

FONDAMENTAUX

DE LA MODÉLISATION

DE MENACES



HÅKAN
GEIJERZ

La guerre est le domaine de l'incertitude. Les trois quarts des éléments sur lesquels se fonde l'action flottent dans le brouillard d'une incertitude plus ou moins épaisse. C'est donc dans ce domaine plus qu'en tout autre qu'une intelligence fine et pénétrante est requise, pour discerner la vérité à la seule mesure de son jugement.

Carl Von Clausewitz, *De La Guerre*, 1832

INTRODUCTION

Nous autres anarchistes ressentons une tension entre les mesures que nous estimons nécessaires à prendre contre toutes les formes de domination et les conséquences que nous sommes prêt·es à risquer dans notre poursuite de nos objectifs. La Répression, de manière explicite quand elle vient d'un État ou de manière implicite quand elle vient d'acteurs non-étatiques, contraint l'ensemble d'actions que l'on est prêt·es à entreprendre, et elle fait ça principalement en imposant un plafond bas au degré d'« extrémité » de ces actions. Les actions extrêmes sont motrices de changement, et les agents de l'État le savent, donc ils font ce qu'ils peuvent pour nous isoler de nos outils les plus puissants.

Mais il est possible de se les réapproprier.

On se tourne vers la Sécurité Opérationnelle (*OpSec*) et la Culture de la Sécurité comme principaux moyens de « nous en tirer à bon compte », ou plus précisément, d'atteindre nos objectifs avec des conséquences minimales. Ces termes sont utilisés de manière assez désinvolte, et les pratiques autour de ceux-ci sont parfois à la fois opaques et dogmatiques même quand ils contribuent à réduire les effets de la répression. Le terme « modélisation de menaces » est lancé à tort et à travers avec encore moins d'explications de ce qu'il signifie ou de comment quelqu'un pourrait la pratiquer.¹ La littérature sur le sujet tends à se focaliser sur la sécurisation de systèmes informatiques d'entreprises contre le risque de piratage, et même si ces textes peuvent nous apprendre des choses intéressantes, ils nécessitent un pas de côté mental et technique significatif pour pouvoir être appliqués à la situation d'une personne radicale lambda. Cette brochure cherche à combler ce gouffre.

C'est par la modélisation de menace que l'on peut justifier les nombreuses pratiques de l'*OpSec* et d'une culture de la sécurité (que l'on appellera « sécurité » dans le reste du texte, pour faire court). Quelqu'un pourrait nous dire de laisser notre téléphone cellulaire à la maison le temps d'une action, et ce n'est pas juste un geste pour faire joli mais plutôt parce que nos téléphones — même éteints — laissent fuiter

1. Quand à ce qui constitue même une « menace », pensez-y comme à quelque chose qu'un ennemi pourrait intentionnellement faire pour vous faire du mal, mais on verra une définition légèrement plus précise plus tard.

des données de localisation. Toutes les pratiques et normes de sécurités devraient découler d'un modèle de menace basé sur des faits, et il devrait y avoir une traçabilité entre les actions (observées ou inférées) d'adversaires et les contre-mesures déployées à l'encontre de ces actions. Certaines pratiques à présent datées sont encore pratiquées par tradition et de nouvelles pratiques devraient être adoptées mais ne le sont souvent pas car des gens ne comprennent pas le paysage de menaces dans lequel ils opèrent. C'est par le processus de modélisation de menace que ces failles sont identifiées et résolues.

Au premier abord, la modélisation de menace peut sembler être un champ d'expertise de niche, mais c'est quelque chose que l'on fait toutes dans la vie de tous les jours. Par exemple il se pourrait que vous trouviez des manières élaborées de rien branler au travail et ce processus de découverte se fait car vous essayez de maximiser votre temps de rien-branler ininterrompus jusqu'à ce que vous vous fassiez chopper. Les habitudes de votre patron, la présence de caméras de sécurité, et de collègues poudrilles sont tous des facteurs pouvant influencer vos actions, et la durée de rien-branler pourrait ainsi changer de temps à autre voire même en fonction du créneau durant lequel vous travaillez. Chaque jour on se demande ce qui pourrait se produire, quelles sont les chances que ça se produise, et ce qu'on peut faire par rapport à ces possibilités, et ensuite on ajuste notre comportement en fonction de.

Cette brochure est écrite pour être accessible à tous le monde, pas juste les gens qui s'intéressent à la sécurité. Elle part du principe que vous n'avez pas de connaissance de la sécurité dans le contexte des mouvement sociaux radicaux, mais les vétérans aguerris pourront aussi trouver une utilité à une discussion plus structurée des pratiques de sécurité qu'ils appliquent déjà. Bien que les principes généraux de la modélisation de menace qui sont discutés ici peuvent être appliqués dans beaucoup de scénarios, cette brochure se focalise spécifiquement sur le fait de résister à une répression de la part de forces de police locales, agences de renseignement et fascistes, qu'ils soient organisés ou solitaires.

Cette brochure n'est pas une autorité unique sur comment quelqu'un devrait construire un modèle de menace ou protocole de sécurité. Chaque personne est différente, chaque milieu a ses propres particularités, et chaque région a ses propres menaces uniques de flics, fascistes et autres méchants ignobles. Toutes ces choses changent avec le temps. Prenez dans ce texte ce qui vous est utile, adaptez-le, et mettez de côté les parties inutiles ou datées.

Un peu de modélisation de menace peut grandement aider à réduire les effets de la répression tout en augmentant la palette de stratégies possibles et disponibles dans notre poursuite de ces objectifs. Elle peut être fait comme exercice solitaire, durant une conversation libre avec des camarades, ou comme élément d'une analyse focalisée en préparation à une action importante. Une fois qu'on a développé cette compétence, on peut avoir plus de confiance dans notre capacité à réduire les effets de la répression et on peut réduire la pression et le temps d'exécution quand on

planifie des actions ou même tout simplement dans notre existence ordinaire de militant·e.

VOCABULAIRE

La modélisation de menace est un processus structuré. Donc nous devons avoir un vocabulaire bien définis afin de nous assurer que nous comprenons les mots que nous utilisons :

Sujet : Un sujet est une personne ou groupe qui pourrait être la cible d'un processus de surveillance, répression ou espionnage. C'est l'entité sur laquelle le modèle de menace se focalise. Vous pourriez être le sujet d'un modèle de menace ainsi que (directement ou indirectement) le sujet du modèle de menace d'une de vos équipes.

Objectif : Un objectif est quelque chose que le sujet souhaite accomplir.² Les objectifs peuvent être des trucs comme « entraver les activités de groupe nazi X » ou « éviter de se faire *doxxer* ».

Stratégie : Une stratégie est un ensemble spécifique d'actions prises par un sujet afin d'accomplir leurs objectifs. C'est le chemin exact, parmi beaucoup de chemins possibles, qui permet d'arriver à un (ou plusieurs) de leurs objectifs possibles. Si vous faites de l'anti-fascisme classique et que votre objectif est d'entraver le militantisme droitard, une stratégie que vous pourriez sélectionner est : mardi prochain on va coller 500 affiches qui dénoncent un facho local à travers le quartier où iel bosse et vit.

Adversaire : Un adversaire est une personne ou groupe qui veut empêcher le sujet d'accomplir ses objectifs. Les adversaires peuvent être internes (poucaves, entrepreneurs politiques carriéristes) ou externes (flics, fachos). Ils peuvent être directs (flics) ou indirects (faux alliés, factions concurrentes).

Capacité : Une capacité est une connaissance, compétence ou un outil qu'un adversaire possède qu'il peut utiliser contre le sujet pour l'empêcher d'arriver à ses objectifs. Les capacités peuvent être un troupeau de motos, la surveillance en direct du trafic internet d'un FAI (fournisseur d'accès internet) ou la capacité à utiliser la violence légale et extra-légale.

Vulnérabilité : Une vulnérabilité est un aspect de la vie du sujet ou d'un protocole de sécurité qui peut être exploité par un adversaire afin d'entraver ses objectifs ou de riposter contre ellui.³ Avoir l'habitude de se vanter est une vulnérabilité car ça peut amener un sujet à fuiter des informations sur des actions passées qui sont

2. Beaucoup de textes sur la modélisation de menace parlent de biens qu'un sujet veut protéger et ça fait sens quand on parle d'objets précieux dans un coffre ou de données dans un réseau informatique. Pour les radical·es, on est généralement moins concerné·es par des biens physiques que par des biens intangibles. C'est pour ça qu'on parlera « d'objectifs ».

3. D'autres textes dans d'autres littératures font une différence entre faiblesses et vulnérabilités et dans le contexte de systèmes informatiques, ça pourrait faire sens. Un programme peut avoir des faiblesses de conception et ça n'en fait pas des vulnérabilités. Ça n'a aucun sens de dire qu'un être humain a des faiblesses de conception dans comment ellui vit sa vie sous un régime répressif.

sensées être secrètes. Ne pas avoir la citoyenneté de l'État sous lequel on réside (i.e. pouvoir être déporté) est aussi une vulnérabilité potentielle.

Menace : Une menace est la possibilité réaliste qu'un adversaire exploite une vulnérabilité. Elle peut être abstraite comme « quelqu'un pourrait hacker notre ordinateur » ou concrète comme « le groupe nazi X va se pointer au prochain événement drag au lieu Y ». L'adjectif réaliste est inclus afin de garder notre processus de recherche dans une perspective relativement étroite. Bien que l'État ait la capacité à faire une frappe de drone sur notre bus du lundi quand on va au taf, si on vit dans un pays du soi-disant occident, les chances pour que ça se produise — présentement — sont quasiment de zéro donc il ne s'agit pas d'une menace réelle.

Impact : L'impact est la mesure des conséquences négatives dans le cas où une vulnérabilité est exploitée. Une vulnérabilité qui pourrait révéler l'information biométrique d'un sujet durant une action pourrait avoir un impact fort (des décennies en prisons). Une vulnérabilité qui pourrait révéler les horaires de travail d'un sujet pourrait avoir un impact faible (la plupart des gens travaillent durant le jour/soir).

Probabilité (d'impact) : La probabilité d'un impact est la mesure qui considère à la fois les chances qu'un adversaire va tenter d'exécuter une menace mais aussi quelles chances il a de réussir.⁴ Des caméras de surveillance sont pratiquement certaines d'être présentes sur des terrains-cibles intéressants et sans des contre-mesures appropriées comme couvrir son visage ou ses tatouages, la présence des caméras a une forte probabilité d'avoir un impact. Un agent de sécurité privée se baladant par hasard dans la même rue que vous pendant que vous êtes en train de graffer un immeuble à peu de chances de se produire même si se faire voir par celui-ci assurait d'être attrapé.

Risque : Le risque est la mesure combinée de l'impact et de sa probabilité. Une vulnérabilité avec un impact extrêmement fort qui a une probabilité faible mais réaliste de se produire pourrait être considérée comme un risque faible, mais une vulnérabilité avec un impact seulement modéré et une forte probabilité pourrait être considéré comme un risque haut.

Contre-mesure : Une contre-mesure est une action prise afin de réduire le risque.⁵ Les contre-mesures pourraient réussir à réduire des probabilités en adressant la vulnérabilité-même ou bien elles pourraient viser à altérer des aires adjacentes de la vie du sujet ou du protocole de sécurité. Les contre-mesures n'ont pas besoin d'amener le risque à 0 pour valoir la peine d'être mises en place.

Protocole de Sécurité : Un protocole de sécurité est l'ensemble de contre-mesures qui sont prises par un sujet par rapport à un ensemble donné d'adversaires

4. Certaines méthodes de modélisation de menace vont utiliser deux mesures (exploitabilité et les probabilités [de tentatives de menace]), mais pour faire simple dans cette brochure on les a réduits à une mesure unique. Si vous voulez les séparer pour faire des échelons de menaces, libre à vous.

5. Dans d'autres textes, on appelle ça des « mitigations » mais étant donné la nature extrêmement active de l'OpSec par rapport au simple fait de mettre à jour un ordinateur, le terme contre-mesures semble plus approprié.

avec des capacités connues ou anticipée.⁶ Il pourrait désigner les stratégies d'OpSec prises spécifiquement durant une action et sa phase préparatoire, ou bien il pourrait désigner les normes qui constituent la culture de sécurité d'un milieu social donné.

Modèle de menace : Le modèle de menace est le résultat du processus de modélisation de menace. C'est un modèle qui énumère les objectifs et stratégies du sujet, ses adversaires et leurs capacités, ainsi que les contre-mesures que le sujet peut utiliser contre elleux. Le modèle de menace informe le protocole de sécurité qui guide les actions du sujet. Parfois le terme est utilisé pour parler d'un adversaire spécifique avec des capacités connues, par exemple « notre modèle de menace est pensé contre les services de renseignement domestique. »

LES BASES

La modélisation de menace est le processus structuré qui consiste à identifier les menaces à tes objectifs et à sélectionner des contre-mesures qui peuvent être déployées contre ces menaces. Le but de la modélisation de menace est d'analyser tes comportements et stratégies afin d'apprendre comment tu es ou pourrait être exposé·e à la Répression. Il y a bien des manières de modéliser des menaces, et elles ont toutes des avantages et des inconvénients. Elles tendent toutes à partager certaines caractéristiques, par exemple elles répondent la plupart du temps aux questions suivantes :⁷

1. Quels sont nos objectifs ?
2. Comment est-ce que quelqu'un pourrait s'y opposer ?
3. Qu'est-ce qu'on peut faire contre ça ?
4. Est-ce que notre modèle de menace est prédictif ? Et est-ce que notre protocole de sécurité est efficace ?

Le dernier point est essentiel car la modélisation de menace n'implique pas uniquement des itérations dans une seule session. Elle est aussi itérative sur une période, au fur et à mesure que ses succès et échecs deviennent apparents une fois qu'on l'a appliqué au monde réel.

Il y a deux résultats désirables à une bonne modélisation de menace : la découverte de nouvelles menaces ou contre-mesures, et une amélioration de la capacité prédictive du modèle.⁸ La modélisation de menace devrait révéler quelque chose de nouveau, même à la plus expérimentée des équipes. Si ce n'est pas le cas (c'est à

6. On appelle ça parfois un « plan de sécurité » mais les plans sont flexibles et souvent négligé sur un coup de tête. Un protocole est quelque chose d'à la fois rigide et demandant, et ces deux aspects sont positifs. S'il est trop rigide il faut le renégocier collectivement, et pas l'ignorer de manière aléatoire sans en parler aux autres.

7. Ceci est une adaptation du Four Question Framework.

8. Notez que l'on utilise pas le terme « précis ». Tous les modèles ont des inexactitudes inscrites intentionnellement parce qu'une précision infinie mènerait à des modèles bien trop complexes pour être compris par les humains. La précision est tout de même importante étant donné qu'un modèle inexacte

dire si vous ne faites qu'écrire ce que vous savez déjà), l'exercice est quand même utile dans le sens où il s'assure que votre modèle de menace et protocole de sécurité sont bel et bien partagés par toutes les personnes de l'équipe. De plus, si aucune nouvelle découverte n'émerge de la modélisation de menace, plus de travail de recherche devrait être effectué afin de découvrir de nouvelles menaces inconnues ou des stratégies supérieures.

Le produit final de la modélisation de menace itérée doit être assez spécifique. « Créer l'anarchie » n'est pas un objectif suffisamment spécifique pour être mis en œuvre, et « quelqu'un pourrait nous espionner » n'est pas une catégorie d'adversaire suffisamment spécifique pour qu'on se défende contre elle. La précision est importante pas seulement parce qu'elle nous informe de contre quoi on doit se défendre à coup de contre-mesures mais aussi ce contre quoi on n'a pas besoin de contre-mesures. Des objectifs et adversaires non-spécifiques peuvent être des points de départ, mais il faut itérer jusqu'à ce qu'ils deviennent spécifiques.

Tu peux modéliser des menaces seul-e. Ça peut t'être utile parce que ça te permet d'identifier les limites de ta tolérance au risque ainsi que le genre d'objectifs que tu veux poursuivre. Avoir des objectifs qui ne s'alignent pas avec ceux de tes camarades peut vouloir dire compromettre certains de tes idéaux ou désirs afin de travailler avec d'autres personnes. Avoir des tolérances divergentes au risque peut signifier que quelqu'un se sentira anxieux-e — possiblement au point de ne pas être fiable — ou que quelqu'un sentira qu'il ne fait pas assez et devrait faire des actions plus audacieuses. Il est souvent plus facile de contempler et méditer en solitaire sur des objectifs et sur notre tolérance personnelle au risque plutôt qu'en face d'un groupe où la pression sociale et la fierté peuvent t'influencer et te pousser à cacher tes vraies préférences. Une fois que tu as ton propre modèle, va voir ton groupe et modélise avec eux. Tu pourras découvrir qu'il te faut trouver un nouveau groupe affinitaire, et il n'y a aucun mal à ça.

L'avantage de modéliser en groupe c'est qu'on peut utiliser les connaissances des autres afin d'informer et enrichir le modèle. Le gros de ce qu'on fait, que ce soit tracter dans la rue ou saboter des équipements de forage, on le fait avec au moins une autre personne. Ainsi il peut y avoir une limite à l'utilité de nos modèles de menace individuels car nos objectifs personnels peuvent être différents de ceux de notre équipe. Le protocole de sécurité qui émerge d'une session collective de modélisation pourrait être quelque chose avec laquelle tu n'es pas à l'aise. Peut-être que tu es d'accord avec les objectifs et stratégies du groupe, mais que leur approche bâclée de la sécurité génère une dose intolérable de risque à tes yeux.

La modélisation de menace demande de la préparation. Si c'est possible il faudrait y dédier une heure si tu comptes le faire seul-e, et deux à quatre heures si tu le fais en groupe. Et il faudrait sans doute faire plusieurs sessions car ça demande un travail de recherche pour remplir les lacunes et manques de connaissances.

qui ne reflète pas du tout la réalité sera incapable de faire des prévisions utiles des actions et réactions d'un adversaire.

Il y a une dépendance circulaire dans la modélisation de menace. Afin de savoir où et comment on peut faire de la modélisation de menace en toute sécurité, il faut déjà avoir un modèle de menace qui réponds à la question : « est-ce qu'on se prendrait de la répression dans la gueule de la part de l'État si on avait cette conversation ? »⁹ Il pourrait y avoir des limitations pré-existantes qui font que les lieux fréquentés par les personnes radicales (ton appart, un squat, un infokiosque) sont déjà sous surveillance et pourraient ainsi ne pas être appropriés pour de telles discussions. Vous pourriez le faire en l'absence de vos outils électroniques dans un café ou un parc où vous n'allez jamais. Pensez à prendre de quoi écrire. Peut-être aussi un briquet en fonction de ce que vous allez modéliser, il vous faudra peut-être détruire le modèle-même et uniquement conserver le protocole de sécurité.

UNE MÉTHODE PARTICULIÈRE

La méthode décrite dans cette section n'est pas « la meilleure » (en partant même du principe qu'une telle chose existe). C'est juste une méthode qui sera discutée avec assez de profondeur pour vous permettre de construire la votre. Je ne vais même pas essayer de la nommer pour éviter de lui donner une importance exagérée par rapport aux méthodes que vous pourriez générer.

Cette méthode est orientée-vers-des-objectifs. Elle se focalise sur ce que vous voulez et applique ensuite des contraintes.¹⁰

IDENTIFIER DES OBJECTIFS

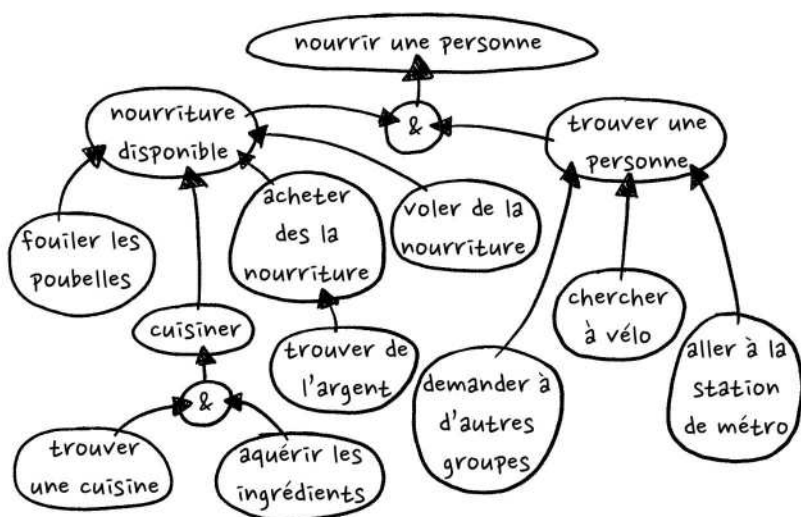
Si vos objectifs et désirs ne sont pas clairs, la première étape est de les matérialiser en quelque chose de plus spécifique. Commencez par un brainstorming ou par faire une carte heuristique. Écrivez toutes vos idées sur un papier, ou bien chaque idée sur une carte séparée. Ne vous inquiétez pas de la faisabilité des idées ou de leurs risques. Commencez juste par écrire. Prenez des objectifs et divisez-les en sous-objectifs ou bien notez les prérequis permettant d'accomplir un objectif. Dessinez des connexions entre des objectifs liés ou bien entre objectifs et sous-objectifs, et si vous utilisez les cartes séparées, agglomérez-les ensemble.

9. Ceci s'applique généralement à tous les adversaires. Si vous modéliser les menaces en jeu dans le fait de quitter une relation abusive, votre partenaire pourrait vous espionner afin de vous empêcher de partir.

10. Ceci est en opposition à, par exemple, une modélisation de menace orientée vers un adversaire qui se focalise sur les adversaires et leurs capacités avant de regarder quels objectifs sont disponibles étant donné les contraintes observées. Une faiblesse de l'approche orientée vers un adversaire est que quelques rares démonstrations de force ou coups de chances dans des enquêtes peuvent nous amener à voir ces moments comme la base des capacités de l'adversaire. Ceci nous isole de manière anticipée de certaines voies d'attaques que l'on aurait laissé ouvertes dans d'autres situations. Comme nous le verrons dans les sections suivantes, il faut évaluer la probabilité d'une menace, et pas juste la possibilité qu'elle se produise dans l'absolu.

Pour certaines personnes, leurs méthodes et objectifs généraux sont déjà rodés parce qu’elles sont dans de mouvements radicaux depuis de nombreuses années. Un groupe affinitaire pourrait opérer selon le principe directeur d’entraver les fascistes dans leur contexte local immédiat. Les contraintes de temps et de matériel pourraient limiter leur travail, donc leurs objectifs vont se voir attribuer une envergure appropriée à ce qu’iels peuvent raisonnablement accomplir tout en réussissant quand même à se nourrir et à trouver un peu de joie de vivre dans ce monde. Ainsi, dans leur situation, identifier des objectifs pourrait être plus proche d’un processus de sélection de cibles.

FIGURE 1 – Diagramme sur comment nourrir des gens à la rue



Les objectifs n’ont pas besoin d’être des objectifs anarchistes « classiques » comme attaquer ceci ou construire cela. En fonction des circonstances personnelles ou identités de quelqu’un, un objectif pourrait tout simplement être de survivre malgré l’oppression qu’iel subit. Ceci pourrait amener à des sous-objectifs comme éviter des interactions avec la police, ne pas attirer l’attention sur soi, ou accumuler assez de ressources pour pouvoir bouger vers un endroit plus sûr.

Parce que l’anarchisme est un mouvement social et pas juste un projet purement individuel, on va implicitement générer des objectifs qui impactent d’autres personnes. Nos stratégies et protocoles de sécurité devraient généralement chercher à protéger les autres de démêlés légaux, de l’incarcération, ou de dommages physiques ou psychologiques. Ces objectifs altruistes existent afin que notre protocole de sécurité ne soient pas seulement centrés autour de « comment est-ce que j’évite de me faire arrêter » mais plutôt « comment est-ce qu’on évitent toutes de

se faire arrêter ».

IDENTIFIER DES PRÉREQUIS AUX OBJECTIFS

Une fois que vous avez vos objectifs, les prérequis doivent être clairement établis. En écrivant les étapes dont vous avez besoin pour atteindre cet objectif, vous pouvez examiner la faisabilité de l'objectif, et avoir les étapes écrites devant vos yeux aidera à trouver des vulnérabilités lors des étapes suivantes de la modélisation de menace. Dans la plupart des cas, il n'y a pas un chemin unique en direction de votre objectif, donc rédigez le plus de chemins possibles pour son accomplissement. Des cartes heuristiques basiques et des diagramme font le taf pour tout ça.¹¹

Dans l'exemple (simplifié) de la Figure 1, il y a deux nécessité à nourrir la personne à la rue : on doit à la fois avoir de la nourriture et être capable de la trouver. Ces deux sous-objectifs peuvent être accomplis de plusieurs manières.

IDENTIFIER ADVERSAIRES ET CAPACITÉS

Identifier des objectifs peut demander beaucoup d'introspection au fur et à mesure qu'on détermine ce qui a de l'importance ou quelles stratégies on considère comme efficaces ou éthiques. Identifier des adversaires et capacités demande un vrai travail de recherche.

Le « bon-sens » nous fait immédiatement penser qu'en tant que membres d'un mouvement social anti-autoritaire, nos le gros de nos adversaires sont les flics locaux, les agences de sécurité nationale, les loups solitaires de droite, et les groupes organisés de fascistes. Le même bon sens nous suggère la forme générale de leurs capacités tels que la collecte de preuves légales, l'écoute téléphonique, ou juste la violence. Les choses qu'on pourrait appeler du « bon sens », pourrait être qualifiées par quelqu'un d'autre de paranoïa délirante, et une autre personne encore pourrait les qualifier de vision très simpliste des méthodes opératoire de l'État et des fascistes. Déterminer qu'est-ce qui, dans toutes ces intuitions, reflète la réalité nous demande de faire des recherches sur chacune de ces affirmations.

Bien qu'on puisse être tenté-e de lister chacune des agences de renseignement spécifique existant dans notre région, il pourrait y avoir peu de différences significatives entre deux agences nationales de renseignement ou deux détachements de police locaux. Il y a seulement une poignée d'unités policières et seulement une poignée de groupes fascistes. Elles peuvent généralement être regroupés dans des grosses catégories telles que :

- État : renseignement domestique, police locale.
- Non-État : Groupes qui frappent les premiers, groupes qui frappent en réponse-à, groupes qui sont de manière effective non-violents.

11. Certaines personnes utilisent les diagrammes os de poisson (Ishikawa), mais je les trouve difficiles à lire, et ils peuvent finir bordéliques.

Cependant avec l'utilisation de partages de données et de *fusion centers*,¹² ou plus généralement la digitalisation du maintien de l'ordre et l'utilisation des ordinateurs pour faire des quantités inhumaines de référencement de données, les capacités pourraient se confondre les unes dans les autres entre agences de renseignement et police locale. Donc ces grandes catégories pourraient, dans les années à venir, cesser d'être des distinctions claires et significatives.¹³

Dans votre analyse, débarrassez-vous de grandes catégories d'adversaires, de manière raisonnable. En tant que gérant-es d'une soupe populaire, vous n'êtes probablement pas sous surveillance active de la part du renseignement domestique ou la cible la plus probable pour une infiltration profonde.¹⁴ Additionnellement, enlevez des entités non-étatiques. Des factions rivales pourraient ne pas s'en prendre les unes aux autres, et dans le cadre de vos actions, leur existence pourrait être sans conséquence. Des gangs ou mafia pourraient vous laisser totalement tranquille tant que vous ne piétez pas sur leur racket, et si ce n'est pas votre but alors vous n'avez potentiellement pas à vous soucier d'eux dans votre modèle de menace.

UTILISER UNE LIBRAIRIE DE MENACES

De grandes catégories d'adversaires sont identifiées afin que des capacités spécifiques puissent leur être assignées. Deux équipes opposées l'une à l'autre mais également opposées à l'État et qui sont prônes au même type d'activités remarquables vont avoir le même adversaire avec les mêmes capacités. Deux équipes dans des parties complètement différentes du monde pourraient faire face à des adversaires quasiment identiques (par exemple deux États non-alliés avec des agences de renseignement domestiques similaires), et en conséquence si ces deux équipes énumèrent les capacités de leurs adversaires, elles pourraient se retrouver avec des listes très similaires. Tout ceci pour dire que bien que la modélisation de menace puisse être réalisée de manière individuelle, des travaux de recherche sur des adversaires peuvent être partagés afin de réduire massivement l'effort qui a tendance à être dupliqué. Les collections de connaissances de ce type sont appelées des librairies de menaces.

Une librairie de menace catégorise et explique les capacités que des adversaires

12. *N.d.t.* : https://en.wikipedia.org/wiki/Fusion_center, des initiatives américaines visant à assouplir la collaboration entre différents corps de renseignement domestique au niveau fédéral; la France n'a pas les mêmes problématiques du fait d'un mille-feuille administratif différent mais il y a tout de même des initiatives similaires en termes d'harmonisation du fichage et du partage d'information entre différentes administrations.

13. Par exemple, le NYPD a le *Department Intelligence Bureau* qui opère hors du cadre de supervision classique et a des actifs de renseignement étranger et des connexions directes à la police étrangère. Est-ce que ça fait d'eux une police locale au sens classique? Une agence de renseignement? Quelque chose d'autre entièrement?

14. Elleux pourraient utiliser une soupe populaire pour gagner un ancrage leur permettant de s'introduire dans le reste du milieu mais la soupe populaire elle-même a peu de chances d'être leur cible principale.

pourraient avoir et les menaces qu'ils pourraient produire. Elle pourrait avoir la forme d'une pile de cartes classées par couleur que l'on garde sous la main pour pouvoir s'en servir hors-ligne/quand on a pas accès à un ordinateur, ou bien elle pourrait être une base de donnée hébergée en ligne quelque part.¹⁵ Un exemple peut-être observé en Table 1.

TABLE 1 – Échantillonnage d'Index d'une Librairie de Menaces

Surveillance Physique	Opération Psychologique	Hacking
Mise sur écoute de ligne téléphoniques fixes	Répandre des rumeurs	Exploits de hardware
Caméras de surveillance (en intérieur et en extérieur)	Infiltrés (effets de second ordre)	Accords avec des fabricants pour implémenter des backdoors
Tracking de localisation de téléphone mobile	Harcèlement en personne	Zero-days dans des librairies open-source

Que vous utilisiez une librairie externe ou que vous créiez la votre, voici quelques éléments qu'une librairie de menace devrait avoir pour être utile.

- Les menaces devraient avoir des noms clairs et spécifiques.
- Les menaces devraient être liées à des menaces apparentés, notamment les menaces adjacentes, parentes et enfants.
- Les menaces qui se sont produites par le passé, même s'il n'y a pas d'occurrence récente connues de ces dernières, devraient être incluses. Elles pourraient se produire à nouveau.
- Les menaces théoriques qui sont à l'horizon même si on sait qu'elles n'existent pas encore, devraient tout de même être incluses (surtout celles qui sont numériques).
- Les menaces devraient lister des exemples clairs de s'être produits. Cela les rends « réelles » et donne aussi des ramifications pour des recherches futures dans le cadre d'études de cas.

15. La librairie de menace du *No Trace Project* (<https://www.notrace.how/threat-library/>) est la seule que j'ai trouvée qui semble être spécifiquement construite avec des radical-es en tête et qui est de qualité correcte.

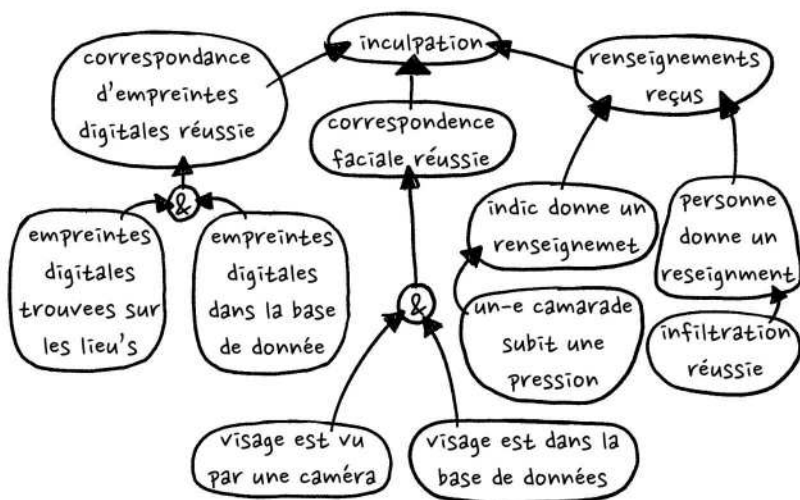
- Les menaces devraient être cartographiées en relation explicite avec des adversaires, avec une idée certaine de la fréquence à laquelle elles sont utilisées par ceux-ci ou bien la fréquence avec laquelle elles les ont amenés à réaliser leurs objectifs (par exemple : avec quelle fréquence est-ce que le hacking d'ordinateurs à aider la police nationale à obtenir une condamnation de leurs cibles ?)

CRÉER UN ARBRE D'ATTAQUE

Tout comme nous, nos adversaires ont des objectifs, des sous-objectifs et des pré-requis à ceux-là.

Les objectifs des flics sont généralement d'« empêcher les anars de faire des trucs » ou d'« arrêter les anars après qu'ils aient fait des trucs ». Un groupe de droitards locaux pourrait avoir pour simple objectif de « tabasser quelques pédés ». Ces objectifs sont vagues, et on pourrait les préciser en des menaces plus spécifiques. Ces menaces en elles-mêmes peuvent être reconsidérées jusqu'à ce qu'on précise comment elles pourraient se produire dans les faits. On appelle arbre d'attaque un diagramme qui montre les pré-requis d'une menace donnée. Il y a un exemple en Figure 2.

FIGURE 2 – Arbre d'attaque pour examiner une action directe (incomplet)



Quand vous créez des arbres d'attaque, vous pouvez annoter chaque nœud avec des informations afin de vous aider à déterminer si une menace spécifique est plausible. Une annotation est une estimation de l'effort dépensé pour chaque pré-requis ou sous-objectif. S'il y a deux manières d'arriver à un sous-objectif, votre adversaire

a des chances de choisir celui qui nécessite le moins d'efforts. Vous pourriez aussi annoter les retours de bâtons possibles.

Un meurtre illégal — ou quasi-légal grâce à l'immunité policière — pourrait avoir des conséquences pour l'individu qui commet l'acte ou l'organisation qui l'a commandité, et vous pourriez estimer qu'éviter les conséquences de leurs actions est quelque chose que les flics locaux vont valoriser. Si votre adversaire est un acteur non-étatique, vous pourriez annoter les sous-objectifs et leur légalité, ou plutôt les chances de se faire poursuivre judiciairement étant donné que les flics laissent parfois des fachos faire le sale boulot à leur place.

Les arbres d'attaque peuvent aider à penser comme votre adversaire afin de vous aider à deviner quelles stratégies iel va sélectionner pour essayer de vous entraver. Ils sont limités par le fait que vous ne savez pas ce que votre adversaire sait ou comment iel évalue certaines situations. En tant qu'anarchistes vous pouvez partir du principe que vous allez vous prendre un certain niveau de répression dans la gueule quand vous tentez des actions criminalisées mais ces actions ne sont pas forcément criminalisées pour des acteurs fascistes. Estimer la complexité de leurs efforts peut aussi être difficile car vous pourriez surestimer les compétences techniques d'un adversaire ou sous-estimer le temps et les ressources dédiées au fait de vous poursuivre. Comme pour le reste de votre modèle de menace, il pourrait y avoir des choses que vous échouez totalement à voir.

Un concept lié à celui d'arbre d'attaque c'est la chaîne d'attaque. L'idée de la chaîne d'attaque est qu'il y a une séquence d'événements qui doivent se produire afin que votre adversaire réussisse à vous empêcher d'atteindre votre objectif. Si vous pouvez interrompre cette séquence à n'importe quelle étape (c'est à dire casser la chaîne) vos contre-mesures seront suffisantes pour vous permettre d'atteindre votre objectif. Vu qu'il pourrait y avoir un certain nombre de séquences, on peut les visualiser comme un arbre. Si l'arbre est précis alors on peut trouver des manières de couper assez de branches afin d'empêcher que l'adversaire atteigne son objectif.

ÉNUMÈRE ET PRIORISE LES MENACES

Une fois que vous avez une liste d'adversaires et une méthode pour lister quelles menaces iels pourraient représenter, il faut trouver une manière de rendre cette liste actionnable. Si votre librairie de menaces est suffisamment détaillée et que vous avez assez fait de recherche sur des équipes similaires, vous aurez une idée des capacités qui sont réellement déployées. Cependant vous devriez garder en tête que votre adversaire n'est pas forcément en train d'utiliser toutes ses capacités afin que ses capacités maximales réelles restent inconnues du public. En ce qui concerne l'État, la police peut faire de la construction parallèle afin que même si on lise de manière très attentive les cas légaux et les preuves présentées, vous pourriez n'apprendre qu'uniquement comment ils *disent* s'être procurés les preuves nécessaires à une condamnation, mais pas comment ils se sont procurés les preuves *dans les faits*.

Après avoir retiré les menaces qui ont une chance faible de se produire, votre liste va probablement quand même être trop grande pour être adressée de manière totale. Pour prioriser des menaces, il vous faut une heuristique. Une méthode consiste à utiliser son « intuition » et à juste les réarranger dans une disposition qui semble correcte. Ceci tout à fait raisonnable.

Une autre méthode consiste à assigner à chaque menace un impact et une probabilité, puis à multiplier ces scores.

Des scores pour la probabilité et l'impact sont visibles sur la Table 2. Les nombres sont biaisés pour mettre plus d'emphasis sur l'impact et surtout sur les impacts sévères. Restez conscientes que le terme « catastrophique » est relatif aux objectifs que l'on considère. « Catastrophique » pour un graff (une amende d'un montant moyen) et « Catastrophique » pour un acte de sabotage (dix ans de prison) peuvent être deux cas de figure extrêmement différent et ne sont pas comparables dans ce contexte.

Si après avoir trié les menaces, vous avez l'impression que le tri n'est « pas correct » vous pouvez les réarranger. Les nombres sont juste là pour vous aider à faire un tri initial.

TABLE 2 – Calculer des scores de risque

	Impact				
	Rien (0)	Mineur (1)	Impact modéré (3)	Major (5)	Catastrophique (10)
Probabilités					
Jamais (0)	0	0	0	0	0
Rare (1)	0	1	3	5	10
Improbable (2)	0	2	6	10	20
Probable (3)	0	3	9	15	30
Pratiquement certain (5)	0	5	15	25	50

$$\text{Risque} = \text{Probabilité} \times \text{Impact}$$

Mêmes les meilleures d'entre nous avec des années d'expérience sont limités dans notre capacité à évaluer précisément le risque, ou plus spécifiquement la probabilité (d'impact). Pour certains types d'actions, les probabilités pourraient toutes être regroupées dans la colonne « rare » ce qui veut dire qu'il nous faut renommer les noms des probabilités, par exemple rare, moins rare, très inhabituel, etc. Comment on décide de ce qui est ou n'est pas à tel ou tel niveau de rareté est, au mieux, basé sur des sentiments d'intuition pas nécessairement très précis. Estimer des risques n'est pas tâche aisée.

IDENTIFIER DES CONTRE-MESURES

Une fois que des menaces ont été identifiées, elles doivent être résolues. Les résoudre ne veut pas dire la même chose que résoudre un problème totalement. On veut plutôt dire qu'on prends la décision consciente sur ce qu'on va faire à propos de la menace spécifique. Les menaces peuvent être considérées comme résolues de quatre manières possibles.

Acceptée — Une menace est acceptée s'il est décidé que rien ne peut être fait à son sujet. La menace est identifiée, des contre-mesures sont considérées et ensuite s'il est déterminé que l'objectif ou la stratégie qui s'appliquent à cette menace sont trop critiques pour être altérées ou retirées, ou que les contre-mesures ne sont pas faisables étant donné les ressources disponibles actuellement, la menace est marquée comme acceptée.

Évitée — Une menace est évitée si un objectif ou une stratégie est complètement retirée des plans d'actions possibles. Les menaces de géolocalisation seront évitées si vous choisissez un objectif qui ne mène jamais à ce genre d'enquête. La menace n'est juste plus là.

Remédié — Une menace est remédié si la probabilité qu'elle ait un impact est réduite. Ses causes premières sont identifiées, et des altérations à la stratégie peuvent empêcher qu'elle ne se produise. La menace d'une arrestation à cause de l'analyse de SMS collectés à travers une surveillance par dragnet peut être remédié en ayant recours à des applications de messagerie chiffrées de bout-en-bout(même si le contenu de ces messages est toujours aussi incriminant). La menace d'une identification par empreintes digitales est remédié par le fait de porter des gants solides et d'essuyer tous les objets utilisés durant une action (car il y aura une équipe de police scientifique, mais ils ne trouveront pas d'empreintes).

Mitigée — Une menace est mitigée si son impact est réduit. Il n'est peut-être pas possible d'éviter la menace ou d'en réduire la probabilité de quelque façon que ce soit, mais la sévérité de son impact pourrait tout de même être réduite. On va mitiger l'impact de l'arrestation d'un membre de l'équipe (et son incarcération potentiel) en se tenant à un strict code de silence, l'impact plus large sur les 6 autres individus dans l'équipe sera réduit. Avoir le contact d'un conseiller légal mitige les effets d'une arrestation car cela garantit que quelqu'un sera là pour vous assister durant la durée de l'enquête.

Il n'y a pas de ligne claire pour délimiter les remédiations et les mitigations, mais il est utile de se rappeler des deux manières de contrer un risque. On peut rendre un risque moins probable de se produire (y remédier) et/ou le rendre moins impactant lorsqu'il se produit (le mitiger).

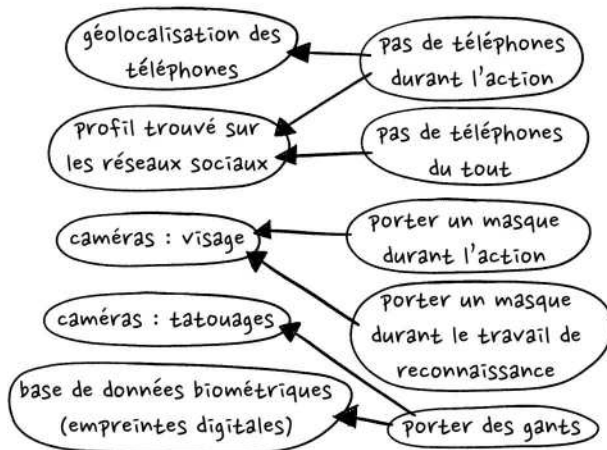
Il y a une cinquième non-résolution qui est de juste faire acte d'ignorance. Une menace est ignorée quand elle n'est pas connue soit par ignorance ou quand le sujet d'un modèle de menace choisit de ne pas l'analyser. Par exemple, un membre d'une équipe qui nourrit des sans-abris pourrait indiquer que des fascistes ont commencé à attaquer les gens à la rue et leurs défenseurs. Si le reste de l'équipe dit « c'est pas

un problème » sans en discuter, ceci est une situation où iels ignorent la menace, choisissant de ne pas l'accepter. Accepter une menace mène à un consentement informé à propos de l'action. Ignorer une menace n'est pas la même chose.

Afin d'identifier des contre-mesures, prenez votre liste classée de menaces et utilisez une carte heuristique afin de brainstorm les manières possibles de les résoudre. Certaines contre-mesures s'appliqueront à plusieurs menaces (pour voir l'exemple, regardez Figure 3). Certaines menaces demanderont que plusieurs contre-mesures soient utilisées les unes avec les autres afin d'être résolues. Certaines contre-mesures pourraient être redondantes et ça peut-être une bonne chose.¹⁶ Si vous avez une librairie détaillée de menaces, vous pourriez sélectionner des contre-mesures depuis celle-ci. Sinon, conservez les idées que vous avez brainstormé durant votre session et mettez à jour votre librairie après coup. Vous pourriez avoir besoin de faire des recherches afin de valider que ces contre-mesures sont suffisamment efficaces.

Dessiner plusieurs lignes entre des nœuds sur une carte heuristique peut mener à un modèle brouillon. Une autre méthode consiste à étiqueter toutes les menaces avec des marques comme M1, M2, etc. ou de petits noms comme M-HACK-TÉLÉPHONE, M-RECONNAISSANCE-FACIALE, etc. Faites la même chose avec les contre-mesures comme C1, C2, etc. ou C-PAS-TÉLÉPHONES, C-MASQUES, etc. Faites une table avec vos menaces et listez tout simplement quelles contre-mesures s'appliquent à quelles menaces (Table 3).

FIGURE 3 – Menaces et Contre-mesures : carte heuristique



Une fois que cette partie de la méthode sera complétée, Vous aurez un modèle de menace complet et l'étape suivante consiste à vous en servir afin d'enrichir un

16. On appelle ça la « défense en profondeur » ou le « modèle du fromage suisse »

TABLE 3 – Menaces et Contre-mesures

Menaces	Contre-mesures
M1	C1, C3
M2	C2
M3	C4, C5, C6
M4	<i>Rien! À travailler!</i>
M5	C2

protocole de sécurité. Si vous n'êtes pas satisfait-es avec le modèle, vous avez peut-être besoin d'itérer.

Quand vous identifiez des contre-mesures vous avez peut-être découvert que vous aviez de nouveaux objectifs à considérer. Pendant que vous identifiez des menaces, vous avez peut-être découvert de nouveaux adversaires à considérer. De même, une fois que vous commencez à produire votre protocole de sécurité, vous vous êtes peut-être rendus compte que votre modèle de menace a des lacunes qui doivent être comblées.

ÉLABORER UN PROTOCOLE DE SÉCURITÉ

Arrivé-es à cette étape, votre modèle de menace compte sans doute beaucoup d'objectifs avec beaucoup de chemins possibles à ceux-ci, et chaque objectif ou tâche pourrait venir avec son lot de menaces. Vous n'avez besoin que d'une stratégie (ou ensemble de chemins) afin d'arriver à votre objectif. Différentes stratégies vous amèneront à produire différents protocoles de sécurité, et chacun aura des risques associés différents. Certains protocoles pourraient être trop encombrants pour être réellement appliqués soit parce que trop complexes à exécuter ou parce qu'ils ralentissent de manière excessive la progression vers un objectif. Certains protocoles pourraient traiter les risques de manière insuffisante.

Choisissez une des stratégies possibles qui vous permet d'arriver à votre objectif, et copiez ses sous-objectifs et menaces sur une feuille séparée, en pensant également à recopier les connexions entre celles-ci. Dans un coin, notez deux numéros pour décrire le risque de la stratégie. Le premier nombre est le plus haut score de risque de toutes les menaces (le maximum). Le second est la somme de tous les risques de toutes les menaces (le total).

$$\text{Risque}_{\max} = \max(M_1, M_2, \dots, M_n)$$

$$\text{Risque}_{\text{total}} = M_1 + M_2 + \dots + M_n$$

Ensuite, parlez de quelles contre-mesures peuvent être appliquées. Écrivez les

contre-mesures et reliez-les à leurs menaces respectives. Pour chaque menace, recalculez ses risques étant donné les contre-mesures qui lui sont appliquées. Sous les deux nombres dans le coin de la page, écrivez les nouveaux scores de risque pour cette stratégie. Si vous avez réussi à réduire le risque total et le risque maximum possible, la stratégie pourrait être acceptable.

Répétez ce processus avec d'autres stratégies possibles. Si vous pensez qu'une stratégie est moins risquée, mais que son nombre de risque est plus élevé qu'un autre, cela ne veut pas dire que votre intuition est fausse. Les nombres ne sont pas du tout absolus. Ils sont là pour que vous vous arrêtiez et que vous réfléchissiez. Si votre intuition n'est pas alignée avec les nombres, c'est que vous devez creuser un peu la question. Essayez de comprendre *pourquoi* les nombres ne sont pas les bons. Peut-être qu'il y a une très grande menace et que l'efficacité de ses contre-mesures est surestimée. Peut-être que c'est le contraire et vous avez assigné bien trop d'impact à quelque chose. Ajustez les scores sur vos stratégies.

Une fois que vous avez sélectionné une stratégie et les contre-mesures associées, votre protocole de sécurité est complété.

APPLICATION ET ÉVALUATION

Appliquez le protocole et faites vos actions. Soyez attentives à ce que tous le monde s'y conforme/l'applique. Ceci demande de la discipline (de ta part) et de la confiance (dans les autres). Si tu n'as pas assez confiance dans les autres pour qu'elleux admettent ne pas suivre le protocole (quand c'est le cas), alors il est possible que tu sois dans une situation où il n'y a pas assez de confiance mutuelle pour pouvoir effectuer l'action en question. Ou peut-être que tu as prévue tout ça dans ton modèle de menace auquel cas des déviations mineures ne vont pas dérailler le projet.

Au fil du temps et surtout après l'action, évaluez si le modèle semblait être aligné avec la réalité ou si le protocole était efficace. Si vous effectuez beaucoup d'actions mineures similaires mais que vous continuez à être entravées durant celle-ci, il est possible que votre modèle de menace manque d'adresser quelque chose d'important. Rassemblez-vous à nouveau pour discuter du modèle quand les échecs deviennent apparents. Des discussions périodiques peuvent être utiles car il pourrait y avoir possibilité à relaxer un peu la sécurité ou bien à ajouter des contre-mesures qui étaient auparavant estimées trop difficiles à appliquer de manière continue. Au minimum, seule ou avec ton équipe, tu devrai refaire ton modèle de menace au moins une fois par an. Les nouvelles technologies, tactiques des adversaires ou contextes politiques nécessitent une réévaluation du modèle de menace.

Il peut être dangereux de partager tes modèle de menace et protocole de sécurité exacts avec d'autres personnes, mais tu devrai essayer d'avoir des conversations générales sur le sujet de la sécurité afin d'apprendre des méthodes et pratiques d'autres équipes. Ceci peut t'alerter sur les nouvelles tendances et pratiques de tes adversaires ou de ton milieu. Tu pourrai te rendre compte que des équipes qui sem-

blaient fiables et sûres ne le sont en fait pas du tout, et vice-versa.

COMMENTAIRES SUR LA MÉTHODE

Après avoir lu la section sur la méthode, vous pourriez être en train de vous dire « putain mec c'est beaucoup trop complexe ». Écrit comme ça, c'est pas faux, mais en réalité c'est juste une heuristique que beaucoup d'entre nous utilisons déjà sans l'articuler. Prendre le temps de ralentir et nommer les différentes étapes que l'on accomplit habituellement en seulement quelques secondes peut donner l'impression que c'est plus complexe que ça ne l'est réellement. Et une fois familiarisées avec des menaces et contre-mesures courantes, la modélisation de menace deviendra très rapide et pourrait demander moins d'utilisation de chiffres pour aider avec la priorisation de risques. J'ai vu des groupes de radicaux vétérans mettre au point de nouveaux protocoles de sécurité en à peine 15 minutes quand les gens étaient bien informés avant de commencer et que les niveaux de tolérances au risque des gens étaient à peu près égaux. La majeure partie de ce temps sert juste à établir si tous le monde est d'accord, et une fois que c'est le cas, écrire le protocole est assez direct.

Cela dit, une modélisation de menace avec ce niveau de spécificité et profondeur peut-être une surenchère pour beaucoup de scénarios. Un processus typique pourrait consister à faire quelques recherches rapides sur des adversaires et d'ensuite se rendre rapidement compte que d'autres équipes qui ont utilisées des stratégies similaires ont fait face à une répression minimale. L'équipe pourrait ensuite décider de prendre des mesures de sécurité standard¹⁷ comme le simple fait de ne pas bavarder des actions en publique ou sur les réseaux sociaux. Tout ça est totalement OK.

Certaines actions demandent plus de focalisation et cette méthode n'est pas applicable à celles-ci. Cela dit ce processus demande de la pratique. Réussir totalement la première fois sur une action majeure est assez improbable. Si tu espères modéliser les menaces qu'implique une action très intéressante tu vas probablement devoir pratiquer cette méthode avec ton équipe sur des actions plus petites et ordinaires afin de t'habituer à utiliser la méthode et à suivre un protocole de sécurité explicite.

EXEMPLES

Pour rendre un peu plus concret le processus de modélisation de menace discuté précédemment, voici quelques exemples. L'un suit une méthode de manière assez stricte, et ce n'est pas parce que tous le monde *doit* le suivre à la lettre mais plutôt parce que c'est un exemple intentionnellement détaillé. Les autres ne le sont pas

17. Standard par rapport à leur contexte spatial et temporel mais aussi par rapport à l'intensité de leurs actions.

pour des raisons qui vont vite devenir claires. Souvenez-vous que bien qu'ils soient dérivés de cas réels, et bien que vous soyez capable d'appliquer des éléments de ça à votre situation, vous devriez éviter d'appliquer directement les modèles résultants à votre situation sans les modifier.

Une indication de plus sur le format : les gens qui travaillent sur ces scénarios auraient beaucoup d'espace et de papier pour faire leurs cartes heuristiques, mais on essaye de tout faire rentrer dans un zine format A5 donc toutes les itérations et notes ne sont pas inclues. Des versions abrégées et finales de leurs modèles sont utilisés pour gagner de l'espace.

EXTINCTEURS DE PNEUS

SCÉNARIO

Une bande d'ami-e fait partie du groupe de jeunesse de la branche locale du Parti Vert. Iels en ont marre de toutes les manœuvres politiciennes et du peu de progrès accomplis par rapport au temps passé en réunions, conférences et manifs. Iels n'ont jamais fait d'action directe, mais après avoir vu quelques posts sur les réseaux sociaux par les Extincteurs de Pneus,¹⁸ iels décident d'essayer leur tactique. La bande s'assied un soir et essaye de comprendre comment la mettre en œuvre.

MODÉLISATION DE MENACE ET PROTOCOLE DE SÉCURITÉ

Au début, le groupe a deux objectifs : dégonfler des pneus de SUV, et ne pas se faire chopper par les flics ou les proprios des voitures (les deux font trop flipper). Ayant lu le site des Extincteurs de Pneus, les ami-es savent qu'elleux ont juste besoin de lentilles vertes pour dégonfler les pneus et de tracts pour communiquer leur message politique. La bande fait un brainstorming et produit un diagramme de leurs objectifs et nécessités (Figure 4). Au fur et à mesure de leur séance de brainstorming, il devient clair que l'un de leurs objectifs n'est pas réel. Elleux s'en foutent de dégonfler des pneus. Leur but c'est décourager les gens de conduire des grands véhicules.

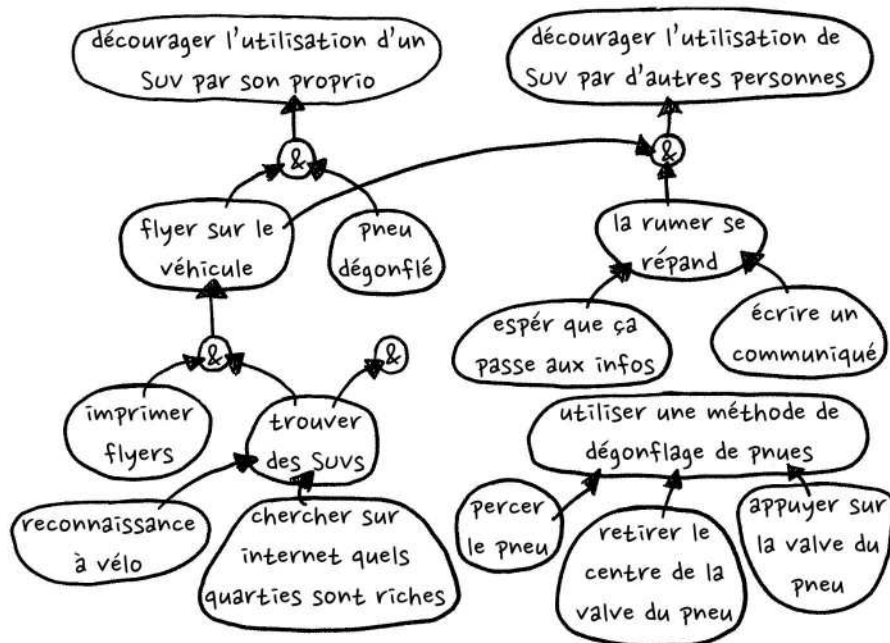
Une fois qu'elleux ont leurs objectifs et sous-objectifs, iels pensent à qui pourrait essayer de les arrêter : la police et des individus « civils » (soit les proprios du véhicule, soit des gens du voisinage). Iels commencent à faire une table de toutes les capacités que ces deux adversaires pourraient avoir (Table 4). Iels utilisent principalement des références de pop culture pour comprendre comment quelqu'un pourrait se faire attraper, et leurs données semblent assez détaillées à leur goût. Tout ça se passe bien jusqu'à ce qu'une personne pointe du doigt que certaines imprimantes laissent des marques à peine visibles sur les pages afin d'identifier quels appareils

18. Vous pouvez les traiter de libéraux si vous voulez mais iels ont inspirés des gens à faire de l'action directe autonome au lieu de juste se faire arrêter de manière performative. Tous le monde commence quelque part.

les ont imprimées.¹⁹ Ils font un peu plus de recherches pour voir s'il reste des données inconnues du genre.

Une fois que leurs menaces sont listées, des scores de risque y sont attribués (??). À partir de là on peut voir que les plus grands risques d'après leurs modèles sont :

FIGURE 4 – Objectifs des Extincteurs de Pneus



1. (50) M-CAMÉRAS/FLICS : la police capture leur image sur vidéo
2. (30) M-PHYSIQUE/CIVILS : Un civil intervient de manière physique
3. (30) M-PHONE/FLICS : les flics vérifient les données téléphoniques
4. (25) M-CAMÉRAS/CIVILS : Un civil vérifie les caméras de sécurité de sa voiture et envoie les données à la police ou bien publie les visages des personnes sur la vidéo sur les réseaux sociaux

Leur risque maximum est de 50 (le plus haut nombre sur l'échelle), et leur risque total après avoir fait la somme des colonnes flics/risque et civils/risque est de 206. Ils commencent à chercher des contre-mesures pour chacune de leurs menaces.

19. Ces marques sont appelées des Codes d'identification de machines et elles sont utilisées depuis les années 80.

TABLE 4 – Adversaires des extincteurs de pneus et leurs capacités

Menaces	Flics	Civils
M-CAMÉRAS	Examinent les vidéos de caméras publiques et caméras de voitures	Examinent les vidéos des caméras de voitures
M-REMARQUER	Peuvent remarquer des gens la nuit	Peuvent aussi remarquer des gens la nuit
M-SOCIAL	Peuvent chercher des posts pertinents en ligne sur les réseaux sociaux	Peuvent aussi chercher sur les réseaux sociaux
M-DOIGT	Collectent des empreintes digitales	—
M-PHYSIQUE	—	Interviennent physiquement
M-PHONE	Requête auprès des base de données de géolocalisation des appareils mobiles	—
M-PRINTER	Query printer dot database	—
M-EMAIL	Get email logs of who sent the communiqué	—
M-WEB	Get internet traffic logs of who sent the communiqué	—
M-IMPRIMANTE	Requête auprès de la base de données d'identification des impressions	—
M-EMAIL	Peuvent obtenir des données sur l'email ayant servis à envoyer le communiqué	—
M-WEB	Peuvent obtenir des données de trafic sur qui a envoyé le communiqué	—

TABLE 5 – Scores de Risques

Menaces	Flics			Civils		
	Proba.	Impact.	Risque	Proba.	Impact.	Risque
M-CAMÉRAS	5	10	50	5	5	25
M-REMARQUÉES	2	3	6	2	1	2
M-SOCIAL	2	3	6	2	3	6
M-DOIGT	3	5	15	0	0	0
M-PHYSIQUE	1	10	10	3	10	30
M-PHONE	3	10	30	0	0	0
M-IMPRIMANTE	2	3	6	0	0	0
M-EMAIL	2	5	10	0	0	0
M-WEB	2	5	10	0	0	0

D'abord en les organisant sous forme de table et en les associant avec les menaces connues (Table 6).

TABLE 6 – Contre-mesures des Extincteurs de Pneus

Menaces	Contre-mesures	Raison
M-CAMÉRAS	C-MASQUE C-CHAPEAU C-ORDINAIRE	Masques rendent non-identifiables Couvrir le reste de la tête aide à rendre non-identifiable Des habits ordinaire rendent moins uniques et donc moins identifiables
M-REMARQUÉES	C-ORDINAIRE C-NUIT	Des habits ordinaire rendent moins remarquables La nuit aide à l'anonymat
M-SOCIAL	C-PAS-SOCIAL	Éviter les réseaux sociaux pour ne pas laisser de traces
M-DOIGT	C-GANTS	Porter des gants : ne pas laisser d'empreintes sur les véhicules
M-PHYSIQUE	C-SILENCE C-ORDINAIRE C-NUIT	Ne pas faire de bruit pour éviter que des gens sortent de chez eux pour voir ce qu'il se passe Avoir l'air ordinaire veut dire qu'ils ne seront remarquées que durant la phase de dégonflage Moins de chances que les gens soient réveillés pour les confronter
M-PHONE	C-PAS-DE-PHONE	Pas de phone : pas de données de géolocalisation
M-PRINTER	C-IMPRIMANTE-PUBLIQUE	Trop de personnes différentes utilisent l'imprimante pour que leurs papiers aident à les retracer
M-EMAIL	C-PAS-D'EMAIL	Pas d'Emails : pas de retraçage d'adresse mail etc
M-WEB	C-TOR	Utiliser Tor pour envoyer le communiqué est suffisamment anonyme

Le groupe propose le protocole de sécurité suivant :

- Imprimer des tracts à leur université en utilisant l'imprimante de l'association des étudiants qui ne nécessite pas d'identifiants.
- Uniquement dégonfler les pneus de manière silencieuse la nuit quand les proprios et les voisins curieux ne sont pas là.
- Porter des vêtements lambda trouvés dans des magasins de reventes qui ne seront portées que durant les actions, et ne pas s'habiller tout en noir parce que ça attirerait trop d'attention.
- Laisser les téléphones à la maison.
- Porter des masques qui couvrent leurs visages pendant l'action.
- Porter des vêtements lambda trou-

Ayant comme but implicite d'inspirer d'autres personnes à suivre leur exemple, le groupe ne met pas dans son protocole une interdiction de parler de leurs actions à leurs amis, en plus de ça ils se disent que des rumeurs et du bouche-à-oreille ne sont pas suffisantes pour amener quelqu'un à enquêter sur eux.

Iels créent ensuite une table de risques afin de voir si les contre-mesures ont eu un impact significatif ou non (Table 7).

TABLE 7 – Scores de risque modifiés

Menaces	Flics			Civils		
	Proba.	Impact.	Risque	Proba.	Impact.	Risque
M-CAMÉRAS	5	3	15	5	2	10
M-REMARQUÉES	1	3	3	2	1	2
M-SOCIAL	0	3	0	0	3	0
M-DOIGT	3	0	0	0	0	0
M-PHYSIQUE	1	10	10	1	10	10
M-PHONE	0	10	0	0	0	0
M-IMPRIMANTE	2	1	1	0	0	0
M-EMAIL	0	5	0	0	0	0
M-WEB	2	5	10	0	0	0

Leur nouveau risque maximum est de 15 et leur nouveau risque total est de 61. C'est une diminution de 70 % de leur risque maximum, et (par coïncidence) c'est un poil au dessus de la diminution de 70 % de leur risque total. Le groupe pense que ces estimations sonnent assez bien et que ces résultats sont raisonnables, et iels décident que le protocole de sécurité est suffisamment efficace.

ANALYSE

Tous les éléments de leur protocole de sécurité vont aider à leur éviter de se faire arrêter, et en fonction de l'effort que la police met dans son enquête sur les actes de vandalisme, ça pourrait suffire. Leur protocole a quelques trucs qui pourraient être améliorés avec des efforts minimes.

Elleux n'ont pas considéré qu'il leur faudrait être prudent-es durant la phase de recherche. Si la bande fait de la reconnaissance de terrain à vélo, il faudrait le faire d'une manière qui cache leur identité au moins marginalement mais qui ne les rendent pas non plus remarquables. Si elleux font des recherches en ligne et utilisent des cartes ou google streetview, iels devraient utiliser Tor afin d'empêcher que leur adresse IP ou leurs cookies ne les relient aux espaces où iels vont accomplir leurs actions. Mais si on reste réaliste, ces contre-mesures ne sont probablement pas nécessaires.

Elleux n'ont pas non plus considéré que voyager depuis et vers leurs cibles pourrait aider à les identifier. Iels pourraient considérer de porter une tenue jusqu'à un parc, se changer dans leurs habits d'action et cacher leurs sacs, avant de continuer vers leur cible. Iels n'ont pas non plus envisagé de porter leur masque à l'approche de leurs cibles. Ceci pourrait les rendre plus remarquable, mais ça empê-

cherait que des caméras de surveillance publiques ne capturent leur visages et ne puissent être utilisées pour les identifier. Des rumeurs pourraient suffire à ce que des flics viennent toquer à leur porte, mais si d'autres contre-mesures sont appliquées avec succès et qu'ils ferment leur gueule quand la police pose des questions, il y a des chances qu'eux s'en sortent sans problèmes. Tout comme avec la prudence vis-à-vis du travail de recherche, tout ça c'est sans doute des contre-mesures non-nécessaires.

Une chose qu'eux n'ont pas considéré c'est comment réduire l'impact d'une confrontation physique. Quelqu'un en colère parce que sa voiture a été « endommagée » pourrait devenir violent. Est-ce que le groupe s'est mis d'accord sur le fait de rester sur place et de se battre si quelqu'un se fait empoigner ? Une erreur commune durant une modélisation de menace est de ne considérer que le « chemin du bonheur » et de développer des contre-mesures uniquement contre les menaces d'un scénario où toutes les choses se déroulent selon le plan. Ils devraient considérer un « chemin triste » et développer des contre-mesures pour les scénarios probables une fois que les choses dérapent.

Pareillement, ils ne se sont pas mis d'accord sur un code du silence. Si une personne se fait chopper, est-ce qu'il va balancer les autres ? Peut-être. Peut-être pas. Ça n'a pas été discuté, et ceci va mener à des tensions si quelqu'un ne se comporte pas comme ce à quoi s'attendaient les autres.

En partant du principe que tout ceci se passe dans le soi-disant Occident, il y a une répression minimale contre les actes mineurs de sabotage,²⁰ et leurs contre-mesures sont probablement suffisantes pour empêcher qu'eux soient identifiés.

UN INFOKIOSQUE MENACÉ

SCÉNARIO

Un petit infokiosque anarchiste appelé The Black Flag (TBF) est géré par un collectif du même nom. Ils vendent des bouquins, donnent des brochures et des stickers, et laissent des gens utiliser leur espace pour des réunions et événements informationnels. TBF s'est pris une descente de police parce qu'il a été accusé de distribuer du matériel séditionnel, même si l'enquête a été stoppée et qu'aucune accusation formelle n'a été émise contre le groupe. TBF a prévu une réunion afin de déterminer comment l'infokiosque devrait continuer ses opérations, et plus spécifiquement quels risques seront associés à quelle décision.

MODÉLISATION DE MENACE ET PROTOCOLE DE SÉCURITÉ

Des membres de TBF commencent leur processus de modélisation de menace en considérant ce que la police fait, plutôt que de contempler leurs objectifs (c'est à

20. À l'exception des États-Unis, où il y a de chances probables de se faire descendre par quelqu'un qui défends sa « propriété » d'une menace « terroriste ».

dire qu'ils font de la modélisation de menace orientée-vers-l'adversaire). Ce qu'ils savent c'est que :

- La police a reçu un «renseignement» (mais elle aurait pu fabriquer entièrement l'existence d'un indic ou du tuyau) à propos de matériel illégal à TBF.
- La police a fait une descente et a saisi du matériel qui pourrait être considéré comme séditieux lorsque présenté au tribunal et qui pourrait impliquer les membres de TBF dans des activités criminelles.
- Une telle descente pourrait se produire à nouveau.
- Les descentes ont fait assez peur à certaines personnes pour qu'ils ne reviennent plus.
- Le kiosque est sans doute sous surveillance (si ce n'était pas déjà le cas).

TBF a eu depuis ses débuts les objectifs explicites de répandre du matériel anarchiste et de procurer un espace où les idées anarchistes pourraient se développer et se répandre. Parce que chaque membre a une idée différente de ce qui constitue « l'anarchisme » il y a eu beaucoup d'idées différentes, certaines en conflits les unes avec les autres. Dans leur poursuite de ces objectifs et de la liberté de chaque individu, ils ont les objectifs implicites de ne pas se faire arrêter, ne pas se prendre de descentes de police, et de ne pas se faire attaquer par des fachos.

Après quelques tours de débats, TBF a quelques propositions à considérer sur comment continuer à opérer. Le collectif semble se diviser en trois factions centrées autour des décisions qu'ils estiment les plus adaptées.²¹

1. **La Faction Continuité** : Rester ouverts exactement comme avant.

- **Positif** : Pas de capitulation face à l'État, permet de continuer à procurer du matériel et de continuer à soutenir les groupes locaux.
- **Négatif** : Risque potentiellement augmenté de descentes futures, accusations, et surveillances d'individus qui utilisent l'espace.

2. **La Faction Pragmatisme** : Rester ouvert, mais changer le contenu des livres, brochures et événements pour qu'ils soient moins susceptible d'être jugés séditieux.

- **Positif** : Risque potentiellement réduit de descentes et de chefs d'accusations, quelques objets encore disponibles (mieux que rien), une stratégie long-terme meilleure que de juste laisser TBF brûler par principe.
- **Négatif** : Laisser l'État dicter quel matériel anarchiste est « acceptable » sans pour autant se débarrasser totalement du risque de descentes de flics.

3. **La Faction Fermeture** : Fermer entièrement l'infokiosque

- **Positif** : De bien plus faibles chances d'être arrêté sur des charges de sédition liées à l'existence de TBF.

21. Ne partons pas du principe qu'eux ont choisis ces noms, j'ai choisis quelque chose de modérément descriptif afin de rendre le récit plus facile à suivre.

- **Négatif** : Perte d'un espace radical, difficulté plus importante à répandre des textes, difficulté plus importante à répandre des idées.

Durant le débat, il devient clair qu'il y a des objectifs en conflit. La faction Continuité pense que le matériel le plus séditionnel est celui qui vaut le plus la peine d'être répandus (risque fort, mais récompense forte). La faction pragmatisme pense qu'un certain matériel n'est pas particulièrement utile à créer un monde anarchiste et est prête à sacrifier sa disponibilité afin de continuer à faire fonctionner un espace pour certaines formes d'anarchisme. La faction fermeture a deux sous-factions : ceux qui veulent éviter tout risque et ceux qui pensent que le collectif serait plus efficace s'il opérait de manière clandestine/couverte.

La réunion se poursuit et après quelques autres réunions il devient clair que TBF a quelques facteurs qui rendent impossible de continuer à fonctionner comme avant :

- La faction Continuité est inébranlable dans son désir de continuer à héberger le matériel le plus séditionnel, peu importe le risque.
- La faction Pragmatisme n'est pas encline à prendre des risques au prix possible de sa liberté, juste pour héberger le matériel de la faction Continuation.
- La faction Fermeture est constituée d'une moitié d'« anarchistes de bon temps » et une autre moitié qui utilisait TBF uniquement lorsqu'il s'alignait avec leurs objectifs, mais qui n'a pas de liens forts avec.

Il semblerait qu'il y ait impasse car le groupe n'arrive pas à s'aligner sur leurs objectifs, et encore moins sur les risques qu'ils sont prêts à tolérer.

ANALYSE

Chaque groupe avait raison pour des raisons différentes. Refuser d'amadouer l'État a ses mérites, et beaucoup d'anarchistes ont vaillamment publié des textes qui les ont fait atterrir en prison. Jouer à l'intérieur des lois de l'État est relativement pragmatique étant donné qu'il est bien plus dur de générer du changement depuis l'intérieur d'une prison²² plutôt qu'avec les « libertés » et « droits » qu'on nous procure sous notre soi-disant démocratie. Fermer TBF et créer des réseaux décentralisés aide à mitiger les problèmes d'organisations inflexibles et démodées (mais mettre un terme au projet pour éviter le moindre risque à tout prix c'est juste de la merde mdr).

Avec cet exemple, il pourrait sembler qu'il s'est produit un échec dans l'exercice de modélisation étant donné qu'il n'y a pas eu de consensus sur les objectifs, menaces et risques et qu'aucun protocole de sécurité n'a été produit. Elleux ne sont même pas arrivés à l'étape des scores de risque ou de sélection de stratégie. Cependant, la modélisation de menaces qu'ils ont amorcés a fait exactement ce qu'ils

22. Je ne dis pas que des prisonniers ne s'organisent pas de manières anarchistes, mais juste que leur activité est bien plus contrainte.

voulaient : elle a révélée qu'il y avait des incompatibilités entre leurs objectifs, et dans ce cas précis les objectifs dérivait de leurs grilles de lecture idéologique. Le groupe devrait probablement se diviser, même si une des factions parmi les deux qui veulent rester ouvertes va sans doute « gagner » et garder le contrôle du lieu. La modélisation de menace est utile pas seulement pour identifier et gérer les risques, mais aussi comme outil pour faciliter les conversations sur ce qu'on croit et sur ce qui est important pour nous.

QUELQUES FARCES CULOTTÉES

SCÉNARIO

Une ville se retrouve au milieu d'une guerre des farces, comprenant beaucoup de factions, y compris la plus grande faction de farceurs, les *anti-farceurs* autoproclamés. Les farces ont escaladées et en même temps les farceurs ont été capturés par les anti-farceurs afin de les empêcher de faire plus de farces. Trop chiant.

Quinn, un-e farceuse autoproclamé-e, décide de rallier une équipe avec ellui afin de faire une farce épique à un farceur notoire pas-encore-choisis. Au fil de conversations prudentes avec des ami-es et par sa connaissance de leurs alignements idéologiques et tolérances au risque, Quinn laisse entendre à quelques personnes de l'existence d'une réunion secrète dans l'espoir d'assembler une petite équipe.

Quinn rencontre l'équipe potentielle dans un parc un vendredi soir puis marche avec elleux sur quelques pâtés de maison avant d'amener tous le monde dans un bar bondé et bruyant et de les mener à une cabine au fond du bar. Quinn expose son plan approximatif sur le fait de faire une farce épique en utilisant ce qui pourrait être décrit comme les mesures de sécurités pratiquement maximales afin de se protéger durant tous le processus de la farce, depuis ses débuts jusqu'à des années après que la farce ait été accomplie. Parce que Quinn a filtré au préalable des membres potentiels, tous le monde est d'attaque, sans surprise, et elleux commencent à modéliser les menaces.

MODÉLISATION DE MENACE ET PROTOCOLE DE SÉCURITÉ

L'équipe commence par établir vite fait les phases nécessaires pour maintenir un haut de niveau de sécurité contre les anti-farceurs et autres factions de farceurs durant la farce qu'iels ont prévues. Iels décident de quatre phases qui se superposent un peu :

1. Planification : La sécurité pour quand et comment iels se rencontrent pour échanger des informations et planifier les étapes suivantes.
2. Préparation : La sécurité pour la collecte d'information et de ressources, avant la farce.
3. Exécution : La sécurité pour la farce en elle-même

4. Dissolution : La sécurité post-farce

L'équipe commence à écrire des idées sur comment iels pourraient se faire chopper durant chacune des phases, des choses qui seraient généralement vraies peu importe la stratégie qu'iels finiraient par choisir. Leur liste générale de menaces est quasi-totalement centrée autour de la surveillance et de la forensique (et en particulier la surveillance digitale). Certaines des menaces existent peut-être déjà pour chacun-e d'entre elleux car iels sont déjà connus comme proches de la mouvance farceuse.

Certaines des menaces inclues :

- Traquer la localisation géographique de quelqu'un via son téléphone mobile.
- Lire les messages de l'équipe via des téléphones ou ordinateurs piratés.
- Surveillance audio et visuelle directe des résidences des individus.
- Surveillance personnelle à travers des agents à pieds ou en véhicules (et aussi les caméras de surveillance publique).
- Poucaves et infiltré-es qui les balanceraient.

Après avoir écrit tout ça, la stratégie de Quinn pour sélectionner des espaces de réunion devient évidente. Les grands espaces bruyants offrent une couverture raisonnable face aux filatures, et ils ne peuvent pas être mis sur écoute à l'avance. Étant donné la possibilité d'être suivis, le groupe s'accorde sur le fait de toujours se mettre en route bien avant les temps futurs de réunion afin de leur donner l'opportunité de faire des drills anti-surveillance²³ durant leur trajet.

Une seconde chose qu'iels notent est que bien que certaines contre-mesures pourraient être appliquées aux différentes sortes de surveillance digitale, elleux choisissent d'opérer comme si iels se trouvaient dans un environnement « *cyber-contraint* » (C'est à dire qu'il est assumé que le risque d'être piraté ou traqué est si important qu'iels ne se permettent d'utiliser aucun outils électroniques) car cela pourrait très bien être le cas. Ceci leur évite le risque d'avoir des communications interceptées, et c'est bien plus simple que de, par exemple, se procurer de manière très prudente des burners et de ne jamais faire d'erreur avec. Tout ça combinés, iels choisissent de sélectionner leur prochain temps et lieu de réunion à la fin de chaque réunion afin que cette information ne soit jamais rendue numérique. Elleux mettent au point une phrase à utiliser pour s'alerter du besoin d'avoir une discussion privée (« J'ai soif. T'as soif? Une pinte? ») au cas où quelqu'un rate une réunion et doit être informé du lieu et de l'heure de la prochaine. Iels mettent en place une deuxième phrase secrète qui mettra un terme à la farce et sa planification si quiconque sent qu'iel est devenu trop directement surveillé-e ou qu'iel est compromi-e d'une manière ou d'une autre. Enfin, iels s'imposent une interdiction de discuter ou même

23. Les « drills » sont un terme dans l'industrie de la surveillance qui désigne des actions mises en œuvre afin de détecter si on est sous surveillance active.

de sous-entendre l'existence de l'équipe. Il est interdit d'en discuter, même de manière indirecte en disant à des ami-es qu'ils ne peuvent pas venir à un rendez-vous parce qu'ils ont « un truc secret à faire ». L'équipe n'existe pas et personne ne devrait même suspecter son existence ou que n'importe lequel de ses membres en fait partie.

Le protocole de sécurité pour la phase préparatoire est :

- Les réunions ont des points de rendez-vous pré-déterminés.
- Pour les réunions, les membres doivent laisser leurs outils électroniques à leur maison ou lieu de travail.
- Elleux doivent faire des drills anti-surveillances sur le chemin pour arriver au point de rendez-vous et elleux doivent arriver de manière ponctuelle.
- Les points de rendez-vous ne sont jamais réutilisés.
- Il y a une phrase secrète permettant de demander à recevoir en personne les détails pour la prochaine réunion.
- Il y a une phrase secrète permettant de mettre prématurément terme à la farce.
- Aucune discussion à propos de l'équipe, même indirecte, pour quelques raisons que ce soit.
- Iels devraient s'habiller de manière lambda afin d'éviter d'attirer l'attention.

L'équipe discute ensuite de la sécurité nécessaire à faire un travail de reconnaissance. Étant donné qu'ils ont déjà collectés des données sur un ensemble variés de ~~connards~~ farceurs, il y a déjà un peu d'info/renseignement sur les ordinateurs et téléphones portables des individus, et avoir accès à ceux-ci est donc acceptable. Pour plus de recherche, par exemple plus de collecte d'information à propos d'un lieu ou d'itinéraires à cartographier, iels se mettent d'accord pour utiliser uniquement Tails²⁴ sur des réseaux wifis publics aléatoires tout en laissant leurs téléphones et autres outils à la maison. Si iels doivent faire de la reconnaissance physique ou de la surveillance, elleux vont aussi laisser leurs outils électroniques à la maison, faire des drills anti-surveillance, et porter des vêtements lambdas qui cachent partiellement leur identité le plus possible sans non plus les rendre suspect-es. La règle « Pas d'outils électroniques » dans ce cas implique aussi les véhicules personnels étant donné que les automobiles modernes sont souvent connectées au réseau de téléphonie mobile afin de recevoir des mises à jour de software ou pour la cartographie intégrée. Bien qu'ils ne savent pas ce que leur farce va être, iels partent du principe qu'il y aura une enquête sur comment elle s'est produite, voire même une tentative de contre-farce, donc pour éviter de laisser des preuves, iels s'accordent à utiliser des procédures de nettoyage pour tous les objets qui seront collectés afin d'éviter les traces ADN ou empreintes digitales qui pourraient être laissées sur ces objets.

24. Le « *Amnesiac Incognito Live System* », une clé USB avec un petit système d'exploitation qui utilise uniquement Tor pour son trafic internet.

L'équipe gémit en parlant de ça parce qu'ils savent que ce sera chiant à mettre en place, mais elleux savent aussi que c'est la chose à faire.

Le protocole de sécurité pour la phase préparatoire est :

- Aucun outil électronique ne doit être utilisé (y compris les voitures) excepté Tails sur des réseaux wifi aléatoires, pour faire des recherches.
- Il faut utiliser des vêtements lambda permettant de cacher son identité, à cause de la surveillance.
- Il faut faire des drills anti-surveillance sur le chemin qui amène vers les cibles candidates et sur le chemin qui part depuis celles-ci.
- Des procédures de nettoyage sont utilisées pour réduire la quantité de preuves et traces laissées sur les objets utilisés durant la phase d'exécution.

L'équipe peut généralement comprendre de quelle sécurité iels vont avoir besoin pour la farce, donc elleux décident de finir la réunion pour se mettre à chercher qui farcer et quelles options sont disponibles. Iels font des plans pour la réunion suivante et partent chacun-e de leur côté.

Plus tard, l'équipe se retrouve après leur collecte d'information et est prête à modéliser les menaces de leur opération. Iels ont une liste de farceurs, leurs adresses personnelles et professionnelles. L'équipe regarde quel genre de farces ont déjà été faites et bien que certaines ont un risque faible, elles semblent aussi ne pas suffisamment ternir la réputation de la cible, et encore une fois, les membres de l'équipe se sont ralliés afin de mettre en œuvre une farce épique. Après avoir pesé quelques options et évalué leur faisabilité et résultats, iels décident de se faufiler dans la maison d'un des farceurs et de la redécorer rapidement.

Elleux débattent de quand cela devrait se produire. La nuit la visibilité est plus faible, mais il y a aussi moins de circulation qui pourrait gêner des poursuivants durant leur fuite et moins de foules dans lesquelles disparaître. Durant la journée, le groupe serait plus facile à remarquer, mais il est aussi plus probable que le farceur soit au travail plutôt qu'à la maison. Les caméras de surveillance semblent avoir une assez bonne résolution et bien fonctionner dans un environnement à faible luminosité ce qui fait que l'avantage nocturne pourrait disparaître. Cela dit la plupart des farces se produisent la nuit et peu de farceurs se font attraper, L'équipe raisonne que la nuit est tout de même le meilleur moment, ou peut-être le matin très très tôt. La nuit procure également l'avantage de permettre de trouver un coin sombre où se changer rapidement sans attirer l'attention. Enfin, afin d'éviter que les caméras (privées ou autres) n'enregistrent leurs mouvements de manière évidente, iels s'accordent à établir un point de rencontre où iels se changent dans les vêtements jetables propres qu'ils ont achetés dans des boutiques de revente avant de faire la dernière partie de leur trajet en direction de leur cible. Après cela, iels se séparent encore une fois et se débarrassent de leur matériel de farce et changent de vête-

ments. Enfin puisqu'il y a des chances qu'ils se fassent attraper durant la farce, ils se mettent d'accord pour débarrasser leurs lieux de résidence et de travail de tout ce qui pourrait être lié à la farce.

Leur protocole pour la phase d'exécution est :

- Nettoyer leurs lieux de vie et de travail de tout objet lié aux farces, avant la farce.
- Converger vers le point de rencontre et mettre leurs vêtements jetables propres.
- Faire la farce de nuit.
- Après la farce, se séparer et disposer des vêtements et autres objets.

L'équipe prépare enfin la phase de dissolution. Étant donné que les anti-farceurs pourraient essayer de les retrouver ou que les farceurs ciblés pourraient essayer de faire des farces en représailles, ils s'accordent à rester séparés et à faire profil bas. Elleux établissent un protocole de vérification si quiconque a l'impression d'être la cible d'une enquête. Le planning implique une vérification 3, 7, 14 et 30 jours après la farce, durée après laquelle elleux se disent que le danger sera passé. Pour finir, parce qu'il y a la possibilité que l'un-e d'elleux puisse un jour être sujet-te à suffisamment de pression pour les balancer, la bande s'accorde pour ne plus jamais discuter de la farce, à part la réunion unique trois jours après la farce qui servira à discuter de ce qu'il s'est passé. La farce n'existera plus que comme un souvenir chaleureux dans leurs cœurs et leurs consciences, qu'elleux ne sont pas restés à rien faire pendant que la guerre des farces se déroulait.

Le protocole pour la phase de dissolution est :

- Une réunion unique pour discuter de comment la farce s'est déroulée.
- Un code du silence : elleux jurent de ne jamais discuter de la farce.
- Un planning fixe de vérifications minimales afin de s'assurer que personne n'est la cible d'une enquête.

ANALYSE

Bon c'est un scénario fictif donc c'est un peu dur de dire si c'est vraiment faisable ou non. C'est aussi trop peu détaillé pour tout couvrir. Par exemple quelles sont exactement les procédures de nettoyage ? Peut-être que c'est un sujet pour une autre brochure. Peut-être qu'ils auraient pu utiliser des téléphones burners ou des radios cryptées pour la phase d'exécution ou la phase préparatoire, mais ça aurait introduit un autre type de complexité et des coûts additionnels pour se procurer ces objets. Les bons modèles suivent souvent le principe KISS : *Keep It Super Simple* (gardez les choses très simples). Ceci réduit les chances d'erreur humaine.

En tant que farceurs de longue haleine, il y a de fortes chances qu'ils soient restés informés sur les pratiques modernes de l'anti-farce et qu'elleux aient fait des farces similaires par le passé. Ils ont méticuleusement examiné tout ce qu'ils

devaient faire et considérées comment iels pourraient faire des erreurs à chaque étape et comment ces erreurs pourraient les faire attraper. Iels n'ont pas modélisés des scores de risque parce qu'il y a tellement de variations que ça n'a pas d'importance, et les conséquences d'une erreur pourraient être si catastrophiques. L'équipe a fait l'effort presque-maximum pour réduire le risque, donc comparer un score avant/après pour réduire le risque est inutile car il n'y a pratiquement rien qu'eux pourraient faire de plus. Chaque chose identifiée comme une menace a été adressée de la manière la plus complète possible. En fait, iels ont peut-être même sur-adressés certaines choses, et il est possible que leur utilisation forcée de Tails ne soit plus stricte que ce qui est nécessaire, tout comme leur utilisation de lieux aléatoires et leurs drills sur le chemin de leurs réunions. Parfois des contre-mesures qui n'offrent qu'un bénéfice minimal, en plus de contre-mesures effectivement bénéfiques, peuvent procurer un confort psychologique, tant qu'elles ne sont pas un fardeau et n'interfèrent pas avec les contre-mesures nécessaires.

Est-ce que tout ça va faire qu'iels ne se font pas chopper ou ne se prennent pas de farces en retour? Comme le dit le dicton, il est possible de ne faire aucune erreur et de quand même perdre. Leurs modèle de menace et protocole de sécurité vont les aider le plus possible, mais elleux savent que tout ceci n'est pas sans risque. C'est le jeu.

QUAND LA MODÉLISATION SE PASSE MAL

Quand on applique n'importe quelle méthode de modélisation de menaces et de création de protocoles de sécurité, il y a des classes récurrentes d'échecs. Certaines sont discutées dans les sections suivantes.

TROP MODÉLISER

Une fois qu'on commence, il peut être tentant de faire le modèle de menace le plus précis possible. On peut se perdre dans les détails et commencer à se focaliser obsessivement sur toutes les possibilités adjacentes possibles au fur et à mesure qu'on cherche à réduire le risque à zéro. Ceci ne sera jamais possible, et il y aura toujours du risque. Cette paralysie peut être surmontée en visant des objectifs qui sont minimalement réprimés et ensuite de persévérer vers des objectifs plus grands. Sticker des endroits et voler dans les magasins peut aider à s'habituer à agir en dehors de la loi. Graffer dans la rue, dérouler une bannière dans un lieu public et dégonfler des pneus de SUV apprend à agir la nuit de manière cachée. Si tu te retrouves pétrifiée en pensant à la force de tes adversaires, envisage d'abord de faire des actions dont on sait qu'elles sont minimalement réprimées.

PROTOCOLES INFAISABLES

On peut créer un modèle de menace pratique et produire un protocole de sécurité efficace qui permettraient que des actions soient entreprises avec une répression minimale. Si le nouveau protocole nécessite bien plus d'efforts que le protocole actuel il y a peu de chances que la mise en place complète du protocole soit faisable. L'oubli et les vieilles habitudes sont de réels problèmes, et tenter de changer totalement notre comportement d'un seul coup n'est généralement pas possible. Il y aura des lacunes et des erreurs. Les protocoles pourraient nécessiter d'être mis en place lentement, soit une pièce à la fois ou bien avec une complexité grandissante, et ceci pourrait vouloir dire sélectionner des objectifs ou tactiques qui présentent moins de menace, pour commencer. Essayer de faire se rencontrer un groupe la nuit dans un endroit et à un moment précis sans téléphone pourrait déboucher sur un résultat médiocre. Peut-être qu'il serait plus facile de s'entraîner à se rencontrer en utilisant ces stratégies afin de s'assurer que les gens sont ponctuel-les et capables de se repérer sans cartes connectées.

MANQUE D'ANTICIPATION

Un autre échec consiste à entraver ou fermer des stratégies futures — ou à créer une quantité intolérable de risque — en sélectionnant des protocoles de sécurité qui laissent à la merci des méthodes actuelles de surveillance ou répression. Des actes de désobéissance civile qui mènent à une arrestation des personnes pourraient rendre plus difficile à ces mêmes personnes d'entreprendre des actions futures due à une pénalité plus importante ou bien au fait que les données biométriques de tel ou tel individu se retrouvent stockées dans une base de donnée de police. On pourrait décider qu'il n'y a pas de risques associés au fait de parler de crimes potentiels parce qu'on ne compte pas les faire réellement... jusqu'à ce qu'on se rende compte un jour que les commettre est la marche à suivre, la chose à faire. Une question que l'on devrait toustes se poser c'est si nos protocoles de sécurité présents vont être dommageables aux nous futur-es si on veut un jour pouvoir aller plus loin que ce qu'on fait aujourd'hui.

Pareillement, une stratégie pourrait demander un protocole de sécurité assez relaxé mais choisir de n'utiliser qu'une sécurité minimale pour les menaces anticipées peut potentiellement contraindre l'action jusqu'à un degré problématique. Amener un téléphone à une manif qui est sensée être calme pourrait t'empêcher de faire des actions radicales si la police devient agressive. Quand on regarde un protocole, on devrait réfléchir à s'il y a des chances (raisonnablement parlant) qu'on veuille éventuellement agir encore plus si la situation change, et si agir d'une telle manière créerait beaucoup de risques à cause de contre-mesures qu'on n'a pas mis en place.

MODÉLISATION AU LIEU D'ACTION

La modélisation de menace est un exercice lent quand elle est entreprise en profondeur mais parfois on doit agir rapidement. On pourrait être au milieu d'un coup d'état, ou bien des fascistes pourraient être en train d'arpenter les rues, et on ne pourra s'appuyer que sur les procédures opératoires standards. Il n'y aura peut-être pas le temps de modéliser la situation, et même une évaluation mentale rapide pourrait nous laisser incapable d'agir si on vise à garder notre niveau total de risque à un niveau confortablement bas. Les marées du risque montent et chercher à maintenir notre niveau actuel visible de sûreté va nous amener, éventuellement, à devenir totalement inactives. Tout ne peut pas être modéliser, et il y a des moments où la seule chose à faire c'est agir de manière décisive, courageuse et rapide. Avoir établis des pratiques de sécurité peut procurer quelque chose vers quoi se tourner rapidement quand on a besoin d'une solution rapide et probablement correcte à une situation nouvelle.

MODÈLES DE MENACE FAITS SUR MESURE

Les crews d'anarchistes se retrouvent souvent à réinventer la roue quand on en vient à la théorie, les tactiques et les stratégies organisationnelles. Ceci n'est pas juste le cas en ce qui concerne la sécurité. Parfois ceci vient d'une certaine naïveté au sens de ne pas avoir été exposé aux idées et textes pertinents. Dans d'autres cas c'est plutôt une conséquence du fait d'être arrogante et de penser que notre équipe est si spéciale que les idées appliquées par d'autres ne pourraient possiblement pas s'appliquer à notre situation unique. Les personnes qui pensent ça vont partir de leur côté et essayer de produire des pratiques et idées à partir de zéro. Sur des sujets discursifs comme « quelle est la nature même de l'anarchisme » ceci pourrait mener à une convergence vers des idées pré-établies. Avec un sujet plus technique comme la sécurité, qui dépend d'une compréhension des technologies sous-jacentes et de l'observation des actions des agences de police, partir de zéro tends à avoir une convergence très lente vers des pratiques établies et vérifiées, quand il y a convergence tout court.

Autant que possible, dépendre de modèles de menace et protocoles de sécurités faits sur mesure devrait être évité. Bien que des menaces spécifiques pour une classe donnée de personnes dans un espace et temps donné puissent être assez différentes d'autres, il y a des recoupements assez significatifs entre eux. En ce qui concerne les technologies d'information comme les ordinateurs, téléphones et l'internet même, il y a une uniformité générale dans les menaces que l'on rencontre à l'échelle globale. Souvent ce qui est nécessaire c'est de voir si des menaces sont vraiment présentes dans une zone et d'ensuite connecter des contre-mesures connues à ces menaces. Les réponses sont souvent assez directes, et des contre-mesures et une sécurité compliquées pourraient sonner super cool et méga-james-bond mais elles sont souvent basées sur de mauvaises compréhensions de comment la police ou les technologies

opèrent, et leur complexité peut devenir source de fierté. Il y a un certain prestige à savoir et à démontrer des connaissances de niche, et des équipes pourraient se sentir supérieures aux autres car elles ont développées des protocoles customisés que personne d'autre n'a.

Vous devriez utiliser, autant que possible, les connaissances des autres, afin d'enrichir votre modèle de menace. Déterminer quels modèles et bibliothèques de menaces intègrent une compréhension précise du monde peut être difficile, ce qui fait de la vérification est une tâche nécessaire récurrente. Vérifier et synthétiser des connaissances existantes est bien plus facile et mène à des modèles et protocoles bien plus simples que d'essayer de tout conjurer depuis la cinquième dimension.

IGNORER LES RISQUES FUTURS

Le monde est en train de devenir de plus en plus dangereux, surtout pour les radicaux. Les droits les plus basiques de manifester que l'État nous offre si généreusement sont lentement en train d'être grignotés, et même la désobéissance civile la plus minimalement dérangeante est lentement mais sûrement criminalisée. De nouvelles menaces à l'action et au militantisme radical apparaissent, la fréquence avec laquelle des menaces existantes sont déployées est en augmentation, et l'impact de pratiquement toutes les menaces est en train de devenir plus significatif. On ne peut pas exagérer le fait qu'ignorer les risques futurs en maximisant la sûreté dans le présent est une conduite désastreuse. Ceci ne veut pas dire qu'on devrait, je cite « tout cramer » demain en agissant de manière extrêmement risquée, mais plutôt que nous sommes aujourd'hui — très probablement — face à bien moins de répression que ce que nos futures-nous devrons affronter dans 10 ou même 5 ans. Construire le protocole de sécurité « parfait » qui navigue lentement et prudemment des eaux troubles pourrait tout de même nous laisser naufragés à la tempête suivante. Éviter le risque maintenant en entreprenant des actions moins fréquentes ou moins intenses veut juste dire reléguer ce risque au futur.

Cet évitement du risque peut se produire de plusieurs manières. On pourrait éviter les « grosses » choses parce qu'elles sont très fortement criminalisées, mais on pourrait aussi se retrouver à tordre nos tactiques existantes pour éviter une répression mineure. On pourrait être légèrement moins ouvertes aux gens inconnues avec l'espoir de décourager les infiltrations, mais ceci nuit à nos capacités présentes et futures. On pourrait éviter les démonstrations de solidarités avec des groupes criminalisés comme les mouvements kurdes afin d'éviter d'être poursuivis pour « soutien à des organisations terroristes » comme c'a récemment été le cas en soi-disant Suède. Éviter les risques est une position privilégiée, et la solidarité veut dire prendre sur soi quelques risques qui sont dirigées vers des groupes marginalisés. La modélisation de menaces nous indique quels risques existent, mais si on vise à minimiser le risque pour nous-mêmes, on a perdu un élément-clé de ce qui fait de l'anarchisme une idéologie qui en vaut la peine : l'altruisme et l'entraide.

Sous-estimer les risques répétés

Les humains ne sont pas très bons aux statistiques. Considérez le scénario suivant.

Une équipe de 6 personnes fait le même type d'action de manière répétée. Chaque fois qu'elleux la font, chaque membre a 0.5 % de chances de se faire attraper. Combien de fois est-ce qu'ils peuvent le faire avant qu'il y ait 50 % de chance que quelqu'un se fasse attraper ?

Vingt-trois.²⁵

Quelque chose avec une probabilité d'impact petite et presque négligeable qui est faite une semaine sur deux mènerait à des chances de 50/50 qu'au moins une d'entre elleux se fasse attraper en moins d'un an. Ceci ne veut pas dire qu'ils sont en sécurité la Nuit 1 et qu'ils se feront définitivement attraper la Nuit 23. Elleux pourraient se faire attraper la Nuit 1 ou 100, mais les chances que ça se produise « exactement la Nuit 1 » ou « jamais avant la Nuit 100 » sont toutes les deux assez basses.

Cet exemple est assez simpliste et on pourrait argumenter que plus on fait quelque chose plus on devient doué à cette chose, mais inversement l'habitude pourrait rendre inattentif-e ou complaisant-e, ou bien les flics pourraient avoir accumulés assez de preuves pour les chopper. Même toute cette discussion de modèles est problématique parce qu'on ne sait pas réellement tout ce qu'il se passe. La majeure partie de ce qu'on fait c'est deviner. On ne sait pas ce qu'est la probabilité réelle d'impact. Dans l'exemple au-dessus, si la probabilité n'est pas de 0.5 % mais de 1 % alors l'équipe ne peut faire l'action que 12 fois avant d'avoir plus de 50 % de chances de se faire attraper. Est-ce que c'est réellement précisément possible d'estimer la différence entre des probabilités de 0.5 % et 1 % ? Probablement pas. Peut-être que la probabilité est totalement différente, comme entre 0.1 % et 3 % (115 actions ou 4 actions avant que les chances ne montent à 50/50, respectivement)

Estimer des probabilités est *très* difficile et ressentir l'accumulation de risque n'est pas très intuitif. En général, c'est probablement utile de partir du principe que quelque chose qui a une probabilité d'impact rare, quand c'est fait de manière répétée, va presque certainement obtenir une probabilité d'impact si on s'attend à ce qu'il y ait 10 occurrences ou plus.²⁶

25. Ceci est une situation triviale de distribution binomiale où on pose la question « quelles sont les chances que ça ne se produise jamais ? » La formule qui nous offre réponse est : $(1 - 0.005)^6 \cdot 23 \approx 0.5007$

26. Pourquoi 10 ? Honnêtement on sort des nombres du chapeau pour tant de choses et on n'aura jamais assez de données pour un modèle « scientifiquement » juste, mais ça paraît raisonnable. 10 c'est « quelques » mais pas « une tonne ».

RÉPÉRCUSSIONS DE LA PART DE PAIRS ET DE L'ÉTAT

Quand on fait du militantisme il y a des règles implicites sur ce qui est considéré comme « correct » dans un milieu donné. Ceci peut porter sur des sujets comme la méthodologie pour organiser des actions, quelles méthodes les collectifs peuvent utiliser pour le consensus, ou quelle sorte de communiqués/appels sont produits quand de nouveaux sujets importants émergent. On est constamment en train de considérer quelles actions prendre — à tous les niveaux — en fonction du jugement de nos congénères. Quand on nous *call-out* — c'est à dire quand on fait face à des répercussions sociales — on se replie sur la théorie ou bien on affirme que ce qu'on a fait est totalement aligné avec le protocole, comme si ça nous exonérait (et parfois c'est le cas).

Ce n'est pas le cas avec la sécurité. En définitive, ce qui est important c'est si nos objectifs sont atteints et si on ne se fait pas prendre. Quand on suit un protocole de sécurité qui est largement utilisé dans notre milieu, si on reçoit l'approbation de nos pairs, et qu'on finit quand même en prison, le fait que nos pairs aient trouvés nos actions impressionnantes n'a aucune importance.

Un protocole de sécurité n'est pas là pour permettre d'esquiver les critiques. Il n'est pas là pour apaiser les demandes des autres. Il existe pour vous garder (relativement) en sûreté pendant que vous réalisez les actions qui vous amèneront à vos objectifs. Ça sonne comme une évidence mais ça peut vraiment être un changement profonds pour beaucoup de personnes. Quelqu'un pourrait se faire chopper et crier que « c'est pas juste » qu'iel se soit fait attraper parce qu'iel a « tout fait comme il fallait ». C'est sans importance si ellui a utilisé le protocole comme tous le monde. Iel s'est fait choppé, ce qui veut dire que le protocole n'était peut-être pas suffisamment sécurisant pour vraiment le protéger ou bien que l'action était inséparable d'une grande quantité de risque. Et encore une fois, ce qui est considéré comme « juste » par nos pairs comparés à ce qui crée réellement de la sécurité pourrait ne pas être aussi fortement corrélé que l'on pense.

Ton protocole de sécurité n'est pas là pour créer une image visant l'approbation des autres. C'est une réelle question de sécurité. N'oublies pas ça.

REMARQUES FINALES

Quand on apprend à dessiner, on ne produit pas immédiatement des illustrations photoréalistes. On fait plutôt des têtes déformées ou des créatures aux proportions étranges face à un paysage bordélique. Avec le temps et la pratique, on apprend à représenter plus précisément nos sujets. Les formes deviennent plus représentatives, et des détails sont ajoutées ou retirées quand il le faut afin d'aider la focalisation.

Quand tu apprend la modélisation de menace, ton modèle pourrait être grossier et simple, mais au fur et à mesure de ta pratique, il devient plus compréhensif et

réaliste. C'est aussi OK de ne jamais développer d'élégance. Comme quiconque qui a déjà joué à pictionary le sait, un dessin rapide qui comporte uniquement les détails les plus importants est suffisant à gagner une manche. Parfois des gens qui ne font que des brouillons dégueulasses réussissent à battre des artistes qui s'attardent à produire des représentations complètes. Un modèle de menace simple qui identifie quelques trucs importants peut être supérieur à une toile monstrueusement complexe de toutes les contingences possibles.

La modélisation de menace est un processus itératif. Peut-être qu'il n'y a qu'une seule itération et qu'on décide que c'est suffisant, ou bien c'est quelque chose vers lequel on revient de manière continue. Les modèles de menaces des autres peuvent informer les nôtres, surtout vis à vis du renseignement qu'ils ont collectés sur des adversaires communs. Tu pourrais affiner ton modèle de menace périodiquement ou après que tu ais observé la moindre déficience. Ton équipe pourrait avoir besoin de revisiter le sien si iels vivent une escalade dans leurs activités ou quand un nouveau membre les rejoint et apporte ses propres observations et questions. Un modèle de menace statique est en soi une menace car au fur et à mesure qu'il expire, il perds en efficacité. Parce que le processus est itératif, il est mieux de travailler dessus ensemble et de l'appliquer au monde réel et d'itérer, plutôt que de passer des mois à essayer de créer le modèle « parfait ».

Il y a quelques extrêmes dont tu devrais te méfier quand tu discutes de modélisation de menace et de sécurité en générale. Les maximalistes de la sécurité partent du principe que sans un modèle de menace robuste et un protocole de sécurité sans failles, n'importe quelle action mènera à une arrestation ou à un emprisonnement. Elleux tendent à partir du principe que tous les adversaires sont ou seront intéressés par toi et finiront par agir en utilisant leurs capacités maximales. Traîner avec des maximalistes peut être anxiogène parce que peu importe ce qu'on fait en termes de sécurité ce n'est jamais assez et leur paranoïa a tendance à entraver l'action. De l'autre côté il y a des minimalistes de la sécurité qui avancent que tous le monde est trop paranoïaque et que même faire de la modélisation est un exercice ridicule. Elleux partent du principe que si enquête il y a, ça sera uniquement la police locale qui est elle-même trop stupide pour découvrir ce que quiconque a prévu de faire. Les nihilistes de la sécurité pensent — comme les maximalistes — que nos ennemis sont inhumainement puissants, mais au lieu d'essayer de gagner la course à la sécurité contre eux, les nihilistes disent qu'aucune quantité de sécurité ne peut surmonter la menace, donc pourquoi même essayer ?

Je fais en sorte de ne pas m'informer sur les activités concrètes des gens, donc je ne peux pas l'affirmer avec certitude, mais sur la base de mon intuition et d'anecdotes il me semblerait que les maximalistes, minimalistes et nihilistes de la sécurité ne font pas grand-chose de conséquent. Les gens qui font des choses réellement intéressantes tendent à avoir une perspective plus nuancée de la sécurité. C'est à dire que vous devriez considérer les opinions extrêmes avec une certaine prudence.

Peut-être ce qui est le plus important c'est que la modélisation de menace et les

protocoles de sécurité générés sont uniquement utiles s'ils peuvent être appliqués à la vie réelle et **s'ils aident à accomplir nos objectifs**. Un protocole de sécurité qui est trop encombrant pour être réellement exécuté est inutile peu importe à quel point il est bien conçu. Une exécution partielle du protocole parce qu'on est trop pris par notre quête d'accomplissement de nos objectifs peut créer plus de risque que ce que l'on avait prévus. Pareillement un protocole facilement conçu et appliqués qui rends totalement impossible d'arriver à ses fins indique une erreur de conception ou une erreur stratégique, ou bien il indique que nos objectifs et notre tolérance au risque ne sont pas compatibles. Ceci pourrait vouloir dire recommencer le processus de modélisation de menace, changer d'objectifs ou trouver des manières de s'habituer au risque et de construire de la tolérance.

Avec un peu de chance, en lisant tout ça vous pouvez voir que des processus structurés pour évaluer les risques existent. Le monde du secret et de l'insurrection dans lequel nous opérons est complexe. Nous sommes constamment enveloppés dans le brouillard de la guerre. La modélisation de menace peut aider à percer ce voile et nous donner quelques perspectives pour que nos actions soient accomplies avec plus de confiance en nous. On pourrait tout de même se faire attraper et il n'y a pas de monde où on évite toutes la répression. Les enjeux sont de taille, mais appliquer un peu de connaissance peut faire pencher la balance dans notre sens.

On demande à nos camarades « est-ce que c'est OK d'utiliser des téléphones portables ? » et elleux répondent « ça dépend de ton modèle de menace » avant de lister des scénarios qui pourraient nous mettre en danger. À partir de ces interactions on a une idée de comment un modèle de menace est utilisé, mais comment ils sont créés est généralement moins clair. On devrait pas laisser toute notre sûreté à la discrétion des autres, et donc on devrait apprendre à gérer les risques auxquels on fait face. Cette brochure couvre les méthodes que tu peux utiliser pour modéliser des menaces par toi-même et comment ces étapes explicites peuvent te mettre, toi et les personnes autour de toi, plus en sûreté. Notre sécurité est à la hauteur de la justesse de nos modèles, donc posons-nous et pensons très fort à comment garder les flics et les fascistes à distance au fur et à mesure qu'on avance vers un monde libéré.

