



**YIĞIN TOPLULUK ÖĞRENME TEMELLİ
OLTALAMA SALDIRILARI TESPİT SİSTEMİ**

YÜKSEK LİSANS TEZİ

Ramazan Sami ÇINAR

Danışman

Dr. Öğr. Üyesi Ahmet Haşim YURTTAKAL

BİLGİSAYAR ANABİLİM DALI

Temmuz 2024

AFYON KOCATEPE ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ

YÜKSEK LİSANS TEZİ

YIĞIN TOPLULUK ÖĞRENME TEMELLİ OLTALAMA
SALDIRILARI TESPİT SİSTEMİ

Ramazan Sami ÇINAR

Danışman

Dr. Öğr. Üyesi Ahmet Haşim YURTTAKAL

BİLGİSAYAR ANABİLİM DALI

Temmuz 2024

BİLİMSEL ETİK BİLDİRİM SAYFASI

Afyon Kocatepe Üniversitesi

Fen Bilimleri Enstitüsü, tez yazım kurallarına uygun olarak hazırladığım bu tez çalışmada;

- Tez içindeki bütün bilgi ve belgeleri akademik kurallar çerçevesinde elde ettiğimi,
- Görsel, işitsel ve yazılı tüm bilgi ve sonuçları bilimsel ahlak kurallarına uygun olarak sunduğumu,
- Başkalarının eserlerinden yararlanılması durumunda ilgili eserlere bilimsel normlara uygun olarak atıfta bulunduğumu,
- Atıfta bulunduğum eserlerin tümünü kaynak olarak gösterdiğimi,
- Kullanılan verilerde herhangi bir tahrifat yapmadığımı,
- Ve bu tezin herhangi bir bölümünü bu üniversite veya başka bir üniversitede başka bir tez çalışması olarak sunmadığımı

beyan ederim.

29 / 07 / 2024

İmza

Ramazan Sami ÇINAR

ÖZET

Yüksek Lisans Tezi

YIĞIN TOPLULUK ÖĞRENME TEMELLİ OLTALAMA SALDIRILARI TESPİT SİSTEMİ

Ramazan Sami ÇINAR

Afyon Kocatepe Üniversitesi

Fen Bilimleri Enstitüsü

Bilgisayar Anabilim Dalı

Danışman: Dr. Öğr. Üyesi Ahmet Haşim YURTTAKAL

Oltalama, masum web kullanıcılarını görsel olarak gerçek muadiline benzeyen sahte bir web sitesi ile tuzağa düşürmeyi amaçlayan bir siber suçtur. Başlangıçta, kullanıcılar çeşitli sosyal ve teknik yönlendirme metodları aracılığıyla oltalama web sitelerine yönlendirilir. Web sitesinin sahteliğinde habersiz olan kullanıcılar kullanıcı kimliği, parola, kredi kartı bilgileri veya banka hesap bilgileri gibi kişisel bilgilerini bu sitelere verebilir. Oltalama saldırırganları bu tür bilgileri kullanır bankalardan para çalmak, bir markanın imajına zarar vermek ve hatta taahhütte bulunmak gibi ağır suçlara karışır. Mevcut literatürde birçok oltalama tespit ve önleme tekniği bulunsa da makine öğrenimi ve derin öğrenme yöntemlerinin ortaya çıkışı siber güvenlik dünyasında bu alanın kapsamını genişletti. Bu nedenle, oltalama web sitelerini verimli bir şekilde tespit eden bir sistemin geliştirilmesine güçlü bir ihtiyaç var. Pek çok çalışmada k en yakın komşu, lojistik regresyon, rasgele orman, destek vektör makineleri, karar ağaçları gibi yöntemler kullanılırken çalışmada yığın topluluk öğrenme yöntemi önerilmiştir. Yığın topluluk öğrenme tabanlı oltalama tespit yöntemi %96,6'lık doğruluk göstermiştir.

2024, VIII + 60 sayfa

Anahtar Kelimeler: Oltalama, Topluluk öğrenme, Torbalama, Arttırma, Oylama.

ABSTRACT

M.Sc. Thesis

STACK ENSEMBLE LEARNING BASED PHISHING ATTACKS DETECTION SYSTEM

Ramazana Sami ÇINAR

Afyon Kocatepe University

Graduate School of Natural and Applied Sciences

Department of Computer

Supervisor: Asst. Prof. Ahmet Haşim YURTTAKAL

Phishing is a cybercrime that aims to trap innocent web users with a fake website that visually resembles its real counterpart. Initially, users are directed to phishing websites through various social and technical redirection methods. Users who are unaware of the fakeness of the website may provide personal information such as user ID, password, credit card information, or bank account information to these sites. Phishing attackers use such information to commit serious crimes such as stealing money from banks, damaging the image of a brand, or even committing a crime. Although there are many phishing detection and prevention techniques in the existing literature, the emergence of intelligent machine learning and deep learning methods has expanded the scope of this field in the cybersecurity world. Therefore, there is a strong need to develop a system that efficiently detects phishing websites. While many studies use methods such as k nearest neighbor, logistic regression, random forest, support vector machines, decision trees, the stack ensemble learning method was proposed in this study. The stack ensemble learning based phishing detection method showed 96.6% accuracy.

2024, VIII + 60 pages

Keywords: Phishing, Ensemble learning, Bagging, Boosting, Voting.

TEŞEKKÜR

Tez çalışmamı hazırlarken zamanını bana ayıran, tecrübesi ve bilgisiyle bana yol gösterip her türlü desteği veren sevgili danışman hocam Dr. Öğr. Üyesi Ahmet Haşim YURTTAKAL'a saygılarımı sunarım.

Tez çalışmam süresince desteklerini esirgemeyen mesai arkadaşlarıma teşekkürlerimi sunarım.

Eğitim hayatım süresince maddi manevi her konuda desteklerini esirgemeyen ve beni bu günlere getiren kıymetli aileme şükranlarımı sunarım.

Ramazan Sami ÇINAR
Afyonkarahisar 2024

İÇİNDEKİLER DİZİNİ

	Sayfa
ÖZET	i
ABSTRACT	ii
TEŞEKKÜR	iii
İÇİNDEKİLER DİZİNİ.....	iv
SİMGELER ve KISALTMALAR DİZİNİ	vi
ÇİZELGELER DİZİNİ.....	vii
RESİMLER DİZİNİ	viii
1. GİRİŞ	1
1.1 Problem Durumu	1
1.2 Tezin Organizasyonu	1
2. LİTERATÜR BİLGİLERİ.....	3
3. MATERYEL ve METOT.....	32
3.1 Yapay Zeka	32
3.2 Makine Öğrenmesi.....	33
3.2.1 Denetimli Öğrenme	34
3.2.2 Denetimsiz Öğrenme	35
3.2.3 Takviyeli Öğrenme	36
3.3 Topluluk Öğrenme.....	37
3.4 Topluluk Öğrenme Yöntemleri.....	38
3.4.1 Bagging.....	38
3.4.2 Boosting.....	38
3.4.3 Stacking	38
3.5 Oltalama Saldırıları.....	39
3.5.1 Oltalama Yöntemleri	40
3.5.1.1 Bağlantı Manipülasyonu.....	40
3.5.1.2 Filtre Atlatma.....	40
3.5.1.3 Web Sayfası Sahteciliği.....	40
3.5.1.4 Gizli Yönlendirme	40
3.5.1.5 Sosyal Mühendislik	41
3.5.2 Oltalama Tespit Yöntemleri	41
3.5.2.1 Kullanıcı Eğitimi	41

3.5.2.2	Liste Temelli Yaklaşım	41
3.5.2.3	Görsel Benzerlik Temelli Yaklaşım	42
3.5.2.4	Davranışsal ve Makine Öğrenmesi Temelli Yaklaşım	42
4.	BULGULAR.....	43
5.	TARTIŞMA ve SONUÇ.....	47
6.	KAYNAKLAR	49
	ÖZGEÇMİŞ.....	60



SİMGELER ve KISALTMALAR DİZİNİ

Kısaltmalar

GPU	Cross-Site Scripting
HTML	HyperText Markup Language
HTTP	Hyper-Text Transfer Protocol
URL	Uniform Resource Locator
URI	Uniform Resource Identifier
GPU	Graphics processing unit
MTA	Mail Transfer Agent
MUA	Mail User Agent
AOL	American Online
OCR	Optical Character Recognition
AC	Associative Classification
MCAC	Multi-label Classifier based Associative Classification
CSS	Cascading Style Sheets
W3C	World Wide Web Consortium
PII	Personally Identifiable Information
DBN	Deep Belief Network
ISP	Internet Service Provider
IP	Internet Protocol Address
SCT	Social Cognitive Theory
CRB	Cyber Risk Belief

ÇİZELGELER DİZİNİ

	Sayfa
Çizelge 4.1 Veri kümesinden örnekler	43
Çizelge 4.2 Veri kümesinin gruplara göre örneklem sayıları.....	44
Çizelge 4.3 Veri kümesinin gruplara göre örneklem sayıları.....	46



RESİMLER DİZİNİ

	Sayfa
Resim 3.1 Derin öğrenme ve yapay zeka arasındaki ilişki	34
Resim 3.2 Denetimli öğrenme şeması.....	35
Resim 3.3 Denetimsiz öğrenme şeması	36
Resim 3.4 Takviyeli öğrenme şeması	37
Resim 3.5 Topluluk öğrenme iş akışı.....	37
Resim 4.1 Veri kümesinin dağılımı	43
Resim 4.2 Önerilen sınıflandırma modeli	44
Resim 4.3 Karışıklık matrisi	45
Resim 4.4 ROC Eğrisi.....	45

1. GİRİŞ

1.1 Problem Durumu

Hızla artan internet kullanımı ve internet üzerinde paylaşılan bilgilerin artması ile birlikte, bu alan suçlularında dikkatini çekmeye başlamıştır. Online bilgilerimizin içine finansal bilgilerimizin de dahil olması ile suçlular için büyük bir sektör oluşturmuştur (Gupta vd. 2016).

Hızlı gelişen internet teknolojileri pek çok kullanıcının aslında gayet karmaşık ve birbirine güven üzerine dayalı olan internet dünyasına dikkat edilmesi gerekirken pek çok hususu öğrenmeden girmesine sebep olmuştur. Gizlilik, veri güvenliği, dolandırıcılığa karşı koyma gibi günlük yaşamda da kaşımaıza çıkan pek çok konuda dikkatli olan kullanıcılar internete hızlı adımlarla girdiği için bu dünyayı gerçek dünyadan bağımsız ve temiz bir alan olarak gördü. Aslında mevcut dünyanın dijital bir kopyası olan internet bilgisiz kullanıcı için gerçek dünyadan daha tehlikeli bir alana oluşturmaya başlamıştır (Basit vd. 2021).

Tehlikeli olmasına rağmen işlerin çoğunluğunun artık online sistemler üzerinden yapılması sebebiyle bu alandan uzak durmak mümkün ve mantıklı değil. Bu sebeple kullanıcı bilinçlendirme ve eğitimlerin yanı sıra hem sistemleri, hem kullanıcıları koruyacak akıllı tespit ve önleme sistemlerine duyulan ihtiyaç artıyor (Garera vd. 2007).

Bu tür sistemler geçmişte kullanılan kelime filtreleri, url listeleri ve davranışsal analiz gibi klasik yöntemler başarılı olamıyor çünkü saldırganlar her alanda olduğu gibi bu alanda da sürekli gelişim göstererek, siber güvenlik alanında çalışanlardan bir adım önde oluyor. Bu sebeple bu tür sistemlerin geleceği yapay zeka destekli, hızlı ve başarı oranı yüksek algoritmalarda görülmektedir (Basit vd. 2021).

1.2 Tezin Organizasyonu

Çalışmamızın ikinci bölümü alana katkı sağlayan literatürün taramasını barındırmaktadır. Materyal ve Metot kısmında makine öğrenmesi, topluluk öğrenme ve topluluk öğrenme

yöntemleri gibi kullanılan yöntemlerin ve kavramların tanımları yapılmıştır. Ayrıca ortalama saldırıları ve ortalama yöntemleri de bu bölümde tanımlanmıştır. Dördüncü bölüm, Bulgular kısmında kullanılan veri setinin nitelikleri, sınıflandırma için kullanılan yöntem ve yöntemin performansı yer almaktadır. Son olarak tartışma ve sonuç kısmında çalışma sonucu ve öneriler yer almaktadır.



2. LİTERATÜR BİLGİLERİ

Yapay zeka literatürünün uzun yılları kapsayan bir tarihi olmasına rağmen son yıllarda hızla artan çalışmalar bu alana olan ilginin ve ihtiyacın artışı göstermektedir. Bu durum düşünülerek çalışmamız için son yıllara ağırlık verilerek ve makine öğrenmesi konusu dikkate alınarak kapsamlı bir literatür taraması yapılmıştır.

Banu ve Banu (2013) çalışmasında günümüzün veri hırsızlığı sahteciliğin en çok kullanılan yöntemlerinden oltalamayı ve oltalamaya karşı koyma yöntemlerini incelemiştir. Kullanıcı adı ve parola bilgilerinin kötüye kullanımın endüstrinin birçok alanını etkilediğini bu sebeple oltalamaya karşı koyma için birçok teknik geliştirildiğini fakat henüz hiçbirinin tam olarak başarıya ulaşmadığını ortaya koymuştur.

Basit vd. (2021) çalışmasında yapay zekâ destekli oltalama saldırısı tespit yöntemlerini, oltalama saldırılarının mevcut zorluklarını ve bu alandaki gelecek yönelimlerini inceliyor. Çalışma zaman içinde oltalama saldırılarının internet kullanıcılarının, hükümetlerin ve hizmet sağlayıcıların karşılaştığı en öne çıkan saldırı tipi olduğunu ortaya koyuyor. Çalışma oltalama saldırıların gerçekleşme şeklini inceledikten sonra, bu tip saldırıları tespit etmek için kullanılan makine öğrenmesi, derin öğrenme, hibrit öğrenme ve senaryo bazlama öğrenme gibi yapay zeka tekniklerinin üzerine bir literatür araştırması sunuyor.

Garera vd. (2007) çalışmasında oltalama saldırısı tespiti ve ölçümü için bir çatı sunuyor. Yazarlar oltalama saldırısını masum kullanıcıların finansal bilgilerini çalmak için sosyal mühendislik teknikleri ve sofistike saldırı vektörlerini birleştirilmesi olarak tanımlıyor. Bu tip saldırılarda genel olarak kullanılan URL'lerin yapısı üzerine bir inceleme yapılıyor. Bu inceleme sonucunda çoğunlukla sayfa ile ilgili başka bir bilgiye ihtiyaç duymadan sitelerin oltalama saldırısına ait olup olmadığını söylemenin mümkün olduğunu ileriye sürüyor. Çalışma oltalama URL'lerini zararsız olanlardan ayıracak bazı özelliklere yer veriyor.

Gupta vd. (2016) çeşitli oltalama saldırılarını avantajları ve dezavantajları açısından analiz ediyor. Oltalama saldırısını kullanıcıların bilgilerine ulaşmak için gerçek bir web sitesinin sahtesini hazırlamak olarak tanımlayan yazarlar çalışmada asıl olarak bir sosyal

mühendislik tekniği olan ortalama saldırıları ve tespit yöntemleri hakkında literatür araştırması ortaya koyuyor. Makale ortalama saldırılarının teorisi ve insan varlığı ile ilgili sorunları üzerinde duruyor. Çalışma sosyal mühendislik, internet güvenliği ve kendisine ait hesap ve finansal bilgiler hakkında bilgi sahibi olmayan kullanıcıları hedefleyen saldırıların genellikle e-posta ve anlık mesajlaşma ile uygulandığını ortaya koyuyor. Bu saldırıların mali kayıplara yol açtığı üzerinde durulan çalışmada, karşı olarak alınan önlemlerden de bahsediliyor. Bu karşı yöntemlerin artan saldırılar karşısında yetersiz kaldığı üzerinde duruluyor. Çalışma kimlik sahteciliği yapanların motivasyonları ve yöntemleri üzerine ayrıntılı bir analiz içeriyor.

Orunsolu vd. (2022) önerdiği yöntem 15 boyutlu özellik seti ile eğitilmiş destek vektör makinesi ve Naive bayes yöntemini içermektedir. Yöntemlerini klasik tespit yöntemlerine göre daha hassas sonuçlara ulaştığını ileri sürmektedirler. Oluşturulan özellik kümesi URL, web sayfası özellikleri, web sayfası davranışlarından çıkarılmıştır.

Almomani vd. (2013) alanın ilerideki çalışmalarına ışık tutmak için mevcut filtre yöntemlerini inceliyor ve karşılaştırıyor. Ortalama, kurumların finansal kaybına, bireysel kullanıcıların ise rahatsızlığına yol açan günümüzde internetin en büyük problemlerinden biri olarak tanımlanıyor. Bu problem için birçok filtreleme yönteminin geliştirildiği fakat yine de tam bir çözüme ulaşamadığı ortaya konuyor. Çalışma saldırıların çeşitli aşamalarında kullanılmak üzere geliştirilmiş ve özellikle makine öğrenmesi üzerine yoğunlaşan filtreler üzerine yoğunlaşıyor.

Alnajim ve Munro (2009) ortalama sahtekarlığının çevrimiçi bankacılık ve e-ticaret kullanıcıları için bir problem olmaya başladığını ortaya koyuyor. Makalede ortalama web sitelerini tespit etmek için APTIPWD ismi verilen bir sistem kullanılıyor ve değerlendiriliyor. Yöntem kullanıcılara bilgi veriyor ve yanlış seçim yapıldığında da kullanıcıyı uyarıyor. Makale, eski bir yöntem olan kullanıcılara olta web siteleri ile ilgili e-posta göndermek yerine kullanılan yaklaşımın karşılaştırıldığında pozitif bir etki görüldüğünü söylüyor. Yazarlar kullanılan yaklaşımın zararsız ve ortalama web sitelerinin değerlendirmesinde kullanıcılara yardımcı olduğunu öne sürüyor.

He vd. (2011) çalışmasında ortalama saldırılarının her sene arttığını ve insanların e-

ticarete olan güvenini sarsan en büyük internet tehdidi olduğunu ifade ediyor. Makalede bir web sayfasının normal veya ortalama sayfası olduğunu davranışsal olarak tespit eden bir yöntem sunuluyor. Bu yöntemin kara liste tabanlı çalışan anti-ortalama araçlarının tespit edemeyeceği yeni ortaya çıkmış sayfaları da tespit edebileceği öne sürülüyor. Çalışmada normal ve ortalama siteleri incelenerek bir web sayfasının 12 özelliği ortaya çıkarılmış. Eğitim için kullanılan destek vektör makinasında girdi olarak normal ve ortalama web sitelerinde oluşan bir veri seti kullanılmış. Son olarak modele bir test veri seti verilmiş. Mevcut modeller ile kıyaslandığında deneysel sonuçlar sunulan ortalama detektörünün yüksek hassasiyet değeri ile birlikte, düşük yanlış pozitiflik ve düşük yanlış negatiflik sunduğu gösterilmiş.

Chen vd. (2020) çalışmasında, çevrimiçi kullanıcıların ortalama saldırılarına karşı algılanan duyarlılıklarını incelemektedir. Bir bireyin ortalamana duyarlılığının yakın zamanda gerçekleşen ortalama karşılaşmalarına göre şekillenebileceğini ve daha da önemlisi, yeni deneyimin duyarlılık üzerindeki etkisinin kullanıcılar arasında heterojen olacağını varsayıldığı gösterilmiş. Çalışmada ortalama tespitinin hem sürecine hem de sonucuna odaklanılmış. Üniversite öğrencilerinden alınan anket verileri, kişinin duyarlılığının, son ortalama karşılaşmasındaki algılama süreci zorluğundan ve algılama sonucu başarısızlıklarından etkilendiğini doğrulanmış. Sonuçlar aynı zamanda, ortalama tespitindeki geçmiş başarı ve ortalamanı duyarsızlaştırma gibi kişisel özelliklerin, yakın zamanda gerçekleşen bir ortalama karşılaşmasının etkilerini düzenlemedeki önemini de ortaya koyuyor. Son olarak, sonuçlar, iki tespit bileşeninin öncüllerini doğrulamanın yanı sıra, tespit süreci zorluğu ile sonuç başarısızlıkları arasındaki ilişkiyi göstermektedir. Makalede ortalama literatürüne katkıda bulunan yeni bilgiler üretildiği ve aynı zamanda uygulayıcıları bilgilendiren yeni bilgiler sağlandığı öne sürülüyor.

Pan ve Ding (2006) çalışmasında belirtildiğine göre literatürde son zamanlarda pek çok ortalama önleme planı önerilmiştir. Çalışma tüm bu önerilere rağmen ortalama saldırılarının azaltılamadığını öne sürüyor. Yazarlar bu durumun temel sebebi olarak saldırganların yöntem değişikliğinin maliyetinin çok düşük olmasını gösteriyor. Çalışma klasik yöntemlerden farklı olarak web sayfalarındaki anormallikleri, yapısal özellikleri ve http işlemleri arasındaki tutarsızlıkları inceleyen bir yöntem öneriyor. Önerilen yöntem kullanıcı deneyimi ve herhangi bir önbilgiden bağımsız çalışıyor. Makalede önerilen

yöntemin düşük hata oranı ve düşük yanlış pozitiflik ile çalıştığı öne sürülmüştür.

Liu vd. (2010) çalışmasında öncelikle ortalama için oluşturulan web sayfalarını buluyor ve ardından kümeleme özellikleri olarak bunların verilen web sayfasıyla olan ilişkilerini araştırıyor. Bulunan ilişkiler; bağlantı ilişkisi, sıralama ilişkisi, metin benzerliği ve web sayfası düzeni benzerliği ilişkisi olarak sıralanmış. Verilen web sayfasının etrafında bir küme olup olmadığını bulmak için bir DBSCAN kümeleme yöntemi kullanılmıştır. Böyle bir küme mevcutsa, söz konusu web sayfasının bir ortalama web sayfası olduğu sonucuna varılmıştır ve ardından ortalama hedefini (yani saldırdığı gerçek web sayfasını) bu kümeden bulmuştur. Eğer bir küme bulunamadıysa, web sayfası gerçek bir web sayfası olarak tanımlanır. Çalışmada test veri kümesi olarak, Phish Tank'tan seçilen 8745 ortalama sayfası kullanılmıştır. Yapılan deneyler, yaklaşımın ortalama hedeflerinin %91,44'ünü başarıyla tanımlayabildiğini gösteriyor. Öne sürülen yöntem için %3,40 olan yanlış alarm oranını test etmek için 1000 adet gerçek web sayfasından oluşan başka bir veri kümesi kullanılmıştır.

Downs vd. (2007) makalede ortalama saldırılarına karşı mücadelenin başarılı olabilmesi için insanların bu tür saldırılara neden aldandığının bulunması gerektiğini söylüyor. Çalışmada ortalama ve gerçek e-postaları doğru tahmin edilme oranını tespit edebilmek için 232 kullanıcıdan oluşan bir anket çalışması içeriyor. Makale, önceki çalışmalarda insanların ortalama saldırılarına açık olabileceğini, risklere karşı farkındalıklarının ve olta e-postaları tanımaya yönelik bildikleri yöntemlerin yeterli olmadığını öne sürdüğünü belirtiyor. Makalede yapılan anketin bu ilişkileri anlamaya yönelik olduğu belirtiliyor. Veriler, URL'leri doğru bir şekilde yorumlayabilmek ve tarayıcılardaki HTTPS sembolünün ne anlama geldiğini anlamak gibi web ortamını daha iyi anlamanın, ortalama saldırılarına karşı daha az güvenlik açığıyla ilişkili olduğunu gösteriyor. Çalışma kullanıcıları yalnızca eğitmek veya ortalama konusunda uyarılarda bulunmak yerine sezgisel yeteneklerini geliştirmenin daha başarılı sonuçlar ortaya çıkardığını öne sürüyor.

Butavicius vd. (2016) çalışmasında üç sosyal mühendislik stratejisinin, kullanıcıların e-postadaki bir bağlantıya tıklamanın ne kadar güvenli olduğuna ilişkin değerlendirmeleri üzerindeki etkisini incelemiştir. İncelenen üç strateji olan otorite, kıtlık ve sosyal kanıttır. Kullanılan e-postalar gerçek ve olta ve hedef odaklı ortalama e-postalardır. Üç stratejiden

yetki kullanımı, kullanıcıları bir e-postadaki bağlantının güvenli olduğuna ikna etmede en etkili strateji olduğu ortaya konulmuş. Makalede oltalama ve hedef odaklı oltalama e-postalarını tespit ederken, kullanıcılar en kötü performansı e-postaların yetki ilkesini kullandığında, en iyi performansı ise sosyal kanıt mevcut olduğunda gösterdiği gösterilmiş. Ayrıca kullanıcıların gerçek e-postalar ile hedef odaklı oltalama e-postaları arasında ayırım yapmakta zorlandığı belirtilmiş. Sonuç olarak karar verme konusunda daha az dürtüsel davranan kullanıcıların genellikle sahte e-postalardaki bir bağlantıyı güvenli olarak değerlendirme olasılıklarını daha düşük olduğu gösterilmiş. Son bölümde eğitim ve öğretime yönelik çıkarımlar tartışılmıştır.

Downs vd. (2006) makalesinde oltalamayı, kullanıcıları gerçek bir web sayfasıymış gibi aldatarak çeşitli bilgilerin talep edilmesi olarak tanımlamıştır. Bu saldırıları engellemenin birinci aşaması olarak kullanıcıların neden bu saldırılara aldandığının bulunması gerektiği öne sürülmüştür. Çalışmada teknoloji konusunda tecrübeli olmayan 20 kullanıcıyla mülakat yapılmış ve oltalama e-postaları karşısında izleyecekleri stratejiler ortaya konulmaya çalışılmıştır. Makale insanların oltama saldırılarına karşı açık olmasının en önemli sebebi olarak risk farkındalığının bu tür e-postaları tanımlamayla ilişkili olmaması gösterilmiştir. Aksine, ortaya çıkan verilerin insanların en aşına oldukları riskleri yönetebildiklerini, ancak alışılmadık risklere karşı ihtiyatlı olacakları gösterilmiş. Çalışmada özet olarak e-postaları değerlendirirken ve kullanıcıların web'de gezinmesine yardımcı olmaya çalışan tarayıcıların sunduğu uyarıları anlamlandırırken insanların farklı seviyelerde başarı ile kullandıkları çeşitli stratejileri araştırılmıştır.

Finn ve Jakobsson (2007) makalesinde dolandırıcılığın toplumun uzun süredir bir parçası olduğundan fakat oltalama ya da diğer bir tabirle otomatik dolandırıcılık yöntemlerinin nispeten yeni olduğundan bahsedilmiştir. Oltalamanın büyük sorunlara yol açabilecek bir yöntem olduğu ortaya konmuştur. Makalede bu sorunun çözümü için tek bir yöntem olamayacağı, sorunun temelde çok yönlü bir toplumsal ve teknolojik sorun olduğu belirtilmiştir. Çalışmada birçok araştırmacının önleyici yöntemlerin nereye odaklanması gerektiği konusunda literatürde artan çalışmaları izlediğini fakat bu çalışmaların başarılı sonuçlanması için gerekli yöntemlerin pek çok zaman etik kurullara takıldığı öne sürülüyor. Makale etik kurulların yaptıkları incelemeleri ve konu hakkındaki yasal düzenlemeleri de inceliyor. Çalışma ayrıca oltalama çalışmalarının doğasında bulunan ve

bilgilendirilmiş onam gerekliliğinin bazı yönlerinden feragat edilmesini gerektiren tartışmalı etik konulara ilişkin bir tartışmayı da içeriyor. Özet olarak makale oltalama deneylerini etik bir şekilde ve etik kurullara rehberlik eden ilke ve düzenlemelere uygun olarak tasarlama ve analiz etme sürecini özetlemektedir.

Verma vd. (2012) çalışmasında oltalama saldırılarının her yıl milyarlarca dolarlık zarara neden olduğu ve internet ekonomisi için ciddi bir tehdit oluşturduğu öne sürülüyor. E-postanın, oltalama saldırıları için hala en yaygın kullanılan araç olduğu belirtilmiştir. Makalede oltalamanın değişmez olan ve bu yöntemleri temel olarak karakterize eden özellikleri kullanarak oltalama e-postalarını tespit etmek için kapsamlı, doğal dil tabanlı bir şema sunulmuştur. Önerilen yöntem bir e-postada bulunan tüm bilgileri, yani başlık, bağlantılar ve gövdedeki metinleri kullanmaktadır. Oltalama e-postalarının kurbanlarda bir eylem gerçekleştirmesini istediği ortadayken, mevcut yöntemlerin bu konuyu göz önünde bulundurmadığı öne sürülmüştür. Yöntem güvenilir ve bilgilendirici e-postaları ayırt etmek için özel olarak tasarlanmış. Bu amaçla, oltalama tespiti için doğal dil tekniklerini kullanılmıştır. Ayrıca, eğer varsa oltalama e-postaları tespit etmek için bağlamsal bilgilerden de yararlanılmış. Bu amaçla oltalama tespiti için kullanıcının e-posta kutusunun bağlamsal sınırları incelenmiş ve bunun oltalama tespiti için önemi gösterilmiştir. Makalede önerilen yöntemin mevcut yöntemlerden daha yüksek performans gösterdiği öne sürülmüş. PhishNet-NLP adı verilen ve önerilen yöntemi kullanan uygulamanın kullanıcının posta aktarım aracı (MTA) ile posta kullanıcı aracı (MUA) arasında çalıştığı ve gelen her e-postayı, gelen kutusuna ulaşmadan önce oltalama saldırıları için işlediği belirtilmiştir.

Baykara ve Gürel (2018) oltalamayı, saldırganın e-posta veya diğer iletişim araçları aracılığıyla gerçek bir kişi/kurumu resmi kişi veya kuruluş olarak tanıtarak taklit ettiği bir siber suç biçimi olarak tanımlamıştır. Oltalamada, kurbanın oturum açma kimlik bilgilerini veya hesap bilgilerini ele geçirmek de dahil olmak üzere çeşitli işlevleri gerçekleştirebilen oltalama e-postaları aracılığıyla kötü amaçlı bağlantılar veya ekler gönderildiği belirtilmiştir. Bu e-postaların para kaybı ve kimlik hırsızlığı nedeniyle mağdurlara zarar verdiği ortaya konmuştur. Çalışmada, oltalama tespit sorunu ve oltalama e-postalarının nasıl tespit edileceği hakkında bilgi veren “Anti Oltalama Simulator” adı verilen bir yazılım geliştirilmiştir. Bu yazılım ile mail içerikleri

incelenerek ortalama ve spam mailler tespit edilmektedir. Veritabanına eklenen spam kelimelerin Bayesian algoritması ile sınıflandırılması sağlanmıştır.

Basnet vd. (2008) ortalamayı, kötü niyetli bir web sitesinin şifreler, hesap ayrıntıları veya kredi kartı numaraları gibi hassas bilgileri elde etmek amacıyla yasal bir sitenin kimliğine bürünmesiyle ortaya çıkan bir kimlik hırsızlığı biçimi olarak tanımlamıştır. Çalışmada bu tür saldırıları tespit etmek için çeşitli ortalama tespit ve önleme yazılımları olduğu belirtilmiştir. Makalede e-posta denemeleri ve web sitelerindeki ortalama içeriklerini tespit etme çabaları sonucunda kimlik avcılarını, mevcut yazılım ve teknikleri atlatmak için yeni ve hibrit teknikler konuları incelenmiştir.

Wenyin vd. (2005) çalışmasında ortalama web sayfalarının tespitine yönelik, görsel benzerliğe dayalı, ortalamaya karşı kurumsal bir çözümün parçası olarak kullanılabilecek bir yaklaşım önermektedir. Gerçek bir web sayfası sahibinin, görsel olarak gerçek web sayfasına benzeyen şüpheli web sayfalarını webde aramak için bu yaklaşımı kullanabileceği belirtilmiştir. Önerilen yöntemde görsel benzerliğin karşılık gelen önceden ayarlanmış eşikten yüksek olması durumunda bir web sayfası ortalama şüphelisi olarak raporlanır. Yapılan deneyler ile yaklaşımın çevrimiçi kullanıma yönelik ortalama web sayfalarını başarıyla tespit edebildiği gösterilmiştir.

Wu vd. (2006) makalesinde o dönem için yaygın olan ve ortalama sayfalarını tespit etmek için yararlı olan tarayıcı güvenlik araç çubuklarını inceliyor. Çalışma araç çubuklarını kullanıcı deneyimi, kullanılabilirlik ve kişisel bilgileri vermeye yönelik gerçekten doğru bilgiler verip vermediği açısından bir inceleme sunuyor. Yapılan incelemede üç araç çubuğu ve tarayıcıların diğer göstergeleri iki kullanıcı çalışması ile incelenerek, bunların tamamının ortalama saldırılarını önlemede etkisiz olduğu tespit edilmiştir. Deneklerden araç çubuğuna dikkat etmeleri istenmesine rağmen birçoğunun başaramadığı; diğerlerinin ise web sayfalarının içeriği gerçek görünüyorsa araç çubuklarının uyarılarını göz ardı ettiği belirtilmiştir. Pek çok kişinin ortalama saldırılarını anlamadığını veya bu tür saldırıların ne kadar karmaşık olabileceğinin farkında olmadığını öne sürülmüştür.

Rekouche (2011) makalesinde ortalama tarihinin, toplu parola çalma saldırılarının görüldüğü AOL'un ortaya çıktığı 1990'ların ortalarına dayandığı belirtilmiştir. Bu

yazılım programlarından ilki olarak, AOHell adı verilen bir program olduğu ve oltalama kelimesinin bu programla ortaya çıktığı öne sürülmüştür. Yazılım, 1995 yılının ocak ayında ortaya çıktığı ve otomatik bir şifre ve kredi kartı çalma mekanizması sağladığı belirtilmiştir. Parola ve bilgi çalma yazılımlarının bilgisayar ağlarının ilk ortaya çıktığı zamanlara dayandığı ifade edilmekle birlikte AOHell'in oltalama sistemi, yapılan ilk otomatik araç olarak tanımlanmıştır. Program, birkaç yıl içinde oluşturulan diğer birçok otomatik oltalama sisteminin oluşturulmasını etkilediği gösterilmiştir. Sonrasında bu programların sayısız oltalama saldırısına yol açan amatörler tarafından kullanıldığı belirtilmiştir. Amatörler tarafından yürütülen bu faaliyetlerin daha sonra AOL dışına çıkarak internet ortamında yayıldığı ve profesyonel suçlular tarafından da kullanıldığı gösterilmiştir. Makalede bazı gençlerin eğlence amaçlı ortaya çıkardığı yöntemlerin günümüzde büyük şirketleri ve ülkeleri dahi etkilediği ifade edilmektedir.

Irani vd. (2008) çalışmasında oltalama e-postalarındaki değişimi, Ağustos 2006'dan Aralık 2007'ye kadar toplanan 380.000'den fazla oltalama mesajından oluşan bir veri seti ile incelemiştir. Bu çalışma sonucunda öncelikle oltalama mesajları şok saldırılar ve şok olmayan saldırılar olarak ikiye ayrılmıştır. Oltalama mesajı üreticilerinin, aynı mesajı yeniden kullanarak oltalama mesajının kullanılabilirliğini artırmaya çalıştığı gösterilmiştir. Bazı durumlarda bu, aynı oltalama mesajının nispeten daha uzun bir süreye yayılmasına (şok olmayan saldırılar) karşı kısa bir süre boyunca büyük miktarda oltalama mesajı gönderilerek (şok saldırı) yapıldığı ifade edilmiştir. İkinci sonuç olarak oltalama saldırılarının özellikleri, geçici özellikler ve yaygın özellikler olarak iki grupta sınıflandırılmıştır. Birkaç saldırıda mevcut olan ve nispeten kısa bir ömre sahip (geçici) özellikler oltalama saldırılarının güçlü göstergeleri; saldırıların çoğunda bulunan ve uzun bir ömre sahip (yaygın) özellikler ise genellikle oltalamanın zayıf özellikleri olarak gösterilmiştir. Bunların kullanılmasına sebep olarak oltalama mesajı üreticilerinin, geçici özelliklerin faydasını zaman içinde sınırlaması ve gerçek mesajlarda da görünen özellikleri seçerek yaygın özelliklerin faydasını sınırlandırdığı öne sürülmüştür. Sonuçlarımız, oltalama mesajlarının anlaşılmasını geliştirmek açısından yararlı olmakla birlikte, daha fazla çalışmaya ihtiyaç duyulduğunu da göstermektedir.

Moore ve Clayton (2007) makalesinde banka ve şirketlerin oltalama siteleri ile, hizmet sağlayıcılardan sahte siteleri yayından kaldırmasını için baskı yaparak mücadele ettiğini

belirtmiştir. Dolandırıcıların ise bu süreç boyunca pek çok parola, kimlik bilgisini ele geçirdiği, süreçlerin uzun olduğu ve bu süreler ve ziyaretçi sayıları analiz edildiğinde web sitelerinin kaldırılmasının ortalama ile mücadele açısından geçerli bir yöntem olduğu fakat kesin çözüm sunmadığı ortaya konmuştur. Farklı yer sağlayıcıların kaldırma süreleri üzerinde yapılan çalışmada sürelerin normal dağılıma uyduğu fakat bazı yer sağlayıcıların çok daha hızlı işlem yaptığı gösteriliyor. Çalışmada örnek olarak ortalama sistemlerinin yarısına denk gelen ve yapısal değişiklikler ile uzun süreler yayında kalan sitelerin bir alt kümesi incelenmiştir. Son olarak, gözlenen ortalama web siteleri nedeniyle bankacılık sektörünün uğradığı toplam zarara ilişkin bir tahmin sunuluyor.

Williams vd. (2018) çalışmasında ortalama e-postalarını, çalışanları kötü amaçlı bağlantılara veya eklere tıklamaya teşvik ederek kuruluşların teknik sistemlerine sızmak için kullanılan bir araç olarak tanımlıyor. Farkındalık kampanyalarının ve ortalama simülasyonlarının kullanılmasına rağmen çalışanlar ortalama e-postalarına karşı savunmasız olduğu tespit ediliyor. Araştırma, hedef odaklı ortalama olarak bilinen hedefli ortalama e-postalarına karşı çalışanların duyarlılığını araştırmak için karma yöntem yaklaşımını kullanıyor. Birinci çalışmada, altı haftalık bir süre boyunca 62.000 çalışana gönderilen dokuz hedef odaklı ortalama simülasyon e-postası, otoritenin varlığı ve aciliyet etkisi tekniklerine göre derecelendirilmiştir. Sonuçlar, yetki ipuçlarının varlığının, kullanıcının bir e-postada yer alan şüpheli bir bağlantıya tıklama olasılığını artırdığını göstermiştir. İkinci çalışmada, çalışma ortamındaki ek faktörlerin çalışanların hedef odaklı ortalama duyarlılığını etkileyip etkilemediğini araştırmak için ikinci bir kuruluştaki altı odak grubu gerçekleştirilmiştir. Makalede bu faktörleri mevcut teorik yaklaşımlarla ilişkili olarak tartışılıyor ve kullanıcı toplulukları için çıkarımlar sunuluyor.

Dunlop vd. (2010) makalesinde ortalama saldırılarını düzenleyen saldırganların sürekli yeni yöntemler geliştirdiğinden bahsediyor. Makalede pek çok programın bu saldırılardan koruduğunu iddia etmesine rağmen, kullandıkları yöntemler gereği sıfırıncı gün yani yeni ortaya çıkmış kampanyalardan koruyamadığı öne sürülmüş. Bu tür saldırılara karşı koruduğunu özellikle belirtenlerin ise hem gerçek hem olta siteleri yanlış etiketlediğinden bahsedilmiş. Çalışma sıfırıncı gün ortalama kampanyaları için yüksek doğruluk sunan bir yöntem sunuyor. Yöntem web sayfasının ekran görüntüsünü alıp OCR ile metne dönüştürdükten sonra Google PageRank algoritmasını kullanıyor. 100 adet gerçek 100

adet olta site üzerinde yapılan test çalışmasında gerçek sitelerin %100, olta sitelerin %98 başarı oranı ile etiketlendiği öne sürülüyor.

Fatima vd. (2019) makalesinde saldırgan yerine sosyal mühendis ifadesini kullanıyor ve günümüzün teknoloji dünyasında bu sosyal mühendislerin insanların zayıflıklarını bulmak için sürekli yeni taktikler geliştirdiğinden bahsediyor. Bu taktikler ile hesap bilgisi, cihazın diğer verileri ele geçirmek üzere insan psikolojisinin hedef alındığı belirtiliyor. Çalışmada ortalama için kullanılan direkt hedefe yönelik yapılan saldırılar için çevrimiçi bilgi ifşasının neden önemli olduğu ve bu konu dikkat edilmesi gereken hususlardan bahsedilmiş.

Bergholz vd. (2008) çalışmasında ortalamanın internet iletişimi ve web ekonomisi için ne kadar önemli olduğundan ve ortalama saldırılarında saldırganların kurbanları parola, hesap bilgileri, kritik başka kişisel verileri açıklamaya zorladıklarından bahsediliyor. Ortalama kampanyalarının sürekli değişmesi ve gelişmesinden dolayı karaliste tabanlı sistemlerin tam olarak etkili olmasının mümkün olmadığı öne sürülüyor. Çalışmada buna dayanarak mevcut e-postalar kullanılarak eğitilen bir sınıflandırıcının yeni e-postaları sınıflandırıp sınıflandıramayacağı üzerine araştırma yapılıyor. Araştırma uyarlanabilir şekilde eğitilmiş Dinamik Markov Zincirleri ve yeni gizli sınıf konusu modelleri tarafından oluşturulan gelişmiş e-posta özelliklerini tavsiye ediyor. Kamuya açık bir test veri setinde, bu özellikleri kullanan sınıflandırıcılar, önceki çalışmalara kıyasla yanlış sınıflandırılan e-postaların sayısını üçte iki oranında azaltabilmektedir. Çalışmanın yapıldığı tarihte yakın zamanda önerilen daha anlamlı bir değerlendirme yöntemini kullanarak, bu sonuçların istatistiksel olarak anlamlı olduğunu gösterilmiş. Ayrıca yöntem, gerçek hayattan veriler içerene, halka açık olmayan bir e-posta veri seti üzerinde de başarıyla test edilmiştir.

Harrison vd. (2016) çalışmasında ortalama saldırılarının mekanizmaları ortaya çıkararak bireylerin ortalama saldırılarına karşı duyarlılığını ortaya koymaya çalışmıştır. Çalışmanın odak noktası e-posta mesajlarının özellikleri, kullanıcıların ortalama ile ilgili bilgi ve deneyimleri ile bunlar arasındaki etkileşim ve son olarak kullanıcıların ortalama e-postalarını bilişsel olarak nasıl işlediğidir. Araştırmada 194 kullanıcı üzerinde gerçek bir ortalama saldırısını simüle eden bir saha deneyi yapılmıştır. Deneyde mesajların

içeriklere manipüle edilerek kullanıcılar ile alakalı ölçümler yapılmıştır. Sonuçta hedeflerin %47 sinin deneyde kullanılan sahte formu doldurduğu tespit edilmiştir. Bir kişinin ortalama karşı yatkınlığının e-posta özelliklerine olan dikkatin azlığı ve sahte mesajın çok detaylandırılması üzerinden tahmin edilebileceği sonucuna varılmıştır. Tehdit ve ödül öğeleri içeren mesajların süreçleri etkilemediği ve farkındalık konusunda da bir etkisinin olmadığı öner sürülüyor. Deneylerde kullanıcıların e-posta sistemi ile ilgili bilgi ve deneyimlerinin ortalama karşı dayanıklılığı arttırdığı gözlemlenmiştir. Çalışma ortalama saldırılarını önlemek için kullanıcı temelli yaklaşımlar öne sürmektedir. Bu yaklaşımlar eğitim gibi daha klasik yöntemlere dayanmaktadır. Çalışma, ortalamanın neden başarılı olduğu, çevrimiçi aldatmaya yatkınlıkta insan faktörünün yanı sıra bu bağlamda etkili karar vermede bilgi işlemenin rolü hakkında yeni bilgiler sunmaktadır. Bulgulara dayanarak yazarlar, ortalamaya ilişkin yaygın yanlış kanıları ortadan kaldırırken bireyleri bu tür çevrimiçi dolandırıcılıktan korumaya yönelik daha etkili eğitim çabalarını tartışıyor.

Mohammad vd. (2014) çalışmasında ortalamayı kullanıcının kullanıcı adları, parolaları ve sosyal güvenlik numarası gibi özel bilgilerini ele geçirmeyi amaçlayan güvenilir bir firmanın web sitesini taklit etme sanatı olarak tanımlıyor. Ortalama web siteleri, içerik bölümlerinin yanı sıra web sitesiyle birlikte sağlanan tarayıcı tabanlı güvenlik göstergelerinde çeşitli ipuçları içerdiğinden bahsediliyor. Çalışma ortalama ile mücadele etmek için çeşitli çözümler önerildiğinin fakat bu tehdidi kökten çözebilecek tek bir sihirli değnek olmadığını öne sürüyor. Makalede ortalama saldırılarını tahmin etmede kullanılabilecek umut verici tekniklerden biri olarak, veri madenciliğini, özellikle de 'sınıflandırma kurallarının çıkarılmasını' öneriyor. Sebep olarak, ortalama önleme çözümlerinin web sitesi sınıfını doğru bir şekilde tahmin etmeyi amaçladığını ve bu amacın veri madenciliği sınıflandırma tekniği hedeflerine tam olarak uyduğu öne sürülüyor. Çalışmada yazarlar, ortalama web sitelerini gerçek olanlardan ayıran önemli özelliklere ışık tutuyor ve kurala dayalı veri madenciliği sınıflandırma tekniklerinin ortalama web sitelerini tahmin etmede ne kadar iyi olduğunu ve hangi sınıflandırma tekniğinin daha güvenilir olduğunun kanıtlandığını değerlendiriyor.

Fette vd. (2007) pek çok kullanıcının sürekli olarak hesap bilgilerini, oturum açma bilgilerini vermek için saldırılarına maruz kaldığını öne sürüyor. Ortalama adı verilen bu

saldırıların çoğunlukla bilgileri toplamak için yapılmış sahte bir sitenin bağlantısını içeren bir e-posta ile başladığından bahsediliyor. Makalede bu saldırıları tespit etmek için, en genel haliyle, elektronik iletişimde kullanıcı hedefli aldatmacayı öne çıkarmak için tasarlanmış bir özellik seti üzerinde makine öğrenmesi uygulanan olan bir yöntem sunuluyor. Bu yöntem, küçük bir değişiklikle ortalama web sitelerinin veya kurbanları bu sitelere yönlendirmek için kullanılan e-postaların tespitine uygulanabiliyor. Yöntem yaklaşık 860 adet ortalama e-postası ve 6950 adet ortalama olmayan e-posta üzerinde test edilmiş ve ortalama e-postalarının %96'sından fazlasını doğru bir şekilde tanımlarken gerçek e-postaların yalnızca %0,1'ini yanlış sınıflandırmış. Çalışmada özellikle saldırıların ve mevcut bilgilerin gelişimsel doğası açısından, aldatmacayı spesifik olarak tanımlamaya yönelik bu tür tekniklerin geleceği hakkında bir tartışma yer alıyor.

Jakobsson (2005) çalışmasında bir web sitesindeki tehdidin görselleştirilmesi ve ölçülmesine olanak tanıyan ve ortalama saldırılarını modellemek ve tanımlamak için araçlar sunuyor. Önerilen model döneminde yeni ortaya çıkmış bir saldırı sistemi olan bağlama duyarlı ortalama saldırıları üzerinde de test edilmiş. Çalışmada bir saldırının risklerini ekonomik analiz yoluyla ölçmek için sunulan modeli kullanmanın yollarını ve açıklanan saldırılara karşı savunma yöntemlerin açıklanıyor.

Bergholz vd. (2010) çalışmasında ortalama e-postalarını, genellikle güvenilir görünen bir kaynaktan gelen, kullanıcıdan parolasını veya diğer gizli bilgilerini girmesinin istendiği bir web sitesine giden bağlantıyı tıklamasını isteyen bir mesaj olarak tanımlıyor. Yazarlara göre ortalama e-postalarının çoğu, finansal kurumlardan para çekmeyi veya özel bilgilere erişmeyi amaçlamaktadır. Ayrıca ortalamanın son yıllarda büyük ölçüde arttığı ve küresel güvenlik ve ekonomi için ciddi bir tehdit oluşturduğu belirtiliyor. Makale ortalama karşı bir dizi olası karşı önlem olduğunu, bunların; kimlik doğrulama protokolleri gibi iletişim odaklı yaklaşımlardan kara listeye alma ve içerik bazlı filtreleme yaklaşımlarına kadar uzandığına değiniyor. Yazarlar, ilk iki yaklaşımın şu anda geniş çapta uygulanmadığını veya eksiklikler gösterdiğini savunuyor. Bu nedenle içerik tabanlı ortalama filtreleri iletişim güvenliğini artırmak için gereklidir ve yaygın olarak kullanılır. Buna göre e-postanın içeriğini ve yapısal özelliklerini yakalayıp bir dizi özellik çıkarılır. Daha sonra istatistiksel bir sınıflandırıcı, spam (yasal, örneğin toplu banka e-postaları), spam veya ortalama olarak etiketlenen bir e-posta eğitim seti üzerinde bu

özellikleri kullanarak eğitilir. Bu sınıflandırıcı daha sonra yeni gelen e-postaların sınıflarını tahmin etmek için bir e-posta akışına uygulanabilir. Makalede özellikle ortalama e-postalarını tanımlamaya çok uygun olan bir dizi yeni özellik açıklanıyor. Bunlar, e-posta konularının düşük boyutlu açıklamalarına yönelik istatistiksel modelleri, e-posta metninin ve dış bağlantıların sıralı analizini, gömülü logoların tespitini ve ayrıca gizli tuzlama göstergelerini içeriyor. Gizli tuzlama, içeriğin okuyucu tarafından algılanamayacak şekilde kasıtlı olarak eklenmesi veya çarpıtılması olarak tanımlanıyor. Araştırmada deneysel değerlendirme için, spam, ortalama ve yasal spam olarak önceden etiketlenmiş geniş ve gerçekçi bir e-posta veri seti kullanılmış. Deneylerde kullanılan yöntemin, ortalama e-postalarını sınıflandırmaya yönelik yayınlanmış diğer yaklaşımlardan daha iyi performans gösterdiği öne sürülüyor. Makalede bu yaklaşımın bir e-posta sağlayıcısının iş akışında pratik uygulaması yönelik sonuçlar tartışılıyor. Araştırmada son olarak filtrelerin nasıl güncellenebileceği ve yeni ortalama türlerine nasıl uyarlanabileceğine dair bir strateji açıklanıyor.

Moghimi ve Varjani (2016) makalesinde internet bankacılığı temel alınarak ortalama saldırılarının tespiti için yeni bir yöntem sunuyor. Kural tabanlı çalışan yöntem, web sayfası kimliğini belirlemek için önerilen iki yeni özellik kümesini kullanıyor. Önerilen özellik kümeleri, sayfa kaynakları kimliğini değerlendirmek için dört özellik ve sayfa kaynağı öğelerinin erişim protokolünü tanımlamak için dört özellik içeriyor. Önerilen ilk özellik setinde içerik ile bir sayfanın URL'i arasındaki ilişkiyi belirlemek için yaklaşık dize eşleştirme algoritmaları kullanılmış. Önerilen özellikler, arama motorlarının sonuçları ve/veya web tarayıcı geçmişi gibi üçüncü taraf hizmetlerden bağımsızdır. Web sayfalarını sınıflandırmak için destek vektör makinesi (SVM) algoritmasını kullanmıştır. Yazarlar deneylerin, önerilen modelin internet bankacılığındaki ortalama sayfalarını %99,14 doğru pozitif ve yalnızca %0,86 yanlış negatif alarmla tespit edebildiğini göstermiştir. Duyarlılık analizinin çıktısı, önerilen özelliklerin geleneksel özellikler üzerindeki önemli etkisini göstermektedir. İlgili bir yöntemi benimseyerek önerilen SVM modelinden gizli bilgiler çıkarılmıştır. Araştırmada önerilen yöntemi daha işlevsel ve kullanımı kolay hale getirmek için çıkarılan kuralları PhishDetector adlı bir tarayıcı uzantısına yerleştirilmiş. Geliştirilen tarayıcı uzantısı internet bankacılığındaki ortalama saldırılarını yüksek doğruluk ve güvenilirlikle tespit edebildiği gösterilmiştir. Makaleye

göre PhishDetector sıfırncı gün ortalama saldırılarını da tespit edebiliyor.

Ludl vd. (2007) makalesinde ortalamay, saldırganların sosyal mühendislik ve web sitesi kimlik sahtekarlığı tekniklerini bir arada kullanarak kullanıcıyı gizli bilgileri ifşa etmesi için kandırdığı bir elektronik çevrimiçi kimlik hırsızlığı yöntemi olarak tanımlıyor. Bu bilgilerin genellikle yasadışı ekonomik kâr elde etmek için kullanıldığı üzerinde duruluyor. Yazarlar, basit olmasına rağmen, ortalama saldırılarının oldukça etkili olduğundan bahsediyor. Makaleye göre başarılı ortalama saldırılarının sayısı sürekli artıyor ve literatürde ortalamaya karşı birçok çözüm öneriliyor. Popüler ve yaygın olarak uygulanan çözümlerden biri olarak, kara listeye dayalı ortalama önleme tekniklerinin tarayıcılara entegrasyonu gösteriliyor. Ancak bu tür kara listeye alma yaklaşımlarının gerçek hayattaki ortalama saldırılarını azaltmada ne kadar etkili olduğu şu anda belirsizdir. Araştırmada, iki popüler ortalama önleme çözümünün etkinliğini analiz etmeye ilişkin bulgular sunuluyor. Üç haftalık bir süre boyunca Google ve Microsoft tarafından tutulan kara listelerin etkinliğini 10.000 ortalama URL'si ile otomatik olarak test edilmiştir. Ayrıca, çok sayıda ortalama sayfasını analiz edilerek, ortalama sayfalarını tanımlamak için kullanılabilecek sayfa özelliklerinin varlığı araştırılmıştır.

Yu vd. (2009)'ne göre ortalama, web suçluları arasında en popüler uygulama haline gelmiştir. Yazarlar ortalama saldırılarının giderek daha sık ve karmaşık hale geldiğini öne sürüyor. Ortalamanın etkisinin, kimlik hırsızlığı ve mali kayıp riskini ortaya çıkardığından ciddi ve önemli olduğu öne sürülüyor. Makalede ortalama için kullanılan en popüler yöntemler ve ortalamayı tespit etmek için geliştirilen PhishCatch algoritması açıklanmaktadır. PhishCatch algoritması, ortalama e-postalarını tespit edecek ve kullanıcıları ortalama e-postaları konusunda uyaracak buluşsal tabanlı bir algoritmadır. Algoritmadaki ortalama filtreleri ve kuralları, ortalama metodolojileri ve taktikleri üzerine yapılan kapsamlı araştırmaların ardından formüle edilmiştir. Test sonrasında PhishCatch algoritmasının yakalama oranının %80, doğruluğunun ise %99 olduğunu tespit edilmiştir. Ayrıca makalede bu algoritmanın geliştirilmesinde kullanılan yaklaşım, uygulama detayları ve test sonuçları tartışılmaktadır.

Han vd. (2016) çalışmasında ortalamay, farkında olmayan kullanıcıları gizli bilgilerini ifşa etmeleri için kandıran bir çevrimiçi kimlik hırsızlığı biçimi olarak tanımlıyor.

Oltalama saldırılarının azaltılması için önemli çaba harcanmasına rağmen, bu saldırıların gerçek yaşamdaki tüm yaşam döngüsü hakkında hala çok az şey bilindiği ancak bu durumun kapsamlı oltalama önleme tekniklerinin geliştirilmesine yönelik önemli bir adım teşkil ettiği öne sürülüyor. Makalede oltalama yöntemlerine yönelik kurbanların mahremiyetini tamamen koruyan yeni bir yaklaşım sunuluyor. Bu tekniği kullanarak, beş aylık bir süre boyunca toplanan verilere dayanarak, oltalama saldırıları, mekanizmaları ve suçluların, mağdurların ve sürece dahil olan güvenlik birimlerinin davranışları hakkında kapsamlı bir değerlendirme yapılıyor. Kullanılan altyapı, saldırganın oltalama sayfalarını güvenliği ihlal edilmiş bir ana makineye kurup test ettiği andan gerçek kurbanlarla ve güvenlik araştırmacılarıyla son etkileşimine kadar olan sürede, bir oltalama saldırısının ilk kapsamlı resminin çizilmesine olanak sağlıyor. Çalışma bu popüler tehdidin süresi ve etkinliğine ilişkin doğru ölçümler sunmakta ve yüzlerce oltalama kampanyasını izleyerek gözlemlenen birçok yeni ve ilginç hususu tartışmaktadır.

Alkhalil vd. (2021) makalesinde internet kullanımının önemli ölçüde artmasıyla birlikte insanlar kişisel bilgilerini giderek daha fazla çevrimiçi paylaştığını söylüyor. Sonuç olarak, çok büyük miktarda kişisel bilgi ve finansal işlem siber suçluların saldırısına açık hale geliyor. Oltalama, suçluların kullanıcıları aldatmasına ve önemli verileri çalmasına olanak tanıyan son derece etkili bir siber suç biçimi olarak tanımlanıyor. Yazarlara göre 1990 yılında bildirilen ilk oltalama saldırısından bu yana, oltalama daha karmaşık bir saldırı vektörüne dönüştü. Şu anda oltalama, İnternet'teki dolandırıcılık faaliyetlerinin en sık görülen örneklerinden biri olarak kabul edilmektedir. Oltalama saldırıları, kurbanları için hassas bilgiler, kimlik hırsızlığı, şirketler ve devlet sırları gibi ciddi kayıplara yol açabiliyor. Makale, oltalamanın mevcut durumunu belirleyerek ve mevcut oltalama tekniklerini gözden geçirerek bu saldırıları değerlendirmeye çalışıyor. Çalışma, oltalama saldırılarını temel oltalama mekanizmalarına ve karşı önlemlere göre sınıflandırmıştır ve oltalamanın uçtan uca yaşam döngüsünün önemini dikkat çekmektedir. Makale, saldırı aşamalarını, saldırgan türlerini, güvenlik açıklarını, tehditleri, hedefleri, saldırı ortamlarını ve saldırı tekniklerini içeren yeni ve ayrıntılı bir oltalama anatomisi önermektedir. Ayrıca önerilen anatomi, araştırmacıların bir oltalama saldırısının süreç yaşam döngüsünü anlamalarına yardımcı olacak ve bu da bu oltalama saldırılarına ve

kullanılan tekniklere ilişkin farkındalığı artıracak bir sistem olarak sunulmaktadır. Yazarlar bu çalışmanın bütünsel bir ortalama önleme sisteminin geliştirilmesine de yardımcı olacağını düşünüyor. Makalede ayrıca bazı önleyici tedbirler araştırılıyor ve yeni stratejiler öneriliyor.

Abdelhamid vd. (2014)'ne göre web sitesi ortalama, günlük olarak gerçekleştirilen çok sayıda çevrimiçi işlem nedeniyle çevrimiçi topluluk için en önemli güvenlik sorunlarından biri olarak kabul edilmektedir. Makalede web sitesi ortalama, çevrimiçi kullanıcılardan kullanıcı adları ve şifreler gibi hassas bilgileri elde etmek için güvenilir bir web sitesini taklit etmek olarak tanımlanmaktadır. Yazarlara göre kara listeler, beyaz listeler ve arama yöntemlerinin kullanılması bu tehdidin riskini en aza indirecek çözüm örnekleridir. Makale ilişkisel sınıflandırma (AC) adı verilen veri madenciliğine dayalı akıllı bir yaklaşımı, ortalama web sitelerini yüksek doğrulukla etkili bir şekilde tespit edebilecek potansiyel bir çözüm olarak gösteriyor. Yapılan deneysel çalışmalara göre, AC sıklıkla basit "Eğer-O halde" kurallarını içeren sınıflandırıcıları yüksek derecede tahmin doğruluğuyla üretiyor. Makalede ortalama sorununa uygulanabilirliğini araştırmak için çok etiketli sınıflandırıcı tabanlı ilişkisel sınıflandırma (MCAC) adı verilen gelişmiş bir AC yöntemini kullanarak web sitesi ortalama sorunu araştırılmış. Araştırmada ortalama web sitelerini gerçek olanlardan ayıran özellikler de belirlenmeye çalışılmış. Ayrıca ortalama sorununu çözmek için kullanılan akıllı yaklaşımlar da araştırılmış. Farklı kaynaklardan toplanan gerçek verileri kullanan deneysel sonuçlar, AC'nin, özellikle MCAC'ın, ortalama web sitelerini diğer algoritmalarla göre daha yüksek doğrulukla tespit ettiğini gösterilmiş. Ayrıca MCAC, diğer algoritmaların bulamadığı yeni gizli bilgileri ürettiği ve bunun sınıflandırıcılarının tahmin performansını geliştirdiği öne sürülmüş.

Shahrivari vd. (2020)'ne göre internet hayatımızın vazgeçilmez bir parçası haline geldi ancak aynı zamanda ortalama gibi kötü amaçlı faaliyetlerin anonim olarak gerçekleştirilmesine de olanak sağlamıştır. Ortalama saldırganları, kişi ve kuruluşlardan hesap ID, kullanıcı adı, şifre gibi bilgileri çalmak için sosyal mühendislik yaparak veya sahte web siteleri oluşturarak kurbanlarını kandırmaya çalışıyorlar. Makaleye göre ortalama sitelerini tespit etmek için pek çok yöntem önerilmiş olsa da ortalama yapanlar bu tespit yöntemlerinden kaçmak için yöntemlerini geliştirmektedirler. Bu kötü amaçlı

etkinliklerin tespitinde en başarılı yöntemlerden biri makine öğrenmesi olarak gösterilmiş. Sebep olarak çoğu ortalama saldırısının, makine öğrenimi yöntemleriyle tanımlanabilecek bazı ortak özelliklere sahip olması öne sürülmüştür. Çalışmada ortalama web sitelerini tahmin etmek için birden fazla makine öğrenimi yönteminin sonuçlarını karşılaştırmıştır.

Khonji vd. (2013) makalesinde ortalama saldırılarının tespitine ilişkin literatürü incelemiştir. Yazarlar ortalama saldırılarının, sistemlerde insan faktöründen kaynaklanan açıkları hedef aldığını belirtiyor. Pek çok siber saldırı, son kullanıcılarda bulunan zayıflıklardan yararlanan mekanizmalar aracılığıyla yayılıyor ve bu da kullanıcıları güvenlik zincirinin en zayıf unsuru haline getiriyor. Yazarlar göre ortalama sorunu geniş kapsamlı bir sorundur ve tüm güvenlik açıklarını etkili bir şekilde azaltacak tek bir sihirli çözüm mevcut değildir. Bu nedenle belirli saldırıları azaltmak için sıklıkla birden fazla teknik uygulanır. Makale yakın zamanda önerilen ortalama azaltma tekniklerinin çoğunu incelemeyi amaçlamıştır. Çalışmada ortalama önleme tekniklerinin çeşitli kategorilerine ilişkin üst düzey bir genel bakış da sunulmaktadır

Jain ve Gupta (2017)'ya göre ortalama, siber dünyanın karşılaştığı en büyük sorunlardan biridir ve hem sektörler hem de bireyler için mali kayıplara yol açmaktadır. Ortalama saldırısının yüksek doğrulukla tespiti her zaman zorlu bir konu olmuştur. Makale görsel benzerliklere dayalı tekniklerin, ortalama web sitelerinin etkili bir şekilde tespit edilmesi için çok faydalı olduğunu öne sürmektedir. Ortalama web siteleri, kullanıcıları doğru web sitesine göz attıklarına inandırmak amacıyla, görünüm olarak ilgili gerçek web sitesine çok benzemektedir. Görsel benzerliğe dayalı ortalama tespit teknikleri, karar vermek için metin içeriği, metin formatı, HTML etiketleri, Basamaklı Stil Sayfası (CSS), resim vb. gibi özellik kümesinden yararlanır. Bu yaklaşımda, şüpheli web sitesini çeşitli özellikler kullanarak ilgili gerçek web sitesiyle karşılaştırır ve benzerlik önceden tanımlanan eşik değerinden büyükse ortalama olarak etiketlenir. Makale ortalama saldırılarının kapsamlı bir analizini, bunların istismarını, ortalama tespitine yönelik güncel görsel benzerliğe dayalı yaklaşımlardan bazılarını ve karşılaştırmalı çalışmasını sunmaktadır.

Aleroud ve Zhou (2017) makalesinde ortalamanın, büyük ölçüde gelişen web, mobil ve sosyal ağ teknolojilerinin etkisiyle çevrimiçi alanda giderek artan bir tehdit haline

geldiğini söylüyor. Yazarlar önceden yapılan ortalama sınıflandırma çalışmalarının esas olarak ortalamanın altında yatan mekanizmalara odaklandığını ancak ortaya çıkan saldırı tekniklerinin, hedeflenen ortamların ve yeni ortalama türlerini azaltmaya yönelik karşı önlemleri göz ardı ettiğini savunuyor. Araştırmada yalnızca e-posta ve web siteleri gibi geleneksel ortamlardan değil, mobil ve sosyal ağ siteleri gibi yeni ortamlarda da geliştirilen ortalama saldırılarını ve anti-ortalama tekniklerini göz önünde bulunduruluyor. Ortalamaya bütünsel bir bakış açısıyla bakılarak; saldırı tekniklerini, karşı önlemleri, hedeflenen ortamları ve iletişim medyasını içeren bir sınıflandırma sunuluyor. Bu sınıflandırmanın yalnızca çeşitli ortamlarda ortalama tespiti ve önlenmesi için etkili tekniklerin tasarımı için rehberlik sağlamakla kalmayacağı, aynı zamanda uygulayıcıların belirli ortalama sorunları türlerini ele almaya yönelik araçları, yöntemleri ve özellikleri değerlendirmesini ve seçmesini kolaylaştıracağı öne sürülüyor.

James (2005) çalışmasında, kimlik avcılarının küresel finans sektörüne karşı başarıyla dolandırıcılık eylemleri gerçekleştirebilmelerini sağlayan teknikleri açıklıyor. Ayrıca bu sömürü sanatını çevreleyen motivasyon, psikoloji ve hukuki yönleri de vurguluyor. Çalışma organize bireylerin kimliklerini ortaya çıkarmak için kullanılan yenilikçi adli tıp tekniklerini ana hatlarıyla anlatıyor ve günümüzde önleme ve yakalamayı bu kadar zorlaştıran hukuki zorluklar konusuna da değiniyor. Hedef kitle bir finans kurumunun üst düzey yöneticiden ortalama internet kullanıcısına ve ileri düzey güvenlik mühendisine kadar geniş bir kapsamda belirlenmiş.

Jansson ve von Solms (2013)'a göre suçlular, çeşitli sosyal mühendislik tekniklerini kullanarak internette birçok soruna yol açıyor ve birçok insanı çeşitli şekillerde dolandırıyor. Bu durum ise çeşitli organizasyonel toplulukları (şirketler, hükümetler) riske sokuyor. Makaleye göre buna bağlı olarak bu tür topluluklardaki insanların siber uzayda aktif olduklarında veya siberle ilgili teknolojilerle uğraşırken kendilerini nasıl koruyacaklarını öğrenmeleri gereklidir. Buradan yola çıkarak eğitimin bu bağlamda büyük bir rol oynadığı ve sonuç olarak birçok insanın güvensiz davranışlarını değiştirmeye yardımcı olabileceği öne sürülüyor. Makale yerleşik eğitimle birlikte ortalama saldırılarını simüle etmenin, kullanıcıların "ortalama saldırılarına" karşı direncini artırmaya katkıda bulunup bulunamayacağını tespit etmeye çalışmaktadır. Bu hedefe ulaşmak amacıyla araştırmada Güney Afrika'daki bir kurumda ortalama çalışması

gerçekleştirilmiş ve sonuçları sunulmuştur.

Dodge vd. (2007)'ne göre kullanıcı güvenliği eğitimi ve öğretimi, bir kuruluşun güvenlik politikasının en önemli yönlerinden biridir. Bu konuyu güçlendirmek için güvenlik tatbikatları hem eğitim hem de endüstri tarafından sıklıkla yapılmaktadır; ancak bu tatbikatlara genellikle istekli katılımcılar katılmaktadır. Çalışmada kullanıcının e-posta ortalama saldırılarına habersiz bir testle yanıt verme eğilimini değerlendirmek için bir alıştırma kullanma konseptini ele alınmış ve bunu uygulama yapılırken üzerinde bazı değişiklikler yapılmış. Makale, kullanıcı bilgileri güvenliği eğitim programının bir yönünün değerlendirilmesini de içerirken, programın oluşturulması ve uygulanması için kullanılan süreci açıklamaktadır. Yapılan değerlendirme işlemi, yanıtları kaydeden ortalama tarzı bir e-postanın gönderildiği bir test olarak açıklanabilir.

Mohebzada vd. (2012) çalışmasında ortalamayı, potansiyel bir kurbanı meşru bir kaynağın veya kuruluşun kimliğine bürünen bir mesajın gönderildiği bir tür sosyal mühendislik tekniği olarak tanımlıyor. Yazarlar göre ortalama saldırıları genellikle hedefleri parola, kredi kartı bilgileri, banka hesap numaraları veya diğer hassas bilgiler gibi gizli bilgileri açığa çıkarmaya teşvik ediyor. Çalışma insan davranışı ve teknolojinin, ortalama saldırılarının eşit derecede önemli iki yönü olarak tanımlıyor. Mevcut ortalama önleme araştırmaları teknoloji cephesine odaklanırken, ortalama saldırılarının insani yönlerine odaklanan çok az sayıda gerçek hayat çalışması yapıldığı üzerinde duruluyor. Makalede bir üniversitenin öğrencileri, mezunları, öğretim üyeleri ve personelinin oluşan 10.000'den fazla topluluk üyesine gerçekleştirilen iki büyük ölçekli gerçek hayattaki ortalama saldırısının sonuçları sunuluyor. Çalışmanın insanlar üzerinde yapılan ilk büyük ölçekli ortalama deneyi olduğu öne sürülüyor. Önceki çalışmaların kullanıcıların demografik özelliklerinin, ortalama saldırılarına karşı en savunmasız kullanıcıları belirlemede yararlı göstergeler olduğunu öne sürdüğü belirtiliyor. Bu çalışmanın sonuçları, kullanıcı demografisinin tek başına kullanıcının ortalama saldırılarına duyarlılığını tahmin edemeyeceğini göstermektedir. Araştırmada, kullanıcıları ortalama riskleri konusunda uyarmanın, daha fazla kullanıcının ortalama saldırısına yanıt vermesini önlemek için tek başına yeterli olmadığını da görüldüğü belirtiliyor. Denekler ortalama e-postalarına yanıt vermemeleri konusunda uyarılmış olsa da pek çok kişi bu uyarıyı dikkate almadığı gösterilmiş. Bulgular gerçekleştirilen iki

gerçek hayattaki ortalama saldırısının ampirik sonuçlarının analizi yoluyla açıklanıyor.

Zieni vd. (2023)'ne göre ortalama, hedeflenen markaların yanı sıra bireyler üzerinde de ciddi etkileri olan bir güvenlik tehditidir. Araştırma bu tür tehditlerin uzun süredir ortalıkta olmasına rağmen hala oldukça aktif ve başarılı olduğu gösteriyor. Yazarlar aslında saldırganların kullandığı taktikleri, saldırıları daha inandırıcı ve etkili kılmak için yıllar içinde sürekli olarak geliştirdiğini gösteriyor. Bu bağlamda ortalama tespitinin birincil öneme sahip olduğu belirtiliyor. Literatürde bu sorunla ve özellikle ortalama web sitelerinin tespitiyle başa çıkabilen çok çeşitli çözümler sunulmaktadır. Makale, ana zorlukları ve bulguları tartışarak bu alandaki en son teknolojinin geniş ve kapsamlı bir incelemesini sunuyor. Daha spesifik olarak, araştırmanın tartışma konusu liste tabanlı, benzerlik tabanlı ve makine öğrenimi tabanlı olmak üzere üç önemli algılama yaklaşımı kategorisi etrafında yoğunlaşmaktadır. Her kategori için literatürde önerilen tespit yöntemlerini, değerlendirmeleri için dikkate alınan veri kümeleriyle birlikte açıklıyoruz ve doldurulması gereken bazı araştırma boşluklarını tartışıyoruz.

Arachchilage vd. (2016) makalesinde ortalamayı, kurbanlarından kullanıcı adı, parola ve çevrimiçi bankacılık bilgileri gibi hassas bilgileri çalmayı amaçlayan bir çevrimiçi kimlik hırsızlığı olarak tanımlamıştır. Yazarlar göre ortalama eğitiminin bu tehditle mücadelede bir araç olarak değerlendirilmesi gerekir. Çalışma bilgisayar kullanıcılarının ortalama saldırılarına karşı kendilerini korumalarına yardımcı olan bir eğitim aracı olarak bir mobil oyun prototipinin tasarımını ve geliştirilmesini içermektedir. Çalışmada sunulan mobil oyun tasarımı, kullanıcıların ortalama tehditlerine karşı kendilerini koruma motivasyonu yoluyla kaçınma davranışlarını geliştirmeyi amaçlamıştır. Geliştirilen mobil oyun prototipi aracılığıyla oyun tasarım çerçevesini değerlendirmek için ön ve son testlerin yanı sıra sesli düşünme çalışması yapılmıştır. Çalışma sonuçları, test sonrası değerlendirmelerde katılımcıların ortalamadan kaçınma davranışlarında önemli bir iyileşme olduğunu göstermiştir. Ayrıca, çalışma bulguları katılımcıların tehdit algısının, koruma etkinliğinin, öz yeterliliğin, algılanan şiddetin ve algılanan duyarlılık unsurlarının tehditten kaçınma davranışını olumlu yönde etkilediği, koruma maliyetinin ise olumsuz yönde etkilediğini göstermektedir.

Alkhozae ve Batarfi (2011) çalışmasında W3C konsorsiyumu, World Wide Web'in

uluslararası standartlar organizasyonu olduğuna değinilmiştir. Bu konu üzerinde durulmasının sebebi W3C'nin birlikte çalışabilirliği artırmak ve web içeriği hakkında fikir birliğini en üst düzeye çıkarmak ve World Wide Web'in çalışmasını sağlayan şeyin önemli kısımlarını tanımlamak için standartlar, spesifikasyonlar ve öneriler geliştirmesidir. Çalışmada ortalama, dolandırıcılık web siteleri aracılığıyla kullanıcının parolalar, kredi kartı numaraları, banka hesap ayrıntıları ve diğer hassas bilgiler gibi kimlik bilgilerini ele geçirmeyi amaçlayan bir tür internet dolandırıcılığı olarak tanımlanmıştır. Makale ortalama sitelerinin web sayfası kaynak kodunda, ortalama web sitelerini gerçek web sitelerinden ayıran ve W3C standartlarına uymayan bazı özellikler barındırdığını ve bu nedenle web sayfasını kontrol ederek ortalama saldırılarını tespit etmenin mümkün olduğunu gösteriyor.

Halevi vd. (2013) makalesinde ortalama saldırılarının çevrimiçi kullanıcılar için giderek artan bir tehdit haline geldiğini belirtiyor. Makalenin yazıldığı dönemde son araştırmaların insanların kendilerine tepki vermesine neden olan faktörlere odaklanmaya başladığı belirtilmiş. Çalışma beş büyük kişilik özellikleri ile e-posta ortalama yanıtları arasındaki ilişkiyi inceliyor. Ayrıca bu faktörlerin, kişisel bilgilerin yayınlanması ve Facebook gizlilik ayarlarının seçilmesi de dahil olmak üzere kullanıcıların Facebook'taki davranışlarını nasıl etkilediğini de inceliyor. Araştırma ödül içerikli ortalama e-postası kullanıldığında cinsiyet ile ortalama e-postasına verilen yanıt arasında güçlü bir ilişki bulunduğunu gösteriyor. Ayrıca, insandaki stabil olmayan ruh halinin bu e-postaya yanıt vermeyle en ilişkili faktör olduğu da bulunmuş. Çalışma aynı zamanda dışa dönük kişilerin hem Facebook'ta daha fazla bilgi paylaşma eğiliminde olduklarını hem de daha az sıkı gizlilik ayarlarına sahip olduklarını, bunun da onların gizlilik saldırılarına karşı duyarlı olmalarına neden olabileceğini göstermektedir. Buna ek olarak çalışma katılımcıların ortalama saldırılarına karşı savunmasız olma tahminleri ile ortalama maruz kalma arasında herhangi bir korelasyon tespit edilmemiş. Buna bağlı olarak ortalama yatkınlığının ortalama risklerine ilişkin farkındalık eksikliğinden kaynaklanmadığı ve ortalama karşı gerçek zamanlı tepkiyi tahmin etmenin zor olduğu gösterilmiştir. Çevrimiçi güvenlik açığına katkıda bulunan özelliklerin daha iyi anlaşılmasının, gelecekte kullanıcıların gizliliğini ve güvenliğini artırmaya yönelik yöntemlerin geliştirilmesine yardımcı olabileceği ifade edilmiştir.

Prakash vd. (2010)'ne göre oltalama, internette hile ve aldatmanın kolay ve etkili bir yoludur. URL'leri kara listeye alma gibi çözümler bir dereceye kadar etkili olsa da bunların kara listeye alınan sitelerle tam eşleşmeye dayanması, saldırganların bu filtreleri atlatmasını kolaylaştırıyor. Çalışmada saldırganların URL'lerde sıklıkla basit değişiklikler (ör. üst düzey etki alanını değiştirmek, .tr yerine .gr kullanmak gibi) kullandığı gözlemlenmiştir. Çalışmada geliştirilen PhishNet adlı uygulama, bu gözlemi iki bileşen kullanarak yapıyor. İlk bileşen, yeni oltalama URL'lerini keşfetmek amacıyla bilinen oltalama sitelerinin basit kombinasyonlarını numaralandırmak için beş buluşsal yöntem öneriyor. İkinci bileşen, bir URL'yi kara listedeki girişlerle tek tek eşleştirilen birden fazla bileşene ayıran yaklaşık bir eşleştirme algoritmasından oluşuyor. Gerçek zamanlı kara listelerle yapılan değerlendirmede, 6000 yeni kara liste girdisinden yaklaşık 18.000 yeni oltalama URL'si keşfedilmiştir. Ayrıca yaklaşık eşleştirme algoritmasının çok az sayıda yanlış pozitif (%3) ve negatif (%5) yol açtığını da gösterilmiştir.

Afroz ve Greenstadt (2011) çalışmasında oltalamayı, kişinin kendisini güvenilir bir varlık olarak sunarak hassas veya özel verileri elde etmesini içeren bir güvenlik saldırısı olarak tanımlıyor. Yazarlara göre oltalama saldırganları genellikle görsel olarak gerçek bir siteye benzeyen web sayfalarını kullanarak kullanıcıların bir sitenin görünümüne olan güvenini suistimal ediyor. Makale oltalamayı tespit etmek için güvenilir web sitelerinin görünüm profillerini kullanan PhishZoo adını verdikleri bir oltalama tespit yaklaşımı önermektedir. Yaklaşım, sıfır gün oltalama saldırılarını ve daha küçük sitelere yönelik hedefli saldırıları sınıflandırabilme avantajıyla kara listeye alma yaklaşımlarına benzer doğruluk (%96) sağlar. Makale bilgisayarlı görme tekniklerinin pratik bir şekilde kullanılmasına yönelik bir performans analizi ve bir çerçeve içermektedir.

Parno vd. (2006)'ne göre oltalama veya web sahtekarlığı büyüyen bir sorundur. Çalışma Oltalama Koruması Çalışma Grubunun (APWG), Ağustos 2005'te neredeyse 14.000 benzersiz oltalama raporu aldığını ve bunun Aralık 2004'teki rapor sayısından %56 daha fazla olduğunu gösteriyor. Çalışma, güvenin müşteri ilişkilerinin temelini oluşturduğundan ve oltalama saldırıları bir kuruma olan güveni zayıflattığından, finansal kurumlar için oltalamanın özellikle sinsi bir sorun olduğunun üzerinde duruyor.

Parmar (2012)'a göre giderek artan sayıda tehditler ve daha katı iş düzenlemeleriyle karşı

karşıya kalan kuruluşlar, BT altyapısında güvenliğin ve uyumluluğun yeterli olmasını sağlama konusunda sürekli olarak zorluk yaşıyor. Çalışmaya göre dolandırıcılık ve hileler pek yeni olmasa da dünyanın internete, e-postaya ve sosyal medyaya bağımlılığının artmasıyla bunların hızı ve erişimi muazzam derecede artmıştır. Özellikle işyerlerinde e-postanın yaygınlaşması sadece işletmelerin başarısını kolaylaştırmakla kalmamış, aynı zamanda önemli güvenlik tehditlerine de kapı açmıştır. Yazarlar hedef odaklı ortalama türündeki saldırıların artması ile kuruluşların kendilerini nasıl koruyacaklarına ilişkin fikirlerini gözden geçirmeleri gerektiği üzerinde duruyor. Çalışma hedef odaklı ortalamanın özellikle insanların güvenliğini sömürmeye yönelik olması sebebiyle sonuçları yıkıcı olabileceği konusuna değiniyor. Çalışma bu duruma cevabın derinlemesine savunma olduğunu ve geleneksel anti-virüs çözümlerine güvenmek, beyaz liste çözümleri ve yedeklilik gibi çözümleri daha fazla kullanmaya yönelmemiz gerektiğini savunuyor.

Zuraiq ve Alkasassbeh (2019)'a göre ortalama, kimlik bilgileri ve kredi kartı bilgileri gibi hassas bilgileri ele geçirmek amacıyla kullanıcıyı sahte web siteleriyle aldatmak gibi sosyal mühendislik tekniklerini kullanan internetteki en yaygın saldırı türlerinden biridir. Yazarlara göre bu bilgiler kötüye kullanılabilir ve bu kullanıcıların büyük mali kayıplara uğramasına neden olabilir. Çalışmaya göre ortalama tespit algoritmaları, kullanıcıları bu tür saldırılardan korumak için etkili bir yaklaşım olabilir. Bu doğrultuda makalede içerik tabanlı, sezgisel tabanlı ve bulanık mantık tabanlı yaklaşımları içeren farklı ortalama tespit yaklaşımları incelenmiştir.

Burns vd. (2019)'ne göre pek çok sektörün üst düzey yöneticileri, hedef odaklı ortalama olarak bilinen sosyal mühendislik saldırılarının kurbanı olmuştur. Sosyal mühendisler, yüksek düzeyde hedeflenmiş e-postalar kullanarak, gerçek aktörleri taklit ederek mağdurları istenmeyen eylemler gerçekleştirmeleri için kandırıyor. Çalışmada hedef odaklı ortalama karşı eğitimlere ışık tutmak için çok yönlü bir deney gerçekleştirilmiştir. Sonuçlar kullanıcıların bireysel kayıp (finansal kayıp gibi) mesajları ile eğitilmesinin eğitimin etkinliğini artırabileceğini göstermektedir. Ek olarak, kurumsal eğitimin doğrudan eğitim almamış kişiler için bile genel hedef odaklı ortalama farkındalığının artmasına yol açabileceğine dair potansiyel kanıtlar bulunmuştur. Makale ayrıca bu sonuçlara rağmen, bireylerin yüksek hedefli, hedef odaklı ortalama saldırılarına karşı hassasiyetinin, uygulayıcılar ve araştırmacılar için sorun olmaya devam ettiğini

belirtmektedir.

Bullee vd. (2017) çalışmasında bir ortalama e-postasının açılış cümlesinin alıcının eylemini nasıl etkilediğini araştırmıştır. Çalışmada metot olarak, 593 çalışana kişisel olarak tanımlanabilir bilgiler sağlamaları istenen iki tür ortalama e-postası gönderilmiştir. E-postaların yarısında kişiselleştirilmiş bir hedef odaklı ortalama e-posta açılışı rastgele olarak kullanılmıştır. Bulgular çalışanların yüzde 19'unun kişisel kimlik bilgilerini genel bir ortalama e-postasında verdiğini, hedef odaklı ortalama durumunda bu oranın yüzde 29 olduğunu göstermektedir. Çalışma müteaddiyin kültürel geçmişe sahip çalışanların, diğerlerine kıyasla PII sağlama olasılıkları daha yüksek olduğunu göstermiştir. E-posta alıcısının kuruluşdaki hizmet yılı dikkate alındığında, talep edilen kişisel bilgilerin sağlanmasında yaşı her hangi bir etkisi olmadığı gözlemlenmiştir. Çalışma ortalama e-postasının açılış cümlesi kişiselleştirildiğinde başarının daha yüksek olduğunu göstermektedir. Ortaya konulan model, ortalama e-postalarından kaynaklanan mağduriyeti iyi bir şekilde açıklamıştır ve uygulayıcıların etkilerini en üst düzeye çıkarmak için farkındalık kampanyalarına odaklanması gerektiğine değinmektedir. Çalışma hedef odaklı ortalamanın dört sosyo-demografik değişken kullanılarak açıklanmayı hedeflemiştir.

Halevi vd. (2015)'ne göre son zamanlarda yapılan araştırmalar, insanların ortalama saldırılarına tepki vermesine neden olan faktörlere odaklanmaya başlamıştır. Yazarlar kullanıcıların kişiliğinin, tutumlarının ve algılanan yeterlilik faktörlerinin kendilerini bu tür bir saldırıya maruz bırakma eğilimlerini nasıl etkilediğini incelemek amacıyla kurumsal ortamlardaki çalışanlara hedef odaklı bir ortalama saldırısı gerçekleştirmiştir. Makale hedef odaklı ortalama saldırılarının, hedeflenen kurban hakkındaki kişisel bilgileri kullandıkları ve hem potansiyel kurbanlar hem de e-posta ortalama filtreleri tarafından tespit edilmeleri açısından zorluk teşkil ettikleri için normal ortalama saldırılarından daha karmaşık olduğunu göstermektedir.

Abad (2005)'a göre ortalama, bir kişiyi saldırganın güvenilir bir varlık olduğuna inandırarak kişisel bilgilerin sahtekarlıkla elde edilmesi olarak tanımlanmaktadır. Yazarlar ortalama saldırılarının daha karmaşık hale geldiğine ve gün geçtikçe arttığına değiniyor. Çalışmaya göre ortalama sorunuyla mücadeleye yönelik etkili stratejiler ve

özmler geliřtirmek iin ortalama camiasının ekonomisinin geliřtiđi altyapının anlaşılması gerekiyor. Bu sebeple arařtırmada ortalama ađlarını ortaya ıkarmak iin kapsamlı alıřmalar yapılmıřtır. Sonu olarak 3.900.000 ortalama e-postası, ortalama ile ilgili 13 nemli sohbet odasından, 13.000 normal sohbet odasından ve altı sohbet ađında yayılan 48.000 kullanıcıdan ve botnetlerde kullanılan 4.400 ele geirilmiş bilgisayardan toplanan 220.000 mesajın ayrıntılı analiz edilmiřtir. Makale, bu arařtırmadan elde edilen bulguların yanı sıra ortalama altyapısının bir analizini sunmaktadır.

Darwish vd. (2012)'ne gre gnmzde siber gvenlik zincirinin en zayıf halkasının bilgisayar kullanıcılarıdır. Sosyal mhendislik saldırıları, bilgisayar kullanıcılarını zel bilgileri sızdırabilecek eylemler gerekleřtirmeye ynlendirmek iin yaygın olarak kullanılmaktadır. Bu tr saldırılar, bilgisayar kullanıcılarını gizli bilgilerini aıđa ıkarmak iin psikolojik olarak maniple eder. Bu nedenle literatrde bilgisayar kullanıcıları, siber gvenlik olayları ile mađdurun gemiři arasındaki iliřkiyi anlamak iin gvenlik arařtırmacıları tarafından pek ok kez incelenmiřtir. Makale kurbanların gemiřleri ile ortalama saldırıları arasındaki iliřkiyi anlamayı amalayan son alıřmaların kapsamlı bir incelemesini sunuyor. alıřma ortalama kurbanlarının zelliklerini, demografik ve kiřilik yapılarını inceliyor ve son olarak bunların bir zetini sunuyor.

Tuđal vd. (2021)'ne gre yksek đretim yerleřkeleri deđiřik ađ yapıları ve kullanıcı kategorilerini ieren yapılardır. Yazarlara gre đretim kalitesini ykseltmek iin yksekđretim kurumlarında bilgi teknolojileri yođun olarak kullanılmaktadır ve bu durum bu ađları daha karmařık hale getirmektedir. Bu tr kurumlar personel ve đrenci verilerini kendi bnyelerinde barındırmaktadır ve kullanılan pek ok sistem mecburi olarak dıř dnyaya aıktır. Bu durum bu tr sistemlere yapılan saldırıları arttırmaktadır. alıřmada siber tehlikelere karřı birok zayıf halkası bulunan niversitelerin neden saldırıların hedefinde olduđu ortaya konmaya alıřılmıřtır. Arařtırmada en sık kullanılan saldırı yntemlerinin neler olduđuna bakılmıştır. Sonu olarak ortaya ıkan zafiyetlerin giderilmesi iin özmler nerilmiřtir. Temelde bilgi gvenliđi politikalarının oluřturulması ve uyulması, siber farkındalık eđitimleri aldırılması, bilgi sistemleri altyapılarının glendirilmesi gerektiđi tespit edilmiřtir. alıřma siber farkındalık eđitimleri iin izlenecek rnek bir yol haritası da nermektedir.

Chiew vd. (2015)'ne göre oltalama, kullanıcıları gizli bilgileri ifşa etmeleri için kandırmakta kullanılan sosyal mühendislik ve web sitesi sahtekarlığı tekniklerini birleştiren bir güvenlik tehdididir. Makalede internet kullanıcılarını oltalama saldırılarından korumak için bir oltalama tespit yöntemi önerilmiştir. Önerilen yöntem özellikle, bir web sitesi verildiğine oltalama olup olmadığını tespit etmeye yönelik tasarlanmıştır. Çalışmada bir web sitesinin gerçek kimliği ile taklit edilen kimliği arasındaki kimlik tutarlılığını belirlemek için logo görseli kullanılmıştır. Önerilen yöntem logo çıkarma ve kimlik doğrulama olmak üzere iki süreçten oluşmaktadır. İlk işlem, bir web sayfasının indirilen tüm görsel kaynaklarından logo görselini algılayacak ve çıkaracaktır. Doğru logo görselini tespit etmek için makine öğrenimi tekniğinden yararlanılmıştır. Çıkarılan logo görseline dayalı olarak ikinci süreçte, taklit edilen kimliği bulmak için Google görsel araması kullanılmıştır. Logo ile alan adı arasındaki ilişki özel olduğundan alan adını kimlik olarak ele almak mantıklıdır. Bu nedenle, Google tarafından döndürülen alan adının sorgu web sitesindeki alan adı ile karşılaştırılması, oltalama olan web sayfasını gerçek bir web sitesinden ayırt etmeyi sağlayacaktır. Yazarlar oltalama web sitesini tespit etmek için logo gibi grafiksel bir öğe kullanmanın etkinliğini ve uygulanabilirliğini göstermiştir.

Medvet vd. (2008)'ne göre oltalama, kullanıcının çevrimiçi bankacılık şifreleri veya kredi kartı numaraları gibi hassas bilgilerini çalmayı amaçlayan bir çevrimiçi dolandırıcılık biçimidir. Bu saldırı biçiminde kurban, gerçek bir sayfayı taklit edecek şekilde saldırgan tarafından hazırlanan bir web sayfasına bu tür bilgileri girmesi için kandırılır. Makale oltalama saldırılarının sayısının giderek arttığına ilişkin dönemdeki istatistiklerin, bu güvenlik sorununun hâlâ ciddi bir çalışmayı hak ettiğini gösteriyor. Makalede şüpheli bir oltalama sayfasını meşru sayfayla görsel olarak karşılaştırmak için yeni bir teknik sunulmuştur. Amaç, iki sayfanın şüphe uyandıracak derecede benzer olup olmadığını belirlemektir. Çalışmada bir oltalama sayfasının yasal bir sayfaya benzer görünmesini sağlamada önemli rol oynayan üç sayfa özelliği belirlenmiş ve değerlendirilmiştir. Bu özellikler, metin parçaları ve bunların stili, sayfaya gömülü resimler ve sayfanın tarayıcı tarafından oluşturulan genel görsel görünümüdür. Sunulan yaklaşımın uygulanabilirliğini doğrulamak için, 41 gerçek dünya oltalama sayfasından ve bunlara karşılık gelen gerçek hedeflerden oluşan bir veri kümesi kullanılarak deneysel bir değerlendirme

gerçekleştirilmiştir. Makalede deney sonuçlarının yanlış pozitifler ve yanlış negatifler açısından etkili sonuçlar sunduğu belirtilmiştir.

Yi vd. (2018)'ne göre web sayfası ortalama, gerçek bir varlığın kimliğine bürünerek kullanıcı adları, parolalar ve kredi kartı bilgileri gibi özel bilgileri çalmayı amaçlar. Saldırgan bunun sonunda bilgilerin ifşa edilmesine ve maddi hasara yol açacaktır. Makale esas olarak ortalama web sitelerini tespit etmek için derin öğrenme çatısının uygulanmasına odaklanmaktadır. Çalışmada ilk olarak web ortalamaya yönelik orijinal özellikler ve etkileşim özellikleri olmak üzere iki tür özellik tasarlanmıştır. Daha sonra Derin İnanç Ağlarına (DBN) dayalı bir tespit modeli sunulmuştur. ISP'den (İnternet Servis Sağlayıcısı) gelen gerçek IP akışlarını kullanan testler, DBN'ye dayalı tespit modelinin yaklaşık %90 gerçek pozitif oranına ve %0,6 yanlış pozitif oranına ulaşabileceğini göstermektedir.

Jakobsson vd. (2007) makalesinde, çeşitli gerçek ve ortalama uyaranları üzerinden logolar, üçüncü taraf onayları ve asma kilit simgeleri gibi çeşitli yaygın güvenlik göstergelerine verilen tepkileri ölçen bir kullanıcı çalışmasının önemli noktalarını incelemektedir. Çalışmanın sesli düşünme aşaması sırasında katılımcılar, e-posta mesajlarına ve web sayfalarına karşı farklı hassasiyetler sergilemişlerdir. Ulaşılan sonuç ortalama e-postalarının ve web sayfalarının orijinal görünmesini sağlayan şeyin analiz edilmesidir. Bu sonuç sadece yazarların bilimsel merakını gidermek ile kalmayıp, aynı zamanda web sayfası tasarımının gereksiz risklerden kaçınarak gerçek sitelerin tasarımına da rehberlik etmektedir. Çalışmadan çıkarılabilecek ikinci bir sonuç ise gerçek içeriğin tüketicilere şüpheli görünmesine neden olan etkenlere ilişkin gözlemlerdir.

Wen vd. (2019) makalesinde ortalama saldırılarının, 2016 ABD başkanlık seçimleri sırasında personeli sahte Google güvenlik e-postaları yoluyla şifrelerini paylaşmaları için kandırdığı ve büyük bir sorun oluşturmaya başladığına değinmektedir. Yazarlara göre bu tür güvenlik açıkları kısmen siber güvenlik konusunda yetersiz ve yorucu kullanıcı eğitiminden kaynaklanmaktadır. İdeal bir dünyada siber güvenliği aktif ve eğlenceli bir şekilde öğreten daha ilgi çekici eğitim yöntemlerine sahip olmamız gerekir. Çalışma bu ihtiyacı karşılamak için, yalnızca ortalama kavramlarını öğretmekle kalmayıp aynı zamanda kullanıcıları kendini savunma konusunda pratik yapmaya teşvik etmek için bir

rol yapma oyununda gerçek ortalama saldırılarını simüle eden What.Hack oyununu sunuyor. Yapılan kullanıcı araştırması tasarlanan oyunun, standart bir eğitim biçimine ve rakip bir eğitim oyunu tasarımına göre performansı artırmada daha ilgi çekici ve etkili olduğunu göstermektedir.

Kwak vd. (2020)'ne göre siber güvenlik eğitim programları, kullanıcıları şüpheli hedef odaklı ortalama e-postalarını bildirmeye teşvik eder ve çoğu ortalama önleme yazılımı, raporlamaya yardımcı olacak arayüzler sağlar. Ancak yapılan araştırmalar raporlamanın düşük olduğunu gösteriyor. Yazarlar çalışmada bunun neden böyle olduğu araştırmıştır. Bu amaçla Sosyal Bilişsel Teori (SCT) ortalama önleme davranışlarına yönelik algılanan öz yeterlilik üçlü faktörlerinin, hedef odaklı ortalama e-postalarının bildirilmesinden beklenen olumsuz sonuçların ve siber güvenlikle ilgili kendi kendini izlemenin bireylerin bu saldırılara maruz kalma olasılıkları üzerindeki etkisini incelemek için kullanılmıştır. Ortalama kurbanları üzerine yapılan dönemdeki son araştırmalara dayanan makale aynı zamanda siber risk inançlarını (CRB'ler) SCT çerçevesine dahil etmiştir. Anket verileri kullanılarak test edilen model, hedef odaklı ortalama e-postalarını bildirme olasılığının, algılanan öz yeterlik, beklenen olumsuz sonuçlar ve siber güvenlik öz izlemesi tarafından artırıldığını ortaya çıkarmıştır. Ayrıca çalışmaya göre, CRB'ler doğrudan üç SCT faktörünü ve dolaylı olarak bireylerin hedef odaklı ortalama e-postalarını bildirme olasılığını etkilemiştir.

Bhardwaj vd. (2020) makalesinde son zamanlarda toplu spam e-postalarından hedefli e-posta ortalama kampanyalarına doğru dramatik bir geçiş yaşandığını öne sürüyor. Bunun sonucunda da bu tür saldırıların küresel çapta kuruluşlara büyük marka, mali ve operasyonel zararlar vermeye başladığı belirtiliyor. Ortalama saldırıları basit, anlaşılır ve maskeleyen yöntemlerini temel alır. Yazarlar göre buradaki temel amaç, SMS, e-posta, WhatsApp ve diğer mesajlaşma servislerini veya bilinen, güvenilir arkadaşlardan geliyormuş gibi görünen sahte telefon çağrılarını kullanarak mümkün olduğunca fazla bilgi elde etmek için hiçbir şeyden haberi olmayan bir kurbanı cezbetmek ve kandırmaktır.

Alsharnouby vd. (2015) çalışmasında iyileştirilmiş tarayıcı güvenlik göstergelerinin ve artan ortalama farkındalığının, kullanıcıların kendilerini bu tür saldırılara karşı daha iyi

koruma becerisine yol açıp açmadığını değerlendirmek için bir kullanıcı araştırması gerçekleştirmiştir. Katılımcılara bir dizi web sitesi gösterilmiş ve ortalama web sitelerini tanımlamaları istenmiştir. Çalışmada web sitelerinin gerçekliğini belirlerken kullanıcıların dikkatini çeken görsel ipuçlarına ilişkin nesnel niceliksel veriler elde etmek için göz takibini kullanılmıştır. Sonuçlar kullanıcıların ortalama web sitelerini tanımlamaya hazırlandıklarında bile yalnızca %53'ünü başarıyla tespit ettiğini ve değerlendirme yaparken genellikle web sitesi içeriğine kıyasla güvenlik göstergelerine bakarak çok daha az zaman harcadıklarını göstermiştir. Ancak tarayıcı olarak kullanılan Chrome öğelerine bakma süresinin, ortalamayı tespit etme yeteneğinin artmasıyla ilişkili olduğunu tespit edilmiştir. Çalışmada ulaşılan diğer bir sonuç kullanıcıların genel teknik yeterliliğinin gelişmiş olması tespit puanlarıyla ilişkili olmadığıdır.

Dhamija vd. (2006)'ne göre kullanıcıları sahte web sitelerinden koruyan sistemler oluşturmak için tasarımcıların hangi saldırı stratejilerinin işe yaradığını ve nedenini bilmesi gerekir. Makale kullanıcıları aldatmada hangi kötü amaçlı stratejilerin başarılı olduğuna dair ampirik kanıtlar sunmaktadır. Öncelikle yakalanan çok sayıda ortalama saldırısını analiz edilmiş ve bu stratejilerin neden işe yarayabileceğine dair bir dizi hipotez geliştirilmiştir. Daha sonra bu hipotezler, 22 katılımcıya 20 web sitesinin gösterildiği ve hangilerinin dolandırıcı olduğunu belirlemelerinin istendiği bir kullanılabilirlik çalışmasıyla değerlendirilmiştir. Katılımcıların %23'ünün adres çubuğu, durum çubuğu ve güvenlik göstergeleri gibi tarayıcı tabanlı ipuçlarına bakmadığını, bunun da %40 oranında yanlış seçimlere yol açtığını tespit edilmiştir. Bunun üzerine çalışmada bazı görsel yanıltma saldırılarının en gelişmiş kullanıcıları bile kandırabildiği belirlenmiştir. Yazarlara göre bu sonuçlar, standart güvenlik göstergelerinin kullanıcıların önemli bir kısmı için etkili olmadığını göstermekte ve alternatif yaklaşımlara ihtiyaç duyulduğunu göstermektedir.

Literatürde yapılan çalışmalar incelendiğinde, makine öğrenmesi yöntemlerinin son yıllarda yapılan çalışmalarda sıklıkla kullanıldığını görmekteyiz. Topluluk öğrenme yöntemlerinin kısmen daha az görüldüğü çalışmalar içinde çalışmamızın eksik bir alanı dolduracağını düşünmekteyiz.

3. MATARYEL ve METOT

1950'lerden bu yana literatüre kazandırılan pek çok çalışma ve yöntem alanın genişlemesine yol açmakta, bu yöntemlerin hibrit kullanımı ise çalışmalarda seçilecek yöntemleri çeşitli hale getirmektedir. Çalışmamızda temel olarak yığın topluluk öğrenme yöntemi seçilmiş olup veri seti olarak Hannousse ve Yahiouche tarafından 2021 yılında geliştirilen veri ambarı kullanılmıştır.

3.1 Yapay Zeka

Günümüzdeki kullanım alanları ve oluşan alt alanları göz önüne alındığında basit bir tanım oluşturmak zor olmak ile birlikte 1955 yılında John McCharty yapay zekayı akıllıymış gibi davranan makinalar olarak açıklamıştır (Ertel 2018). Başka bir tanım ise mevcut durumda insanların iyi yaptığı işleri makinaların aynı düzeyde yapmasıdır (Rich 1983). Buradaki tanım zorluğu aslen zeka kelimesinin tanımlanmasının zorluğundan kaynaklanmaktadır. Anlamak ya da bilgi almak ve unutmamak, deneyimleyerek öğrenmek olarak tanımlanabilecek zeka, makinalara uygulanmaya çalışılsa da sözlük anlamı ile uygulamak mümkün değildir. Bu tanımlı yaparken yaşadığımız güçlük makinaların zeki olup olmadığını anlama noktasında da bize güçlük çıkarmaktadır. Bir makinanın zeki olup olmadığını ölçmek için geliştirilen en bilinen yöntem adını geliştiricisi Alan Turingden alan Turing Testidir. Bu test Alan Turing'in imitation game adını verdiği bir oyunu temel alır. Bu testte bir insanın kendisine bilgi verilmeden test edilecek makine ve bir insanla görüşmeye alınır. Makina insan zekasını taklit etmeye çalışır ve eğer mülakatı yapan hangisinin makine olduğunu anlayamazsa makine testi geçmiş kabul edilir (Fetzer 1990).

Kavramın ortaya çıkmasından itibaren 60 yıldır sürdürülen çalışmalar AlphaGo'nun dünya şampiyonu bir go oyuncusuna karşı 2016 yılında galip gelmesi ile birlikte alana olan ilgide arttı. Ekonomik birçok getirinin yanında yaşamın pek çok alanında değişimler görülmeye başlamıştır. İş gücünden tasarruf, yapay zekaya bağlı yeni iş alanları ve toplumsal değişim rahatça gözlenebilir durumdadır (Zhang ve Lu 2021).

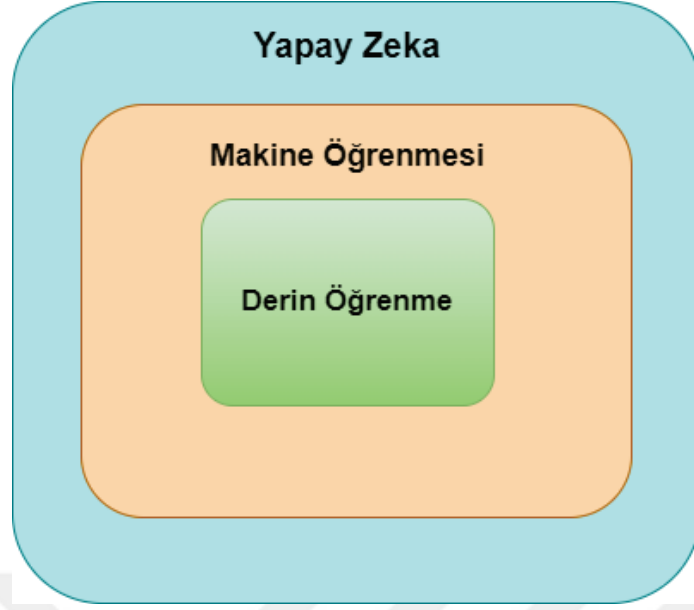
Politik atmosfer izlendiğinde pek çok ülkenin yeni bir yarış alanı olan yapay zeka

konusunda politikalar geliřtirdiđi g r lebilir. Her biri bir  lke ekonomisine eřit teknoloji devleri de bu yarıřtan geri kalmadı ve yeni yapay zeka hizmetleri ortaya  ıkardı veya mevcut  r nlerine yapay zekayı entegre etti (Zhang ve Lu 2021).

60 yıldır s regelen  alıřmalar belirli aralıklarla bilgi iřlem sistemlerinin iřlem kapasitelerine takılmıştır ve bu durum alandaki geliřimi yavaşlatmıştır. IBM Deep Blue gibi  nemli atılımlar ile ařılan bu sorunlar, son yıllarda artan GPU kapasiteleri ile birlikte tekrar hız kazanmıştır (Zhang 2019).

3.2 Makine  ğrenmesi

Makine  ğrenmesi,  evredeki ortamdan  ğrenerek insan zekasını taklit etmek  zere tasarlanmış hesaplamalı algoritmaların geliřen bir dalıdır. Bunlar, b y k verinin yeni  ađında iřleri y r ten ara lar olarak kabul edilir. Makine  ğrenmesine dayalı teknikler, desen tanıma, bilgisayarlı g r , uzay aracı m hendisliđi, finans, eđlence ve hesaplamalı biyolojiden biyomedikal ve tıbbi uygulamalara kadar  eřitli alanlarda bařarıyla uygulanmıştır. Kanserli hastaların yarısından fazlası tedavilerinin bir par ası olarak iyonlařtırıcı radyasyon, diđer adıyla radyo terapi, alır ve bu lokal hastalığın ileri evrelerinde ana tedavi y ntemidir. Radyoterapi, yalnızca kons ltasyondan tedaviye kadar olan d nemi deđil, aynı zamanda hastaların re ete edilen radyasyon dozunu aldıklarından ve iyi yanıt verdiklerinden emin olmak i in bunun  tesine uzanan geniř bir s re  k mesini i erir. Bu s re lerin karmařıklık dereceleri deđiřebilir ve karmařık insan-makine etkileřimlerinin ve karar almanın birkaç ařamasını i erebilir, bu da dođal olarak radyasyon fiziđi kalite g vencesi, tedavi planlaması, g r nt  destekli radyoterapi, solunum hareketi y netimi, tedavi tepki modellemesi ve sonu  tahmini dahil ve bunlarla sınırlı olmamak  zere bu s re leri optimize etmek ve otomatikleřtirmek i in makine  ğrenimi algoritmalarının kullanımını zorunlu kılmaktadır. Makine  ğrenmesi algoritmalarının mevcut bađlamdan  ğrenme ve g r lmeyen g revlere genelleme yapma yeteneđi, radyoterapi uygulamasının hem g venliđinde hem de etkinliđinde iyileřtirmelere izin vererek daha iyi sonu lara yol a acaktır. Buradan anlařılacađı  zere Makine  ğrenmesi sadece biliřim alanında deđil, sađlık gibi insan hayatını etkileyen pek  ok alanda da bařarılı uygulama alanları bulmaktadır (El Naqa ve Murphy 2015).



Resim 3.1 Derin öğrenme ve yapay zeka arasındaki ilişki

3.2.1 Denetimli Öğrenme

Denetimli öğrenme, makine öğreniminde çok sayıda araştırma faaliyetine yol gösterir ve birçok denetimli öğrenme tekniği multimedya içeriğinin işlenmesinde alanında uygulama alanı bulmuştur. Denetimli öğrenmenin tanımlayıcı özelliği, açıklamalı eğitim verilerinin kullanılabilir olmasıdır. Bu metodun adı öğrenme sistemine eğitim örnekleriyle ilişki kurması için etiketler konusunda talimat veren bir 'gözetimci' fikrinden yola çıkılarak verilmiştir. Bu etiketler çoğunlukla, sınıflandırma problemlerindeki sınıf etiketleridir. Denetimli öğrenme algoritmaları, bu eğitim verilerinden modeller türetir ve bu modeller, diğer etiketlenmemiş verileri sınıflandırmak için kullanılabilir (Cunningham vd. 2008).

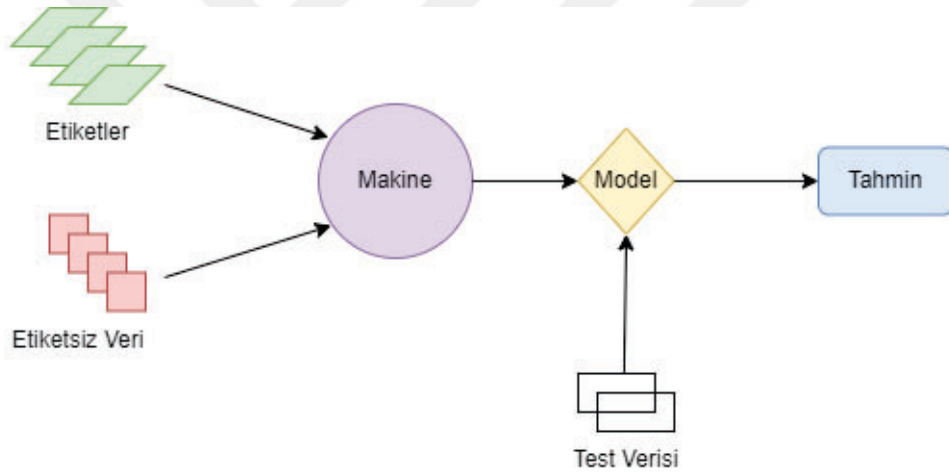
Denetimli öğrenme sınıflandırma problemleri için en çok tercih edilen yöntemdir. Bu yöntemde özellikleri verilen nesnenin etiketlerini tahmin eden bir tahminci geliştirilmesi gerekmektedir. Daha sonra algoritma girdi verileri ile birlikte düzeltilmiş çıktıları alır ve gerçek çıktıları ile düzeltilmiş çıktıları karşılaştırarak hataları bulur. Modeli bu sonuçlara göre değiştirir (Nasteski 2017).

Denetimli öğrenmede görevler ikiye ayrılır: sınıflandırma ve tekrarlama.

Sınıflandırmada etiketler ayrıkken, tekrarlama etiketler sürekli. Algoritma eğitim sürecinde verilen x diyeceğimiz veriler arasında ayrımı yapar. Bu aşamada denetimli öğrenme algoritması tahmini modeli oluşturur. Eğitim sürecinden sonra son model test verisindeki yeni bir x verisinin etiketlerini tahmin etmeye çalışır. Çıktı olan y 'nin tipine bağlı olarak:

- Eğer y kategorik çıktı değerlerine sahipse (yani tam sayı ise), y yi tahmin etme işlemi sınıflandırma olarak,
- Eğer y ondalıklı değerler sahipse, y 'yi tahmin etme işlemi tekrarlama,

olarak adlandırılır (Nasteski 2017).



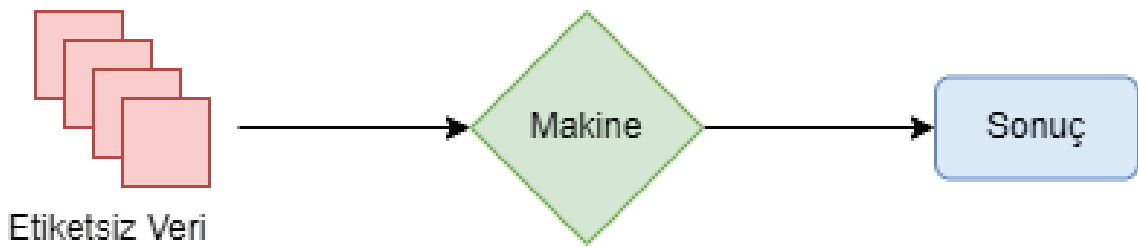
Resim 3.2 Denetimli öğrenme şeması

3.2.2 Denetimsiz Öğrenme

Denetimsiz öğrenmede makine sadece $x_1, x_2, \dots$, girdileri alır ancak çevresinden ne denetimli hedef çıktıları ne de ödüller elde eder. Çevresinden herhangi bir geri bildirim almadığı göz önüne alındığında makinenin ne öğrenebileceğini hayal etmek biraz gizemli görünebilir. Ancak, makinenin amacının karar alma, gelecekteki girdileri tahmin etme, girdileri başka bir makineye etkili bir şekilde iletme vb. için kullanılabilecek girdi temsilleri oluşturmak olduğu fikrine dayalı denetimsiz öğrenme için genel bir çerçeve geliştirmek mümkündür. Bir anlamda, denetimsiz öğrenme, saf yapılandırılmamış

gürültü olarak kabul edilebilecek olanın ötesinde verilerde desenler bulmak olarak düşünülebilir. Gözetimsiz öğrenmenin iki çok basit klasik örneği kümeleme ve boyut azaltmadır (Ghahramani 2004).

Denetimsiz öğrenmede aslen yapılan iş verinin olasılıksal modelinin öğrenilmesidir. Bir denetim veya geri besleme olmadığında dahi verilen yeni bir x verisi için olasılıksal dağılım önceden verilen veriler üzerinden çıkarılabilir. Bu sebeple denetimsiz öğrenme borsa veya hava durumu tahmini için uygun bir yöntemdir (Ghahramani 2004).



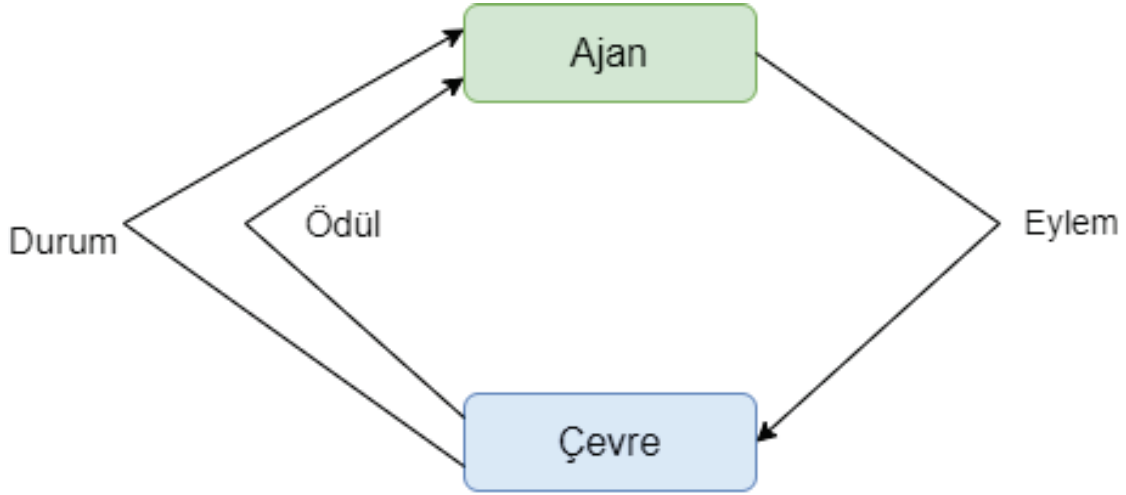
Resim 3.3 Denetimsiz öğrenme şeması

Bu tür bir model aykırı durumların tespiti veya izleme işlerinde kullanılabilir. X in bir tesiste kurulan yeni bir sensör, $P(x)$ inde normal çalışan bir tesisten alınan verilerden üretilen model olduğunu düşünürsek, $P(x)$ yeni sensörden alınan verileri değerlendirmek için kullanılabilir. Eğer elde edilen olasılık normal olmayan şekilde düşükse burada ya modelin zayıf olduğu yada tesisin normal işlemediği düşünülür (Ghahramani 2004).

3.2.3 Takviyeli Öğrenme

Takviyeli öğrenme, bilgisayarların ve insanlar gibi doğal sistemlerin hem etiketli hem de etiketsiz verilerin varlığında nasıl öğrendiklerinin incelenmesiyle ilgilenen bir öğrenme paradigmasıdır. Geleneksel olarak öğrenme ya tüm verilerin etiketsiz olduğu denetimsiz paradigmadır ya da tüm verilerin etiketli olduğu denetimli paradigmadır incelenmiştir. Takviyeli öğrenmenin amacı, etiketli ve etiketsiz verilerin birleştirilmesinin öğrenme davranışını nasıl değiştirebileceğini anlamak ve böyle bir kombinasyondan yararlanan algoritmalar tasarlamaktır. Takviyeli öğrenme, etiketli veriler kıt veya pahalı olduğunda gözetimli öğrenme görevlerini geliştirmek için kolayca bulunabilen etiketsiz verileri kullanabildiği için makine öğrenmesi ve veri madenciliğinde büyük ilgi görmektedir.

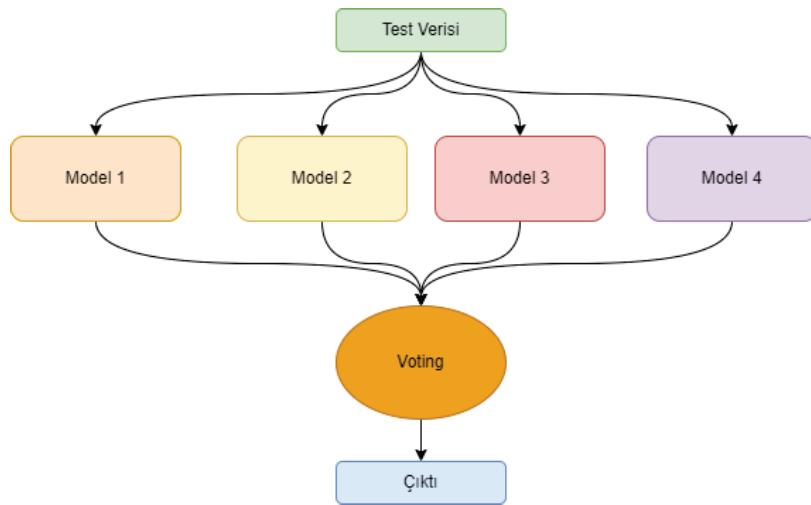
Takviyeli öğrenme ayrıca, girdinin çoğunun kendiliğinden etiketsiz olduğu insan kategori öğrenmesini anlamak için nicel bir araç olarak potansiyel göstermektedir (Zhu ve Goldberg 2009).



Resim 3.4 Takviyeli öğrenme şeması

3.3 Topluluk Öğrenme

Çoklu sınıflandırıcı sistem ve komite tabanlı öğrenme olarak da bilinen topluluk öğrenmesi, bir öğrenme problemini çözmek için birden fazla öğrenciyi eğitir ve bir araya getirir. İş akışında gösterildiği gibi önce var olan öğrenme algoritmaları ile bir grup bağımsız öğrenci eğitilir, daha sonra çeşitli stratejiler ile çıktılar birleştirilir (Zhou 2021).



Resim 3.5 Topluluk öğrenme iş akışı

3.4 Topluluk Öğrenme Yöntemleri

3.4.1 Bagging

Bagging bir temel öğrenme algoritması çağırarak her biri farklı bir önyükleme örneğinden bir dizi temel öğrenciyi eğitir. Bir önyükleme örneği eğitim veri kümesinin değiştirilerek ve alt örnekleme yapılarak elde edilir. Burada bir örneğin boyutu, eğitim veri kümesinin boyutuyla aynıdır. Bu nedenle, bir önyükleme örneği için bazı eğitim örnekleri görünebilir ancak bazıları görünmeyebilir; burada bir örneğin en az bir kez görünme olasılığı yaklaşık 0,632'dir. Temel öğrencileri elde ettikten sonra, Bagging bunları çoğunluk oylama birleştirir ve en çok oy alan sınıf tahmin edilir (Breiman 1996). Bagging'in bir çeşidi olan Random Forests'in bugüne kadarki en güçlü topluluk yöntemlerinden biri olarak kabul edildiğini belirtmekte fayda vardır (Breiman 2001).

3.4.2 Boosting

Boosting aslında birçok çeşidi olduğu için bir algoritma ailesidir. Örnek olarak AdaBoost algoritması ele alınmaktadır. Algoritma ilk olarak, tüm eğitim örneklerine eşit ağırlıklar atar. t -inci öğrenme turundaki ağırlıkların dağılımını D_t olarak gösteririz. Algoritma, eğitim veri kümesinden ve D_t 'den, temel öğrenme algoritmasını çağırarak bir temel öğrenci h_t : $X \rightarrow Y$ üretilir. Daha sonra, h_t 'yi test etmek için eğitim örneklerini kullanılır ve yanlış sınıflandırılan örneklerin ağırlıkları artırılır. Böylece, güncellenmiş bir ağırlık dağılımı D_{t+1} elde edilir. Eğitim veri kümesinden ve D_{t+1} 'den AdaBoost, temel öğrenme algoritmasını tekrar çağırarak başka bir temel öğrenci üretir. Böylece bu süreç T kez tekrarlanır, her birine tur denir ve son öğrenci, T temel öğrencinin ağırlıklı çoğunluk oylamasıyla türetilir. Burada öğrencilerin ağırlıkları eğitim süreci sırasında belirlenir. Uygulamada, temel öğrenme algoritması, ağırlıklı eğitim örneklerini doğrudan kullanabilen bir öğrenme algoritması olabilir. Aksi takdirde ağırlıklar, eğitim örneklerini ağırlık dağılımı D_t 'ye göre örnekleme suretiyle kullanılabilir (Freund ve Schapire 1997).

3.4.3 Stacking

Yığınlamanın tipik bir uygulamasında, farklı öğrenme algoritmaları kullanılarak eğitim

veri kümesinden bir dizi birinci seviye bireysel öğrenen üretilir. Bu bireysel öğrenenler daha sonra meta-öğrenen olarak adlandırılan ikinci seviye bir öğrenen tarafından birleştirilir. Yığınlamanın bilgi birleştirme yöntemleriyle yakın bir ilişkisi olduğu açıktır (Wolpert 1992).

3.5 Oltalama Saldırıları

van der Merwe vd. (2005) oltalamayı bir kullanıcıyı kişisel, finansal veya parola verilerini göndermesi için kandırmak amacıyla mevcut bir web sayfasının bir kopyasını oluşturmayı içeren sahtekarlık faaliyeti olarak tanımlıyor. Bu tanım oltalamayı, kullanıcıyı banka bilgileri ve kredi kartı numaraları gibi hassas bilgileri ifşa etmeye kandırma girişimi olarak tanımlıyor; bunun için kullanıcıya sahte web sitesine yönlendiren kötü amaçlı bağlantılar gönderiliyor. Literatürde birçok çalışma ise tek saldırı vektörü olarak e-postaları gösteriyor. Örneğin, PhishTank oltalamayı “genellikle e-posta yoluyla yapılan, kişisel bilgilerinizi çalmaya yönelik hileli girişim” olarak tanımlıyor (Kirda ve Kruegel 2005). Oltalama için yapılan bir açıklamada oltalama “kullanıcılardan çevrimiçi bankacılık şifreleri ve kredi kartı bilgileri gibi hassas bilgileri çalmayı amaçlayan bir tür çevrimiçi kimlik hırsızlığı” olarak tanımlanıyor (Kirda ve Kruegel 2005). Bazı tanımlar, sosyal ve teknik becerilerin bir arada kullanılmasının altını çiziyor. Örneğin, APWG oltalamayı “tüketicilerin kişisel kimlik verilerini ve finansal hesap kimlik bilgilerini çalmak için hem sosyal mühendislik hem de teknik hileler kullanan bir suç mekanizması” olarak tanımlıyor (“APWG | Phishing Activity Trends Reports”, t.y.). Ayrıntılı bir başka tanım ise oltalamayı “bir saldırının, kimlik avcısı olarak da bilinir, güvenilir veya kamuya açık bir kuruluştan gelen elektronik iletişimleri otomatik bir şekilde taklit ederek meşru kullanıcıların gizli veya hassas kimlik bilgilerini hileli bir şekilde almaya çalıştığı bir sosyal mühendislik biçimi” olarak tanımlamaktadır. Bu tür iletişimler çoğunlukla kullanıcıları sahte web sitelerine yönlendiren ve daha sonra söz konusu kimlik bilgilerini toplayan e-postalar aracılığıyla yapılır (Jakobsson ve Myers 2006).

3.5.1 Oltalama Yöntemleri

3.5.1.1 Bağlantı Manipülasyonu

Oltalama esas olarak bağlantılarla ilgilidir. Bir URL'i gerçek bir URL gibi göstermek için bazı akıllıca yöntemler vardır. Bir yöntem, kötü amaçlı URL'leri web sitelerinde adları olan köprü metinleri olarak sunmaktır. Başka bir yöntem, gerçek bir URL gibi görünecek yanlış yazılmış URL'ler kullanmaktır, örneğin ghoogle.com. Bahsedilen bağlantı manipülasyon yöntemlerine kıyasla tanınması çok daha zor olan bir yanlış yazılarak oluşturulan manipülasyon çeşidine IDN Spoofing denir; burada saldırganlar, İngilizce karşılıkları yerine Kiril alfabesindeki "c" veya "a" gibi İngilizce bir karaktere tıpatıp benzeyen İngilizce olmayan bir karakter kullanırlar (Bhavsar vd. 2018).

3.5.1.2 Filtre Atlama

Bu yöntemde sahtekarlar web sitelerinin içeriğini resimlerle gösterir veya Adobe-Flash kullanırlar. Bu da bazı sahtekarlık tespit yöntemleri tarafından tespit edilmelerini zorlaştırır. Bu tür saldırılardan kaçınmak için optik karakter tanıma kullanmak gerekir (Lam vd. 2009).

3.5.1.3 Web Sayfası Sahteciliği

Bu tür saldırılarda, oltalama hedef web sitesinin JavaScript kodunu manipüle ederek gerçek bir web sitesi üzerinden gerçekleştirilir. Cross-site scripting olarak da bilinen bu tür saldırıları tespit etmek çok zordur çünkü kurban meşru web sitesini kullanmaktadır (Shahrivari vd. 2020).

3.5.1.4 Gizli Yönlendirme

Bu saldırı, OAuth 2.0 ve OpenID protokolünü kullanan web sitelerini hedef alır. Kullanıcılar meşru bir web sitesine token erişimi vermeye çalışırken, tokenlerini kötü amaçlı bir servise verirler. Ancak, bu yöntem teknik olarak başka bir çok tedbir olması sebebiyle düşük önem arz ettiğinden fazla ilgi görmemektedir (Shahrivari vd. 2020).

3.5.1.5 Sosyal Mühendislik

Bu tür ortalama saldırıları sosyal etkileşim yoluyla gerçekleştirilir. Kullanıcıları kandırmak ve güvenlik bilgilerini vermek için psikolojik hileler kullanır. Saldırıları çok aşamalı olarak gerçekleştirilir. İlk olarak, kimlik avcısı saldırı için gereken hedeflerin potansiyel zayıf noktalarını araştırır. Daha sonra, kimlik avcısı hedefin güvenini kazanmaya çalışır ve en sonunda hedefin önemli bilgileri ifşa ettiği bir durum ortaya çıkarılır. Bazı sosyal mühendislik ortalama yöntemleri olarak; yemleme, korkutma yazılımı, bahane uydurma ve hedef odaklı ortalama sayılabilir (Krombholz vd. 2015).

3.5.2 Ortalama Tespit Yöntemleri

3.5.2.1 Kullanıcı Eğitimi

Kullanıcıları ve şirket çalışanlarını eğitmek ve onları ortalama saldırıları konusunda uyarmak ortalama saldırılarını önlemede etkili bir yöntemdir. Literatürde kullanıcıları eğitmek için birden fazla yöntem önerilmiştir. Birçok araştırma, kullanıcıların ortalama ile gerçek web sitelerini ayırt etmelerine yardımcı olmak için en etkili yaklaşımın etkileşimli öğretim olduğu sonucuna varmıştır (Kumaraguru vd. 2009). Kullanıcı eğitimi etkili bir yöntem olmasına rağmen, yine de insan hataları meydana gelmektedir ve insanlar eğitimlerini unutmaya eğilimlidir. Eğitim ayrıca önemli miktarda zaman gerektirir ve teknik olmayan kullanıcılar tarafından pek beğenilmez (Dhamija vd. 2006).

3.5.2.2 Liste Temelli Yaklaşım

Ortalama tespiti için yaygın olarak kullanılan yöntemlerden biri, web tarayıcılarına entegre edilmiş kara liste tabanlı ortalama önleme yöntemlerini kullanmaktır. Bu yöntemler, gerçek web sitelerinin adını içeren beyaz liste ve kötü amaçlı web sitelerinin kaydını tutan kara liste olmak üzere iki tür liste kullanır. Genellikle kara listeler ya kullanıcı geri bildirimleri ya da başka bir ortalama tespit şeması kullanılarak oluşturulan üçüncü taraf bildirimleri aracılığıyla oluşturulur. Bazı çalışmalar, kara liste tabanlı ortalama önleme yaklaşımlarının ilk kontrol sırasında kötü amaçlı web sitelerinin %90'ını tespit edebileceğini göstermiştir (Ludl vd. 2007).

3.5.2.3 Görsel Benzerlik Temelli Yaklaşım

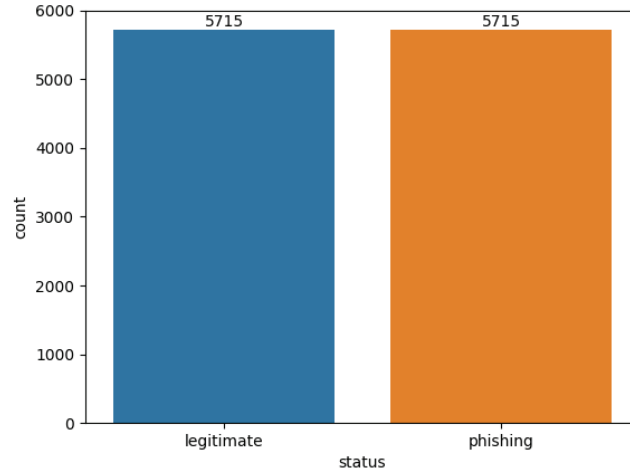
İnsanların gerçek bir web sitesi kullandıklarına inandırılırken gerçekte ise kötü amaçlı bir web sitesinde bir form doldurmalarının başlıca nedenlerinden biri, ortalama web sitesinin görünümünün hedeflenen gerçek web sitesine tam olarak benzemesidir. Bu tür saldırıları belirlemek için yöntem olarak; metin içeriğini, metin biçimini, HTML'i, CSS'i ve web sayfalarının görüntüleri analiz edilerek görsel benzerlikleri tespit edilir (Afroz ve Greenstadt 2011). Ayrıca ortalama tespitini bir görüntü eşleştirme sorunu olarak ele alan ayırt edici anahtar nokta özellikleri önerilmektedir. Görsel benzerliğe dayalı yaklaşımların bazı sınırlamaları vardır; örneğin bir web sitesinin içeriğini kullanan yöntemler, metin yerine görüntü kullanan web sitelerini tespit edemez. Görüntü eşleştirme yöntemlerini kullanan yöntemler çok zaman alıcıdır ve yeterli veri toplamak zordur (Chen vd. 2009).

3.5.2.4 Davranışsal ve Makine Öğrenmesi Temelli Yaklaşım

Makine öğrenmesi yöntemlerinin, spam e-postaları veya ortalama web siteleri gibi kötü amaçlı etkinlikleri veya eserleri sınıflandırmak için güçlü bir araç olduğu kanıtlanmıştır. Bu yöntemlerin çoğu eğitim verisi gerektirir, literatürde bir makine öğrenimi modelini eğitmek için birçok ortalama web sitesi örneği bulunmaktadır. Bazı makine öğrenimi yöntemleri, bir web sitesinin anlık görüntüsünü analiz ederek görme tekniklerini kullanır (Rao ve Ali 2015). Bazı yöntemler ise ortalama tespiti için web sitesinin içeriğini ve özelliklerini kullanır. Ortalama web sitelerini tespit etmek için birden fazla makine öğrenimi yöntemi kullanılmıştır; bunlardan bazıları Lojistik regresyon, karar ağacı, random forest, Ada boost, SVM, KNN, sinir ağları, gradyan güçlendirme ve XGBoost'tur (Shahrivari vd. 2020).

4. BULGULAR

Oltalama saldırıları, siber saldırganların kullanıcıların kimlik bilgileri gibi kişisel verilerini ele geçirmede kullandıkları saldırı türlerinden birisidir. Günümüzde insanların internet üzerinden yürüttüğü faaliyetlerin artması, oltalama saldırılarının artmasına neden olmuştur. Öte yandan oltalama saldırılarının tespiti zor, uzmanlık gerektiren ve karmaşık bir işlemdir. Intel tarafından yürütülen bir araştırmada güvenlik uzmanlarının %97'sinin oltalama saldırılarını gerçek e-postalardan ayırt etmekte zorlandığını ortaya koymuştur. Bu yüzden bu çalışmada, yapay zeka tabanlı otomatik ve hızlı oltalama saldırıları tespit sistemi geliştirilmiştir. Uygulama açık kaynak kodlu Python ortamında geliştirilmiştir. Hannousse ve Yahiouche tarafından 2021 tarafından geliştirilen veri ambarı kullanılmıştır. Veri kümesi 87 özelliğe sahip 11430 adet web adresi içermektedir. Sınıflandırma etiketine göre verilerin dağılımı Resim 4.1'de verilmiştir.



Resim 4.1 Veri kümesinin dağılımı

Veri kümesi dengeli olup %50'si zararlı, %50'si zararsızdır. Veri kümesinden bazı örnekler Çizelge 4.1'de verilmiştir.

Çizelge 4.1 Veri kümesinden örnekler

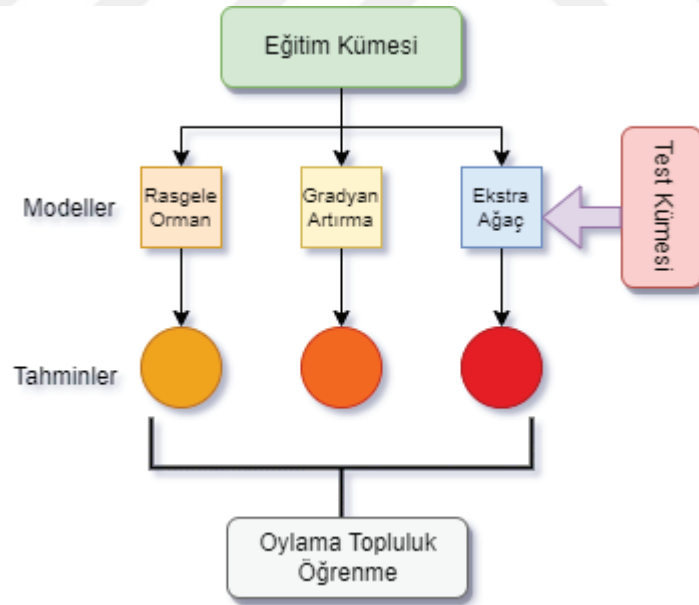
Web adresi	Etiket
http://www.crestonwood.com/router.php	Zararsız
http://shadetretechnology.com/V4/validation/a111aedc8ae390eabcfa130e041a10a4	Zararlı

Web adreslerinden elde edilen 87 adet özellikten 56'sı web adresi yapısından, 24'ü sayfaların içeriğinden, 7'si harici servislerin sorgulanmasıyla elde edilmiştir. Veri kümesinin %80'i eğitim için, %20'si test için ayrılmıştır. Buna göre eğitim aşamasından önce gruplarda oluşan örneklem sayıları Çizelge 4.2'de verilmiştir.

Çizelge 4.2 Veri kümesinin gruplara göre örneklem sayıları

Grup	Zararlı	Zararsız	Toplam
Eğitim	4581	4563	9144
Doğrulama	1134	1152	2286
Toplam	5715	5715	11430

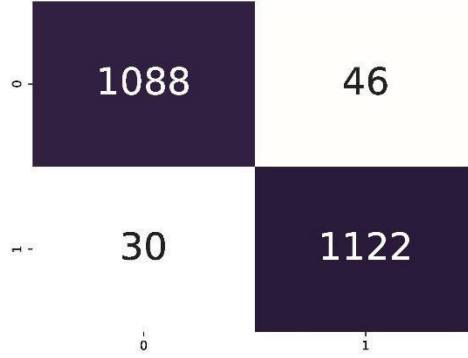
Geliştirilen sistemin sınıflandırma aşamasında mevcut algoritmalar kıyaslanarak en yüksek doğruluğu veren rasgele orman algoritması, gardiyan artırma algoritması ve ekstra ağaç algoritmasından oluşan oylama tabanlı topluluk öğrenme tabanlı bir makine öğrenmesi önerilmiştir. Bu doğrultuda önerilen sistemin mimari yapısı Resim 4.2'de verilmiştir.



Resim 4.2 Önerilen sınıflandırma modeli

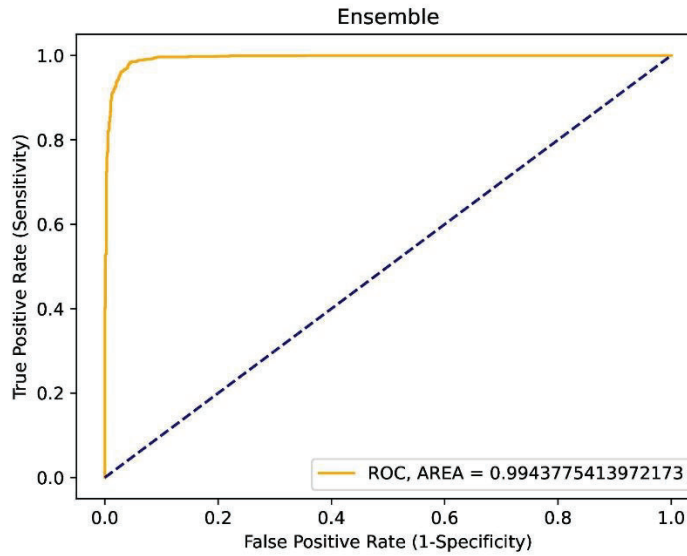
Topluluk öğrenmenin temelinde yer alan makine öğrenmesi algoritmalarında varsayılan hiper parametreler kullanılırken oylama aşamasının parametresi olarak soft oylama ve uniform ağırlıklandırma kullanılmıştır. Sonuçlara ilişkin değerlendirmeler karışıklık

matrisi ile değerlendirilmiştir. Resim 4.3’de sonuçlara ilişkin karışıklık matrisi verilmiştir.



Resim 4.3 Karışıklık matrisi

Grafikte satırlar doğru değerleri gösterirken sütunlar tahmin edilen değerleri göstermektedir. Buna göre 1134 zararsız web adresinden 1088 doğru tahmin edilirken 1152 zararlı web adresinden 1122’si doğru tahmin edilmiştir. Toplamda 2286 test verisinden 2210 tanesi doğru sınıflandırılarak %96.6 doğruluk elde edilmiştir. Önerilen yöntem ayrıca ROC eğrisi ve Eğri altında kalan alan değerleriyle de değerlendirilmiş ve %99.43 eğri altında kalan alan oranına ulaşmıştır. Resim 4.4’de ROC eğrisi verilmiştir.



Resim 4.4 ROC Eğrisi

Önerilen yöntem diğer klasik makine öğrenmesinden elde edilen sonuçlarla da karşılaştırılmıştır. Çizelge 4.3’de karşılaştırmalı analiz sunulmuştur. Önerilen yöntem diğer klasik tüm makine öğrenmesi algoritmalarından daha iyi performans göstermiştir.

Çizelge 4.3 Veri kümesinin gruplara göre örneklem sayıları

Algoritma	Doğruluk	Kesinlik	Duyarlılık	F1 Skor
K En Yakın Komşu	0.948	0.935	0.964	0.950
Lojistik Regresyon	0.947	0.943	0.955	0.949
Rasgele Orman	0.965	0.964	0.966	0.965
Destek Vektör Makineleri	0.960	0.952	0.970	0.961
Karar Ağaçları	0.941	0.937	0.949	0.943
Gaussian Naive Bayes	0.685	0.620	0.970	0.756
Lineer Diskriminant	0.931	0.930	0.934	0.932
AdaBoost	0.951	0.943	0.961	0.952
Gradyan Artırma	0.960	0.952	0.970	0.961
Ekstra Ağaç	0.964	0.955	0.975	0.965
Önerilen Yöntem	0.966	0.960	0.975	0.967

5. TARTIŞMA ve SONUÇ

İnternet dünyası ve bu dünyanın en önemli parçalarından olan finansal teknolojiler insan yaşamının önemli bir kısmını kapsamaya başlamıştır. Banka şubelerinin artık nakit para dahi bulundurmağı günümüz dünyasında, gerçek yaşamdaki suç tekniklerinin dijital dünyaya aktarılmaması beklenemez. Hatta gerçek dünyadaki yöntemlerin aktarılmasının yanı sıra bu dünyaya has tekniklerin gelişmesi de beklenecek bir durumdur.

Yapılan işleri kolaylaştıran internet uygulamaları dikkatli olmayı da beraberinde getiriyor. Kullanıcı dikkatinin pek çok finansal kaybı önleyebileceğı yaşanan vakaların analizlerinde görülebilmekle birlikte, her internet kullanıcısından bu dikkati beklemek doğru olmaz. İnternetin ilk dönemlerinde her kullanıcı kurslar vasıtası ile minimum bir bilgi düzeyine ulaşmak ile birlikte, ortalama saldırıları da günümüzdeki kadar yaygın ve karmaşık değildi.

Mobil cihazların gelişimi ile birlikte kullanımı artan çevrimiçi uygulamalar bilinçsiz kullanıcılarında bu dünyaya hızlı bir adım atmasına sebep olmuştur. Bu hızlı giriş saldırganların gözünden yeni bir gelir kapısıyken, siber güvenlik alanında çalışanlar içinse zorlu bir yolun açılması olarak görülebilir.

Sürekli gelişmesi sebebiyle insan çabası ile engellenmesi mümkün olmayan bu saldırılar günümüzün yükselen trendi olan yapay zeka uygulamaları ile daha başarılı ve efektif olarak engellenebilir.

Bilgi işlem sektöründe artan iş gücü ihtiyacı, insan ile ilgili pek çok sorunu da beraberinde getirdiğı ve insanın hata yapmaya açık olması sebebiyle bu tür yüksek hacimli verinin analizini gerektiren işlerin yapay zeka uygulamaları ile çözümünün hem kısa, hem uzun vadede ekonomik ve başarılı sonuçlar getireceğini değerlendirmekteyiz.

Zararlı web sayfalarında oluşan veri seti üzerinde yapılan çalışmanın %96.6 oranında başarı göstermesi, sistemin geliştirilerek diğer ortalama yöntemleri (e-posta, sosyal mühendislik gibi) üzerinde de başarı göstereceğı düşünülmektedir.

Ayrıca çalışma tek başına kullanılan klasik modellerin yerine yığın topluluk öğrenme modellerinin daha başarılı olabileceği yönünde de göstergeler barındırmaktadır. Bu durum bulgular bölümünde yapılan karşılaştırma ile ortaya konmuştur. İleriki çalışmalarda sınıflandırma modellerinde yapılacak değişiklikler ile sistem başarımının geliştirilebileceği düşünülmektedir.

Çalışmamızda kullanılan topluluk öğrenme sistemleri ile var olan öğrenme yöntemleri birleştirilerek mevcut sistemlerden daha isabetli sonuçlar elde edilmeye çalışılmıştır. Önerilen sistemin kullanılan mevcut sistemlere entegre edilmesi ile ortalama saldırılarının daha isabetli tespiti ve önlenmesinin sağlanacağı ve modelin gelişime açık olması sebebiyle bu sonuçlarda süreklilik sağlanacağı değerlendirilmektedir.

6. KAYNAKLAR

- Abad C, 2005, "The Economy of Phishing: A Survey of the Operations of the Phishing Market". *First Monday*, Eylül. <https://doi.org/10.5210/fm.v10i9.1272>.
- Abdelhamid N, Ayesh A, Thabtah F, 2014, "Phishing detection based Associative Classification data mining". *Expert Systems with Applications* 41 (13): 5948-59. <https://doi.org/10.1016/j.eswa.2014.03.019>.
- Afroz S, Greenstadt R, 2011, "PhishZoo: Detecting Phishing Websites by Looking at Them". İçinde *2011 IEEE Fifth International Conference on Semantic Computing*, 368-75. <https://doi.org/10.1109/ICSC.2011.52>.
- Aleroud A, Zhou L, 2017, "Phishing environments, techniques, and countermeasures: A survey". *Computers & Security* 68 (Temmuz):160-96. <https://doi.org/10.1016/j.cose.2017.04.006>.
- Alkhalil Z, Hewage C, Nawaf L, Khan I, 2021, "Phishing Attacks: A Recent Comprehensive Study and a New Anatomy". *Frontiers in Computer Science* 3 (Mart). <https://doi.org/10.3389/fcomp.2021.563060>.
- Alkhozae M G, Batarfi O A, 2011, "Phishing Websites Detection Based on Phishing Characteristics in the Webpage Source Code" 1 (6).
- Almomani A, Gupta B B, Atawneh S, Meulenberg A, Almomani E, 2013, "A Survey of Phishing Email Filtering Techniques". *IEEE Communications Surveys & Tutorials* 15 (4): 2070-90. <https://doi.org/10.1109/SURV.2013.030713.00020>.
- Alnajim A, Munro M, 2009, "An Anti-Phishing Approach that Uses Training Intervention for Phishing Websites Detection". İçinde *2009 Sixth International Conference on Information Technology: New Generations*, 405-10. <https://doi.org/10.1109/ITNG.2009.109>.
- Alsharnouby M, Alaca F, Chiasson S, 2015, "Why phishing still works: User strategies for combating phishing attacks". *International Journal of Human-Computer Studies* 82 (Ekim):69-82. <https://doi.org/10.1016/j.ijhcs.2015.05.005>.
- "APWG | Phishing Activity Trends Reports", t.y. Erişim 28 Temmuz 2024. <https://apwg.org/trendsreports/>.

- Arachchilage N A G, Love S, Beznosov K, 2016, "Phishing threat avoidance behaviour: An empirical investigation". *Computers in Human Behavior* 60 (Temmuz):185-97. <https://doi.org/10.1016/j.chb.2016.02.065>.
- Banu D M N, Banu S M, 2013, "A Comprehensive Study of Phishing Attacks" 4.
- Basit A, Zafar M, Liu X, Javed A R, Jalil Z, Kifayat K, 2021, "A Comprehensive Survey of AI-Enabled Phishing Attacks Detection Techniques". *Telecommunication Systems* 76 (1): 139-54. <https://doi.org/10.1007/s11235-020-00733-2>.
- Basnet R, Mukkamala S, Sung A H, 2008, "Detection of Phishing Attacks: A Machine Learning Approach". İçinde *Soft Computing Applications in Industry*, editör Bhanu Prasad, 373-83. Berlin, Heidelberg: Springer. https://doi.org/10.1007/978-3-540-77465-5_19.
- Baykara M, Gürel Z Z, 2018, "Detection of phishing attacks". İçinde *2018 6th International Symposium on Digital Forensic and Security (ISDFS)*, 1-5. <https://doi.org/10.1109/ISDFS.2018.8355389>.
- Bergholz A, Chang J H, Paass G, Reichartz F, Strobel S, 2008, *Improved Phishing Detection using Model-Based Features*.
- Bergholz A, De Beer J, Glahn S, Moens M-F, Paaß G, Strobel S, 2010, "New Filtering Approaches for Phishing Email". *Journal of Computer Security* 18 (1): 7-35. <https://doi.org/10.3233/JCS-2010-0371>.
- Bhardwaj A, Sapra V, Kumar A, Kumar N, Arthi S, 2020, "Why is phishing still successful?" *Computer Fraud & Security* 2020 (9): 15-19. [https://doi.org/10.1016/S1361-3723\(20\)30098-1](https://doi.org/10.1016/S1361-3723(20)30098-1).
- Bhavsar V, Kadlak A, Sharma S, 2018, "Study on Phishing Attacks". *International Journal of Computer Applications* 182 (Aralık):27-29. <https://doi.org/10.5120/ijca2018918286>.
- Breiman L, 1996, "Bagging Predictors". *Machine Learning* 24 (2): 123-40. <https://doi.org/10.1007/BF00058655>.
- , 2001, "Random Forests". *Machine Learning* 45 (1): 5-32. <https://doi.org/10.1023/A:1010933404324>.

- Bullee J-W, Montoya L, Junger M, Hartel P, 2017, "Spear phishing in organisations explained". *Information & Computer Security* 25 (5): 593-613. <https://doi.org/10.1108/ICS-03-2017-0009>.
- Burns A J, Johnson M E, Caputo D D, 2019, "Spear phishing in a barrel: Insights from a targeted phishing campaign". *Journal of Organizational Computing and Electronic Commerce* 29 (1): 24-39. <https://doi.org/10.1080/10919392.2019.1552745>.
- Butavicius M, Parsons K, Pattinson M, McCormac A, 2016, "Breaching the Human Firewall: Social engineering in Phishing and Spear-Phishing Emails". arXiv. <https://doi.org/10.48550/arXiv.1606.00887>.
- Chen K-T, Chen J-Y, Huang C-R, Chen C-S, 2009, "Fighting Phishing with Discriminative Keypoint Features". *IEEE Internet Computing* 13 (3): 56-63. <https://doi.org/10.1109/MIC.2009.59>.
- Chen R, Gaia J, Rao H R, 2020, "An examination of the effect of recent phishing encounters on phishing susceptibility". *Decision Support Systems* 133 (Haziran):113287. <https://doi.org/10.1016/j.dss.2020.113287>.
- Chiew K L, Chang E H, Sze S N, Tiong W K, 2015, "Utilisation of website logo for phishing detection". *Computers & Security, Secure Information Reuse and Integration & Availability, Reliability and Security 2014*, 54 (Ekim):16-26. <https://doi.org/10.1016/j.cose.2015.07.006>.
- Cunningham P, Cord M, Delany S J, 2008, "Supervised Learning". İçinde *Machine Learning Techniques for Multimedia: Case Studies on Organization and Retrieval*, editör Matthieu Cord ve Pádraig Cunningham, 21-49. Cognitive Technologies. Berlin, Heidelberg: Springer. https://doi.org/10.1007/978-3-540-75171-7_2.
- Darwish A, Zarka A E, Aloul F, 2012, "Towards understanding phishing victims' profile". İçinde *2012 International Conference on Computer Systems and Industrial Informatics*, 1-5. <https://doi.org/10.1109/ICCSII.2012.6454454>.
- Dhamija R, Tygar J D, Hearst M, 2006, "Why phishing works". İçinde *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, 581-90. CHI

- '06. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/1124772.1124861>.
- Dodge R C, Carver C, Ferguson A J, 2007, "Phishing for user security awareness". *Computers & Security* 26 (1): 73-80. <https://doi.org/10.1016/j.cose.2006.10.009>.
- Downs J S, Holbrook M B, Cranor L F, 2006, "Decision strategies and susceptibility to phishing". İçinde *Proceedings of the second symposium on Usable privacy and security*, 79-90. SOUPS '06. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/1143120.1143131>.
- Downs J S, Holbrook M, Cranor L F, 2007, "Behavioral response to phishing risk". İçinde *Proceedings of the anti-phishing working groups 2nd annual eCrime researchers summit*, 37-44. eCrime '07. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/1299015.1299019>.
- Dunlop M, Groat S, Shelly D, 2010, "GoldPhish: Using Images for Content-Based Phishing Analysis". İçinde *2010 Fifth International Conference on Internet Monitoring and Protection*, 123-28. <https://doi.org/10.1109/ICIMP.2010.24>.
- El Naqa I, Murphy M J, 2015, "What Is Machine Learning?" İçinde *Machine Learning in Radiation Oncology: Theory and Applications*, editör Issam El Naqa, Ruijiang Li, ve Martin J. Murphy, 3-11. Cham: Springer International Publishing. https://doi.org/10.1007/978-3-319-18305-3_1.
- Ertel W, 2018, *Introduction to Artificial Intelligence*. Springer.
- Fatima R, Yasin A, Liu L, Wang J, 2019, "How Persuasive Is a Phishing Email? A Phishing Game for Phishing Awareness". *Journal of Computer Security* 27 (6): 581-612. <https://doi.org/10.3233/JCS-181253>.
- Fette I, Sadeh N, Tomasic A, 2007, "Learning to detect phishing emails". İçinde *Proceedings of the 16th international conference on World Wide Web*, 649-56. WWW '07. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/1242572.1242660>.
- Fetzer J H, 1990, "What Is Artificial Intelligence?" İçinde *Artificial Intelligence: Its Scope and Limits*, editör James H. Fetzer, 3-27. Studies in Cognitive Systems. Dordrecht: Springer Netherlands. https://doi.org/10.1007/978-94-009-1900-6_1.

- Finn P, Jakobsson M, 2007, "Designing ethical phishing experiments". *IEEE Technology and Society Magazine* 26 (1): 46-58. <https://doi.org/10.1109/MTAS.2007.335565>.
- Freund Y, Schapire R E, 1997, "A Decision-Theoretic Generalization of On-Line Learning and an Application to Boosting". *Journal of Computer and System Sciences* 55 (1): 119-39. <https://doi.org/10.1006/jcss.1997.1504>.
- Garera S, Provos N, Chew M, Rubin A D, 2007, "A framework for detection and measurement of phishing attacks". İçinde *Proceedings of the 2007 ACM workshop on Recurring malware*, 1-8. WORM '07. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/1314389.1314391>.
- Ghahramani Z, 2004, "Unsupervised Learning". İçinde *Advanced Lectures on Machine Learning: ML Summer Schools 2003, Canberra, Australia, February 2 - 14, 2003, Tübingen, Germany, August 4 - 16, 2003, Revised Lectures*, editör Olivier Bousquet, Ulrike von Luxburg, ve Gunnar Rätsch, 72-112. Lecture Notes in Computer Science. Berlin, Heidelberg: Springer. https://doi.org/10.1007/978-3-540-28650-9_5.
- Gupta S, Singhal A, Kapoor A, 2016, "A literature survey on social engineering attacks: Phishing attack". İçinde *2016 International Conference on Computing, Communication and Automation (ICCCA)*, 537-40. <https://doi.org/10.1109/CCAA.2016.7813778>.
- Halevi T, Lewis J, Memon N, 2013, "Phishing, Personality Traits and Facebook". arXiv. <https://doi.org/10.48550/arXiv.1301.7643>.
- Halevi T, Memon N, Nov O, 2015, "Spear-Phishing in the Wild: A Real-World Study of Personality, Phishing Self-Efficacy and Vulnerability to Spear-Phishing Attacks". SSRN Scholarly Paper. Rochester, NY. <https://doi.org/10.2139/ssrn.2544742>.
- Han X, Kheir N, Balzarotti D, 2016, "PhishEye: Live Monitoring of Sandboxed Phishing Kits". İçinde *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*, 1402-13. CCS '16. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/2976749.2978330>.
- Harrison B, Svetieva E, Vishwanath A, 2016, "Individual processing of phishing emails:

- How attention and elaboration protect against phishing”. *Online Information Review* 40 (2): 265-81. <https://doi.org/10.1108/OIR-04-2015-0106>.
- He M, Horng S-J, Fan P, Khan M K, Run R-S, Lai J-L, Chen R-J, Sutanto A, 2011, “An efficient phishing webpage detector”. *Expert Systems with Applications* 38 (10): 12018-27. <https://doi.org/10.1016/j.eswa.2011.01.046>.
- Irani D, Webb S, Giffin J, Pu C, 2008, “Evolutionary study of phishing”. İçinde *2008 eCrime Researchers Summit*, 1-10. <https://doi.org/10.1109/ECRIME.2008.4696967>.
- Jain A K, Gupta B B, 2017, “Phishing Detection: Analysis of Visual Similarity Based Approaches”. *Security and Communication Networks* 2017 (1): 5421046. <https://doi.org/10.1155/2017/5421046>.
- Jakobsson M, 2005, “Modeling and Preventing Phishing Attacks”. İçinde *Financial Cryptography and Data Security*, editör Andrew S. Patrick ve Moti Yung, 3570:89-89. Lecture Notes in Computer Science. Berlin, Heidelberg: Springer Berlin Heidelberg. https://doi.org/10.1007/11507840_9.
- Jakobsson M, Myers S, 2006, *Phishing and Countermeasures: Understanding the Increasing Problem of Electronic Identity Theft*. John Wiley & Sons.
- Jakobsson M, Tsow A, Shah A, Blevis E, Lim Y-K, 2007, “What Instills Trust? A Qualitative Study of Phishing”. İçinde *Financial Cryptography and Data Security*, editör Sven Dietrich ve Rachna Dhamija, 356-61. Berlin, Heidelberg: Springer. https://doi.org/10.1007/978-3-540-77366-5_32.
- James L, 2005, *Phishing Exposed*. Elsevier.
- Jansson K, Solms R von, 2013, “Phishing for phishing awareness”. *Behaviour & Information Technology* 32 (6): 584-93. <https://doi.org/10.1080/0144929X.2011.632650>.
- Khonji M, Iraqi Y, Jones A, 2013, “Phishing Detection: A Literature Survey”. *IEEE Communications Surveys & Tutorials* 15 (4): 2091-2121. <https://doi.org/10.1109/SURV.2013.032213.00009>.
- Kirda E, Kruegel C, 2005, “Protecting users against phishing attacks with AntiPhish”.

- İçinde *29th Annual International Computer Software and Applications Conference (COMPSAC'05)*, 1:517-524 Vol. 2. <https://doi.org/10.1109/COMPSAC.2005.126>.
- Krombholz K, Hobel H, Huber M, Weippl E, 2015, "Advanced social engineering attacks". *Journal of Information Security and Applications*, Special Issue on Security of Information and Networks, 22 (Haziran):113-22. <https://doi.org/10.1016/j.jisa.2014.09.005>.
- Kumaraguru P, Cranshaw J, Acquisti A, Cranor L, Hong J, Blair M A, Pham T, 2009, "School of phish: a real-world evaluation of anti-phishing training". İçinde *Proceedings of the 5th Symposium on Usable Privacy and Security*, 1-12. SOUPS '09. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/1572532.1572536>.
- Kwak Y, Lee S, Damiano A, Vishwanath A, 2020, "Why do users not report spear phishing emails?" *Telematics and Informatics* 48 (Mayıs):101343. <https://doi.org/10.1016/j.tele.2020.101343>.
- Lam I-F, Xiao W-C, Wang S-C, Chen K-T, 2009, "Counteracting Phishing Page Polymorphism: An Image Layout Analysis Approach". İçinde *Advances in Information Security and Assurance*, editör Jong Hyuk Park, Hsiao-Hwa Chen, Mohammed Atiquzzaman, Changhoon Lee, Tai-hoon Kim, ve Sang-Soo Yeo, 270-79. Berlin, Heidelberg: Springer. https://doi.org/10.1007/978-3-642-02617-1_28.
- Liu G, Qiu B, Wenyin L, 2010, "Automatic Detection of Phishing Target from Phishing Webpage". İçinde *2010 20th International Conference on Pattern Recognition*, 4153-56. <https://doi.org/10.1109/ICPR.2010.1010>.
- Ludl C, McAllister S, Kirda E, Kruegel C, 2007, "On the Effectiveness of Techniques to Detect Phishing Sites". İçinde *Detection of Intrusions and Malware, and Vulnerability Assessment*, editör Bernhard M. Hämmerli ve Robin Sommer, 20-39. Berlin, Heidelberg: Springer. https://doi.org/10.1007/978-3-540-73614-1_2.
- Medvet E, Kirda E, Kruegel C, 2008, "Visual-similarity-based phishing detection". İçinde *Proceedings of the 4th international conference on Security and privacy in*

- communication networks*, 1-6. SecureComm '08. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/1460877.1460905>.
- Merwe A van der, Seker R, Gerber A, 2005, "Phishing in the system of systems settings: mobile technology". İçinde *2005 IEEE International Conference on Systems, Man and Cybernetics*, 1:492-498 Vol. 1. <https://doi.org/10.1109/ICSMC.2005.1571194>.
- Moghimi M, Varjani A Y, 2016, "New rule-based phishing detection method". *Expert Systems with Applications* 53 (Temmuz):231-42. <https://doi.org/10.1016/j.eswa.2016.01.028>.
- Mohammad R M, Thabtah F, McCluskey L, 2014, "Intelligent Rule-Based Phishing Websites Classification". *IET Information Security* 8 (3): 153-60. <https://doi.org/10.1049/iet-ifs.2013.0202>.
- Mohebzada J G, Zarka A E, Bhojani A H, Darwish A, 2012, "Phishing in a university community: Two large scale phishing experiments". İçinde *2012 International Conference on Innovations in Information Technology (IIT)*, 249-54. <https://doi.org/10.1109/INNOVATIONS.2012.6207742>.
- Moore T, Clayton R, 2007, "Examining the impact of website take-down on phishing". İçinde *Proceedings of the anti-phishing working groups 2nd annual eCrime researchers summit*, 1-13. eCrime '07. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/1299015.1299016>.
- Nasteski V, 2017, "An overview of the supervised machine learning methods". *HORIZONS.B* 4 (Aralık):51-62. <https://doi.org/10.20544/HORIZONS.B.04.1.17.P05>.
- Orunsolu A A, Sodiya A S, Akinwale A T, 2022, "A predictive model for phishing detection". *Journal of King Saud University - Computer and Information Sciences* 34 (2): 232-47. <https://doi.org/10.1016/j.jksuci.2019.12.005>.
- Pan Y, Ding X, 2006, "Anomaly Based Web Phishing Page Detection". İçinde *2006 22nd Annual Computer Security Applications Conference (ACSAC'06)*, 381-92. <https://doi.org/10.1109/ACSAC.2006.13>.

- Parmar B, 2012, "Protecting against spear-phishing". *Computer Fraud & Security* 2012 (1): 8-11. [https://doi.org/10.1016/S1361-3723\(12\)70007-6](https://doi.org/10.1016/S1361-3723(12)70007-6).
- Parno B, Kuo C, Perrig A, 2006, "Phoolproof Phishing Prevention". İçinde *Financial Cryptography and Data Security*, editör Giovanni Di Crescenzo ve Avi Rubin, 1-19. Berlin, Heidelberg: Springer. https://doi.org/10.1007/11889663_1.
- Prakash P, Kumar M, Kompella R R, Gupta M, 2010, "PhishNet: Predictive Blacklisting to Detect Phishing Attacks". İçinde *2010 Proceedings IEEE INFOCOM*, 1-5. <https://doi.org/10.1109/INFCOM.2010.5462216>.
- Rao R S, Ali S T, 2015, "A Computer Vision Technique to Detect Phishing Attacks". İçinde *2015 Fifth International Conference on Communication Systems and Network Technologies*, 596-601. <https://doi.org/10.1109/CSNT.2015.68>.
- Rekouche K, 2011, "Early Phishing". arXiv. <https://doi.org/10.48550/arXiv.1106.4692>.
- Rich E, 1983, *Artificial Intelligence*. McGraw-Hill.
- Shahrivari V, Darabi M M, Izadi M, 2020, "Phishing Detection Using Machine Learning Techniques". arXiv. <https://doi.org/10.48550/arXiv.2009.11116>.
- Tuğal İ, Almaz C, Sevi M, 2021, "Üniversitelerdeki Siber Güvenlik Sorunları ve Farkındalık Eğitimleri". *Bilişim Teknolojileri Dergisi* 14 (3): 229-38. <https://doi.org/10.17671/gazibtd.754458>.
- Verma R, Shashidhar N, Hossain N, 2012, "Detecting Phishing Emails the Natural Language Way". İçinde *Computer Security – ESORICS 2012*, editör Sara Foresti, Moti Yung, ve Fabio Martinelli, 824-41. Berlin, Heidelberg: Springer. https://doi.org/10.1007/978-3-642-33167-1_47.
- Wen Z A, Lin Z, Chen R, Andersen E, 2019, "What.Hack: Engaging Anti-Phishing Training Through a Role-playing Phishing Simulation Game". İçinde *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*, 1-12. CHI '19. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/3290605.3300338>.
- Wenyin L, Huang G, Xiaoyue L, Min Z, Deng X, 2005, "Detection of phishing webpages based on visual similarity". İçinde *Special interest tracks and posters of the 14th*

- international conference on World Wide Web*, 1060-61. WWW '05. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/1062745.1062868>.
- Williams E J, Hinds J, Joinson A N, 2018, "Exploring susceptibility to phishing in the workplace". *International Journal of Human-Computer Studies* 120 (Aralık):1-13. <https://doi.org/10.1016/j.ijhcs.2018.06.004>.
- Wolpert D H, 1992, "Stacked generalization". *Neural Networks* 5 (2): 241-59. [https://doi.org/10.1016/S0893-6080\(05\)80023-1](https://doi.org/10.1016/S0893-6080(05)80023-1).
- Wu M, Miller R C, Garfinkel S L, 2006, "Do security toolbars actually prevent phishing attacks?" İçinde *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, 601-10. CHI '06. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/1124772.1124863>.
- Yi P, Guan Y, Zou F, Yao Y, Wang W, Zhu T, 2018, "Web Phishing Detection Using a Deep Learning Framework". *Wireless Communications and Mobile Computing* 2018 (1): 4678746. <https://doi.org/10.1155/2018/4678746>.
- Yu W D, Nargundkar S, Tiruthani N, 2009, "PhishCatch - A Phishing Detection Tool". İçinde *2009 33rd Annual IEEE International Computer Software and Applications Conference*, 2:451-56. <https://doi.org/10.1109/COMPSAC.2009.175>.
- Zhang C, 2019, "Research on the Fluctuation and Factors of China TFP of IT Industry". *Journal of Industrial Integration and Management*, Aralık. <https://doi.org/10.1142/S2424862219500131>.
- Zhang C, Lu Y, 2021, "Study on artificial intelligence: The state of the art and future prospects". *Journal of Industrial Information Integration* 23 (Eylül):100224. <https://doi.org/10.1016/j.jii.2021.100224>.
- Zhou Z-H, 2021, "Ensemble Learning". İçinde *Machine Learning*, editör Zhi-Hua Zhou, 181-210. Singapore: Springer. https://doi.org/10.1007/978-981-15-1967-3_8.
- Zhu X, Goldberg A B, 2009, *Introduction to Semi-Supervised Learning*. Synthesis Lectures on Artificial Intelligence and Machine Learning. Cham: Springer International Publishing. <https://doi.org/10.1007/978-3-031-01548-9>.

Zieni R, Massari L, Calzarossa M C, 2023, “Phishing or Not Phishing? A Survey on the Detection of Phishing Websites”. *IEEE Access* 11:18499-519. <https://doi.org/10.1109/ACCESS.2023.3247135>.

Zuraiq A A, Alkasassbeh M, 2019, “Review: Phishing Detection Approaches”. İçinde *2019 2nd International Conference on new Trends in Computing Sciences (ICTCS)*, 1-6. <https://doi.org/10.1109/ICTCS.2019.8923069>.

