


```
</lan>  
</lans>  
</cluster>
```

Le certificat serveur et la configuration du cluster doivent contenir la même liste de valeurs (noms DNS et/ou adresses IP). Si ce n'est pas le cas, la console web SafeKit et les commandes distribuées ne fonctionneront pas correctement.

Pour vérifier que cela est bien le cas :

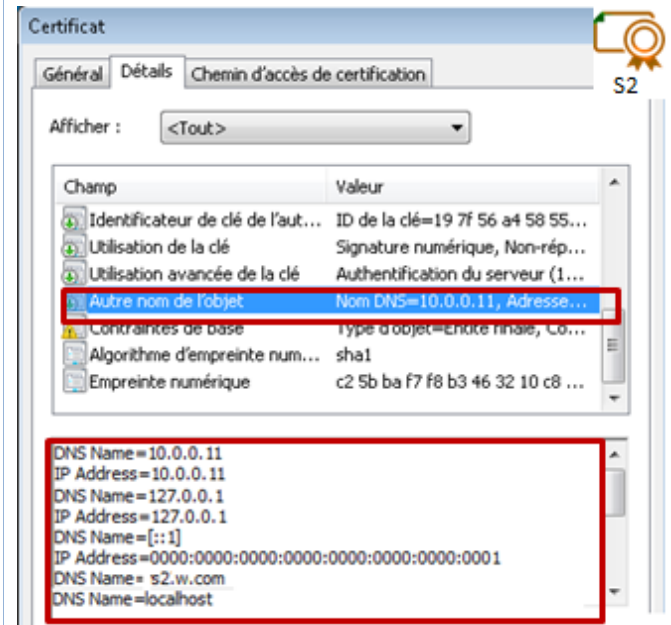
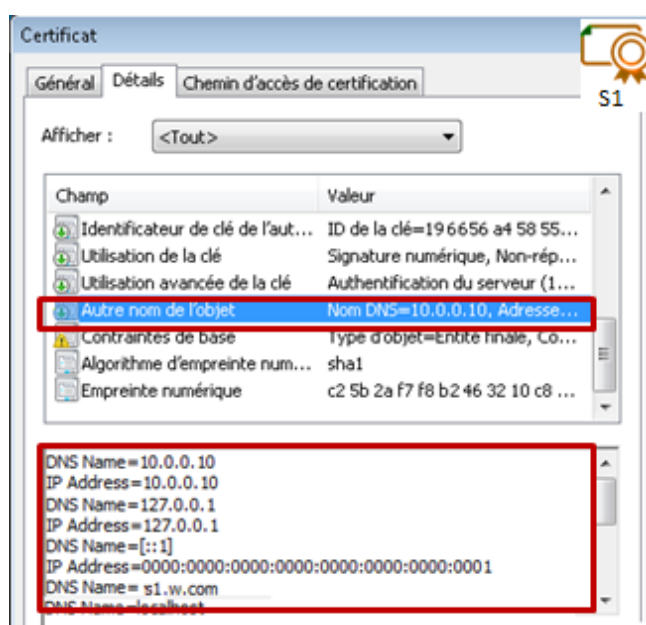
1. Copier le fichier .crt (ou .cer) sur une station de travail Windows
2. Double cliquer sur le fichier pour l'ouvrir avec *Extension noyau de chiffrement*
3. Cliquer sur l'onglet *Détails*
4. Vérifier le contenu du champ *Autre nom de l'objet* (Subject Alternative Name)



Si vous préférez utiliser la ligne de commande, exécutez sur chaque nœud :

```
SAFE/web/bin/openssl.exe x509 -text -noout -in SAFE/web/conf/server.crt
```

Et vérifier le contenu de la valeur après Subject Alternative Name



⇒ localhost et 127.0.0.1 doivent être présents

⇒ avec SafeKit <= 7.5.2.9, le nom du serveur doit être présent

11.3.2.2 Récupérer et installer le certificat CA

11.3.2.2.1 Récupérer le fichier du certificat

Vous devez récupérer le certificat de l'Autorité de Certification CA (chaîne de certificats pour la CA racine et les intermédiaires, s'il y en a) utilisé pour générer les certificats serveur de S1 et S2. Le format attendu est le suivant :

 CA cacert.crt	Le certificat de l'Autorité de Certification CA utilisée pour générer les certificats serveur. ⇒ fichier pour le certificat X509 dans le format PEM La chaîne de certificats pour la CA racine et les intermédiaires, s'il y en a	 S1 S2 Certificats serveur pour S1 et S2
---	--	---

Si vous rencontrez des difficultés pour récupérer ce fichier depuis votre PKI, vous pouvez le construire à l'aide de la procédure décrite en 7.18. [page 131](#).

11.3.2.2.2 Installer le fichier dans SafeKit

Installer le certificat comme suit (SAFE=C:\safekit en Windows si System Drive=C: ; et SAFE=/opt/safekit en Linux) :

 CA cacert.crt	Sur S1 et S2 : ⇒ copier cacert.crt dans SAFE/web/conf/cacert.crt
---	--

Vous pouvez contrôler l'installation avec :

```
cd SAFE/web/bin
checkcert -t CA
```

Cette commande retourne en échec si une erreur est détectée.

Vous devez également vérifier que le fichiers cacert.crt contient bien la chaîne de certificats pour les autorités de certification racine et intermédiaires.

11.3.2.3 Configurer et redémarrer le service web HTTPS

Pour activer la configuration HTTPS (SAFE=C:\safekit en Windows si System Drive=C: ; et SAFE=/opt/safekit en Linux):

 httpd.webconsolessl.conf	Sur S1 et S2 : ⇒ copier SAFE/web/conf/httpd.webconsolessl.conf dans SAFE/web/conf/ssl/httpd.webconsolessl.conf
---	--

	Sur S1 et S2 : ⇒ exécuter <code>safekit webserver restart</code>
--	---

11.3.2.4 Changer les règles du pare-feu

Vous pouvez exécuter le script `firewallcfg` pour appliquer les règles pour SafeKit au pare-feu du système d'exploitation (en Windows, Microsoft Windows Firewall ; en Linux, `firewalld` ou `iptables`).

Pare-feu	Sur S1 et S2 : ⇒ aller dans le répertoire <code>SAFEBIN</code> (=SAFE/privatebin) ⇒ exécuter <code>firewallcfg add</code>
----------	---

N'exécutez pas cette commande si vous souhaitez configurer vous-même le pare-feu ou si vous utilisez un pare-feu différent de celui du système. Pour la liste des processus et ports de SafeKit, voir 10.3 [page 160](#).

11.4 Configuration de l'authentification utilisateur

Mettez en œuvre l'une des méthodes suivantes pour l'authentification utilisateur :

- ⇒ 11.4.1 « Configuration l'authentification à base de fichier » [page 196](#)
- ⇒ 11.4.2 « Configuration de l'authentification à base de serveur LDAP/AD » [page 199](#)
- ⇒ 11.4.3 « Configuration de l'authentification à base de certificat client avec la PKI SafeKit » [page 202](#)
- ⇒ 11.4.4 « Configuration de l'authentification à base de certificat client avec une PKI externe » [page 209](#)

A l'issue de cette configuration, vous pouvez utiliser la console web sécurisée.

11.4.1 Configuration l'authentification à base de fichier

L'authentification à base de fichier peut être appliquée en HTTP ou HTTPS. Elle repose sur les fichiers suivants :

 user.conf	Contrôle d'accès à partir du fichier des utilisateurs
 group.conf	Configuration facultative pour restreindre le rôle de l'utilisateur. Si le fichier <code>group.conf</code> n'est pas présent, tous les utilisateurs authentifiés auront le rôle Admin.

11.4.1.1 Gérer les utilisateurs et groupes

Les utilisateurs et groupes doivent être identiques sur S1 et S2, ainsi que les mots de passe. Ils sont définis par les fichiers `user.conf` et `group.conf` dans le répertoire `SAFE/web/conf` (`SAFE=C:\safekit` en Windows si `%SYSTEMDRIVE%=C: ;` et `SAFE=/opt/safekit` en Linux).



Pendant l'initialisation de la configuration par défaut, décrite dans 11.2.1 page 183, l'utilisateur nommé `admin` a été créé et est donc présent dans `user.conf`. Vous pouvez décider de le supprimer si vous en créez d'autres.

⇒ Création d'un utilisateur

Les utilisateurs sont créés avec la commande `SAFE/web/bin/htpasswd`.

Par exemple, pour ajouter le nouvel utilisateur `manager` et lui affecter son mot de passe à `managerpassword`, exécutez :

```
SAFE/web/bin/htpasswd -b SAFE/web/conf/user.conf manager managerpassword
```

Le nouvel utilisateur est inséré dans le fichier `SAFE/web/conf/user.conf` :



```
admin:$2y$05$0PquL6Z2Y78QcXpHIako.O58Z6lWfa5A86XD.eCbEnbRcguJln9Ce
manager:$apr1$U2GLivF5$x39WKmSpq6BGmLybESgNV1
operator1:$apr1$DetdwaZz$hy5pQzpU1Pny3qsXrIS/z1
operator2:$apr1$ICiZv2ru$wRkc3BclBhXzc/4llofocl
```

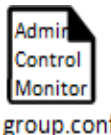
⇒ Affecter un rôle au nouvel utilisateur

Par défaut, tous les utilisateurs ont le rôle Admin. Si vous souhaitez affecter des rôles différents en fonction des utilisateurs, vous devez créer le fichier `SAFE/web/conf/group.conf` et y définir le rôle de chaque utilisateur. Ce fichier peut contenir les 3 groupes : Admin, Control, Monitor. Les utilisateurs auront le rôle correspondant au groupe auquel ils appartiennent.



Chaque ligne débute par le nom du groupe, suivi de :, suivi de la liste des utilisateurs (dont le séparateur est l'espace). Voir l'exemple ci-dessous.

Par exemple, pour affecter le rôle Control au nouvel utilisateur `manager` :



```
Admin : admin
Control : manager
Monitor : operator1 operator2
```



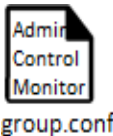
Si vous activez la gestion de rôle, vous devez insérer l'utilisateur `admin` dans `group.conf`. Sinon, cet utilisateur ne sera plus opérationnel.

⇒ Supprimer un utilisateur, ...

Exécuter `htpasswd -?` Pour lister toutes les options de gestion des utilisateurs.

11.4.1.2 Installer les fichiers

Installer les fichiers comme décrit ci-dessous (`SAFE=C:\safekit` en Windows si `%SYSTEMDRIVE%=C: ;` et `SAFE=/opt/safekit` en Linux):

	Sur S1 et S2 : ⇒ copier <code>user.conf</code> dans <code>SAFE/web/conf/user.conf</code>
	Sur S1 et S2, si les groupes sont définis : ⇒ copier <code>group.conf</code> dans <code>SAFE/web/conf/group.conf</code>

Ces fichiers doivent être identiques sur tous les nœuds.

11.4.1.3 Configurer et redémarrer le service web

Pour activer l'authentification à base de fichier (`SAFE=C:\safekit` en Windows si `%SYSTEMDRIVE%=C: ;` et `SAFE=/opt/safekit` en Linux) :

	Sur S1 et S2 : ⇒ éditer le fichier <code>SAFE/web/conf/httpd.conf</code> ⇒ décommenter <code>usefile</code> <pre>Define usefile</pre> ⇒ vérifier que <code>useldap</code> est commenté <pre># Define useldap</pre> ⇒ vérifier que seul <code>httpadminport</code> est défini <pre>httpadminport (9010) #httpcontrolport (9010) #httpmonitorport (9010)</pre>
	Sur S1 et S2 : ⇒ exécuter <code>safekit webserver restart</code>

11.4.1.4 Tester la console web et la commande distribuée

La configuration est terminée ; vous pouvez maintenant vérifier qu'elle est opérationnelle :

⇒ Tester la console web

1. Démarrer un navigateur web

2. Le connecter à l'URL <http://servername:9010> (où `servername` est l'adresse IP ou le nom d'un nœud SafeKit). Si HTTPS est configuré, il y a une redirection automatique sur <https://servername:9453>.
1. Dans la page de connexion, saisir le **Nom utilisateur** et le **Mot de passe**, puis cliquer sur **Connecter**. Avec la configuration par défaut de SafeKit 7.5, vous pouvez vous connecter avec l'utilisateur `admin` en donnant le mot de passe que vous lui avez attribué lors de l'initialisation.
2. La page chargée contient uniquement les onglets autorisés en fonction du rôle de l'utilisateur. Si les groupes n'ont pas été définis, tous les utilisateurs ont le rôle Admin.

⇒ Tester une commande distribuée

1. Se loguer sur S1 ou S2 en tant que administrateur/root
2. Ouvrir un terminal (PowerShell, shell, ...)
3. Aller dans le répertoire `SAFE`
4. Exécuter `safekit -H "*" level`
qui doit retourner le résultat de la commande `level` sur tous les nœuds

11.4.2 Configuration de l'authentification à base de serveur LDAP/AD

L'authentification LDAP/AD peut être appliquée en HTTP ou HTTPS. Elle repose sur :

	Contrôle d'accès à partir d'un compte LDAP/AD associé à l'utilisateur
	Configuration facultative des groupes LDAP/AD pour restreindre le rôle de l'utilisateur. Lorsque les groupes ne sont pas définis, tous les utilisateurs authentifiés ont le rôle Admin.



Sur certaines distributions Linux (telles que RedHat 8 et CentOS 8), le démarrage du service web échoue lorsqu'il est configuré avec l'authentification LDAP/AD. Dans ce cas, appliquer la solution décrite dans [SK-0092](#).

Appliquer les étapes décrites ci-dessous après avoir vérifié que S1 et S2 peuvent bien se connecter au port du domaine contrôleur LDAP (par défaut est 389).

11.4.2.1 Créer les utilisateurs et groupes

Si nécessaire, demandez à l'administrateur LDAP de créer les utilisateurs de la console web SafeKit.

Si vous souhaitez restreindre les accès en fonction des rôles, demandez à l'administrateur LDAP de créer les groupes Admin, Control, Monitor et d'affecter les utilisateurs au groupe adéquate. Si les groupes ne sont pas définis, tous les utilisateurs auront le rôle Admin.

11.4.2.2 Configurer et redémarrer le service web

Pour activer l'authentification LDAP/AD (`SAFE=C:\safekit` en Windows si `%SYSTEMDRIVE%=C: ;` et `SAFE=/opt/safekit` en Linux) :

	Sur S1 et S2 : Initialiser l'authentification pour la commande distribuée. Cela peut avoir déjà été fait si vous avez initialisé la configuration par défaut après l'installation de SafeKit. Sinon : ⇒ exécuter <code>webservercfg -rcmdpasswd pwd</code> où <code>pwd</code> est le mot de passe pour l'utilisateur privé <code>rcmdadmin</code>. Vous n'avez pas besoin de le mémoriser.
 httpd.conf	Sur S1 et S2 : ⇒ éditer le fichier <code>SAFE/web/conf/httpd.conf</code> ⇒ commenter <code>usefile</code> <pre># Define usefile</pre> ⇒ décommenter <code>uselldap</code> <pre>Define uselldap</pre> ⇒ vérifier que seul <code>httpadminport</code> est défini <pre>httpadminport (9010) #httpcontrolport (9010) #httpmonitorport (9010)</pre> ⇒ décommenter les lignes suivantes et remplacer les valeurs en gras par celles correspondant à la configuration de votre service LDAP/AD : <pre>Define binddn "CN=bindCN,OU=bindOU1,OU=bindOU2,DC=domain,DC=fq,DC=dn"</pre> <pre>Define bindpwd "Password0"</pre> <pre>Define searchurl "ldap://ldaporad.fq.dn:389/OU=searchou,DC=domain,DC=fq,DC=dn ?sAMAccountName,memberOf?sub?(objectClass=*)"</pre> ✓ les variables <code>binddn</code> et <code>bindpwd</code> doivent contenir les identifiants d'un compte possédant les droits de recherche dans le répertoire LDAP ✓ la variable <code>searchurl</code> définit l'url de recherche (au sens RFC2255) permettant d'authentifier l'utilisateur  Note <code>CN</code>: common name <code>OU</code>: organization unit

	DC: domain component (one field for each part of the FQDN) Si aucun groupe n'est défini, tous les utilisateurs authentifiés auront le rôle Admin.
	Sur S1 et S2 : Pour activer la gestion de groupes : ⇒ éditer le fichier <code>SAFE/web/conf/httpd.conf</code> ⇒ décommenter les lignes suivantes et remplacer les valeurs en gras par celles correspondant à la configuration de votre service LDAP/AD : <pre>Define admingroup "CN=Group1CN, OU=Group1OU1, OU=Group1OU2, DC=domain, DC=fq, DC=dn" Define controlgroup "CN=Group2CN, OU=Group2OU1, OU=Group2OU2, DC=domain, DC=fq, DC=dn" Define monitorgroup "CN=Group3CN, OU=Group3OU1, OU=Group3OU2, DC=domain, DC=fq, DC=dn"</pre> Les utilisateurs appartenant au groupe <code>admingroup</code>, <code>controlgroup</code> ou <code>monitorgroup</code>, auront respectivement les rôles Admin, Control et Monitor. Pour une configuration plus avancée, voir la documentation du service web Apache (voir http://httpd.apache.org).
	Sur S1 et S2 : ⇒ exécuter <code>safekit webserver restart</code>

11.4.2.3 Tester la console web et la commande distribuée

La configuration est terminée ; vous pouvez maintenant vérifier qu'elle est opérationnelle :

⇒ Tester la console web

3. Démarrer un navigateur web
4. Le connecter à l'URL `http://servername:9010` (où `servername` est l'adresse IP ou le nom d'un nœud SafeKit). Si HTTPS est configuré, il y a une redirection automatique sur `https://servername:9453`.
5. Dans la page de connexion, saisir le Nom utilisateur et le Mot de passe, puis cliquer sur `Connecter`
6. La page chargée contient uniquement les onglets autorisés en fonction du rôle de l'utilisateur. Si les groupes n'ont pas été configurés, tous les utilisateurs ont le rôle Admin.

⇒ Tester une commande distribuée

1. Se loguer sur S1 ou S2 en tant que administrateur/root
2. Ouvrir un terminal (PowerShell, shell, ...)
3. Aller dans le répertoire `SAFE`
4. Exécuter `safekit -H "*" level`

qui doit retourner le résultat de la commande `level` sur tous les nœuds

11.4.3 Configuration de l'authentification à base de certificat client avec la PKI SafeKit

Une fois la configuration de HTTPS terminée, comme décrit dans 11.3.1 [page 187](#), vous pouvez poursuivre avec la configuration de l'authentification à base de certificats clients :

1. La configuration du service web doit être modifiée pour activer l'authentification des certificats du client.
2. L'assistant de configuration des certificats client assiste dans la création et le téléchargement des certificats client pour les importer sur le poste de travail de l'utilisateur de la console. Pour cela, appliquer les sections de 11.4.3.2 à 11.4.3.6 sur chaque poste de travail des utilisateurs concernés.



Important

Vérifier que l'horloge système est réglée à la date et l'heure courante sur tous les clients. Les certificats portant une date de validité, une différence de date entre les systèmes peut avoir pour effet de les invalider.

Les assistants de configuration associés à la PKI SafeKit gèrent la création de tous les certificats nécessaires à la mise en place de l'authentification par certificat client pour la console web et les commandes distribuées :

 CLCA	Le certificat de l'Autorité de Certification CLCA utilisée pour générer les certificats client
 ADMIN CONTROL MONITOR	Les certificats client permettant d'assurer l'identité de l'utilisateur de la console et de restreindre ses accès en fonction de son rôle
 ADMIN1 ADMIN2	Les certificats client de S1 et de S2 permettant de s'assurer de l'identité du nœud qui exécute une commande distribuée
 PROXY1 PROXY2	Les certificats client de S1 et de S2 permettant de s'assurer de l'identité du nœud qui accède à la console en mode proxy

Cette implémentation correspond à celle supportée depuis SafeKit 7.4. Elle a été modifiée depuis SafeKit 7.5 mais uniquement pour la configuration avec une PKI externe.

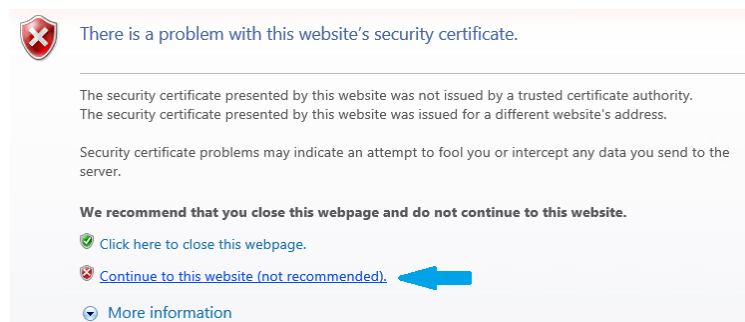
11.4.3.1 Configurer et redémarrer le service web

Pour activer l'authentification à base de certificat client (`SAFE=C:\safekit` en Windows si `%SYSTEMDRIVE%=C:` ; et `SAFE=/opt/safekit` en Linux) :

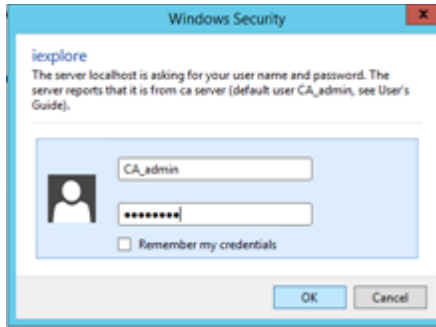
 httpd.conf	Sur S1 et S2 : ⇒ éditer le fichier <code>SAFE/web/conf/httpd.conf</code> ⇒ commenter <code>usefile</code> et <code>useldap</code> <pre># Define useldap ... # Define usefile</pre> ⇒ vérifier que seul <code>httpadminport</code> est défini <pre>httpadminport (9010) #httpcontrolport (9010) #httpmonitorport (9010)</pre>
	Sur S1 et S2 : ⇒ exécuter <code>safekit webserver restart</code>

11.4.3.2 Démarrer l'assistant de configuration de la console HTTPS

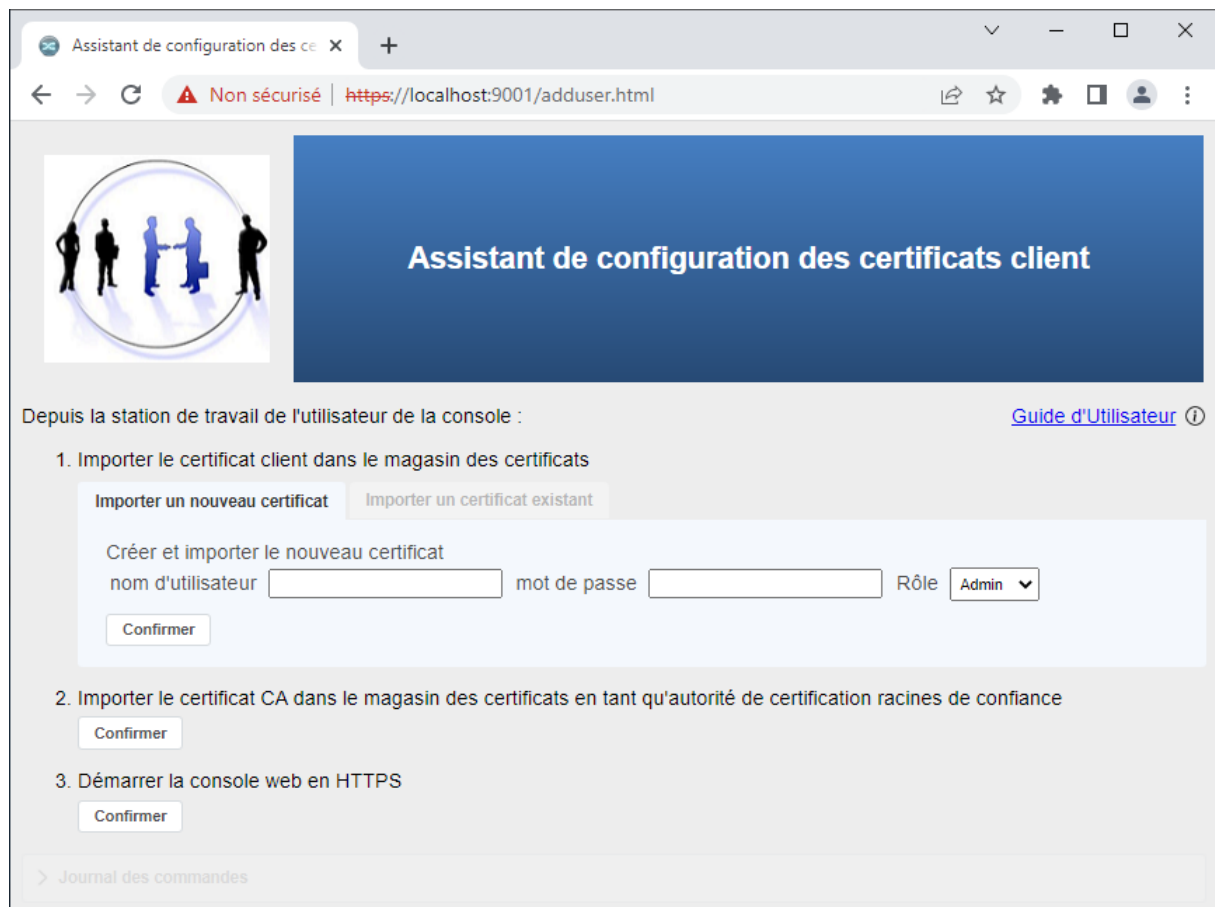
- ⇒ Se connecter sur le poste de travail de l'utilisateur qui aura accès à la console
- ⇒ Lancer un navigateur avec l'URL `https://firstserver:9001/adduser.html` où `firstserver` est l'adresse IP du premier serveur configuré pour HTTPS
- ⇒ Le certificat associé au service web CA est auto-signé. Par conséquent, le navigateur affiche une alerte de sécurité indiquant que le certificat est invalide. Vous devez cliquer sur `Continue to this website (not recommended)` pour poursuivre



- ⇒ À l'invite de connexion, entrez `CA_admin` le nom d'utilisateur et le mot de passe que vous avez spécifié lors du démarrage du service web CA (par exemple, `PasW0rD`)



L'assistant de configuration des certificats client est affiché :



11.4.3.3 Créer et/ou télécharger les certificats client

Créez un nouveau certificat client pour l'utilisateur s'il n'existe pas déjà.

⇒ Créer un nouveau certificat



- ⇒ Remplir les champs **nom d'utilisateur**, **mot de passe** et sélectionner le rôle dans le formulaire. Notez que le nom d'utilisateur doit être unique.
- ⇒ Cliquer sur **Confirmer**

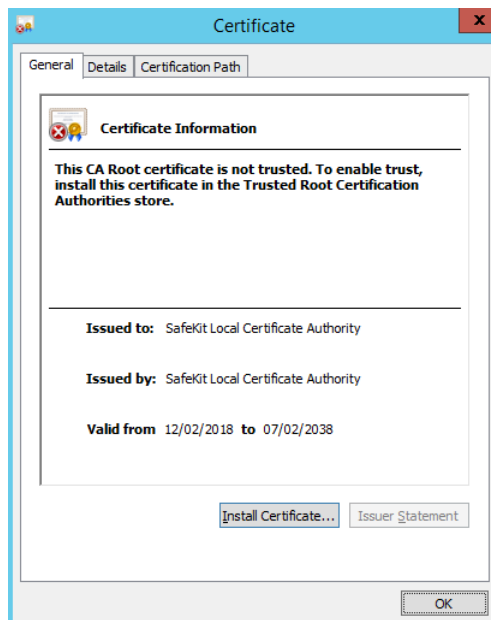
Une fois le traitement du formulaire terminé, le certificat client généré (le fichier `user_Admin_administrator.p12`) est téléchargé par le navigateur

Une fois le certificat client téléchargé (le nouveau ou l'existant), l'importer dans le magasin des certificats de la station de travail de l'utilisateur. L'accès à la console web SafeKit est interdit tant que le certificat client n'a pas été importé.

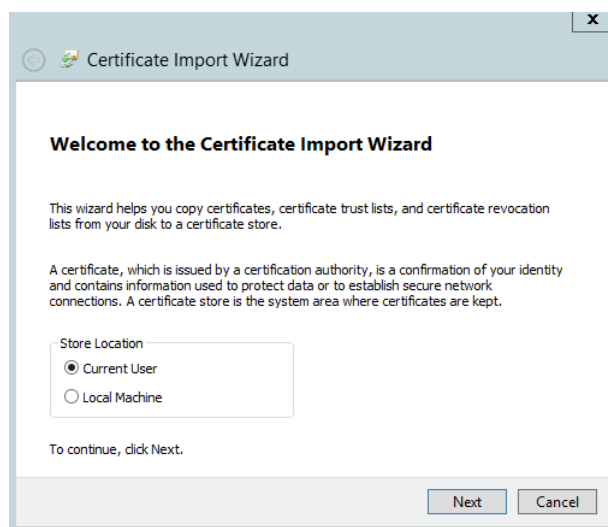
11.4.3.4 Importer le certificat client dans le magasin personnel

La procédure d'importation dépend du navigateur web et/ou du système d'exploitation. La procédure ci-dessous est décrite pour Windows.

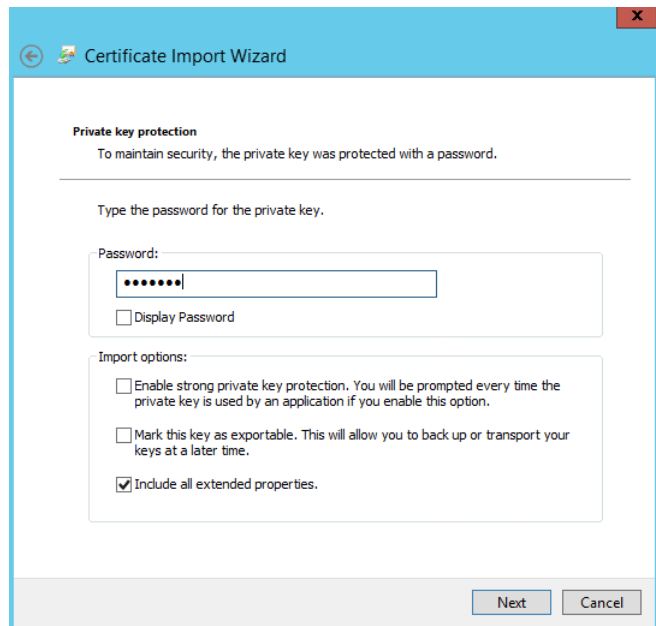
- ⇒ Cliquer sur le fichier `.p12` téléchargé (par exemple `user_Admin_administrator.p12`) pour ouvrir la fenêtre **Certificate**. Cliquer ensuite sur le bouton **Install Certificate**



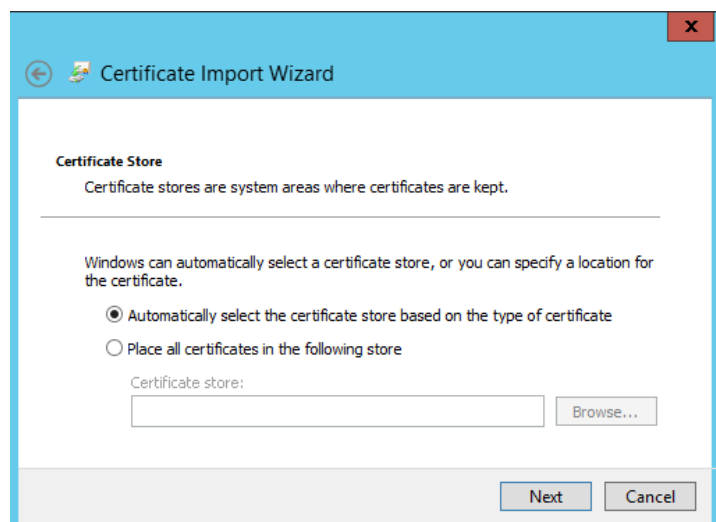
- ⇒ L'assistant d'importation de certificat s'ouvre. Sélectionner **Current User** puis cliquer sur le bouton **Next**
- ⇒ Poursuivre jusqu'à la demande de la saisie du mot de passe qui protège le certificat



- ⇒ Saisir le mot de passe. Il s'agit du mot de passe donné lors de la création du certificat décrite ci-dessus



- ⇒ Poursuivre en laissant l'assistant choisir automatiquement le magasin où importer le certificat. Il s'agit du magasin Personnel



- ⇒ Enfin terminer l'importation du certificat

11.4.3.5 Importer le certificat CA en tant qu'Autorité de certification racines de confiance

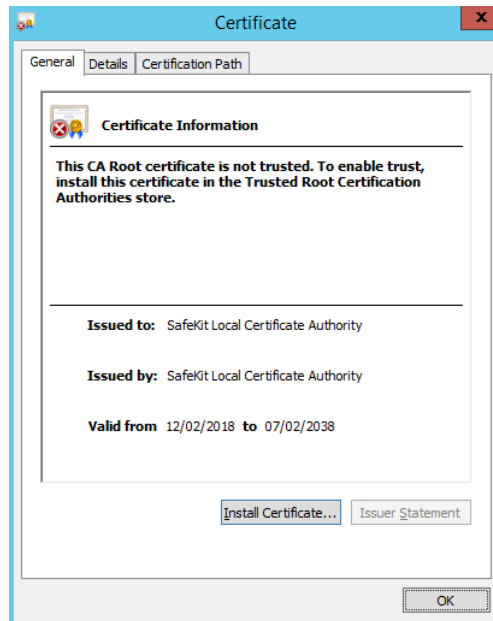
Le navigateur émet des alertes de sécurité lorsque vous vous connectez à la console web tant que ce certificat n'a pas été importé. La procédure d'importation dépend du navigateur web et/ou du système d'exploitation. La procédure ci-dessous est décrite pour Windows.

- ⇒ Cliquer sur **Confirmer** pour télécharger le certificat CA

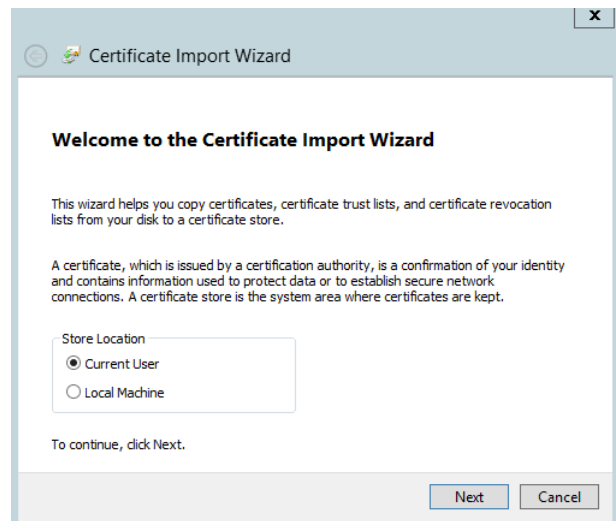
2. Importer le certificat CA dans le magasin des certificats en tant qu'autorité de certification racines de confiance

Confirmer

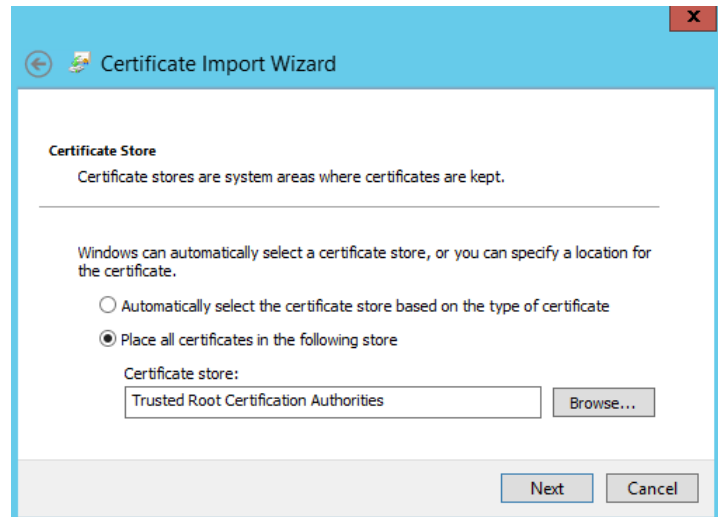
- ⇒ Cliquer sur le fichier `cacert.crt` téléchargé pour ouvrir la fenêtre Certificate. Cliquer ensuite sur le bouton Install Certificate



- ⇒ L'assistant d'importation de certificat s'ouvre. Sélectionner Current User puis cliquer sur le bouton Next



- ⇒ Parcourir les magasins pour sélectionner le magasin Trusted Root Certification Authorities. Puis cliquer sur le bouton Next



- ⇒ Enfin terminer l'importation du certificat

11.4.3.6 Tester la console web et la commande distribuée

La configuration est terminée ; vous pouvez maintenant vérifier qu'elle est opérationnelle :

- ⇒ Tester la console web

Une fois les certificats importés sur son poste de travail, l'utilisateur peut commencer à utiliser la console web.

1. Cliquer sur **Confirmer** pour démarrer la console

3. Démarrer la console web en HTTPS

Confirmer

Ou

1. Démarrer un navigateur web
2. Le connecter à l'URL `http://servername:9010` (où `servername` est l'adresse IP ou le nom d'un nœud SafeKit). Comme HTTPS est configuré, il y a une redirection automatique sur `https://servername:9453`.
3. En fonction des navigateurs et des serveurs, l'utilisateur peut avoir à sélectionner le certificat client à utiliser (le champ `friendly-name` du certificat est affiché)
4. La page chargée contient uniquement les onglets autorisés en fonction du rôle de l'utilisateur

- ⇒ Tester une commande distribuée

1. Se loguer sur S1 ou S2 en tant que administrateur/root
2. Ouvrir un terminal (PowerShell, shell, ...)
3. Aller dans le répertoire `SAFE`
4. Exécuter `safekit -H "*" level`
qui doit retourner le résultat de la commande `level` sur tous les nœuds

11.4.4 Configuration de l'authentification à base de certificat client avec une PKI externe

Une fois la configuration de HTTPS terminée, comme décrit dans 11.3.2 [page 192](#), vous pouvez poursuivre avec la configuration de l'authentification à base de certificats clients. Celle-ci repose sur la présence d'un ensemble de certificats énumérés ci-dessous.

 CLCA	Le certificat de l'Autorité de Certification CLCA utilisée pour générer les certificats client
 USER	Les certificats client permettant d'assurer l'identité de l'utilisateur de la console et de restreindre ses accès en fonction de son rôle ⇒ Certificat personnel déployé dans l'entreprise comme identité numérique de l'utilisateur Ou ⇒ Certificat dédié généré pour accéder à la console
 ADMIN1 ADMIN2	Les certificats client de S1 et de S2 permettant de s'assurer de l'identité du nœud qui exécute une commande distribuée ⇒ Certificat du serveur quand il peut être utilisé aussi en tant que certificat client Ou ⇒ Certificat dédié généré pour exécuter les commandes distribuées
 PROXY1 PROXY2	Les certificats client de S1 et de S2 permettant de s'assurer de l'identité du nœud qui accède à la console en mode proxy. Ils sont construits à partir de <code>admin1</code> et <code>admin2</code>.

Suivez les étapes suivantes pour mettre en œuvre l'authentification à base de certificats clients générés avec votre PKI. La configuration HTTPS, décrite en 11.3.2 [page 192](#), doit avoir été effectuée au préalable.

11.4.4.1 Récupérer et installer les certificats client pour la commande distribuée

11.4.4.1.1 Utiliser les certificats serveurs

Les certificats serveurs sont ceux qui ont été générés pour mettre en œuvre la configuration HTTPS en 11.3.2 [page 192](#).

Pour vérifier que ces certificats peuvent être utilisés comme certificats clients, lisez leur contenu. Voici la procédure à suivre sous Windows :

1. Copier le fichier `.crt` (ou `.cer`) sur une station de travail Windows
2. Double cliquer sur le fichier pour l'ouvrir avec *Extension noyau de chiffrement*

3. Cliquer sur l'onglet **Détails**

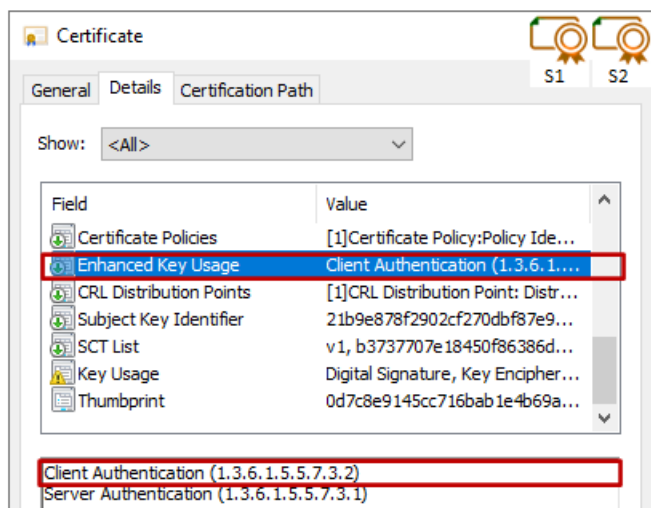
4. Vérifier le champ **Utilisation avancée de la clé (Enhanced Key Usage)**



Si vous préférez les lignes de commande, exécutez sur chaque nœud :

```
SAFE/web/bin/openssl.exe x509 -text -noout -in SAFE/web/conf/server.crt
```

Et recherchez la valeur **TLS Web Client Authentication** dans le champ **X509v3 Extended Key Usage**.



Notez la valeur du sous-champ **CN (Common Name)** incluse dans le champ **Objet (Subject)** pour configurer les rôles plus tard.

Quand ce champ contient la valeur **Authentification du client (Client Authentication)**, ces certificats peuvent être utilisés comme certificats client pour la commande distribuée :

ADMIN1 admin1.crt admin1.key	Le certificat client de S1 pour l'autoriser à exécuter des commandes distribuées. ✓ copier server1.crt dans admin1.crt ✓ copier server1.key dans admin1.key
ADMIN2 admin2.crt admin2.key	Le certificat client de S2 pour l'autoriser à exécuter des commandes distribuées. ⇒ copier server2.crt dans admin2.crt ⇒ copier server2.key dans admin2.key

11.4.4.1.2 Utiliser des certificats client dédiés

Lorsque les certificats serveur ne peuvent pas être utilisés, vous devez obtenir de nouveaux certificats de client de votre PKI avec le format attendu décrit ci-dessous :

 ADMIN1	Le certificat client de S1 pour l'autoriser à exécuter des commandes distribuées.
 ADMIN2	Le certificat client de S2 pour l'autoriser à exécuter des commandes distribuées.
admin1.crt admin2.crt	⇒ fichier pour le certificat X509 dans le format PEM Le champ Utilisation avancée de la clé (Enhanced Key Usage) contient la valeur Authentification du client (Client Authentication). Le sous-champ CN (Common Name) incluse dans le champ Objet (Subject) contient le nom du serveur.  Notez la valeur de CN pour configurer les rôles plus tard. Important ⇒ la clé privée, *non cryptée*, associée aux certificats admin1.crt/admin2.crt

11.4.4.1.3 Installer les fichiers dans SafeKit

Installer les fichiers correspondants aux certificats comme suit (`SAFE=C:\safekit` en Windows si `System Drive=C:` ; et `SAFE=/opt/safekit` en Linux) :

 ADMIN1 admin1.crt admin1.key	Sur S1 : ➔ copier admin1.crt dans SAFE/web/conf/admin.crt ➔ copier admin1.key dans SAFE/web/conf/admin.key
 ADMIN2 admin2.crt admin2.key	Sur S2 : ➔ copier admin2.crt dans SAFE/web/conf/admin.crt ➔ copier admin2.key dans SAFE/web/conf/admin.key

 PROXY proxy.crtkey	Sur S1 et S2 : Construire le fichier <code>SAFE/web/conf/proxy.crtkey</code> comme suit : ⇒ convertir <code>admin.key</code> au format <code>rsa</code> en exécutant <pre>SAFE/web/bin/openssl rsa -in SAFE/web/conf/admin.key -out SAFE/web/conf/rsa-admin.key</pre> ⇒ concaténer les fichiers <code>admin.crt</code> et <code>rsa-admin.key</code> dans <code>proxy.crtkey</code> à l'aide d'un éditeur ou d'une commande
--	---

Vous pouvez contrôler les certificats installés avec :

```
cd SAFE/web/bin
checkcert -t client
```

Cette commande retourne en échec si une erreur est détectée.

11.4.4.2 Récupérer et importer les certificats client pour la console web

11.4.4.2.1 Utiliser les certificats personnels

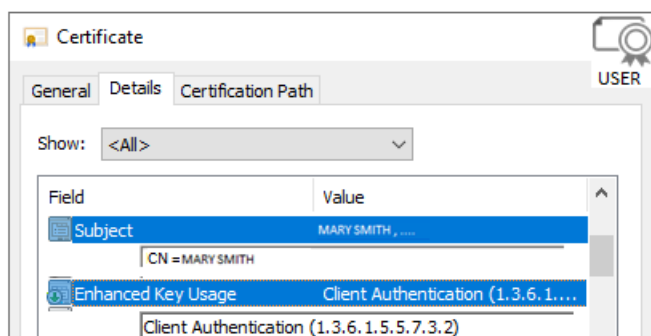
Dans certaines entreprises, chaque utilisateur dispose d'un certificat personnel comme identifiant numérique, qui est présent dans le magasin de certificats sur le poste de travail de l'utilisateur.

 USER Certificat personnel	Le certificat personnel est utilisé pour assurer l'identité de l'utilisateur dans la console.
--	--

Pour vérifier que ce certificat peut être utilisé comme certificat client pour la console web, lisez son contenu. Voici la procédure à suivre sous Windows :

1. Connectez-vous sur la station de travail de l'utilisateur
2. Ouvrez une console PowerShell
3. Exécutez `certmgr`
4. Localisez le certificat dans le magasin `Certificats - Utilisateur actuel\Personnel\Certificats`, puis faites un clic droit dessus. Si l'utilisateur possède plusieurs certificats, sélectionner celui ayant « **Authentification du client** » comme « **Rôles prévus** » et dont la « **Date d'expiration** » n'est pas passée
5. Cliquer sur le certificat avec le bouton droit puis « **Ouvrir** ». Cela ouvre la fenêtre de certificat
6. Vérifiez le contenu du champ `Utilisation avancée de la clé (Enhanced Key Usage)`

Ci-dessous l'exemple du certificat personnel de Mary Smith :



Notez la valeur du sous-champ CN (Common Name) incluse dans le champ Objet (Subject) pour configurer les rôles plus tard.

Comme le champ Key Usage contient la valeur Client Authentication, ce certificat personnel peut être utilisé comme certificat client pour la console. Dans ce cas, il n'est pas nécessaire de l'importer puisqu'il est déjà présent dans le magasin de certificats du poste de travail de l'utilisateur.

11.4.4.2.2 Utiliser des certificats client dédiés

Lorsque les certificats personnels ne peuvent pas être utilisés, vous devez obtenir un nouveau certificat client de votre PKI avec le format attendu décrit ci-dessous :



USER

user.p12

Le certificat client pour authentifier l'utilisateur de la console web

⇒ fichier pour le certificat X509 dans le format PKCS#12

Le champ Utilisation avancée de la clé (Enhanced Key Usage) contient la valeur Authentification du client (Client Authentication). Le sous-champ CN (Common Name) incluse dans le champ Objet (Subject) contient le nom de l'utilisateur.



Important

Notez la valeur de CN pour configurer les rôles plus tard.

Pour chaque utilisateur de la console web, vous devez ensuite importer son certificat dans son magasin de certificats comme suit :

1. Connectez-vous sur la station de travail de l'utilisateur
2. Cliquez sur le fichier `user.p12` pour ouvrir la fenêtre **Certificate**
3. Cliquez sur le bouton **Install Certificate**
L'assistant d'importation de certificat s'ouvre
4. Sélectionnez **Current User** puis cliquez sur le bouton **Next**
5. Poursuivez en laissant l'assistant choisir automatiquement le magasin où importer le certificat
Il doit s'agir du magasin **Personnel**.
6. Terminez l'importation du certificat

11.4.4.3 Récupérer, installer et importer le certificat CLCA

11.4.4.3.1 Récupérer le fichier du certificat

Vous devez récupérer ce certificat de votre PKI avec le format attendu :

 CLCA clcacert.crt	Le certificat de l'Autorité de Certification CLCA utilisée pour générer les certificats client. ⇒ fichier pour le certificat X509 dans le format PEM La chaîne de certificats pour la CA racine et les intermédiaires, s'il y en a Si différents CLCAs sont utilisés pour générer les certificats client, le fichier <code>clcacert.crt</code> doit contenir la concaténation de ces différents CLCAs.	 ADMIN1 ADMIN2 Certificats client pour la commande distribuée  USER Certificats client pour les utilisateurs de la console web
--	--	--

Si vous rencontrez des difficultés pour récupérer ce fichier depuis votre PKI, vous pouvez le construire à l'aide de la procédure décrite en 7.18. [page 131](#).



Si votre PKI utilise la même Autorité de Certification pour émettre des certificats serveur et client, les fichiers `cacert.crt` et `clcacert.crt` sont identiques. Le fichier `cacert.crt` a été récupéré et installé durant la procédure de configuration HTTPS (voir 11.3.2.2 [page 195](#)).

11.4.4.3.2 Installer le fichier dans SafeKit

Installer le certificat comme suit (`SAFE=C:\safekit` en Windows si System Drive=C: ; et `SAFE=/opt/safekit` en Linux) :

 CLCA clcacert.crt	Sur S1 et S2: ⇒ copier <code>clcacert.crt</code> dans <code>SAFE/web/conf/clcacert.crt</code>
--	---

Vous pouvez contrôler les certificats installés avec :

```
cd SAFE/web/bin
checkcert -t CLCA
```

Cette commande retourne en échec si une erreur est détectée.

Vous devez également vérifier que le fichier `clcacert.crt` contient bien la chaîne de certificats pour les autorités de certification racine et intermédiaires.

11.4.4.3 Importer le certificat dans le magasin des certificats de l'utilisateur

Tant que le certificat de l'autorité de certification n'a pas été importé, le navigateur émet des alertes de sécurité lorsque l'utilisateur se connecte à la console web avec son certificat client. Si l'importation n'a pas déjà été faite, appliquez la procédure ci-dessous en Windows :

1. Connectez-vous sur la station de travail de l'utilisateur
2. Cliquez sur le fichier `clacert.crt` pour ouvrir la fenêtre **Certificate**
3. Cliquez sur le bouton **Install Certificate**
L'assistant d'importation de certificat s'ouvre
4. Sélectionnez **Current User** puis cliquez sur le bouton **Next**
5. Poursuivez l'assistant et sélectionnez le magasin **Trusted Root Certification Authorities**
6. Terminez l'importation du certificat

11.4.4.4 Configurer les rôles

Le certificat client est utilisé pour authentifier l'utilisateur de la console ou de la commande distribuée. Un rôle doit lui être attribué pour définir les actions autorisées.

Ceci doit être défini dans le fichier `sslgroup.conf` qui peut contenir les 3 groupes Admin, Control, Monitor. Les utilisateurs de ces groupes auront les rôles correspondants.



Chaque ligne débute par le nom du groupe, suivi de :, suivi de la liste des utilisateurs (dont le séparateur est l'espace). Voir l'exemple ci-dessous.

 sslgroup.conf	⇒ éditer le fichier <code>sslgroup.conf</code> ⇒ assigner un rôle à chacun des certificats client 1. récupérer la valeur du sous-champ CN (Common Name) incluse dans le champ Objet (Subject) du certificat 2. insérer CN avec le rôle souhaité ✓ pour les certificats de la console, cela peut-être n'importe lequel des rôles Admin, Control ou Monitor ✓ pour les certificats de la commande distribuée, le rôle doit être Admin
	Sur S1 et S2: ⇒ copier <code>sslgroup.conf</code> dans <code>SAFE/web/conf/sslgroup.conf</code>

Dans l'exemple suivant, `s1.w.com` et `s2.w.com` sont la valeur CN des certificats de commande distribuée ; les autres noms sont la valeur CN des certificats de console :

 Admin Control Monitor sslgroun.conf	<pre>Admin : "s1.w.com" "s2.w.com" "MARY SMITH" admin Control: "NAD ROU" "DAVID JOHNS" manager Monitor : monitor</pre>
--	--

11.4.4.5 Configurer et redémarrer le service web

Pour activer l'authentification à base de certificat client (SAFE=C:\safekit en Windows si %SYSTEMDRIVE%=C: ; et SAFE=/opt/safekit en Linux) :

 httpd.conf	Sur S1 et S2 : ⇒ éditer le fichier SAFE/web/conf/httpd.conf ⇒ commenter usefile et useldap <pre># Define useldap ... # Define usefile</pre> ⇒ vérifier que seul httpadminport est défini <pre>httpadminport (9010) #httpcontrolport (9010) #httpmonitorport (9010)</pre>
	Sur S1 et S2 : ⇒ exécuter <code>safekit webserver restart</code>

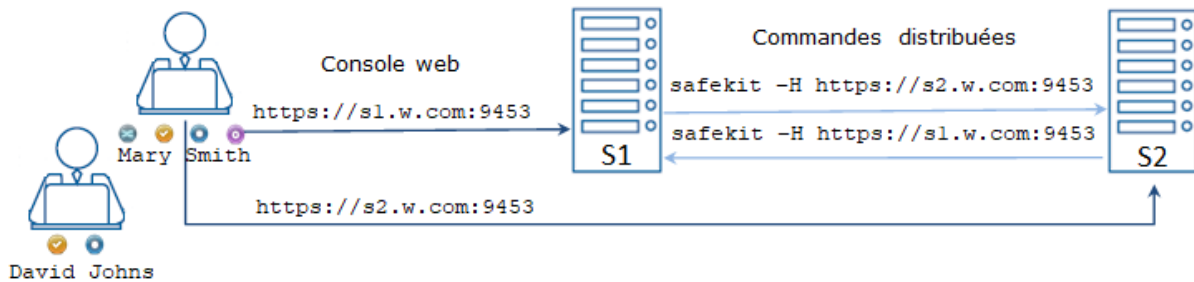
11.4.4.6 Tester la console web et la commande distribuée

La configuration est terminée ; vous pouvez maintenant vérifier qu'elle est opérationnelle :

- ⇒ Tester la console web
 1. Démarrer un navigateur web
 2. Le connecter à l'URL `http://servername:9010` (où `servername` est l'adresse IP ou le nom d'un nœud SafeKit). Comme HTTPS est configuré, il y a une redirection automatique sur `https://servername:9453`.
 3. En fonction des navigateurs et des serveurs, l'utilisateur peut avoir à sélectionner le certificat client à utiliser (le champ `friendly-name` du certificat est affiché)
 4. La page chargée contient uniquement les onglets autorisés en fonction du rôle de l'utilisateur
- ⇒ Tester une commande distribuée
 1. Se loguer sur S1 ou S2 en tant que administrateur/root
 2. Ouvrir un terminal (PowerShell, shell, ...)
 3. Aller dans le répertoire SAFE
 4. Exécuter `safekit -H "*" level`
qui doit retourner le résultat de la commande `level` sur tous les nœuds

11.5 Exemple de configuration de HTTPS et de l'authentification par certificat personnel

Cette section est une synthèse de la configuration avec un PKI externe. Elle montre la configuration de HTTPS et de l'authentification basée sur les certificats personnels des utilisateurs, pour l'exemple suivant :



Cette configuration simplifiée n'est disponible que sous certaines conditions décrites ci-dessous. Pour tous les autres cas, veuillez vous référer à 11.3.2 [page 192](#) pour la configuration HTTPS et à 11.4.4 [page 209](#) pour la configuration de l'authentification utilisateur basée sur des certificats client.

11.5.1 Vérifier les prérequis

⇒ Configuration du cluster SafeKit

Configurer le cluster, dans `SAFEVAR/cluster/cluster.xml`, comme suit :

```
<?xml version="1.0" encoding="UTF-8"?>
<cluster>
<lans>
  <lan name="default" console="on" framework="on">
    <node name="s1" addr="s1.w.com"/>
    <node name="s2" addr="s2.w.com"/>
  </lan>
</lans>
</cluster>
```

⇒ Certificats serveur

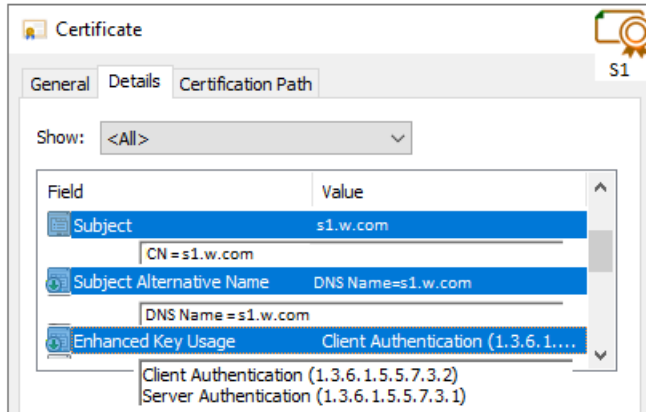
Demander à votre PKI, pour chaque nœud SafeKit, les fichiers pour le certificat et la clé associée.

Consulter le contenu du certificat :

- ✓ vérifier que le champ Utilisation avancée de la clé (Enhanced Key Usage) contient bien Authentification du client (Client Authentication)
- ✓ vérifier que la valeur de `addr` dans `cluster.xml` est bien présente dans le champ Autre nom de l'objet (Subject Alternative Name)
- ✓ Avec SafeKit `<= 7.5.2.9`, vérifier que le nom du serveur est bien présent dans le champ Autre nom de l'objet (Subject Alternative Name)
- ✓ Noter la valeur du CN dans le champ Objet (Subject), qui sera utilisée plus loin pour remplir le fichier `the sslgroup.conf`

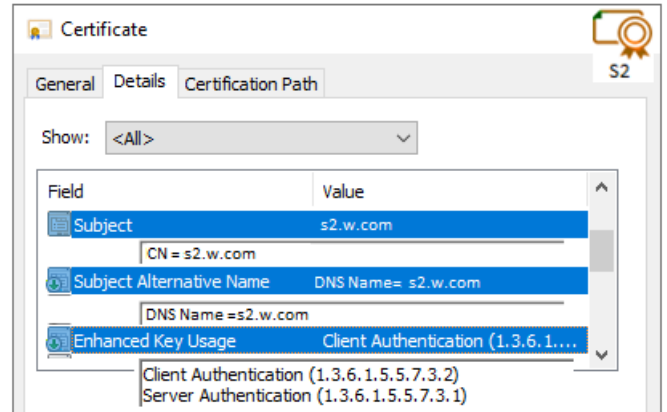
Certificat serveur de S1

Fichiers `server1.crt` et `server1.key`



Certificat serveur de S2

Fichiers `server2.crt` et `server2.key`



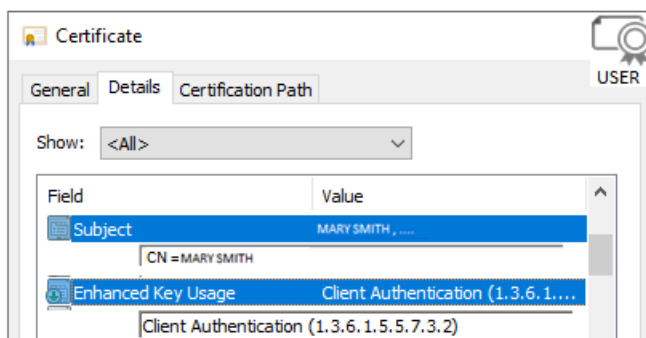
⇒ Certificats personnels des utilisateurs

Ils devraient déjà être présents dans le magasin des certificats de la station de travail de chaque utilisateur (`certmgr.msc`) sous « Certificats - Utilisateur actuel\Personnel\Certificats »

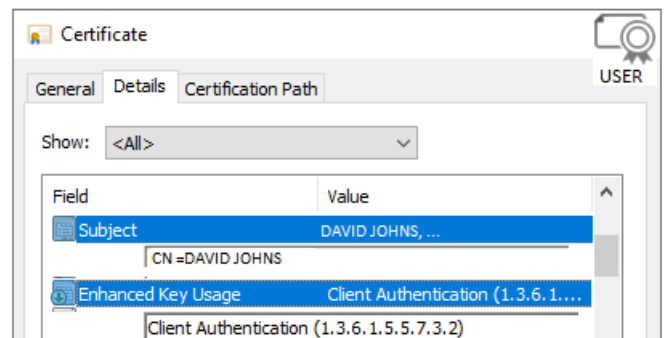
Consulter le contenu du certificat :

- ✓ Vérifier que le champ Utilisation avancée de la clé (Enhanced Key Usage) contient bien Authentification du client (Client Authentication)
- ✓ Noter la valeur du CN dans le champ Objet (Subject), qui sera utilisée plus loin pour remplir le fichier `the sslgroup.conf`
- ✓

Certificat personnel de Mary Smith



Certificat personnel de David Johns



11.5.2 Configuration de HTTPS et de l'authentification à base de certificat personnel

Appliquer les étapes suivantes pour configurer HTTPS et mettre en œuvre l'authentification à base de certificat personnel.

11.5.2.1 Récupérer et installer les certificats dans SafeKit

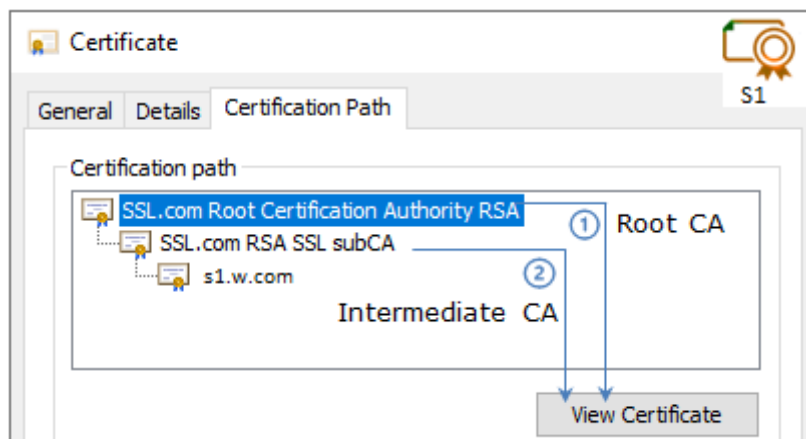
 S1 server1.crt server1.key	Sur S1 : ⇒ Certificat serveur ✓ copier server1.crt dans SAFE/web/conf/server.crt ✓ copier server1.key dans SAFE/web/conf/server.key ⇒ Certificat client pour les commandes distribuées ✓ copier server1.crt dans SAFE/web/conf/admin.crt ✓ copier server1.key dans SAFE/web/conf/admin.key
 S2 server2.crt server2.key	Sur S2 : ⇒ Certificat serveur ✓ copier server2.crt dans SAFE/web/conf/server.crt ✓ copier server2.key dans SAFE/web/conf/server.key ⇒ Certificat client pour les commandes distribuées ✓ copier server2.crt dans SAFE/web/conf/admin.crt ✓ copier server2.key dans SAFE/web/conf/admin.key
 PROXY proxy.crtkey	Sur S1 et S2 : ⇒ Certificat client pour le mode proxy de la console Construire le fichier SAFE/web/conf/proxy.crtkey comme suit : ✓ convertir admin.key au format rsa en exécutant <pre>SAFE/web/bin/openssl rsa -in SAFE/web/conf/admin.key -out SAFE/web/conf/rsa-admin.key</pre> ✓ concaténer les fichiers admin.crt et rsa-admin.key dans proxy.crtkey à l'aide d'un éditeur ou d'une commande



CA
cacert.crt

Récupérer le certificat de l'autorité de certification utilisé pour émettre les certificats du serveur (la chaîne de certificats du CA racine et des intermédiaires, s'il y en a).

Si vous ne l'avez pas, vous pouvez le construire comme indiqué ci-dessous :



- ⇒ « Afficher le certificat » racine et intermédiaires pour les exporter dans un fichier au format « Codé à base 64 X.509 (.cer). ».
- ⇒ Concaténer 1, 2 dans cacert.crt

Sur S1 et S2 :

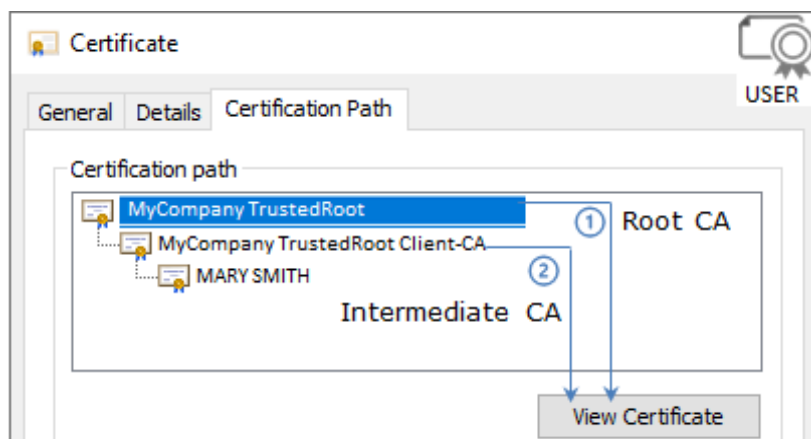
- ⇒ copier cacert.crt dans SAFE/web/conf/cacert.crt



CLCA
clcacert.crt

Récupérer le certificat de l'autorité de certification utilisé pour émettre les certificats personnels (la chaîne de certificats du CA racine et des intermédiaires, s'il y en a).

Si vous ne l'avez pas, vous pouvez le construire comme indiqué ci-dessous :



- ⇒ « Afficher le certificat » racine et intermédiaires pour les exporter dans un fichier au format « Codé à base 64 X.509 (.cer). »
- ⇒ concaténer 1, 2 dans personalcacert.crt

	Sur S1 et S2 : ⇒ concaténer <code>cacert.crt</code> et <code>personalcacert.crt</code> dans <code>SAFE/web/conf/clcacert.crt</code>
--	--

11.5.2.2 Configurer les rôles

 sslgroup.conf	Sur S1 et S2 : ⇒ éditer le fichier <code>SAFE/web/conf/sslgroup.conf</code> ⇒ insérer les lignes suivantes <pre>Admin:"s1.w.com" "s2.w.com" "MARY SMITH" Control:"DAVID JOHNS"</pre>
--	--

11.5.2.3 Configurer et redémarrer le serveur web

 httpd.conf	Sur S1 et S2 : ⇒ éditer le fichier <code>SAFE/web/conf/httpd.conf</code> ⇒ commenter <code>usefile</code> et <code>uselldap</code> <pre># Define uselldap ... # Define usefile</pre> ⇒ vérifier que seul <code>httpadminport</code> est défini <pre>httpadminport (9010) #httpcontrolport (9010) #httpmonitorport (9010)</pre>
 httpd.webconsoleless.conf	Sur S1 et S2 : ⇒ copier <code>SAFE/web/conf/httpd.webconsoleless.conf</code> dans <code>SAFE/web/conf/ssl/httpd.webconsoleless.conf</code>
	Sur S1 et S2 : ⇒ exécuter <code>safekit webserver restart</code>

11.5.2.4 Changer les règles du pare-feu

Pare-feu	Sur S1 et S2 : ⇒ aller dans le répertoire <code>SAFEBIN</code> ⇒ exécuter <code>firewallcfg add</code>
----------	--

11.5.3 Tester la console web et la commande distribuée

La configuration est terminée ; vous pouvez maintenant vérifier qu'elle est opérationnelle :

⇒ Tester la console web

1. Démarrer un navigateur web
2. Le connecter à l'URL <http://s1.w.com:9010/> ou <http://s2.w.com:9010/>. Comme HTTPS est configuré, il y a une redirection automatique sur <https://servername:9453>.
3. En fonction des navigateurs et des serveurs, l'utilisateur peut avoir à sélectionner le certificat client à utiliser (le champ `friendly-name` du certificat est affiché).
4. La page chargée contient uniquement les onglets autorisés en fonction du rôle de l'utilisateur
 - ✓ Mary Smith, avec le rôle Admin, accède aux onglets  Configuration,  Contrôle,  Supervision et  Configuration Avancée
 - ✓ David Johns, avec le rôle Control, accède uniquement aux onglets  Contrôle et  Supervision

⇒ Tester une commande distribuée

1. Se loguer sur S1 ou S2 en tant que administrateur/root
2. Ouvrir un terminal (PowerShell, shell, ...)
3. Aller dans le répertoire `SAFE`
4. Exécuter `safekit -H "*" level`
qui doit retourner le résultat de la commande `level` sur tous les nœuds

11.6 Configuration avancée avec la PKI SafeKit

11.6.1 Configuration rapide de HTTPS en ligne de commandes

Tout d'abord, choisissez parmi les nœuds composant le cluster SafeKit, un nœud pour agir en tant que serveur d'autorité de certification. Le nœud sélectionné sera appelé ci-après le `serveur CA`. Les autres nœuds de cluster sont appelés `serveur non-CA`. Appliquez ensuite séquentiellement chacune des sous-sections suivantes pour activer la configuration HTTPS préconfigurée pour sécuriser la console web SafeKit.



Vérifier que l'horloge système est réglée à la date et l'heure courante sur tous les clients et les serveurs. Les certificats portant une date de validité, une différence de date entre les systèmes peut avoir pour effet de les invalider.

11.6.1.1 Démarrer le service web CA sur le serveur CA

Appliquer la même procédure que celle décrite en 11.3.1.1 [page 188](#), pour démarrer le service web CA (service `safecaserv`).

11.6.1.2 Générer les certificats sur le serveur CA

Pendant cette étape, l'environnement de génération des certificats est mis en place ; les certificats de l'autorité de certification et du serveur CA sont générés et installés à l'emplacement attendu par la configuration HTTPS ; les trois certificats clients sont également créés.



Vérifier que l'horloge système est réglée à la date et l'heure courante sur le serveur.

Par défaut, le certificat serveur inclut toutes les adresses IP et les noms DNS définis localement. Ils sont répertoriés dans les fichiers `SAFE/web/conf/ipv4.json`, `SAFE/web/conf/ipv6.json` et `SAFE/web/conf/ipnames.json`. Ces fichiers sont générés par la commande qui démarre le service web CA, appelée à l'étape précédente.



Si vous souhaitez accéder à la console web depuis un nom DNS ou une adresse IP non répertorié, modifiez le fichier correspondant pour insérer la nouvelle valeur avant d'exécuter la commande `initssl`. Cela est nécessaire par exemple pour accéder depuis internet à un cluster SafeKit dans le cloud, quand les serveurs ont une adresse publique mappée sur une adresse privée.

Sur le serveur CA :

- ⇒ Se connecter en tant qu'administrateur/root et ouvrir une fenêtre d'invite de commandes
- ⇒ Aller dans le répertoire `SAFE/web/bin`
- ⇒ Exécuter la commande :

```
./initssl ca
```

Cette commande crée un certificat CA avec un « subject name » par défaut. Pour spécifier sa valeur, exécuter plutôt la commande étendue :

```
./initssl ca "/O=My Company/OU=My Entity/CN=My Company Private  
Certificate Authority"
```

A l'invite, entrer le mot de passe qui protégera chacun des trois certificats clients :

```
Enter the password for the Admin role pkcs12 file
(../conf/ca/private/user_Admin_administrator.p12) twice:
```

```
Enter Export Password: pwd1
```

```
Verifying - Enter Export Password: pwd1
```

```
Enter the password for the Control role pkcs12 file
(../conf/ca/private/user_Control_manager.p12) twice:
```

```
Enter Export Password: pwd2
```

```
Verifying - Enter Export Password: pwd2
```

```
Enter the password for the Monitor role pkcs12 file
(../conf/ca/private/user_Monitor_operator.p12) twice:
```

```
Enter Export Password: pwd3
```

```
Verifying - Enter Export Password: pwd3
```



Le mot de passe saisi sera demandé au moment de l'importation du certificat client sur la station cliente.

- ⇒ Copier les certificats clients devant être exportés (fichiers .p12) dans le répertoire `../conf/ca/certs`

11.6.1.3 Générer les certificats sur le serveur non-CA

Pendant cette étape, le certificat du serveur local est généré, les certificats signés sont téléchargés depuis le serveur CA, et enfin les certificats sont installés à l'emplacement prévu par la configuration HTTPS.

Appliquer en séquence la procédure suivante sur chaque serveur non-CA :

- ⇒ Se connecter en tant qu'administrateur/root et ouvrir une fenêtre d'invite de commandes
- ⇒ Aller dans le répertoire `SAFE/web/bin`
- ⇒ Répertorier les noms DNS et adresses IP du serveur

Par défaut, le certificat serveur inclut toutes les adresses IP et les noms DNS définis localement. Ils sont répertoriés dans les fichiers `SAFE/web/conf/ipv4.json`, `SAFE/web/conf/ipv6.json` et `SAFE/web/conf/ipnames.json`. Pour générer ces fichiers, exécutez la commande :

- ✓ en Linux

```
./getipandnames
```

Cette commande utilise sur la commande `host` délivrée avec le package `bind-utils`. Installez-le si nécessaire ou remplissez manuellement les noms DNS dans le fichier `SAFE/web/conf/ipnames.json`.

- ✓ en Windows

```
./getipandnames.ps1
```



Important

Si vous souhaitez accéder à la console web depuis un nom DNS ou une adresse IP non répertorié, modifiez le fichier correspondant pour insérer la nouvelle valeur avant d'exécuter la commande `initssl`. Cela est nécessaire par exemple pour accéder depuis internet à un cluster SafeKit dans le cloud, quand les serveurs ont une adresse publique mappée sur une adresse privée.

⇒ Exécuter la commande :

```
./initssl req https://CAserverIP:9001 CA_admin
```

A chaque fois que cela est demandé, entrer le mot de passe qui a été utilisé au moment du démarrage du service web CA du serveur CA (`PasWOrD`). Pour ne pas avoir à saisir le mot de passe, exécuter plutôt la commande :

```
./initssl req https://CAserverIP:9001 CA_admin:PasW0rD
```

`CAserverIP` est l'adresse IP ou le nom DNS du serveur CA.



Note

Si la commande retourne l'erreur "Certificate is not yet valid", cela signifie que les horloges des deux serveurs ne sont pas synchronisées. Pour corriger le problème, il faut changer la date et l'heure des serveurs puis réexécuter la commande `initssl`.

11.6.1.4 Reconfigurer le service web en HTTPS sur les serveurs CA et non-CA

Une fois les certificats générés sur le serveur CA et sur chaque serveur non-CA, le service web SafeKit (service `safewebserver`) peut être configuré pour HTTPS. Appliquez la procédure décrite en 10.6.4 [page 174](#), sur **tous** les serveurs.

11.6.1.5 Configurer le pare-feu sur les serveurs CA et non-CA

Une fois le service Web SafeKit configuré en HTTPS, les communications réseau peuvent être ouvertes en configurant le pare-feu comme décrit en 10.3 [page 160](#).

11.6.1.6 Utiliser la console web SafeKit en HTTPS

5. Télécharger depuis le serveur CA, les certificats clients (`.p12` files) et le certificat du serveur CA (`cacert.crt` file), localisés dans `SAFE/web/conf/ca/certs`
6. Importer le certificat client comme décrit en 11.4.3.4 [page 205](#)
7. Importer le certificat CA comme décrit en 11.4.3.5 [page 206](#)

11.6.1.7 Arrêter le service web CA sur le serveur CA

Une fois tous les serveurs et clients configurés, il est recommandé d'arrêter le service web CA (service `safecaserv`) sur tous les serveurs. Cela limite le risque d'accès accidentel ou malveillant à l'assistant de configuration HTTPS. Pour arrêter le service web CA en ligne de commande :

- ⇒ Se connecter en tant qu'administrateur/root et ouvrir une fenêtre d'invite de commandes
- ⇒ Aller dans le répertoire `SAFE/web/bin`
- ⇒ Exécuter la commande `./stopcaserv`



Note

En Windows, cette commande supprime également l'entrée du service `safecaserv` pour empêcher son démarrage accidentel par la suite.

En Linux, le port 9001 est automatiquement fermé sur le pare-feu local.

11.6.1.8 Arrêter le service web CA sur le serveur CA

Une fois tous les serveurs et clients configurés, il est recommandé d'arrêter le service web CA (service `safecaserv`) sur le serveur CA. Cela limite le risque d'accès accidentel ou malveillant à l'assistant de configuration HTTPS.

- ⇒ Se connecter en tant qu'administrateur/root et ouvrir une fenêtre d'invite de commandes
- ⇒ Aller dans le répertoire `SAFE/web/bin`
- ⇒ Exécuter la commande `./stopcaserv`



En Windows, cette commande supprime également l'entrée du service `safecaserv` pour empêcher son démarrage accidentel par la suite.

En Linux, le port 9001 est automatiquement fermé sur le pare-feu local.

Cette étape n'est pas obligatoire, mais en production il est préférable de ne pas laisser accessible les fichiers sensibles.

Les fichiers présents sur le serveur CA, dans le répertoire `SAFE/web/conf/ca` (en particulier les clés privées sous `SAFE/web/conf/ca/private/*.keys`) doivent être sauvegardés dans un espace de stockage sûr et détruits sur le serveur. Ces fichiers devront être restaurés à leur emplacement initial si le serveur CA est à nouveau nécessaire (par exemple pour sécuriser un nouveau serveur SafeKit).

Pour les serveurs non-CA, il faut sauvegarder et détruire les fichiers présents sous le répertoire `SAFE/web/conf/ca`.

11.6.2 Renouvellement des certificats

Chaque certificat possède une date d'expiration. Par défaut, la date d'expiration du certificat CA est fixée à 10 ans après la date d'installation. Par défaut, la date d'expiration des certificats serveur et client est fixée à 5 ans après la date de demande de certificat.

Lorsque le certificat client est expiré, la console web ne peut se connecter aux serveurs SafeKit. Si le certificat serveur est expiré, la console web émet une alerte lors de la connexion au serveur. Une fois le certificat CA expiré, il sera impossible pour le service web SafeKit de valider les certificats présentés.

Il est possible de renouveler les certificats ou de régénérer une requête de création de certificat à partir des clés privées utilisées précédemment. Cette procédure n'est pas explicitée dans ce document. Il est proposé à la place de créer de nouveau jeu de certificats, pour remplacer les anciens :

- ⇒ supprimer le répertoire `web/conf/ca` sur tous les serveurs, y compris le serveur CA
- ⇒ supprimer les certificats des magasins des stations de travail des utilisateurs
- ⇒ réappliquer complètement les procédures décrites en 11.3.1 [page 187](#) et 11.4.3 [page 202](#)

11.6.3 Révocation des certificats

Il est possible de modifier la configuration des serveurs web SafeKit pour utiliser une liste de révocation de certificats (CRL). Cette procédure n'est pas explicitée dans ce document. Reportez-vous à la documentation apache et openssl.

Vous pouvez sinon créer un nouveau jeu de certificats, y compris celui de l'autorité de certification, pour remplacer le précédent. Cela a pour effet de révoquer les anciens certificats, car le certificat de l'autorité de certification a changé.

11.6.4 Commandes pour la génération des certificats

Les commandes sont localisées et doivent être exécutées depuis le répertoire SAFE/web/bin.

	Paramètres <Subject>: le sujet du certificat du CA, qui identifie le propriétaire du CA. Exemple <pre>initssl ca "/O=My Company/OU=My Unit/CN=My Company Private Certificate Authority"</pre> Description Tous les chemins d'accès ci-dessous sont relatifs au répertoire SAFE/web. Cette commande crée le répertoire conf/ca utilisé par les commandes openssl. Les certificats générés sont stockés sous conf/ca/certs ; les clés privées générées sont stockées sous conf/ca/private.  Note Habituellement, il est préférable de protéger les clés privées par un mot de passe. Cela impliquant une configuration plus complexe, ce n'est pas mis en place. Si besoin, voir la documentation d'Apache et d'OpenSSL. ⇒ Création du certificat CA <code>conf/ca/certs/cacert.crt</code> et de la clé associée <code>conf/ca/private/cacert.key</code> ⇒ Création du certificat du serveur <code>conf/ca/certs/server_localca.crt</code> et de la clé associée <code>conf/ca/private/server_localca.key</code> ⇒ Création du certificat client pour les commandes distribuées sur le cluster <code>conf/ca/certs/user_Admin_system.crt</code> et de la clé associée <code>conf/ca/private/user_Admin_system.key</code> ⇒ Création des trois certificats clients (correspondant aux trois rôles Admin, Control, Monitor) et des fichiers pkcs12 associés. Dans cette phase, le script demande d'entrer deux fois le mot de passe qui sera demandé à l'importation du fichier pkcs12. Les fichiers générés sont : ✓ <code>conf/ca/private/user_Admin_administrator.p12</code> ✓ <code>conf/ca/private/user_Control_manager.p12</code> ✓ <code>conf/ca/private/user_Monitor_operator.p12</code> Les certificats clients sont utilisés pour authentifier le client lorsqu'il se connecte au service web en HTTPS. Pour pouvoir connecter la console web, le certificat client correspondant au rôle désiré doit d'abord être importé dans le magasin des certificats du navigateur web. ⇒ Copie le certificat du CA, du serveur ainsi que les certificats clients dans le répertoire conf
--	---

<pre>initssl req <url> <user>[:<password>]]</pre>	Paramètres <url>: Url du service web du serveur CA (https://192.168.0.1:9001) <user>, <password>: utilisateur et mot de passe protégeant l'accès au service web. La valeur par défaut pour <user> est CA_admin. La valeur pour <password> doit être celle affectée au moment du lancement du service web. Si ce champ n'est pas donné en argument, le script demandera sa saisie. Exemple <pre>initssl req https://192.168.0.1:9001 CA_admin:PasW0rD</pre> Description Tous les chemins d'accès ci-dessous sont relatifs au répertoire SAFE/web. <hostname> est le nom du serveur local. ⇒ Création d'une requête de certificat pour la génération du certificat serveur. La requête contient toutes les adresses IP et noms DNS associés au serveur local. La requête de certificat est stockée dans conf/ca/private/server_<hostname>.csr et la clé associée dans conf/ca/private/server_<hostname>.key. ⇒ Création d'une requête de certificat pour la génération du certificat client avec le rôle Admin (nécessaire pour les commandes distribuées sur le cluster). La requête de certificat est stockée dans conf/ca/private/user_Admin_<hostname>.csr et la clé associée dans conf/ca/private/user_Admin_<hostname>.key. ⇒ Téléchargement du certificat du CA depuis le serveur CA ⇒ Téléchargement depuis le serveur CA, des certificats signés qui ont été générés à partir des requêtes construites précédemment ⇒ Installation des certificats et des clés ⇒ Contrôle de validité des certificats
<pre>initssl req</pre>	Paramètres None Description Tous les chemins d'accès ci-dessous sont relatifs au répertoire SAFE/web. La commande se termine après avoir généré les requêtes de certificats pour : ⇒ le serveur local (conf/ca/private/server_<hostname>.csr) ⇒ le certificat client avec le rôle Admin (conf/ca/private/user_Admin_<hostname>.csr) Ces requêtes sont stockées dans le format base64 pour pouvoir être transmises à un CA externe tel Microsoft Active Directory Certificate Services (voir la documentation de Microsoft).

<pre>makeusercert <name> <role></pre>	Paramètres <name> correspond au champ CN du sujet du certificat, habituellement le nom de l'utilisateur du sujet <role> correspond au rôle lors de l'utilisation de la console web. Les valeurs valides sont Admin ou Control ou Monitor. Exemples <pre>makeusercert administrator Admin makeusercert manager Control makeusercert operator Monitor</pre> Description Tous les chemins d'accès ci-dessous sont relatifs au répertoire SAFE/web. Création d'une requête de certificat client (ainsi que le certificat, le fichier pkcs12, et la clé associée si la commande est exécutée sur le serveur CA) pour <name> et <role>. Lors de la génération du fichier pkcs12, le script demande d'entrer deux fois le mot de passe qui sera utilisé pour protéger son accès. La clé privée non cryptée est stockée dans conf/ca/private/user_<role>_<name>.key file. Quand ils ont générés, le certificat est stocké dans conf/ca/certs/user_<role>_<name>.crt et le pkcs12 dans conf/ca/private/user_<role>_<name>.p12.
---	--

11.6.5 Service web du serveur de CA

La configuration du service web du serveur de CA se trouve dans le fichier SAFE/web/conf/httpd.caserv.conf.

Ce service implémente une sous-partie des fonctionnalités d'un PKI :

⇒ Le certificat CA est accessible depuis l'URL

`https://CAserverIP>:9001/<certificate name>.crt`

Le téléchargement de ce fichier ne nécessite pas de s'authentifier.

⇒ Les requêtes de certificats sont soumises par l'envoi d'un POST à l'URL `https://<CA server IP>:9001/caserv`

Les arguments sont :

action = signrequest

name = <certificate name>

servercsr = <file content of the server certificate request>

or

usercsr = <file content of the client certificate request>

12.Cluster.xml pour la configuration d'un cluster SafeKit

- ⇒ 12.1 « Le fichier `cluster.xml` » [page 231](#)
- ⇒ 12.2 « Configuration du cluster SafeKit » [page 235](#)

Depuis SafeKit 7.2, SafeKit utilise le fichier de configuration : `cluster.xml`. Ce fichier définit tous les serveurs qui composent le cluster SafeKit ainsi que l'adresse IP (ou nom) de ces serveurs sur les réseaux utilisés pour communiquer avec les nœuds du cluster. Ce fichier permet aussi de spécifier l'usage des réseaux :

- ✓ un réseau de type framework (`framework="on"`) est un réseau utilisé pour les communications internes au framework de SafeKit. Il s'agit des communications globales au cluster et internes aux modules ; ces communications étant encryptées. Ce type de réseau est aussi utilisé pour l'exécution des commandes distribuées. Vous devez définir au moins un réseau de type framework qui inclut tous les nœuds du cluster. Il est recommandé de définir plusieurs réseaux de type framework pour tolérer au moins une défaillance réseau.
- ✓ un réseau de type console (`console="on"`) est un réseau sur lequel la console web SafeKit peut se connecter pour la configuration et l'administration du cluster et des modules. Ce type de réseau doit obligatoirement inclure tous les nœuds qui composent le cluster SafeKit. Vous pouvez définir plusieurs réseaux de type console en fonction des besoins d'administration et de la topologie réseau.

Par défaut, les réseaux sont utilisés à la fois par la console et le framework. Toutes les communications

12.1 Le fichier `cluster.xml`

Chaque réseau (`lan`) possède un nom logique qui sera utilisé dans la configuration des modules pour nommer les réseaux (à condition que le réseau soit configuré avec `framework="on"`) :

- ⇒ dans la section heartbeat pour un module miroir (voir section 13.3 [page 241](#))
- ⇒ dans la section lan pour un module ferme (voir section 13.4 [page 243](#))

Pour chaque réseau, il peut être spécifié s'il peut être utilisé par la console et/ou le framework. Par défaut, un réseau peut être utilisé à la fois par la console et le framework (`console="on" framework="on"`).

Le nom du nœud est utilisé par le service d'administration de SafeKit (`safeadmin`) pour identifier de manière unique un nœud SafeKit. Vous devez toujours utiliser le même nom pour désigner le même serveur sur les différents réseaux. Ce nom est aussi utilisé par la console web SafeKit lors de l'affichage du nom du nœud.

12.1.1 Cluster.xml exemple

Dans l'exemple ci-dessous, les 2 réseaux peuvent être utilisés à la fois par la console et le framework (par défaut : `console="on" framework="on"`).

```
<cluster>
  <lans>
    <lan name="default">
      <node name="node1" addr="192.168.1.67"/>
      <node name="node2" addr="192.168.1.68"/>
      <node name="node3" addr="192.168.1.69"/>
      <node name="node4" addr="192.168.1.70"/>
    </lan>
    <lan name="repli">
      <node name="node1" addr="10.0.0.1"/>
      <node name="node2" addr="10.0.0.2"/>
      <node name="node3" addr="10.0.0.3"/>
      <node name="node4" addr="10.0.0.4"/>
    </lan>
  </lans>
</cluster>
```

Dans l'exemple ci-dessous, le réseau `private` ne peut être utilisé par la console puisqu'il n'inclut pas tous les nœuds du cluster. Il s'agit par exemple d'un lien de réplication dédié pour un module de type miroir.

```
<cluster>
  <lans>
    <lan name="default">
      <node name="node1" addr="192.168.1.67"/>
      <node name="node2" addr="192.168.1.68"/>
      <node name="node3" addr="192.168.1.69"/>
      <node name="node4" addr="192.168.1.70"/>
    </lan>
    <lan name="repli" console="off">
      <node name="node1" addr="10.0.0.1"/>
      <node name="node2" addr="10.0.0.2"/>
    </lan>
  </lans>
</cluster>
```

Dans l'exemple ci-dessous, le réseau `public` est utilisé uniquement pour l'administration via la console. Il s'agit par exemple d'un réseau public sur lequel l'administrateur ne souhaite pas voir transiter les communications internes au cluster SafeKit.

```
<cluster>
  <lans>
    <lan name="default">
      <node name="node1" addr="192.168.1.67"/>
      <node name="node2" addr="192.168.1.68"/>
    </lan>
    <lan name="public" framework="off">
      <node name="node1" addr="node1.mydomain.com"/>
      <node name="node2" addr="node2.mydomain.com"/>
    </lan>
  </lans>
</cluster>
```

Dans l'exemple ci-dessous, un seul réseau est utilisé, mais dans une configuration NAT (Network address translation). Pour chaque nœud du cluster deux adresses doivent être

définies : L'adresse locale (celle de l'interface locale) et l'adresse externe (celle connue des autres nœuds).

```
<cluster>
  <langs>
    <lan name="default">
      <node name="node1" addr="server1.dns.name" laddr="10.0.0.1"/>
      <node name="node2" addr="server2.dns.name" laddr="10.0.0.2"/>
    </lan>
  </langs>
</cluster>
```

Notes :

- ✓ Tous les nœuds doivent pouvoir communiquer avec les autres via les adresses externes.
- ✓ Si un LAN NATé est utilisé comme LAN console, la console Web doit avoir accès aux nœuds via les adresses externes.
- ✓ La configuration d'un cluster avec adresses Natés ne peut être effectué via la console Web. L'interface en ligne de commande doit être utilisée (voir section 12.2.2 page 236).

12.1.2 Cluster.xml syntaxe

```
<cluster>
  <langs [port="4800"]>
    <lan name="lan_name" [console="on|off"] [framework="on|off"]
[command="on|off"] >
      <node name="node_name" addr="IP1_address"|"IP1_name"
[ laddr="local_IP1_address" ]/>
      <node name="node_name" addr="IP2_address"|"IP2_name"
[ laddr="local_IP2_address" ]/>
      ...
    </lan>
    ...
  </langs>
</cluster>
```

12.1.3 <langs>, <lan>, <node> attributs

<langs	Début de la définition des nœuds SafeKit et de la topologie réseau
[port="xxxx"]	Port UDP sur lequel le protocole membership échange. Valeur par défaut: 4800
[pulse="xxxx"]	Période d'émission des messages du protocole d'appartenance. Un pulse élevé utilise moins de bande passante, mais entraîne un délai de réaction plus long.
[mlost_count="xx"]	Nombre de périodes écoulées sans réception de message avant d'élire un nouveau leader.
[slost_count="xx"]	Nombre de périodes écoulées sans réception de message avant de déclarer un follower hors ligne.

<lan	Définition d'un réseau sur lequel s'exécute le protocole membership. Au moins un réseau. Autant de sections qu'il y a de LANs entre les serveurs.
name="lan name"	Nom logique unique pour le réseau Ce nom est utilisé dans la configuration du module pour définir les réseaux utilisés.
framework="on" "off"	Mettre <code>framework="off"</code> pour ne pas utiliser ce réseau pour les communications du framework SafeKit. Dans ce cas, ce réseau ne peut être utilisé dans la configuration d'un module. Par défaut, <code>framework = "on"</code> . Vous pouvez définir plusieurs sections <lan> avec <code>framework="on"</code> ou <code>framework="off"</code> . Il faut au moins une section <lan> avec <code>framework="on"</code> , qui inclut tous les nœuds du cluster. Valeur par défaut : <code>on</code>
console="on" "off"	Mettre <code>console="off"</code> pour ne pas utiliser ce réseau pour connecter la console web SafeKit. Par défaut, <code>console="on"</code> . Quand <code>console="on"</code> , la section <lan> doit inclure tous les nœuds qui composent le cluster. Vous pouvez définir plusieurs sections <lan> avec <code>console="on"</code> ou <code>console="off"</code> . Il faut au moins une section <lan> avec <code>console="on"</code> si vous souhaitez utiliser la console web. Valeur par défaut : <code>on</code>
command="on" "off"	Mettre <code>command="on"</code> pour utiliser ce réseau pour l'exécution des commandes distribuées sur le cluster. Dans ce cas, la section <lan> doit inclure tous les nœuds qui composent le cluster et avoir aussi <code>framework="on"</code> . Il faut définir une unique section <lan> avec <code>command="on"</code> . Quand cet attribut n'est pas positionné, c'est la première section <lan> avec <code>framework="on"</code> qui est utilisée pour l'exécution des commandes distribuées sur le cluster. Valeur par défaut : <code>off</code>
<node	Définition d'un nœud dans le réseau. Positionner autant de nœuds qu'il y a de serveurs dans le cluster SafeKit (au moins 2).
name="node name"	Nom unique logique pour le serveur SafeKit Vous devez toujours utiliser le même nom pour désigner le même serveur sur les différents réseaux.

<pre>addr= "IP_address" "IP_name"</pre>	Adresse IPv4 ou IPv6, ou nom du nœud tel qu'il est connu par les autres nœuds dans le LAN (préférer une adresse IP pour être indépendant d'un serveur de noms DNS). Pour des configurations NAT, c'est l'adresse extérieure qui doit être indiquée. Lors de la définition d'une adresse IPv6, utiliser le format littéral : l'adresse est placée entre crochets (par exemple [2001::7334])
<pre>laddr= "local_IP_address"</pre>	Adresse IP locale dans le LAN. A utiliser uniquement pour les configurations NAT.

12.2 Configuration du cluster SafeKit

12.2.1 Configuration avec la console web de SafeKit

La console web de SafeKit fournit une interface graphique pour l'édition du fichier `cluster.xml` et appliquer la configuration sur tous les nœuds qui composent le cluster SafeKit. Pour la description complète, voir la section 3.2 [page 37](#).

- Cliquer sur > Configuration du cluster pour ouvrir le panneau. La liste des nœuds du cluster s'affiche.
- Editer la configuration. Dans le mode d'édition Simple, vous ne pouvez éditer que le réseau de connexion de la console. Passer dans le mode d'édition Avancé pour éditer tous les réseaux.
- Cliquer sur le bouton Appliquer pour sauvegarder la configuration et générer de nouvelles clés de chiffrement.
- Dans le mode d'édition Simple, ce bouton n'est actif que si des changements ont été faits. Si vous souhaitez régénérer de nouvelles clés ou réappliquer la configuration sur tous les nœuds sans changer la configuration, basculer dans le mode d'édition Avancé, puis cliquer sur le bouton Appliquer.

Configuration du cluster - cluster1

Attention : cette configuration est exploitée par tous les modules. Par conséquent, sa modification peut impacter plusieurs modules et nécessiter leur arrêt et reconfiguration

Appliquer

Recharger

Mode d'édition : ☐ Simple ☒ Avancé

Paramètres de sécurité

Nœuds du cluster

Cliquez sur Nouveau nœud pour définir les serveurs qui composent le cluster

-

centos7-test3 / Linux 3.10.0-1160.11.1.el7.x86_64 / SafeKit 7.5.0.9 / configuration modifiée et appliquée le 2021-06-16 10:05:04

adresse

10.0.0.103

nom

centos7-test3

-

centos7-test4 / Linux 3.10.0-1160.21.1.el7.x86_64 / SafeKit 7.5.0.9 / configuration modifiée et appliquée le 2021-06-16 10:05:04

adresse

10.0.0.104

nom

centos7-test4

12.2.2 Configuration en ligne de commandes

Les commandes équivalentes pour configurer le cluster SafeKit avec une nouvelle clé de chiffrement sont :

⇒ `safekit cluster config [<filepath>]`

où `filepath` est le chemin d'accès au nouveau `cluster.xml`

quand `filepath` n'est pas passé en argument, la configuration courante est conservée et seule la clé d'encryption est régénérée

⇒ `safekit -H "*" -G`

la configuration est appliquée sur les nœuds du cluster

Les commandes équivalentes pour reconfigurer sans clé de chiffrement sont :

⇒ `safekit cluster delkey`

⇒ `safekit -H "*" -G`

Les commandes équivalentes pour régénérer des clés de chiffrement et les prendre en compte sont :

⇒ `safekit cluster genkey`

⇒ `safekit -H "*" -G`

Pour la description complète des commandes, se référer à la section 9.3 [page 148](#).

12.2.3 Changements de configuration

Lorsque la configuration du cluster SafeKit est modifiée, la nouvelle configuration doit impérativement être appliquée sur tous les serveurs qui composent le cluster. Si la configuration n'est appliquée que sur un sous-ensemble des nœuds présents dans la configuration du cluster, seul le sous-ensemble des nœuds sera en mesure de communiquer. Cela peut avoir pour conséquence de perturber le fonctionnement des modules installés sur les serveurs. Pour rétablir un fonctionnement correct, vous devez réappliquer la configuration sur tous les nœuds du cluster comme décrit précédemment.



Il est possible d'afficher la configuration courante en exécutant la commande `safekit cluster confinfo` sur chaque nœud (voir section 9.3 [page 148](#)). Quand la configuration est correcte, cette commande retourne sur tous les nœuds, la même liste de nœuds et la même signature de configuration.

Changer la configuration du cluster peut aussi avoir un impact important sur la configuration des modules car les noms logiques de réseau `<lan>` sont utilisés dans les configurations de modules. Tout changement de configuration déclenche une mise à jour des modules démarrés pour prendre en compte ces modifications. Cela peut conduire à un arrêt de ces modules en cas d'incompatibilité (par exemple sur destruction d'un réseau alors qu'il est utilisé par un module). Aussi, il faut être prudent lors de la modification de la configuration du cluster quand des modules sont en cours de fonctionnement.

13. Userconfig.xml pour la configuration d'un module

- ⇒ 13.1 « Macro définition (<macro> tag) » [page 238](#)
- ⇒ 13.2 « Module ferme ou miroir (<service> tag) » [page 238](#)
- ⇒ 13.3 « Heartbeats (<heart>, <heartbeat > tags) » [page 241](#)
- ⇒ 13.4 « Topologie d'une ferme (<farm>, <lan> tags) » [page 243](#)
- ⇒ 13.5 « Adresse IP virtuelle (<vip> tag) » [page 245](#)
- ⇒ 13.6 « Réplication de fichiers (<rfs>, <replicated> tags) » [page 252](#)
- ⇒ 13.7 « Activer les scripts utilisateurs (<user>, <var> tags) » [page 270](#)
- ⇒ 13.8 « Hostname virtuel (<vhost>, <virtualhostname> tags) » [page 271](#)
- ⇒ 13.9 « Détection de la mort de processus ou de services (<errd>, <proc> tags) » [page 273](#)
- ⇒ 13.10 « Checkers (<check> tags) » [page 279](#)
- ⇒ 13.11 « TCP checker (<tcp> tags) » [page 280](#)
- ⇒ 13.12 « Ping checker (<ping> tags) » [page 281](#)
- ⇒ 13.13 « Interface checker (<intf> tags) » [page 282](#)
- ⇒ 13.14 « IP checker (<ip> tags) » [page 283](#)
- ⇒ 13.15 « Checker personnalisé (<custom> tags) » [page 285](#)
- ⇒ 13.16 « Module checker (<module> tags) » [page 286](#)
- ⇒ 13.17 « Splitbrain checker (<splitbrain> tag) » [page 288](#)
- ⇒ 13.18 « Failover machine (<failover> tag) » [page 289](#)



A chaque fois que vous modifiez `userconfig.xml`, vous devez appliquer la nouvelle configuration sur les serveurs avec :

- ✓ la console web/ Configuration Avancée / Modules installés/ module/ Appliquer la configuration
- ✓ ou la console web/ Configuration/ sur le module/ Editer la configuration
- ✓ ou la commande `safekit config -m AM`

Il est possible d'appliquer une nouvelle configuration alors que le module est démarré, mais uniquement dans l'état `ALONE` (vert) et `WAIT` (rouge). Cette fonctionnalité est appelée *configuration dynamique*. Uniquement un sous-ensemble restreint de la configuration peut être modifié dynamiquement. Quand la nouvelle configuration ne peut être appliquée, un message d'erreur est affiché. Les attributs de configuration pouvant être changés dynamiquement sont listés ci-après.



Exemple de userconfig.xml

```
<safe>
  <!-- Insérer ci-dessous les tags <macro> <service> -->
</safe>
```

13.1 Macro définition (<macro> tag)

13.1.1 <macro> Exemple

```
<macro name="ADDR1" value="aa.bb.com"/>
```

Un exemple d'utilisation de macro est donné dans 15.3 page 305.

13.1.2 <macro> Syntaxe

```
<macro
  name="identifiant"
  value="value"
/>
```

13.1.3 <macro> Attributs

<macro	
name="identifiant"	Une chaîne de caractères qui identifie la macro.
value="value"	La valeur qui remplacera chaque occurrence de %identifiant% dans la suite de userconfig.xml.
/>	



La syntaxe %identifiant% peut être aussi utilisée dans userconfig.xml pour récupérer la valeur d'une variable d'environnement de nom identifiant. En cas de conflit, c'est la valeur de macro qui prime.

13.2 Module ferme ou miroir (<service> tag)

13.2.1 <service> Exemple

Exemple pour un module miroir :

```
<service mode="mirror"
defaultprim="alone" maxloop="3" loop_interval="24" failover="on">
  <!-- Insérer ci-dessous les tags <heartbeat> <rfs> <vip> <user> <vhost>
<errd> <check> <failover> -->
</service>
```

Exemple pour un module ferme :

```
<service mode="farm" maxloop="3" loop_interval="24">
  <!-- Insérer ci-dessous les tags <farm> <vip> <user> <vhost> <errd> <check>
<failover> -->
</service>
```

Des exemples de définition de `<service>` sont donnés pour un module miroir dans 15.1 [page 302](#) et pour un module ferme dans 15.2 [page 303](#).

13.2.2 <service> Syntaxe

```
<service mode="mirror"|"farm"|"light"
  [boot="off"|"on"|"auto"|"ignore"]
  [boot_delay="0"]
  [failover="on"|"off"]
  [defaultprim="alone"|"server_name"|"lastprim"]
  [maxloop="3"] [loop_interval="24"]
  [automatic_reboot="off"|"on"]>
</service>
```



Seuls les attributs `boot`, `maxloop`, `loop_interval` et `automatic_reboot` peuvent être modifiés dynamiquement.

13.2.3 <service> Attributs

<code><service</code>	Première section à définir dans un module
<code>mode=</code> <code>"mirror" </code> <code>"farm" </code> <code>"light"</code>	<code>mirror</code> pour un module miroir. Le protocole de synchronisation entre les 2 serveurs est défini en 13.3 page 241. Voir le module applicatif <code>mirror.safe</code> en tant qu'exemple. <code>farm</code> pour un module ferme. Le protocole de synchronisation entre les serveurs est défini en 13.4 page 243. Voir le module applicatif <code>farm.safe</code> en tant qu'exemple. <code>light</code> pour une configuration sur un seul serveur avec une détection d'erreur logicielle et un redémarrage local seulement.
<code>[boot=</code> <code>"on" </code> <code>"off" </code> <code>"auto" </code> <code>"ignore"]</code>	Si <code>on</code>, le module est automatiquement démarré au boot de la machine. Si <code>off</code>, le module n'est pas démarré au boot de la machine. Si <code>auto</code>, le module est démarré automatiquement au boot de la machine s'il était démarré avant le reboot. Avant SafeKit 7.5, la configuration du démarrage au boot du module se faisait avec la commande <code>safekit boot -m AM on off</code> (qui devait être exécutée sur chaque nœud). Si vous préférez continuer à utiliser cette commande, supprimez l'attribut <code>boot</code> ou affectez-lui la valeur <code>ignore</code> (valeur par défaut). Le module ne sera pas démarré au boot, à moins que la commande <code>safekit boot -m AM on</code> ne soit exécutée. L'état de la configuration du démarrage au boot est visible dans la ressource <code>usersetting.boot</code>. L'état des ressources est visible dans la console web/Contrôle /Sélection du nœud/onglet Ressources ; via la commande <code>safekit state -m AM -v</code> Valeur par défaut : <code>off</code>
<code>[boot_delay="0"]</code>	Le délai, en secondes, avant le démarrage du module au boot.

	Valeur par défaut : 0 (pas de délai)
[failover= "on" "off"]	Pour un module miroir seulement. Si <code>on</code>, reprise automatique sur le serveur secondaire lorsque le serveur primaire est arrêté. Si <code>off</code>, lorsque le serveur primaire est arrêté, le serveur secondaire se met en attente (pas de reprise automatique). Seule la commande <code>safekit prim</code> peut forcer le démarrage du serveur secondaire en primaire. Pour une description, voir la section 5.7 page 106. Valeur par défaut : <code>on</code>
[defaultprim= "alone" "server_name" "lastprim"]	Pour un module miroir seulement. <code>defaultprim</code> décide quel serveur parmi 2 serveurs est le serveur primaire par défaut pour un module applicatif. Cette option est utile quand un module est <code>ALONE</code> sur un serveur et que le module est démarré sur l'autre serveur. Avec <code>defaultprim="alone"</code>, le module <code>ALONE</code> devient <code>PRIM</code> alors que le module redémarré devient <code>SECOND</code>. Valeur recommandée pour éviter les basculements d'application juste après la réintégration. Avec <code>defaultprim="server_name"</code>, lorsque le module tourne sur deux serveurs, le serveur <code>PRIM</code> est celui indiqué dans <code>defaultprim</code>. Cette valeur peut être utile dans les architectures actif/actif (voir section 1.5.1 page 20) ou N-1 (voir section 1.5.2 page 21). Avec <code>defaultprim="lastprim"</code>, le serveur redémarré redevient <code>PRIM</code> s'il était <code>PRIM</code> avant son dernier arrêt. Valeur par défaut : <code>alone</code>
[maxloop="3"]	Nombre de détections d'erreur successives avant arrêt. Cet attribut définit le maximum de <code>restart</code> ou <code>stopstart</code> appelés par les détecteurs d'erreur avant d'arrêter localement SafeKit. Ce compteur est réinitialisé à sa valeur initiale à l'expiration du <code>timeout loop_interval</code> ou lors d'une commande manuelle <code>safekit start, restart, swap, stopstart...</code> Noter qu'une commande <code>safekit</code> émise par un détecteur avec l'option <code>-i identity</code> décrémente le compteur, alors qu'une commande manuelle sans cette option ne décrémente pas le compteur. Pour plus d'informations, voir section 13.18.4 page 290.  Note La valeur de cet attribut peut être modifiée dynamiquement. Depuis SafeKit 7.5, <code>maxloop</code> est représenté par la ressource <code>heart.stopstartloop</code>. Sa valeur courante correspond à la date à laquelle le compteur a été initialisé (sous la forme d'un timestamp Epoch Unix) ; et sa date d'affectation correspond soit à son

	initialisation, soit à un rebouclage (<code>stopstart</code>, <code>restart</code>). Consulter l'historique des ressources pour visualiser chaque incrémentation du compteur. Valeur par défaut : 3
<pre>[loop_interval ="24"]</pre>	Intervalle de temps, en heures, sur lequel <code>maxloop</code> s'applique. Affectez sa valeur à 0 pour désactiver le compteur <code>maxloop</code>. Valeur par défaut : 24 heures  La valeur de cet attribut peut être modifiée dynamiquement.
<pre>[automatic_reboot ="off" "on"]</pre>	Si positionné à <code>on</code>, reboot sur <code>safekit stopstart</code> au lieu de stopper puis redémarrer. Valeur par défaut : <code>off</code>  La valeur de cet attribut peut être modifiée dynamiquement.

13.3 Heartbeats (<heart>, <heartbeat > tags)

Les heartbeats doivent être utilisés seulement avec les modules miroirs. La topologie d'une ferme est décrite dans la section 13.4 [page 243](#).

Le mécanisme basique pour synchroniser deux serveurs et détecter les défaillances d'un serveur est le heartbeat (battement de cœur), qui consiste en un échange de petits paquets UDP entre les serveurs. Normalement, on met autant de heartbeats qu'il y a de réseaux connectant les 2 serveurs. Dans une situation normale, les 2 serveurs échangent leurs états (`PRIM`, `SECOND`, les états des ressources) à travers les heartbeats et synchronisent ainsi les procédures de démarrage arrêt applicatif.

Si tous les heartbeats sont perdus, cela signifie que l'autre serveur est en panne. Le serveur local décide de devenir `ALONE`. Bien que non obligatoire, il est préférable d'avoir deux voies de heartbeats sur deux réseaux différents afin de distinguer une panne réseau d'une panne serveur et d'éviter le cas du splitbrain.

13.3.1 <heart> Exemple

```
<heart>
  <heartbeat name="default" ident="Hb1" />
  <heartbeat name="net2" ident="Hb2" />
</heart>
```

Un exemple de configuration de heartbeats avec de multiples voies est donné en 15.4 [page 306](#).

13.3.2 <heart> Syntaxe

```
<heart
```

```
[port="xxxx"] [pulse="700"] [timeout="30000"]
[permanent_arp="on"]
>
<heartbeat
  [port="xxxx"] [pulse="700"] [timeout="30000"] name="network" [ident="name"]
>
[<!-- syntaxe pour SafeKit < 7.2 -->
  <server addr="IP1_address" />
  <server addr="IP2_address" />
]
</heartbeat>
...
</heart>
```



Le tag `<heart>` et son sous-arbre peuvent être entièrement modifiés dynamiquement.

13.3.3 `<heart>`, `<heartbeat attributs`

<code><heart</code>	
<code>[port="xxxx"]</code>	Port UDP sur lequel les heartbeats sont échangés. Valeur par défaut : dépend de l'id du module applicatif. Rendu par la commande <code>safekit module getports</code> .
<code>[pulse="700"]</code>	Délai en milliseconde entre l'émission de 2 paquets de <code>heartbeat</code> . Valeur par défaut : 700 ms
<code>[timeout="30000"]</code>	Timeout de détection en ms de la perte d'un <code>heartbeat</code> . Valeur par défaut : 30 000 ms
<code><heartbeat</code>	Définition d'un <code>heartbeat</code> . Il y a autant de sections <code><heartbeat></code> qu'il y a de voies de heartbeats (réseaux connectant les serveurs). Au moins 1 <code>heartbeat</code> .
<code>[port="xxxx"]</code>	Redéfinition du port de <code>heartbeat</code> . Par défaut le même que celui dans <code><heart></code> .
<code>[pulse="700"]</code>	Redéfinition du délai en milliseconde entre l'émission de 2 paquets de <code>heartbeat</code> . Par défaut le même que celui dans <code><heart></code> .
<code>[timeout="30000"]</code>	Redéfinition du timeout de détection en ms de la perte d'un <code>heartbeat</code> . Par défaut le même que celui dans <code><heart></code> .
<code>name="network"</code>	Nom du réseau utilisé par le <code>heartbeat</code> . "network" doit être le nom d'un réseau défini dans la configuration du cluster SafeKit (voir section 12 page 231).

<code>[ident="name"]</code>	Donne le nom qui sera utilisé pour ce heartbeat dans la console web ainsi que pour la ressource interne correspondante, i.e. : <code>heartbeat.name</code> peut être utilisé dans la failover machine (voir section 13.18 page 289). Si l'attribut <code>ident</code> n'est pas défini la valeur de l'attribut <code>name</code> est utilisée.  Important Si vous positionnez un heartbeat <code>ident="flow"</code>, le flux de réplication sera positionné automatiquement sur la même voie. Si vous positionnez <code>ident="flow"</code> sans configuration <code><rfs></code>, le démarrage du module bloque dans l'état <code>WAIT</code>.
<code>[permanent_arp="on" "off"]</code>	Régulièrement, <code>heart</code> positionne un ARP permanent pour ses voies de heartbeats. Cette procédure peut geler <code>heart</code> sur certains systèmes (Linux). Dans ce cas, positionner à "off" et affecter l'ARP permanent sur les voies de heartbeats au boot. Sur Linux, ceci peut être fait en ajoutant la commande suivante dans un script exécuté au boot : <pre>arp -s hostname hw_addr</pre> Valeur par défaut : on
<code><server addr="IP1_address" /></code>	Définition de l'adresse ip du serveur pour ce heartbeat. Le tag <code><server></code> était utilisé dans l'ancienne syntaxe de configuration (avant SafeKit 7.2). Il est supporté pour assurer la compatibilité ascendante, mais ne doit pas être utilisé pour la configuration de nouveaux modules.  Important Vous ne devez pas utiliser dans le même <code>userconfig.xml</code>, la syntaxe de SafeKit 7.1 et celle introduite depuis SafeKit 7.2.

13.4 Topologie d'une ferme (<farm>, <lan> tags)

Le mécanisme basique pour synchroniser une ferme de serveurs est un protocole de groupe qui détecte automatiquement les membres disponibles dans le groupe (protocole membership).

13.4.1 <farm> Exemple

```
<farm>
  <lan name="default" />
  <lan name="net2" />
</farm>
```

Pour des exemples de configuration `<farm>`, voir section 15.5 [page 306](#).

13.4.2 <farm> Syntaxe

```
<farm [port="xxxx"]>
  <lan name="network">
    [<!-- syntaxe pour SafeKit < 7.2 -->
```

```
<node name="server1" addr="IP1_address"|"IP1_name" />
<node name="server2" addr="IP2_address"|"IP2_name" />
]
</lan>
...
</farm>
```



Le tag `<farm>` et son sous-arbre **ne peuvent pas** être modifiés dynamiquement.

13.4.3 `<farm>`, `<lan>` Attributs

<code><farm</code>	
<code>[port="xxxx"]</code>	Port UDP sur lequel le protocole membership échange. Valeur par défaut : dépend de l'id du module applicatif. Rendu par la commande <code>safekit module getports</code>
<code>[pulse= "xxx"]</code>	Période d'émission des messages du protocole d'appartenance. Un pulse élevé utilise moins de bande passante, mais entraîne un délai de réaction plus long.
<code>[mlost_count= "xx"]</code>	Nombre de périodes écoulées sans réception de message avant d'élire un nouveau leader.
<code>[slost_count= "xx"]</code>	Nombre de périodes écoulées sans réception de message avant de déclarer un follower hors ligne.
<code><lan</code>	Définition d'un lan sur lequel s'exécute le protocole membership. Au moins un lan doit être défini. Autant de sections qu'il y a de réseaux entre les serveurs.
<code>name="network"</code>	Nom du réseau utilisé. <code>network</code> doit être le nom d'un réseau défini dans la configuration du cluster SafeKit (voir section 12 page 231).
<code>< node name="identity" addr= "IP1_address" /></code>	Adresse IP et nom du nœud dans ce lan. Le tag <code><node></code> était utilisé dans l'ancienne syntaxe de configuration (avant SafeKit 7.2). Il est supporté pour assurer la compatibilité ascendante, mais ne doit pas être utilisé pour la configuration de nouveaux modules.  Vous ne devez pas utiliser dans le même <code>userconfig.xml</code> , la syntaxe de SafeKit 7.1 et celle introduite depuis SafeKit 7.2.

13.5 Adresse IP virtuelle (<vip> tag)



Si vous installez plusieurs modules applicatifs sur le même serveur, l'adresse IP virtuelle doit être différente pour chaque module.

13.5.1 <vip> Exemple dans une architecture ferme

L'exemple ci-dessous configure l'équilibrage de charge, à destination du port 80 et de l'adresse IP virtuelle, entre les nœuds d'un cluster sur site :

```
<vip>
  <interface_list>
    <interface check="on" arpreroute="on" arpinterval="60" arpelapse="10">
      <virtual_interface type="vmac_directed">
        <virtual_addr addr="192.168.1.222" where="alias" check="on"/>
      </virtual_interface>
    </interface>
  </interface_list>
  <loadbalancing_list>
    <group name="FarmProto">
      <rule port="80" proto="tcp" filter="on_port"/>
    </group>
  </loadbalancing_list>
</vip>
```

Voir aussi l'exemple en 15.2 [page 303](#).

13.5.2 <vip> Exemple dans une architecture miroir

L'exemple ci-dessous configure l'adresse IP virtuelle sur le nœud primaire d'un cluster sur site :

```
<vip>
  <interface_list>
    <interface check="on" arpreroute="on">
      <real_interface>
        <virtual_addr addr="192.168.1.222" where="one_side_alias"
check="on"/>
      </real_interface>
    </interface>
  </interface_list>
</vip>
```

Voir aussi l'exemple en 15.1 [page 302](#).

13.5.3 Alternative à <vip> pour des serveurs dans des réseaux IP différents

La configuration d'une adresse IP virtuelle avec une section <vip> dans `userconfig.xml` requiert des serveurs dans le même réseau IP (reroutage réseau et équilibrage de charge effectués au niveau 2).

Si les serveurs se trouvent dans des réseaux IP différents, la section <vip> ne peut pas être configurée. Dans ce cas, une alternative consiste à configurer l'adresse IP virtuelle dans un équilibreur de charge. L'équilibreur de charge achemine les paquets vers les adresses IP physiques des serveurs en testant le statut d'une URL nommée vérificateur d'état (health check) et géré par SafeKit.

SafeKit fournit donc un vérificateur d'état pour chaque module. Vous devez configurer le test de vérification dans le load balancer avec :

- ⇒ le protocole HTTP
- ⇒ le port 9010, port du service web de SafeKit
- ⇒ l'URL `/var/modules/AM/ready.txt` où AM est le nom du module

Pour un module miroir, le test retourne :

- ⇒ OK, qui signifie que l'instance est saine, quand le module est dans l'état  PRIM (vert) ou  ALONE (vert)
- ⇒ NOT FOUND, qui signifie que l'instance est hors service, dans tous les autres états

Pour un module ferme, le test retourne :

- ⇒ OK, qui signifie l'instance est saine, quand le module est dans l'état  UP (vert)
- ⇒ NOT FOUND, qui signifie que l'instance est hors service, dans tous les autres états

Une autre alternative consiste à ce que vous implémentiez une configuration DNS spéciale et une commande de redirection DNS insérée dans les scripts de redémarrage de SafeKit.

13.5.4 <vip> Syntaxe

13.5.4.1 Partage de charge réseau dans une architecture ferme

```
<vip [tcprset="off"|"on"]>
  <interface_list>
    <interface
      [check="off"|"on"]
      [arpreroute="off"|"on"]
      [arpinterval="60"]
      [arpelapse="1200"]
    >
    <virtual_interface
      [type="vmac_directed"|"vmac_invisible"]
      [addr="xx:xx:xx:xx:xx"]
    >
    <virtual_addr
      addr="virtual_IP_name"|"virtual_IP_address"
      [where="alias"]
      [check="off"|"on"]
      [connections="off"|"on"]
    />
    ...
  </virtual_interface>
</interface>
</interface_list>
<loadbalancing_list>
  <group name="group_name"
    <cluster>
      <host name="node_name" power="integer" />
      ...
    </cluster>
  <rule
    [virtual_addr="*"|"virtual_IP_name"|"virtual_IP_address"]
    [port="*"|"value"]
    proto="udp"|"tcp"
```

```

        filter="on_addr"|"on_port"|"on_ipid"
    />
    ...
</group>
...
</loadbalancing_list>
</vip>

```



Le tag <vip> et son sous-arbre **ne peuvent pas** être modifiés dynamiquement.

13.5.4.2 Basculement réseau dans une architecture miroir

Pour un cluster sur site :

```

<vip [tcpreset="off"|"on"]>
  <interface_list>
    <interface
      [check="off"|"on"]
      [arpreroute="off"|"on"]
      [arpinterval="60"]
      [arpelapse="1200"]
    >

    <real_interface>
      <virtual_addr
        addr="virtual_IP_name"|"virtual_IP_address"
        where="one_side_alias"
        [check="off"|"on"]
        [connections="off"|"on"]
      />
      ...
    </real_interface>
  </interface>
  ...
</interface_list>
</vip>

```

13.5.5 <interface_list>, <interface>, <virtual_interface>, <real_interface>, <virtual_addr> Attributs

<vip>	
[tcpreset="off" "on"]	Avant de déconfigurer l'adresse IP virtuelle, les connexions l'ayant comme IP source, sont rompues. La rupture de connexion est désactivée en positionnant cet attribut à <code>off</code> . Valeur par défaut : <code>on</code>
<interface_list>	

<code><interface</code>	Définition des adresses IP virtuelles sur une interface. Mettre autant de sections <code><interface></code> que vous avez d'interfaces réseau à configurer.
<code>[check="off" "on"]</code>	Positionner un checker sur l'interface réseau. Le module est mis dans l'état <code>WAIT</code> tant que l'interface est down. Le nom du checker d'interface est <code>intf.<network_IP_mask> (intf.192.168.0.0)</code> . Valeur par défaut : <code>on</code> Pour plus d'informations, voir section 13.13 page 282 .
<code>[arpreroute="off" "on"]</code>	Broadcast de gratuitous ARP pour le reroutage des adresses IP virtuelles définies dans les sections <code><real_interface></code> . Valeur par défaut : <code>off</code>
<code>[arpinterval="60"]</code>	Temps en seconde entre 2 gratuitous ARP. Valeur par défaut : <code>60 s</code>
<code>[arpelapse="1200"]</code>	Temps total pendant lequel des gratuitous ARP sont émis. Valeur par défaut : <code>1200 s</code>
<code>[name="interface name"]</code>	Linux seulement. Vous pouvez spécifier le nom de l'interface. Exemple, positionner <code>name="bond0"</code> Par défaut, SafeKit détecte l'interface réseau à configurer à partir des adresses IP virtuelles configurées sur cette interface.

13.5.5.1 `<virtual_interface>`, `<virtual_addr>` Attributs dans une architecture ferme

A utiliser pour les modules ferme avec partage de charge sur l'IP virtuelle :

<code><virtual_interface</code>	Définition des adresses IP virtuelles configurées sur une interface Ethernet.
<code>type=</code> <code>"vmac_directed" </code> <code>"vmac_invisible"</code>	<code>vmac_directed</code> : Associe l'adresse MAC de l'un des serveurs à l'adresse IP virtuelle, comme pour le reste du trafic normal. Voir la description en 13.5.7.3 page 252 <code>vmac_invisible</code> : adresse MAC virtuelle jamais visible dans les entêtes Ethernet pour permettre le broadcast des switchs. Nécessite le support du mode promiscuous. Voir la description en 13.5.7.2 page 251 Note : configuration possible pour un module miroir et pour un reroutage transparent sans gratuitous ARP.
<code>[addr="xx:xx:xx:xx:xx"]</code>	Unicast virtual MAC adresse.

	Si non positionné, par défaut concaténation de "5A:FE" (Safe) et de la 1 ^{ère} adresse IP virtuelle en hexadécimal. Ignoré lorsque type="vmac_directed"
<virtual_addr	Définition d'une adresse IP virtuelle. Mettre autant de lignes <virtual_addr> qu'il y a d'adresses IP virtuelles à configurer sur l'interface.
addr="virtual_IP_name" "virtual_IP_address"	Nom ou adresse IP virtuelle (préférer une adresse IP pour être indépendant de la panne du serveur de nom). Adresse IPv4 ou IPv6.
where="alias"	L'adresse IP virtuelle est définie sur tous les serveurs de la ferme en alias. Note : Dans le cas particulier d'une configuration d'un module miroir avec VMAC mettre ici where="one_side_alias".
[check="off" "on"]	Positionner un IP checker sur l'adresse virtuelle. Le module exécute un stopstart quand l'IP virtuelle est détruite. Le nom de l'IP checker est ip.<virtual addr value> (ip.192.168.1.99). Valeur par défaut : on Pour plus d'informations, voir section 13.14 page 283 .
[connections="off" "on"]	Active le comptage du nombre de connexions actives sur l'adresse virtuelle. Ce nombre est stocké dans la ressource nommée connections.<addr value> (par exemple : connections.192.168.1.99) qui est affectée toutes les 10 secondes. Cette valeur est fournie à un titre indicatif uniquement. Valeur par défaut : off
netmask="defaultnetmask"	Linux et IPV4 seulement Par défaut, prend le netmask de l'interface. A positionner si l'interface a plusieurs netmasks.
</virtual_interface>	

13.5.5.2 <real_interface>, <virtual_addr> Attributs dans une architecture miroir

A utiliser pour les modules miroir avec basculement de l'IP virtuelle :

<real_interface>	Définition d'adresses IP virtuelle associée avec l'adresse MAC réel de l'interface.
<virtual_addr	Définition d'une adresse IP virtuelle. Mettre autant de lignes <virtual_addr> qu'il y a d'adresses IP virtuelles à configurer sur l'interface.

addr= "virtual_IP_name" "virtual_IP_address"	Nom ou adresse IP virtuelle (préférer une adresse IP pour être indépendant de la panne du serveur de nom). Adresse IPv4 ou IPv6.
where="one_side_alias"	Adresse mise en alias sur le serveur PRIM ou ALONE.
[check="off" "on"]	Positionner un IP checker sur l'adresse virtuelle. Le module exécute un <code>stopstart</code> quand l'IP virtuelle est détruite. Le nom de l'IP checker est <code>ip.<addr value></code> (<code>ip.192.168.1.99</code>). Valeur par défaut : <code>on</code> Pour plus d'informations, voir section 13.14 page 283 .
[connections="off" "on"]	Active le comptage du nombre de connexions actives sur l'adresse virtuelle. Ce nombre est stocké dans la ressource nommée <code>connections.<virtual addr value></code> (par exemple : <code>connections.192.168.1.99</code>) qui est affectée toutes les 10 secondes. Cette valeur est fournie à un titre indicatif uniquement. Valeur par défaut : <code>off</code>
netmask="defaultnetmask"	Linux et IPV4 seulement Par défaut, prend le netmask de l'interface. A positionner si l'interface a plusieurs netmasks.
</real_interface>	

13.5.6 <loadbalancing_list>, <group>, <cluster>, <host> Attributs

Pour des exemples de load balancing, voir section 15.5 [page 306](#).

A utiliser avec un module ferme

<loadbalancing_list>	
<group	Définition d'un groupe de load balancing. Mettre autant de sections qu'il y a de groupes : un exemple est donné en 15.5.3 page 308 .
name="group_name"	Nom du groupe de load balancing.
<cluster	Définition des serveurs et des poids. Sans la section <cluster>, les règles s'appliquent sur tous les serveurs de la ferme.
<host	Définition d'un nœud dans le groupe
name = "node_name"	Nom du nœud utilisé. <code>node_name</code> doit être le nom d'un serveur défini dans la configuration du cluster SafeKit (voir section 12 page 231).

<code>power = "value"</code>	Poids du nœud dans le groupe. Peut-être égal à 0 si l'on ne veut aucun trafic sur le nœud. Pour plus d'informations, voir 13.5.7.4 page 252 .
<code></cluster></code>	
<code><rule</code>	Définition d'une règle de load balancing dans le groupe. Autant de lignes que de règles de load balancing.
<code>[virtual_addr= "*" "virtual_IP_address" "virtual_IP_name"]</code>	Adresses IP virtuelles concernées par le load balancing. Par défaut toutes : *
<code>[port="*" "value"]</code>	Port TCP ou UDP sur lequel s'applique la règle de load balancing Par défaut tous les ports : *
<code>proto="udp" "tcp" "arp"</code>	<code>proto="udp"</code> : règle de load balancing UDP. <code>proto="tcp"</code> : règle de load balancing TCP. <code>proto="arp"</code> : règle de load balancing pour le protocole de résolution IP<->MAC.
<code>filter="on_addr" "on_port" "on_ipid"</code>	<code>filter="on_addr"</code> La règle de load balancing est réalisée sur l'adresse IP client en entrée. <code>filter="on_port"</code> La règle de load balancing est réalisée sur le port client en entrée. Voir l'exemple en 15.5.1 page 306 . <code>filter="on_ipid"</code> La règle de load balancing est réalisée sur l'ip_id en entrée. Utile pour UDP seulement (voir l'exemple en 15.5.2 page 307).

13.5.7 <vip> Description

13.5.7.1 <vip> prérequis

Voir les prérequis réseau décrits en 2.3.2 [page 30](#).

13.5.7.2 Qu'est-ce que le type "vmac_invisible" ?

La configuration `type="vmac_invisible"` associe une adresse MAC virtuelle à l'adresse IP virtuelle. Avec une adresse MAC virtuelle, les paquets émis vers l'adresse IP virtuelle sont reçus par tous les serveurs. Dans un module noyau, chaque serveur décode le paquet réseau et l'accepte ou le rejette. Après une sélection à très bas niveau des paquets réseau, l'application sur le serveur gère seulement le trafic réseau sélectionné par le module noyau.

Le mécanisme d'adresse MAC virtuelle consiste à associer une adresse MAC unicast dite virtuelle à l'adresse IP virtuelle. Quand un routeur ou une machine réseau recherche

l'adresse IP virtuelle, les serveurs SafeKit répondent avec l'adresse MAC virtuelle (via le protocole standard). Cependant, chaque serveur utilise son adresse MAC physique pour communiquer. Ainsi, l'adresse MAC virtuelle est invisible et non localisable par les switchs Ethernet. Par défaut, les switchs émettent ce type de paquet sur tous leurs ports (flooding), ceux-ci sont alors reçus par tous les serveurs de la ferme.

Avec la technologie d'adresse MAC virtuelle, le reroutage en cas de panne et de reprise est immédiat. Tous les équipements conservent l'association adresse IP virtuelle, adresse MAC virtuelle dans leur cache ARP.

Pour tester une adresse MAC virtuelle sur votre réseau, faire d'abord le test de compatibilité décrit en 4.3.7 [page 88](#).

13.5.7.3 Qu'est-ce que le type "vmac_directed" ?

La configuration `type= "vmac_directed"` modifie le fonctionnement du filtre. Dans ce mode, il n'y a pas de MAC virtuelle ; vu de l'extérieur, l'adresse IP virtuelle se comporte comme une adresse IP normale du point de vue de la résolution IP<->MAC.

Le module noyau est chargé de filtrer et transmettre les paquets entrant au serveur désigné par l'algorithme de partage de charge.

Le mode "vmac_directed" introduit un délai pour les clients ayant résolu l'adresse IP virtuelle sur l'adresse MAC d'un serveur qui est devenu indisponible. Ceci est comparable à ce qui se passe dans le cas `<real_interface>`. Les autres clients ne sont pas affectés.

Pour minimiser ce délai en IPV4, positionner `arpreroute="on"` sur l'interface correspondante, et régler les paramètres `arpelapse` et `arpinterval`.

Ipv6 possède un mécanisme interne et ne nécessite pas de configuration particulière.

13.5.7.4 Comment fonctionne le load balancing ?

Dans un module kernel, l'algorithme de load balancing est réalisé par filtrage sur les l'identité des paquets en réception. Cette identité est définie par configuration dans `userconfig.xml` : adresse IP client, port client ... (i.e. : load balancing de niveau 4). L'identité est passée dans une table de hachage (une bitmap de 256 bits) qui indique si le paquet doit être accepté ou rejeté sur le serveur. Un seul filtre accepte le paquet dans la ferme de serveurs. Quand un serveur est défaillant, le protocole membership reconfigure les filtres pour redistribuer le trafic du serveur défaillant sur les serveurs disponibles.

Chaque serveur peut avoir un poids (=1, 2...) et prendre plus ou moins de trafic. Le poids est mis en œuvre par le nombre bits à 1 dans la table de hachage (la bitmap de 256 bits).

Un exemple de bitmap est donné en 4.3.5 [page 86](#).

13.6 Réplication de fichiers (<rfs>, <replicated> tags)

S'applique à un module miroir seulement.

Sur Linux, vous devez définir la même valeur pour les uid/gid sur les deux nœuds pour la réplication des permissions sur les fichiers. Lors de la réplication d'un point de montage du système de fichiers, vous devez appliquer une procédure spéciale décrite en 13.6.4.2 [page 261](#).

Sur Windows, il est vivement recommandé d'activer le journal USN sur le lecteur contenant le répertoire répliqué, comme décrit en 13.6.4.3 [page 263](#).



Si vous exécutez plusieurs modules simultanément, les répertoires répliqués doivent être différents pour chaque module.

13.6.1 <rfs> Exemple

Exemple en Windows :

```
<rfs async="second" >
  <replicated dir="c:\safedir" mode="read_only"/>
</rfs>
```

Exemple en Linux :

```
<rfs async="second" >
  <replicated dir="/safedir" mode="read_only"/>
</rfs>
```

Voir l'exemple de flux de réplication dédié décrit en 15.4 [page 306](#).

13.6.2 <rfs> Syntaxe

```
<rfs
  [acl="on"|"off"]
  [async="second"|"none"]

  [iotimeout="nb seconds"]
  [roflags="0x10"|"0x10000"]
  [locktimeout="100"]
  [sendtimeout="30"]

  [nbrei="3"]
  [ruzone_blocksize="8388608"]
  [namespacepolicy="0"|"1"|"3"|"4"]
  [reitimeout="150"]
  [reicommith="0"]
  [reidetail="on"|"off"]
  [allocthreshold="0"]
  [nbremconn ="1"]

  [checktime="220000"]
  [checkintv="120"]

  [nfsbox_options="cross"|"nocross"]
  [scripts="off"]
  [reiallowdbw="20000"]
  [syncdelta="nb minutes"]
  [syncat="planification de la synchronisation"]
>

[<flow name="network" >
  [<!-- syntaxe pour SafeKit < 7.2 -->
    <server addr="IP_address_1" />
    <server addr="IP_address_2" />
  ]
</flow>]

<replicated dir="absolute path of a directory"
```

```
[mode="read_only"]>
<tocheck path="relative path of a file or subdir" />
<notreplicated path="relative path of a file or subdir" />
<notreplicated regxpath="regular expression on relative path of a file or
subdir" />
...
</replicated>
</rfs>
```



Seuls les attributs `async`, `nbrei`, `reitimeout` et `reidetail` du tag `<rfs>` peuvent être modifiés par une configuration dynamique. Le tag `<flow>`, qui décrit le flux de réplication, peut également être changé dynamiquement.

13.6.3 <rfs>, <replicated> Attributs

<rfs	
[mountoversuffix = "suffix"]	Sur Linux uniquement. À la configuration du module miroir, le répertoire répliqué <code>/a/dir</code> est renommé en <code>/a/dirsuffix</code>. Le répertoire /a/dir est créé et c'est : ⇒ un point de montage vers <code>/a/dirsuffix</code> lorsque le module est démarré ⇒ un lien vers <code>/a/dirsuffix</code> lorsque le module est arrêté Par défaut le suffix est « <code>_For_SafeKit_Replication</code> » S'il y a une défaillance matérielle, le lien symbolique n'est pas restauré. Dans ce cas, vous devez le restaurer manuellement. Restriction Vous ne pouvez pas spécifier directement une racine de système de fichiers comme répertoire répliqué (car le renommage du répertoire racine ne fonctionne pas). Le contournement consiste à une manipulation du fichier <code>fstab</code> tel qu'expliquée dans un KB sur https://support.evidian.com. Lorsque le module est démarré, NE PAS ACCEDER les fichiers dans <code>/a/dirsuffix</code>, car les modifications ne seront pas répliquées et le système deviendra incohérent. TOUJOURS ACCEDER les fichiers via <code>/a/dir</code>.

[acl= "on" "off"]	Active la réplication des ACLs sur les fichiers et répertoires. Valeur par défaut : <code>off</code>  Restrictions sur Windows La réplication des ACLs ne fonctionnera pas si le compte SYSTEM n'a pas les droits "Full control" sur toute la forêt répliquée. Le service <code>safeadmin</code> s'exécute dans le compte SYSTEM. Les ACLs sont répliqués littéralement (SID), sans translation, donc les ACLs "local users/groups" ne sont pas utilisables sur le serveur distant. Le cryptage et la compression des fichiers ne sont pas supportés.
[async= "second" "none"]	Positionner <code>async="second"</code> améliore les performances de la réplication de fichiers : les opérations d'écriture répliquées sont mises en cache sur le serveur secondaire et les acquittements sont envoyés plus rapidement au serveur primaire. Positionner <code>async="none"</code> assure plus de sécurité : les opérations d'écriture répliquées sont mises sur disque avant d'envoyer les acquittements au serveur primaire. Avec <code>async="second"</code>, en cas de double panne simultanée des 2 serveurs <code>PRIM</code> et <code>SECOND</code>, si le serveur <code>PRIM</code> ne peut pas redémarrer, alors le serveur <code>SECOND</code> n'a pas les données à jour sur son disque. Il y a perte de données si on force le serveur <code>SECOND</code> à redémarrer en primaire avec la commande <code>prim</code>. Valeur par défaut : <code>second</code>  La valeur de cet attribut peut être modifiée dynamiquement.
[packetsize]	Linux seulement. Taille maximale en octet des paquets de réplication NFS. Elle doit être inférieure ou égale à la taille maximale des paquets supportée par le serveur NFS des 2 serveurs. Quand cet attribut est affecté dans la configuration, il est utilisé pour affecter <code>rsize</code> et <code>wsize</code> au montage NFS. Par défaut, la taille est celle du serveur NFS.
[reipacketsize= "8388608"]	Taille maximale en octets des paquets de réintégration. En Linux, cette taille doit être inférieure ou égale à <code>packetsize</code>. Valeur par défaut en Linux : valeur de <code>packetsize</code> si elle est affectée dans la configuration et est <code>< 8388608</code>; sinon <code>8388608</code> Valeur par défaut en Windows : <code>8388608</code> octets

[ruzone_blocksize="8388608"]	Taille en octet d'une zone pour la bitmap de modification d'un fichier. Ça doit être un multiple de l'attribut <code>reipacketsize</code>. Valeur par défaut : valeur de <code>reipacketsize</code> si elle est affectée dans la configuration ; sinon 8388608
[iotimeout]	Windows seulement. Timeout en seconde sur les IO gérées dans le filtre file system Windows. Si une IO ne peut pas être répliquée et si le timeout du filtre expire, alors le serveur <code>PRIM</code> devient <code>ALONE</code>. Si non positionné, la valeur par défaut est calculée dynamiquement.
[roflags="0x10" "0x10000"]	Windows seulement. Pour garantir la cohérence des données répliquées sur les 2 serveurs, la modification des répertoires/fichiers répliqués ne doit avoir lieu que sur le serveur <code>PRIM</code>. Si des modifications ont lieu sur le serveur <code>SECOND</code>, celles-ci sont notifiées dans le journal du module avec l'identification du processus responsable afin que l'administrateur puisse corriger cette anomalie. C'est le comportement avec <code>roflags="0x10"</code>. Depuis SafeKit 7.4.0.31, le module peut en plus être arrêté sur le serveur <code>SECOND</code> en définissant <code>roflags="0x10000"</code>. Valeur par défaut : 0x10
[locktimeout="100"]	Timeout en secondes des requêtes répliquées. Si une requête ne peut pas être traitée dans cet intervalle de temps, le serveur <code>PRIM</code> devient <code>ALONE</code> Valeur par défaut : 100 secondes
[sendtimeout="30"]	Depuis SafeKit> 7.4.0.5 Timeout en secondes pour l'envoi de paquets TCP au nœud distant. Si l'envoi du paquet n'a pas pu être effectué dans le délai imparti, le serveur <code>PRIM</code> devient <code>ALONE</code>. Augmentez cette valeur en cas de réseau lent. Valeur par défaut : 30 secondes  Note Dans SafeKit 7.4.0.5, la valeur par défaut était de 120 secondes.
[nbrei="3"]	Nombre de threads de réintégration s'exécutant en parallèle pour resynchroniser les fichiers. Valeur par défaut : 3  Note La valeur de cet attribut peut être modifiée dynamiquement.
[namespacepolicy="0" "1" "3" "4"]	En Windows, avec l'option <code>namespacepolicy="1"</code>, la réintégration par zones ne peut être assurée après l'arrêt propre du module, puis reboot du serveur.

	Pour que ce cas soit supporté en Windows, il faut positionner <code>namespacepolicy="3"</code>. Cette option exploite l'USN journal du volume qui contient le répertoire répliqué (voir la commande <code>fsutil usn</code> pour la création du journal). Malgré cette option, la réintégration complète doit être appliquée lorsque : ⇒ l'USN journal associé au volume a été détruit/recréé par l'administrateur ⇒ une discontinuité dans le journal est détectée Lorsque la synchronisation par zones n'est pas possible (lors de la première réintégration ou lorsque les zones ne sont pas disponibles), les fichiers devant être synchronisés sont entièrement copiés. Si cette réintégration ne se termine pas, la suivante copiera à nouveau ces fichiers. Pour éviter cela, définissez <code>namespacepolicy = "4"</code>. Cette option active également la vérification de journal USN en Windows. Utiliser <code>namespacepolicy="0"</code> pour désactiver la synchronisation par zones sur Windows ou Linux. Valeur par défaut : 4 pour SafeKit > 7.4.0.5 (non supporté dans les versions antérieures)
<code>[reitimeout="150"]</code>	Timeout en seconde des requêtes de réintégration. Ce timeout peut être augmenté pour éviter les arrêts de réintégration à cause d'une machine primaire chargée. Valeur par défaut : 150 secondes  La valeur de cet attribut peut être modifiée dynamiquement.
<code>[reicommit="0"]</code>	Linux seulement. Positionner <code>reicommit="nb blocks"</code> pour commiter sur disque tous les <code>(nb blocks)*reipacketsize</code> lors de la réintégration d'un fichier (en plus du commit réalisé à la fin de la copie). Ceci peut aider la réintégration de gros fichiers mais ralentit le temps de réintégration global. Valeur par défaut : 0 signifie pas de commit intermédiaire.
<code>[reidetail="on" "off"]</code>	Journal détaillé de la réintégration. Valeur par défaut : <code>off</code>  La valeur de cet attribut peut être modifiée dynamiquement.
<code>[allocthreshold="0"]</code>	Windows seulement. Taille en Go pour appliquer la politique d'allocation avant réintégration. Quand <code>allocthreshold > 0</code>, activation de l'allocation rapide de l'espace disque pour les fichiers à réintégrer sur le nœud

	secondaire. Cette fonctionnalité permet, quand le fichier est très gros (> 200 Go) et pas encore complètement recopié, d'éviter le timeout de l'écriture du primaire en fin de fichier. Depuis SafeKit 7.4.0.64, la politique d'allocation a changé et est appliquée : ⇒ pour les nouveaux fichiers (fichiers n'existant pas sur la secondaire quand la réintégration commence) ⇒ pour une synchronisation de type <code>full</code> (par exemple, lors de la 1ère réintégration ou quand le secondaire est démarré avec <code>safekit second fullsync</code>) ⇒ quand la taille du fichier sur la primaire est <code>>= allocthreshold</code> (taille en Go) Valeur par défaut : 0 (qui désactive la fonctionnalité)
<code>[nbremconn="1"]</code>	Nombre de connexions TCP entre les nœuds primaire et secondaire. Cette valeur peut être augmentée pour améliorer le débit de réplication et de synchronisation lorsque le réseau présente une latence élevée (dans le cloud, par exemple). Valeur par défaut : 1
<code>[checktime="220000"]</code>	Linux seulement. Timeout en millisecondes pour une requête null qui vérifie le bon fonctionnement local de la réplication de fichiers. Commande <code>stopstart</code> si le timeout est atteint. Valeur par défaut : 220000
<code>[checkintv="120"]</code>	Linux seulement. Intervalle en secondes entre 2 requêtes null. Valeur par défaut : 120 secondes
<code>nfsbox_options="cross" "nocross"</code>	Windows seulement. Cette option spécifie la politique à appliquer lorsqu'un reparse point du type <code>MOUNT_POINT</code> est présent dans l'arborescence répliquée. Cette politique est globale à tous les répertoires répliqués. Dans NTFS, les reparse point de type <code>MOUNT_POINT</code> représentent : ✓ soit un point de montage NTFS (par exemple <code>D:\directory</code>) ✓ soit un "directory junction" NTFS (en quelque sorte un lien symbolique vers une autre partie du système de fichiers) Avec l'option <code>nfsbox_options="cross"</code>, les points de montages sont évalués avant réplication et le contenu de la cible du point de montage est répliqué, ce qui rend le point de montage équivalent à un répertoire normal. Ce comportement est utile lorsque le répertoire répliqué est la racine d'un système de fichier (par exemple <code>D:\</code>). C'est le comportement par défaut. Lorsque <code>nfsbox_options="nocross"</code>, les points de montages ne sont pas évalués et sont répliqués en tant que point de montage

	(fichier de type « reparse point »). Le contenu de la cible du point de montage n'est pas répliqué ou réintégré lorsque le point de montage est sollicité. Ce comportement est utile lorsque la cible du point de montage est située dans un autre répertoire répliqué (fichier de type « junctions »). Par exemple, les bases de données PostgreSQL utilisent des fichiers de ce type. Valeur par défaut : <code>cross</code>
<pre>[scripts= "on" "off"]</pre>	<code>scripts="on"</code> active les callback vers les scripts <code>_rfs_*</code> utilisés pour mettre en œuvre une réplication de données externe (voir le module Linux DRBD.safe pour plus d'information) Valeur par défaut : <code>off</code>
<pre>[reiallowedbw= "20000"]</pre>	Quand cet attribut est défini, il spécifie la bande passante maximum susceptible d'être utilisée par la phase de réintégration (par exemple 20000 KB/s), en kilo octets par secondes (KB/s). Etant donné l'implémentation retenue, une fluctuation de +/- 10% de la bande passante réellement utilisée est observable.  Note La bande passante utilisée par la réplication n'est pas affectée par ce paramètre. Par défaut : l'attribut n'est pas défini et la bande passante utilisée par la réintégration n'est pas limitée
<pre>[syncdelta="nb minutes"]</pre>	Quand cet attribut est ≤ 1, l'attribut est ignoré et la politique par défaut de démarrage et de reprise sur panne est appliquée : seul le serveur avec les données à jour peut démarrer en primaire ou effectuer une reprise sur panne. Quand cet attribut est > 1, la politique par défaut de démarrage et de reprise sur panne est modifiée. Le serveur avec des données non à jour peut devenir primaire mais uniquement si le temps écoulé depuis sa dernière synchronisation est inférieur à la valeur de <code>syncdelta</code> (en minutes). Valeur par défaut : 0 minutes
<pre>[syncat="planific ation de la synchronisation"]</pre>	Valeur par défaut : réplication temps réel et synchronisation automatique (pas de planification) Utiliser l'attribut <code>syncat</code> pour planifier la synchronisation des répertoires répliqués sur le nœud secondaire à des dates/heures données. Pour plus de détails, voir 13.6.4.10 page 270. Le module doit être démarré pour activer cette fonctionnalité. Une fois synchronisé, le module se bloque dans l'état <code>WAIT</code> (rouge) jusqu'à la prochaine synchronisation. La planification est basée sur le gestionnaire de tâches du système d'exploitation : ⇒ en Windows, la tâche est définie comme tâche système ⇒ en Linux, la tâche est insérée dans la <code>crontab</code> de l'utilisateur <code>safekit</code> Vous devez simplement configurer <code>syncat</code> avec la syntaxe du gestionnaire de tâche du système. Par exemple, pour une synchronisation quotidienne, après minuit :

	⇒ en Windows <pre>syncat="/SC DAILY /ST 00:01:00"</pre> ⇒ en Linux <pre>syncat="01 0 * * *"</pre>  Voir la documentation de <code>crontab</code> en Unix et de <code>schtasks.exe</code> en Windows, pour une description complète de la syntaxe  Si la configuration ou la planification ne fonctionnent pas correctement, vérifier d'abord les erreurs de syntaxe ou d'utilisation du gestionnaire de tâches du système.
<pre>[<flow name="network" > <server addr="IP_1" /> <server addr="IP_2" /> </flow>]</pre>	Ancienne configuration préservée pour la compatibilité ascendante. Quand cette section n'est pas définie, le flux de réplication passe par le heartbeat défini avec <code>ident="flow"</code> s'il y en a un, sinon par la voie du 1^{er} heartbeat (pour la description des heartbeats, voir section 13.3 page 241). Si vous utilisez cette configuration, il faut assurer la cohérence avec la définition d'un heartbeat avec <code>ident=flow</code> car des règles de failover par défaut sont définies (décrites en 13.18.5 page 291).  Le sous-arbre <code><flow></code> peut être modifié dynamiquement, pour changer le flux de réplication par exemple. L'attribut <code>name</code> de <code><flow></code> nommé le réseau utilisé pour le flux de réplication. <code>network</code> doit être le nom d'un réseau défini dans la configuration du cluster global (voir section 12 page 231). Le tag <code><server></code> était utilisé dans l'ancienne syntaxe de configuration (avant SafeKit 7.2). Il est supporté pour assurer la compatibilité ascendante, mais ne doit pas être utilisé pour la configuration de nouveaux modules.  Vous ne devez pas utiliser dans le même <code>userconfig.xml</code>, la syntaxe de SafeKit 7.1 et celle introduite depuis SafeKit 7.2.
<code><replicated</code>	Définition des répertoires répliqués Mettre autant de sections que de répertoires à répliquer
<code>dir="/abs_path"</code>	Path absolu du répertoire à répliquer.
<code>[mode="read_only"]</code>	Accès en read-only sur la machine secondaire pour éviter la corruption.
<code><notreplicated path="relative" /></code>	Path relatif d'un fichier ou d'un sous répertoire d'un répertoire répliqué. Le fichier (ou sous répertoire) est non répliqué. Autant de lignes que de fichiers ou sous répertoire à non répliquer.
<code><notreplicated regexpath="expre</code>	Linux seulement.

<pre>ssion régulière" /></pre>	Expression régulière pour définir les fichiers ou sous répertoires non répliqués. Exemple (pour plus d'information, voir "man regex"): <pre><replicated dir="/safedir"> <notreplicated regexp=".*\.tmp" /> </replicated></pre> Dans cet exemple, /safedir/conf/config.tmp et /safedir/log.tmp sont non répliqués alors que /safedir/conf/config.tmp.bak est répliqué.
<pre><tocheck path="relative" /></pre>	Path relatif d'un fichier ou d'un sous répertoire dans un répertoire répliqué. Vérifier sa présence avant de démarrer. Evite un démarrage sur un répertoire vide. Mettre autant de lignes que nécessaire.

13.6.4 <rfs>Description

13.6.4.1 <rfs> prérequis

Voir les prérequis décrits en 2.2.4 [page 29](#).

En Windows, activez le journal USN sur le lecteur qui contient les répertoires répliqués afin d'activer la réintégration par zones, après redémarrage du serveur, à condition que le module ait été arrêté proprement.

13.6.4.2 <rfs> Linux

Sur Linux, l'interception des données répliquées est basée sur un montage NFS local. Et le flux de réplication entre les serveurs est basé sur le protocole NFS v3 / TCP.

Le montage NFS des répertoires répliqués à partir de clients Linux externe n'est pas supporté. En revanche, le montage d'autres répertoires peut être réalisé avec les commandes standards.

Procédure pour répliquer un point de montage

Quand un répertoire répliqué est un point de montage, la configuration du module échoue avec l'erreur suivante :

Erreur: périphérique ou ressource occupée

Dans la suite, nous prenons l'exemple du module PostgreSQL qui définit en tant que répertoires répliqués /var/lib/pgsql/var et /var/lib/pgsql/data. Le fichier userconfig.xml du module contient :

```
<rfs ... >
  <replicated dir="/var/lib/pgsql/var" mode="read_only" />
  <replicated dir="/var/lib/pgsql/data" mode="read_only" />
</rfs>
```

Ces répertoires sont des points de montage comme le montre le résultat de la commande `df -H`. La commande retourne par exemple :

```
/dev/mapper/vg01-lv_pgs_var ... /var/lib/pgsql/var
/dev/mapper/vg02-lv_pgs_data ... /var/lib/pgsql/data
```

Vous devez appliquer la procédure suivante pour configurer le module avec la réplication de ces répertoires.



C'est la même procédure pour tous les points de montage qui doivent être répliqués.

- ⇒ démonter les systèmes de fichiers en exécutant :

```
umount /var/lib/pgsql/var
umount /var/lib/pgsql/data
```

- ⇒ configurer le module en exécutant :

```
/opt/safekit/safekit config -m postgresql
```

La configuration se termine avec succès.

- ⇒ vérifier l'existence des liens symboliques créés lors de la configuration en exécutant `ls -l /var/lib`. La commande retourne :

```
lrwxrwxrwx 1 root root var -> var_For_SafeKit_Replication
lrwxrwxrwx 1 root root data -> data_For_SafeKit_Replication
```

- ⇒ éditer `/etc/fstab` et modifier les 2 lignes :

```
/dev/mapper/vg01-lv_pgs_var /var/lib/pgsql/var ext4...
/dev/mapper/vg02-lv_pgs_data /var/lib/pgsql/data ext4...
```

par

```
/dev/mapper/vg01-lv_pgs_var
/var/lib/pgsql/var_For_SafeKit_Replication ext4...

/dev/mapper/vg02-lv_pgs_data
/var/lib/pgsql/data_For_SafeKit_Replication ext4..
```

- ⇒ monter les systèmes de fichiers en exécutant :

```
mount /var/lib/pgsql/var_For_SafeKit_Replication
mount /var/lib/pgsql/data_For_SafeKit_Replication
```



Appliquez cette procédure sur les deux nœuds si les répertoires répliqués sont des points de montage sur les deux nœuds. Une fois appliquée, vous pouvez utiliser le module comme d'habitude : par exemple `safekit start stop` etc ...



Pour empêcher le démarrage du module quand le répertoire est non monté et vide, vous pouvez insérer dans `userconfig.xml` la vérification de la présence d'un fichier dans le répertoire répliqué. Exemple pour `/var/lib/pgsql/var` (faire de même pour `/var/lib/pgsql/data` en testant un fichier toujours présent dans ce répertoire) :

```
<replicated dir="/var/lib/pgsql/var" mode="read_only">
  <tocheck path="postgresql.conf" />
</replicated>
```

A la déconfiguration du module (ou désinstallation du package SafeKit), vous devez appliquer la procédure inverse pour restaurer l'état initial :

- ⇒ démonter les systèmes de fichiers en exécutant :

```
umount /var/lib/pgsql/var_For_SafeKit_Replication
```

```
umount /var/lib/pgsql/data_For_SafeKit_Replication
```

⇒ déconfigurer le module en exécutant `/opt/safekit/safekit deconfig -m postgresql`

⇒ éditer `/etc/fstab` pour y restaurer son état initial

⇒ monter les systèmes de fichiers en exécutant :

```
mount /var/lib/pgsql/var
```

```
mount /var/lib/pgsql/data
```

13.6.4.3 <rfs> Windows

Sur Windows, l'interception des données est basée sur un filtre file system. Et le flux de réplication entre les serveurs est basé sur le protocole NFS v3 / TCP.

Certains anti-virus peuvent empêcher le fonctionnement correct de la réplication.

Sur Windows, il est possible de monter à distance un répertoire répliqué. Si vous voulez pouvoir monter avec le nom et non l'adresse IP virtuelle, vous devez positionner les valeurs suivantes dans les bases de registre des deux serveurs SafeKit :

```
[HKEY_LOCAL_MACHINE\SYSTEM\CurrentControlSet\Control\Lsa]
```

```
"DisableLoopbackCheck"=dword:00000001
```

```
[HKEY_LOCAL_MACHINE\SYSTEM\CurrentControlSet\Services\lanmanserver\parameters] "DisableStrictNameChecking"=dword:00000001
```

En Windows, pour activer la réintégration par zones après le redémarrage du serveur, lorsque le module a été correctement arrêté, le composant <rfs> utilise le journal NTFS USN pour vérifier que les informations enregistrées sur les zones sont toujours valables après le redémarrage. Lorsque le contrôle réussit, la réintégration par zones peut être appliquée sur le fichier ; sinon, le fichier doit être recopié dans sa totalité.

Par défaut, seul le lecteur système a un journal USN actif. Si les répertoires répliqués sont situés sur un lecteur différent du lecteur système, vous devez créer le journal (avec commande `fsutil usn`). Voir [SK-0066](#) pour un exemple.

13.6.4.4 <rfs> Réplication et reprise sur panne

Avec la réplication de fichiers, l'architecture miroir est particulièrement adaptée à la haute disponibilité des applications base de données avec des données critiques à protéger contre les pannes. En effet, les données du serveur secondaire sont fortement synchronisées avec celles du serveur primaire. Le serveur est dit à jour et seul un serveur à jour peut démarrer en primaire ou effectuer une reprise sur panne

Si la disponibilité de l'application est plus critique que la synchronisation des données, la politique par défaut peut être relâchée pour autoriser un serveur non à jour à devenir primaire mais uniquement si la date de la dernière synchronisation est inférieure à un délai configurable. Cela est configuré avec l'attribut `syncdelta` du tag <rfs> dont la valeur est exprimée en minutes :

⇒ `syncdelta <= 1`

L'attribut est ignoré et la politique par défaut de démarrage en primaire et de reprise sur panne est appliquée. La valeur par défaut 0.

⇒ `syncdelta > 1`

Si le serveur à jour ne répond pas, le serveur non à jour peut devenir primaire mais uniquement si le temps écoulé depuis la dernière synchronisation est inférieur à la valeur de `syncdelta` (en minutes).

Cette fonctionnalité est implémentée à l'aide de :

⇒ la ressource `rfs.synced`

Quand `syncdelta` est > 1 , la gestion de la ressource `rfs.synced` est activée. Cette ressource est dans l'état `UP` si les données répliquées sont cohérentes et le temps écoulé depuis la dernière synchronisation est inférieur à la valeur de `syncdelta`.

⇒ Le checker `syncedcheck`

Quand `syncdelta` est > 1 , ce checker est activé. Il affecte la valeur de la ressource `rfs.synced`.

⇒ La règle de failover `rfs_forceuptodate`

Quand `syncdelta` est > 1 , la règle de failover suivante est valide :

```
rfs_forceuptodate:      if (heartbeat.* == down && cluster() == down &&  
rfs.synced == up && rfs.uptodate == down) then rfs.uptodate=up;
```

Cette règle provoque le démarrage en primaire du serveur lorsque le serveur à jour ne répond pas, et à condition que ce serveur soit isolé et considéré synchronisé en fonction de la valeur de `syncdelta`.

13.6.4.5 <rfs> Vérification de la réplication

Vous pouvez vérifier que les fichiers sont identiques sur le primaire et le secondaire avec la commande suivante à passer sur la machine `SECOND` : `safekit rfsverify -m AM`. Exécuter `safekit rfsverify -m AM > log` pour rediriger la sortie de la commande dans un fichier nommé `log`.

La sortie de la commande est un journal similaire à celui de la réintégration dans lequel sont indiqués les fichiers à recopier (donc différents).

Quand sur la primaire, il y a de l'activité sur les répertoires répliqués, il se peut qu'une anomalie soit détectée alors qu'il n'y a pas de différence entre les fichiers. Cela se produit dans les cas suivants :

- ⇒ sur Windows à cause des modifications faites sur disque avant d'être répliquées,
- ⇒ avec `async="second"` (défaut) car les lectures peuvent dépasser les écritures.

Pour vérifier s'il y a vraiment une incohérence, vous devez relancer la commande sur le serveur secondaire en s'assurant qu'il n'y plus d'activité sur le serveur primaire.

Sur Windows, certains fichiers modifiés avec l'option `SetvalidData` sont systématiquement détectés différents car ils sont étendus sans reset des données : le contenu en lecture des zones étendues est le contenu aléatoire du disque au moment de la lecture.



Il est fortement recommandé d'exécuter cette commande uniquement lorsqu'il n'y a pas d'accès aux répertoires répliqués sur le primaire.

13.6.4.6 <rfs> Fichiers modifiés depuis la dernière synchronisation

Avant de démarrer le serveur secondaire, il peut être utile d'évaluer le nombre de fichiers et la quantité de données qui ont été modifiés sur le serveur primaire depuis l'arrêt du serveur secondaire. Cette fonctionnalité est fournie en exécutant la commande suivante sur le serveur ALONE : `safekit rfsdiff -m AM`. Exécuter `safekit rfsdiff -m AM > log` pour rediriger la sortie de la commande dans un fichier nommé `log`.

Cette commande exécute des vérifications en ligne du contenu des fichiers réguliers du module AM. Elle analyse l'arborescence répliquée entièrement et affiche le nombre de fichiers qui ont été modifiés ainsi que la taille qui doit être recopiée. Elle affiche également une estimation du temps total de réintégration. Ceci n'est qu'une évaluation car seuls les fichiers réguliers sont analysés et d'autres modifications peuvent se produire jusqu'à ce que la synchronisation soit exécutée par le serveur secondaire.

Cette commande doit être utilisée avec précaution sur un serveur en production car elle entraîne une surcharge sur le serveur (pour la lecture, avec verrouillage, de l'arborescence et des fichiers). En Windows, le renommage des fichiers peut échouer pendant cette évaluation.



Il est fortement recommandé d'exécuter cette commande uniquement lorsqu'il n'y a pas d'accès aux répertoires répliqués.

13.6.4.7 <rfs> Bande passante de réplication et de réintégration

Le composant de réplication collecte sur le serveur PRIM la bande passante utilisée par les opérations d'écritures de réplication et de réintégration.

Deux ressources (`rfs_bandwidth.replication` et `rfs_bandwidth.reintegration`), exprimées en kilo octet par seconde (KB/s), reflètent la bande passante moyenne utilisée respectivement par la réplication et la réintégration durant les 3 dernières secondes.

Si la charge de réplication est très active en écriture, une saturation du lien réseau peut se produire lors de la phase de réintégration, entraînant un ralentissement significatif de l'application. Dans ce cas, l'attribut `<rfs> reallowedbw` peut être utilisé pour limiter la bande passante utilisée par la phase de réintégration (voir section 13.6.3 page 254). Il faut cependant considérer que la limitation de la bande passante de réintégration allongera la durée de la phase de réintégration.

Depuis SafeKit 7.5, il y a 2 nouvelles ressources qui reflètent la bande passante réseau (en KOctets/sec), utilisée entre les processus `nfsbox`, qui s'exécutent sur chaque nœud pour implémenter la réplication et la réintégration :

- ⇒ `rfs.netout_bandwidth` est la bande passante utilisée en sortie
- ⇒ `rfs.netin_bandwidth` est la bande passante utilisée en entrée

Vous pouvez observer la valeur de `rfs.netout_bandwidth` sur le primaire ou de `rfs.netin_bandwidth` sur le secondaire pour connaître le taux de modification au moment de l'observation (écriture, création, suppression, ...). L'historique des valeurs de la ressource donne un aperçu de son évolution dans le temps.

La valeur de la bande passante dépend de l'activité applicative, système et réseau. Sa mesure n'est disponible qu'à titre d'information.

13.6.4.8 <rfs> Synchronisation par date

Depuis SafeKit 7.2, SafeKit offre la nouvelle commande `safekit secondforce -d date -m AM` qui force le module AM à démarrer comme secondaire après avoir copié uniquement les fichiers modifiés après la date spécifiée.



Cette commande doit être utilisée avec précautions car la synchronisation ne copiera pas les fichiers modifiés avant la date spécifiée. Il incombe à l'administrateur de s'assurer que ces fichiers sont cohérents et à jour.

La date est dans le format YYYY-MM-DD[Z] ou "YYYY-MM-DD hh:mm:ss[Z]" ou YYYY-MM-DDThh:mm:ss[Z], où :

- YYYY-MM-DD indique l'année, le mois et le jour
- hh:mm:ss indique l'heure, les minutes et secondes
- Z indique que la date est exprimée dans le fuseau horaire UTC ; s'il n'est pas spécifié, la date est exprimée dans le fuseau horaire local

Par exemple :

- `safekit secondforce -d 2016-03-01 -m AM` copie uniquement les fichiers modifiés après le 1er Mars 2016
- `safekit secondforce -d "2016-03-01 12:00:00" -m AM` copie uniquement les fichiers modifiés après le 1er Mars 2016 à 12h, heure locale
- `safekit secondforce -d 2016-03-01T12:00:00Z -m AM` copie uniquement les fichiers modifiés après le 1er Mars 2016 à 12h, dans le fuseau horaire UTC

Cette commande peut être utile dans le cas suivant :

- le module est arrêté sur le serveur primaire et une sauvegarde des données répliquées est effectuée (sur un lecteur amovible par exemple)
- le module est arrêté sur le serveur secondaire et les données répliquées sont restaurées à partir de la sauvegarde. Il peut s'agir du premier démarrage ou de la réparation du serveur secondaire.
- le module est démarré sur le serveur primaire qui devient `ALONE`
- le module est démarré sur le secondaire avec la commande `safekit secondforce -d date -m AM` où la date est la date de sauvegarde

Dans ce cas, seuls les fichiers modifiés depuis la date de la sauvegarde seront copiés (dans leur totalité), au lieu de la copie complète de tous les fichiers.



En Windows, la date de modification du fichier sur le serveur secondaire est affectée lorsque le fichier est copié par le processus de réintégration. Par conséquent, `safekit secondforce -d date -m AM`, dont la date est antérieure à la dernière réintégration sur ce serveur, n'a aucun intérêt.

13.6.4.9 <rfs> Synchronisation externe

Lors de la première synchronisation, tous les fichiers répliqués sont copiés dans leur totalité du nœud principal vers le nœud secondaire. Lors des synchronisations suivantes, nécessaires lors du redémarrage du nœud secondaire, seules les zones modifiées des fichiers sont recopiées. Lorsque les répertoires répliqués sont volumineux, la première

synchronisation peut prendre beaucoup de temps en particulier si le réseau est lent. C'est pourquoi, depuis SafeKit> 7.3.0.11, SafeKit fournit une nouvelle fonctionnalité pour synchroniser un grand volume de données qui doit être utilisée conjointement avec un outil de sauvegarde.

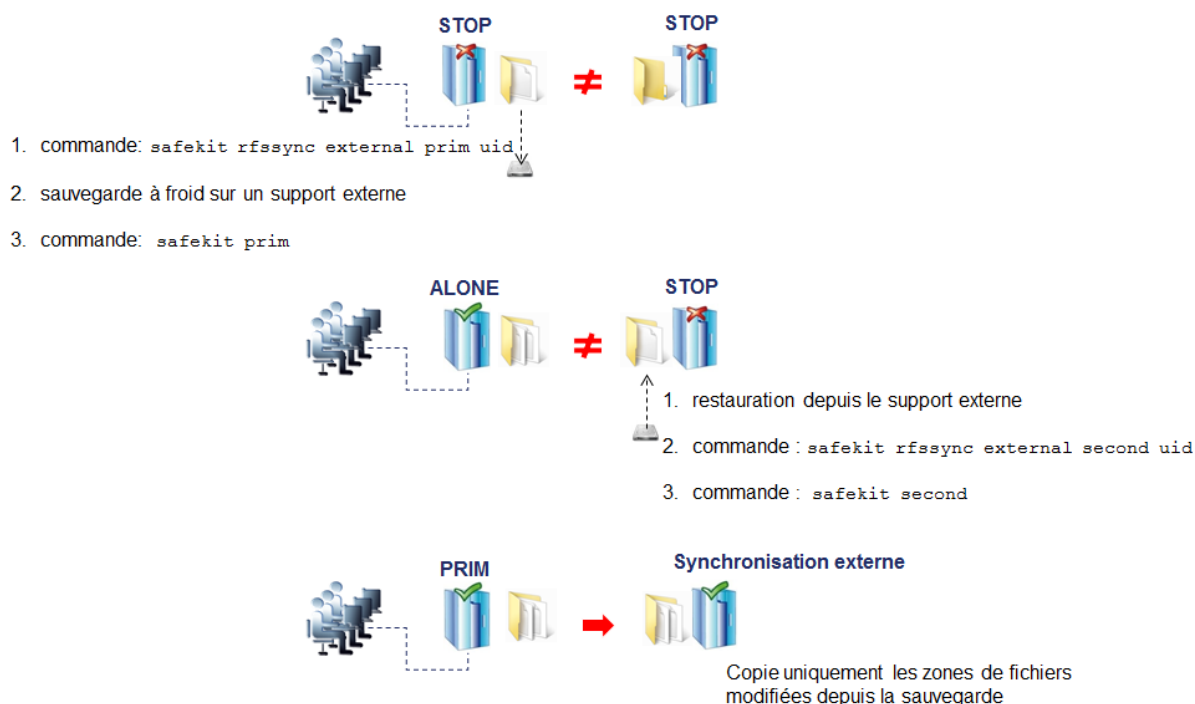
Sur le nœud principal, il suffit d'effectuer une sauvegarde des répertoires répliqués et de passer la politique synchronisation au mode externe. La sauvegarde est transportée (en utilisant un lecteur externe par exemple) et restaurée sur le nœud secondaire, qui est aussi configuré pour effectuer une synchronisation externe. Lorsque le module est démarré sur le nœud secondaire, il recopie uniquement les zones de fichiers modifiées sur le nœud principal depuis la sauvegarde.

La synchronisation externe repose sur une nouvelle commande `safekit rfssync` qui doit être appliquée sur les deux nœuds afin de positionner le mode de synchronisation à `external`. Cette commande prend comme arguments :

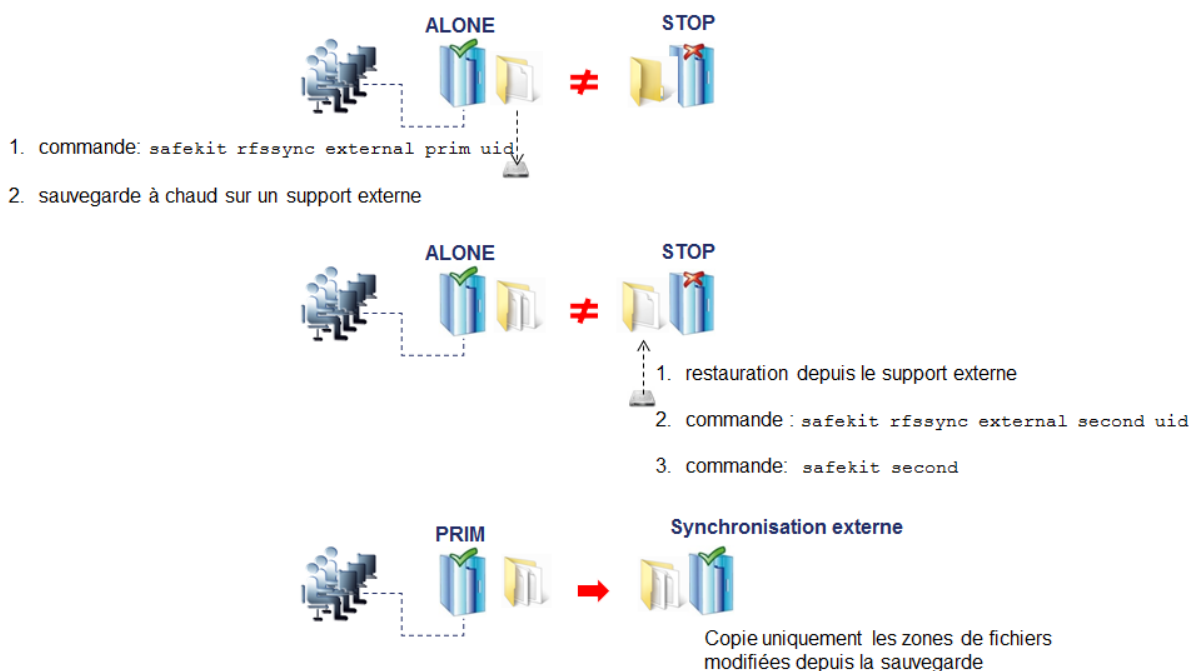
- le rôle du nœud (`prim` | `second`)
- un identificateur unique (`uid`)

Procédure de synchronisation externe

La procédure de synchronisation externe, décrite ci-dessous, est la procédure à appliquer dans le cas d'une sauvegarde à froid des répertoires répliqués. Dans ce cas, l'application doit être arrêtée et toute modification des répertoires répliqués est interdite jusqu'au démarrage du module, et de l'application, en `ALONE` – vert. L'ordre des opérations doit être strictement respecté.



La procédure de synchronisation externe, décrite ci-dessous, est la procédure à appliquer dans le cas d'une sauvegarde à chaud des répertoires répliqués. Dans ce cas, le module est `ALONE` – vert ; l'application est démarrée et les modifications du contenu des répertoires répliqués sont autorisées. L'ordre des opérations doit être strictement respecté.



Commande `safekit rfssync`

<pre>safekit rfssync external prim <uid> [-m AM]</pre>	Positionne la politique de synchronisation à <code>external</code>. Elle est identifiée par la valeur de <code>uid</code> (max 24 char). Le nœud est le primaire, source pour la synchronisation des données.
<pre>safekit rfssync external second <uid> [-m AM]</pre>	Positionne la politique de synchronisation à <code>external</code>. Elle est identifiée par la valeur de <code>uid</code> (max 24 char). Le nœud est le secondaire, destination de la synchronisation des données.
<pre>safekit rfssync -d prim <uid> [-m AM] safekit rfssync -d second <uid> [-m AM]</pre>	Désactive le contrôle des modifications des répertoires répliqués entre le moment de la sauvegarde à froid/la restauration et le démarrage du module.  Cette option doit être utilisée avec précautions puisque la synchronisation externe peut dans ce cas ne pas détecter toutes les modifications à recopier.
<pre>safekit rfssync full [-m AM]</pre>	Positionne la politique de synchronisation à <code>full</code>. Celle-ci entraîne la recopie de tous les fichiers dans leur totalité à la prochaine synchronisation.
<pre>safekit rfssync</pre>	Affiche la politique de synchronisation courante.

Internes

La politique de synchronisation est représentée par des ressources du module : `usersetting.rfssyncmode`, `usersetting.rfssyncrole`, `usersetting.rfssyncuid` et `rfs.rfssync` :

⇒ `usersetting.rfssyncmode="default"`
(`usersetting.rfssyncrole="default"`, `usersetting.rfssyncuid="default"`)

Ces valeurs sont associées à la politique de synchronisation standard, celle appliquée par défaut. Elle consiste à ne copier que les zones modifiées des fichiers. Quand cette stratégie ne peut être appliquée, les fichiers modifiés sont recopiés dans leur totalité.

⇒ `usersetting.rfssyncmode="full"`
(`usersetting.rfssyncrole="default"`, `usersetting.rfssyncuid="default"`)

Ces valeurs sont associées à la politique de synchronisation `full`. Elle est appliquée :

- au premier démarrage du module après sa première configuration
- sur commandes `safekit (safekit second fullsync ; safekit rfssync full ; safekit primforce ; safekit config ; safekit deconfig)`
- sur changement d'appariement du module

La politique de synchronisation `full` entraîne la recopie de tous les fichiers dans leur totalité à la prochaine synchronisation.

⇒ `usersetting.rfssyncmode="external"`, `usersetting.rfssyncrole="prim | second"` and `usersetting.rfssyncuid="uid"`

Ces valeurs sont associées à la politique de synchronisation `external` affectée avec les commandes `safekit rfssync external prim uid` ou `safekit rfssync external second uid`. La prochaine synchronisation appliquera la politique de synchronisation externe.

⇒ `rfs.rfssync="up | down"`

Cette ressource vaut `up` uniquement lorsque la politique de synchronisation, définie par les ressources précédentes, peut être appliquée.

Quand la politique de synchronisation n'est pas celle par défaut, celle-ci repasse automatiquement dans le mode par défaut une fois la synchronisation appliquée avec succès.

Dans certains cas, la synchronisation externe ne peut être appliquée et le nœud secondaire s'arrête avec une erreur indiquée dans le journal du module. Dans cette situation, il faut soit :

- ⇒ compléter la procédure de synchronisation externe si celle-ci n'a pas été effectuée dans sa totalité sur les 2 nœuds
- ⇒ réappliquer complètement la procédure de synchronisation sur les 2 nœuds
- ⇒ appliquer la politique de synchronisation `full` (commande `safekit rfssync full`)

- ⇒ appliquer la synchronisation par date, en utilisant la date de la sauvegarde (voir section 13.6.4.8 [page 266](#)). Contrairement à la synchronisation externe, la synchronisation par date va copier dans leur totalité (et non par zones) les fichiers modifiés sur le nœud primaire.

13.6.4.10 <rfs> Synchronisation planifiée

Par défaut, SafeKit offre la réplication de fichiers en temps réel et une synchronisation automatique. Si la charge est importante sur le nœud primaire ou si le réseau a une forte latence, il peut être préférable d'accepter que le nœud secondaire ne soit pas fortement synchronisé avec le nœud primaire. Pour cela, vous pouvez utiliser l'attribut `syncat` pour planifier une synchronisation régulière des répertoires répliqués sur le nœud secondaire. Le module doit être démarré pour activer cette fonctionnalité. Une fois synchronisé, le module se bloque dans l'état `WAIT` (rouge) jusqu'à la prochaine synchronisation planifiée. Cette fonctionnalité est implémentée avec :

- ⇒ la ressource `rfs.syncat` affectée à `up` aux dates planifiées et affectée à `down` une fois le nœud secondaire synchronisé
- ⇒ la règle de failover `rfs_syncat_wait` qui bloque le nœud secondaire dans l'état `WAIT` (rouge) jusqu'à ce que la ressource `rfs.syncat` soit `up`

Si vous souhaitez forcer la synchronisation en dehors des dates planifiées, il faut exécuter la commande `safekit set -r rfs.syncat -v up -m AM` quand le module est dans l'état `WAIT` (rouge).

La configuration de `syncat` se fait simplement en utilisant la syntaxe du gestionnaire de tâches du système d'exploitation : `crontab` en Linux et `schtasks.exe` en Windows (voir section 13.6.3 [page 254](#)).

13.7 Activer les scripts utilisateurs (<user>, <var> tags)

Cette section décrit uniquement les options de configuration du tag `<user>`. Pour une description complète des scripts, voir la section 14 [page 293](#).

13.7.1 <user> Exemple

```
<user logging="userlog" >
  <var name="VARENV" value="V1" />
</user>
```

Voir un exemple en 15.1 [page 302](#).

13.7.2 <user> Syntaxe

```
<user
  [nicestoptimeout="300"]
  [forcestoptimeout="300"]
  [logging="userlog"|"none"]
  [userlogsize="2048"]
>
  <var name="ENVIRONMENT_VARIABLE_1" value="VALUE_1" />
  ...
</user>
```



Le tag `<user>` et son sous-arbre peuvent être entièrement modifiés dynamiquement.

13.7.3 `<user>`, `<var>` Attributs

<code><user</code>	
<code>[nicestoptimeout="300"]</code>	Timeout en secondes pour exécuter les scripts <code>stop_xx</code> . Valeur par défaut : 300 secondes
<code>[forcestoptimeout="300"]</code>	Timeout en secondes pour exécuter les scripts <code>stop_xx -force</code> Valeur par défaut : 300 secondes
<code>[logging="userlog" "none"]</code>	<code>logging="userlog"</code> : les messages stdout et stderr de l'application démarrée dans les scripts redirigés dans le journal applicatif dans le fichier <code>SAFEVAR/modules/AM/userlog.ulong</code> où AM est le nom du module (SAFEVAR=C:\safekit\var sur Windows et /var/safekit sur LINUX). Depuis SafeKit 7.4.0.19, l'extension du nom de fichier du journal applicatif a changé. Le fichier s'appelle dorénavant <code>userlog.ulong</code> au lieu de <code>userlog.AM</code> <code>logging="none"</code> : les messages stdout et stderr de l'application démarrée dans les scripts ne sont pas logués Valeur par défaut : <code>userlog</code>
<code>[userlogsize="2048"]</code>	Taille limite en KO du journal applicatif. Au démarrage du module, le fichier est reseté si sa taille a dépassé la limite. Valeur par défaut : 2048 KO
<code><var</code> <code> name="ENV_VARIABLE_1"</code> <code> value="VALUE_1" /></code>	Variable d'environnement et sa valeur exportée avant l'exécution des scripts. Mettre autant de lignes que de variables.

13.8 Hostname virtuel (`<vhost>`, `<virtualhostname>` tags)

13.8.1 `<vhost>` Exemple

```
<vhost>
  <virtualhostname name="vhostname" envfile="vhostenv" />
</vhost>
```

Voir l'exemple en 15.6 [page 309](#).

13.8.2 <vhost> Syntaxe

```
<vhost>
  <virtualhostname
    name="virtual_hostname"
    envfile="path_of_a_file"
    [when="prim"|"second"|"both"]
  />
</vhost>
```



Le tag <vhost> et son sous-arbre **ne peuvent pas** être modifiés dynamiquement.

13.8.3 <vhost>, <virtualhostname> Attributs

<vhost>	
<virtualhostname	
name="virtual_hostname"	Définition du nom virtuel.
envfile="path_of_envfile"	Chemin d'un fichier d'environnement généré automatiquement par SafeKit à la configuration. Si le chemin du fichier est relatif, le fichier est généré dans les fichiers d'environnement du module, i.e. : SAFEUSERBIN Ce fichier est utilisé dans les scripts pour positionner le hostname virtuel. Voir le module <code>vhost.safe</code> livré avec le package Linux et Windows.
[when="prim" "second" "both"]	Définis quand le hostname virtuel doit être rendu. Par défaut, <code>prim</code> signifie quand le module est primaire (<code>PRIM</code> ou <code>ALONE</code>).
/>	
</vhost>	

13.8.4 <vhost> Description

Certaines applications ont besoin de voir le même hostname quel que soit le serveur d'exécution (typiquement, elles stockent le hostname dans un fichier répliqué). Le hostname virtuel peut être présenté à ces applications alors que les autres applications voient le hostname physique des serveurs.

⇒ Sur Linux

La mise en œuvre est basée sur la variable d'environnement `LD_PRELOAD` : les fonctions `gethostname` et `uname` sont surchargées.

⇒ Sur Windows

La mise en œuvre est basée sur la variable d'environnement `CLUSTER_NETWORK_NAME_` : les fonctions de l'API `name query` (`GetComputerName`, `GetComputerNameEx`, `gethostname`) sont surchargées.

Pour utiliser vhost avec un service, utiliser les commandes `vhostservice`
`<service>` [`<file>`] avant/après le démarrage/arrêt du service dans les scripts utilisateurs.

Pour un exemple complet, voir 15.6 [page 309](#).

13.9 Détection de la mort de processus ou de services (`<errd>`, `<proc>` tags)



La section `<errd>` nécessite d'avoir défini la section `<user/>`

13.9.1 `<errd>` Exemple

13.9.1.1 Surveillance de processus

Linux and Windows `myproc` est le nom de la commande associée au processus à surveiller :

```
<errd>
  <proc name="myproc" atleast="1" action="restart" class="prim" />
</errd>
```

Linux uniquement (pour SafeKit > 7.2.0.29), `oracle_.*` est une expression régulière sur le nom de la commande associée au processus à surveiller :

```
<errd>
  <proc name="oracle" nameregex="oracle_.*" atleast="1" action="restart"
class="prim"/>
</errd>
```

Voir l'exemple en 15.7 [page 311](#).

13.9.1.2 Surveillance de service

`myservice` est le nom du service Windows (pour safekit > 7.3) ou du service systemd Linux (pour safekit > 7.4.0.19) à surveiller :

```
<errd>
  <proc name="myservice" service="yes" action="restart" class="prim" />
</errd>
```

13.9.2 `<errd>` Syntaxe

```
<errd
  [polltimer="10"]
>
  <proc name="command name and/or resource name for the monitored process or
service"
    [service="no|yes"]
    [nameregex="regular expression on the command name"]
    [argregex="regular expression on process arguments, including command
name"]
    atleast="1"
    action="stopstart"|"restart"|"stop"|"executable_name"
    class="prim"|"both"|"pre"|"second"|"sec"|"othername"]
    [start_after="nb polling cycles"]
```

```
[atmax="-1"]
/>
...
</errd>
```



Le tag `<errd>` et son sous-arbre peuvent être entièrement modifiés dynamiquement.

13.9.3 `<errd>`, `<proc>` Attributs

<code><errd</code>	
<code>polltimer="30"</code>	Temps de polling, en secondes, entre 2 surveillances des processus. Valeur par défaut : 30 secondes
<code><proc</code>	Définition du processus à surveiller. Autant de lignes que de processus. Une ressource est associée à chaque <code><proc></code> et est nommée <code>proc.<valeur de l'attribut name></code> (par exemple <code>proc.process_name</code>). La ressource vaut <code>up</code> lorsque la condition de surveillance est vraie ; vaut <code>down</code> sinon.
<code>name="command_name"</code>	<code>name</code> est le nom de la commande associée au processus à surveiller. C'est aussi le nom de la ressource associée au processus surveillé. Au maximum 15 caractères pour Linux (le nom de la commande peut être tronqué) ; 63 pour Windows. Exemple : sur LINUX, <code>name="vi"</code> et sur Windows <code>name="notepad.exe"</code>. En Windows. Le nom est automatiquement converti en minuscule. Pour retrouver le nom des commandes associées aux processus, voir les commandes décrites en 13.9.4 page 277. Ou <code>name="command_name"</code> <code>nameregex="regular expression on the command name"</code> <i>Linux uniquement</i> <code>nameregex</code> est une expression régulière sur le nom de la commande pour sélectionner le processus à surveiller. <code>name</code> est le nom de la ressource associée au processus surveillé. Comme les expressions régulières sont définies dans le fichier XML <code>userconfig.xml</code>, certains caractères ne peuvent pas être utilisés '<code><</code>' ou '<code>></code>'.

Ou <pre>name="service_name" service="yes"</pre>	Exemple : <code>nameregex = "oracle _.*" name = "oracle"</code> surveille les processus oracle dont le nom de la commande respecte l'expression régulière La ressource associée est <code>proc.oracle</code> L'attribut <code>nameregex</code> est facultatif <code>name</code> est le nom du service à surveiller. C'est aussi le nom de la ressource associée au service surveillé. Au maximum, 63 caractères. Exemples : <pre>name="W32Time" service="yes" surveille le service de Temps Windows.</pre> <pre>name="ntpd" service="yes" surveille le service de Temps Linux (systemd ntpd.service) .</pre> L'attribut <code>service</code> est facultatif et sa valeur par défaut est : <code>no</code>
<pre>class= "prim" "both" "pre" "second" "sec" "othername"</pre>	Le processus appartient à une classe. La surveillance démarre avec la commande <code>safekit errd enable classname -m AM</code>. Les fonctions <code>enable/disable</code> des classes <code>prim</code>, <code>both</code>, <code>pre</code>, <code>second</code>, <code>sec</code> sont automatiques et faites par SafeKit dans le composant <code><user/></code> après/avant <code>start_prim/stop_prim</code>, <code>start_both/stop_both</code>, <code>start_second/stop_second</code>, <code>start_sec/stop_sec</code>. Pour une description des scripts, voir 14 page 293. Avec un autre nom de classe, les fonctions <code>enable/disable</code> doivent être faites explicitement après/avant le démarrage/arrêt des processus de la classe.
<pre>[argregex="regular expression on process arguments"]</pre>	Expression régulière sur la liste des arguments du processus incluant le nom de l'exécutable. Paramètre optionnel. Pour retrouver la liste des arguments d'un processus, utiliser les commandes décrites en 13.9.4 page 277. Exemples sur Linux avec l'éditeur <code>vi</code> sur <code>myfile</code> ("man <code>regex</code>" pour plus d'information) : <pre>name="vi" argregex=".*myfile.*" name="vi" argregex="/myrep/myfile.*" name="vi" argregex="/myrep/myfile"</pre> Exemples Windows avec l'éditeur Notepad sur <code>myfile</code> ("class <code>CatlRegExp</code>" pour plus d'information): <pre>name="notepad.exe" argregex=".*myfile.*"</pre>

	<pre>name="notepad.exe" argregex="c:\\myrep\\myfile.*" name="notepad.exe" argregex="c:\\myrep\\myfile"</pre>  Important Comme les expressions régulières sont définies dans le fichier XML <code>userconfig.xml</code>, certains caractères ne peuvent pas être utilisés '<' ou '>'.
<code>atleast="1"</code>	Nombre minimum de processus qui doivent s'exécuter. Si le minimum est atteint, SafeKit déclenche l'action. Exemple: <code>name="oracle" argregex=".*db1.*"</code> <code>atleast="1"</code> signifie qu'une action est déclenchée si 0 processus s'exécute sur l'instance <code>oracle/db1</code>. S'il est positionné à -1, ce critère n'est pas pris en compte. Valeur par défaut : 1
<pre>action= "restart" "stopstart" "stop" "noaction" "executable_name"</pre>	Action (ou handler) à exécuter sur le module. <code>noaction</code> réalise juste un logging de message, <code>restart</code> un redémarrage local et <code>stopstart</code> un basculement. Les commandes <code>restart/stopstart</code> incrémente le compteur <code>maxloop</code> pour vérifier que l'on n'est pas sur une faute reproductible. Pour la description de <code>maxloop</code>, voir 13.2 page 238. Pour un handler, soit mettre le chemin absolu du handler, soit mettre le chemin relatif au répertoire bin du module ("<code>SAFE/modules/AM/bin/</code>"). Nous conseillons un chemin relatif avec un handler défini dans le module. Avec un handler spécial, une classe avec un nom propre doit être définie. Pour un handler spécial en Linux, insérer à la fin du script, <code>exit 0</code> Pour un handler spécial en Windows, mettre à la fin <code>%SAFEBIN%\exitcode 0</code>. Sinon, c'est un échec et SafeKit exécute <code>stopstart</code> sur le module. Avec un handler spécial, le compteur <code>maxloop</code> n'est pas incrémenté. Utiliser la commande : <pre>safekit incloop -m AM -i <handler name></pre> Cette commande incrémente le compteur et retourne 1 quand la limite est atteinte. Voir l'exemple décrit en 15.7 page 311. Valeur par défaut : <code>stopstart</code>
<code>start_after=[nb polling cycles]</code>	Sans le paramètre <code>start_after</code>, la surveillance des processus est effective immédiatement. Sinon elle est retardée de $(n-1) * polltimer$ secondes où :

	⇒ <code>n</code> est la valeur de <code>start_after</code> ⇒ <code>polltimer</code> est le paramètre de polling d'<code>errd</code> (30 secondes par défaut) Par exemple si <code>start_after="3"</code>, la surveillance est retardée de 60 secondes $((3-1)*30)$. Le paramètre <code>start_after</code> est utile si le processus prend un certain temps à démarrer. Valeur par défaut : 0
Paramètres avancés	
<code>atmax="-1"</code>	Nombre maximum de processus qui peuvent s'exécuter. Si le maximum est atteint, SafeKit déclenche l'action. Avec <code>atmax="0"</code>, une action est déclenchée à chaque démarrage du processus. Valeur par défaut : -1 critère non pris en compte
<code></errd></code>	

13.9.4 <errd> Commandes



Si la commande est utilisée dans un script utilisateur, alors la variable d'environnement `SAFEMODULE` est positionnée et le paramètre "-m AM" n'est pas nécessaire

<pre>safekit -r errdpoll_running</pre>	Stocke son résultat dans le fichier <code><SAFEVAR>/errdpoll_reserrd</code> (<code>SAFEVAR = /var/safekit</code> sur Linux ou <code>c:\safekit\var</code> sur Windows) avec une ligne pour chaque processus : <code><pid> <command name> <command full name and arguments list> (parent=<parent pid>)</code> En Windows, le nom de la commande est affiché en minuscule. Utile pour déterminer le nom des processus à surveiller et leurs arguments pour une configuration <code><errd></code>
<pre>safekit errd disable "classname" -m AM</pre>	Suspend la surveillance des processus de la classe <code>classname</code> (pour le module applicatif AM). Doit être explicitement appelé dans les scripts <code>stop_...</code> avant d'arrêter l'application, pour des processus dans des classes différentes de <code>prim</code>, <code>both</code>, <code>second</code>, <code>sec</code>.

<pre>safekit errd enable "classname" -m AM</pre>	Redémarre la surveillance des processus de la classe <code>classname</code> (pour le module applicatif <code>AM</code>). Doit être explicitement appelé dans les scripts <code>start_...</code> après le démarrage de l'application, pour des processus dans des classes différentes de <code>prim</code>, <code>both</code>, <code>second</code>, <code>sec</code>.
<pre>safekit errd suspend -m AM</pre>	Suspend la surveillance des processus du module <code>AM</code> (sauf les processus SafeKit). Utile lorsque l'on stoppe manuellement l'application et que l'on ne veut pas de détection.
<pre>safekit errd resume -m AM</pre>	Redémarre la surveillance des processus du module <code>AM</code>
<pre>safekit errd list -m AM</pre>	Liste tous les processus du module <code>AM</code> surveillés incluant les processus SafeKit. La liste peut être également lue dans <code>SAFEVAR/modules/AM/errdlist</code>.
<pre>safekit kill -name="process_name" [-argregex="..."] -level="kill_level"</pre>	Le composant <code><errd></code> doit s'exécuter. Tue le(s) processus identifié(s) par le nom et les arguments. <code>level="test"</code> : affiche seulement la liste des processus <code>level="terminate"</code> : kill les processus en Windows <code>level="9"</code> : envoie le signal SIGKILL aux processus en Linux <code>level="15"</code> : envoie le signal SIGTERM aux processus en Linux Exemples Windows ("class CatlRegExp" pour plus d'information): <pre>safekit kill -name="notepad.exe" -argregex=".*myfile.*" -level="terminate"</pre> <pre>safekit kill -name="notepad.exe" -argregex="c:\\myrep\\myfile.*" -level="terminate"</pre> Exemples Linux ("man regex" pour plus d'information) : <pre>safekit kill -name="vi" -argregex=".*myfile.*" -level="9"</pre> <pre>safekit kill -name="vi" -argregex="/myrep/myfile.*" -level="9"</pre>

13.10 Checkers (<check> tags)

SafeKit apporte des checkers avec des règles de failover par défaut (voir section 13.18.5 page 291). Les checkers sont :

- ⇒ 13.11 « TCP checker (<tcp> tags) » page 280.
- ⇒ 13.12 « Ping checker (<ping> tags) » page 281.
- ⇒ 13.13 « Interface checker (<intf> tags) » page 282.
- ⇒ 13.14 « IP checker (<ip> tags) » page 283
- ⇒ 13.15 « Checker customisé (<custom> tags) » page 285.
- ⇒ 13.16 « Module checker (<module> tags) » page 286
- ⇒ 13.17 « Splitbrain checker (<splitbrain> tag) » page 288

13.10.1 <check> Exemple

Tous les checkers se définissent dans une seule section <check> :

```
<check>
  <!-- Insérer ci-dessous les tags <tcp> <ping> <intf> <ip> <custom> <module>
        <splitbrain> -->
</check>
```

13.10.2 <check> Syntaxe

```
<check>
  <tcp ...>
    <to .../>
  </tcp>
  ...
  <ping ...>
    <to .../>
  </ping>
  ...
  <intf ...>
    <to .../>
  </intf>
  ...
  <ip ...>
    <to .../>
  </ip>
  ...
  <custom .../>
  ...
  <module ...>
    [<to .../>]
  </module>
  ...
  <splitbrain .../>
</check>
```



Le tag <check> et son sous-arbre peuvent être entièrement modifiés dynamiquement.

13.11 TCP checker (<tcp> tags)



Par défaut, un checker <tcp> réalise un redémarrage local du module lorsque le service TCP est down.

13.11.1 <tcp> Exemple

```
<check>
  <tcp ident="Rltest" when="prim" >
    <to addr="R1" port="80"/>
  </tcp>
</check>
```



Insérer le tag <tcp> dans la section <check> si celle-ci est déjà définie.

Voir l'exemple en 15.8 [page 313](#).

13.11.2 <tcp> Syntaxe

```
<tcp
  ident="tcp_checker_name"
  when="prim|second|both|pre"
>
  <to
    addr="IP_address" or "name_to_check"
    port="TCP_port_to_check"
    [interval="10"]
    [timeout="5"]
  />
</tcp>
```

13.11.3 <tcp> Attributs

<tcp	Positionner autant de sections <tcp> qu'il y a de checkers <tcp>.
ident="tcp_checker_name"	Nom du checker TCP.
when="prim second both"	Utiliser cette valeur pour tester un service TCP interne à l'application Suivant cette valeur, le checker est démarré/arrêté après/avant les scripts <code>start_prim/stop_prim</code>, <code>start_second/stop_second</code>, <code>start_both/stop_both</code>. Action en cas de défaillance: <code>restart</code> du module (voir les règles de failover par défaut décrites en 13.18.5 page 291). Chaque restart incrémente le compteur <code>maxloop</code> (voir section 13.2.3 page 239).

when="pre"	Utiliser cette valeur pour tester un service TCP externe à l'application Le checker est démarré/arrêté après/avant les scripts prestart/poststop. Vous devez ajouter une règle de failover spéciale pour un tel checker. Typiquement: <code>external_tcp_service: if (tcp.tcp_checker_name == down) then wait();</code> Cette règle réalise un stopwait et met le module dans l'état <code>WAIT</code> tant que le service TCP externe ne répond pas (pour plus d'information, voir 13.18 page 289). Chaque stopwait incrémente le compteur <code>maxloop</code> (voir section 13.2.3 page 239).
<to	
addr="IP_@" or "name"	Adresse IP ou nom à checker (ex : 127.0.0.1 pour un service local). Adresse IPv4 ou IPv6.
port="value"	Port TCP port à checker
[interval="10"]	Temps, en secondes, entre 2 pollings. Valeur par défaut : 10 secondes
[timeout="5"]	Délai en secondes pour l'acceptation de la connexion. Valeur par défaut : 5 secondes
</tcp>	

13.12 Ping checker (<ping> tags)



Par défaut, un <ping> checker met le module dans l'état `WAIT` et attend que ping redevienne up.

13.12.1 <ping> Exemple

```
<check>
  <ping ident="testR2" >
    <to addr="R2"/>
  </ping>
</check>
```



Important

Insérer le tag <ping> dans la section <check> si celle-ci est déjà définie.

Voir l'exemple en 15.9 [page 313](#).

13.12.2 <ping> Syntaxe

```
<ping
  ident="ping_checker_name"
  [when="pre"]
```

```
>
<to
  addr="IP_address" or "name_to_check"
  [interval="10"]
  [timeout="5"]
/>
</ping>
```

13.12.3 <ping> Attributs

<ping	Mettre autant de section <ping> qu'il y a de ping checkers.
ident="ping_checker_name"	Nom du checker ping
[when="pre"]	Valeur par défaut Le checker est démarré/arrêté après/avant les scripts prestart/poststop. La règle de failover par défaut réalise un stopwait qui met le module dans l'état <code>WAIT</code> tant que le ping ne répond pas (pour plus d'informations, voir 13.18 page 289). Chaque stopwait incrémente le compteur <code>maxloop</code> (voir 13.2.3 page 239).
<to	
addr="IP_@ or name"	Adresse IP ou nom à checker Adresse IPv4 ou IPv6.
[interval="10"]	Intervalle de temps, en secondes, entre 2 pollings. Valeur par défaut : 10 secondes
[timeout="5"]	Délai en secondes sur la réponse au ping. Valeur par défaut : 5 secondes
</ping>	

13.13 Interface checker (<intf> tags)



Par défaut, un checker <intf> arrête le module et attend que l'interface réseau soit up.

13.13.1 <intf> Exemple

```
<check>
  <intf ident="test_eth0">
    <to local_addr="192.168.1.10"/>
  </intf>
</check>
```



Insérer le tag `<intf>` dans la section `<check>` si celle-ci est déjà définie.

Voir l'exemple en 15.10 [page 313](#).

13.13.2 `<intf>` Syntaxe

```
<intf
  ident="intf_checker_name"
  [when="pre"]

>
  <to
    local_addr="interface_physical_IP_address"/>
</intf>
```

13.13.3 `<intf>` Attributs

<code><intf</code>	 Les sections <code><intf></code> sont automatiquement générées lorsque <code><interface check="on"></code> est positionné (voir section 13.5 page 245).
<code>ident="intf_checker_name"</code>	Le nom de l'interface checker (adresse IP du réseau).
<code>[when="pre"]</code>	Valeur par défaut Le checker est démarré/arrêté après/avant les scripts <code>prestart/poststop</code>. La règle de failover par défaut réalise un stopwait qui met le module dans l'état <code>WAIT</code> tant que le ping ne répond pas (pour plus d'informations, voir 13.18 page 289). Chaque stopwait incrémente le compteur <code>maxloop</code> (voir 13.2.3 page 239).
<code><to local_addr="IP_@ /></code>	Adresse IP associée à l'interface à tester. Adresse IPv4 ou IPv6.
<code></intf></code>	

13.14 IP checker (`<ip>` tags)

En Windows et en Linux, ce checker vérifie qu'une adresse IP est bien configurée localement ; en Windows, il détecte en plus les conflits sur cette adresse.



Par défaut, un checker `<ip>` exécute un stopstart local du module lorsque l'adresse testée est down.

13.14.1 <ip> Exemple

```
<check>
  <ip ident="ip_check" >
    <to addr="192.168.1.10" />
  </ip>
</check>
```



Important

Insérer le tag <ip> dans la section <check> si celle-ci est déjà définie.

Voir l'exemple en 15.11 [page 314](#).

13.14.2 <ip> Syntaxe

```
<ip
  ident="ip_checker_name"
  [when="prim"]
>
  <to
    addr="IP_address" or "name_to_check"
    [interval="10"]
  />
</ip>
```

13.14.3 <ip> Attributs

<ip	Mettre autant de section <ip> qu'il y a de checkers ip.
ident="ip_checker_name"	Nom du checker ip (dont l'état est affiché par la commande <code>safekit state -v -m AM</code>). Chaque checker doit avoir un nom différent.
[when="prim"]	Valeur par défaut Le checker est démarré/arrêté après/avant les scripts <code>prestart/poststop</code> . La règle de failover par défaut réalise un <code>stopstart</code> du module (pour plus d'informations, voir 13.18 page 289). Chaque <code>stopstart</code> incrémente le compteur <code>maxloop</code> (voir 13.2.3 page 239).
<to	
addr="IP_@ or name"	adresse IP locale ou nom DNS à tester. Adresse IPv4 ou IPv6.
[interval="10"]	Intervalle de temps, en secondes, entre 2 tests. Valeur par défaut : 10 secondes
</ip>	

13.15 Checker personnalisé (<custom> tags)

Un checker personnalisé est un programme (script ou autre) que vous développez pour tester une ressource, l'application, Il s'agit d'une boucle qui effectue un test suivant une période adéquate. Son rôle est de positionner une ressource à "up" ou "down". Puis, une règle de failover dans `userconfig.xml` décide l'action à exécuter lorsque la ressource est down

13.15.1 <custom> Exemple

```
<check>
  <custom ident="AppChecker" when="prim" exec="mychecker" />
</check>
```



Important

Insérer le tag `<custom>` dans la section `<check>` si celle-ci est déjà définie. Définir en plus le checker personnalisé ainsi que la règle de failover associée.

Pour des exemples complets, voir les sections 15.12 [page 315](#).

13.15.2 <custom> Syntaxe

```
<custom
  ident="custom_checker_name"
  when="pre|prim|second|both"
  exec="executable_path"
  arg="executable_arguments"
/>
```

13.15.3 <custom> Attributs

<code><custom</code>	Positionner autant de sections <code><custom></code> qu'il y a de checkers personnalisés.
<code>ident="custom_checker_name"</code>	Nom du checker personnalisé Un checker personnalisé positionne sa ressource avec la commande <code>safekit set -r custom.custom_checker_name -v up down</code>.
<code>when="pre"</code>	Le checker est démarré/arrêté après/avant les scripts <code>prestart/poststop</code>. Vous devez écrire une règle de failover qui doit réaliser un <code>stopwait</code>. Typiquement: <pre>wait_custom_checker: if (custom.custom_checker_name == down) then wait();</pre> Elle met le module dans l'état <code>WAIT</code> tant que la ressource est down. Noter que l'état initial de la ressource est <code>init</code> et que la failover machine reste dans l'état <code>WAIT</code> tant que ressource n'est pas positionnée à <code>up</code> (pour plus d'information, voir 13.18 page 289). Chaque <code>stopwait</code> incrémente le compteur <code>maxloop</code> (voir 13.2.3 page 239).

<code>when="prim" "second" "both"</code>	Suivant cette valeur, le checker est démarré/arrêté après/avant les scripts <code>start_prim/stop_prim</code>, <code>start_second/stop_second</code>, <code>start_both/stop_both</code>. Vous devez écrire une règle de failover qui doit réaliser un restart, stopstart ou stop. Typiquement: <pre>restart_custom_checker: if (custom.custom_checker_name == down) then restart();</pre> Pour plus d'information, voir 13.18 page 289. Chaque restart ou stopstart incrémente le compteur <code>maxloop</code> (voir 13.2.3 page 239).
<code>exec="executable_path"</code>	Définit le chemin de l'exécutable du checker personnalisé (un script ou un binaire). Lorsque le chemin est relatif, le custom checker est recherché dans <code>SAFEUSERBIN</code>. Mettre dans ce cas votre binaire dans <code>SAFE/modules/AM/bin/</code> (pour plus d'information, voir 10.1 page 157). Nous conseillons un chemin relatif avec un exécutable dans le module. En Windows, l'exécutable peut être un binaire ou un script ps1, vbs ou cmd En Linux, l'exécutable peut être un binaire ou un script shell
<code>arg="executable_arguments"</code>	Définit les arguments à passer au checker personnalisé.

13.16 Module checker (<module> tags)

Le checker de module teste la disponibilité d'un autre module. Le checker est démarré/arrêté après/avant les scripts `prestart/poststop`. Lorsque le checker de module détecte qu'un module externe est down, la failover machine réalise un stopwait et met le module local en `WAIT` jusqu'à ce que le module externe soit de nouveau détecté up. Le checker de module effectue également une action stopstart lorsqu'il détecte que le module externe a été redémarré (soit par un restart, un stopstart ou un failover).

Chaque stopwait ou stopstart incrémente le compteur `maxloop` (voir 13.2.3 [page 239](#)).

Le checker de module récupère l'état du module en se connectant au service web de SafeKit qui s'exécute sur le serveur sur lequel le module est activé (voir 10.6 [page 170](#)).

13.16.1 <module> Exemple

Exemple utilisant la configuration par défaut du service web de SafeKit (protocole : HTTP, port : 9010) :

```
<check>
  <module name="mirror">
    <to addr="Mlhost" />
  </module>
</check>
```

Exemple utilisant la configuration sécurisée du service web de SafeKit (protocole : HTTPS, port : 9453) :

```
<check>
  <module name="mirror">
    <to addr="Mlhost" port="9453" secure="on"/>
  </module>
</check>
```



Insérer le tag <module> dans la section <check> si celle-ci est déjà définie.

Important

Pour d'autres exemples, voir les sections 15.3 [page 305](#) et 15.13 [page 317](#).

13.16.2 <module> Syntaxe

```
<module
  [ident="module_checker_name"]
  name="external_module_name">
  [<to
    addr=" IP_@ or name the Safekit server running the external module"
    [port=port of the SafeKit web server"]
    [interval="10"]
    [timeout="5"]
  />]
</module>
```

13.16.3 <module> Attributs

<module	Positionner autant de sections <module> qu'il y a de checkers de module.
name="external_module_name"]	Nom du checker de module.
[ident="module_checker_name"]	Nom du module externe à checker. Par défaut external_module_name_<IP_@ or name of the server >.
[<to	Définition du/des serveurs exécutant le module externe. Par défaut, le serveur local.
addr="IP_@ or name of the server"	Adresse IP ou nom du module externe. Adresse IPv4 ou IPv6.
[port=port du service web de SafeKit"]	Port du service web SafeKit pour le checking. Par défaut, 9010.
[interval="10"]	Intervalle de temps, en secondes, entre 2 pollings. Valeur par défaut : 10 secondes

<code>[timeout="5"]</code>	Délai en secondes pour la réception d'une réponse à la requête check. Valeur par défaut : 5 secondes
<code>[secure="on" "off"]</code>	Utilisation du protocole HTTP (<code>secure="off"</code>) ou HTTPS (<code>secure="on"</code>) Par défaut : <code>off</code>
<code>/>]</code>	
<code></module></code>	

13.17 Splitbrain checker (<splitbrain> tag)

SafeKit offre un checker de splitbrain à utiliser pour des architectures miroir. Le splitbrain est une situation où, à la suite d'une panne matérielle ou logicielle, les deux nœuds SafeKit deviennent primaires, chaque nœud croyant que l'autre ne fonctionne plus. Ceci implique que l'application tourne sur les deux nœuds et, lorsque la réplication est active, les données peuvent se trouver dans un état incohérent.

Pour prévenir ce phénomène, le checker de splitbrain, sur détection d'isolation réseau entre les serveurs, sélectionne un unique nœud pour devenir primaire. L'autre nœud devient non à jour et se bloque dans l'état `WAIT` :

⇒ Jusqu'à ce qu'il reçoive à nouveau les heartbeats de l'autre serveur

Ou

⇒ Si l'administrateur le juge opportun et s'assure que l'autre serveur n'est plus dans l'état primaire, il peut forcer son démarrage en primaire (par l'exécution des commandes `safekit stop -m AM` puis `safekit prim -m AM`).

L'élection du serveur primaire repose sur le ping d'un composant externe, appelé **witness**. Le réseau doit être configuré de telle manière qu'en cas d'isolation réseau, un seul des deux serveurs a accès au witness (c'est celui-ci qui sera élu primaire). Dans le cas contraire, les deux nœuds deviendront primaires.



Le ping entre les 2 nœuds et avec le witness doit être ouvert.

13.17.1 <splitbrain> Exemple

```
<check>
  <splitbrain ident="SBtest" exec="ping" arg="192.168.1.100" />
</check>
```



Insérer le tag `<splitbrain>` dans la section `<check>` si celle-ci est déjà définie.

13.17.2 <splitbrain> Syntaxe

```
<splitbrain
  ident="witness"
  exec="ping"
  arg="witness IP address "
/>
```

13.17.3 <splitbrain> Attributs

<splitbrain	Une seule section <splitbrain>.
ident="witness"	Nom de la ressource affichée par la commande <code>safekit state -v -m AM.splitbrain.witness</code> représente l'état du witness.
[when="pre"]	Valeur fixe. Le checker est démarré/arrêté après/avant les scripts <code>prestart/poststop</code> . Sur détection de splitbrain, le serveur pour lequel <code>splitbrain.witness="up"</code> devient primaire. Sur l'autre, pour lequel <code>splitbrain.witness="down"</code> , il y a affectation de la ressource <code>splitbrain.uptodate</code> à "down". La règle de failover par défaut réalise un stopwait qui met le module dans l'état <code>WAIT</code> (pour plus d'information, voir 13.18 page 289). Chaque stopwait incrémente le compteur <code>maxloop</code> (voir 13.2.3 page 239).
exec="ping"	Valeur fixe. Utilise ping pour tester l'accessibilité au witness et affecte la ressource <code>splitbrain.witness</code> .
arg="IP_@ or name"	Adresse IP ou nom du witness Adresse IPv4 ou IPv6.
</splitbrain>	

13.18 Failover machine (<failover> tag)

SafeKit apporte des checkers (interface réseau, ping, tcp, custom, module). Un checker vérifie l'état d'une ressource (par défaut toutes les 10 secondes) et positionne son état à up ou down (voir section 13.10 [page 279](#)). La machine de failover évalue régulièrement (par défaut toutes les 5 secondes) l'état global de toutes les ressources et applique une action en fonction des règles de failover écrites dans un langage simple.

Dans un module ferme, la machine de failover ne fonctionne que sur les états locaux des ressources alors que dans un module miroir, elle peut fonctionner sur les états locaux et les états distants des ressources. Comme les états des ressources transitent par les voies de heartbeats, il est conseillé d'avoir plusieurs voies de heartbeats (voir section 13.3 [page 241](#)).

13.18.1 <failover> Exemple

```
<failover>
  <![CDATA[
    ping_failure: if (ping.testR2 == down) then stopstart();
  ]]>
</failover>
```

Voir aussi les règles de failover par défaut décrites en 13.18.5 [page 291](#).

13.18.2 <failover> Syntaxe

```
<failover [extends="yes"] [period="5000"] [handle_time="15000"]>
  <![CDATA[
    label: if (expression) then action;
    ...
  ]]>
</failover>
```



Le tag <failover> et son sous-arbre **ne peuvent pas** être modifiés dynamiquement.

13.18.3 <failover> Attributs

<failover	
[extends="yes" "no"]	Avec extends="yes", les nouvelles règles de failover étendent les règles par défaut (voir 13.18.5 page 291). Avec extends="no", les règles par défaut sont écrasées (à éviter). Valeur par défaut : yes
[period="5000"]	Période entre 2 évaluations. Valeur par défaut : 5000 millisecondes (5 secondes)
[handle_time="15000"]	Une action (stop() stopstart() wait() restart() swap()) doit être stable pendant handle_time (en millisecondes) avant d'être appliquée. Valeur par défaut : 15000 millisecondes (15 secondes). handle_time doit être un multiple de period.

13.18.4 <failover> Commandes

<pre>safekit set [-m AM] -r resource_class.resource_id -v resource_state</pre>	L'état d'une ressource est positionné par un checker avec cette commande Exemples : <pre>safekit set -r custom.myresource -v up safekit set -r custom.myresource -v down</pre>
<pre>[-n] [-l]</pre>	Depuis SafeKit 7.5, chaque affectation des principales ressources est mémorisée dans un journal afin de conserver l'historique leur état.

	Utiliser <code>-n</code> pour désactiver la journalisation de l'affectation en cours ou <code>-l</code> pour la forcer
<code>safekit stopwait -i "identity"</code>	Equivalent à la commande <code>wait()</code> de la failover machine (voir 13.18 page 289). Avec <code>stopwait</code>, (1) <code>poststop</code> et <code>prestart</code> scripts ne sont pas appelés et (2) les checkers de type <code>when="pre"</code> ne sont pas stoppés.

Les autres commandes `restart()`, `stopstart()`, `stop()`, `swap()` sont équivalentes aux commandes de contrôle (avec le paramètre `-i identity` en plus) décrites en 9.4 page 150.



`maxloop / loop_interval / automatic_reboot` s'appliquent si le paramètre `-i identity` est passé aux commandes (voir 13.2 page 238). Ce qui est le cas lorsque les commandes sont appelées à partir de la failover machine.

13.18.5 Règles de failover

Les règles de failover par défaut pour les checkers sont :

```
<failover>
<![CDATA[
/* rule for module checkers */
module_failure: if (module.? == down) then wait();
/* rule for interface checkers */
interface_failure: if (intf.? == down) then wait();
/* rule for ping checkers */
ping_failure: if (ping.? == down) then wait();
/* rule for tcp checkers */
tcp_failure: if (tcp.? == down) then restart();
/* rule for ip checkers */
ip_failure: if (ip.? == down) then stopstart();
/* rules for splitbrain */
splitbrain_failure: if (splitbrain.uptodate == down) then wait();
]]>
</failover>
```

Elles sont définies dans le fichier `SAFE/private/conf/include/failover.xml`. Il existe en plus des règles de failover dédiées à la gestion de la réplication.

La commande `WAKEUP` est automatiquement générée lorsqu'aucune action `wait()` n'est applicable.



Depuis SafeKit 7.5, les règles de failover utilisent une nouvelle syntaxe et les règles pour la réplication sont définies dans un autre fichier `SAFE/private/conf/include/rfs.xml`.

En complément des règles par défaut, l'utilisateur peut définir ses propres règles de (pour un custom checker par exemple) en respectant la syntaxe suivante :

```
label: if ( expression ) then action;
```

with:

```
⇒ label ::= string
⇒ action ::= stop() | stopstart() | wait() | restart() | swap()
⇒ expression ::= ( expression )
| ! expression
| expression && expression
| expression || expression
| expression == expression
| expression != expression
| resource ::= [local. | remote.] 0/1resource_class.resource_id
| resource_state
```

La syntaxe pour désigner les ressources est la suivante :

```
resource ::= [local. | remote.] 0/1resource_class.resource_id (default: local)
resource_class ::= ping | intf | tcp | ip | custom | module | heartbeat | rfs
resource_id ::= * | ? | name
resource_state ::= init | down | up | unknown
```

init	Etat spécial d'initialisation tant que le checker n'a pas démarré. Si une ressource dans l'état init est testé dans une règle de failover, cette règle n'est pas évaluée tant que la ressource est dans l'état init.
up	Etat OK
down	Etat KO
unknown	Etat inconnu pour une ressource distante (module arrêté sur l'autre nœud)

14. Scripts applicatifs pour la configuration du module

- ⇒ 14.1 « Liste des scripts » [page 293](#)
- ⇒ 14.2 « Automate d'exécution des scripts » [page 295](#)
- ⇒ 14.3 « Variables d'environnement et arguments passés aux scripts » [page 296](#)
- ⇒ 14.4 « Commandes spéciales SafeKit pour les scripts » [page 296](#)

Pour activer l'appel des scripts utilisateurs, le tag `<user>` doit être présent dans `userconfig.xml` (voir 13.7 [page 270](#)).

Les scripts doivent être des exécutables :

- ✓ en Windows, un exécutable avec l'extension et le type : `.cmd`, `.vbs`, `.ps1`, `.bat` ou `.exe`
- ✓ en Linux, tout type d'exécutable

Chaque fois que vous modifiez les scripts, vous devez réappliquer la configuration sur les serveurs (avec la console ou la commande SafeKit).

Des exemples de scripts sont donnés en 15.1 [page 302](#) pour un module miroir, et en 15.2 [page 303](#) pour un module ferme.



Au moment de la configuration, les scripts sont copiés de `SAFE/modules/AM/bin` dans le répertoire d'exécution `SAFE/private/modules/AM/bin` (=SAFEUSERBIN, ne pas modifier les scripts à cet endroit).

14.1 Liste des scripts

Ci-dessous la liste des scripts pouvant être définis par l'utilisateur. Les scripts essentiels sont les scripts start/stop qui démarrent et arrêtent l'application au sein du module.

14.1.1 Scripts start/stop

<code>start_prim</code> <code>stop_prim</code>	Scripts pour un module miroir. Démarre l'application sur le serveur <code>ALONE</code> ou <code>PRIM</code>
<code>start_both</code> <code>stop_both</code>	Scripts pour un module ferme. Démarre l'application sur tous les serveurs <code>UP</code> de la ferme. Dans le cas spécial où ils existent dans un module miroir, ils sont exécutés dans les états <code>PRIM</code> , <code>SECOND</code> ou <code>ALONE</code>

start_second stop_second	Scripts spéciaux pour un module miroir. Exécutés sur le serveur SECOND  Quand le serveur SECOND devient primaire, stop_second suivi de start_prim est exécuté
start_sec stop_sec	Scripts spéciaux pour un module miroir. Voir l'automate d'exécution des scripts décrit en 14.2 page 295
stop_[both, prim, second, sec] [force]	Scripts pour tous les modules. Ces scripts sont appelés 2 fois : une fois pour un arrêt propre (sans le paramètre force en premier argument) et une fois pour un arrêt forcé (avec le paramètre force en premier argument)
prestart poststop	Scripts pour tous les modules. Exécutés au tout début du démarrage du module et à la fin. Par défaut, prestart contient stop_sec, stop_second, stop_prim, stop_both pour arrêter l'application avant de démarrer le module sous contrôle SafeKit
transition	Scripts pour tous les modules. Ce script est appelé sur changement d'état décrit en 14.2 page 295

14.1.2 Autres scripts

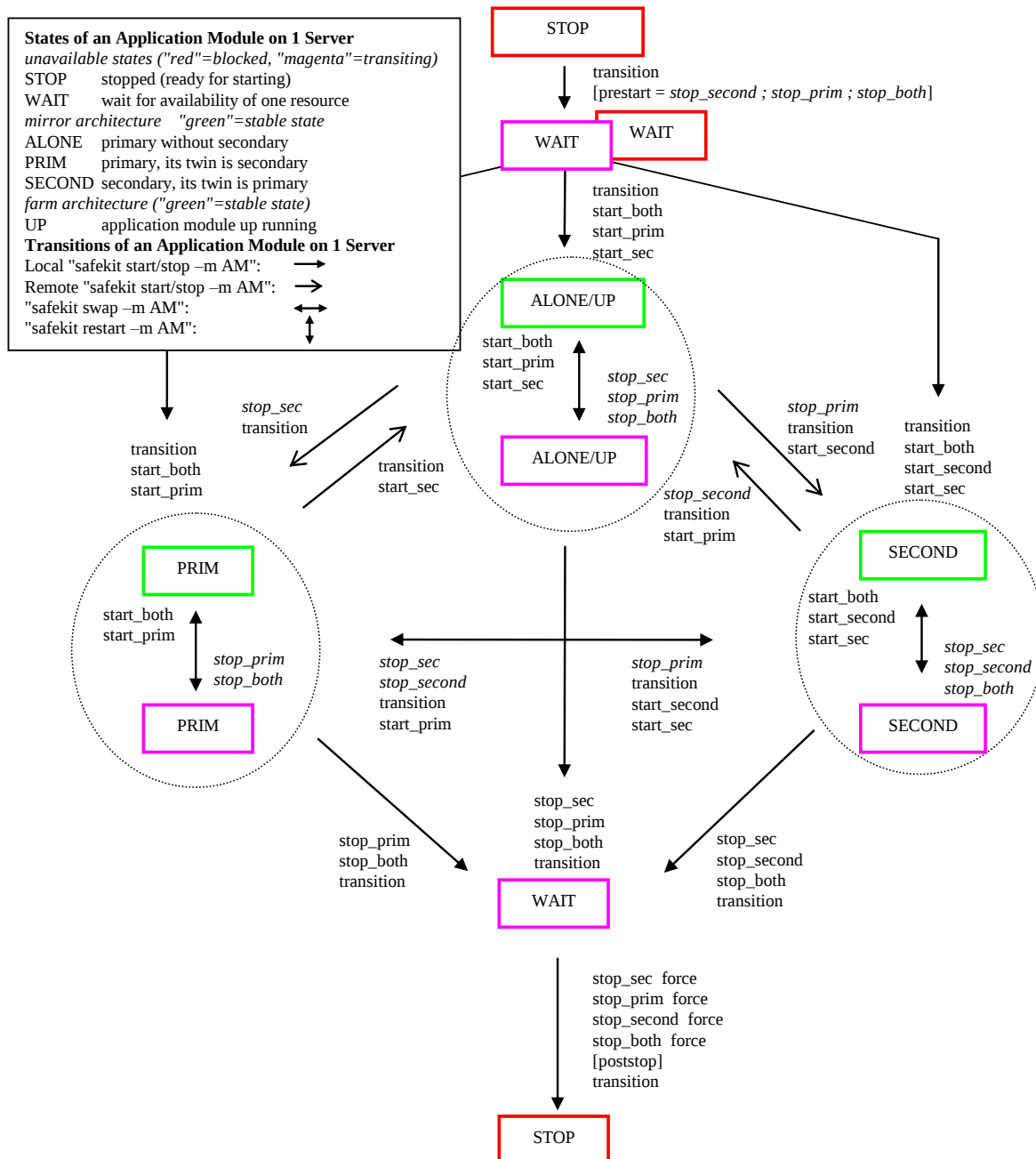
config	config est appelé lors de l'exécution de la commande safekit config -m AM. Vous pouvez exécuter dans ce script des commandes de configuration applicative.
deconfig	deconfig est appelé lors de l'exécution de la commande safekit deconfig -m AM, celle-ci étant automatiquement effectuée lors de la désinstallation du module Dans ce script, vous devez « défaire » les actions effectuées dans le script config.
confcheck	confcheck est appelé lors de l'exécution de la commande safekit confcheck -m AM. Vous pouvez exécuter dans ce script des opérations de contrôle de changement de configuration de l'application.
state	state est appelé lors de l'exécution de la commande safekit state -v -m AM.
level	level est appelé lors de l'exécution de la commande safekit level -m AM command.

14.2 Automate d'exécution des scripts



Exemple : au démarrage, la première transition de STOP vers WAIT appelle le script `transition STOP WAIT`.

La plupart du temps les scripts stop sont appelés 2 fois (sans l'option `force` puis avec l'option `force`). Dans ce cas, le nom du script est en italique.



14.3 Variables d'environnement et arguments passés aux scripts

Tous les scripts sont appelés avec 3 paramètres :

- ✓ l'état courant (STOP, WAIT, ALONE, PRIM, SECOND, UP)
- ✓ le prochain état (STOP, WAIT, ALONE, PRIM, SECOND, UP)
- ✓ l'action (start, stop, stopstart ou stopwait)

Les scripts de type `stop` sont appelés 2 fois :

- ✓ une première fois pour un arrêt propre de l'application
- ✓ une deuxième fois avec l'argument `force` pour un arrêt forcé (avec `force` comme premier argument)

Les variables d'environnement utilisables à l'intérieur des scripts sont :

- ✓ `SAFE`, `SAFEBIN`, `SAFEUSERBIN`, `SAFEUSERVAR` (voir 10.1 page 157)
- ✓ toutes les variables définies dans le tag `<user>` de `userconfig.xml` (voir 13.7 page 270).

14.4 Commandes spéciales SafeKit pour les scripts

Les commandes spéciales sont sous `SAFE/private/bin`. Les commandes peuvent être appelées directement dans les scripts utilisateurs avec :

- ⇒ `%SAFEBIN%\specialcommand` en Windows
- ⇒ `$SAFEBIN/specialcommand` en Linux

En dehors des scripts, il faut utiliser la commande `safekit -r`.

<pre>safekit -r <special command> [<args>]</pre>	<code><special command> <args></code> exécuté dans l'environnement SafeKit. Lorsque le path absolu n'est pas spécifié, la commande est recherchée dans <code>SAFEBIN=SAFE/private/bin</code>.  A utiliser lorsque les commandes spéciales sont passées hors script SafeKit
--	--

14.4.1 Commandes pour Windows

14.4.1.1 Commandes `sleep`, `exitcode`, `sync`

- ⇒ `%SAFEBIN%\sleep.exe <timeout value in seconds>`

A utiliser dans les scripts `stop` car `net stop service` n'est pas synchrone.

- ⇒ `%SAFEBIN%\exitcode.exe <exit value>`

Pour sortir d'un script avec une valeur d'erreur

- ⇒ `%SAFEBIN%\sync.exe \\.\<<drive letter:>`

Pour synchroniser les caches file system

14.4.1.2 Commande namealias

```
%SAFEBIN%/namealias [-n | -s ] <alias name>
```

-n pour ajouter un nouveau nom NetBIOS en alias (`start_prim`) ou -s pour supprimer un alias NetBIOS (`stop_prim`)

Vous pouvez utiliser la commande SafeKit `netnames` (ou la commande Windows `nbtstat`) pour lister les informations NetBIOS.

14.4.2 Commandes pour Linux

14.4.2.1 Gestion de la crontab

<code>\$SAFEBIN/gencron</code>	
<code>[del add]</code>	del pour désactiver des entrées dans <code>stop_prim</code> (en insérant des commentaires) ou add pour activer des entrées dans <code>start_prim</code> (en retirant les commentaires).
<code><user name></code>	Nom de l'utilisateur dans la crontab.
<code>[all <command name>]</code>	S'applique à toutes les entrées ou au nom de la commande
<code>-c "<comment>"</code>	Entête du commentaire qui sera insérée.

Par exemple, pour désactiver/activer l'entrée de la crontab :

```
5 0 * * * $HOME/bin/daily.job >> $HOME/tmp/out 2>&1
```

Insérer dans `stop_prim` :

```
$SAFEBIN/gencron del admin daily.job -c "SafeKit configuration for $SAFEMODULE"
```

Et insérer dans `start_prim` :

```
$SAFEBIN/gencron add admin daily.job -c "SafeKit configuration for $SAFEMODULE"
```

14.4.2.2 Commande bounding

```
$SAFEBIN/boundcmd <timeout value> <command path> [<args>]
```

Limite le temps d'exécution d'une commande

`boundcmd` retourne l'exit code de la commande si la commande se termine et 2 sinon.

Par exemple, pour flusher des données sur disque avec un timeout de 30 secondes :

```
$SAFEBIN/boundcmd 30 /bin/sync 1>/dev/null 2>&1
```

14.4.3 Commandes pour Windows et Linux

<pre>safekit -r processtree list kill ...</pre>	Liste les processus en cours d'exécution sous forme d'arbre d'appel (excepté pour l'option <code>all</code>) et arrête les processus (option <code>kill</code>) ⇒ <code>safekit -r processtree list all</code> Liste tous les processus en cours d'exécution. ⇒ <code>safekit -r processtree list <process command name></code> Liste les processus dont le nom de commande est celui spécifié. ⇒ <code>safekit -r processtree kill <process command name></code> Liste et arrête les processus en cours d'exécution dont le nom de commande est celui spécifié. ⇒ <code>safekit -r processtree list kill <process command name> all <regular expression on the full command - path and arguments></code> Liste (et arrête) les processus dont le nom de commande est celui spécifié et dont les arguments correspondent à l'expression régulière. Exemples en Windows ("class CatlRegExp" pour plus d'informations) <pre>safekit -r processtree kill notepad.exe ".*myfile.*" safekit -r processtree list all "mirror"</pre> Exemples en Linux ("man regex" pour plus d'informations) : <pre>safekit -r processtree kill vi ".*myfile.*" safekit -r processtree list all "mirror"</pre>
<pre>safekit incloop -m AM -i <handler name></pre>	SafeKit fournit un compteur <code>maxloop</code>, du nombre de <code>restart</code> et <code>stopstart</code> du module sur détection d'erreur. Le module est arrêté lorsque ce compteur atteint la valeur <code>maxloop</code> sur la période <code>loop_interval</code>. Lors de l'exécution d'un handler spécial, le compteur <code>maxloop</code> n'est pas incrémenté. Pour l'incrémenter, utilisez la commande : <pre>safekit incloop -m AM -i <handler name></pre> La commande augmente le compteur <code>maxloop</code> pour le module <code>AM</code> et retourne 1 lorsque la limite est atteinte.
<pre>safekit resetloop -m AM [-i <handler name>]</pre>	Réinitialise le compteur <code>maxloop</code> pour le module <code>AM</code> (affectation de la valeur à 0)
<pre>safekit checkloop -m AM</pre>	Pour tester le compteur <code>maxloop</code> pour le module <code>AM</code>, utilisez la commande <code>safekit checkloop -m AM</code> ⇒ elle retourne 0 lorsque la limite <code>maxloop</code> n'est pas atteinte ou si la dernière incrémentation a eu lieu en dehors de la période <code>loop_interval</code>

⇒	elle retourne 1 quand la limite <code>maxloop</code> a été atteinte dans la période <code>loop_interval</code>
---	--

15.Exemples de userconfig.xml et scripts utilisateurs

- ⇒ 15.1 « Exemple du module générique miroir avec `mirror.safe` » [page 302](#)
- ⇒ 15.2 « Exemple du module générique ferme avec `farm.safe` » [page 303](#)
- ⇒ 15.3 « Un module ferme dépendant d'un module miroir » [page 305](#)
- ⇒ 15.4 « Exemple d'un flux de réplication dédié » [page 306](#)
- ⇒ 15.5 « Exemples de partage de charge dans un module ferme » [page 306](#)
- ⇒ 15.6 « Exemple d'un hostname virtuel avec `vhost.safe` » [page 309](#)
- ⇒ 15.7 « Détection de la mort de processus avec `softerrd.safe` » [page 311](#)
- ⇒ 15.8 « Exemple d'un checker TCP » [page 313](#)
- ⇒ 15.9 « Exemple d'un checker ping » [page 313](#)
- ⇒ 15.10 « Exemple d'un checker d'interface réseau » [page 313](#)
- ⇒ 15.11 « Exemple d'IP checker » [page 314](#)
- ⇒ 15.12 « Exemple d'un checker personnalisé avec `customchecker.safe` » [page 315](#)
- ⇒ 15.13 « Exemple d'un checker de module avec `leader.safe` et `follower.safe` » [page 317](#)

Certains exemples décrits sont tirés des modules livrés avec le package SafeKit, sous `SAFE/Application_Modules`. Vous pouvez les installer avec la console web (voir 3.7.4 [page 64](#)) pour examiner en détail leur contenu.

D'autres exemples d'intégration sont décrits sous <https://www.evidian.com/fr/produits/haute-disponibilite-logiciel-clustering-application/configuration-cluster-basculement/>.



Les `.safe` sont plate-forme dépendants et, par conséquent, différents en Windows et Linux.

Tous les exemples utilisent la configuration du cluster SafeKit suivante (voir 12 [page 231](#)) :

```
<cluster>
  <lan name="net3">
    <node name="node1" addr="10.1.0.2"/>
    <node name="node2" addr="10.1.0.3"/>
    <node name="node3" addr="10.1.0.3"/>
  </lan>

  <lan name="default" console="off">
    <node name="node1" addr="192.168.1.1"/>
    <node name="node2" addr="192.168.1.2"/>
  </lan>
  <lan name="repli" console="off">
    <node name="node1" addr="10.0.0.2"/>
  </lan>
</cluster>
```

```
<node name="node2" addr="10.0.0.3"/>
</lan>
</lans>
</cluster>
```

15.1 Exemple du module générique miroir avec `mirror.safe`

Ci-dessous le fichier de configuration et les scripts utilisateurs du module miroir générique, `mirror.safe`, pour Windows. Pour Linux, se référer au module `mirror.safe` livré avec le package Linux.

conf/userconfig.xml – voir 13 [page 237](#)

```
<!-- Mirror Architecture with Real Time File Replication and Failover -->
<!DOCTYPE safe>
<safe>
  <service mode="mirror" defaultprim="alone" maxloop="3" loop_interval="24"
failover="on">
    <heart pulse="700" timeout="30000">
      <heartbeat name="default" ident="flow"/>
    </heart>
    <rfs async="second" acl="off" locktimeout="200" nbrei="3"
iotimeout="300">
      <replicated dir="c:\test1replicated" mode="read_only"/>
      <replicated dir="c:\test2replicated" mode="read_only"/>
    </rfs>
    <vip>
      <interface_list>
        <interface check="on" arpreroute="on">
          <real_interface>
            <virtual_addr addr="192.168.4.10" where="one_side_alias"/>
          </real_interface>
        </interface>
      </interface_list>
    </vip>
    <user nicestoptimeout="300" forcestoptimeout="300" logging="userlog"/>
  </service>
</safe>
```

bin/start_prim.cmd – voir 14 [page 293](#)

```
@echo off

rem Script called on the primary server for starting application services

rem For logging into SafeKit log use:
rem "%SAFE%\safekit" printi | printe "message"

rem stdout goes into Application log
echo "Running start_prim %*"

set res=0

rem Fill with your services start call
rem net start "myservice" /Y

set res=%errorlevel%
```

```
if %res% == 0 goto end

:stop
"%SAFE%\safekit" printe "start_prim failed"

rem uncomment to stop SafeKit when critical
rem "%SAFE%\safekit" stop -i "start_prim"

:end
```

bin/stop_prim.cmd - voir 14 [page 293](#)

```
@echo off
rem Script called on the primary server for stopping application services

rem For logging into SafeKit log use:
rem "%SAFE%\safekit" printi | printe "message"

rem -----
rem
rem 2 stop modes:
rem
rem - graceful stop
rem   call standard application stop with net stop
rem
rem - force stop (%1=force)
rem   kill application's processes
rem
rem -----

rem stdout goes into Application log
echo "Running stop_prim %*"

set res=0

rem default: no action on forcestop
if "%1" == "force" goto end

rem Fill with your service(s) stop call
rem net stop "myservice" /Y

rem If necessary, uncomment to wait for the stop of the services
rem "%SAFE\BIN%\sleep" 10

if %res% == 0 goto end

"%SAFE%\safekit" printe "stop_prim failed"
:end
```

15.2 Exemple du module générique ferme avec `farm.safe`

Ci-dessous le fichier de configuration et les scripts utilisateurs pour le module ferme générique, `farm.safe`, pour Windows. Pour Linux, se référer au module `farm.safe` livré avec le package Linux.

conf/userconfig.xml - voir 13 [page 237](#)

```
<!-- Farm Architecture with Load-Balancing and Failover -->
<!DOCTYPE safe>
<safe>
  <service mode="farm" maxloop="3" loop_interval="24">
    <farm>
      <lan name="default" />
    </farm>
    <vip>
      <interface_list>
        <interface check="on" arpreroute="on">
          <virtual_interface type="vmac_directed">
            <virtual_addr addr="192.168.4.20" where="alias"/>
          </virtual_interface>
        </interface>
      </interface_list>
      <loadbalancing_list>
        <group name="FarmProto">
          <rule port="9010" proto="tcp" filter="on_port"/>
        </group>
      </loadbalancing_list>
    </vip>
    <user nicestoptimeout="300" forcestoptimeout="300" logging="userlog"/>
  </service>
</safe>
```

bin/start_both.cmd - voir 14 [page 293](#)

```
@echo off

rem Script called on all servers for starting applications

rem For logging into SafeKit log use:
rem "%SAFE%\safekit" printi | printe "message"

rem stdout goes into Application log
echo "Running start_both %*"

set res=0

rem Fill with your services start call
rem net start "myservice" /Y

set res=%errorlevel%

if %res% == 0 goto end

:stop
set res=%errorlevel%
"%SAFE%\safekit" printe "start_both failed"

rem uncomment to stop SafeKit when critical
rem "%SAFE%\safekit" stop -i "start_both"
:end
```

bin/stop_both.cmd - voir 14 [page 293](#)

```
@echo off
```

```

rem Script called on all servers for stopping application

rem For logging into SafeKit log use:
rem "%SAFE%\safekit" printi | printe "message"

rem -----
rem
rem 2 stop modes:
rem
rem - graceful stop
rem   call standard application stop with net stop
rem
rem - force stop (%1=force)
rem   kill application's processes
rem
rem -----

rem stdout goes into Application log
echo "Running stop_both %"

set res=0

rem default: no action on forcestop
if "%1" == "force" goto end

rem Fill with your services stop call
rem net stop "myservice" /Y

rem If necessary, uncomment to wait for the stop of the services
rem "%SAFEBIN%\sleep" 10

if %res% == 0 goto end

"%SAFE%\safekit" printe "stop_both failed"
:end

```

15.3 Un module ferme dépendant d'un module miroir

Dans l'exemple ci-dessous, le module ferme ne peut démarrer qu'à condition que le module miroir soit opérationnel. Cette architecture peut être utilisée pour lier un module ferme IIS à un module miroir Microsoft SQL server. Elle est basée sur la configuration d'un module checker dans le module ferme. Pour plus détails, voir 13.16 [page 286](#).

farm/conf/userconfig.xml - voir 13 [page 237](#)

```

...
<!-- Checker Configuration: module dependency to mirror + local TCP checker
-->
<check>
  <module name="mirror">
    <to addr="192.168.1.31"/>
  </module>
</check>
...

```



La dépendance des modules peut être utilisée lorsque vous déployez des modules ferme et miroir sur le même cluster SafeKit ou lorsque vous déployez des modules ferme et miroir sur deux clusters différents.

15.4 Exemple d'un flux de réplication dédié

L'attribut `ident="flow"` sur le heartbeat, permet d'identifier le flux de réplication. Pour plus d'information, voir 13.6 [page 252](#).

conf/userconfig.xml – voir 13 [page 237](#)

```
...
<heart>
  <heartbeat name="default" />
  <!-- 2nd heartbeat special for dedicated replicated network -->
  <heartbeat name="repli" ident="flow" />
</heart>
...
```

15.5 Exemples de partage de charge dans un module ferme

15.5.1 Exemple d'un load balancing TCP

Avec le fichier de configuration **userconfig.xml** suivant, vous définissez une ferme de 3 serveurs avec partage de charge et reprise sur les ports TCP 9010 (service web SafeKit), 23 (Telnet), 80 (HTTP), 443 (HTTPS), 8080 (HTTP proxy) and 389 (LDAP).



Avec HTTP et HTTPS, le partage de charge est défini sur l'adresse IP client ("`on_addr`") et non sur le port TCP client ("`on_port`"), pour assurer que le même client est toujours connecté sur le même serveur sur plusieurs connexions TCP (service à état versus service sans état ; voir la description en 1.4 [page 19](#))

conf/userconfig.xml – voir 13 [page 237](#)

```
<!DOCTYPE safe>
<safe>
  <service mode="farm">
    <farm>
      <lan name="net3" />
    </farm>
    <vip>
      <interface_list>
        <interface check="on" arpreroute="on">
          <virtual_interface type="vmac_directed">
            <virtual_addr addr="192.168.1.50" where="alias" />
          </virtual_interface>
        </interface>
      </interface_list>
      <loadbalancing_list>
        <group name="tcpservices" >
          <cluster>
            <host name="node1" power="1" />
            <host name="node2" power="1" />
          </cluster>
        </group>
      </loadbalancing_list>
    </vip>
  </service>
</safe>
```

```

        <host name="node3" power="1" />
    </cluster>
    <rule port="9010" proto="tcp" filter="on_port" />
    <rule port="23" proto="tcp" filter="on_port" />
    <rule port="80" proto="tcp" filter="on_addr" />
    <rule port="443" proto="tcp" filter="on_addr" />
    <rule port="8080" proto="tcp" filter="on_addr" />
    <rule port="389" proto="tcp" filter="on_port" />
</group>
</loadbalancing_list>
</vip>
</service>
</safe>

```

15.5.2 Exemple de load balancing UDP

Avec le fichier **userconfig.xml** suivant, vous définissez une ferme de 3 serveurs avec partage de charge et reprise sur les services UDP 53 (DNS) et 1645 (RADIUS).

conf/userconfig.xml - voir 13 [page 237](#)

```

<!DOCTYPE safe>
<safe>
<service mode="farm">
<farm>
    <lan name="net3" />
</farm>
<vip>
    <interface_list>
        <interface check="on">
            <virtual_interface type="vmac_invisible">
                <virtual_addr addr="192.168.1.50" where="alias" />
            </virtual_interface>
        </interface>
    </interface_list>
    <loadbalancing_list>
        <group name="udpservices" >
            <cluster>
                <host name="node1" power="1" />
                <host name="node2" power="1" />
                <host name="node3" power="1" />
            </cluster>
            <rule port="53" proto="udp" filter="on_ipid" />
            <rule port="1645" proto="udp" filter="on_ipid" />
        </group>
    </loadbalancing_list>
</vip>
</service>
</safe>

```



Avec **on_ipid**, le load balacing est réalisé sur le champ "IP identifier" dans l'entête IP du paquet. Le load balancing fonctionne même si le client présente la même adresse IP client et le même port en entrée.

15.5.3 Exemple d'un load balancing multi-groupes

Avec le fichier `userconfig.xml` suivant, vous définissez une ferme de 3 serveurs avec une priorité pour le trafic HTTP pour le serveur 1, HTTPS pour le serveur 2 et proxy HTTP pour le serveur 3.

`conf/userconfig.xml` - voir [13 page 237](#)

```
<!DOCTYPE safe>
<safe>
<service mode="farm">
  <farm>
    <lan name="net3" />
  </farm>
  <vip>
    <interface_list>
      <interface check="on">
        <virtual_interface type="vmac_invisible">
          <virtual_addr addr="192.168.1.50" where="alias" />
        </virtual_interface>
      </interface>
    </interface_list>
    <loadbalancing_list>
      <group name="http_service" >
        <cluster>
          <host name="node1" power="3" />
          <host name="node2" power="1" />
          <host name="node3" power="1" />
        </cluster>
        <rule port="80" proto="tcp" filter="on_addr" />
      </group>
      <group name="https_service" >
        <cluster>
          <host name="node1" power="1" />
          <host name="node2" power="3" />
          <host name="node3" power="1" />
        </cluster>
        <rule port="443" proto="tcp" filter="on_addr" />
      </group>
      <group name="httpproxy_service" >
        <cluster>
          <host name="node1" power="1" />
          <host name="node2" power="1" />
          <host name="node3" power="3" />
        </cluster>
        <rule port="8080" proto="tcp" filter="on_addr" />
      </group>
    </loadbalancing_list>
  </vip>
</service>
</safe>
```

15.6 Exemple d'un hostname virtuel avec vhost.safe

Le module de démonstration `vhost.safe` montre comment positionner un hostname virtuel (voir 13.8 [page 271](#)).

conf/userconfig.xml – voir 13 [page 237](#)

```
...
<vhost>
  <virtualhostname name="virtualname" envfile="vhostenv.cmd" />
</vhost>
...
```

En plus de cette configuration, il faut exécuter des commandes spéciales dans les scripts utilisateur. Ci-dessous l'exemple des scripts Windows. Pour Linux, se référer au `vhost.safe` livré avec la package Linux.

bin/start_prim.cmd – voir 14 [page 293](#)

```
@echo off

rem Script called on the primary server for starting application services

rem For logging into SafeKit log use:
rem "%SAFE%\safekit" printi | printe "message"

rem stdout goes into Application log
echo "Running start_prim %*"

rem Set virtual hostname
CALL "%SAFEUSERBIN%\vhostenv.cmd"

rem Next commands use the virtual hostname
FOR /F %x IN ('hostname') DO SET servername=%x
echo "hostname is "%servername%"

rem WARNING: previous virtual hostname setting is insufficient to change the
hostname for services
rem If one service needs the virtual hostname, you need also to uncomment the
rem following

rem "%SAFE%\private\bin\vhostservice" SERVICE_TO_BE_DEFINED

set res=0

rem Fill with your services start call

set res=%errorlevel%

if %res% == 0 goto end

:stop
"%SAFE%\safekit" printe "start_prim failed"

rem uncomment to stop SafeKit when critical
rem "%SAFE%\safekit" stop -i "start_prim"
:end
```

bin/stop_prim.cmd - voir 14 [page 293](#)

```
@echo off
rem Script called on the primary server for stopping application services

rem For logging into SafeKit log use:
rem "%SAFE%\safekit" printi | printe "message"

rem -----
rem
rem 2 stop modes:
rem
rem - graceful stop
rem   call standard application stop with net stop
rem
rem - force stop (%1=force)
rem   kill application's processes
rem
rem -----

rem stdout goes into Application log
echo "Running stop_prim %"

set res=0

rem Reset virtual hostname
CALL "%SAFEUSERBIN%\vhostenv.cmd"

rem Next commands use the real hostname
FOR /F %x IN ('hostname') DO SET servername=%x
echo "hostname is "%servername%

rem default: no action on forcestop
if "%1" == "force" goto end

rem Fill with your services stop call

rem If necessary, uncomment to wait for the stop of the services
rem "%SAFEBIN%\sleep" 10

if %res% == 0 goto end

"%SAFE%\safekit" printi "stop_prim failed"

:end
rem WARNING: if the virtual hostname was set for services in start_prim.cmd,
rem uncomment the following to restore the real hostname in last stop phase :

rem "%SAFE%\private\bin\vhostservice" SERVICE_TO_BE_DEFINED
```

15.7 Détection de la mort de processus avec `softerrd.safe`

Le module `softerrd.safe` est un module de démonstration de la surveillance de la mort de processus (pour plus d'information, voir 13.9 [page 273](#)).

Le module surveille la présence des processus :

- ⇒ `mybin` et `myappli` démarrés/arrêtés sur le nœud primaire avec `start_prim/stop_prim`
- ⇒ `myotherbin` démarré/arrêté sur le nœud secondaire avec `start_second/stop_second`

La détection de l'arrêt de :

- ⇒ `mybin` provoque un `restart` du module
- ⇒ `myappli` provoque l'exécution d'un handler spécial `restart_myappli.cmd`. Ce script incrémente le compteur `maxloop` et redémarre le processus `myappli`
- ⇒ `myotherbin` provoque un `stop` du module

Les tests consistent à tuer les processus `mybin`, `myotherbin` ou `myappli` avec la commande `safekit kill`.

Ci-dessous un extrait de `softerrd.safe` pour Windows. Pour Linux, regardez celui livré avec le package Linux.

`conf/userconfig.xml` - voir 13 [page 237](#)

```
...
    <errd>
      <proc name="mybin.exe" atleast="1" action="restart" class="prim"/>
      <proc name="myotherbin.exe" atleast="1" action="stop"
class="second"/>
      <proc name="myappli.exe" atleast="1" action="restart_myappli"
class="myappli"/>
    </errd>
...
```

`bin/start_prim.cmd` - voir 14 [page 293](#)

Noter l'appel à `%SAFE%\safekit errd enable myappli` pour commencer la surveillance des processus avec `class="myappli"`

```
@echo off

%SAFE%\safekit printi "start mybin"
start %SAFEUSERBIN%\mybin.exe 10000000

%SAFE%\safekit printi "start myappli"
start %SAFEUSERBIN%\myappli.exe 10000000
%SAFE%\safekit errd enable myappli

:end
```

bin/stop_prim.cmd - voir 14 [page 293](#)

Noter l'appel à %SAFE%\safekit errd disable myappli pour arrêter la surveillance des processus avec class="myappli"

```
@echo off
rem default: no action on forcestop
if "%1" == "force" goto end

%SAFE%\safekit printi "stop mybin"
%SAFE%\safekit kill -level="terminate" -name="mybin.exe"

%SAFE%\safekit printi "stop myappli"
%SAFE%\safekit errd disable myappli
%SAFE%\safekit kill -level="terminate" -name="restart_myappli.cmd"
%SAFE%\safekit kill -level="terminate" -name="myappli.exe"

:end
```

bin/restart_myappli.cmd

Noter l'incrémentation du compteur de rebouclage et l'arrêt du module quand maxloop est atteint

```
@echo off

rem Template for script called by errd on error detection instead of standard
restart
%SAFE%\safekit printi "restart_myappli"

rem first disable monitoring of the application
%SAFE%\safekit errd disable myappli

rem increment loop counter
%SAFE%\safekit incloop -i "restart_myappli"
if %errorlevel% == 0 goto next
rem max loop reached
%SAFE%\safekit stop -i "restart_myappli"
%SAFEBIN%\exitcode 0

:next
rem max loop not reached : go on restarting the application
%SAFE%\safekit printi "Restart myappli"
%SAFE%\safekit kill -level="terminate" -name="myappli.exe"
start %SAFEUSERBIN%\myappli.exe 1000000

rem finally, enable monitoring of the application
%SAFE%\safekit errd enable myappli
```

15.8 Exemple d'un checker TCP

Le checker tcp teste le service web Apache (pour plus d'information, voir 13.11 [page 280](#)). L'action par défaut quand le service TCP est down est de redémarrer le module (voir 13.18.5 [page 291](#)).

conf/userconfig.xml - voir 13 [page 237](#)

```
...
  <check>
    <tcp
      ident="Apache_80"
      when="both"
    >
      <to
        addr="172.21.10.5"
        port="80"
        interval="120"
        timeout="5"
      />
    </tcp>
  </check>
...
```

15.9 Exemple d'un checker ping

Le checker ping suivant teste un routeur d'adresse IP 192.168.1.1 (pour plus d'information, voir 13.12 [page 281](#)). L'action par défaut lorsque le routeur est down et de stopper localement le module et d'attendre que le ping redevienne up (voir 13.18.5 [page 291](#)).

conf/userconfig.xml - voir 13 [page 237](#)

```
...
<check >
  <ping ident="router">
    <to addr="192.168.1.1"/>
  </ping>
</check>
...
```

15.10 Exemple d'un checker d'interface réseau

Ci-dessous, l'exemple d'une configuration d'interface checker générée automatiquement lorsque l'option `<interface check="on">` est définie (voir 13.5 [page 245](#)). Dans le fichier de configuration, l'adresse virtuelle est définie comme suit :

conf/userconfig.xml - voir 13 [page 237](#)

```
...
<vip>
  <interface_list>
    <interface check="on">
      <real_interface>
        <virtual_addr addr="192.168.1.32" where="one_side_alias"/>
      </real_interface>
    </interface>
  </interface_list>
...
```

```
</vip>  
...
```

L'action par défaut lorsque l'interface est down et de stopper localement le module et d'attendre que l'interface redevienne up (voir 13.18.5 [page 291](#)).

⇒ Configuration générée en Windows

```
<check>  
  <intf when="pre" ident="192.168.1.0"  
    intf="{8358A0EE-2F3F-4FEE-A33B-EDC406C0C858}">  
    <to local_addr="192.168.1.228"/>  
  </intf>  
</check>
```

Où {8358A0EE-2F3F-4FEE-A33B-EDC406C0C858} est l'identité de l'interface réseau avec le réseau 192.168.1.0 et avec 192.168.228 comme première adresse IP (safekit -r vip_if_ctrl -L).

⇒ Configuration générée en Linux

```
<check>  
  <intf when="pre" ident="192.168.1.0" intf="eth2">  
    <to local_addr="192.168.1.20"/>  
  </intf>  
</check>
```

Où eth2 est l'identité de l'interface réseau avec le réseau 192.168.1.0 et avec 192.168.228 comme première adresse IP (ifconfig -a ou ip addr show).

Pour plus de détails, voir 13.13 [page 282](#).

15.11 Exemple d'IP checker

Ci-dessous, l'exemple d'une configuration d'IP checker générée automatiquement lorsque l'option `<virtual_addr check="on" ...>` est positionnée (voir 13.5 [page 245](#)). Dans le fichier `userconfig.xml`, l'adresse virtuelle est définie comme suit :

conf/userconfig.xml - voir 13 [page 237](#)

```
...  
<vip>  
  <interface_list>  
    <interface check="on" arpreroute="on">  
      <real_interface>  
        <virtual_addr addr="192.168.1.99" where="one_side_alias" check="on"/>  
      </real_interface>  
    </interface>  
  </interface_list>  
</vip>  
...
```

L'action par défaut lorsque le checker détecte que l'adresse n'est plus configurée est d'exécuter un stopstart du module (voir 13.18.5 [page 291](#)).

⇒ Configuration générée

```
<check>
<ip ident="192.168.1.99" when="prim">
<to addr="192.168.1.99"/>
</ip>
</check>
```

Pour plus de détails, voir 13.14 [page 283](#).

15.12 Exemple d'un checker personnalisé avec `customchecker.safe`

Le module `customchecker.safe` est un module de démonstration d'un checker personnalisé (pour plus d'information, voir section 13.15 [page 285](#)).

- ⇒ Le checker teste la présence d'un fichier sur le serveur primaire (`when="prim"`). La ressource associée s'appelle `custom.checkfile` (`ident="checkfile"`). Elle est affectée à `up` quand le fichier est présent à `down` sinon.
- ⇒ La règle de failover associée (configuré dans `<failover>`), nommée `custom_failure`, provoque le `restart` du module si le fichier est absent (voir 13.18.5 [page 291](#)). Cet exemple peut servir de base pour l'écriture de votre propre checker.

conf/userconfig.xml - voir 13 [page 237](#)

```
...
    <check>
      <custom ident="checkfile" exec="checker.ps1"
arg="c:\safekit\checkfile" when="prim"/>
    </check>
    <user>
    </user>
    <failover>
      <![CDATA[
        custom_failure:
        if( custom.checkfile == down ) then restart();
      ]]>
    </failover>
  ...
```

bin/checker.ps1

Noter l'appel à `safekit set -r custom.checkfile -m AM` pour affecter l'état de la ressource (`up` or `down`)

```
param([Parameter(Mandatory = $true, ValueFromPipeLine = $true,
position=1)][String]$ModName,
      [Parameter(Mandatory = $true, ValueFromPipeLine = $true,
position=2)][String]$RName,
      [Parameter(Mandatory = $true, ValueFromPipeLine = $true,
position=3)][String]$Arg1Value,
      [Parameter(Mandatory = $false, ValueFromPipeLine = $false,
position=4)][String]$Grace="2",
      [Parameter(Mandatory = $false, ValueFromPipeLine = $false,
position=5)][String] $Period="5")
```

```
)
# return up on success | down on failure
Function test([String]$Arg1Value)
{
    $res="down"
    # Replace the following by your test
    if (Test-Path "$Arg1Value")
    {
        $res="up"
    }
    return $res
}

$customchecker=$MyInvocation.MyCommand.Name
$safekit="$env:SAFE/safekit.exe"
$safebin="$env:SAFEBIN"
$gracecount=0
$prevrstate="unknown"
# wait a little
Start-Sleep $Period

while ($true){
    Start-Sleep $Period
    $rstate = test($Arg1Value)
    if($rstate -eq "down"){
        $gracecount+=1
    }else{
        $gracecount = 0
        if($prevrstate -ne $rstate){
            & $safekit set -r "$RName" -v $rstate -i
            $customchecker -m $ModName
            $prevrstate = $rstate
        }
    }
    if($gracecount -ge $Grace){
        if($prevrstate -ne $rstate){
            & $safekit set -r "$RName" -v $rstate -i
            $customchecker -m $ModName
            $prevrstate = $rstate
        }
        $gracecount = 0
    }
}
```

L'exécutable associé au checker est automatiquement appelé avec au moins 2 arguments :

- ⇒ Le 1^{er} argument est le nom module
- ⇒ Le 2^{ème} est le nom de la ressource à affecter

Si la configuration <custom> contient l'attribut `arg`, sa valeur est passée comme arguments suivants.

Le script `checker` est écrit en respectant les précautions suivantes :

- ⇒ La ressource n'est affectée que si sa valeur a changé

- ⇒ Quand la ressource est `down`, le checker consolide cet état (`grace` fois) avant de l'affecter. Cela peut permettre d'éviter de fausses détections d'erreur.



A chaque fois que vous modifiez le fichier de configuration ou le script, vous devez réappliquer la nouvelle configuration.

15.13 Exemple d'un checker de module avec `leader.safe` et `follower.safe`

Cet exemple présente 2 modules de démonstration : `leader.safe` et `follower.safe`.

- ⇒ Le module `leader` définit les ressources partagées par tous les `followers` comme l'adresse IP virtuelle ou les répertoires répliqués.
- ⇒ Les modules `followers` contiennent le démarrage et l'arrêt de plusieurs applications isolées dans différents modules. Chaque module `follower` peut être arrêté et démarré indépendamment sans que les autres modules soient arrêtés et sans déconfigurer les ressources du module `leader`.

Le module `leader` est un miroir : il inclut dans ses scripts le démarrage/arrêt des modules `followers`.

Chaque module `follower` est un module `light` avec les scripts de démarrage/arrêt d'une application et la détection d'erreur. Chaque module `follower` est dépendant des défaillances du module `leader` avec le checker de module suivant :

`follower/conf/userconfig.xml` - voir 13 [page](#) 237

```
...
<check>
  <module name="leader"/>
</check>
...
```

Il s'agit d'une configuration synthétique à la place de :

```
...
<check>
  <module name="leader">
    <to addr="127.0.0.1" port="9010"/>
  </module>
</check>
...
```



Si vous avez modifié le port d'écoute du service web de SafeKit (voir 10.6 [page](#) 170), remplacez la configuration synthétique par la configuration complète et changez la valeur du port.

16.Cluster SafeKit dans le cloud

- ⇒ 16.1 « Cluster SafeKit dans Amazon AWS » [page 319](#)
- ⇒ 16.2 « Cluster SafeKit dans Microsoft Azure » [page 325](#)
- ⇒ 16.3 « Cluster SafeKit dans Google GCP » [page 330](#)

Vous pouvez installer, configurer et administrer des modules SafeKit qui s'exécutent sur des serveurs virtuels dans les clouds Microsoft Azure, Amazon AWS et Google GCP plutôt que sur des serveurs physiques sur site. Cela nécessite un minimum de paramétrage du cloud et/ou des serveurs, en particulier pour mettre en œuvre une adresse IP virtuelle. Ces paramétrages sont automatiquement réalisés par les modèles AWS CloudFormation et Azure Resource. Ces modèles offrent une solution très rapide et simple pour installer et préconfigurer un cluster SafeKit dans les clouds AWS et Azure.

Pour une mise en œuvre rapide, se référer à :

- ⇒ [cluster miroir dans AWS](#) ou [cluster ferme dans AWS](#)
- ⇒ [cluster miroir dans Azure](#) ou [cluster ferme dans Azure](#)
- ⇒ [cluster miroir dans GCP](#) ou [cluster ferme dans GCP](#)

16.1 Cluster SafeKit dans Amazon AWS

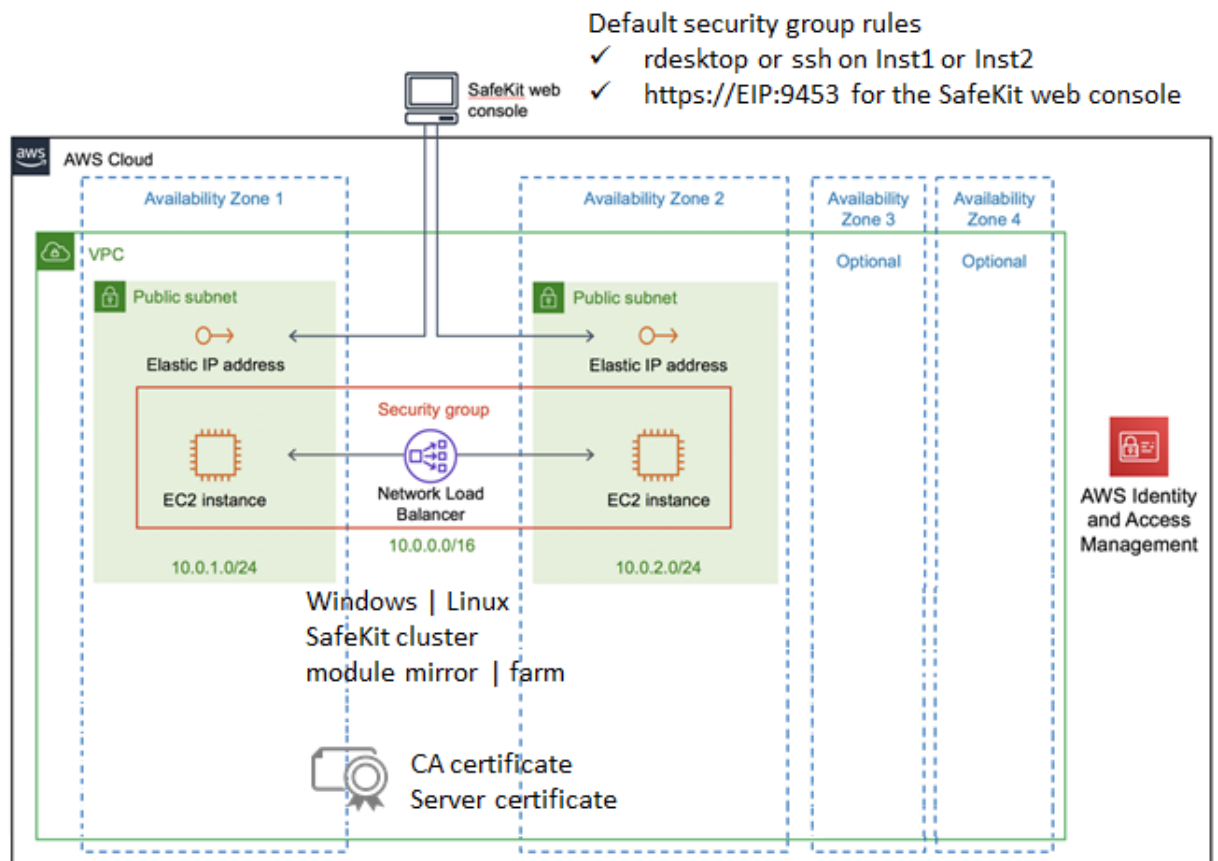
Dans ce qui suit, nous supposons que vous êtes familiers avec :

- ⇒ Amazon Elastic Compute Cloud (Amazon EC2) qui offre des capacités de calcul dans le cloud Amazon Web Services (AWS). Pour plus d'informations sur les fonctionnalités d'Amazon EC2, consultez la page du [produit Amazon EC2](#).
- ⇒ AWS CloudFormation qui aide à déployer des instances et des applications dans Amazon EC2. Cela permet de gagner beaucoup de temps et d'efforts d'installation : le temps libéré pour la gestion des ressources EC2 peut être exploité pour se consacrer aux applications qui s'exécutent dans AWS.

16.1.1 Installer un cluster SafeKit avec le modèle AWS CloudFormation pour SafeKit

SafeKit fournit un modèle AWS CloudFormation pour AWS QuickStart qui constitue un moyen simple et rapide pour implémenter un cluster SafeKit. SafeKit offre 2 modèles :

- ⇒ un modèle pour déployer un cluster miroir, avec des paramètres spécifiques décrits en 16.1.3 [page 322](#)
- ⇒ un modèle pour déployer un cluster ferme, avec des paramètres spécifiques décrits en 16.1.4 [page 323](#)



- ⇒ il déploie deux instances EC2 (jusqu'à quatre) dans la même région mais réparties sur plusieurs zones de disponibilité. Vous pouvez choisir le type d'instance et le système d'exploitation (Windows 2016 ou CentOS 7)
- ⇒ il configure un réseau virtuel (virtual private cloud, VPC)
 - ✓ il configure une adresse publique (Elastic IP, EIP) et une adresse privée pour chaque instance
 - ✓ il configure les groupes de sécurité AWS pour autoriser uniquement les connexions à distance de l'administrateur (bureau à distance pour Windows, ssh pour Linux) et de la console Web SafeKit (https://EIP:9453).
- ⇒ il configure un load-balancer AWS pour la mise en œuvre de l'IP virtuelle d'un module miroir ou ferme selon le modèle choisi
- ⇒ il exécute toutes les opérations pour que SafeKit soit prêt à l'emploi:
 - ✓ il installe le package SafeKit
 - ✓ il renseigne la configuration du cluster SafeKit et l'applique à tous les nœuds (voir 12 page 231).
 - ✓ il applique la configuration HTTPS pour sécuriser la console Web SafeKit (voir 11.6.1 page 223)
 - ✓ il installe, configure et démarre un module miroir ou ferme selon le modèle choisi

À la fin du déploiement du modèle AWS CloudFormation pour SafeKit, il vous suffit de connecter un navigateur web à l'URL indiquée dans le résultat de l'application du modèle. Cet URL permet d'accéder à l'assistant de configuration décrit en 11.4.3 [page 202](#). Appliquer la procédure décrite pour ensuite utiliser la console web SafeKit sécurisée depuis votre navigateur (connecté à `https://EIP:9453`, où EIP correspond à l'adresse publique d'un nœud du cluster). Vous pouvez ensuite exploiter SafeKit comme sur une installation sur site en installant, configurant et administrant un module miroir ou ferme dans le cloud AWS.

16.1.2 Installer un cluster SafeKit en dehors du modèle AWS CloudFormation pour SafeKit

Vous pouvez mettre en œuvre SafeKit dans des instances AWS créées en dehors du modèle AWS CloudFormation pour SafeKit. Dans ce cas, avant de pouvoir mettre en œuvre un module SafeKit, l'administrateur doit effectuer manuellement les paramétrages d'AWS, des instances et de SafeKit. Il faut en plus des configurations spécifiques en fonction du module que vous souhaitez mettre en œuvre :

- ⇒ pour un cluster miroir, voir 16.1.3 [page 322](#)
- ⇒ pour un cluster ferme, voir 16.1.4 [page 323](#)

Paramétrage d'AWS

Il faut paramétrer AWS pour :

- ⇒ associer des adresses publiques à chaque nœud si vous souhaitez les administrer avec la console web SafeKit depuis internet
- ⇒ configurer les groupes de sécurité associés au(x) réseau(x) pour ouvrir les communications du framework SafeKit et de la console web SafeKit. Les ports à ouvrir sont décrits en 10.3.3.2 [page 163](#)
- ⇒ utiliser un réseau à large bande passante et à faible latence si la réplication temps réel est utilisée dans un module miroir

Paramétrage des instances

Sur chaque instance, il faut en plus :

- ⇒ installer le package SafeKit
- ⇒ appliquer la configuration HTTPS pour sécuriser la console Web SafeKit (voir 11.6.1 [page 223](#))

Paramétrage de SafeKit

Enfin il faut saisir la configuration du cluster SafeKit et l'appliquer à tous les nœuds (voir 12 [page 231](#)). Pour chaque réseau, il peut être spécifié s'il peut être utilisé par la console et/ou le framework. Par défaut, un réseau peut être utilisé à la fois par la console et le framework (`console="on" framework="on"`).

Dans le cas du réseau publique accessible depuis internet, il est préférable de ne pas l'utiliser pour les communications du framework SafeKit mais uniquement pour la console (`console="on" framework="off"`).

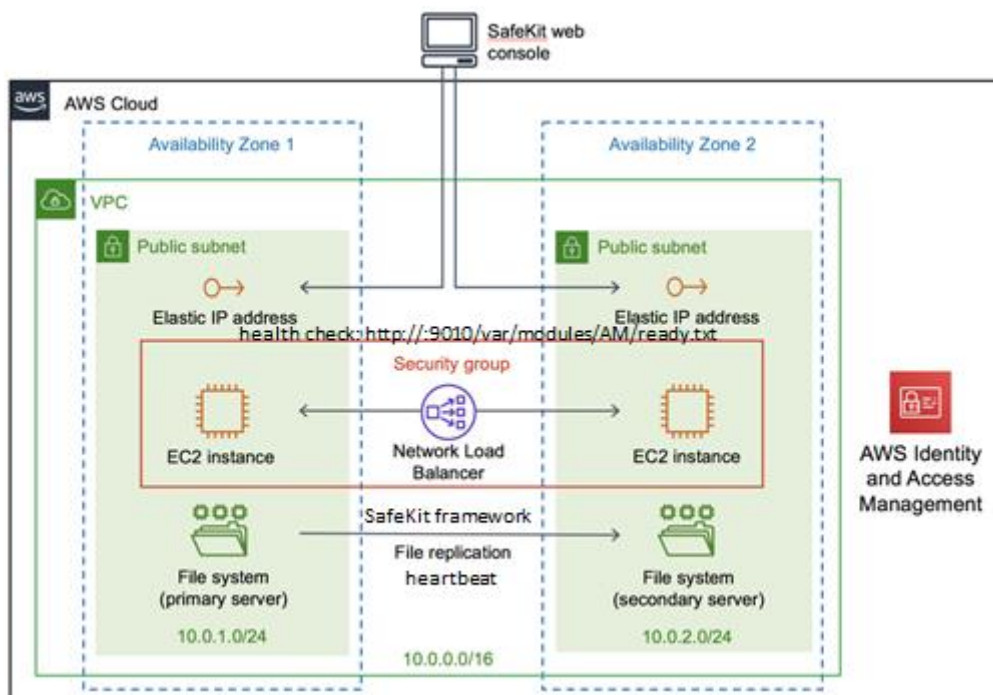
Par exemple, le fichier de configuration du cluster SafeKit serait :

```
<cluster>
<lans>
<lan name="Public" console="on" framework="off">
<node name="Server1" addr="18.214.97.59"/>
<node name="Server2" addr="52.5.205.73"/>
</lan>
<lan name="Private" console="on" framework="on">
<node name="Server1" addr="10.0.11.10"/>
<node name="Server2" addr="10.0.12.10"/>
</lan>
</lans>
</cluster>
```

Le premier `lan` défini est utilisé uniquement pour la console web de SafeKit ; le deuxième est aussi utilisé pour les communications du framework SafeKit entre les nœuds du cluster.

16.1.3 Cluster miroir dans AWS

Les fonctionnalités du module miroir sont opérationnelles dans le cloud AWS (réplication de fichiers temps réel, reprise sur panne, détection de mort de processus, checkers, ...), à l'exception du basculement d'adresse IP virtuelle. A la place, vous pouvez configurer un module miroir sur le cluster et utiliser le produit Elastic Load Balancing d'AWS (voir [Produits Elastic Load Balancing](#) dans AWS) en le configurant de façon à diriger tout le trafic vers le nœud primaire. L'adresse IP et/ou un nom DNS associés au load balancer, jouent le rôle d'IP virtuelle. Le modèle AWS CloudFormation pour SafeKit configure le load balancer et effectue tous les paramétrages nécessaires. Il vous suffit de modifier la règle de load-balancing et le groupe de sécurité réseau pour votre application.



Si vous mettez en place le module miroir en dehors du modèle AWS CloudFormation pour SafeKit, vous devez configurer vous-même le load balancer et le groupe de sécurité AWS.

Pour le load balancer, vous devez :

- ⇒ spécifier les règles pour votre application
- ⇒ définir dans le groupe cible du trafic les nœuds du cluster SafeKit
- ⇒ définir le test de `vérification de l'état`. Ce test permet de vérifier si l'instance est dans un état sain ou non

Le load balancer achemine le trafic uniquement vers les instances saines. Il reroute le trafic vers l'instance lorsque celle-ci a été restaurée dans un état sain.

SafeKit fournit un testeur de `vérification de l'état` pour chaque module. Vous devez configurer le test de vérification dans le load balancer avec :

- ⇒ le protocole HTTP
- ⇒ le port 9010, port du service web de SafeKit
- ⇒ l'URL `/var/modules/AM/ready.txt` où AM est le nom du module

Pour un module miroir, le test retourne :

- ⇒ OK, qui signifie l'instance est saine, quand le module est dans l'état  PRIM (vert) ou  ALONE (vert)
- ⇒ NOT FOUND, qui signifie que l'instance est hors service, dans tous les autres états

Le groupe de sécurité doit au minimum être configuré pour autoriser les communications pour les protocoles et ports :

- ⇒ UDP - 4800 pour le service `safeadmin` (entre les nœuds du cluster SafeKit)
- ⇒ UDP - 8888 pour le heartbeat du module (entre les nœuds du cluster SafeKit)
- ⇒ TCP - 5600 pour la réplication temps réelle du module (entre les nœuds du cluster SafeKit)
- ⇒ TCP - 9010 pour la console web SafeKit en HTTP
TCP - 9453 pour la console web SafeKit en HTTPS
- ⇒ TCP - 9001 pour configurer la console web SafeKit en HTTPS



Les valeurs de ports du module dépendent de son id (voir 10.3.3.2 [page 163](#)). Les valeurs ci-dessus sont pour le premier module installé.

16.1.4 Cluster ferme dans AWS

La plupart des fonctionnalités du module ferme sont opérationnelles dans le cloud AWS (détection de mort de processus, checker, ...), à l'exception du partage de charge sur l'adresse IP virtuelle. A la place, vous pouvez configurer un module ferme sur le cluster et utiliser le produit Elastic Load Balancing d'AWS (voir [Produits Elastic Load Balancing](#) dans AWS). L'adresse IP et/ou un nom DNS associés au load balancer, jouent le rôle d'IP virtuelle. Le modèle AWS CloudFormation pour SafeKit configure le load balancer et

- ⇒ TCP – 9010 pour la console web SafeKit en HTTP
- TCP – 9453 pour la console web SafeKit en HTTPS
- ⇒ TCP – 9001 pour configurer la console web SafeKit en HTTPS

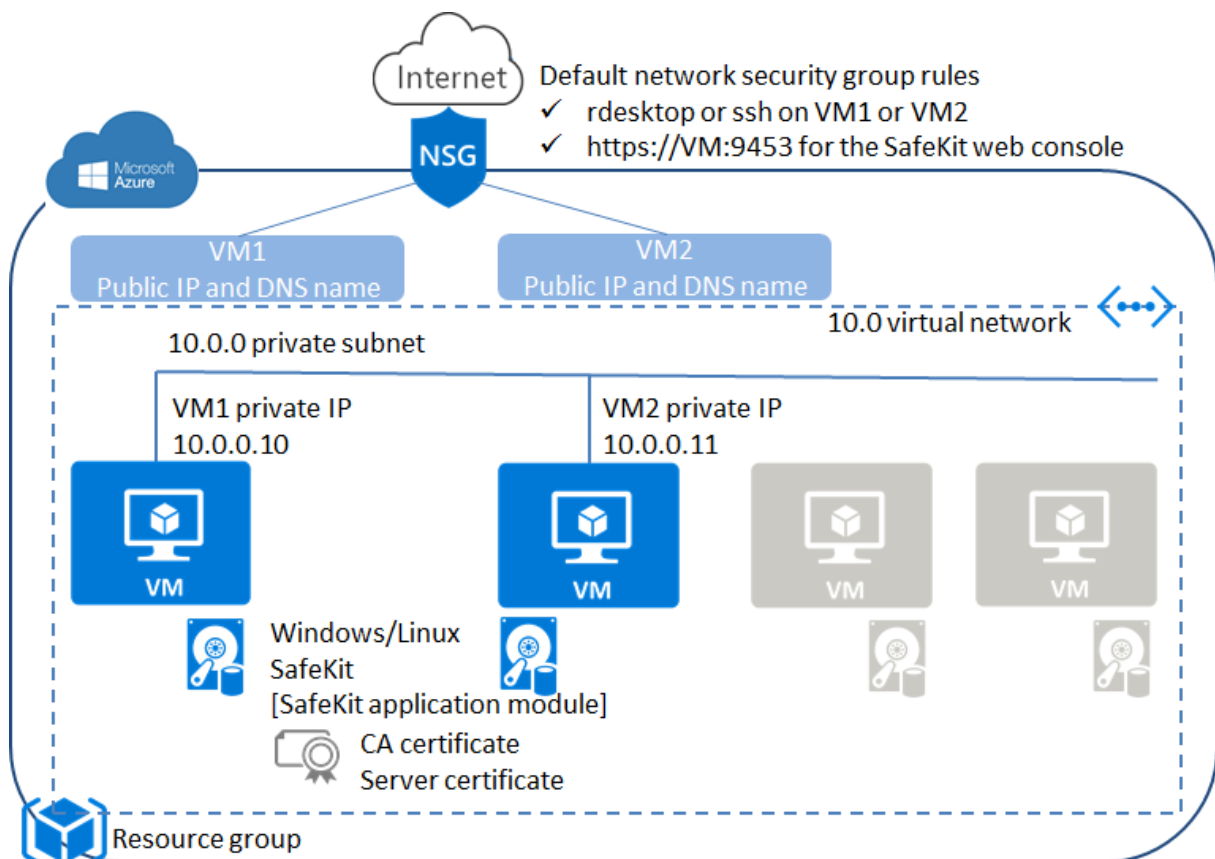
16.2 Cluster SafeKit dans Microsoft Azure

Dans la suite, nous supposons que vous êtes familiers avec Microsoft Azure, qui est un service de cloud computing créé par Microsoft pour créer, tester, déployer et gérer des applications et des services à travers un réseau mondial de centres de données Microsoft. Pour plus d'informations sur les fonctionnalités et l'utilisation d'Azure, voir le [portail Microsoft Azure](#).

16.2.1 Installer un cluster SafeKit avec le modèle Azure Resource pour SafeKit

SafeKit fournit un modèle Azure Resource qui constitue un moyen simple et rapide pour implémenter un cluster SafeKit. SafeKit offre 2 modèles :

- ⇒ un modèle pour déployer un cluster miroir, avec des paramètres spécifiques décrits en 16.2.3 [page 327](#)
- ⇒ un modèle pour déployer un cluster ferme, avec des paramètres spécifiques décrits en 16.2.4 [page 329](#)



- ⇒ il déploie un groupe de ressources qui contient toutes les ressources nécessaires à l'implémentation du cluster SafeKit dans un emplacement
- ⇒ le groupe de ressources est composé de deux machines virtuelles (jusqu'à quatre). Vous pouvez choisir le système d'exploitation (Windows 2016, Linux

CentOS 7). Chaque machine virtuelle s'exécute dans une zone de disponibilité différente et peut être accédée depuis internet via une adresse IP publique ou un nom DNS

- ⇒ il configure un réseau privé pour les communications du framework SafeKit
- ⇒ il configure le groupe de sécurité réseau pour autoriser uniquement les connexions à distance de l'administrateur (bureau à distance pour Windows, ssh pour Linux) et de la console Web SafeKit (<https://DNS:9453>)
- ⇒ il configure un load-balancer Azure pour la mise en œuvre de l'IP virtuelle du module miroir ou ferme selon le modèle choisi
- ⇒ il exécute toutes les opérations pour que le module SafeKit soit prêt à l'emploi:
 - ✓ il installe le package SafeKit
 - ✓ il renseigne la configuration du cluster SafeKit et l'applique à tous les nœuds (voir 12 [page 231](#)).
 - ✓ il applique la configuration HTTPS pour sécuriser la console Web SafeKit (voir 11.6.1 [page 223](#))
 - ✓ il installe, configure et démarre un module miroir ou ferme selon le modèle choisi

À la fin du déploiement du modèle Azure Resource pour SafeKit, il vous suffit de connecter un navigateur web à l'URL indiquée dans le résultat de l'application du modèle. Cet URL permet d'accéder à l'assistant de configuration décrit en 11.4.3 [page 202](#). Appliquer la procédure décrite pour ensuite utiliser la console web SafeKit sécurisée depuis votre navigateur (connecté à <https://DNS:9453>, où DNS correspond au nom DNS d'un nœud du cluster). Vous pouvez ensuite exploiter SafeKit comme sur une installation sur site en installant, configurant et administrant un module miroir ou ferme dans le cloud Azure.

16.2.2 Installer un cluster SafeKit en dehors du modèle Azure Resource pour SafeKit

Vous pouvez mettre en œuvre SafeKit dans des machines virtuelles Azure créées en dehors du modèle Azure Resource pour SafeKit. Dans ce cas, avant de pouvoir mettre en œuvre un module SafeKit, l'administrateur doit effectuer manuellement les paramétrages d'Azure, des machines virtuelles et de SafeKit. Il faut en plus des configurations spécifiques en fonction du module que vous souhaitez mettre en œuvre :

- ⇒ pour un cluster miroir, voir 16.2.3 [page 327](#)
- ⇒ pour un cluster ferme, voir 16.2.4 [page 329](#)

Paramétrage d'Azure

Il faut paramétrer Azure pour :

- ⇒ associer des adresses IP publiques (éventuellement nom DNS) à chaque machine virtuelle si vous souhaitez les administrer avec la console web SafeKit depuis internet
- ⇒ configurer, si nécessaire, le groupe de sécurité réseau pour ouvrir les communications du framework SafeKit et de la console web SafeKit. Les ports à ouvrir sont décrits en 10.3.3.2 [page 163](#)

- ⇒ utiliser un réseau à large bande passante et à faible latence si la réplication temps réel est utilisée dans un module miroir

Paramétrage des machines virtuelles

Sur chaque machine virtuelle, il faut en plus :

- ⇒ installer le package SafeKit
- ⇒ appliquer la configuration HTTPS pour sécuriser la console Web SafeKit (voir 11.6.1 [page 223](#))

Paramétrage de SafeKit

Enfin il faut saisir la configuration du cluster SafeKit et l'appliquer à tous les nœuds (voir 12 [page 231](#)). Pour chaque réseau, il peut être spécifié s'il peut être utilisé par la console et/ou le framework. Par défaut, un réseau peut être utilisé à la fois par la console et le framework (`console="on" framework="on"`). Dans le cas du réseau publique accessible depuis internet, il est préférable de ne pas l'utiliser pour les communications du framework SafeKit mais uniquement pour la console (`console="on" framework="off"`).

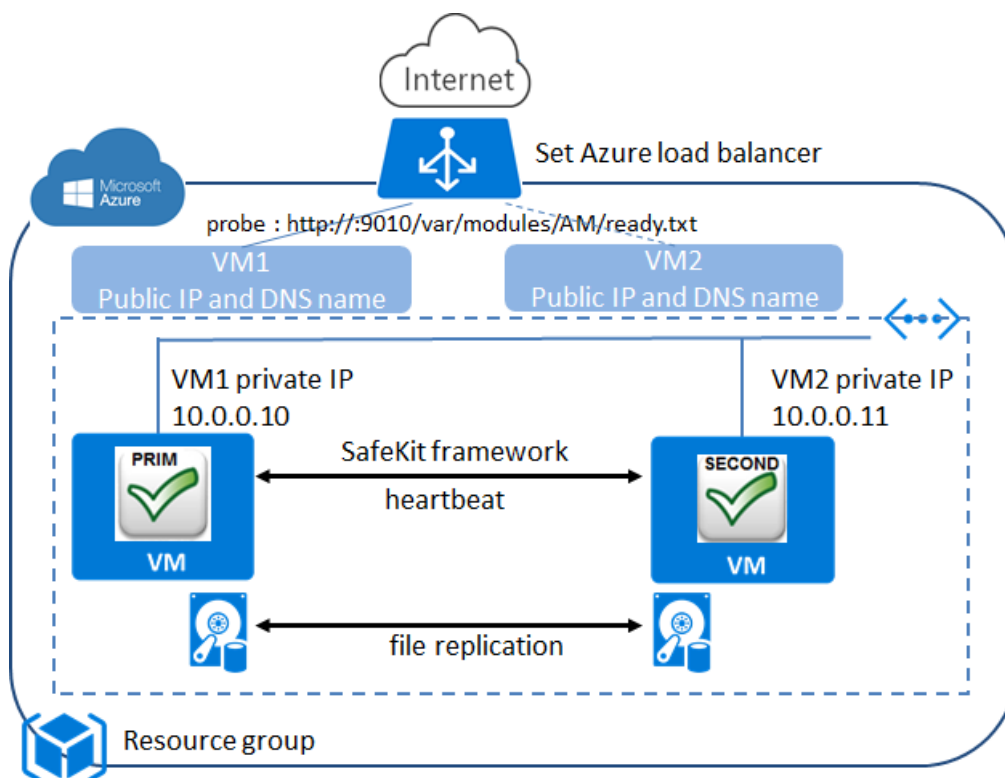
Par exemple, le fichier de configuration du cluster SafeKit serait :

```
<cluster>
<lans>
<lan name="Public" console="on" framework="off">
<node name="Server1" addr="centosazurlinvm1.westeurope.cloudapp.azure.com"/>
<node name="Server2" addr="centosazurlinvm2.westeurope.cloudapp.azure.com"/>
</lan>
<lan name="Private" console="on" framework="on">
<node name="Server1" addr="10.0.0.10"/>
<node name="Server2" addr="10.0.0.11"/>
</lan>
</lans>
</cluster>
```

Le premier `lan` défini est utilisé uniquement pour la console web de SafeKit ; le deuxième est aussi utilisé pour les communications du framework SafeKit entre les nœuds du cluster.

16.2.3 Cluster miroir dans Azure

Les fonctionnalités du module miroir sont opérationnelles dans le cloud Azure (réplication de fichiers temps réel, reprise sur panne, détection de mort de processus, checkers, ...), excepté le basculement d'adresse IP virtuelle. A la place, vous pouvez configurer un module miroir sur le cluster et utiliser load balancer proposé par Azure (voir [Load Balancer](#) dans Azure) en le configurant de façon à diriger tout le trafic vers le nœud primaire. L'adresse IP associée au load balancer, jouent le rôle d'IP virtuelle. Le modèle Azure Resource pour SafeKit configure le load balancer et effectue tous les paramétrages nécessaires. Il vous suffit de modifier la règle de load-balancing et le groupe de sécurité réseau pour votre application.



Si vous mettez en place le module miroir en dehors du modèle Azure Resource pour SafeKit, vous devez configurer vous-même le load balancer et le groupe de sécurité Azure.

Pour le load balancer, vous devez :

- ⇒ spécifier les règles pour votre application
- ⇒ définir dans le backend pool les nœuds du cluster SafeKit
- ⇒ définir une sonde d'intégrité. Cette sonde permet de vérifier si l'instance est dans un état sain ou non

Le load balancer achemine le trafic uniquement vers les instances saines. Il reroute le trafic vers l'instance lorsque celle-ci a été restaurée dans un état sain.

SafeKit fournit une sonde d'intégrité pour chaque module. Vous devez configurer la sonde d'intégrité dans le load balancer avec :

- ⇒ le protocole HTTP
- ⇒ le port 9010, port du service web de SafeKit
- ⇒ l'URL `/var/modules/AM/ready.txt` où AM est le nom du module

Pour un module miroir, le test retourne :

- ⇒ OK, qui signifie l'instance est saine, quand le module est dans l'état PRIM (vert) ou ALONE (vert)
- ⇒ NOT FOUND, qui signifie que l'instance est hors service, dans tous les autres états

Le groupe de sécurité réseau doit au minimum être configuré pour autoriser les communications pour les protocoles et ports :

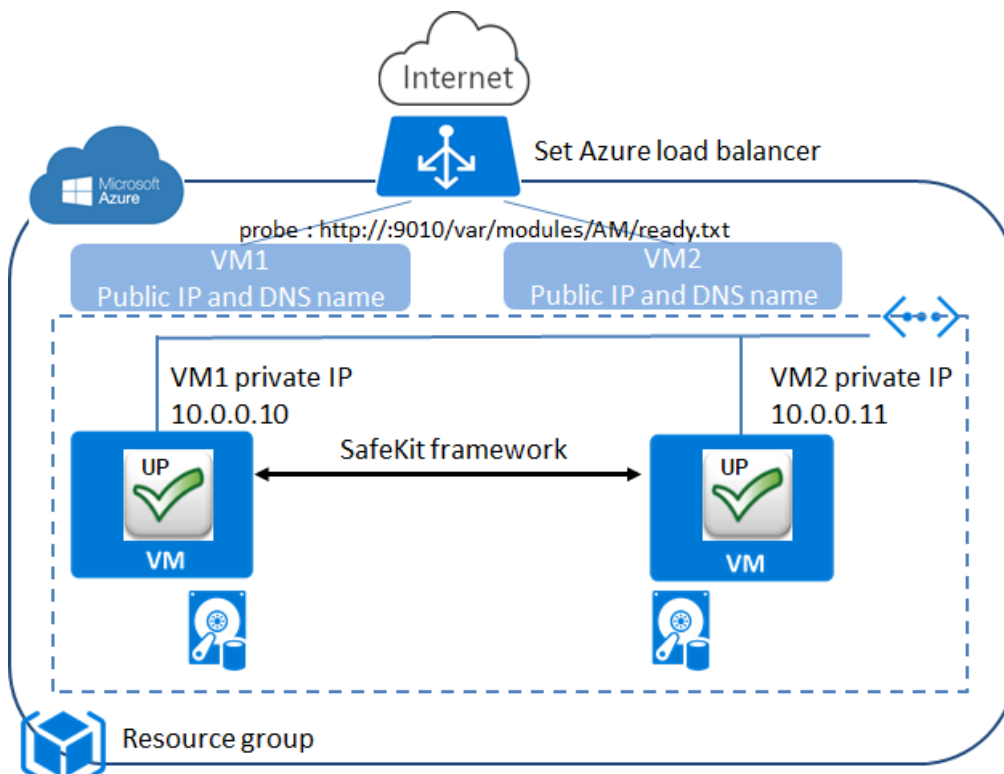
- ⇒ UDP - 4800 pour le service `safeadmin` (entre les nœuds du cluster SafeKit)
- ⇒ UDP - 8888 pour le heartbeat du module (entre les nœuds du cluster SafeKit)
- ⇒ TCP - 5600 pour la réplication temps réelle du module (entre les nœuds du cluster SafeKit)
- ⇒ TCP - 9010 pour la console web SafeKit en HTTP
- TCP - 9453 pour la console web SafeKit en HTTPS
- ⇒ TCP - 9001 pour configurer la console web SafeKit en HTTPS



Les valeurs de ports du module dépendent de son id (voir 10.3.3.2 [page 163](#)). Les valeurs ci-dessus sont pour le premier module installé.

16.2.4 Cluster ferme dans Azure

La plupart des fonctionnalités du module ferme sont opérationnelles dans le cloud Azure (détection de mort de processus, checker, ...), à l'exception du partage de charge sur l'adresse IP virtuelle. A la place, vous pouvez configurer un module ferme sur le cluster et utiliser load balancer proposé par Azure (voir [Load Balancer](#) dans Azure). L'adresse IP associée au load balancer, jouent le rôle d'IP virtuelle. Le modèle Azure Resource pour SafeKit configure le load balancer et effectue tous les paramétrages nécessaires. Il vous suffit de modifier la règle de load-balancing et le groupe de sécurité réseau pour votre application.



Si vous mettez en place le module ferme en dehors du modèle Azure Resource pour SafeKit, vous devez configurer vous-même le load balancer et le groupe de sécurité Azure.

Pour le load balancer, vous devez :

- ⇒ spécifier les règles pour votre application
- ⇒ définir dans le backend pool les nœuds du cluster SafeKit
- ⇒ définir une sonde d'intégrité. Cette sonde permet de vérifier si l'instance est dans un état sain ou non

Le load balancer achemine le trafic uniquement vers les instances saines. Il reroute le trafic vers l'instance lorsque celle-ci a été restaurée dans un état sain.

SafeKit fournit une sonde d'intégrité pour chaque module. Vous devez configurer la sonde d'intégrité dans le load balancer avec :

- ⇒ le protocole HTTP
- ⇒ le port 9010, port du service web de SafeKit
- ⇒ l'URL `/var/modules/AM/ready.txt` où AM est le nom du module

Pour un module ferme, le test retourne :

- ⇒ OK, qui signifie l'instance est saine, quand le module est dans l'état  UP (vert)
- ⇒ NOT FOUND, qui signifie que l'instance est hors service, dans tous les autres états

Le groupe de sécurité doit au minimum être configuré pour autoriser les communications pour les protocoles et ports :

- ⇒ UDP - 4800 pour le service `safeadmin` (entre les nœuds du cluster SafeKit)
- ⇒ TCP - 9010 pour la console web SafeKit en HTTP
TCP - 9453 pour la console web SafeKit en HTTPS
- ⇒ TCP - 9001 pour configurer la console web SafeKit en HTTPS

16.3 Cluster SafeKit dans Google GCP

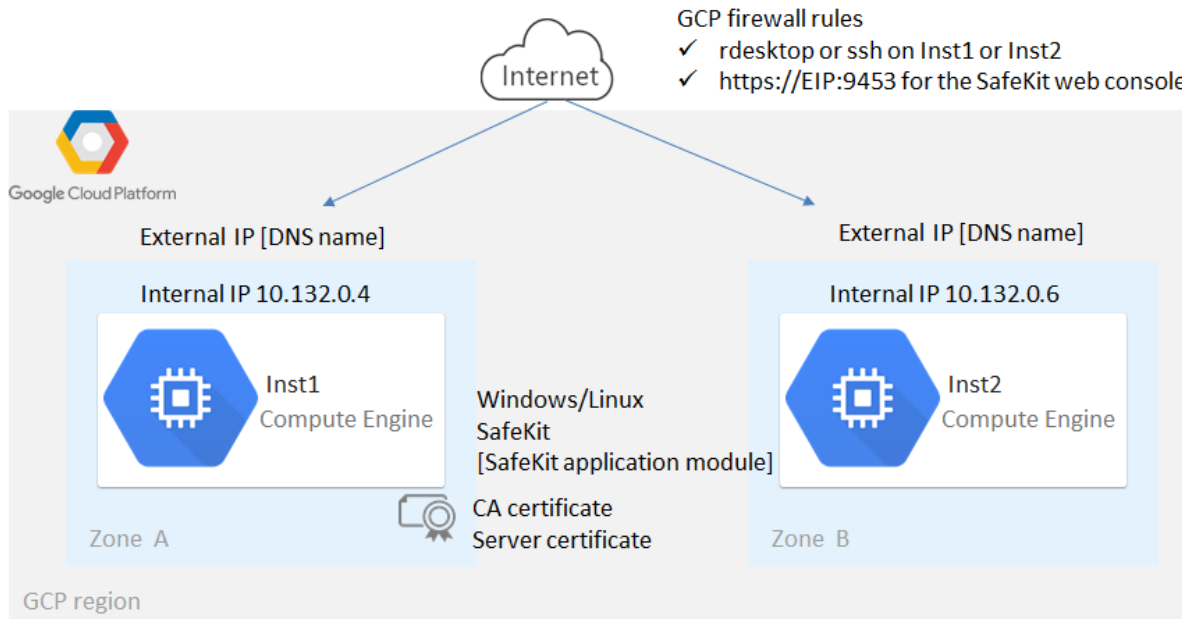
Dans la suite, nous supposons que vous êtes familiers avec Google Cloud Platform (GCP), fournisseur des machines virtuelles (VM) qui s'exécutent dans les centres de données innovants et sur le réseau de fibre optique mondial de Google. Pour plus d'informations sur ses fonctionnalités et son utilisation, voir la documentation [Google Cloud Computing](#).

16.3.1 Installer un cluster SafeKit avec la solution SafeKit pour le Google Marketplace

SafeKit fournit des solutions dans le Marketplace de Google qui constituent un moyen simple et rapide pour implémenter un cluster SafeKit. SafeKit offre 4 solutions :

- ⇒ 2 solutions pour déployer un cluster miroir (une solution pour Windows, l'autre pour Linux), avec des paramètres spécifiques décrits en 16.3.3 [page 333](#)
- ⇒ 2 solutions pour déployer un cluster ferme (une solution pour Windows, l'autre pour Linux), avec des paramètres spécifiques décrits en 16.3.4 [page 334](#)

Voir le guide [Startup Guide](#) pour une description complète du déploiement de ces solutions.



- ⇒ il déploie deux instances de VM dans la même région mais réparties sur deux zones de disponibilité. Vous pouvez choisir le type d'instance et le système d'exploitation (Windows 2019 ou CentOS 7)
- ⇒ il utilise un réseau virtuel (virtual private cloud, VPC) attaché au projet dans lequel la solution est déployée
 - ✓ il configure une adresse publique (External IP, EIP) et une adresse privée pour chaque instance
 - ✓ il configure le pare-feu pour autoriser uniquement les connexions à distance de l'administrateur (bureau à distance pour Windows, ssh pour Linux) et de la console Web SafeKit (https://EIP:9453).
- ⇒ il configure un load-balancer GCP pour la mise en œuvre de l'IP virtuelle d'un module miroir ou ferme selon la solution choisie
- ⇒ il exécute toutes les opérations pour que SafeKit soit prêt à l'emploi:
 - ✓ il inclut dans l'image de l'instance le package SafeKit
 - ✓ il renseigne la configuration du cluster SafeKit et l'applique à tous les nœuds (voir 12 [page 231](#)).
 - ✓ si HTTPS a été sélectionné au déploiement, il applique la configuration HTTPS pour sécuriser la console Web SafeKit (voir 11.6.1 [page 223](#)),
 - ✓ il installe, configure et démarre un module miroir ou ferme selon le modèle choisi

À la fin du déploiement de la solution SafeKit pour le Google Marketplace, il vous suffit de suivre les recommandations décrites dans « Suggested next steps ». Celles-ci doivent être appliquées dans le cas où vous avez choisi HTTPS pour le mode d'accès à la console web. Il s'agit de connecter un navigateur web à l'URL indiqué pour ouvrir l'assistant de configuration décrit en 11.4.3 [page 202](#). Appliquer la procédure décrite pour ensuite utiliser la console web SafeKit sécurisée depuis votre navigateur (connecté à https://EIP:9453, où EIP correspond à l'adresse publique d'un nœud du cluster). Vous

pouvez ensuite exploiter SafeKit comme pour une installation sur site en installant, configurant et administrant un module miroir ou ferme dans le cloud Google GCP.

16.3.2 Installer un cluster SafeKit en dehors de la solution SafeKit pour le Google Marketplace

Vous pouvez mettre en œuvre SafeKit dans des instances de machines virtuelles Google. Dans ce cas, l'administrateur doit effectuer manuellement les paramétrages de Google Compute Engine, des instances et de SafeKit. Il faut en plus des configurations spécifiques en fonction du module que vous souhaitez mettre en œuvre :

- ⇒ pour un cluster miroir, voir 16.3.3 [page 333](#)
- ⇒ pour un cluster ferme, voir 16.3.4 [page 334](#)

Paramétrage du GCP

Il faut paramétrer le GCP pour :

- ⇒ associer des adresses publiques (External IP), ou noms DNS, à chaque nœud si vous souhaitez les administrer avec la console web SafeKit depuis internet
- ⇒ configurer les règles de pare-feu pour le réseau Virtual Private Cloud (VPC) pour ouvrir les communications du framework SafeKit et de la console web SafeKit. Les ports à ouvrir sont décrits en 10.3.3.2 [page 163](#)
- ⇒ utiliser un réseau à large bande passante et à faible latence si la réplication temps réel est utilisée dans un module miroir

Paramétrage des instances

Sur chaque instance, il faut en plus :

- ⇒ installer le package SafeKit
- ⇒ appliquer la configuration HTTPS pour sécuriser la console Web SafeKit (voir 11.6.1 [page 223](#))

Paramétrage de SafeKit

Enfin il faut saisir la configuration du cluster SafeKit et l'appliquer à tous les nœuds (voir 12 [page 231](#)). Pour chaque réseau, il peut être spécifié s'il peut être utilisé par la console et/ou le framework. Par défaut, un réseau peut être utilisé à la fois par la console et le framework (`console="on" framework="on"`). Dans le cas du réseau publique accessible depuis internet, il est préférable de ne pas l'utiliser pour les communications du framework SafeKit mais uniquement pour la console (`console="on" framework="off"`).

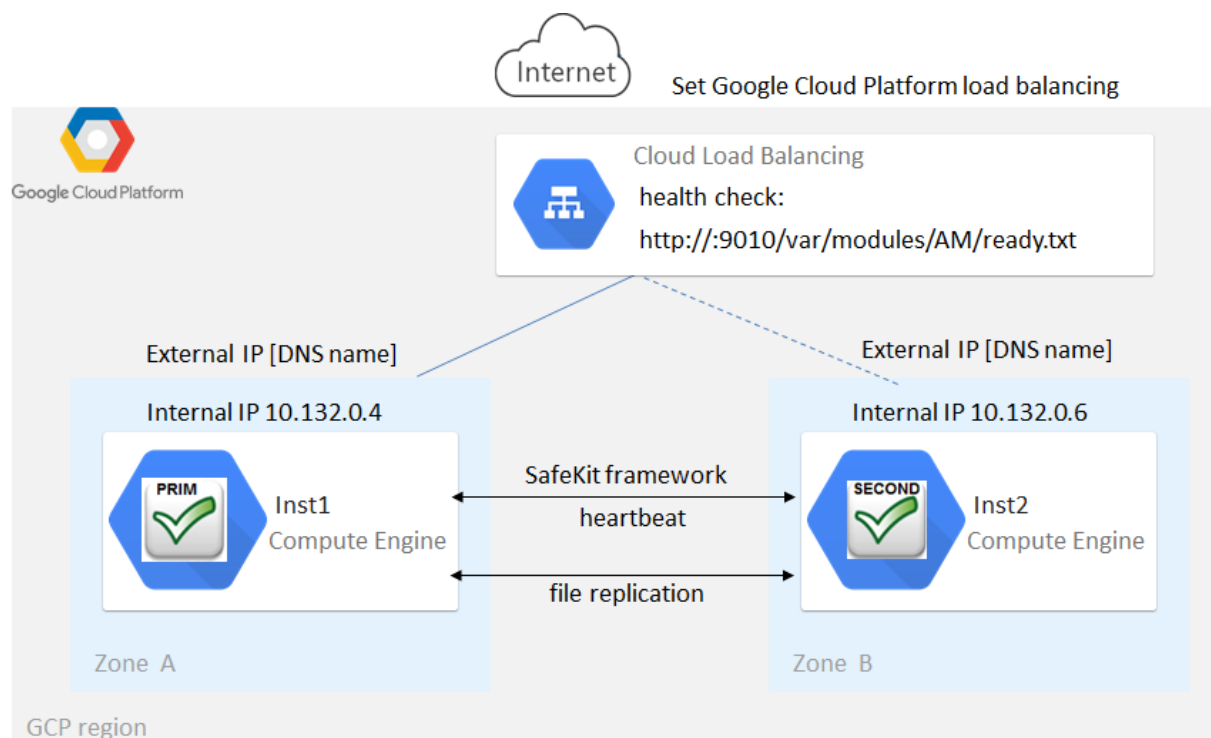
Par exemple, le fichier de configuration du cluster SafeKit serait :

```
<cluster>
<lans>
<lan name="Public" console="on" framework="off">
<node name="Server1" addr="18.214.97.59"/>
<node name="Server2" addr="52.5.205.73"/>
</lan>
<lan name="Private" console="on" framework="on">
<node name="Server1" addr="10.0.11.10"/>
<node name="Server2" addr="10.0.12.10"/>
</lan>
</lans>
</cluster>
```

Le premier `lan` défini est utilisé uniquement pour la console web de SafeKit ; le deuxième est aussi utilisé pour les communications du framework SafeKit entre les nœuds du cluster.

16.3.3 Cluster miroir dans GCP

Les fonctionnalités du module miroir sont opérationnelles dans GCP (réplication de fichiers temps réel, reprise sur panne, détection de mort de processus, checkers, ..) à l'exception du basculement d'adresse IP virtuelle. A la place, vous pouvez configurer un module miroir sur le cluster et utiliser le load balancing GCP (voir [Load Balancer](#) de GCP) en le configurant de façon à diriger tout le trafic vers le nœud primaire. L'adresse IP associée au load balancer, jouent le rôle d'IP virtuelle. La solution SafeKit pour le Google Marketplace configure le load balancer et effectue tous les paramétrages nécessaires. Il vous suffit de modifier la règle de load-balancing et le groupe de sécurité réseau pour votre application.



Si vous mettez en place le module miroir en dehors de la solution du Google Marketplace, vous devez configurer vous-même le load balancer et le pare-feu réseau de Google Cloud Platform.

Pour le load balancer, vous devez :

- ⇒ spécifier les règles pour votre application
- ⇒ définir comme cibles du trafic les nœuds du cluster SafeKit
- ⇒ définir le test de vérification de l'état. Ce test permet de vérifier si l'instance est dans un état sain ou non

Le load balancer achemine le trafic uniquement vers les instances saines. Il reroute le trafic vers l'instance lorsque celle-ci a été restaurée dans un état sain.

SafeKit fournit un testeur de vérification de l'état pour chaque module. Vous devez configurer le test de vérification dans le load balancer avec :

- ⇒ le protocole HTTP
- ⇒ le port 9010, port du service web de SafeKit
- ⇒ l'URL `/var/modules/AM/ready.txt` où AM est le nom du module

Pour un module miroir, le test retourne :

- ⇒ OK, qui signifie l'instance est saine, quand le module est dans l'état  PRIM (vert) ou  ALONE (vert)
- ⇒ NOT FOUND, qui signifie que l'instance est hors service, dans tous les autres états

Le pare-feu du réseau doit au minimum être configuré pour autoriser les communications pour les protocoles et ports :

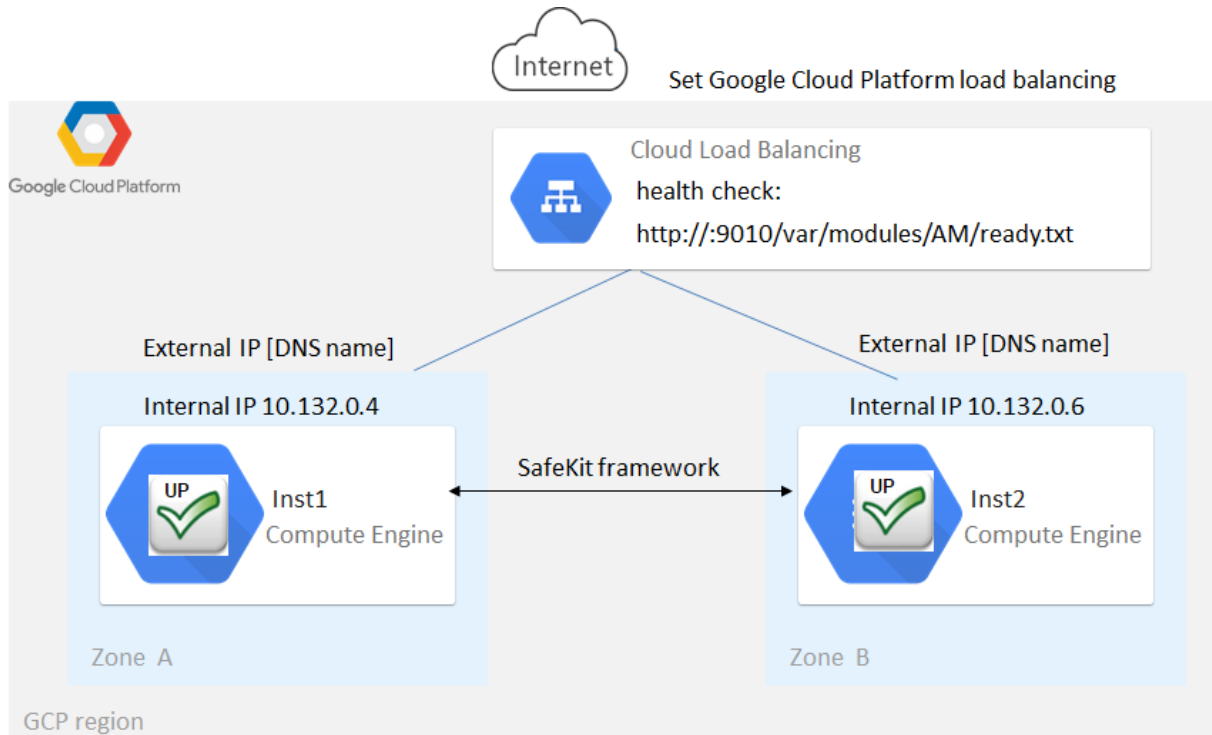
- ⇒ UDP - 4800 pour le service `safeadmin` (entre les nœuds du cluster SafeKit)
- ⇒ UDP - 8888 pour le heartbeat du module (entre les nœuds du cluster SafeKit)
- ⇒ TCP - 5600 pour la réplication temps réelle du module (entre les nœuds du cluster SafeKit)
- ⇒ TCP - 9010 pour la console web SafeKit en HTTP
TCP - 9453 pour la console web SafeKit en HTTPS
- ⇒ TCP - 9001 pour configurer la console web SafeKit en HTTPS



Les valeurs de ports du module dépendent de son id (voir 10.3.3.2 [page 163](#)). Les valeurs ci-dessus sont pour le premier module installé.

16.3.4 Cluster ferme dans GCP

La plupart des fonctionnalités du module ferme sont opérationnelles dans GCP (détection de mort de processus, checker, ...), à l'exception du partage de charge sur l'adresse IP virtuelle. A la place, vous pouvez configurer un module ferme sur le cluster et utiliser le load balancing GCP (voir [Load Balancer](#) de GCP). L'adresse IP associée au load balancer, jouent le rôle d'IP virtuelle. La solution SafeKit pour le Google Marketplace configure le load balancer et effectue tous les paramétrages nécessaires. Il vous suffit de modifier la règle de load-balancing et le groupe de sécurité réseau pour votre application.



Si vous mettez en place le module ferme en dehors de la solution du Google Marketplace, vous devez configurer vous-même le load balancer et le pare-feu réseau de Google Cloud Platform.

Pour le load balancer, vous devez :

- ⇒ spécifier les règles pour votre application
- ⇒ définir comme cibles du trafic les nœuds du cluster SafeKit
- ⇒ définir le test de vérification de l'état. Ce test permet de vérifier si l'instance est dans un état sain ou non

Le load balancer achemine le trafic uniquement vers les instances saines. Il reroute le trafic vers l'instance lorsque celle-ci a été restaurée dans un état sain.

SafeKit fournit un testeur de vérification de l'état pour chaque module. Vous devez configurer le test de vérification dans le load balancer avec :

- ⇒ le protocole HTTP
- ⇒ le port 9010, port du service web de SafeKit
- ⇒ l'URL `/var/modules/AM/ready.txt` où AM est le nom du module

Pour un module ferme, le test retourne :

- ⇒ OK, qui signifie l'instance est saine, quand le module est dans l'état UP (vert)
- ⇒ NOT FOUND, qui signifie que l'instance est hors service, dans tous les autres états

Le pare-feu du réseau doit au minimum être configuré pour autoriser les communications pour les protocoles et ports :

- ⇒ UDP - 4800 pour le service `safeadmin` (entre les nœuds du cluster SafeKit)

- ⇒ TCP – 9010 pour la console web SafeKit en HTTP
- TCP – 9453 pour la console web SafeKit en HTTPS
- ⇒ TCP – 9001 pour configurer la console web SafeKit en HTTPS

17. Logiciels tiers

SafeKit utilise les logiciels tiers listés ci-dessous. Pour les détails des licences, se référer aux liens indiqués ou aux fichiers de licence répertoriés sous `SAFE/licenses` (`SAFE=/opt/safekit` en Linux et `SAFE=C:\safekit` en Windows si `%SYSTEMDRIVE%=C:`).

libxml	http://xmlsoft.org MIT license - http://www.xmlsoft.org/FAQ.html#License Utilisé par le framework SafeKit
libxslt	http://xmlsoft.org/XSLT/ MIT license - https://gitlab.gnome.org/GNOME/libxslt/blob/master/Copyright Utilisé par le framework SafeKit
Net-SNMP	http://net-snmp.sourceforge.net BSD like and BSD license - http://www.net-snmp.org/about/license.html Utilisé par l'agent SNMP de SafeKit
HTTP server	https://httpd.apache.org/ Apache license - https://www.apache.org/licenses/LICENSE-2.0 Utilisé par la console web SafeKit, l'ancienne console java, les commandes distribuées et le module checker
APR	https://apr.apache.org/ Apache license - https://www.apache.org/licenses/LICENSE-2.0 Utilisé par le serveur Apache HTTP
PCRE	http://www.pcre.org/ BSD license - https://www.pcre.org/licence.txt Utilisé par le serveur Apache HTTP
libexpat	https://github.com/libexpat/libexpat BSD license - https://github.com/libexpat/libexpat/blob/master/expat/COPYING Utilisé par le serveur Apache HTTP
cURL	http://curl.haxx.se Curl license - https://github.com/curl/curl/blob/master/docs/LICENSE-MIXING.md Utilisé par les commandes distribuées et le module checker
cgic	ANSI C library for CGI Programming Cgic license - Credits and License Terms Utilisé par le serveur Apache HTTP de SafeKit
OpenSSL	http://www.openssl.org

	dual OpenSSL and SSLeay license - https://www.openssl.org/source/license.html
	Utilisé pour sécuriser la console web SafeKit, les commandes distribuées et le module checker
Lua	http://www.lua.org MIT license - https://www.lua.org/license.html Utilisé par la commande <code>safekit config</code> et la console web
Info-ZIP	http://info-zip.org BSD like license - http://infozip.sourceforge.net/license.html Utilisé pour le pack/unpack d'un template .safe
libnet	Packet Construction and Injection Libnet license - license Utilisé par <code>arpreroute</code> et <code>ping</code>

SafeKit utilise les logiciels tiers suivants uniquement pour la console Web :

jquery	http://jquery.org/ MIT license - https://jquery.org/license/ jQuery est une librairie javascript compacte, performante et riche en fonctionnalités
jquery-ui	http://jqueryui.com/ MIT license - https://github.com/jquery/jquery-ui/blob/master/LICENSE.txt jQuery UI étend jQuery avec des fonctionnalités de gestion de l'interface utilisateur, la définition de widgets et de thèmes
jquery-lang	https://github.com/Irrelon/jquery-lang-js MIT license - https://github.com/Irrelon/jquery-lang-js/blob/master/js/jquery-lang.js Utilisé pour la traduction des messages affichés dans la console web
jquery-ui-tabs.paging	jquery-ui-tabs-paging - MIT license Utilisé pour ajuster les tabulations en fonction de la taille de la fenêtre
codemirror	http://codemirror.net/ MIT-style license - https://github.com/codemirror/CodeMirror/blob/master/LICENSE Editeur texte développé par Marijn Haverbeke
codemirror-ui	https://github.com/jagthedrummer/codemirror-ui MIT License - https://github.com/jagthedrummer/codemirror-ui/blob/master/LICENSE

Simple interface graphique, basée sur le widget codemirror et développée par Jeremy Green

bootstrap icons <https://icons.getbootstrap.com/> - MIT license –
<https://github.com/twbs/bootstrap/blob/v4.0.0/LICENSE>

Remerciements à iTweek (<http://itweek.deviantart.com/>) pour les icônes « Knob buttons toolbar icons ».

Index des messages du journal du module

"Action ..."

"Action forcestop called by web@<IP>/SYSTEM/root", 118, 150
"Action prim called by web@<IP>/SYSTEM/root", 101, 150
"Action primforce called by SYSTEM/root", 108
"Action restart called by web@<IP>/SYSTEM/root", 77, 83, 118, 150
"Action restart|stopstart called by customscript", 97, 122, 150
"Action restart|stopstart called by errd", 90, 122, 150
"Action restart|stopstart from failover rule tcp_failure", 91, 122, 150
"Action second called by web@<IP>/SYSTEM/root", 101, 150
"Action shutdown called by SYSTEM", 80, 89, 147
"Action start called at boot time", 80, 81, 89, 147
"Action start called automatically", 90, 91, 97
"Action start called by web@<IP>/SYSTEM/root", 76, 83, 118, 150
"Action stop called by web@<IP>/SYSTEM/root", 76, 83, 118, 150
"Action stopstart called by failover-off", 106, 150
"Action stopstart called by modulecheck", 95, 150
"Action stopstart called by web@<IP>/SYSTEM/root", 118, 150
"Action stopstart from failover rule customid_failure", 97, 122, 150
"Action swap called by web@<IP>/SYSTEM/root", 77, 118, 150
"Action wait from failover rule customid_failure", 96, 121
"Action wait from failover rule tcpid_failure", 92, 121
"Action wait from failover rule degraded_server", 105
"Action wait from failover rule interface_failure", 93, 121
"Action wait from failover rule module_failure", 95, 121
"Action wait from failover rule notuptodate_server", 103, 121
"Action wait from failover rule ping_failure", 94, 121
"Action wait from failover rule splitbrain_failure", 121

Réplication et réintégration

"Copied <reintegration statistics>", 79
"Data may be inconsistent for replicated directories (stopped during reintegration)", 108
"Data may not be uptodate for replicated directories (wait for the start of the remote server)", 101, 103, 121
"If you are sure that this server has valid data, run safekit prim to force start as primary", 101, 103, 121

"If you are sure that this server has valid data, run safekit primforce to force start as primary", 108

"Reintegration ended (synchronize)", 79

"Updating directory tree from /replicated", 79

Load-balancing

"farm load: 128/256 (group FarmProto)" , 111, 86, 87

"farm membership: node1 (group FarmProto)", 86, 87

"farm membership: node1 node2 (group FarmProto)" , 111, 86, 87

"farm membership: node2 (group FarmProto)", 87

"Local state ..."

"Local state `ALONE` green", 100, 76, 82

"Local state `PRIM` green", 100, 76

"Local state `SECOND` green", 100, 76

"Local state `UP` green", 110, 111

"Local state `WAIT` red", 121, 106

"Remote state ..."

"Remote state `ALONE` green", 100, 82

"Remote state `PRIM` green", 100, 76

"Remote state `SECOND` green", 100, 76

"Remote state `UNKNOWN` grey", 81, 82

"Resource ..."

"Resource custom.id set to down by customscript", 96, 97, 121, 122

"Resource custom.id set to up by customscript", 96

"Resource heartbeat.0 set to down by `heart`", 81, 82

"Resource heartbeat.flow set to down by `heart`", 81, 82

"Resource intf.ip.0 set to down by intfcheck", 93, 121

"Resource intf.ip.0 set to up by intfcheck", 93

"Resource module.othermodule_ip set to down by modulecheck", 95, 121

"Resource module.othermodule_ip set to up by modulecheck", 95

"Resource ping.id set to down by pingcheck", 94, 121

"Resource ping.id set to up by pingcheck", 94

"Resource rfs.degraded set to up by `nfsadmin`", 105

"Resource tcp.id set to down by tcpcheck", 91, 92, 121, 122

"Resource tcp.id set to up by tcpcheck", 92

"Script ..."

"Script start_prim", 293, 76, 77, 80, 81

"Script stop_prim", 293, 76, 80, 82

"Script start_both", 293, 83, 89

"Script stop_both", 293, 83

"Transition ..."

"Transition RESTART|STOPSTART from failover rule customid_failure", 97

"Transition STOPSTART from failover-off", 106

"Transition SWAP from defaultprim", 107

"Transition SWAP from SYSTEM", 77

"Transition WAIT_TR from failover rule customid_failure", 96

"Transition WAIT_TR from failover rule interface_failure", 93

"Transition WAKEUP from failover rule Implicit_WAKEUP", 92, 93, 94, 95, 96

Autres messages

"Begin of Swap", 77, 107

"End of stop", 76, 83, 80, 89

"event atleast on proc appli.exe", 90, 122

"Failover-off configured", 106

"Previous halt unexpected", 81, 89

"Reason of failover: no heartbeat", 81

"Reason of failover: remote stop", 76, 80

"Requested prim start aborted ", 108

"Split brain recovery: exiting alone", 82

"Split brain recovery: staying alone", 82

"Stopping loop", 123, 90, 91, 92, 93, 94, 95, 96, 97, 122

"Virtual IP <ip 1.10 of mirror> set", 78

"Virtual IP <ip1.20 of farm> set", 84

Index

Architectures

miroir, ferme... - 15
cloud - 319

Installation

installation, upgrade... - 25

Console

configuration, contrôle... - 35
sécurisation (https,...) - 181

Configuration avancée

cluster.xml - 231
userconfig.xml - 237
scripts de redémarrage - 293
exemples - 301

Administration

mirror - 99
farm - 109
avancée - 157
ligne de commande - 145

Support

tests - 73
problèmes - 113
site support - 137
messages du journal - 341

Autres

table des matières - 5
logiciels tiers - 337

