

MEGA-RCD: Um Método de Gerenciamento de Ameaças em Rede de Comunicação de Dados

Miquelângelo de Souza Dias, Cesar A. C. Marcondes e Adilson Marques da Cunha
Instituto Tecnológico de Aeronáutica (ITA), São José dos Campos/SP - Brasil

Resumo—No mundo moderno, principalmente a partir de 2020, período da pandemia e pós-pandemia do Vírus COVID-19 (SARS-CoV-2), vem ocorrendo um crescimento de ameaças cibernéticas devido ao aumento na utilização de Portais da Internet, aplicativos de compra *online* e outras tecnologias com vulnerabilidades, visando oportunizar e intensificar acessos, via Internet, aos produtos e serviços. Neste contexto, ameaças originadas de Agentes Maliciosos para explorar vulnerabilidades em produtos vêm representando o principal foco deste trabalho de pesquisa, ainda em fase inicial de investigação envolvendo a concepção, modelagem e implementação de um Método de Gerenciamento de Ameaças em Rede de Comunicação de Dados denominado MEGA-RCD, com o propósito de mitigar vulnerabilidades, como medida preventiva contra ataques cibernéticos. Os resultados iniciais desta pesquisa apontam para a concepção de um método de identificação de vulnerabilidades composto por Plataforma de Orquestração integrada à Inteligência de Ameaças. Assim, propiciando a mitigação de vulnerabilidades.

Palavras-Chave—Cibersegurança, Inteligência de Ameaças, Verificação de Vulnerabilidades.

I. INTRODUÇÃO

A partir do primeiro Semestre de 2020, período compreendido pela pandemia e pós-pandemia do Vírus COVID-19 (SARS-CoV-2), foi intensificado o uso de sistemas online no mundo, como uma forma de proteção de contatos humanos e espalhamentos do vírus.

Nesse novo cenário, houve um aumento exponencial do uso de sistemas computacionais ligados à Portais da Internet, Aplicativos de Compras *Online*, Plataformas de Ensino à Distância, Acessos à Redes corporativas, entre outros.

Entretanto, ao mesmo tempo que o uso de Sistemas *Online* foi incrementado, tal migração repentina propiciou um aumento na frequência e intensidade de violações de segurança. Isto ocorreu, devido ao fato de que todos os processos e preocupações com segurança, como acessos remotos e processos de desenvolvimento seguro, não são devidamente levados à risca.

Além dessa situação única devido a pandemia, outro movimento, envolvendo novas fontes de vulnerabilidades, foi devido aos processos de digitalização de governos.

Um desses processos tem sido a Estratégia de Governo Digital para o período de 2020 a 2022, envolvendo a previsão de utilização de tecnologias emergentes e a criação ou migração de ambientes e sistemas locais para plataformas em nuvem, disponibilizando serviços para a sociedade como um todo. Em outras palavras, um aumento de conectividade, uso de rede, uso de sistemas *online* intensivo em futuro próximo [1] [2].

Dentro desse contexto digital, os chamados "Agentes Maliciosos", grupos ou indivíduos com intenções malignas, vêm realizando diversas ameaças no espaço cibernético, explorando

um volume cada vez maior de vulnerabilidades, sofisticando as suas táticas e métodos e ocasionando sérios riscos e prejuízos à governos, sociedades e empresas [1].

Exemplos recentes, no ambiente brasileiro, mostram empresas energéticas tendo que pagar resgates para liberar seus dados, chantagens e até plataformas de mensagem em redes sociais sendo invadidas para propagação de golpes.

Desse modo, torna-se necessário tratar eficientemente, com boa assertividade e de maneira automatizada, um número cada vez maior de crimes cibernéticos, originados a partir de explorações de vulnerabilidades.

Esta temática representou uma das motivações deste trabalho de pesquisa para se reduzir ameaças cibernéticas, através da concepção, modelagem e implementação de método de melhoria do processo de decisão com inteligência artificial e mitigação automática de vulnerabilidades conhecidas, pela integração com sistemas atuadores em redes.

Em resumo, nosso artigo tem a pretensão posicional de apresentar ideias acerca do método e sua arquitetura, bem como detalhes do fluxo de informações, algoritmos pretendidos para melhorar o desempenho e a apresentação de resultados iniciais em *notebooks*. As contribuições do trabalho situam-se na área de orquestração de segurança, em especial, e enriquecimento de dados para melhoria da tomada de decisão, tratando especificamente as vulnerabilidades de rede.

Nossa hipótese de pesquisa é que os sistemas de proteção não têm acompanhado em velocidade e volume, os atacantes, devido a enxurrada de vulnerabilidades percebidas semanalmente. E a rede tornou-se um campo comum para a chegada de vulnerabilidades e, portanto, o local mais apropriado para se desenvolver um ambiente automático de proteção.

Para isso, faz-se necessário melhorar decisões, por meio de classificação de padrões e ranqueamento de ameaças, bem como estabelecimento de algum nível de proteção automática, em tempo real. Portanto, o nosso método denominado **MEGA-RCD - MÉtodo de Gerenciamento de Ameaças em Rede de Comunicação de Dados** foi elaborado para melhorar o tratamento dessas questões.

O restante do artigo encontra-se organizado da seguinte forma: na segunda seção, faremos uma apresentação de conceitos em segurança cibernética para descrever como as ameaças são coletadas e organizadas em termos de inteligência e como sistemas podem atuar de maneira semi ou automatizada para contê-las.

Em seguida, na terceira seção, apresentaremos trabalhos da literatura científica correlatos ao nosso, que capturamos através de uma Revisão Sistemática. A quarta seção descreve o método **MEGA-RCD**, apresentando detalhes de sua arquitetura, o fluxo de decisão e mensagens a atuadores em tempo real. Finalmente, na quinta e última seção, apresentamos os resultados iniciais mostrando-se uma prova de conceito em *notebooks* do método sendo desenvolvido, as principais

E-mail para correspondência com autor: miquelangelo@ita.br. Agradecimentos ao Exército Brasileiro e ao Instituto Tecnológico de Aeronáutica (ITA) por todo apoio prestado a esta pesquisa.

conclusões, recomendações e sugestões para trabalhos futuros.

II. FUNDAMENTAÇÃO TEÓRICA

A área de segurança cibernética é bastante vasta em nomenclatura, definições, métodos, técnicas e métricas, portanto, cabe dar importância a alguns temas cruciais para o entendimento do trabalho. Neste sentido, nesta seção, abordaremos as áreas de segurança, inteligência de ameaças e plataformas de orquestração que nos nortearam para a concepção da proposta.

A. Segurança Cibernética

A ideia de se prover Segurança Cibernética consiste em propiciar a proteção do espaço cibernético contra atividades que culminem em ferir os princípios fundamentais da Segurança da Informação, tais como: confidencialidade, integridade, disponibilidade e autenticidade de serviços e dados armazenados, processados ou transmitidos. Este papel demanda uma visão atenta para gerenciar mudanças contínuas em múltiplas áreas políticas, tecnológicas, educacionais e legais, a fim de se manter segurança nacional, economia e liberdade [3] [1].

Desse modo, medidas têm sido adotadas para a segurança de operações digitais, melhorando-se requisitos de princípios de segurança. Tais medidas podem ser divididas em três grupos de funcionalidades, complementares entre si, caracterizadas por ações de prevenção, detecção e resposta à incidentes de rede, conforme as definições mostradas na Tabela I [4].

TABELA I
MEDIDAS DE SEGURANÇA CIBERNÉTICA.

Termo	Descrição
Prevenção	mecanismo que visa proteger contra ataques, através de medidas que busquem antecipar possíveis ataques, podendo fornecer autenticação, controle de acesso, políticas e segmentação de rede.
Deteção	estratégia que objetiva identificar comportamentos anômalos e ataques ao sistema.
Resposta	mecanismo que tem a finalidade de mitigar o impacto de ataques correntes, preferencialmente de forma automatizada através de gatilhos, ou restaurar o sistema.

Detalhando essas ações, começaremos especificando a prevenção de exploração de vulnerabilidades. Esta constitui-se em uma tarefa crítica de manutenção de segurança de operações. É de conhecimento geral que vulnerabilidades de bibliotecas e sistemas, quando descobertas, geram relatórios de Vulnerabilidades e Exposições Comuns (em inglês, *Common Vulnerability and Exposure* - CVE).

Esses relatos integram bases de vulnerabilidades, como a Base Nacional de Vulnerabilidades (em inglês, *National Vulnerability Database* - NVD), mantido pelo Instituto Nacional de Padrões e Tecnologias Americano (em inglês, *National Institute of Standards and Technology* - NIST) e outra base é mantida pela *MITRE Corporation* do governo americano. Tais bases são utilizadas pela comunidade em geral para verificação pró-ativa de riscos em seus sistemas [5].

Esses CVE representam um dicionário público de vulnerabilidades conhecidas, com o propósito de identificar e divulgar vulnerabilidades de versões específicas de software [6]. Existem outras taxonomias para retratar vulnerabilidades como [7] [8]:

- Táticas, técnicas e conhecimentos comuns de atacantes - (em inglês, *Adversarial Tactics, Techniques, and Common*

Knowledge - ATT&CK), que fornece uma coleção de atores e suas táticas conhecidas (10 categorias de táticas) e técnicas pós-comprometimento para alcançar seus objetivos;

- Enumeração e Classificação de Padrão de Ataque Comum - (em inglês, *Common Attack Pattern Enumeration and Classification* - CAPEC), que fornece uma coleção das técnicas e métodos mais comuns usados em ataques cibernéticos, incluindo resumos, pré-requisitos de ataque e soluções de contramedidas para os padrões de ataque mais comuns;
- Enumeração de configuração comum - (em inglês, *Common configuration enumeration* - CCE), que fornece identificadores exclusivos para problemas de configuração do sistema relacionados à segurança, a fim de melhorar o fluxo de trabalho, facilitando a correlação rápida e precisa dos dados de configuração em várias fontes e ferramentas de informação;
- Sistema de Pontuação de Vulnerabilidade Comum - (em inglês, *Common Vulnerability Scoring System* - CVSS), enquanto poucos fornecem apenas uma medida de impacto, alguns (*Virtual Data Bases* - VDBs) fornecem informações de impacto e risco para cada vulnerabilidade;
- Enumeração de fraqueza comum - (em inglês, *Common Weakness Enumeration* - CWE), que oferece uma categorização de vulnerabilidade em formato de hierarquia; e
- Enumeração de Plataforma Comum - (em inglês, *Common Platform Enumeration* - CPE), que é uma enumeração de plataforma de suporte padrão. Correlacionar CPE com informações de vulnerabilidade permite mapear vulnerabilidades para suas entidades afetadas.

Nessas atividades de prevenção, organizações utilizam-se de variadas soluções de segurança para prevenir ataques e reduzir consequências associadas à vulnerabilidades e ameaças. Algumas dessas soluções se constituem em antivírus, *Firewall* e Sistema de Prevenção de Intrusão (em inglês, *Intrusion Prevention System* - IPS) [9].

Outra atividade associada a prevenção tem sido o Teste de Intrusão (em inglês, *Pentesting*), como uma técnica que envolve simulação de ataques reais para avaliação de vulnerabilidades em sistemas computacionais [10].

A partir dessas definições, a pesquisa vem se restringindo em conceber um método para identificar e mitigar, de forma automatizada, vulnerabilidades exploradas por Ameaças originadas de Agentes Maliciosos, como medida preventiva contra ataques cibernéticos. Para atingir estes objetivos o método proposto deve utilizar-se de Inteligência de Ameaças (em inglês, *Threat Intelligence* - TI) e estar envolto em uma Plataforma de Orquestração de Segurança.

B. Inteligência de Ameaças

O conceito de Inteligência de Ameaças vem crescendo em importância, por proporcionar um meio para se receber e compartilhar técnicas, táticas e procedimentos realizados por agentes maliciosos em evento cibernético, a fim de facilitar a identificação e mitigação de ameaças. Deste modo, as Tecnologias da Informação - TIs, através de compartilhamentos de informações a respeito de ataques cibernéticos, podem auxiliar na prevenção de outros ataques, na medida em que

representam conhecimentos baseados em evidências [11] [12]. A Inteligência de Ameaças também vem sendo considerada como um processo de obtenção de informações provenientes de diversas fontes sobre ameaças cibernéticas, a fim de enriquecimento da base de conhecimento, proporcionando um melhor entendimento de correlação de eventos, prováveis atores envolvidos, alvos de interesse de agentes maliciosos, entre outras informações [12] [7].

Ferramentas de TI, como a Plataforma de Compartilhamento de Ameaças (em inglês, *Malware Information Sharing Platform - MISP*) [13], vêm propiciando inúmeras informações tais como: base de dados de eventos, com incidentes, atacantes, vulnerabilidades utilizadas, entre outras; e indicadores de compromissos e regras para equipamentos de segurança em redes.

Outras técnicas para obtenção de Inteligência de Ameaças são conhecidas como ferramentas de Inteligência de fonte aberta (em inglês, *Open Source Intelligence - OSINT*), consistindo em coletar, processar e correlacionar informações públicas de fontes abertas como: redes sociais, médias, fóruns e *blogs*, dados governamentais e comerciais, além de informações obtidas através de buscas na Internet, com termos chave, endereços de IP, serviços e sistemas operacionais, representando uma ferramenta para identificação de vulnerabilidades [14].

Assim, as TIs vêm propiciando informações relevantes, atualizadas e, em alguns casos, com aplicabilidades imediatas para equipamentos de segurança de rede. Estas informações proporcionam, para um método de gerenciamento de ameaças, oportunidades de mitigar vulnerabilidades na infraestrutura, reduzindo a probabilidade de comprometimento inicial.

C. Plataformas de Orquestração

Uma Plataforma de Orquestração de Segurança, Automação e Resposta (em inglês, *Security Orchestration, Automation and Response - SOAR*) consiste de um conjunto de soluções de softwares compatíveis que possibilitem a macro gestão de diversos ativos e tecnologias, automação de coleta e resposta a incidentes de rede.

Essa plataforma pode incluir recursos de gestão de segurança, análise de dados de múltiplas fontes para proporcionar automatização de fluxos de trabalhos de análises e relatórios para equipes de segurança.

Ela pode oferecer configurações de segurança através de Inteligência de Ameaças para outros dispositivos, como *firewalls*, Sistema de Detecção de Intrusão (em inglês, *Intrusion Detection System - IDS*), Sistema de Prevenção de Intrusão (em inglês, *Intrusion Prevention System - IPS*) e Software de Gerenciamento de Informações e Eventos de Segurança (em inglês, *Security Information and Event Management - SIEM*).

Com essa plataforma, a partir de soluções SOAR, espera-se responder adequadamente a eventos de segurança, e a aprimorar a eficácia das operações no cenário digital [1].

Uma Segurança Orquestrada, também, pode pressupor planejamento, integração, cooperação e coordenação de ferramentas de segurança com o intuito de produzir ações automáticas em resposta a eventos de segurança. Neste sentido, uma Plataforma de Orquestração de Segurança pode ser composta por unidades de Unificação, Orquestração e Automação, com seus respectivos módulos, conforme mostrado na Figura 1, proporcionando

funcionalidades que objetivam, principalmente, interligar a lacuna entre detecção e remediação de eventos de segurança [9].

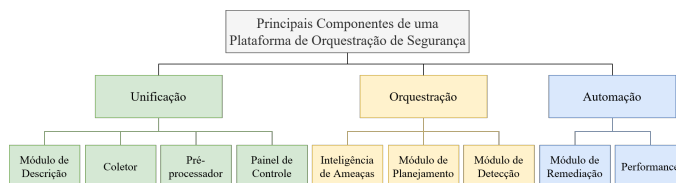


Fig. 1. Componentes de uma Plataforma SOAR [9].

Uma Plataforma de Orquestração pode proporcionar informações mais relevantes e com maior detalhamento utilizando ferramentas que possibilitem o enriquecimento de eventos em Plataformas de TI. Este enriquecimento pode ser proveniente da integração com outras bases de conhecimento, através de um Método de Gerenciamento de Ameaças.

III. PRINCIPAIS TRABALHOS CORRELATOS ENCONTRADOS

Na condução deste trabalho de pesquisa foi realizado um levantamento do estado da arte, utilizando-se um Método de Revisão de Literatura e a Síntese dos Principais Trabalhos Correlatos Encontrados na revisão descritos a seguir.

A. Método Utilizado para Revisão da Literatura

O método de levantamento bibliográfico utilizado adotou mapeamento de literatura empregada por [9], consistindo em seis passos: I) Identificação da pesquisa; II) Estratégia de pesquisa; III) Critérios de Elegibilidade; IV) Seleção dos estudos; V) Extração de informações; e VI) Síntese.

Neste mapeamento, utilizou-se as seguintes palavras-chave para identificação dos trabalhos: 1. critical infrastructure AND (vulnerabilit* AND (analys* OR Assessment OR characterization OR scan* OR tool)) AND (threat AND (shar* OR intelligent)); 2. (network AND (Security OR “Vulnerabilit* Tool” OR scan*)) AND (shar* AND (threat OR risk OR intelligent OR information)) AND (risk AND (Classification OR reduction OR mitigation)); e 3. (threat AND (shar* OR intelligent)) AND (test AND (intrusion OR penetration)). Após a execução dos cinco passos iniciais, aplicou-se os seguintes critérios de aceitação para trabalhos correlatos: 1. constituir ferramenta, método, *framework* ou produto; 2. ter aplicabilidade, direta ou indireta, em Rede de Comunicação de Dados de Infraestrutura Crítica; e 3. envolver medidas preventivas de Segurança Cibernética.

B. Síntese dos Principais Trabalhos Correlatos Encontrados

A partir da aplicação do Método de Revisão Sistemática da Literatura utilizando Inteligência de Ameaças associada à Automação, foram estudados com maior profundidade, os seguintes trabalhos correlatos, com suas respectivas contribuições.

[15] concebe um *framework* composto por 3 fases: Teste de Intrusão de Caixa Preta, Teste de Intrusão de Caixa Branca e Implementação de Contramedidas, com o propósito de identificar e aplicar contramedidas em vulnerabilidades e ameaças de Segurança Cibernética em Computação em Nuvem, incluindo ameaças na rede de informação.

Em sua abordagem, [16] adota um esquema de gerenciamento seguro dividido em três partes: busca de vulnerabilidades, coleta de logs e análise de correlação, a fim de identificar, nos ambientes de Computação em Nuvem, potenciais perigos e ataques de agentes maliciosos, emitindo alerta neste último caso.

Considerando os modernos sistemas em rede, [17] aponta uma Análise de Segurança abrangente por meio de *Unified Vulnerability Risk Analysis Module* (UV-RAM). Esta arquitetura proposta constitui-se de três funcionalidades principais: descobrir caminhos de ataque, incorporar vulnerabilidades na Análise de Segurança e formular estratégias de mitigação para fortalecer o sistema em rede.

No contexto de sistemas de infraestruturas críticas, [18] aperfeiçoa seu método de Avaliação de Risco, abordado em outro artigo, de forma a mapear etapas de avaliação e mitigação contidas no Guia NIST para Segurança de Sistemas de Controle Industrial (ICS), Publicação Especial 800-82r2. Para fins de exemplificação, este trabalho apresenta um arquitetura fictícia, com a aplicação do método proposto em um estudo de caso, denominado System Under Study (SUS).

Neste ponto de nossa pesquisa, constatou-se que todas as abordagens estudadas possuíam algumas lacunas importantes a serem preenchidas para a melhoria do processo de filtragem de alarmes e posterior falta de automação da proteção, a partir das fontes de inteligência. Encaixando-se, portanto, este trabalho de pesquisa nestes aspectos.

IV. MEGA-RCD: O MÉTODO PROPOSTO

Uma tarefa crítica para a manutenção de sistemas tem sido a prevenção de exploração de vulnerabilidades [5]. Além disso, com o desenvolvimento de novas tecnologias, os aumentos de conectividade de sistemas e de suas complexidades têm tornado a prevenção de exploração de vulnerabilidades ainda mais difícil [11]. Esta exploração de vulnerabilidades vem gerando incidentes cibernéticos, que ocorrem devido às instituições terem dificuldades de prover Segurança Cibernética aos seus ativos computacionais [9].

Como medida preventiva, vêm sendo utilizadas ferramentas de verificação de vulnerabilidades, que permitem detectar falhas em diferentes partes e tipos de aplicativos. Destes, existem dois modos distintos de busca: interna e/ou externa [8].

Deste modo, o Método de Gerenciamento de Ameaças em Rede de Comunicação de Dados (MEGA-RCD) emergiu da necessidade em se identificar vulnerabilidades de forma automatizada e integrada a uma Plataforma de Inteligência de Ameaças, a fim de aumentar a Segurança Cibernética de sistemas.

Essa demanda por identificar vulnerabilidades representa um dos maiores riscos de segurança do Projeto de Segurança Web Aberto (em inglês, *Open Web Application Security Project* - OWASP). Nossa observação chave é que existe ganhos na integração de prevenção automatizada de incidentes cibernéticos através de Plataformas de Orquestração de Segurança combinadas com a gestão de Inteligência de Ameaças [19].

O MEGA-RCD, que se encontra ainda em fase preliminar de elaboração, consiste de 5 passos, integráveis a uma Plataforma de Orquestração Automatizada e Resposta de Segurança e a sistemas de busca de vulnerabilidades.

Ele envolve coleta e processamento de bases de Inteligência de Ameaças e fontes estruturadas e semiestruturadas de vulnerabilidades, que visam possibilitar os processos de identificação e mitigação de ameaças no espaço cibernético e, posteriormente, apresentar uma Consciência Situacional de ameaças existentes.

O método MEGA-RCD possui os passos descritos a seguir: Coleta de Dados; Normalização; Processamento e Mitigação; Conformidade; e Consciência Situacional, conforme seu modelo conceitual mostrado na Figura 2 e suas dependências.

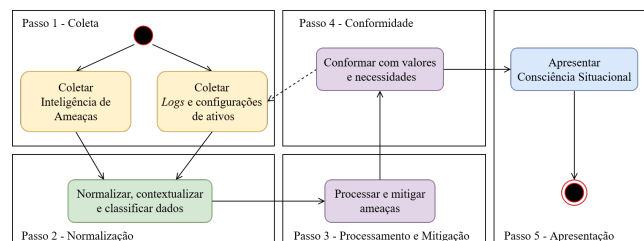


Fig. 2. Modelo Conceitual do MEGA-RCD.

A. Passo 1 - Coleta

Este passo representa a atividade de obtenção de informações de diversas fontes, através de ferramentas de OSINT, busca de vulnerabilidades, entre outras. Ele visa a identificação de sistemas operacionais e outras fontes de vulnerabilidades bem como a obtenção de informações de ameaças cibernéticas para uma Inteligência de Ameaças customizada com a instituição.

A sua validação experimental vem sendo realizada com o uso da ferramenta PyMISP, para coleta de Inteligência de ameaças, conforme mostrado na Figura 3. Neste passo, utiliza-se a interface de busca de CVE da CIRCL, endereço “<https://cve.circl.lu/api/cve/>”, trazendo-se detalhes a respeito das vulnerabilidades conhecidas e dos locais onde residem as sincronizações com as Bases de Vulnerabilidades NVD e CIRCL, em sua totalidade, para a obtenção da Base de Fraquezas conhecidas e da Base de códigos de exploração de vulnerabilidades, também conhecida como *Exploit Database*.

Outro procedimento, que vem sendo experimentado neste passo, consiste em buscas pró-ativas e automáticas de vulnerabilidades em sistemas, adquiridas pela inteligência, envolvendo alvos para as ferramentas NMAP e OpenVas. A pretensão arquitetural é de que tais componentes sejam realizados em contêineres e executados de forma remota, via protocolo SSH, possibilitando sua flexibilização de localização na infraestrutura de rede, seja interna ou externa.

Os dados coletados neste passo até o momento se encontram armazenados em variáveis ou arquivos temporários em formato csv, a serem consumidos pelos passos posteriores. O gargalo nos experimentos encontra-se relacionado à velocidade da rede, dado que sempre tratamos com dados atualizados em cada ciclo de execução.

B. Passo 2 - Normalização, Contextualização e Classificação

Este passo compreende em se normalizar os dados coletados de Inteligência de Ameaças e Bases de Conhecimento, ora armazenados em variáveis ou arquivos temporários, correlacionar e enriquecer a Base de Inteligência de Ameaças, através de ferramenta PyMISP.

Passo 1: Coleta

```
***
```

MISP

Busca por IDs de CVE na Plataforma de Inteligência de Ameaças MISP, com "Metasploit exploits" disponíveis.

```
[ ]: relative_path = 'attributes/restSearch'
body = {
    "returnFormat": "csv",
    "type": "vulnerability"
}
dados = misp.py_misp.direct_call(relative_path, body)
misp_dados = pd.read_csv(StringIO(dados))
#misp_dados.to_csv(out_misp, index=False)
```

```
[ ]: #misp_dados = pd.read_csv(out_misp)
misp_dados
```

Bases de Vulnerabilidades

Atualizar Bases de Vulnerabilidades

```
***
```

Busca de detalhes das IDs, extraídas do MISP

```
[ ]: list_cve = []
data_cve = []
y = 0
try:
    dados = pd.read_csv(out_file_cve)
    for x in dados.index:
        list_cve.append(dados.iloc[x].loc['id'])
except IOError:
    print("I/O error")

[ ]: # Search CVE via https://cve.circl.lu/
#misp_dados = pd.read_csv('misp.csv')
for x in misp_dados.index:
    try:
        if(misp_dados.iloc[x].loc['value'] not in list_cve):
            print( str(x) + " - ", end="", flush=True)
            w = (cve.search_cve_data(misp_dados.iloc[x].loc['value']))
            if(w != None):
                data_cve.append(w)
                list_cve.append(misp_dados.iloc[x].loc['value'])
            if((x % 300) == 0):
                nome = 'test_' + str(y) + '.csv'
                mf = pd.DataFrame.from_dict(data_cve)
                mf.to_csv(nome, index=False, mode="a") #, header=False
                data_cve = []
                y = y + 1
    except IOError:
        print("I/O error")
```

Busca de Informações

```
[ ]: # Declaração do(s) alvo(s):
target = ['192.168.0.1']

[ ]: # NMAP
cmd1 = 'nmap -p0-65535 -T5 -A -Pn -sV -sT --script=vuln --min-rate 1000 -oX'
for alvo in target:
    cmd = cmd1 + ' ' + alvo + '.xml ' + alvo
    os.system(cmd)

[ ]: # OpenVias
ssh_stdin, ssh_stdout, ssh_stderr = ssh.exec_command(cmd2)
```

Passo 2: Normalização, contextualização e classificação

Alguns aspectos de normalização são realizados durante a coleta, afim de simplificar o código.

```
***
```

Enriquecer MISP

```
[ ]: # Adicionar tag
misp.py_misp.add_tag({"name": "Exploit Available",
    "colour": "#22681c",
    "exportable": "true",
    "org_id": "0",
    "user_id": "0",
    "hide_tag": "false"})

# Aplicar tag em atributos ou eventos no MISP
name_tag = "Exploit Available"
for x in misp_dados.index:
    y = misp_dados.iloc[x].loc['uid']
    misp.py_misp.tag(misp_entity=y, tag=name_tag)
```

```
***
```

Busca por vulnerabilidades

```
[3]: from bs4 import BeautifulSoup
port_list = []
proto_list = []
info_list = []
cpe_list = []

[4]: # A partir da correlacionamento das informações obtidas em Busca de Informações e na
# Base de Inteligência de Ameaças, obtém-se as vulnerabilidades conectadas.
for alvo in target:
    cmd = 'target.' + alvo + '.xml'
    with open(cmd, "r") as f:
        nmap_result = f.read()
        bs = BeautifulSoup(nmap_result, 'lxml')
        ports = bs.find_all('port')
        for idx, port in enumerate(ports):
            port_list.append(str(port.attrs['portid']))
            proto_list.append(port.attrs['protocol'])
            for obj_child in port.contents:
                if obj_child.name == 'service':
                    if obj_child.cpe != None:
                        cpe_list.append(str(obj_child.cpe))
                    else:
                        cpe_list.append('unknown')
                if 'os_type' in obj_child.attrs:
                    os_type = obj_child.attrs['os_type']
                    temp_info = ''
                    if 'name' in obj_child.attrs:
                        temp_info += obj_child.attrs['name'] + ' - '
                    if 'product' in obj_child.attrs:
                        temp_info += obj_child.attrs['product'] + ' - '
                    if 'version' in obj_child.attrs:
                        temp_info += obj_child.attrs['version'] + ' - '
                    if 'extrainfo' in obj_child.attrs:
                        temp_info += obj_child.attrs['extrainfo']
                    if temp_info != '':
                        info_list.append(temp_info)
                    else:
                        info_list.append('unknown')
```

```
***
```

```
[5]: data = {'port': port_list, 'protocol': proto_list, 'info': info_list, 'cpe': cpe_list}
pd.DataFrame.from_dict(data)
```

	port	protocol	info	cpe
0	22	tcp	ssh - OpenSSH - 7.9p1 Debian 10-deb10u2 - prot...	<cpe>cpe:/a:openssh:openssh:7.9p1</cpe>
1	2377	tcp	swarm -	unknown
2	3306	tcp	mysql - MySQL - 5.7.30 -	<cpe>cpe:/a:mysql:mysql:5.7.30</cpe>
3	7946	tcp	unknown -	unknown
4	8000	tcp	nagios-nrca - Nagios NSCA -	unknown
5	8080	tcp	http - Tomado httpd - 6.0.3 -	<cpe>cpe:/a:tomadotomado:tomado:6.0.3</cpe>
6	8085	tcp	http - Apache httpd - 2.4.25 - (Debian)	<cpe>cpe:/a:apache:http_server:2.4.25</cpe>
7	8090	tcp	http - Apache httpd - 2.4.29 - (Ubuntu)	<cpe>cpe:/a:apache:http_server:2.4.29</cpe>
8	9000	tcp	clitester -	unknown

Fig. 3. Validação Experimental MEGA-RCD.

Através da coleta de informações dos alvos e das informações enriquecidas na Base de Inteligência de Ameaças, propicia-se condições para identificação de vulnerabilidades nos alvos, em especial neste passo, focaremos em ameaças de redes de computadores. Neste sentido, pretende-se utilizar métodos de classificação para caracterização de severidade e atualidade das vulnerabilidades dos alvos. Ao final deste passo, os dados normalizados e contextualizados encontram-se publicados na Base de Inteligência de Ameaças e os demais acham-se armazenados em variáveis.

C. Passo 3 - Mitigação de ameaças

Neste passo, apontam-se medidas para sanar as vulnerabilidades. Pretende-se realizá-lo, inicialmente, verificando-se a disponibilidade de processos de mitigação conhecidos, através de Processamento de Linguagem Natural (*Natural Language*

Processing - NLP) nos textos contidos nos CVEs e, caso não seja encontrado, sugere-se, em trabalhos futuros, a coleta de dados em bases não estruturadas, como *Blogs* e *Twitter*.

Medidas de mitigação iniciais envolvem a classificação e verificação de pertinência em processos de atualização de segurança de pacotes, criação ou edição de regras de segurança ou desabilitar/substituir serviços vulneráveis.

Um algoritmo de aprendizagem de máquina promissor para essa tarefa tem sido a árvore de decisão, por facilitar o entendimento de tomadas de decisão, porém, pretende-se verificar também outros métodos. O desfecho deste passo se dá no envio de comandos, via acesso remoto, ao alvo e aos equipamentos de segurança, bem como o reporte de vulnerabilidades encontradas, com a classificação de processos de mitigação, a árvore de abordagem (vetor de acesso) das vulnerabilidades do alvo, entre outras informações.

D. Passo 4 - Conformidade

Após o processo de mitigação, faz-se necessária a verificação de sensibilização das medidas mitigatórias, tanto nas funcionalidades exercidas pelo alvo, quanto na correção de vulnerabilidades. Assim, este passo consisti em realizar a atualização e comparação dos dados de coleta dos alvos, antes e depois do passo três, por meio das ferramentas utilizadas no passo um.

Neste passo, caso haja uma vulnerabilidade apontada como explorável, é realizado o Testes de Intrusão, por meio de ferramentas como o *framework* de código aberto Metasploit, utilizado por Takaesu, em sua ferramenta DeepExploit [20]. Ao término deste passo, ocorre a atualização de variáveis e Bases de Inteligência de Ameaças, bem como a pontuação de processos mitigatórios considerados eficazes.

E. Passo 5 - Apresentação

Finalmente, apresenta-se um relatório das atividades desenvolvidas pelo método, bem como o detalhamento dos sistemas alvos com respectivos relatórios e árvores de falhas, constando as vulnerabilidades restantes e os vetores de acesso.

Assim, este passo proporciona a visualização do nível de comprometimento dos alvos propostos da infraestrutura. Deste modo, facilitando-se tomadas de decisão em relação à segurança cibernética.

V. CONCLUSÃO

Esta pesquisa objetiva investigar o desenvolvimento de um modelo de gerenciamento contínuo de ameaças para Redes de Comunicação de Dados (RCD), visando aumentar sua eficiência, segurança, consciência situacional e reduzir ameaças cibernéticas.

Neste sentido, esta pesquisa, que ainda se encontra em estágio posicional, evidencia aplicabilidade prática das seguintes contribuições: 1) MEGA-RCD - Um Método de Gerenciamento de Ameaças em RCD; 2) Orquestração e automatização de mitigação de vulnerabilidades com Inteligência de Ameaças; 3) Redução de tempo de mitigação de vulnerabilidades conhecidas; 4) Uma forma de se enriquecer os dados da base de Inteligência de Ameaças de modo automatizado; 5) Uma forma de se depreender quais dos ativos possuem vulnerabilidades exploráveis por ferramentas; 6) A possibilidade de sua padronização para emprego real, com uma capacitação mínima, em vários níveis de necessidades; e 7) A possibilidade de sua aplicação prática com regras de segurança.

Durante os testes práticos desta pesquisa em *notebook*, notou-se a possibilidade de refinamento de algumas investigações adicionais nas áreas de identificação de Indicadores de Compromisso, possibilitando-se assim melhorias na detecção de intrusões. Recomenda-se a utilização de outras bases de conhecimento estruturadas para o enriquecimento da Plataforma de Inteligência de Ameaças e a realização de estudos de caso com o MEGA-RCD em sistemas em produção.

Como sugestões para trabalhos futuros, ao longo desta pesquisa, foram encontradas algumas questões ainda em aberto, julgadas pertinentes, porém que fogem ao seu escopo, como por exemplo, no passo de mitigação de ameaças, onde se sugere a utilização de bases não estruturadas de conhecimento para o enriquecimento da base de Inteligência de Ameaças com o *Twitter* e *Blogs*.

Outra sugestão para trabalhos futuros seria no sentido de se melhorar a Consciência Situacional com avaliação de riscos envolvendo Inteligência Artificial, envolvendo-se Lógica Fuzzy, *Machine Learning*, entre outras técnicas.

REFERÊNCIAS

- [1] BRASIL, "Decreto nº 10.222, de 5 de fevereiro de 2020. Aprova a Estratégia Nacional de Segurança Cibernética," pp. 6–17, 2020.
- [2] —, "Decreto nº 10.332, de 28 de abril de 2020. Institui a Estratégia de Governo Digital para o período de 2020 a 2022," pp. 6–8, 2020.
- [3] —, "Portaria nº 93, de 29 de setembro de 2019. Aprova o Glossário de Segurança da Informação," pp. 3–10, 2019.
- [4] J. Giraldo, E. Sarkar, A. A. Cardenas, M. Maniatakos, and M. Kantarcioglu, "Security and Privacy in Cyber-Physical Systems: A Survey of Surveys," *IEEE Design and Test*, vol. 34, no. 4, pp. 7–17, 2017.
- [5] D. Gonzalez, H. Hastings, and M. Mirakhorli, "Automated Characterization of Software Vulnerabilities," *Proceedings - 2019 IEEE International Conference on Software Maintenance and Evolution, ICSME 2019*, pp. 135–139, 2019.
- [6] D. Adinolfi and A. Singleton, "CVE IDs and How to Get Them," p. 20, 2017. [Online]. Available: <https://cve.mitre.org/CVEIDsAndHowToGetThem.pdf>
- [7] V. Mavroeidis and S. Bromander, "Cyber Threat Intelligence Model: An Evaluation of Taxonomies, Sharing Standards, and Ontologies within Cyber Threat Intelligence," in *2017 European Intelligence and Security Informatics Conference (EISIC)*, vol. 2017-Janua. IEEE, sep 2017, pp. 91–98. [Online]. Available: <http://ieeexplore.ieee.org/document/8240774/>
- [8] K. Kritikos, K. Magoutis, M. Papoutsakis, and S. Ioannidis, "A survey on vulnerability assessment tools and databases for cloud-based web applications," *Array*, vol. 3-4, no. June, p. 100011, 2019. [Online]. Available: <https://doi.org/10.1016/j.array.2019.100011>
- [9] C. Islam, M. A. Babar, and S. Nepal, "A multi-vocal review of security orchestration," *ACM Computing Surveys*, vol. 52, no. 2, 2019.
- [10] J. d. C. Martins, "MAVARS - Um Método Ágil Para Validação de Requisitos de Segurança," Dissertação de Mestrado, Instituto Tecnológico de Aeronáutica, 2018.
- [11] W. Tounsi and H. Rais, "A survey on technical threat intelligence in the age of sophisticated cyber attacks," *Computers and Security*, vol. 72, pp. 212–233, 2018. [Online]. Available: <https://doi.org/10.1016/j.cose.2017.09.001>
- [12] N. Moustafa, E. Adi, B. Turnbull, and J. Hu, "A New Threat Intelligence Scheme for Safeguarding Industry 4.0 Systems," *IEEE Access*, vol. 6, pp. 32 910–32 924, 2018. [Online]. Available: <https://ieeexplore.ieee.org/document/8374422/>
- [13] C. Wagner, A. Dulaunoy, G. Wagener, and A. Iklody, "Misp: The design and implementation of a collaborative threat intelligence sharing platform," in *Proceedings of the 2016 ACM on Workshop on Information Sharing and Collaborative Security*. ACM, 2016, pp. 49–56.
- [14] J. Pastor-Galindo, P. Nespoli, F. Gomez Marmol, and G. Martinez Perez, "The Not Yet Exploited Goldmine of OSINT: Opportunities, Open Challenges and Future Trends," *IEEE Access*, vol. 8, pp. 10 282–10 304, 2020. [Online]. Available: <https://ieeexplore.ieee.org/document/8954668/>
- [15] R. M. Jabir, S. I. R. Khanji, L. A. Ahmad, O. Alfandi, and H. Said, "Analysis of cloud computing attacks and countermeasures," in *2016 18th International Conference on Advanced Communication Technology (ICACT)*. IEEE, jan 2016, pp. 117–123. [Online]. Available: <http://ieeexplore.ieee.org/document/7423296/>
- [16] P. Y. Wang and M. Q. Hong, "A Secure Management Scheme Designed in Cloud," *Proceedings - 2nd IEEE International Conference on Big Data Security on Cloud, IEEE BigDataSecurity 2016, 2nd IEEE International Conference on High Performance and Smart Computing, IEEE HPSC 2016 and IEEE International Conference on Intelligent Data and S*, pp. 158–162, 2016.
- [17] J. B. Hong and D. S. Kim, "Discovering and Mitigating New Attack Paths Using Graphical Security Models," *Proceedings - 47th Annual IEEE/IFIP International Conference on Dependable Systems and Networks Workshops, DSN-W 2017*, pp. 45–52, 2017.
- [18] A. A. Jillepalli, F. T. Sheldon, D. C. de Leon, M. Haney, and R. K. Abercrombie, "Security management of cyber physical control systems using NIST SP 800-82r2," in *2017 13th International Wireless Communications and Mobile Computing Conference (IWCMC)*. IEEE, jun 2017, pp. 1864–1870. [Online]. Available: <http://ieeexplore.ieee.org/document/7986568/>
- [19] The OWASP Foundation, "OWASP Top Ten," 2017. [Online]. Available: <https://owasp.org/www-project-top-ten/>
- [20] I. Takaesu, "DeepExploit," 2019. [Online]. Available: https://github.com/130-bbr-bbq/machine_learning.security/tree/master/DeepExploit