

RFC 7835 : Locator/ID Separation Protocol (LISP) Threat Analysis

Stéphane Bortzmeyer
<stephane+blog@bortzmeyer.org>

Première rédaction de cet article le 30 avril 2016

Date de publication du RFC : Avril 2016

<https://www.bortzmeyer.org/7835.html>

Le protocole LISP (à ne pas confondre avec le langage de programmation) est un protocole à séparation de l'identificateur et du localisateur <<https://www.bortzmeyer.org/separation-identificateur-localisateur.html>>. Ces protocoles présentent des défis de sécurité particuliers et ce RFC est donc consacré à une étude détaillée des menaces et risques associés à LISP.

LISP est normalisé dans le RFC 9300¹. Il peut être utilisé dans un réseau interne ou bien sur l'Internet public et ce RFC étudie sa sécurité dans ce dernier cas, évidemment plus difficile.

Le RFC est en trois parties : la section 2 détaille le modèle de menaces (qui est l'attaquant), la section 3 expose les techniques que peut utiliser l'attaquant, la section 5 décrit les solutions possibles. Il s'agit bien de regarder les attaques génériques contre LISP, pas celles contre une mise en œuvre particulière de LISP, qui peut évidemment avoir ses bogues spécifiques.

Commençons donc par le modèle de menace (section 2). On suppose un attaquant situé quelque part sur l'Internet, éventuellement en plusieurs points de celui-ci. Il peut être sur le chemin des paquets ("*on path*") ou pas ("*off path*"). Le premier est évidemment le plus dangereux. S'il est attaquant actif, il peut modifier des paquets et même faire des attaques de l'Homme du Milieu. Si l'attaquant est en dehors du chemin, cela va être plus difficile pour lui. Il ne peut pas modifier les paquets échangés, il peut par contre envoyer des paquets mais, pour que ceux-ci soient acceptés, il faudra qu'ils correspondent à ce qu'attendent les parties légitimes, ce qui fait que l'attaquant devra souvent deviner des informations et compter sur la chance.

1. Pour voir le RFC de numéro NNN, <https://www.ietf.org/rfc/rfcNNN.txt>, par exemple <https://www.ietf.org/rfc/rfc9300.txt>

Autre façon de classer les attaquants, il y a ceux qui sont internes à un site LISP et ceux qui sont externes. Les attaquants de l'intérieur sont plus dangereux car la machine qu'ils contrôlent (suite à une trahison ou un piratage) est a priori jugée de confiance.

Il faut aussi distinguer les attaquants permanents ("*live*") des attaquants ponctuels ("*time-shifted*"). L'attaquant permanent est celui qui reste actif pendant toute la durée de l'attaque. Dès qu'il est neutralisé, l'attaque stoppe. L'attaquant ponctuel peut lancer une attaque, puis se retirer et les effets de l'attaque vont continuer. Ce genre d'attaques est évidemment bien plus dangereux.

L'attaquant peut viser le système de contrôle ("*control plane*") de LISP ou bien son système de transmission des données ("*data plane*"). Un exemple typique est le mécanisme de correspondance ("*mapping*") qui permet aux routeurs LISP de trouver le localisateur (le RLOC) en échange de l'identificateur (l'EID). Ce mécanisme de correspondance fait partie du système de contrôle. Si un attaquant veut détourner les paquets, il n'est pas obligé de trouver une faille dans le routage lui-même : s'il arrive à modifier les correspondances, il obtiendra le même résultat. (Comme le ferait une attaque DNS dans l'Internet classique.)

Plusieurs types d'attaques peuvent être faites contre un protocole réseau comme LISP. Il y a les attaques par rejeu, où un attaquant capture des paquets légitimes pour les rejouer intacts plus tard. Ces attaques marchent souvent, même quand les paquets sont protégés par de la cryptographie (l'attaquant n'a pas besoin de modifier les paquets, ni même de les comprendre). Il y a les manipulations de paquets (l'attaquant prend un paquet légitime et le réinjecte). Il y a la suppression complète des paquets.

Plus complexe et plus vicieux, il y a l'usurpation ("*spoofing*"). L'attaquant injecte des paquets prétendant venir d'une autre machine. C'est quelque chose qu'on voit souvent sur l'Internet <<https://www.bortzmeyer.org/usurpation-adresse-ip.html>>. LISP est une technologie de tunnel et l'usurpation peut donc porter sur deux endroits : l'adresse externe, celle que voient les routeurs IP ordinaires (c'est le RLOC), et l'adresse interne, celle du paquet encapsulé (c'est l'EID). Identifier un usurpateur est bien plus compliqué lorsque des tunnels sont en jeu !

L'attaquant n'est pas forcément un usurpateur. Il peut très bien dire la vérité sur son adresse ("*rogue attack*") par exemple parce qu'il se moque d'être identifié, ou bien parce qu'il est un zombie.

Autre type d'attaque, évidemment, les attaques par déni de service, où l'attaquant ne cherche pas à prendre le contrôle du système, mais à le paralyser (ou à le ralentir).

Parfois, la cible immédiate de l'attaque n'est pas la principale victime : dans les attaques par réflexion <<https://www.bortzmeyer.org/attaques-reflexion.html>>, un attaquant utilise un tiers pour renvoyer ses paquets vers la vraie victime. C'est surtout utilisé en combinaison avec l'amplification : lorsque la réponse est plus grosse que la question, un attaquant peut amplifier son trafic. En outre, pour LISP, le système de transmission est certainement bien plus rapide que celui de contrôle et il y a donc en théorie une possibilité d'attaquer le second avec le premier.

Et, bien sûr, il y a les attaques d'espionnage passif, où des grandes oreilles écoutent votre trafic pour vous surveiller, sans eux-mêmes envoyer de paquets. Grâce à Edward Snowden, l'ampleur de ce type d'attaques par les États est désormais bien connu (RFC 7258).

Bon, assez de théorie, comment fait-on avec LISP quand on est un méchant et qu'on veut effectuer une des attaques en question ? Rentrons désormais dans les détails techniques (qui nécessitent de connaître un peu LISP : relisez les RFC 9300, RFC 7215, RFC 6832, RFC 9301 et RFC 9302). D'abord,

le « glanage » (*"gleaning"*), la collecte passive d'informations sur les correspondances identificateur-localisateur (RFC 6830, section 6). Il peut être utilisé pour empoisonner le cache des correspondances d'un routeur : le méchant envoie un paquet LISP fabriqué, les routeurs innocents l'observent et enregistrent la correspondance entre l'EID et le RLOC dans leur cache et paf, le méchant a pu détourner le trafic futur. C'est un exemple d'attaque ponctuelle : l'attaquant envoie juste un paquet puis plus rien, mais l'effet persiste pendant toute la durée de vie de l'information dans le cache (quelques secondes, si on suit fidèlement les conseils du RFC 6830, section 6.2).

Autre exemple d'empoisonnement d'un routeur LISP avec de fausses informations, le champ LSB (*"Locator Status Bits"*), qui indique l'état (joignable ou pas) des machines situées sur le site de départ du paquet (RFC 6830, section 6.3). En envoyant un paquet usurpé avec un faux LSB, on peut tromper des routeurs innocents. On peut par exemple mettre tous les bits à zéro (qui signifie que le préfixe est injoignable par ce routeur), menant à une attaque par déni de service ou bien les mettre tous à zéro sauf un, menant à une surcharge de cet unique routeur. C'est également une attaque ponctuelle : ses effets se feront sentir même si l'attaquant est neutralisé.

Un point important de LISP est le test de la « joignabilité » (*"reachability"*, RFC 6830, section 6.3). LISP vérifie que la machine qui prétend être joignable, et par tel routeur, l'est effectivement. Un des outils pour cela est un numnique <https://www.bortzmeyer.org/nonce.html> envoyé à l'autre machine et qu'elle doit renvoyer. Un attaquant qui arriverait à tricher avec la joignabilité pourrait pousser un routeur LISP d'entrée du tunnel (un ITR) à changer de routeur de sortie (l'ETR) au profit d'un routeur qui ne marche pas. Cela réaliserait une attaque par déni de service. Même s'il n'arrive pas à faire changer de routeur, il pourrait perturber le routage en envoyant des tas de numniques différents. (Rassurez-vous, cette attaque est plus dure à réussir qu'il ne semble, voyez la section 5.)

LISP permet que plusieurs espaces d'adressage coexistent, différenciés par leur *"instance ID"* indiquée dans l'en-tête (RFC 6830, section 5.5). Elle n'est pas authentifiée et un attaquant peut donc envoyer des paquets à une autre *"instance ID"* que la sienne.

LISP permet de l'interconnexion avec des réseaux non-LISP (RFC 6832). Ces mécanismes ont des attaques très similaires à celles qu'on peut faire contre LISP lui-même. Par exemple, un paquet avec une adresse source usurpée peut être transmis sur l'Internet par la passerelle LISP.

Mais les pires attaques se situeront sans doute sur le système de correspondance (*"mapping"*). Le point central de tous les systèmes à séparation de l'identificateur et du localisateur <https://www.bortzmeyer.org/separation-identificateur-localisateur.html> est de découpler deux fonctions qui, dans IP, sont confondues. Cela suit un grand principe de l'informatique : « tout problème peut être résolu en ajoutant un degré d'indirection ». Sauf que cette séparation se paie en sécurité : il faut bien, à un moment, faire correspondre identificateur et localisateur et, là, cette nouvelle fonction, la correspondance, offre de nouvelles possibilités d'attaque.

Ainsi, les messages Map-Request (RFC 6830, section 6) peuvent être utilisés de plusieurs façons. Envoyés en masse, à un système de contrôle qui est bien moins rapide que le système de transmission, ils peuvent saturer les routeurs. Comme la réponse (Map-Reply) est plus grosse que la question, ils peuvent servir à des attaques par amplification.

Et si un attaquant réussit à fabriquer de fausses réponses (Map-Reply) et à les faire accepter ? Il devra pour cela mettre le numnique correct dans ses paquets. Comme il fait 64 bits, si on a suivi les bons principes du RFC 4086 pour le générer, un attaquant qui n'est pas sur le chemin (et doit donc deviner le numnique) n'a que peu de chances de réussir. Mais attention, le numnique n'est pas une signature : il indique juste que le paquet est bien une réponse à une question posée. Si l'attaquant peut modifier

les paquets, il peut empoisonner le cache de correspondance du routeur. Variante de cette attaque, si le méchant est un routeur légitime ("*rogue attack*"), il peut répondre, mais avec de fausses réponses. Il ne peut pas annoncer des préfixes quelconques mais il peut annoncer des préfixes plus généraux que ceux qu'il gère réellement (par exemple annoncer qu'il gère 192.0.2.0/24 alors qu'il ne contrôle que 192.0.2.0/26), ce qui lui permet de détourner le trafic.

Enfin, il y a les "*Map-Register*", ces messages envoyés par les routeurs aux serveurs de correspondance pour les informer des préfixes gérés (RFC 9301, section 4.2). Ils sont authentifiés donc, normalement, un attaquant quelconque ne peut pas les usurper. Mais il reste des attaques résiduelles comme l'annonce de préfixes plus généraux que les siens.

Au passage, un problème de sécurité qu'on oublie parfois est celui de la vie privée (section 4). Les correspondances LISP sont publiques (comme les tables BGP dans l'Internet classique) et il ne faut donc pas oublier qu'on ne peut pas participer à LISP discrètement.

Et pour finir, les solutions (section 5). Comment faire pour éviter ces menaces ? Le RFC donne des conseils très généraux (« déployer soigneusement », « appliquer les règles habituelles de sécurité ») mais aussi des indications plus précises et spécifiques à LISP. Clairement, le plan de contrôle est la partie la plus sensible. Il est donc important d'utiliser les techniques d'authentification décrites dans le RFC 9301 (voir notamment sa section 6). Des extensions de sécurisation de LISP sont en cours de développement.

D'autre part, les conseils de sécurité du RFC 9300 doivent être suivis. Par exemple, quand ce RFC écrit qu'il faut limiter le rythme des *Map-Request*, cela doit être appliqué.

L'information obtenue par examen des paquets (comme le glanage cité plus haut) n'est évidemment pas de confiance : à n'utiliser que pour le flot de données qui contenait cette information, ou alors à vérifier avant. (Voir aussi « "*How to mitigate the effect of scans on mapping systems*" » <<http://inl.info.ucl.ac.be/publications/how-mitigate-effect-scans-mapping-systems>> ».)