024-09-01T23:59:50.550074+00:00	2024-09-02T00:00:01.413132+00:00	180.117.252.74	CN	True	[/ip cloud print]

22031 rows × 6 columns

Návrh postupu riešenia

**Analýza
podobných
prác**

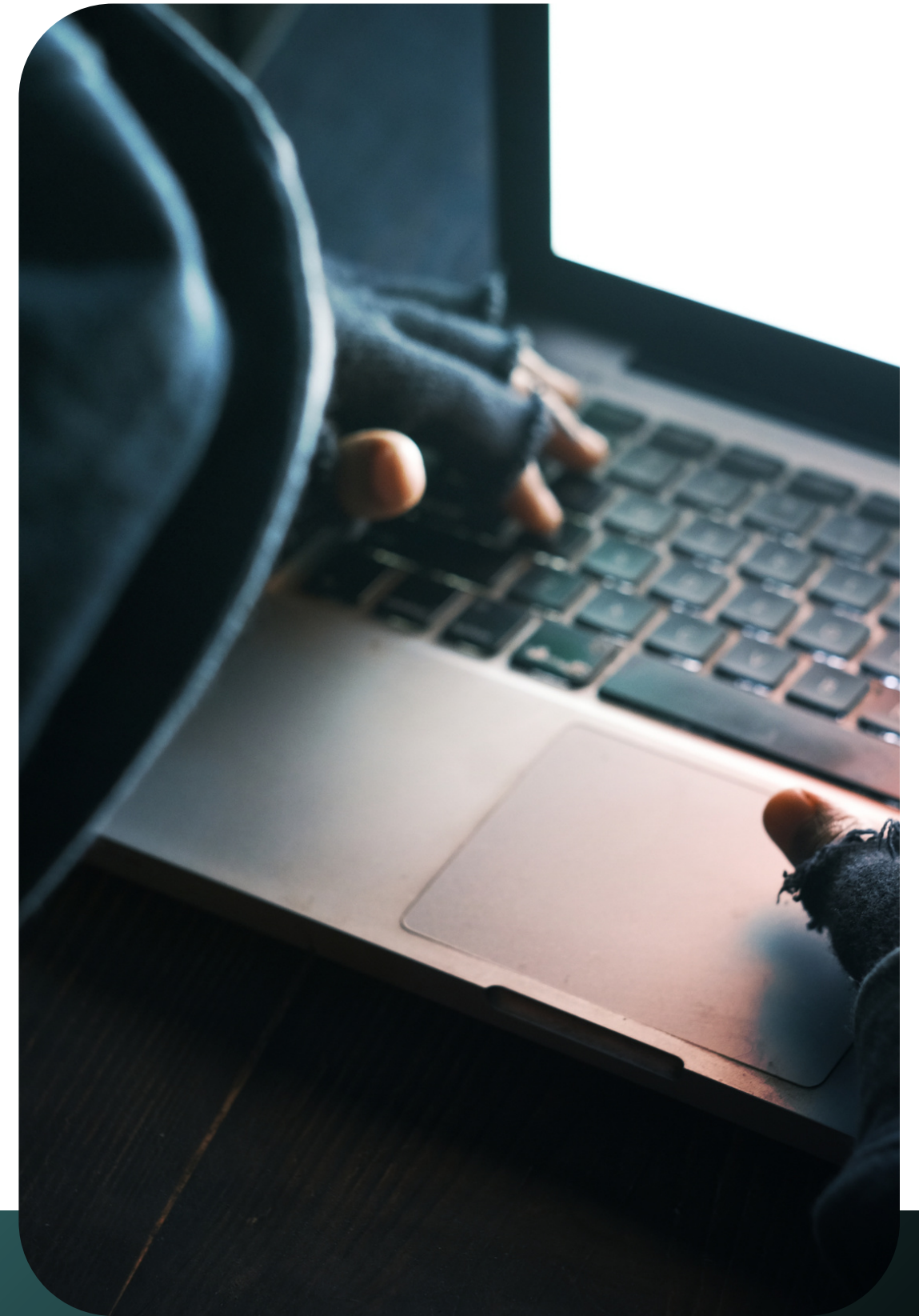
**Predspracovanie
a analýza dát**

**Možnosť
použitia iných
dát, alebo
požiadaviek**

**Analýza
možných
prípadov**

Use case

Predikcia správania útočníka



Materiály

ChatGPT for digital forensic investigation: The good, the bad, and the unknown

- Vplyv ChatGPT (GPT-4) na digitálnu forenznú analýzu,
- LLM modely: ChatGPT, BERT, Bard, LLaMA – generujú odpovede na základe veľkých dát,
- experimenty: Skúmanie ChatGPT pri úlohách:
 - porozumenie artefaktov,
 - vyhľadávanie dôkazov,
 - generovanie kódu,
 - detekcia anomálií,
 - riešenie incidentov,
 - vzdelávanie.

Large Language Models in Cybersecurity: State-of-the-Art

- Vplyv veľkých jazykových modelov (LLMs) na kybernetickú bezpečnosť,
- kybernetická bezpečnosť: pomaly prijíma strojové učenie, doteraz odolná voči dátovým riešeniam
 - preskúmanie a kategorizácia obranných a útočných aplikácií LLM v kybernetickej bezpečnosti,
 - identifikácia výskumných medzier.

Large Language Models in Cybersecurity

- Riziká veľkých jazykových modelov (LLMs) v kybernetickej bezpečnosti a ich zmiernenie,
- varovať vývojárov LLM pred rizikami a poskytnúť im nástroje na ich zníženie
- Predstihnúť útočníkov predpovedaním ich krokov,
- hrozby LLM pre kyberbezpečnosť, vrátane nových útokov,
- predpovede vývoja LLM a regulačné možnosti,
- techniky na bezpečný vývoj a nasadenie LLM,
- návrh bezpečnej integrácie LLM a zhrnutie ich dvojznačnosti v kybernetickej bezpečnosti.

Ciele

01

Preštudovať a spracovať prehľad existujúcich prístupov k analýze bezpečnostných údajov pomocou veľkých jazykových modelov.

02

Navrhnuť a implementovať vhodnú metódu na analýzu bezpečnostných údajov z honeypotov pomocou veľkých jazykových modelov.

03

Porovnať výkonnosť navrhnutého riešenia s inými dostupnými štúdiami.

Ďakujem za pozornosť