



MYTHES ET LEGENDES DES TIC ***TOME 2***

10 décembre 2018

Collection ATENA



Une réalisation de Forum ATENA avec la collaboration de *(par ordre alphabétique)* :

Aurélie Banck, Laurent Bloch, Jean Marie Corriere, Louis Derathe, Jean-Baptiste Fahy, [Alessandro Fiorentino](#), Jean-Denis Garo, Jean-Yves Gresser, Ivan de Lastours, Hervé Lehning, Jean-Marc Do Livramento, Francesca Musiani, Louis Pouzin, Nicolas Ruff, Agnes Teissier, Viken Toramania

Livre collectif sous la direction de Gérard Peliks

Les ajouts depuis mai 2018 apparaissent en bleu

Copyright forum ATENA – Voir en dernière page les droits de reproduction

INTRODUCTION

Ce document est le début du tome 2 du livre collectif "Mythes et légendes des TIC" développé dans le cadre de l'association Forum ATENA.

Si vous désirez obtenir la version complète PDF du tome 1 du livre "Mythes et légendes des TIC", il est téléchargeable en :

<https://www.forumaterna.org/mythes-et-l%C3%A9gendes-des-tic-tome-1>

Si vous désirez commander sa version papier c'est sur :

<http://www.lulu.com/product/couverture-souple/mythes-l%C3%A9gendes-des-tic/15739496>

Si vous êtes intéressés d'être tenus au courant de ses développements, voire si vous désirez devenir un des auteurs, demandez-le moi par e-mail.

Gérard Peliks

gerard.peliks@forumaterna.org

Président de l'atelier sécurité et VP de Forum ATENA

SOMMAIRE

Table des matières

TAPEZ LE TITRE DU CHAPITRE (NIVEAU 1)	1
MYTHES ET LEGENDES DES TIC TOME 2.....	1
INTRODUCTION	2
L'ETHIQUE.....	5
QUI AURAIT ENVIE DE TRAVAILLER DANS LES METIERS DE LA SECURITE DU NUMERIQUE ?	6
L'IMPASSE DES NOMS DE DOMAINE	12
1° PARTIE : ASPECTS INFORMATION ET SYSTEMES D'INFORMATION	14
MYTHES ET LEGENDES DU SYSTEME D'INFORMATION (SI).....	15
MYTHES ET LEGENDES DU REGLEMENT GENERAL SUR LA PROTECTION DES DONNEES.....	23
MYTHES ET LEGENDES DU PRIVACY BY DESIGN	32
MYTHES ET LEGENDES DU MULTILINGUISME ET DU PLURILINGUISME DANS L'INTERNET	37
MYTHES ET LEGENDES DE LA BLOCKCHAIN	43
MYTHES ET LEGENDES DU BIG DATA.....	47
MYTHES ET LEGENDES DE LA 4G	54
MYTHES ET LEGENDES DU PEER-TO-PEER	70
MYTHES ET LEGENDES DE LA MOBILITE EN ENTREPRISE	76
MYTHES ET LEGENDES DE LA VIDEOCONFERENCE.....	78
2° PARTIE : ASPECTS SECURITE ET SURETE	82
MYTHES ET LEGENDES DES CHIFFREMENTS HOMOMORPHES.....	83
MYTHES ET LEGENDES DES APT	87
MYTHES ET LEGENDES DE L'ANALYSE DE RISQUE.....	92
MYTHES ET LEGENDES DE LA GUERRE DANS LE CYBERESPACE	98
MYTHES ET LEGENDES SUR LA CONFIANCE NUMERIQUE	101
MYTHES ET LEGENDES DE LA CONFIANCE DANS LA PROTECTION DES DONNEES	106
3° PARTIE : ASPECTS PHYSIQUES.....	108
MYTHES ET LEGENDES DES DANGERS SANITAIRES DE LA TELEPHONIE MOBILE	109
4° PARTIE : ASPECTS METIERS.....	115
MYTHES ET LEGENDES DE LA RESPONSABILITE SOCIETALE DE L'ENTREPRISE INFORMATIQUE QUI GENERE DE LA VALEUR.....	116
MYTHES ET LEGENDES DU SOCIAL SHOPPING	121
ACRONYMES.....	124

GLOSSAIRE	126
POUR ALLER PLUS LOIN DANS LA CONNAISSANCE DES TIC	129
WEBOGRAPHIE :	132
A PROPOS DES AUTEURS	134

L'ETHIQUE



QUI AURAIT ENVIE DE TRAVAILLER DANS LES METIERS DE LA SECURITE DU NUMERIQUE ?

Gérard Peliks, Président de l'atelier sécurité et VP de Forum ATENA

Personne n'a envie de travailler dans les métiers de la sécurité du numérique, surtout pas les femmes qui seraient d'ailleurs incompétentes car leurs gènes ne les prédisposent pas à envisager des carrières uniquement techniques. Cela coupe déjà les métiers de la sécurité du numérique de 50% de la population.

UN METIER PUREMENT TECHNIQUE, APRES DES ETUDES INGRATES

Le métier de la sécurité du numérique est donc éminemment un métier technique. Dès le plus jeune âge, il faut aimer les affres des mathématiques et de la physique, puis se diriger après le bac « S » vers les prépas scientifiques, pour réussir après deux ou trois ans d'un travail de forçat, un improbable concours d'entrée dans une école d'ingénieurs, ou de tenter sa chance dans les écoles d'ingénieurs qui recrutent leurs élèves dès le bac... « S ». Cela impose un travail assidu pendant cinq ou six années de sa jeunesse. Pourquoi alors ne pas opter plutôt pour des prépas commerciales ou sciences politiques, ou pour la FAC ? Qui est intéressé par les filières d'écoles d'ingénieurs et même par un enseignement scientifique et technique ? Presque personne, et c'est normal, car qui pourrait avoir envie de passer sa jeunesse à essayer de comprendre la complexité du monde, et surtout celle du cybermonde pour avoir de bonnes notes ? On préfère travailler au soleil que dans les nuages. Le Cloud ? Non, trop fumeux, pas pour moi.

Cela coupe le métier de la sécurité du numérique de 90% de la population qui s'engage vers les études de l'enseignement du supérieur, et qui se dirige plutôt vers le droit, les sciences humaines ou la médecine.

Son diplôme brillamment obtenu, vers quels métiers se diriger ? Celui de la sécurité du numérique ? Surtout pas ! On n'a pas subi tant d'années d'études pour coder toute sa vie en Python, en Java ou en C++, ou pire, pour voir défiler des lignes d'alerte sur une console, du matin au soir, et pire encore, du soir au matin, avec pour seul rôle de réagir en cas de soupçon d'avoir repéré une ligne traduisant peut-être le début d'une cyberattaque, tout en étant menacé d'être mis à la porte si on signale un « faux positif » !

Cela coupe le métier de la sécurité du numérique de la quasi-totalité de la population active. Qui pourrait accepter d'aller vers cette profession ingrate et pas vraiment reconnue par son employeur ? De plus, c'est un métier peu valorisant côté business, et donc côté avenir.

UN METIER PEU VALORISANT, MAL PAYE, A RISQUE ET SANS AVENIR

Pas valorisant côté business et côté avenir, le métier de la sécurité du numérique ? En effet la sécurité présente un coût pour l'entreprise, toujours trop élevé, et absolument pas productif, qui ne rapporte rien et qui coûte beaucoup. Ah si on pouvait s'en passer ! En voilà une belle variable d'ajustement à la baisse pour pouvoir augmenter le budget des activités qui rapportent, comme par exemple le marketing et le commercial, voire pour attribuer plus de budget à la technique si les clients jugent que les produits que l'entreprise fabrique ne donnent pas entièrement satisfaction. La destruction ou la compromission ou pire l'indisponibilité des

informations de l'organisation ? Non ça ne peut nous arriver, nous ne sommes pas une organisation qui travaille dans un domaine stratégique, et nous ne nous connaissons pas d'ennemis. Le RGPD ? C'est trop compliqué et c'est pour les autres, et on n'a pas le temps ni les moyens de s'y intéresser.

Qui reste-t-il ? Une seule personne, le RSSI (Responsable de la Sécurité du Système d'Information). Le RSSI a accepté ce métier à risques, peu rémunéré, très méprisé, et qui en plus produit toute sortes d'interdictions d'utiliser le cyberspace comme on le voudrait pour être efficace dans son travail. Et de plus, le RSSI c'est celui qui multiplie les gênes en imposant des mots de passe trop compliqués, trop longs, trop difficiles à mémoriser, et de plus qui doivent trop souvent être changés, et pas toujours au même moment alors qu'il en faut un différent pour chaque application et qu'il ne faut pas les mettre sur des post-it ! Le RSSI, c'est celui qui accède à vos fichiers à caractère personnel avec la complicité de la Direction.

Pas étonnant dans ces conditions que le RSSI est celui qui mange tout seul à la cantine !!!

C'ETAIT JUSTE POUR DECRIRE QUELQUES MYTHES ET LEGENDES

Pas réjouissant ce qui précède pour le métier de la sécurité du numérique ? En effet, sauf que... **tout ce qui précède est entièrement faux et digne du domaine des légendes urbaines. Le métier de la sécurité du numérique n'est absolument pas ça, et c'est même tout le contraire !**

Le ton pour cet ouvrage collectif « Mythes et légendes des TIC » est donné. Maintenant venons en à la réalité

LA SECURITE DU NUMERIQUE, UN METIER POUR LES HOMMES... ET POUR LES FEMMES

Balayons déjà la légende : La sécurité du numérique est un métier d'homme. S'il est vrai qu'aujourd'hui les femmes ne représentent qu'une trop faible partie des experts sécurité, il n'y a aucune raison pour qu'elles se privent de ces métiers si intéressants. Mesdames, les métiers de la sécurité du numérique vous attendent et vous ne le regretterez pas. Des associations existent pour inciter les femmes à aller vers ces métiers. Citons le CEFCS (Cercle des Femmes pour la CYberSécurité) www.cefcys.fr dont les deux principales composantes de la sécurité, le côté technique représenté par sa présidente, Nacira Salvan, docteure en informatique et le côté juridique représenté par sa vice-présidente Maître Isabelle Landreau, docteure en droit, sont présentes parmi d'autres membres éminents de cette association.

Donc ça nous fait pour faire bref et synthétique, 100% de la population susceptible d'être intéressée par les métiers de la sécurité du numérique.

LES METIERS DE LA SECURITE DU NUMERIQUE NE SONT PAS QUE TECHNIQUES

Deuxième légende urbaine, les métiers de la sécurité du numérique ne permettent d'accéder qu'à des carrières uniquement techniques. C'est faux bien sûr, et on peut même remarquer que la sécurité emmène de plus en plus de questions juridiques. Au siècle dernier, dans les événements sécurité, il n'y avait dans l'assistance que des participants d'origine technique et les sujets purement techniques étaient traités par des experts forcément techniques. Ce n'est plus le cas aujourd'hui. Dans les salons et les conférences sécurité, on parle de droit, d'intelligence économique, de géopolitique et de divers sujets intéressant les sciences humaines. C'est dire que la sécurité s'est ouverte à toutes sortes de carrières et que l'expert

sécurité doit avoir des connaissances techniques certes, mais aussi juridiques, intelligence économique, gestion de projet...

Ce qui remet en selle tous les étudiants qui s'engagent vers des études de l'enseignement du supérieur, mais pas forcément dans des écoles d'ingénieurs, des instituts et facultés de sciences.

L'ENSEIGNEMENT DE LA SECURITE DU NUMERIQUE EST INGRAT ?

Balayons aussi le fait que s'engager dans des études scientifiques, voire dans des prépas et des écoles d'ingénieurs est un lourd tribut à payer à sa jeunesse pour prétendre accéder à des fonctions intéressantes. Les sciences exactes, dont les maths, en général, ce n'est que quand on les voit de loin qu'elles peuvent sembler rébarbatives. C'est au contraire un domaine merveilleux à explorer, où le génie humain montre qu'il n'a pas de limites. Cela implique un travail personnel conséquent ? Oui, mais on n'a rien pour rien et quand vous êtes admis dans une filière qui fera de vous un collaborateur apprécié et plein d'avenir, quel plaisir d'avoir acquis les connaissances de base pour faciliter votre carrière !

LE REFERENT SECURITE, UNE PROFESSION PROCHE DE LA DIRECTION

Le métier du numérique, c'est de surveiller ce qui passe sur votre écran, en tant que travailleurs dans un « Secure Operation Center », ou de contrôler ce que font les autres employés et d'être l'empêcheur de tourner en rond ? Non, le RSSI, on parle de plus en plus de CISO (Chief Information Security Officer) car la sécurité n'est pas celle seulement de son système d'information mais surtout de celle de son Information, est le garant de la sécurité des données numériques de l'organisation, c'est celui qui diminue les risques qui pèsent sur l'information des organisations, de toutes les organisations y compris de celles qui ne pensent pas être menacées. La bonne place pour un responsable sécurité du numérique n'est pas d'être englouti dans la direction informatique mais d'être très proche de la Direction générale, voire d'avoir une place de choix dans le COMEX où sont prises les décisions stratégiques. Cela implique donc beaucoup de responsabilités, un travail très varié et une rémunération qui n'a rien à envier par exemple aux métiers de la finance.

LE RGPD EMMENE DE NOUVEAUX METIERS

Et soit dit en passant puisqu'on a parlé du RGPD, ce Règlement Général pour la Protection des Données personnelles concerne toutes les organisations, dans les 29 pays d'Europe et dans les autres pays quand la sécurité des données personnelles d'un citoyen européen est mise en cause. Le RGPD est là pour établir la confiance de tous dans la protection des données personnelles qui sont très menacées aujourd'hui parce qu'elles sont volées par dizaines de millions. Et de la confiance des clients, des prospects et des employés naît la productivité et favorise l'accroissement du business. Le RGPD induit d'autres métiers comme celui de DPO (Data Protection Officer).

Ce qui peut donner du travail très intéressant à 100% de la population, à condition que celles et ceux qui veulent faire carrière dans les nombreux métiers de la sécurité, qu'ils soient techniques ou sciences humaines, qu'on soit un homme, qu'on soit une femme, soient bien formés.

LES LABELLISATIONS POUR BIEN CHOISIR SA FORMATION

Il existe une labellisation gérée par l'ANSSI, la labellisation SecNumEdu attribuée à des formations qui traitent spécifiquement de sécurité du numérique. Et parce que cette sécurité est l'affaire de tous, il existe une autre labellisation attribuée aux formations qui ne sont pas

spécifiquement des formations de sécurité du numérique mais qui intègrent suffisamment de modules de sécurité pour former les cadres de demain, conscients des périls du cyberspace. C'est la labellisation CyberEdu, attribuée par l'Association CyberEdu créée en 2016 par l'ANSSI.

COMMENTAIRES

Voici des commentaires sur ce chapitre écrits par Jean-Marc Do Livramento (voir son chapitre sur les mythes et légendes de la 4G)

Balayons déjà la légende : La sécurité du numérique est un métier d'homme.

S'il est vrai qu'aujourd'hui les femmes ne représentent qu'une trop faible partie des experts sécurité, il n'y a aucune raison pour qu'elles se privent de ces métiers si intéressants. Mesdames, les métiers de la sécurité du numérique >vous attendent et vous ne le regretterez pas.

Oui, c'est vrai. 100% des femmes qui font de la sécurité sont... des femmes. Pour autant, faut-il les pousser à faire de la sécurité par la force des baïonnettes ? Laissons-les choisir leur métier et ne pas vilipender celles qui ne veulent pas faire de sécurité. Aujourd'hui, 98% des emplois sont ouverts aux femmes (quelques emplois notamment dans le domaine militaire restent fermés) et toutes les femmes peuvent les choisir si elles le souhaitent. "Les métiers d'homme" est un débat du millénaire précédent et il n'y a plus que les hommes pour le remettre sur la table.

Deuxième légende urbaine, les métiers de la sécurité du numérique ne permettent d'accéder qu'à des carrières uniquement techniques. C'est faux bien sûr, et on peut même remarquer que la sécurité emmène de plus en plus de questions juridiques. Au siècle dernier, dans les événements sécurité, il n'y avait dans l'assistance que des participants d'origine technique >et les sujets purement techniques étaient traités par des experts forcément techniques. Ce n'est plus le cas aujourd'hui.

Disons plutôt que la sécurité a intégré le champ plus vaste de l'Intelligence Economique et c'est ce champ plus vaste qui n'est pas uniquement technique. L'IE intègre les activités de veille et toutes les activités qui participent à la compréhension de l'environnement de l'entreprise. Certaines veilles sont technologiques et d'autres non (juridique, financière, sociétale, ...). L'IE intègre les activités d'action sur l'environnement de l'entreprise (lobbying, influence, ...). Certaines actions nécessitent des compétences techniques par exemple lorsqu'il faut bâtir un argumentaire visant à faire évoluer la réglementation. L'IE intègre enfin les activités de protection des intérêts de l'entreprise (sécurité informatique, sécurité des biens et des personnes, défense devant les tribunaux, contre-influence, ...). La sécurité informatique reste de la technique tandis que la défense devant les tribunaux reste du juridique et la contre influence reste de la communication.

Même s'il est possible de passer d'un champ de l'IE à un autre au sein d'une carrière (et c'est une bonne chose), n'invertissons pas l'organisation des entreprises : l'IE n'est pas un sous-ensemble de la sécurité. Si on parle d'IE dans des conférences sécurité, c'est par opportunité. Il est plus vendeur aujourd'hui d'organiser des conférences sécurité que des conférences sur l'IE (qui, en France, garde malgré tout son côté sulfureux).

Balayons aussi le fait que s'engager dans des études scientifiques, voire dans des prépas et des écoles d'ingénieurs est un lourd tribut à payer à sa jeunesse pour prétendre accéder à des fonctions intéressantes. Les sciences exactes, dont les maths, en général, ce n'est que quand on les voit de loin qu'elles peuvent sembler rébarbatives. C'est au contraire un domaine merveilleux à explorer, où le génie humain montre qu'il n'a pas de limites. Cela implique un travail personnel conséquent ? Oui, mais on n'a rien pour rien et quand vous êtes admis dans une filière qui fera de vous un collaborateur apprécié et plein d'avenir, quel plaisir d'avoir acquis les connaissances de base pour faciliter votre carrière !

La caste des matheux n'est pas encore morte !!!" En clair, il faut absolument souffrir pour accéder au Paradis de la sécurité ?

C'est pour cette raison que des étudiants se détournent des écoles d'ingénieurs et des études scientifiques. Pour ma part, j'ai fait des études scientifiques avec pas mal de maths. Mon métier consiste à concevoir des réseaux télécoms sécurisés pour entreprises et opérateurs télécom (je ne fabrique pas des équipements dans lesquels je dois coder des algorithmes de chiffrement). Honnêtement, depuis 25 ans que j'exerce ce métier, je ne pense pas que le niveau de math dont j'ai eu besoin ait dépassé le programme de collège.

Les métiers de la sécurité manquent d'effectifs. Arrêtons de filtrer inutilement sur les maths. N'y a-t-il que Xavier Niel pour comprendre cela ? cf. son discours lors de la création de l'école 42 visant à former des développeurs qui se sont révélés non compatibles avec le système scolaire français.

D'autant plus que le bon sens ne s'apprend pas en math sup et en math spé. Combien d'ingénieurs proposent des solutions techniques infaillibles mais impossible à déployer et exploiter opérationnellement si tant est qu'elles soient rentables économiquement ? Quelques exemples comme le D2MAC Packet ou l'ATM devraient nous interpellier. A côté de cela, un Martin Bouygues gère son Groupe depuis 30 ans avec juste son bac en poche (mais c'est peut-être un bac C...).

Plutôt que de harceler les filles bonnes en maths pour qu'elles intègrent la filière sécurité, on ferait mieux d'ouvrir les recrutements. Pas besoin de savoir résoudre des équations de Maxwell pour filtrer un protocole dans un firewall ou pour identifier les biens support dans une EBIOS.

La bonne place pour un responsable sécurité du numérique est d'être très proche de la Direction générale

Oui, c'est vrai. Mais il ne faut pas survendre la perspective. Le CSO deviendra Directeur et intégrera le COMEX au terme d'une longue carrière c'est à dire 30 ans après ses copains de lycée qui, eux, auront fait HEC ou ESSEC. Et je doute que son bonus sera du même pourcentage que celui du Directeur Commercial. Il faut être réaliste. Cela a été rappelé par un patron de COMEX lors d'une conférence d'Orange Cyber Défense fin 2017. Les COMEX n'aiment pas les CSO car ceux-ci n'ont pas une approche business de leur métier. Je le vois régulièrement. Ils sont encore dans une logique de "ce qu'il ne faut pas faire du point de vue utilisateur" ou de "ce qu'il faut dépenser avant que l'utilisateur puisse faire son métier". Ils ne sont pas encore dans la logique "je sécurise le métier de l'utilisateur en me débrouillant pour ne pas être sur le chemin critique de son projet".

Je pense que cela viendra avec le temps. Aujourd'hui, ce pas de la logique "client interne" commence à être franchi par les DSI. J'en vois de plus en plus qui se positionnent comme équipier projet d'un business à créer et non comme ceux qui annoncent que le projet business ne pourra se faire que lorsqu'ils auront déployé en 2023 la release 35 de l'usine à gaz informatique et pas avant parce que le budget n'a pas été intégré dans le Plan Opérationnel 2016. Un jour viendra où les CSO aussi comprendront cela.

D'ici là, l'ingénieur sécurité va commencer par surveiller des réseaux (jusqu'à ce que l'IA lui pique son job) parce que c'est en commençant par faire du terrain que l'on acquiert le bon sens client qui lui manque cruellement à sa sortie d'école.

L'IMPASSE DES NOMS DE DOMAINE

Louis Pouzin, EUROLINC

L'année 2012, avec ses nouvelles extensions, aussi bien nommées *vani-TLD*, a remis à la mode le sujet des noms de domaine, qui ronronnait quelque peu. Au début du siècle on nous racontait les merveilles du DNS qui prenait des noms faciles à retenir, au lieu de numéros anonymes. C'était un peu agaçant de mettre *maconnerie* pour maçonnerie, et puis y avait-il un tiret ou un sous-tiret, ou rien, *com* ou *fr*. Heureusement Google est arrivé, avec sa mémoire d'éléphant, plus besoin de tâtonner, un clic et c'est parti. Qui s'amuse encore à taper des noms de domaine ?

On croyait aussi pouvoir se donner de jolis noms, coquins ou caressants, sérieux ou surprenants. Pas de chance, déjà pris. Sinon, vite les prendre en *com*, *net*, *org*, *biz*, *fr*, *eu*, à titre défensif. Gare aux pirates qui vont déposer avant vous des noms que vous avez imprudemment lâchés en réunion.

L'ICANN s'est embarquée dans un processus inexploré donc inégalement prévisible. Ce serait une erreur de braquer les projecteurs uniquement sur les nouvelles extensions, car les mœurs actuelles font la part belle aux anomalies.

POURQUOI L'ICANN

Pourquoi l'ICANN ? Parce qu'elle a été imposée par le gouvernement US. Les arguments sensés justifier cet organisme sont la coordination d'un certain nombre de paramètres techniques, la litanie sécurité–stabilité–résilience du DNS, et la promotion de la compétition.

- La coordination de paramètres techniques est un classique dans de nombreuses professions, notamment dans les télécommunications. L'UIT est un organisme des Nations Unies, ses dirigeants sont élus par ses membres, qui en votent le programme d'activités. L'ICANN est un organisme californien de droit privé, et n'a pas de membres. Ses dirigeants sont cooptés par ceux en place. 80% des standards internet (RFC) sont produits par des industriels US.
- La sécurité du DNS est restée fallacieuse (faille Kaminsky) jusqu'au déploiement de DNSsec, coordonné par l'ICANN, qui par ailleurs n'a qu'un rôle passif. Les opérateurs de racines sont des institutions indépendantes qui gèrent leurs serveurs en bons professionnels, sans avoir besoin de l'ICANN. D'où une bonne stabilité. La résilience est très bonne du fait de la redondance surabondante des copies de racine. Le maillon faible est la liaison avec Verisign, fournisseur unique des mises à jour, par contrat avec le Département du Commerce, c.a.d. le gouvernement US.
- La compétition sur le marché des noms de domaine existe au niveau des registres d'un même gTLD. Chaque gTLD est un marché captif contrôlé par un registre. Avec *.com* et *.net* Verisign contrôle plus de 80% de la clientèle des gTLD. L'ICANN est un monopole mondial. Le contrat avec l'ICANN permet à Verisign d'augmenter systématiquement les tarifs de location des noms (au profit de l'ICANN), alors que l'expansion du marché devrait justifier le contraire. Il y a là une situation permanente de conflit d'intérêt et d'abus de position dominants institutionnalisés. La Commission

Européenne, qui s'attaque à Google, Intel, ou Microsoft pour des raisons similaires, ne voit rien d'anormal dans le racket de l'ICANN.

DOGMES ET MYTHES

Puisqu'il n'y a pas de réaction musclée contre le monopole ICANN, pourquoi ne continuerait-elle pas à en profiter ? Ce n'est bien évidemment pas le langage convenable pour les utilisateurs. On met donc l'accent sur la nécessité incontestable d'une racine unique seule à même de garantir une conversion sûre d'un nom en adresse IP. Aucune preuve technique n'est apportée, c'est même le contraire qui est prouvé. Les 1550 et plus réseaux de mobiles enregistrés par l'UIT ont chacun leur annuaire pour convertir un numéro d'appel unique en identifiant physique localisant un abonné où qu'il soit. L'internet ne gère pas la mobilité et la population d'internautes n'est que la moitié de celle des abonnés mobiles. L'architecture du DNS est largement obsolète de nos jours, mais le dogme fonctionne toujours.

Non seulement la racine de l'internet doit être unique, mais aussi un nom de vani-TLD. Cette contrainte révèle l'amalgame entre le souci de maintenir le monopole ICANN et de créer une nouvelle écurie de marques déposées en concurrence avec l'OMPI (Organisation Mondiale de la Propriété Industrielle). Toutefois les déposants de marques OMPI disposent d'options géographiques et de classes d'usages dont est dépourvue l'ICANN.

Dans la mesure où les vani-TLD seraient surtout prisés par les grandes marques le système pourrait tenir la route sur une clientèle de niche pour qui quelques M\$ sont une miette d'un budget publicitaire confortable. Combien de clients ? Mettons 10 à 20000 sur 10 ans. C'est évidemment un pactole pharaonique propre à faire fantasmer les dirigeants de l'ICANN et leurs amis. À supposer qu'ils soient lucides, quels noms de domaines comptent ils offrir au tiers-état ?

Diverses prévisions donnent des chiffres de 9 à 10 milliards pour la population humaine vers 2030. Donc autant de noms de domaines pour simplifier. En parallèle il faut aussi prévoir des noms de sociétés, clubs, etc. Mettons, sans base scientifique, autant que d'individus, soit au total 20 milliards. Combien faut-il de caractères pour que chacun ait un nom unique ?

Avec un alphabet de 26 glyphes (ascii) il faudrait 8 caractères, mais les noms ne seraient que des chaînes aléatoires. Pour qu'elles aient un sens souhaité par l'utilisateur il faudrait compter 70 à 80 caractères.

Conclusion: à l'avenir les noms de domaine deviendront des numéros de 12 à 15 chiffres, permettant une certaine structuration. Ils n'auront pas plus de valeur qu'un identifiant de carte crédit. On écrira des livres sur la formidable arnaque inventée par l'ICANN.

Toutefois des annuaires localisés et conçus pour une clientèle spécifique verront le jour ici ou là.

1° PARTIE : ASPECTS INFORMATION ET SYSTEMES D'INFORMATION



MYTHES ET LEGENDES DU SYSTEME D'INFORMATION (SI)

Laurent Bloch, membre de l'Institut de l'Economie, précédemment DSI de l'université Paris-Dauphine et responsable de l'informatique scientifique de l'Institut Pasteur

Le *système d'information* (SI), qu'il ne faut pas confondre avec les *systèmes d'exploitation* (Linux, Windows, Android, macOS...), est un territoire partagé entre les informaticiens, chargés de le faire fonctionner, et les gens des « métiers » (finance, budget, comptabilité, ressources humaines, etc.), qui l'utilisent dans leur travail. Toute la difficulté de la réalisation d'un SI d'entreprise réside dans l'effort que ces acteurs aux métiers et aux cultures différentes doivent faire pour comprendre leurs objectifs et leurs contraintes respectifs. (*Achevé de rédiger le 16 janvier 2018*).

MYTHE N° 1 :

LE SYSTEME D'INFORMATION C'EST FINI

Le Système d'information et ses agents (pour ne pas dire ses soutiers) sont rarement populaires dans les entreprises et autres organisations, tout le monde aimerait bien s'en débarrasser, et souvent il ne leur est possible d'échapper à ce sort cruel qu'en veillant à conserver un certain pouvoir de nuisance.

Jusqu'à tout récemment, l'arme fatale pour la destruction du SI et surtout de ses séides était l'externalisation, mais les dernières années ont enrichi l'arsenal : Software as a service, Cloud, micro-services, sans oublier l'« informatique des utilisateurs » et les « logiciels de gestion intégrée », déjà anciens. On pourra bien sûr mobiliser tout ce vocabulaire pour rédiger les lettres de licenciement des informaticiens, mais cela ne suffira pas à produire les documents budgétaires et comptables, ni les bases de données RH, clients, fournisseurs, ni l'annuaire du personnel, objets certes désuets et fastidieux, mais dont il est à craindre qu'ils ne soient encore indispensables pour les années à venir.

Il m'a été donné l'honneur d'être le premier directeur du système d'information (DSI) d'une prestigieuse université parisienne. J'avais postulé à ce poste en partie parce que mes expériences antérieures m'avaient inspiré une certaine aversion pour le SI, et que je voulais en faire l'expérience. Je dus venir à résipiscence : sans SI, impossible d'administrer les dossiers et les parcours d'une dizaine de milliers d'étudiants et d'un millier d'enseignants de statuts divers et variés, ni de répondre aux exigences de plus en plus tatillonnes au fil des ans de la bureaucratie ministérielle, non plus que de maintenir à jour un annuaire de tous ces gens, ce qui permet de leur faire parvenir du courrier électronique, de leur délivrer des diplômes, de leur ouvrir un accès à l'espace numérique de travail qui contient les textes de leurs cours et de leurs travaux pratiques, ainsi qu'au système d'examens en ligne, en évitant toutefois que les étudiants n'aient un accès trop facile à la base de données de scolarité, qui contient les résultats d'examens.

Comme un de mes fidèles lieutenants lors de cette expédition dans la jungle me l'a fait remarquer, la propriété première et primordiale d'un SI, c'est l'unicité, et elle n'est pas facile à obtenir. Si l'on n'y prend garde, de petits SI dissidents naissent à tout moment et dans les endroits les plus bizarres, et si, par défaut de vigilance, on ne les élimine pas immédiatement, une fois qu'ils ont poussé quelques tentacules ici ou là, on aura le plus grand mal à s'en

débarrasser. Et ils seront des vecteurs d'incohérences irréductibles. Bien sûr, il y a des endroits où la cohérence n'est pas utile, et où il ne faut donc pas l'imposer, les activités de recherche par exemple, ou plutôt, les chercheurs ont leurs propres impératifs de cohérence, sans lien avec le SI, sauf pour les contrats et le financement. Mais pour la rémunération des personnels et le règlement des factures, la cohérence est de mise, parce qu'incohérence signifie malversation, détournement.

MYTHE N° 2 :

LES APPLICATIONS DE GESTION C'EST TRES SIMPLE

Combien de fois n'ai-je pas entendu cette assertion de la bouche de brillants chercheurs, éventuellement issus des meilleures écoles de la République : « Un logiciel de comptabilité, je vous écris ça en une après-midi » (sic, et de quelqu'un de vraiment très bien par ailleurs, qui passait huit heures par jour à écrire des programmes statistiques assez subtils).

Les praticiens des sciences à taux élevé de mathématiques, ou les ingénieurs qui appliquent ces sciences à des questions pratiques, ont souvent tendance à penser que leurs problèmes sont bien plus difficiles que ceux des comptables et autres gestionnaires, ils ont sans doute raison, mais de cela ne découle pas que l'informatisation de leurs problèmes soit plus difficile. Michel Volle m'a fait remarquer un jour que ce qui était difficile dans un problème de physique, par exemple, c'étaient les aspects physique et mathématique, mais justement la programmation informatique de la solution pouvait tirer profit du formalisme mathématique rigoureux élaboré à cet effet.

L'ingénieur qui conçoit un objet matériel dispose d'un ensemble d'outils conceptuels bien établis, notamment les lois de la physique et de la chimie, sur lesquels il peut s'appuyer pour élaborer son projet, l'expliquer à ses collègues et à son employeur, éventuellement par des dessins ou des maquettes, etc.

L'architecte de système d'information ne dispose de rien de tel, la matière dont il doit tirer la substance de son projet est mouvante, floue, variable. La complexité du processus d'informatisation est sous-estimée par les responsables à la formation insuffisante. Une cause récurrente d'échec de l'informatisation réside dans la difficulté éprouvée par beaucoup de dirigeants à se faire une idée raisonnable de ce qu'ils peuvent attendre de l'informatique, de ce que cela va coûter, et surtout de l'implication personnelle exigible de chacun pour que cela réussisse.

La difficulté principale pour construire un SI de gestion, c'est de savoir ce que l'on veut faire (question en général assez facile pour un logiciel de calcul pour la physique ou les statistiques). Cette difficulté se décompose en sous-problèmes : trouver un interlocuteur ; faire le tri entre plusieurs interlocuteurs divergents ; s'adapter à la versatilité des interlocuteurs ; s'adapter à la versatilité du législateur, dont les décisions doivent être prises en considération même si elles remettent en cause des choix antérieurs. Ensuite on peut s'attaquer aux problèmes techniques, qui ne sont pas forcément simples, mais dont la solution ne mènera sans doute pas au Turing Award.

MYTHE N° 3 :

LES SYSTEMES INTEGRES C'EST FINI, VIVENT LES MICRO-SERVICES !

Les SI de gestion, comme les autres applications informatiques, suivent les modes : bases de données, SQL (Structured Query Language), progiciels, ERP (Enterprise Resource Planning, ou Progiciel de Gestion Intégré), UML (Unified Modeling Language), méthodes agiles, et

maintenant micro-services, sans oublier que tout va se faire avec des smartphones. Toutes ces techniques peuvent mener à un bon résultat, ou à un mauvais résultat, ou à l'absence de résultat, cela ne dépend pas de la technique choisie, ni d'ailleurs de la quantité d'argent consacrée au projet, mais bien plutôt de la compétence et de l'implication des acteurs.

J'ai eu l'expérience d'un projet de SI ressources humaines (RH : cette locution m'écorche la bouche autant que le clavier mais il faut s'exprimer ainsi pour être compris) où le DRH n'était pas persuadé de l'utilité d'une base de données unique et cohérente, puisque, disait-il, chacun de ses agents avait dans son PC les feuilles Excel qui allaient bien, et que c'était bien suffisant. Inutile de dire que les feuilles Excel étaient incohérentes, peu à jour, souvent fausses mais fausses différemment d'un poste à l'autre, et que le projet SIRH était ainsi voué à un échec inéluctable. Sans SIRH, pas d'annuaire du personnel, donc pas de système d'identification et d'authentification unique, pas de gestion des habilitations, bref pas de SI du tout. Il a fallu que je convainque le DG de donner l'ordre au DRH d'agir dans le sens requis, ce qui a été possible grâce à un adjoint du DRH compréhensif et diligent (mais je crains que ce résultat n'ait été que temporaire et que l'on ne soit revenu à Excel, c'est-à-dire à la négation de tout SI).

À quel point est-ce grave de ne pas avoir de SI ? Il faut un SI sobre, pour ne pas dire minimum, mais pour certaines fonctions c'est indispensable. Toujours dans la prestigieuse université évoquée ci-dessus, j'ai eu à piloter la mise en application de la loi « Libertés et responsabilités des universités » (LRU, dite « loi Valérie Pécresse ») et des « responsabilités et compétences élargies » (RCE) qui en découlent. Cette application est essentiellement une question de SI. Parmi les RCE en question figure la rémunération des personnels, qui auparavant relevait pour l'essentiel de l'administration centrale du ministère. Il allait donc falloir que l'université sache quels étaient ses salariés, et combien elle les payait : aussi bizarre que cela puisse paraître, personne n'en avait la moindre idée. La recherche de cette information m'a permis de découvrir plusieurs canaux de rémunération, dont tout avait été fait pour qu'il soit impossible de les réunir, afin que plusieurs rémunérations versées à la même personne, sur des comptes bancaires différents (à sa demande), ne puissent pas être rapprochées afin de connaître sa rémunération totale. Le rétablissement d'une vue d'ensemble a fait accuser la DSI de « flicage », puisqu'aussi bien la connaissance par l'employeur de la somme qu'il verse à un de ses salariés serait un empiètement inadmissible sur les libertés individuelles. Mais ce travail de mise en cohérence exigea plusieurs mois d'ingénieur.

MYTHE N° 4 :

IL SUFFIT D'ACHETER LE BON LOGICIEL

Toujours à l'université, la DSI fit l'emplette d'un logiciel de gestion financière et comptable basé sur un célèbre ERP du commerce ; nous disposions déjà d'un logiciel de gestion de la scolarité. Puisqu'ERP signifie en français « Progiciel de gestion intégré » (PGI), beaucoup croyaient (dont certains enseignants de Système d'information) que la gestion de l'université serait « intégrée » ipso facto.

En fait, une fois les logiciels installés, le plus dur reste à faire, parce que ces logiciels sont bien incapables d'inventer par eux-mêmes la liste des cours et des diplômes délivrés par l'université, ni les comptes fournisseurs, ni encore moins les règles de gestion qui s'appliquent à chacun de ces objets.

Les règles de gestion : dans un autre établissement, j'ai eu à piloter le passage à l'Euro. Lorsque l'on convertit des comptes de francs en euros, il y a un problème d'arrondi dans les calculs, cela s'appelle chez les comptables le problème des « rompus ». Trois solutions sont possibles : convertir les lignes détails et rétablir le bon total à la fin, ne convertir que le total, ouvrir pour

chaque ligne détail une ligne supplémentaire pour y inscrire la différence. Les comptables responsables, dont le chef était la personne la mieux payée de toute l'entreprise, n'ont jamais su répondre à la question du choix parmi ces trois possibilités. Finalement, entre informaticiens, nous avons choisi la solution 2, la plus paresseuse, mais ce genre de situation est ce qu'il y a de plus difficile pour un architecte de SI : trouver un interlocuteur qui comprend de quoi il s'agit, alors que c'est censé être son travail. Autrement, les logiciels, on arrive toujours à les faire marcher, à condition de ne pas avoir été submergé par les demandes d'adaptations spécifiques idiotes issues de la consultation des utilisateurs.

MYTHE N° 5 :

IL N'Y A QU'À EXTERNALISER LES DÉVELOPPEMENTS !

De nos jours la construction d'un SI de gestion ne demande généralement pas le développement d'un logiciel spécifique à partir d'une feuille blanche, on part généralement de systèmes de gestion de bases de données (SGBD) et d'ERP existants, mais ce n'est pas pour cela qu'il ne reste pas beaucoup de travail à faire : paramétrage, définition des règles de gestion, description des entités de l'entreprise dans son vocabulaire propre (comme nous l'a appris Michel Volle), création d'interfaces avec les autres systèmes de l'entreprise. Tous ces travaux sont des développements, généralement lourds et complexes. Pourquoi ne pas s'en débarrasser en sous-traitant le tout à des prestataires extérieurs ? Beaucoup d'entreprises, surtout françaises, cèdent à cette tentation. Voici ci-dessous quelques raisons d'être prudent.

Pour que cela se passe bien, il faut avoir au sein de l'entreprise de vraies compétences informatiques, c'est-à-dire des ingénieurs capables de comprendre et de critiquer tous les aspects du développement même s'ils ne le font pas eux-mêmes. Ce qui signifie bien entendu que leur activité n'est pas limitée à la rédaction de cahiers des charges et à l'administration de contrats, mais qu'il leur arrive de développer réellement, et pas tous les dix ans. Incidemment, et contrairement à la doxa répandue par les SSII et la presse spécialisée, il n'y a pas de honte à entreprendre des développements avec ses propres personnels.

Un bon équilibre doit être trouvé entre les compétences internes et externes. Si les effectifs de développeurs internes descendent en dessous d'un seuil de 40% sur l'ensemble des projets, la situation devient dangereuse. Il s'agit bien sûr ici de vraies compétences en informatique, pas en chefferie de projet ou en comptabilité. Ce ratio de 40% doit être adapté aux circonstances : en situation de croissance d'activité il convient d'augmenter la part des effectifs internes pour ne pas être pris de court devant des besoins nouveaux, par exemple jusqu'à 60%, alors que si l'on prévoit une baisse d'activité il peut être judicieux d'accroître le recours à l'externalisation pour éviter d'avoir plus tard à organiser des plans sociaux. Mais l'idée est de maintenir une situation où l'entreprise garde le contrôle de son SI.

Il faut mettre en pratique au moins une idée : la spécification du logiciel évoluera et s'affinera au fur et à mesure du développement, il est inutile, voire nuisible, de fixer trop de choses au départ. La rédaction de l'appel à concurrence et le choix d'un éventuel prestataire dépendent fortement de ce principe.

Le client réel, celui qui sera le bénéficiaire du système à réaliser, doit s'impliquer fortement et concrètement dans le projet. S'il ne le fait pas, il doit renoncer à son droit de parole au profit d'une maîtrise d'ouvrage déléguée qui, elle, devra prendre part activement au développement. Si ni l'une ni l'autre de ces conditions ne peut être remplie, le projet sera un échec, inéluctablement.

MYTHE N° 6 :

CONFIER LE SI AUX METIERS !

Pourquoi y aurait-il une DSI (ou DN pour être à la mode) qui régnerait sur le SI ? Les gens des métiers (finances, RH, communication, etc.) savent mieux que personne ce qu'il leur faut, et comme de toute façon tout le travail sera confié à des prestataires extérieurs, inutile de s'encombrer de ces informaticiens caractériels.

J'ai eu l'occasion d'observer la mise en œuvre de cette politique chez un de mes employeurs. La direction financière et la DRH avaient obtenu d'avoir leur propre équipe SI, qui traitait directement avec ses prestataires. Bien sûr, les équipes SI avaient été recrutées sur un critère simple : ne rien connaître à l'informatique, parce que c'est toujours très pénible de discuter avec des gens compétents, ils soulèvent toujours des tas de problèmes ennuyeux auxquels on ne comprend rien, tandis qu'entre incompetents on s'entend toujours bien mieux.

Les prestataires ne tardent pas à repérer, parmi leurs clients, ceux qui sont incompetents : c'est avec eux qu'ils vont faire leur marge, parce qu'avec les compétents ils couvrent à peine les frais. Alors avec ceux qui ne voient rien ils chargent la barque au maximum : effectifs pléthoriques de soi-disant ingénieurs qu'ils envoient se former aux frais du client et dans ses locaux, par exemple.

Dans ce cas précis nous avons pu estimer le surcoût de projets confiés aux « métiers », par comparaison avec des projets comparables gérés par la DSI : le facteur multiplicatif était 4 (quatre). Mais bon, quand on aime on ne compte pas.

MYTHE N° 7 :

L'ANALYSE DE L'EXISTANT

Comment le maître d'ouvrage, angoissé à juste titre par l'idée d'avoir à rédiger des spécifications pour un système destiné à être mis en place plusieurs années plus tard, va-t-il essayer de rédiger un cahier des charges à l'épreuve de tous les événements prévisibles ou imprévisibles ? C'est là qu'il tombe le plus souvent dans l'ornière que lui conseillent d'ailleurs tous les bons auteurs : l'« analyse de l'existant ». Cette activité en apparence innocente rassure les chefs de projet angoissés devant les incertitudes de l'avenir et les subtilités de technologies qu'ils ne connaissent souvent pas suffisamment. Or l'existant c'est l'organisation issue du passé que l'on veut justement remplacer. Il serait idiot de l'ignorer totalement, mais lui consacrer trop de temps peut avoir une influence néfaste. Pire : la rédaction des spécifications est souvent confiée à des personnes issues du terrain, ancrées dans les pratiques du passé sur lesquelles elles n'ont pas un regard suffisamment critique. Dans ce cas, et il est courant, le cahier des charges va consister en l'énumération de ces procédures, qu'il faudrait justement réformer. La démarche du projet, loin de déboucher sur une rénovation du système d'information, va durcir et ossifier des pratiques qui n'étaient jusque-là que de mauvaises habitudes et dont l'énoncé va désormais être gravé dans le marbre du cahier des charges. Le projet aura débouché sur un renforcement de la bureaucratie et des dispositifs paralysants.

MYTHE N° 8 :

L'EXPRESSION DES BESOINS

Une autre mauvaise pratique prend généralement son essor à peu près au même moment dans la chronologie du projet : le « recensement (ou l'expression) des besoins ». Le lecteur pourrait croire à un paradoxe : il est évident que l'on ne construit pas un système d'information juste pour le plaisir, et qu'il est préférable d'avoir une idée de l'objectif poursuivi. Mais trop souvent

la pratique qui se donne cours sous ce prétexte ne sert en rien à éclairer la cible visée, mais plutôt à recueillir les opinions et desiderata hétéroclites de personnes plus ou moins reliées à l'activité concernée, et à embrouiller la marche à suivre dans un galimatias de motivations privées ou collectives plus ou moins incohérentes. De toute façon, lorsque viendra le jour de la recette, ces motivations auront complètement changé et les personnes consultées (si elles n'ont pas de toute façon changé d'affectation dans l'entreprise) ne seront pas satisfaites. Bien sûr ce tableau peut varier selon le contexte : dans une entreprise de taille moyenne dirigée par son propriétaire, la personne habilitée à énoncer le besoin est bien identifiée et, surtout, il y a de bonnes chances pour qu'elle agisse de façon conséquente par rapport aux décisions qu'elle aura prises, puisque c'est son argent qui sera engagé — et s'il n'en est pas ainsi, la conséquence sera en quelque sorte morale, parce que celui qui aura commis l'erreur en sera la victime. À l'autre bout du spectre, dans une administration publique en proie à la concertation compulsive, savoir ce que l'on veut vraiment faire est à peu près impossible, et il faut reconnaître que certaines commissions de fonctionnaires risquent d'être prêtes à ne pas reculer devant des décisions incohérentes, d'autant plus que leurs membres n'encourent aucun risque personnel pour les inconvénients qui pourraient en découler. Il est d'ailleurs reconnu depuis longtemps par les spécialistes des organisations que la prise de décision collégiale, loin d'engendrer la modération des décisions, suscite la montée aux extrêmes par un emballement rhétorique que ne vient contrebalancer aucune angoisse devant des responsabilités futures qui seront de toute façon assumées par d'autres ; les deux attitudes en jeu dans ce phénomène sont l'illusion d'unanimité ou au contraire la polarisation conflictuelle, on en trouvera une analyse bien documentée sous la plume de Christian Morel. Et pour des exemples récents, qu'il suffise de prononcer les mots, chers au contribuable, Louvois ou SIRHEN...

La démarche d'expression des besoins peut être fautive de façon encore plus patente : considérons un établissement de recherche de taille moyenne, composé d'une direction générale qui regroupe quelques centaines de personnes et des unités de recherche qui regroupent quelques milliers de chercheurs. Les départements de la direction générale, pour accomplir leur mission d'administration de la recherche, désirent, et c'est peut-être légitime, se doter d'un système d'information pour savoir ce que font tous ces chercheurs. Si on les interroge pour qu'ils expriment leurs besoins, nul doute que l'on obtienne une longue liste d'informations, sûrement pertinentes d'ailleurs. Mais qui va devoir fournir ces informations ? Les chercheurs, déjà submergés de tâches administratives dont l'utilité leur semble obscure, et qui risquent de mal accueillir ces nouveaux formulaires à remplir à propos de tout et du reste. L'expression des besoins sera ici franchement nuisible.

MYTHE N° 9 :

L'EXHAUSTIVITE DE L'ANALYSE

Les deux erreurs ci-dessus, conjuguées entre elles, peuvent être utilement complétées par une attitude intellectuelle qui semble de prime abord assez saine, mais qui va se révéler comme une nouvelle erreur : le souci d'exhaustivité. Prenons comme exemple un projet de sécurité informatique destiné à évaluer et à améliorer la disponibilité d'un système d'information. Le responsable du projet pourra par exemple s'inspirer de la méthode EBIOS élaborée en France par la Direction centrale de la sécurité des systèmes d'information (DCSSI, devenue ANSSI). Une telle méthode le conduira à dresser une liste des risques, à associer chacun de ces risques à des vulnérabilités, et à envisager les contre-mesures qu'il peut élaborer pour s'en prémunir, selon une formule pleine de bon sens et d'utilité :

risque = menace × vulnérabilité × sensibilité / contre-mesure

Cette conceptualisation est intéressante, mais elle peut engendrer la tentation de dresser une liste de risques ou une liste de vulnérabilités que l'on placera dans la colonne de gauche d'un tableau, afin d'en remplir la colonne de droite avec les contre-mesures appropriées.

Pourquoi cette démarche est-elle maladroite ? Parce que les risques et les menaces sont nombreux et souvent inconnus, alors que le répertoire des contre-mesures possibles est beaucoup plus réduit ; souvent, cela peut se résumer à cinq ou six grands thèmes : plan de sauvegarde des données, plan de reprise d'activité (PRA), amélioration du stockage, aménagement d'un site de secours avec duplication des données à distance, administration correcte des serveurs (fermeture des services inutiles, séparation des privilèges, application des correctifs de sécurité, surveillance des journaux), sécurisation du réseau (pare-feu, authentification forte, fermeture des services inutiles), sécurisation physique des locaux. Il est donc plus simple et plus efficace de partir de la table inverse de la précédente : mettre les contre-mesures dans la colonne de gauche, et énumérer dans la colonne de droite les risques éliminés par elles, ce qui évitera de payer un consultant pendant des mois pour élaborer la liste des centaines de risques plus ou moins réels que l'on peut envisager.

Cette démarche prétendue exhaustive qui consiste à examiner de longues listes d'items pour cocher des cases dans la colonne d'en face, c'est la façon de travailler des ordinateurs, rapides, systématiques et mécaniques ; l'intérêt du travail d'un professionnel compétent par rapport à l'ordinateur, c'est que l'homme réfléchit, a des idées et des connaissances, il est capable d'en tirer parti pour éliminer les cas non pertinents et regrouper les cas similaires de façon à leur appliquer le même traitement, bref il est capable d'avoir une vue synthétique des choses. Réduire le travail humain à des procédures de type informatique est tout à fait contre-productif.

MYTHE N° 10 :

LE SCHEMA DIRECTEUR

Les trois erreurs ci-dessus peuvent se combiner avec d'autres pour engendrer une méta-erreur encore plus stérilisante : la rédaction d'un Schéma Directeur du Système d'Information. Dans les administrations publiques, commettre cette erreur est une obligation réglementaire. Que le lecteur ne se méprenne pas : mener une réflexion sur les objectifs à moyen terme qu'il semble raisonnable de fixer au Système d'Information ne peut être qu'une bonne pratique, et le faire régulièrement une saine discipline. Mais ce n'est pas de cette oreille que l'entendent les instances de contrôle de l'administration, pour qui l'édification d'un Système d'Information ne saurait être qu'une charge. D'ailleurs le système de comptabilité publique n'offre aucun moyen de prendre en considération la valeur du capital dont la puissance publique se sera dotée en contrepartie de cette dépense. Dans ces conditions, le Schéma Directeur a surtout le rôle de digue destinée à fixer des limites durables aux projets informatiques, et par là à limiter la dépense. Le contribuable qui lit ces lignes pourra penser que limiter les dépenses des administrations n'est peut-être pas une mauvaise chose. Sans doute, mais dans un domaine où l'innovation technique est non seulement foisonnante, mais encore porteuse de réductions de coût soutenues, figer pour plusieurs années le périmètre des développements et, surtout, le catalogue des technologies mises à contribution pour les entreprendre, c'est se donner la certitude de payer trop cher, avec un surcoût de plus en plus important au fur et à mesure que passent les années. Nous avons là un nouvel exemple d'une pratique qui semble de bon sens, alors que sa mise en œuvre sans prise en compte des caractéristiques réelles des objets auxquels elle s'applique, par défaut de compétence technique, aboutit très souvent (et pas seulement dans l'administration) à un résultat opposé à celui qui était espéré.

Évidemment, pour la bureaucratie tout ira très bien, les choses seront bien en ordre, même si elles ont coûté en fait beaucoup trop cher, mais comment le savoir, puisque l'on n'y connaît rien ?

MYTHES ET LEGENDES DU REGLEMENT GENERAL SUR LA PROTECTION DES DONNEES

Aurélien Banck, Juriste en droit de la protection des données, ex juriste à la CNIL

Le Règlement général sur la protection des données personnelles (RGPD) adopté le 27 avril 2016 a fait couler beaucoup d'encre. A quelques mois de son entrée en application (le 25 Mai 2018), faisons le tour de ses mythes et légendes ¹ (*Achévé de rédiger le 16 janvier 2018*).

MYTHE N° 1 :

LE REGLEMENT NE CONCERNE QUE LES CITOYENS EUROPEENS !

L'une des nouveautés du règlement européen est son effet extraterritorial. En effet, ce texte est non seulement applicable aux entreprises et organismes établis sur le territoire de l'Union Européenne (UE) mais également « à des personnes concernées qui se trouvent sur le territoire de l'Union ». Il est, dès lors, applicable à des organismes non établis dans l'UE si les traitements qu'ils mettent en œuvre portent sur une offre de biens ou de services à des personnes dans l'Union ou sur le suivi de leur comportement. C'est donc l'ensemble des résidents européens qui bénéficient de cette protection et non les seuls citoyens européens. Ainsi, un américain vivant à Paris bénéficiera de cette protection ; inversement, le fait de traiter des données de citoyens français installés aux Etats Unis n'aura pas pour effet d'entraîner l'application du règlement européen.

MYTHE N° 2 :

LE REGLEMENT NE S'APPLIQUE QU'AUX SOCIETES EUROPEENNES !

Comme indiqué précédemment, le Règlement est applicable aux traitements de données personnelles effectués par un responsable du traitement ou un sous-traitant établis sur le territoire de l'UE mais également à certains traitements de données visant à offrir des biens ou des services à des personnes situées en Europe ou à suivre leur comportement. Dès lors, une société américaine ou indienne qui mettrait en œuvre ce type de traitement serait soumise au Règlement, et ce même si elle n'est pas établie en Europe.

MYTHE N° 3 :

LE REGLEMENT INTERDIT LA COLLECTE DE DONNEES PERSONNELLES !

Le RGPD n'interdit pas, de manière générale, de collecter ou de traiter des données à caractère personnel. Il fixe un cadre à respecter qui repose sur une approche par les risques. Ainsi, plus les opérations de traitements sont risquées, c'est à dire susceptibles d'avoir un impact pour les droits et libertés des personnes concernées, plus les mesures techniques et organisationnelles à mettre en place seront exigeantes. En fait, il ne comporte que peu d'interdictions formelles, et généralement assorties d'exceptions visant à autoriser le traitement dans des circonstances exceptionnelles (on peut citer le cas des catégories particulières de données mentionnées à l'article 9).

¹ Une partie de cet article a été publié dans la Revue Banque, n°810.

MYTHE N° 4 :

UN RESPONSABLE DU TRAITEMENT OU UN SOUS-TRAITANT SERA REGULE PAR UNE SEULE AUTORITE DE CONTROLE !

Le règlement européen prévoit la possibilité de désigner une autorité chef de file mais uniquement pour ce qui concerne les traitements transfrontaliers. En outre, l'article 55 précise que « *lorsque le traitement est effectué par des autorités publiques ou des organismes privés agissant sur la base de l'article 6, § 1, point c) (respect d'une obligation légale) ou e) (exécution d'une mission d'intérêt public), l'autorité de contrôle de l'Etat membre concerné est compétente. Dans ce cas, l'article 56 (concernant l'autorité chef de file) n'est pas applicable* ». Les traitements mis en œuvre pour satisfaire une obligation légale étant relativement courant, la simplification promise et devant résulter de la désignation de cette autorité chef de file semble réduite.

Ces dispositions peuvent même aboutir pour une entité qui ne serait pas l'établissement principal² ou qui ne serait pas établie dans le ressort de l'autorité chef de file à relever de la compétence de plusieurs autorités de contrôle : l'une pour ses traitements transfrontaliers par l'intermédiaire de son établissement principal, et l'autre pour les traitements nécessaires au respect d'une obligation légale, qui dépendraient de facto de la compétence de l'autorité de son lieu d'établissement. Cette situation pourrait s'avérer fréquente dans le cadre des groupes de société où certains traitements sont mis en œuvre par la maison mère et déployés dans l'ensemble de ces filiales. Elle pourrait donc à l'usage se révéler un facteur de complexification de la gouvernance des données.

Par ailleurs, les autres autorités de contrôle restent compétentes pour traiter les plaintes qui leurs sont adressées. En cas de dépôt d'une plainte auprès d'une autorité locale, celle-ci devra interroger l'autorité chef de file afin de savoir si elle souhaite régler le litige. Or, si l'autorité chef de file ne souhaite pas traiter le cas, il incombera à l'autorité initialement saisie de le faire.

MYTHE N° 5 :

LE CONSENTEMENT EST NECESSAIRE POUR TOUS LES TRAITEMENTS DE DONNEES !

Comme la loi Informatique et Libertés, le règlement prévoit que chaque traitement de données personnelles repose sur un ou plusieurs fondements dont la liste est énumérée à l'article 6 du RGPD. Le consentement est l'un de ces fondements au même titre que le respect d'une obligation légale, l'exécution de mesures contractuelles ou l'intérêt légitime du responsable du traitement. Dès lors, le consentement n'est pas un fondement exclusif. Il est à manier avec précaution car il doit être libre, spécifique et informé et résulter d'une action non ambiguë de la personne concernée. La condition de liberté peut en particulier susciter des interrogations, le règlement s'y attachant particulièrement (cf. considérant 43). Il ne doit pas être forcé. Ainsi, dans l'hypothèse où des données seraient nécessaires à une prestation de service, leur collecte ne serait pas libre. En effet, en cas de refus, la personne ne pourrait bénéficier du service en question. Dès lors, il semble opportun de privilégier un autre fondement.

² Critère utilisé pour déterminer quelle est l'autorité chef de file.

MYTHE N° 6 :

LA DESIGNATION D'UN DPO (DATA PROTECTION OFFICER) EST OBLIGATOIRE POUR LES ENTREPRISES EMPLOYANT PLUS DE 250 COLLABORATEURS !

Les critères à prendre en compte pour la désignation obligatoire d'un DPO ont été discutés tout au long des débats entourant l'adoption du RGPD, évoluant d'une notion quantitative (le nombre de salariés du responsable du traitement ou de personnes concernées par le traitement) à une notion plus qualitative relative à la sensibilité des traitements mis en oeuvre. Consciente du fait que le nombre de personnes concernées ou de personne mettant en oeuvre le traitement est indifférent aux risques qu'il comporte (une violation de données exposant des milliers de personnes pouvant survenir dans n'importe quel contexte), le législateur européen a finalement retenu un critère qualitatif. Ainsi, doivent désigner un DPO :

- les responsables de traitement et les sous-traitants donc l'activité de base consistent en des opérations de traitement qui du fait de leur nature, portée et/ou finalités impliquent un suivi régulier des personnes concernées,
- ou consistent en un traitement à grandes échelles de données sensibles au sens de l'article 9 (données de santé, etc.) et de données relatives à des condamnations pénales.

MYTHE N° 7 :

LE RGPD SUPPRIME L'OBLIGATION DE NOTIFICATION DES TRAITEMENTS : DONC LA CHARGE ADMINISTRATIVE SERA REDUITE !

A cette date, l'ensemble des traitements de données à caractère personnel doivent faire l'objet d'une déclaration ou d'une autorisation de la Commission nationale de l'Informatique et des Libertés. Le règlement supprime cette notification qu'il remplace par l'obligation de tenir un registre et de réaliser des analyses d'impact pour les traitements de données les plus sensibles (dont la liste figure à l'article 35). En outre, il intègre en droit positif la notion d'Accountability : chaque responsable du traitement doit être en mesure de prouver à tout moment qu'il respecte l'ensemble des obligations du RGPD. Dès lors, la documentation à produire en interne, à maintenir et à actualiser apparaît plus importante que la simple déclaration à effectuer auprès de la CNIL.

A noter, que l'obligation de tenir un registre incombe également au sous-traitant et au représentant désignés par les responsables de traitement établis hors Union européenne. Une exception a cependant été introduite pour les entreprises employant moins de 250 salariés qui sont exonérées de cette obligation sous réserve qu'elles ne mettent pas en oeuvre de traitements de données comportant des risques pour les personnes ou des données sensibles, ou s'il n'est pas occasionnel³.

MYTHE N° 8 :

LA CONFORMITE AU REGLEMENT EST UN PROBLEME IT !

Les questions de sécurité et d'architecture informatique doivent impérativement être abordées dans le cadre du Règlement. Toutefois, limiter la question de la conformité au RGPD à la sphère IT équivaut à ignorer tout un pendant du texte relatif aux mesures organisationnelles. A l'image de l'article 24 concernant l'obligation d'Accountability, le RGPD associe

³ Cette notion de traitement occasionnel pourrait réduire substantiellement la portée de cette exception. Elle n'est cependant pas définie dans le RGPD.

systématiquement « mesures techniques et organisationnelles » visant à s'assurer et à être en mesure de démontrer à tout moment que le traitement est effectué conformément au texte. Les processus et procédures doivent donc également être revues afin de garantir la conformité. L'aspect organisationnel est très important notamment car il touche à la formation et que de nombreuses violations de données résultent simplement d'un problème humain.

MYTHE N° 9 :

TOUTES LES DONNEES PERSONNELLES DOIVENT ETRE CHIFFREES POUR ETRE CONFORME AU REGLEMENT EUROPEEN !

Le Règlement repose sur une approche par les risques, y compris en matière de sécurité. Ainsi, l'article 32 prévoit que « compte tenu de l'état des connaissances, des coûts de mise en œuvre et de la nature, de la portée, du contexte et des finalités du traitement ainsi que des risques » qu'il comporte, le responsable du traitement et le sous-traitant « mettent en œuvre les mesures techniques et organisationnelles appropriées afin de garantir un niveau de sécurité adapté au risque ». Le chiffrement des données fait partie des mesures citées par le Règlement au titre des mesures techniques. Toutefois, il n'est pas exclusif. Des mesures de pseudonymisation, d'anonymisation ou d'autres moyens permettant de garantir la confidentialité, l'intégrité, la disponibilité des données peuvent également être mises en place. Le chiffrement est donc à manier avec précaution en raison de son coût et de son impact potentiel sur la disponibilité des informations.

MYTHE N° 10 :

LES VIOLATIONS DE DONNEES DOIVENT ETRE SYSTEMATIQUEMENT NOTIFIEES A L'AUTORITE DE CONTROLE !

L'article 33 du RGPD prévoit qu'en cas de violation de donnée, le responsable du traitement la notifie à l'autorité de contrôle dans les meilleurs délais « à moins que la violation en question ne soit pas susceptible d'engendrer un risque pour les droits et libertés des personnes physiques ». Dès lors, seules les violations susceptibles d'engendrer un risque pour les personnes doivent être portées à la connaissance de l'autorité de contrôle.

De la même manière, seules les violations susceptibles d'« engendrer **un risque élevé** pour les droits et libertés d'une personne physique » doivent être communiquées à la personne concernée. Ainsi, cette notification ne sera pas nécessaire si :

- « le responsable du traitement a mis en œuvre les mesures de protection techniques et organisationnelles appropriées et ces mesures ont été appliquées aux données à caractère personnel affectées par ladite violation, en particulier les mesures qui rendent les données à caractère personnel incompréhensibles pour toute personne qui n'est pas autorisée à y avoir accès, telles que le chiffrement
- le responsable du traitement a pris des mesures ultérieures qui garantissent que le risque élevé pour les droits et libertés des personnes concernées [...] n'est plus susceptible de se matérialiser ».

Dans l'hypothèse où une communication individuelle exigerait des efforts disproportionnés, une communication publique peut être envisagée.

Par contre, toutes les violations de données doivent être documentées conformément à l'article 33, §5. Cette documentation doit mentionner les faits, les effets et les mesures prises pour y remédier.

MYTHE N° 11 :

TOUT PROFILAGE NECESSITE UN CONSENTEMENT !

L'article 22 du RGPD précise que les personnes concernées ont le droit pas être soumises à une décision automatisée produisant des effets juridiques la concernant ou l'affectant de manière significative. Une disposition similaire figurait déjà l'article 15 de la Directive 95/46/CE même si elle ne visait pas explicitement le profilage⁴. Le RGPD prévoit certaines exceptions notamment lorsque la décision est nécessaire à la conclusion ou à l'exécution d'un contrat, si elle est prévue par une disposition législative (sous réserve que celle-ci prévoit des mesures appropriées pour la sauvegarde des droits des personnes) et si le responsable du traitement a recueilli le consentement de la personne concernée. Dès lors, seules les mesures de profilage **produisant des effets juridiques pour la personne ou l'affectant de manière similaire** et qui ne pourraient bénéficier de ces exceptions doivent faire l'objet d'un recueil de consentement.

Le débat se déplace donc sur la notion d'effet juridique pour la personne concernée. On peut imaginer qu'une mesure de profilage conduisant à l'exclusion d'une personne du bénéfice d'un droit ou d'un contrat pourrait rentrer dans cette catégorie. Cette question est tout à fait prégnante en matière de marketing où l'objet de la segmentation peut se limiter à proposer à un client un produit en fonction de ses goûts et ne conduit pas nécessairement à un mécanisme d'exclusion. Ce type de profilage pourrait ainsi être soumis à un droit d'opposition, conformément aux dispositions de l'article 21.

MYTHE N° 12 :

LE CONSENTEMENT DES PARENTS EST SYSTEMATIQUEMENT REQUIS EN CAS DE COLLECTE DE DONNEES RELATIVES A DES ENFANTS !

Le RGPD prévoit des dispositions spécifiques pour protéger les personnes vulnérables en particulier les enfants, c'est-à-dire les individus âgés de moins de 16 ans. Ainsi, lorsque le consentement d'un mineur est requis « en ce qui concerne l'offre directe de services de la société de l'information aux enfants », le traitement de leurs données ne sera licite que si l'enfant est âgé d'au moins 16 ans. Dans le cas contraire, il faudra recueillir le consentement ou l'autorisation du titulaire de l'autorité parentale.

Cet article n'est donc pas applicable à tout traitement de données concernant des enfants mais uniquement à ceux concernant l'offre directe de services de la société de l'information. Cette notion de « services de la société de l'information » n'est pas nouvelle en droit communautaire et couvre tout service fourni, normalement contre rémunération, à distance au moyen d'équipement électronique de traitement (y compris la compression numérique) et de stockage des données, à la demande individuelle d'un destinataire de services⁵. Concrètement cette situation vise majoritairement le cas des réseaux sociaux.

⁴ A noter que la France lors de la transposition de la Directive 95/46/CE n'a pas intégré cette disposition parmi les droits des personnes concernées comme le prévoyait la Directive mais dans les conditions relatives à la licéité des traitements.

⁵ Cette notion est définie à l'article 1^{er}, §2 de la Directive 98/34/CE du 22 juin 1998 prévoyant une procédure d'information dans le domaine des normes et réglementations techniques et des règles relatives aux services de la société de l'information, modifiée par la directive 98/84/CE 20 novembre 1998 concernant la protection juridique des services à accès conditionnel et des services d'accès conditionnel. La Directive 2000/31/CE relative à certains aspects juridiques des services de la société de l'information, et notamment du commerce électronique, dans le marché intérieur y fait également référence.

Le consentement des parents n'est donc pas requis pour tout traitement de données. Ainsi, le considérant 38 précise que « le consentement du titulaire de la responsabilité parentale ne devrait pas être nécessaire dans le cadre de services de prévention ou de conseil proposés directement à un enfant ».

A noter que les Etats Membres peuvent choisir d'abaisser cette limite d'âge dans le cadre de leur loi nationale, sous réserve que l'âge minimum ne soit pas inférieur à 13 ans.

MYTHE N° 13 :

IL N'EST PAS NECESSAIRE DE CONCLURE UN CONTRAT AVEC UN SOUS-TRAITANT CAR CELUI-CI EST DIRECTEMENT RESPONSABLE DEVANT LA CNIL !

Le Règlement crée une responsabilité autonome du sous-traitant qui peut désormais voir sa responsabilité engagée dans les mêmes conditions qu'un responsable du traitement (procédures et sanctions identiques).

Toutefois, l'article 28 du RGPD précise que « le traitement par un sous-traitant est régi par un contrat ou un autre acte juridique au titre du droit de l'Union ou du droit d'un Etat membre, qui lie le sous-traitant à l'égard du responsable du traitement ». La responsabilité directe du sous-traitant n'a donc pas pour effet de se substituer au contrat devant être conclu entre les parties.

Au contraire, l'article 28 du RGPD vient préciser les différentes mentions que doit comporter le contrat dont « l'objet et la durée du traitement, la nature et la finalité du traitement, le type de données à caractère personnel et les catégories de personnes concernées, et les obligations et les droits du responsable du traitement ». Cette nouvelle responsabilité n'a pas, non plus, pour effet d'exonérer le responsable du traitement de ses obligations vis-à-vis de son sous-traitant, notamment de son obligation de contrôle de la prestation dont le contour se dessine de plus en plus au travers des décisions de la CNL (notamment Hertz et Darty).

MYTHE N° 14 :

LES DONNEES DE MES CLIENTS SONT STOCKEES DANS LE CLOUD, DONC SEUL MON PRESTATAIRE DOIT ETRE CONFORME AU RGPD !

L'organisme qui externalise tout ou partie d'un traitement doit s'assurer que son sous-traitant respecte les dispositions du RGPD et ce avant même d'entrée en relation avec lui. L'article 28 du RGPD précise, en effet, que le responsable du traitement « fait uniquement appel à des sous-traitants qui présentent des garanties suffisantes quant à la mise en oeuvre de mesures techniques et organisationnelles appropriées de manière à ce que le traitement réponde aux exigences du présent règlement ». Le sous-traitant doit donc respecter les dispositions du Règlement.

Toutefois, le responsable du traitement devra également respecter les obligations qui lui incombent y compris en ce qui concerne le traitement externalisé. Il devra l'inscrire dans son registre conformément aux dispositions de l'article 30, documenter et notifier les éventuelles violations de données l'impactant, etc. Il devra, en outre, respecter les dispositions figurant à l'article 28 afin de permettre à son prestataire d'assurer sa mission, en particulier lui communiquer des instructions documentées. Une opération d'externalisation ne dispense donc pas un responsable du traitement de son obligation de respecter le RGPD.

MYTHE N° 15 :

MON PRESTATAIRE EST « CERTIFIÉ CNIL », DONC JE N'AI RIEN À CRAINDRE !

Le Règlement prévoit la mise en place de codes de conduite, de certifications, de labels et de marques en matière de protection des données. L'application d'un code de conduite approuvé conformément aux dispositions de l'article 40 ou d'un mécanisme de certification prévu à l'article 42 « peut servir d'élément pour démontrer le respect de certaines exigences » comme celles relatives à la sécurité des données, aux transferts, à l'Accountability ou au principe du Privacy by design et by default. Il s'agit donc d'un indice de conformité.

L'application de codes de conduite ou de mécanisme de certification sera pris en compte par les autorités de contrôle pour déterminer s'il y a lieu ou non d'imposer une amende administrative et pour déterminer son montant (article 83, §2).

La CNIL a déjà développé quelques labels pour les procédures d'audit et de formation, un label gouvernance et un label pour les coffres forts électroniques. Les prestataires labellisés sont donc peu nombreux et pour des prestations très particulières, qui ne sauraient s'entendre comme une labellisation générale des produits et services délivrés par un prestataire.

Même en cas de développement et de généralisation des codes de conduite et des certifications, le statut de ces sceaux de confiance devra être précisé notamment pour déterminer leur impact en termes de responsabilité (par exemple un responsable du traitement peut-il constater la conformité de son sous-traitant sur la base de sa seule certification ? etc.).

MYTHE N° 16 :

LE DROIT À L'OUBLI EST UN DROIT ABSOLU !

Déjà consacré dans l'affaire Google contre Spain⁶, le RGPD introduit un droit à l'oubli au bénéfice des personnes concernées. Souvent identifié comme l'une des nouveautés majeures du règlement, il convient d'examiner ses conditions d'exercice.

L'article 17 fixe une liste de cas dans lesquels le responsable du traitement a l'obligation d'effacer des données. On constate notamment que ce droit s'applique lorsque les données « ne sont plus nécessaires au regard des finalités pour lesquelles elles ont été collectées » ou si les données « ont fait l'objet d'un traitement illicite ». Or, ces deux hypothèses constituent des manquements à d'autres dispositions du règlement. En effet, l'article 5 précise que les données *sont* « conservées sous une forme permettant l'identification des personnes concernées pendant une durée n'excédant pas celle nécessaire au regard des finalités pour lesquelles elles sont traitées ». Chaque responsable du traitement doit donc fixer et d'appliquer des durées de conservation adaptée. En outre, les traitements illicites sont proscrits car par définition, ils ne respectent pas les dispositions du règlement. Ces deux cas d'application du droit à l'effacement semblent donc plus résulter de la mise en oeuvre de mesures correctives relatives à des traitements qui violent les droits et libertés de la personne concernée.

Le retrait du consentement et l'exercice fructueux d'un droit d'opposition sont deux autres cas dans lesquels la personne concernée peut solliciter l'oubli. Or, en l'espèce, l'effacement est la conséquence logique de l'exercice de ces deux autres droits et non un droit autonome. Un cinquième cas est celui dans lequel les données « doivent être effacées pour respecter une obligation légale ». Or, le législateur a plutôt tendance à fixer une durée de conservation des données (ce

⁶ “Google Spain SL and Google Inc. v Agencia Española de Protección de Datos (AEPD) and Mario Costeja González” (C-131/12).

qui peut engendrer une obligation de suppression à échéance) qu'une obligation d'effacement. En tout état de cause, une telle disposition aurait dû être prise en compte dans la définition de la politique de durées de conservation mentionnée supra. Ce cas se rapproche donc plus de l'application d'une mesure corrective.

Enfin, le droit à l'effacement est ouvert lorsque les données ont été collectées « *dans le cadre de l'offre de services de la société de l'information visée à l'article 8* ». Pour faire simple, il s'agit des réseaux sociaux et des mineurs, c'est sûrement ici que réside l'apport majeur du règlement.

Comme mentionné ci-dessus, le droit à l'oubli est donc loin d'être absolu, surtout à la lecture de la longue série d'exceptions qui doivent également être combinées avec ces conditions d'application et qui incluent l'exercice du droit à la liberté d'expression et d'information.

MYTHE N° 17 :

LE DROIT A LA PORTABILITE CONCERNE TOUS LES TRAITEMENTS DE DONNEES !

L'article 20 précise que la personne concernée a le droit de recevoir communication des données qui la concerne dans un « *format structuré, couramment utilisé et lisible par machine* » lorsque le traitement automatisé est fondé sur le recueil de son consentement ou sur un contrat. Dès lors, l'ensemble des traitements mis en œuvre par un responsable ne rentre pas dans le périmètre des traitements éligibles au droit à la portabilité, notamment ceux mis en œuvre sur la base d'un intérêt légitime ou visant à satisfaire une obligation légale.

La détermination des fondements de chaque traitement appelle donc à l'adoption d'une approche globale afin d'en évaluer les conséquences en termes de droits des personnes.

MYTHE N° 18 :

JE SUIS PROPRIETAIRE DES DONNEES A CARACTERE PERSONNEL DE MES CLIENTS !

Le droit à la protection des données relève des droits fondamentaux prévus par la Charte européenne des droits de l'homme. Il s'agit donc d'un droit de la personnalité, consubstantiel de l'individu à qui les données se rattachent et qui ne peut être cédé. Seules les personnes concernées ont le droit de décider de l'utilisation qui sera faite des données les concernant.

Leur simple collecte par un organisme n'a donc pas pour effet de lui permettre d'en acquérir la propriété. Ces données peuvent toutefois être protégées par le droit *sui generis* du producteur de la base de données sous réserve que celui-ci soit en mesure de démontrer un investissement relatif à l'élaboration et au maintien de cette base.

MYTHE N° 19 :

LE MONTANT DES SANCTIONS ADMINISTRATIVES EST PLAFONNE A 10 OU 20 MILLIONS D'EUROS !

L'article 83 fixe le montant des sanctions administratives prévues en cas de violations du RGPD, celui-ci variant en fonction de la gravité des manquements. Le texte fixe une liste de violation faisant l'objet « *d'amendes administratives pouvant s'élever jusqu'à 10 000 000 euros, ou dans le cas d'une entreprise, jusqu'à 2% du chiffre d'affaires annuel mondial total de l'exercice précédent, le montant le plus élevé étant retenu* ». La rédaction est la même pour la seconde série de manquements identifiés comme plus grave et exposant leur auteur à une amende pouvant aller jusqu'à 20 000 000 d'euros ou 4% du chiffre d'affaire. Le montant des sanctions pour une

entreprise dépend donc directement de son chiffre d'affaire et n'est pas plafonné. Pour Facebook, cela pourrait par exemple représenter 990 millions d'euros⁷ !

Quant à savoir quel sera le chiffre d'affaire à prendre en compte pour déterminer le montant de la sanction notamment en cas de contentieux impliquant plusieurs sous-traitants ou responsables de traitement, il semble de bon sens d'affirmer que chacun sera individuellement responsable à la hauteur de son chiffre d'affaire pour la partie des manquements qui lui serait imputable.

MYTHE N° 20 :

LES DONNEES A CARACTERE PERSONNEL QUE J'AI COLLECTEES AVANT LE 25 MAI 2018 NE SONT PAS CONCERNEES PAR LE REGLEMENT !

L'application dans le temps de ces dispositions suscitent de nombreuses interrogations. Le considérant 171 vient préciser que « les traitements déjà en cours à la date d'application du présent règlement devraient être mis en conformité avec celui-ci dans un délai de deux ans après son entrée en vigueur ». L'article 99 précise que le Règlement entre en vigueur le 21^{ème} jour suivant sa publication au Journal officiel de l'Union européenne et sera applicable à partir du 25 mai 2018. Les données collectées avant le 25 mai 2018 sont donc bien soumises au RGPD.

⁷ 4% d'un chiffre d'affaires de 27,638 milliards \$ en 2016, soit 24,790 milliards d'euros

MYTHES ET LEGENDES DU PRIVACY BY DESIGN

Alessandro Fiorentino, Ambassadeur du Privacy by Design, Consultant au cabinet Infhotep

Certains en parlent comme d'une approche, un concept prometteur, un ensemble de principes ou encore un article du règlement général sur la protection des données. Il représente bien plus.

Le concept de Privacy by Design est un ensemble de principes et de pratiques qui régissent les rapports entre les hommes et l'information.

Il s'agit d'un paradigme, une manière d'appréhender les choses, une vision cohérente du système qui repose sur une base définie, une matrice disciplinaire, un modèle théorique.

C'est une forme de rail de la conception dont les lois sont destinées aux concepteurs pour exprimer la façon dont un système d'informations ab initio doit être conçu et pensé dans ses grandes lignes afin d'assurer la conformité informatique et libertés.

MYTHE N° 1 :

LE PRIVACY BY DESIGN EST UN NOUVEAU CONCEPT NE PAR LE BIAIS DE L'ARTICLE 25 DU REGLEMENT GENERAL SUR LA PROTECTION DES DONNEES !

C'est faux, ce concept fut créé au Canada dans les années quatre-vingt-dix par Ann CAVOUKIAN Ph. D dans le cadre de ses fonctions de Commissaire à l'information et à la protection de la vie privée de l'Ontario.

À cette époque la protection des données à caractère personnel n'était pas encore une notion synonyme d'enjeu économique. Portée par la défense des libertés fondamentales de l'individu, Ann CAVOUKIAN présente ce concept de Privacy by Design comme une démarche qui vise à assurer la protection de la vie privée.

Cette démarche impose d'intégrer à toute technologie exploitant des données à caractère personnel des dispositifs techniques de protection de la vie privée dès sa conception et de s'y conformer tout le long du cycle de vie des données.

La démarche repose sur sept principes fondamentaux, et représente une solution permettant aux technologies d'évoluer sans porter atteinte à la vie privée des individus, une coexistence saine et pérenne du numérique et des individus.

Les sept principes fondamentaux

1. Prendre des mesures proactives et non réactives, des mesures préventives et non correctives.
2. Assurer la protection implicite de la vie privée.
3. Intégrer la protection de la vie privée dans la conception des systèmes et des pratiques.
4. Assurer une fonctionnalité intégrale selon un paradigme à somme positive et non à somme nulle
5. Assurer la sécurité de bout en bout, pendant toute la période de conservation des renseignements.
6. Assurer la visibilité et la transparence.

7. Respecter de la vie privée des utilisateurs.

MYTHE N° 2 :

LE PRIVACY BY DESIGN EST UN CONCEPT CANADIEN QUE L'EUROPE N'A JAMAIS SOUTENU !

C'est faux, c'est une recommandation du parlement européen à l'attention du Conseil des ministres sur le renforcement de la sécurité et des libertés fondamentales sur Internet publiée le 26 mars 2009 qui marque pour la première fois l'engagement et le soutien de l'Europe dans la promotion du concept de Privacy by Design « *la protection des données et de la vie privée devrait être introduite dès que possible dans le cycle de vie des nouveaux développements technologiques, assurant aux citoyens un environnement convivial* ».

En mai 2010 la prise de position de Peter HUSTINX, contrôleur européen de la protection des données, sur la nécessité d'intégrer les principes de Privacy by Design dans les technologies de l'information et de la communication marque à son tour un tournant décisif dans la reconnaissance juridique du paradigme.

Peter HUSTINX soulignait que le cadre juridique actuel imposait déjà l'obligation de mettre en œuvre des technologies renforçant la vie privée via le considérant 46 et l'article 17 de la Directive 95/46/CE du Parlement européen et du Conseil, du 24 octobre 1995, relative à la protection des personnes physiques à l'égard du traitement des données à caractère personnel et à la libre circulation de ces données.

MYTHE N° 3 :

LE PRIVACY BY DESIGN EST UNE DEMARCHE REPOSANT UNIQUEMENT SUR DES DISPOSITIFS TECHNIQUES !

Si nous devons retenir l'essentiel, le concept de Privacy by Design implique de mettre en œuvre des mesures techniques et organisationnelles que nous devons déterminer en amont pour encadrer la protection des données et fournir aux personnes concernées un niveau de protection optimum sans aucune intervention de leur part.

Il ne s'agit donc pas uniquement de mesures techniques, la majorité des principes fondamentaux sur lesquels repose le concept, peuvent être couverts grâce à la mise en œuvre de mesures organisationnelles.

L'objectif est de faire preuve de proactivité et donc d'anticiper les incidents d'atteinte à la vie privée avant qu'ils ne se produisent par le biais de mesures organisationnelles et techniques.

Le Privacy by Design place au premier plan l'importance d'agir en amont.

Le concept de Privacy by Design est une forme d'indicateur temporel, il situe dans le temps la place de la protection des données à caractère personnel.

À mon sens, le plus important est de garder à l'esprit que ce doit être une démarche globale. La mise en œuvre de pseudonymisation ou d'un système de chiffrement sont certes des éléments substantiels mais doivent faire partie d'un tout. La mise en œuvre d'une démarche Privacy by Design doit être transverse au sein de l'organisme afin d'être intégrée aux différents processus métiers.

Sur le plan humain, le premier prérequis indispensable à la mise en œuvre optimale du Privacy by Design est le développement d'une culture « Privacy » au sein de l'organisme concerné.

Si une entreprise ou un organisme souhaite s'inscrire dans une démarche responsable et durable, il est impératif que l'ensemble de son personnel à tous les niveaux hiérarchiques soit mobilisé.

Lorsque l'on parle de protection de données à caractère personnel, chacun d'entre nous a un rôle à jouer, les responsabilités de chacun doivent être clairement délimitées.

Ce qu'il faut retenir c'est qu'il n'y a pas de schéma organisationnel de gouvernance bon ou mauvais. Quelle que soit l'organisation envisagée, elle doit tenir compte et répondre à des spécificités structurelles de l'organisme. Cette gouvernance doit s'inscrire dans l'histoire et la propre culture de l'organisme. La seule constante est le point de départ, pour fonctionner, cette démarche doit être initiée sous l'impulsion de la direction comme un pilier des valeurs de l'organisme.

MYTHE N° 4 :

LES TECHNOLOGIES RENFORÇANT LA VIE PRIVÉE APPELEES EGALEMENT LES : « PETS » (PRIVACY ENHANCING TECHNOLOGIES) SONT UNE EXTENSION DU PRIVACY BY DESIGN !

Dans la pratique ce n'est pas faux, mais dans la théorie c'est n'est pas vrai. À l'origine du concept de Privacy by Design on trouve les « technologies renforçant la protection de la vie privée » (PET, Privacy Enhancing Technologies). Le Privacy by Design est une évolution des PETs.

Au cours du séminaire « *Respect de la vie privée dès la conception* » qui se tenait à Madrid le 2 novembre 2009 le contrôleur européen de la protection des données Peter HUSTINX a précisé que ce terme de PET est apparu pour la première fois en 1995 dans le rapport « *Technologies renforçant la protection de la vie privée : le chemin vers l'anonymat* ». Ce rapport était le fruit d'un projet commun mené par les autorités hollandaises pour la protection des données et le Commissariat à l'information et à la protection de la vie privée de l'Ontario. Ann CAVOUKIAN coté canadien et John BORKING du côté hollandais ont à ce moment joué un rôle-clé. Ils présentaient une nouvelle approche de protection de la vie privée.

Les PETs sont à l'origine du principe de « minimisation des données », c'est sur la base de ce principe que s'est développé ce concept de Privacy by Design.

Dans la pratique c'est quoi ces PETs ?

Les PETs (Privacy Enhancing Technologies) sont des outils clés, ce sont des applications et des mécanismes intégrés aux services, aux plateformes en ligne et aux systèmes d'information qui permettent de protéger les données personnelles et de contrôler l'usage qui en est fait. Ces mécanismes permettent notamment aux utilisateurs de contrôler les données qu'ils partagent, de les minimiser ou de les rendre anonymes.

Dans certains cas, ils peuvent également négocier les termes et conditions de traitement de leurs données par les applications ou services en ligne. Les PETs ont vocation à couvrir un ou plusieurs principes fondamentaux sur lesquels repose le concept de Privacy by Design ou encore répondre à des exigences de protection des données par le biais de dispositifs techniques au sein même du système d'information.

Le respect des durées de conservation est un des exemples les plus pertinents, un PET pour répondre à cette exigence consisterait à cartographier les catégories des données présentes au sein du système d'information ainsi que les durées de conservation inhérentes à chacune d'elle, afin d'automatiser la purge des données le moment venu afin d'assurer le cycle de vie des données jusqu'à leur destruction.

Nous pourrions également envisager un PET qui permettrait de définir le nombre maximal d'enregistrement pouvant faire l'objet d'une extraction pour limiter les potentielles fuites de données.

Le chiffrement des données est également un des dispositifs techniques clés dans la mise en œuvre d'une démarche Privacy by Design. Pour beaucoup de DSI, le chiffrement des données est synonyme de perte de performance, et comme dirait mon ami Gérard PELIKS, cette idée fait partie des « mythes et légendes des TIC ». Il existe des techniques comme le chiffrement phonétique qui permettent de délimiter en amont le périmètre de l'information qui sera déchiffrée pour répondre à une requête afin de conserver les performances attendues grâce à un déchiffrement périmétrique.

In fine, l'objectif est de relier l'ensemble des PETs mis en place pour donner naissance à un IPM (Integrated Privacy Minder) intégré au sein même de la structure du système et qui sera présent dans l'ensemble de l'arborescence ou des couches du SI. Il s'agit d'un ange gardien numérique garant de la notion d'« adequacy » (adéquation de la « privacy »). C'est un PET dont la plupart des autres seront dépendants.

Il a vocation à répondre à l'ensemble des exigences des différents principes de protection des données à caractère personnel ainsi qu'à anticiper toutes les dérives potentielles et les risques d'exploitations abusives des données.

La mise en œuvre de ces technologies renforçant la vie privée représente une réelle valeur ajoutée pour assurer la sécurité juridique de l'organisme et assurer le facteur confiance des usagers ou des clients. Néanmoins conformément à l'article 25, ces mesures devront rester appropriées par rapport aux traitements concernés.

MYTHE N° 5 :

LA MISE EN ŒUVRE DU PRIVACY BY DESIGN SE LIMITE A RESPECTER LE CONSIDÉRANT 78 DU REGLEMENT GENERAL SUR LA PROTECTION DES DONNEES !

Dans la théorie ce n'est pas faux, mais dans la pratique c'est n'est toujours pas vrai. Si nous nous attachons à relever les éléments importants présents dans le considérant 78, le respect de ce considérant consiste à assurer la transparence, le contrôle, la minimisation de la collecte des données, la pseudonymisation et la sécurité. Ces éléments sont redondants avec beaucoup de principes fondamentaux Informatique et Libertés et avec les principes fondamentaux sur lesquels repose le concept de Privacy by Design.

La transparence traduit le sixième principe du concept de Privacy by Design : « Assurer la visibilité et la transparence ».

Le contrôle traduit le principe d'autodétermination informationnel, l'un des objectifs principaux de la protection des données dès la conception et est également une interprétation du septième principe : « Respecter de la vie privée des utilisateurs ».

La minimisation de la collecte des données traduit le second principe fondamental : « Assurer la protection implicite de la vie privée ».

La pseudonymisation et la sécurité traduisent le premier principe : « Prendre des mesures proactives et non réactives, des mesures préventives et non correctives », le troisième principe : « Intégrer la protection de la vie privée dans la conception des systèmes et des pratiques » et le cinquième principe : « Assurer la sécurité de bout en bout, pendant toute la période de conservation des renseignements ».

Néanmoins il faut garder à l'esprit que la pseudonymisation n'est qu'une mesure possible, elle est citée dans le considérant 78 comme un exemple.

MYTHE N° 6 :

LES PRIVACY BY DESIGN ET LE PRIVACY BY DEFAULT SONT DEUX CONCEPTS DIFFERENTS !

Cette phrase est vraie, si l'on se limite à l'interprétation du législateur. En effet, l'article 25 alinéa 1 présente le principe de « Privacy by Design », l'alinéa 2 quant à lui présente le principe de « Privacy by Default ».

Néanmoins si l'on revient aux origines du concept c'est faux.

Le principe de « Privacy by Default » correspond au second principe fondamental et appartient intrinsèquement au concept de « Privacy by Design ». Cette dissociation de ce principe par le biais de l'alinéa 2 a certainement un objectif pédagogique afin de faciliter l'appréhension d'un principe peu maîtrisé.

MYTHE N° 7 :

LE PRIVACY BY DEFAULT TRADUIT UNIQUEMENT LE PRINCIPE DE MINIMISATION !

Assurer la protection implicite de la vie privée

« On peut être sûr d'une chose : la protection intégrée de la vie privée est implicite. Elle vise à procurer le maximum de vie privée en veillant à ce que les renseignements personnels soient systématiquement protégés au sein des systèmes informatiques ou dans le cadre des pratiques internes. Donc, la vie privée d'un particulier est protégée même si ce dernier ne pose aucun geste, car la protection de la vie privée est intégrée dans le système, implicitement. » Ann CAVOUKIAN

Il s'agit dans un premier temps d'assurer la minimisation de la collecte des données à caractère personnel, c'est vrai.

Mais il s'agit également d'offrir le maximum de vie privée à l'utilisateur, la protection de la vie privée n'est pas optionnelle, en effet un utilisateur doit bénéficier d'une protection maximale sans aucune intervention de sa part.

La notion de « protection implicite » définit la protection des données à caractère personnel comme une convenance qui vise à protéger l'intérêt individuel de chacun, une forme de savoir-vivre ou de bienséance numérique. »

MYTHES ET LEGENDES DU MULTILINGUISME ET DU PLURILINGUISME DANS L'INTERNET

Jean-Yves Gresser, Société française de terminologie

Pour le Conseil de l'Europe :

- le «multilinguisme» renvoie à la présence, dans une zone géographique déterminée – quelle que soit sa taille – à plus d'une «variété de langues», c'est-à-dire de façons de parler d'un groupe social, que celles-ci soient officiellement reconnues en tant que langues ou non. À l'intérieur d'une telle zone géographique, chaque individu peut être monolingue et ne parler que sa propre variété de langue ;

- le «plurilinguisme» se rapporte au répertoire de langues utilisées par un individu ; il est donc, en un sens, le contraire du multilinguisme⁸. Ce répertoire englobe la variété de langue considérée comme «langue maternelle» ou «première langue», ainsi que toute autre langue ou variété de langue, dont le nombre peut être illimité. Ainsi, certaines zones géographiques multilingues peuvent être peuplées à la fois de personnes monolingues et de personnes plurilingues.⁹

Notions contraires, vraiment, ou voisines ? Ma thèse est que le multilinguisme ne peut fonctionner sans une part importante de plurilinguisme.

Le multilinguisme, c'est la capacité pour tout internaute de pouvoir communiquer dans sa (ou ses) langue(s), notamment d'accéder aux informations de la toile (www) ou d'autres sous-réseaux. Le premier défi est de pouvoir formuler une requête :ou ses) langue(s). L'exemple suivant m'a été donné par mon ami Jérôme Trollet (JT), ancien président de la la CSTIC¹⁰ :

Si un Indien m'écrit à cette adresse ग्यारह, सड़क नाइट हत्या कर दी पेरिस फ्रांस pour signifier 11, rue d'Assas Paris France, je doute que les services de La Poste (j'ai pourtant une considération amicale pour mon facteur !) réussissent à me faire parvenir sa lettre. Il en est ainsi pour l'internet ; mon ami Indien souhaiterait accéder à Wikipédia pour savoir ce que signifie le mot « ontologie » : il doit en premier lieu « se rendre chez Wikipédia » et taper वेब.विकिपीडिया.कॉम pour www.wikipedia.com avant de rechercher le terme आंटलजी pour « Ontologie ».

Sinon le danger est que chaque locuteur ait tendance à rester cantonné dans son propre univers linguistique, alors qu'un des grands enjeux planétaires est le partage des savoirs et l'accès aux services ou produits au niveau mondial. En bref, l'idéal serait que l'internaute puisse être plurilingue sans faire d'effort.

⁸ C'est moi qui souligne (JYG)

⁹ Voir http://www.coe.int/t/dg4/linguistic/division_fr.asp

¹⁰ Commission spécialisée de terminologie et de néologie de l'informatique et des composants électroniques, du dispositif dit « d'enrichissement de la langue française ».

MYTHE N° 1 :

LE MULTILINGUISME, C'EST POUVOIR DONNER UN NOM DE DOMAINE DANS SA LANGUE.

Pour reprendre l'exemple de JT, l'enjeu porte sur deux aspects : l'accès (l'adressage) et la compréhension (le contenu). Lorsque l'on parle de langue sur l'internet, c'est d'abord pour évoquer l'accès, accès où les noms de domaines jouent un rôle important sinon essentiel. Depuis des décennies, le sujet fait, à juste titre, l'objet de nombreux débats. Le dernier en date concerne l'offre de nouveaux suffixes.

Le premier enjeu est de pouvoir créer un nom de domaine, quel qu'il soit, c'est à dire qui ne soit pas déjà enregistré, et de le garder. Cela implique le versement d'une redevance modique à un organisme d'enregistrement.

Le second enjeu est de se protéger de l'usurpation d'un nom que l'on considère comme propriété intellectuelle. Cet enjeu concerne surtout les entreprises disposant d'une marque reconnue. C'est d'abord l'affaire des juristes et des tribunaux de commerce. Avec la multiplication des suffixes cela risque de se compliquer, mais ni plus ni moins que dans la pratique actuelle : des entreprises peuvent porter le même nom si elles opèrent dans des domaines différents.

Ces enjeux sont importants¹¹ mais ni plus ni moins que dans les pratiques antérieures à l'internet. Ils en masquent d'autres d'une portée plus étendue.

Le multilinguisme est effectivement de pouvoir utiliser sa propre langue, avec ses propres graphies, si elle en a, pour désigner un domaine. Mais cela ne suffit pas : il faut que tout internaute puisse accéder à ce domaine même si sa langue n'est pas la langue d'origine. Le multilinguisme est insuffisant sans plurilinguisme.

La question est double :

- comment une langue peut-elle être présente sur la toile ?
- comment peut-elle être véhicule de savoirs au niveau mondial (ou régional) ?

MYTHE N° 2 :

UNE LANGUE DOIT S'ECRIRE POUR ETRE VEHICULE DE SAVOIR SUR LA TOILE.

C'est encore largement vrai mais les contenus sonores ou visuels non textuels sont maintenant fréquents. Par ailleurs, la reconnaissance et la synthèse vocale se développent. Elles sont encore limitées aux grandes langues véhiculaires donc écrites, mais pourraient s'étendre aux autres ou à des formes dialectales non écrites de langue écrite (chinois, arabe etc.).

Il n'en reste pas moins que l'accès à une forme écrite représente, pour une langue ou un dialecte, l'accès à un niveau plus étendu de communication : accessibilité et passage vers d'autres langues.

¹¹ Voir entretien avec L. Pouzin

MYTHE N° 3 :

EN 2012, TOUS LES TEXTES REDIGES DANS LES LANGUES ECRITES CONTEMPORAINES PEUVENT S'ECRIRE ET SE LIRE NUMERIQUEMENT.

Certes, on s'en rapproche mais il reste encore des progrès à faire pour les langues qui n'ont pas pour base l'alphabet latin, et celles qui n'ont pas non plus d'écriture linéaire (écritures indiennes, coréen etc.).

La production numérique des textes dans ces langues est un préalable. Celle-ci requiert l'existence d'un codage adéquat des caractères, d'un clavier ou de tout autre système de saisie de caractère, de composition numérique des chaînes lexicales ou textuelles. Cela dépasse l'internet. C'est d'abord l'affaire des fournisseurs de claviers (réels ou virtuels¹²) et de ceux de logiciels de traitement de texte. Hors les grandes langues véhiculaires, c'est encore la grande débrouille.

Le traitement des requêtes repose sur la lecture et l'interprétation, autrement dit le décodage des caractères et du sens. L'interprétation est précédée ou non d'une traduction. Celle-ci est indispensable pour l'accès aux savoirs etc. dont la langue d'origine n'est pas la langue de requête. Mais il se peut aussi que les savoirs etc. d'une langue soient indexés dans une langue différente.

MYTHE N° 4 :

IL SUFFIT QU'UNE LANGUE SOIT ECRITE ET LUE NUMERIQUEMENT POUR QUE LES SAVOIRS ETC. SOIENT FACILEMENT ACCESSIBLES.

Première étape sur laquelle nous n'insisterons pas : ces savoirs etc. doivent être numérisés et stockés sur un média accessible. C'est pour chaque langue un projet colossal. Chacun y a son rôle, collectif ou individuel, public ou privé.

La deuxième étape est leur indexation et/ou leur référencement. Cela n'est pas indispensable mais reste fondamental pour la facilité et la pertinence d'accès. Indexation ou référencement font couramment appel à des référentiels lexicaux (mots ou groupes de mots, traités comme chaînes de caractères simplement listés ou regroupés en thesaurus), plus rarement terminologiques (terme = désignation + définition) ou même onto-terminologiques (termes et modèles sémantiques).

MYTHE N° 5 :

IL EXISTE DES REFERENTIELS LINGUISTIQUES MULTILINGUES¹³.

Les référentiels linguistiques sont généraux (dictionnaires) ou propres à un domaine particulier (vocabulaires, lexiques¹⁴). La plupart sont monolingues. Il existe pour la traduction (exercice qui va d'une langue source à une langue cible) des dictionnaires bilingues.

Ces dictionnaires sont d'une grande utilité mais ils se limitent souvent au mot à mot. Quant on a testé la traduction automatique du japonais sur Google on en connaît les limites.

Certains dictionnaires en ligne, notamment de chinois, fournissent heureusement de nombreux exemples, comme les dictionnaires classiques de latin ou de grec.

¹² Je salue à cette occasion le travail de Lexilogos.

¹³ A ce niveau, les adjectifs plurilingues et multilingues sont synonymes.

¹⁴ En anglais : glossary.

Dans certains domaines comme la médecine, il existe des référentiels terminologiques d'une langue construits à partir de référentiels d'une autre langue. Ces référentiels sont fondamentaux mais ils ont aussi leurs limites.

Les référentiels qui traitent plus de deux langues n'existent que des contextes particuliers : instances internationales (Nations unies, Commission européenne, UIT) ou Etats fédéraux multilingues (confédération helvétique, Canada) pour un nombre plus ou moins grand de langues avec vocation universelle ou restreinte aux activités de l'instance.

Il existe quelques bases onto-terminologiques « fourre-tout »¹⁵ dans lesquelles les modèles sont accumulés dans plusieurs langues sans souci particulier de cohérence. Ces bases rappellent la tentative de DARPA de constituer à la fin des années 2000 un référentiel universel d'ontologies. Sans un minimum de cohérence, sémantique ou même instrumentale, celui-ci a vite perdu de son intérêt.

Dans des domaines très spécialisés (Continuité d'activité, Forum tripartite) il existe des référentiels qui couvrent une dizaine de langues pour des vocabulaires de l'ordre de la 100 e de mots.

MYTHE N° 6 :

LES REFERENTIELS LINGUISTIQUES SONT OUVERTS, C'EST-A-DIRE D'ACCES LIBRE ET GRATUIT.

En fait, il y en a de quatre types : ceux produits par des organismes publics ou privés et ceux d'accès gratuit ou payant. La tendance pour ceux du secteur public est la gratuité mais ce n'est pas la règle générale. Par ailleurs, la distinction public/privé n'est pas cohérente d'un pays à l'autre, non plus que la politique de certains types d'organismes : normalisation, tutelle sectorielle etc.

D'une façon ou d'une autre, les travaux terminologiques ou même encyclopédiques, car il s'agit bien de cela, doivent être financés. Le cas de Wikipédia montre que compter sur toutes les bonnes volontés du monde n'est pas suffisant.

Cela passe nécessairement par des passionnés de la langue et du savoir mais cela suffit-il ? Nous savons que la langue est au cœur de multiples enjeux mais quel projet peut-il être suffisamment parlant (c'est le cas de le dire) et motivant pour les entreprises et le grand public ?

N'y a-t-il de recours que dans la puissance régaliennne ou dans une Université consciente de sa mission de création et de transmission des savoirs ? L'exemple de l'U. de Chicago...

MYTHE N° 7

A PARTIR DU MOMENT OU JE PEUX FORMULER UNE REQUETE DANS MA LANGUE, J'AI ACCES A TOUS LES SAVOIRS DE LA TOILE.

Même si c'est déjà un pas important de pouvoir utiliser sa langue sur la toile¹⁶. Cela ne suffit pas. Comme nous l'avons dit l'interprétation d'une requête se limite souvent à une opération logique sur les chaînes de caractères sensées constituer des mots ou des groupes lexicaux. Le moteur de recherche va traduire mot à mot et essayer de localiser les occurrences des équivalents ainsi trouvés. Le résultat peut être déroutant. La probabilité de trouver la bonne

¹⁵ ISOCat...

¹⁶ Voir colloque 2011 du CNRS.

information à partir d'une langue étrangère se réduit fortement par rapport à celle d'une requête formulée dans la langue cible.

Par ailleurs, les moteurs actuels butent déjà sur le nombre des éléments à référencer et les internautes sont facilement décontenancés devant les résultats de requêtes fournissant des centaines, sinon des milliers de références. Google et consors ne nous sauverons pas sur ce terrain.

MYTHE N° 8 :

UNE LANGUE DOIT POUVOIR TOUT DIRE.

Sinon elle meurt. C'est du moins la thèse de M. Serres. Effectivement pour rester vivante une langue doit pouvoir nommer, classer les phénomènes nouveaux et favoriser tout processus d'invention ou de création. Mais peut-elle vraiment TOUT nommer ? Chaque langue véhicule une vision du monde qui lui est largement propre. La différence est moins sensible entre langues d'une même aire (ou ère) de civilisation, encore moins entre langues parentes comme les langues latines mais, même dans ce cas, il est des subtilités ou des faux amis qu'il vaut mieux connaître dans la vie pratique.

Mais le multilinguisme vu comme la compréhension de l'analyse d'une situation ou d'un phénomène dans une autre langue peut faire ressortir des éléments critiques, imperceptibles dans une autre. Certains le comparent à l'apport de la 3D dans une image plate. Il peut s'agir de beaucoup plus.

Ce genre d'expérience nous l'avons vécu à l'intérieur de certaines langues ou au moment de l'abandon de l'utilisation du latin des langues modernes à de grands tournants de notre histoire scientifique ainsi en chimie lors du passage de l'alchimie à la chimie moderne, ou de la médecine descriptive à la médecine systémique. Pensons aussi aux concepts du taoïsme inspirant la physique des particules.

Ces expériences, nous pouvons, sur la toile, les vivre de manière quasi-instantanée en accédant aux savoirs originaux ou en essayant d'adapter nos savoirs, produits, services à d'autres modes de pensée et de vivre.

MYTHE N° 9 :

NOUS SERONS SAUVES PAR LES REFERENTIELS EN ANGLAIS ET PAR LES INITIATIVES PRIVEES.

Ces référentiels sont importants mais, pour les raisons exposées à propos du mythe n°9 - limitation à une vision particulière du monde- ils ne peuvent constituer la base unique des travaux d'indexation et de classification.

Par ailleurs, pour des raisons liées essentiellement au fonctionnement social étatsunien, ces référentiels sont payants et il est difficile de les enrichir.

Ainsi les référentiels de médecine sont l'objet d'un bras de fer permanent entre certains auteurs et adaptateurs. Dans ce domaine, on pourrait, sans doute, s'accommoder d'un enjeu économique : il s'agit, certes, de diagnostic et de thérapie mais de répertoires et non de mise en œuvre. Il n'empêche que déjà à ce niveau l'enjeu peut se chiffrer en milliers sinon millions de vies !

MYTHE N° 10 :
NOUS SERONS SAUVES PAR WIKIPEDIA

Hélas pour le multilinguisme, Wikipedia et de nombreuses entreprises similaires sont monolingues. Cela entraîne, outre les défauts bien connus de l'entreprise et en partie à cause de ces défauts, la constitution de savoirs inhomogènes d'un univers linguistique à un autre, susceptibles d'entretenir les différences plutôt que de les atténuer.

On en revient d'abord à la pertinence des articles, que l'on peut étendre à celle de toute information circulant sur la toile, et à celle de la traduction ou de l'adaptation des savoirs.

Mais ce n'est pas la seule piste : les référentiels onto-terminologiques privilégient le sens, les « concepts » par rapport aux mots. Ils sont donc susceptibles d'être mieux compris d'une langue à l'autre... Mais ils nécessitent un travail d'un genre nouveau... Un des grands défis du XXIe ?

MYTHES ET LÉGENDES DE LA BLOCKCHAIN

Gérard Peliks, VP de Forum ATENA (avril 2016)

MYTHE N° 1 :

LA BLOCKCHAIN, C'EST LES BITCOINS

Il est difficile de répondre : « *non pas du tout* » à ce mythe, tant il est partagé par ceux qui ont quelques notions sur le sujet. Mais Blockchain, Bitcoin, est-ce bien pareil ? En fait, non, ce sont deux notions différentes mais intimement liées. **La blockchain est une technologie, le bitcoin est une cryptomonnaie.** Il ne peut y avoir de bitcoins sans blockchain, mais la blockchain peut soutenir des cryptomonnaies autres que les bitcoins. Il existe des centaines d'autres monnaies virtuelles, basées sur des blockchains, mais il est vrai que le bitcoin est de loin la cryptomonnaie la plus utilisée. Il existe aussi des blockchains qui soutiennent des éléments qui ne sont pas des cryptomonnaies, comme par exemple les contrats intelligents qui reposent sur une blockchain nommée Ethereum, dont l'unité monétaire est l'éther. Tout comme aujourd'hui pour le Cloud, il existera des blockchains publiques, comme celle qui gère les bitcoins, des blockchains privées et des blockchains hybrides.

Pour constituer une chaîne de blocs, d'où le nom de blockchain, il faut que certaines personnes, appelées les mineurs, créent des blocs qui vont s'ajouter aux blocs déjà existants dans la blockchain. Ils ne font pas ça que pour le plaisir. Ils utilisent des ressources importantes en calcul et y passent du temps, donc ils s'attendent à être rémunérés pour leurs actions. Là, les bitcoins interviennent pour rémunérer les mineurs pour que la blockchain puisse se développer. La rémunération des mineurs en bitcoins est d'ailleurs la seule possibilité de créer les nouveaux bitcoins.

MYTHE N° 2 :

LE BITCOIN A UNE REALITE MATERIELLE



On représente souvent un bitcoin par une pièce de monnaie frappée d'un grand « B » avec deux petits traits en haut et deux petits traits en bas. De plus on parle de machines distributrices où l'on entre des euros ou des dollars et on obtient en échange des bitcoins et inversement. Ceci rend l'apparence physique du bitcoin crédible, mais dans la réalité, il n'en est rien.

Détenir des bitcoins, en pratique, c'est détenir **une clé privée**, c'est-à-dire un élément numérique d'une certaine longueur, constitué de bits à 0 et de bits à 1, ou si vous préférez d'une suite de chiffres et de lettres, qu'il faut absolument conserver et tenir secrète.

À quoi sert au juste cette clé privée ? À signer électroniquement ses transactions payées en bitcoins, pour que la transaction est, et reste intègre, définitive et incorruptible. Tout le monde pourra vérifier que la transaction a été faite, puisque chaque transaction entrée dans la blockchain est visualisable en clair dans un des blocs de cette blockchain. Ceci atteste qu'à un moment donné, un même bitcoin n'est détenu que par une seule personne. La blockchain des bitcoins contient tous les blocs entrés depuis sa création. Cette blockchain a été créée par Satoshi Nakamoto en 2009. Ce personnage est certainement bien réel, et détient un pourcentage non négligeable des bitcoins en circulation, mais son nom est sans doute un mythe, ou au moins un pseudonyme ?

Nous ne rentrerons pas ici dans les détails de la signature électronique et du chiffrement à clé publique (dit encore chiffrement asymétrique), utilisés par la technologie de la blockchain. Retenez juste que si vous perdez votre clé privée, vous aurez perdu en même temps tous vos bitcoins, car seul celui qui possède cette clé privée, donc en principe vous, peut dépenser vos bitcoins. Si quelqu'un prend connaissance de votre clé privée, il pourra faire les transactions à votre place (en signant électroniquement pour vous), donc c'est comme s'il vous avait volé vos bitcoins. Précisons ici que dans le grand livre, appelé en anglais le ledger, les transactions sont signées mais anonymes. On peut juste savoir, en lisant les blocs de la chaîne, que tel bitcoin appartient, ou a appartenu à tel possesseur de la clé secrète qui a servi à signer la transaction, mais sans nommer le possesseur. Comme avec la monnaie fiduciaire, on ne connaît pas l'identité de celui qui possède ou a possédé tel bitcoin. Votre nom n'est pas sur les billets de 5, 10 ou 20 euros que vous possédez ? Pour les bitcoins, c'est pareil.

Mais rassurez-vous, pour utiliser vos bitcoins, vous n'avez aucun besoin de maîtriser des connaissances en cryptologie. Vous avez juste besoin de posséder un porte-monnaie numérique, qui est un programme que vous téléchargez pour mener des transactions avec des bitcoins. C'est ce programme qui s'occupe de la création et de l'utilisation des clés cryptographiques. Et votre porte-monnaie peut être anonyme, ce qui a parfois déjà donné un aspect sulfureux au bitcoin accusé d'être un vecteur de blanchiment de l'argent mal acquis.

Donc un bitcoin n'a qu'une réalité virtuelle, mais quand même il est une réalité numérique, matérialisée par des transactions électroniques qui sont inscrites dans les blocs de la blockchain. Il est utilisé par celui qui en est propriétaire parce que celui-ci possède la clé privée qui prouve sa propriété quand il l'utilise. Un bitcoin n'est donc en aucun cas une pièce de monnaie physique avec un « B » et des petites barres verticales parallèles.

Pour être virtuelle, ou plus précisément numérique, la monnaie bitcoin est loin d'être une monnaie qui n'existe pas. Aujourd'hui, 15 millions de bitcoins sont en circulation, capitalisés à 5,7 milliards d'euros, et de plus en plus de commerces les acceptent en paiement.

Ça vous paraît extraordinaire d'utiliser une monnaie qui n'a pas d'apparence physique ? Mais vous le faites pourtant depuis longtemps si vous avez accumulé des miles dans les compagnies aériennes ou des points dans les supermarchés. Et sans aller jusque-là, vos transactions bancaires, n'est-ce pas des écritures purement numériques ?

MYTHE N° 3 :

LA BLOCKCHAIN EST UN LIVRE NUMERIQUE DISTRIBUE A TOUS LES UTILISATEURS

La blockchain est un grand livre de comptes (un « ledger », en anglais). Créé en 2009, pour les bitcoins, ce livre de comptes ne fait que grandir à la vitesse d'un bloc toutes les 10 minutes. Il a atteint aujourd'hui une taille de 62 Go (62 milliards d'octets).

Mais non, ce grand livre n'est pas distribué, comme on le croit souvent, à tous les utilisateurs. C'est la confiance en ce système qui est distribuée à tous les utilisateurs, c'est du moins le but visé par cette technologie qui n'implique aucune autorité centralisatrice et garante des transactions. D'ailleurs, aimeriez-vous voir arriver sur votre smartphone ou votre tablette un fichier de 62 Go, juste parce que vous êtes un adepte des paiements par bitcoins ? De plus, ce fichier doit être synchronisé constamment par « Peer to Peer », puisque un nouveau bloc se crée toutes les 10 minutes, et tous les grands livres doivent être identiques, donc synchronisés entre eux. Alors gardez votre puissance de calcul limitée et votre faible espace numérique disponible sur votre smartphone ou votre tablette pour votre usage personnel, en oubliant le grand livre. Ce n'est pas sur votre smartphone ou votre tablette que ça se passe.

En fait, les grands livres résident, tous identiques après synchronisation, sur des serveurs utilisés par des personnes appelées les mineurs. Il y a environ 8000 serveurs dans le monde, qui s'occupent de faire fonctionner la blockchain des bitcoins. Les mineurs ont à leur disposition la capacité disque et la puissance de calcul qu'il faut, mais pas vous. Vous avez, en tant qu'utilisateur de bitcoins, juste accès à ces grands livres, situés sur les 8000 serveurs, où toutes les transactions sont en clair et peuvent être lues.

MYTHE N° 4 :

JE PEUX DEVENIR UN MINEUR ET GAGNER BEAUCOUP DE BITCOINS

Oui, on ne naît pas mineur, on le devient, alors pourquoi pas vous ?

D'autant plus qu'il est très lucratif d'être mineur. 25 bitcoins sont créés et distribués toutes les dix minutes, 24 heures sur 24, 7 jours sur 7, aux mineurs qui ont la chance d'avoir été choisis parce qu'ils ont résolu, dans les dix minutes, un problème mathématique, qui n'a rien de génial ni de très glorieux à avoir été résolu, mais qui est très coûteux en ressources de calcul. Le mineur gagnant valide les transactions qu'il entre dans le bloc qu'il vient de créer. Il détient alors les 25 bitcoins qui viennent d'être créés en rémunération de son travail. C'est d'ailleurs ainsi et ainsi seulement que de nouveaux bitcoins sont créés. Le mineur peut aussi demander à être rémunéré, en plus, en percevant un pourcentage du montant en bitcoins de la transaction qu'il valide en la rentrant dans le bloc le plus récent de la blockchain.

Le problème est justement de pouvoir résoudre ce fameux problème mathématique, toujours renouvelé, et qui devient de plus en plus coûteux en ressources avec le temps, car la difficulté de résolution du calcul s'adapte à l'accroissement de la puissance des serveurs disponible et même la précède. Cette résolution réclame aujourd'hui une puissance de calcul très grande, et vous n'en disposez pas, sauf si vous travaillez dans un organisme tel que, par exemple, dans un énorme Data Center. Et encore on peut douter que le Data Center vous prête ses très gros serveurs pour effectuer les calculs de minage pour vous permettre de gagner des bitcoins. Donc les mineurs sont ceux qui peuvent résoudre le problème posé. On compte environ 100 000 mineurs dans le monde, qui se regroupent pour disposer de la puissance de calcul indispensable pour être éligibles pour créer un bloc et gagner les bitcoins en récompense. Insistons encore sur le fait que c'est la seule méthode de création de nouveaux bitcoins.

Domage pour vous, car 25 bitcoins, au cours aujourd'hui d'environ 400 euros pièce, ça fait quand même un gain de 10 000 euros toutes les 10 minutes, imaginez ce que vous pourriez faire avec ! Bon, ici on combat les mythes, mais on n'interdit pas de rêver.

MYTHE N° 5 :

LA BLOCKCHAIN ENTRAINERAIT LA DISPARITION DE NOMBREUX METIERS

La blockchain, c'est comme la marche vers la numérisation des organisations, ça fait évoluer certains emplois, ça en détruit d'autres et ça en crée de nouveaux. La blockchain étant une innovation de rupture, mieux vaut surveiller cette technologie pour ne pas être pris de court. Comme on dit, mieux vaut penser le changement que changer le pansement.

La principale et géniale innovation de la blockchain est de faire reposer la confiance sur la multitude des utilisateurs plutôt que sur un tiers de confiance unique. Par exemple, une transaction immobilière aujourd'hui repose sur un notaire, une transaction financière repose sur un banquier, une monnaie repose sur un Etat qui exerce sa puissance régaliennne, un certificat numérique, pour ne pas oublier de parler aussi de cryptologie, repose sur une signature d'un tiers de confiance reconnu.

Mais pourquoi la confiance devrait-elle reposer sur une entité unique ? En fait, ce n'est pas indispensable, et c'est justement à ça que répond la blockchain. Avec cette technologie, la confiance repose sur un livre de compte infalsifiable, distribué vers de nombreux endroits et géré par personne en particulier. La confiance ne se décrète pas, elle s'établit, basée sur des sentiments. Vous préférez faire confiance à une entité centralisatrice parce qu'avec elle au moins on sait à qui s'adresser ? Oui mais si cette entité centralisatrice venait à disparaître ? Cette disparition n'est pas envisageable avec une blockchain répartie sur plusieurs milliers de serveurs qui sont synchronisés pour posséder la même chaîne de blocs, consultable par tous les utilisateurs, et infalsifiable.

Alors la blockchain ubérise-t-elle le métier de notaire ? Ubérise-t-elle le métier de banquier ? Ubérise-t-elle la prérogative régaliennne d'un Etat qui garantit sa monnaie ? La blockchain ira-t-elle jusqu'à ubériser Uber ... ? À voir. Il est probable que ces métiers devront tenir compte de cette technologie révolutionnaire et donc on peut penser raisonnablement que ces métiers vont être profondément modifiés pour continuer à exister. D'ailleurs certains l'ont bien compris, en particulier les banques, les assurances, l'administration lancent des projets pour mieux comprendre les implications des blockchains et étudient les opportunités de maîtriser cette technologie plutôt que d'avoir à la subir. Et c'est pourquoi, la blockchain est une technologie de rupture.

MYTHE N° 6 :

LA BLOCKCHAIN EST ECOLOGIQUE

Pourquoi écologique ? Parce qu'elle permettra, avec les bitcoins, entités numériques, d'économiser le métal avec lequel sont faites les pièces de monnaie ? D'économiser le papier avec lequel sont faits les billets de banque et les « contrats intelligents » dans le cas de la blockchain ethereum ?

Et bien c'est tout le contraire. La blockchain, en tout cas la blockchain des bitcoins telle qu'elle fonctionne aujourd'hui, s'avère être une catastrophe écologique ! Nous avons évoqué le fait que, pour être élus comme créateur d'un nouveau bloc de la blockchain, et récolter quelques bitcoins en paiement, les mineurs doivent avoir trouvé la solution à un problème posé, pas génial du tout dans sa résolution, mais extrêmement consommateur en ressources de calcul. Pour résoudre ce problème, il faut utiliser de très grandes capacités de calculs, et on peut ajouter les ressources pour synchroniser en pair à pair les serveurs. Et pour faire tourner les moyens de calcul, il faut de l'électricité.

Aujourd'hui, la consommation d'électricité nécessaire, rien que pour avoir le droit de créer de nouveaux blocs dans la blockchain, est de l'ordre de celle d'un pays comme le Danemark !

Produire de l'électricité tel qu'on le fait aujourd'hui n'est pas vraiment bon pour la planète, et bien évidemment, s'il n'y avait plus d'électricité produite, il n'y aurait plus de bitcoins ni de blockchain.

MYTHES ET LEGENDES DU BIG DATA

Ivan de Lastours, Direction de l'innovation de l'Institut Mines-Télécom

INTRODUCTION

"90% de la donnée mondiale a été générée ces deux dernières années et 90% des données utilisées aujourd'hui seront différentes de celles utilisées dans deux années" S. Gold, IBM.

Cette citation met en avant une réflexion globale sur le stockage, la consommation et l'exploitation des données...

Cette approche a un nom "Le Big Data", mais qu'est-ce donc? Une science? Un argument marketing? Nous tentons ici d'éclaircir ce "Big" sujet et d'en cerner un peu mieux les contours. (Achevé de rédiger en janvier 2013).

MYTHE N° 1 :

LE BIG DATA, C'EST LE CLOUD COMPUTING 2.0

Impossible que vous n'ayez pas entendu prononcer ce terme depuis moins de deux semaines. Tout le monde en parle, même votre voisine de palier, et c'est de plus en plus la vraie tarte à la crème des recruteurs...fini le Cloud, place au Big Data !

Mais qu'est-ce donc ? « Le Big Data est l'ensemble des datas qui dépassent les capacités de traitement des systèmes conventionnels de bases de données. La donnée est trop volumineuse, change trop vite ou ne peut être stockée dans un schéma classique de base de données. Pour exploiter cette data, une nouvelle approche doit être utilisée : le Big Data. » Ed Dumbill de la société O'reilly

Quelle est son origine ? Les géants du web (Google, Amazon et Facebook, LinkedIn, ebay) ont fortement contribué à l'émergence du Big Data avec le lancement d'algorithmes (MapReduce), de framework (Hadoop) et d'outils spécifiques (outils de requêtage Hive développé par Facebook par exemple) pour utiliser celui-ci.

Mais concrètement qu'est-ce que c'est ? Du hardware ? Du software ? Les deux mon capitaine...

Si certains le définissent comme les 4 « V » Volumétrie, Vitesse, Variété et Valeur des données ; le « Big Data » représente encore plus.

Il s'agit d'un écosystème beaucoup plus complet qui met en jeu les éléments suivants:

- De la data...
 - Structurée/non structurée
 - Générée par l'humain/générée par la machine
 - Statique/dynamique
- Une plateforme dédiée (ou louée)
 - Architecture, infrastructure et matériel informatique (stockage, accès, réseau)
 - Frameworks et outils logiciels (Hadoop, Cloudera...)
 - Ressources humaines (ingénieurs, « data scientists »)
- Des procédures d'analyses des données
 - Algorithme, automatisation

- Requête SQL (Structured Query Language, langage informatique effectuant des opérations sur des bases de données relationnelles)
- Business Intelligence
- Un usage « business » efficace
 - Remontée d'informations pertinentes
 - Suivi des "metrics"
 - Prise de décision

Le Big Data n'est donc pas seulement le stockage et l'analyse d'octets. Il s'agit d'exploiter et de valoriser les données de l'institution/l'entreprise dans le cadre de son exercice.

MYTHE N° 2 :

MON ENTREPRISE UTILISE DÉJÀ LE BIG DATA ET CE DEPUIS PLUSIEURS ANNÉES

De nombreux spécialistes autoproclamés du Big Data ont émergé ces trois dernières années mais certains n'ont jamais mené en interne le moindre projet Big Data ! N'est pas spécialiste qui veut.

Au contraire certaines entreprises spécialisées dans l'exploitation et les traitements de données font quasiment déjà du « Big Data » sans même le savoir.

Il ne suffit pas d'avoir des données pour faire du Big Data, ni d'avoir une distribution Hadoop (pour rappel Hadoop est un kit de composants logiciels structurels écrit en Java, il permet la création d'applications distribuées et scalables. Il permet aux applications de travailler avec des milliers de nœuds et des pétaoctets de données).

J'ai déjà de la Business Intelligence ai-je besoin de Big Data ?

Pas nécessairement, si mon système est correctement dimensionné pour mes besoins je n'ai aucune utilité à mettre en place du Big Data même de manière complémentaire.

Que faut-il dans une entreprise/institution pour mettre en place du Big Data ?

- Un jeu de données exploitables
- Une infrastructure de base de données (une BDD relationnelle peut même faire l'affaire en phase de test), cette infrastructure peut être louée
- Un cas d'étude exploitable (l'élément le plus bloquant aujourd'hui pour faire du Big Data)
- Une équipe pour extraire et traiter les données (les fameux « Data scientists »)
- Une diffusion et une exploitation des résultats des requêtes pour une utilisation interne ou une revente du résultat (il faut une finalité sinon le Big Data reste à l'état d'expérimentation en interne)

Le cas d'étude ou cas d'usage « business case » est indispensable et est en général ce qui manque le plus aujourd'hui. Il s'agit d'avoir une vision claire de l'utilisation du Big Data en interne (données en entrée pour quels résultats en sortie, quels objectifs...)

Je ne fais pas de Big Data...dois-je m'y mettre ?

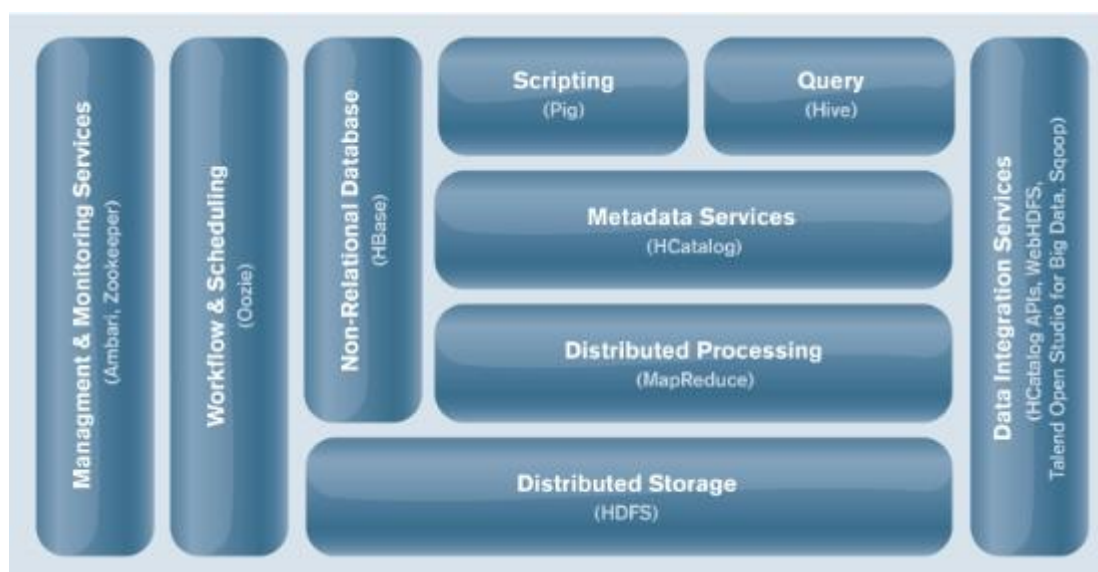
Pas obligatoirement, là aussi cela varie au cas par cas. Si mes volumes de données augmentent ou que je souhaite exploiter des données non structurées qui pourraient me donner un avantage, il vaut mieux commencer par acquérir une expertise en interne. Le risque d'attendre

est que la concurrence développe une approche Big Data et que je doive me mettre à niveau d'un seul coup, ce qui risque de me coûter très cher (solutions clés en main IBM, SAP...)

Je veux faire du Big Data chez moi dans mon garage, de quoi ai-je besoin très concrètement ? Que dois-je installer ?

- Une base de données NoSQL (a d'abord signifié Non-SQL puis Not Only SQL littéralement pas seulement SQL, exemple : Hbase)
- Un framework (un kit de développement) compatible avec l'approche MapReduce (Hadoop)
- Une capacité de stockage importante et un système de fichier adapté (HDFS, en physique/location)
- Des serveurs...beaucoup de serveurs (en physique/location)
- De la capacité de traitement (par du langage R par exemple)
- Des outils de requêtage et d'extraction (Hive)
- Des outils de scripting (langage de programmation haut niveau comme Pig)
- Un module de monitoring/contrôle du système (Zookeeper)
- Un module NLP Natural Language Toolkit si nécessaire (module d'exploitation de données textuelles)
- Un module de Machine Learning (module d'apprentissage de traitement automatique de données comme Mahout)
- Des outils de visualisation (outils graphiques de visualisation des données et des résultats comme Tableau)
- Des outils d'acquisition (outil d'acquisition des données comme Google Refine)
- La liste n'est pas exhaustive...

Fig 1 - Exemple de plateforme Big Data basé sur une technologie Hadoop (Source : Hortonworks)



MYTHE N° 3 :

LE BIG DATA MARQUE UNE RUPTURE TECHNOLOGIQUE DANS L'HISTOIRE DE L'INFORMATIQUE

Le Big Data se base sur des technologies (MapReduce, Hadoop, Hive..) issues des géants du web (Google, Amazon et facebook, LinkedIn, ebay...). Mais cela est-il une rupture technologique ?

Clairement non, en tout cas en termes de technologie pure.

Si il n'y pas de rupture, pourquoi ne l'a-t-on pas fait avant ?

Il n'y avait pas de tels besoins. Les schémas relationnels des bases de données (type SQL) ont dominé le marché à partir des années 60 et convenaient pleinement aux usages.

L'émergence des 4 « V » des données (Volumétrie, Vitesse, Variété et Valeur) a créée les premiers cas d'usage du Big Data. Les premiers exposés à ces problématiques étaient les acteurs du web qui traitaient des volumes massifs de données en quasi temps réel...

Pour résumer, l'émergence du Big Data est la concomitance de 4 effets :

- L'effet Hardware
 - La baisse importante du coût du hardware dans son ensemble et sa « commoditisation » (serveurs, stockage, réseau...)
- L'effet Outils/Framework/Software
 - Le développement par les acteurs du web de systèmes open source de répartition de la charge sur des architectures distribuées et ce sur de très gros volumes (Projet Hadoop et Mapreduce) ont permis de valider d'excellentes performances sur l'exploitation de Big Data (performances meilleures et coûts réduits comparativement à des systèmes de base de données relationnels)
- L'effet Data
 - L'émergence de plus en plus rapide de volumes de données plus importants et plus hétérogènes. Des contenus non structurés pour la plupart, peu adaptés aux bases de données classiques, mais regorgeant d'informations clés pour les entreprises et les administrations
- L'effet Marché
 - L'énorme potentiel économique et scientifique du Big Data a été largement mis en valeur (étude Mckinsey « Big data: The next frontier for innovation »)

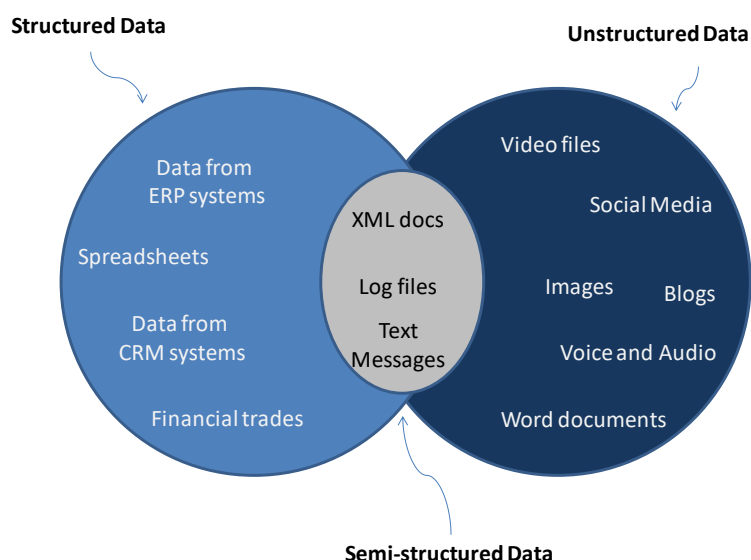
Tous ces effets ont abouti à l'émergence et la consolidation du « Big Data ». Il n'y pas eu de rupture technologique pure, même si on peut parler d'une rupture en terme d'usage des données.

MYTHE N° 4:

LE BIG DATA, SIGNIFIE T'IL LA FIN DES BASES DE DONNEES RELATIONNELLES ?

Si les données structurées et non structurées ne se traitent pas de la même façon, elles ne sont aucunement antinomiques, elles sont au contraire complémentaires.

Fig 2 – Variété des types et sources de données (en évolution constante) Source : JP Morgan

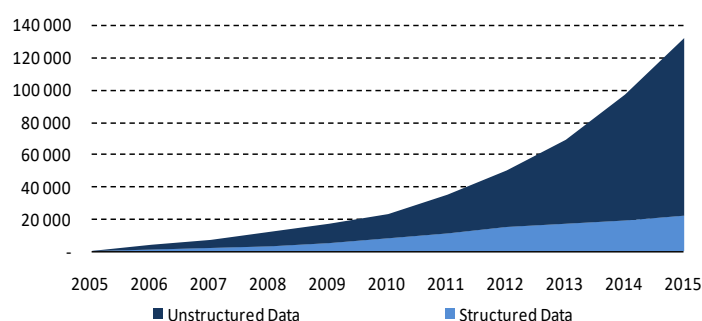


Le Big Data et les bases de données NoSQL (ce qui signifie aujourd'hui Not Only SQL) sont une offre complémentaire des bases de données relationnelles existantes.

En effet, il faut bien faire la différence entre la donnée, l'information structurée et le savoir (interprétation de la donnée et de l'information par l'humain).

La base de données relationnelle permet de traiter de la donnée qualifiée indispensable pour l'exploitation d'autres données moins structurées.

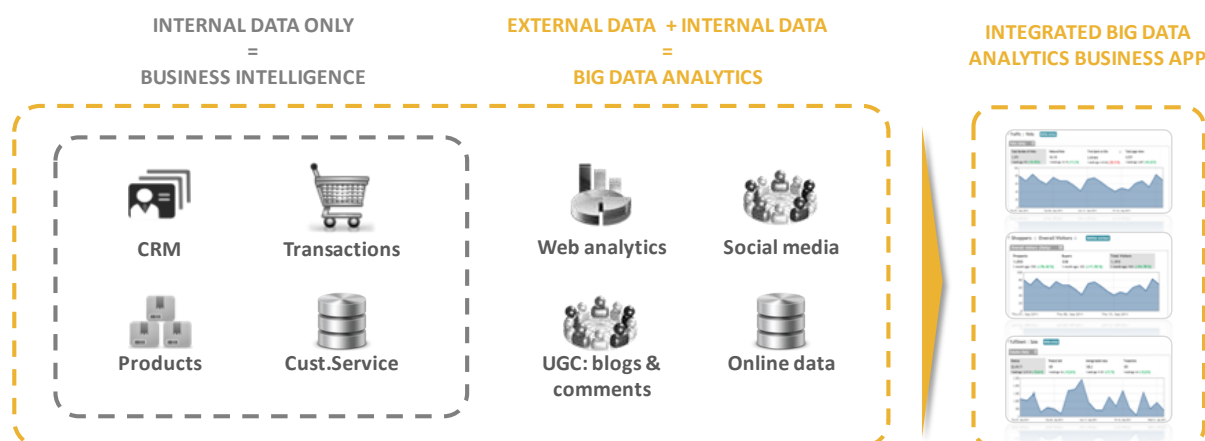
Fig 3 –Evolution des capacités de stockage des entreprises en petabytes entre données structures et non structures (Source CS)



Sur ce schéma on voit bien que la donnée structurée est donc compatible avec des bases de données non-relationnelles et qu'elle est loin de disparaître, même si elle représente moins de volume de données que les données non structurées.

NoSQL regroupe un ensemble de SGBD (systèmes de gestion de base de données) qui ne repose pas sur une architecture classique relationnelle. Il existe cependant des passerelles entre les deux systèmes et ces derniers sont plus complémentaires qu'opposés.

Fig 4 –Le Big Data est l'association de « Business Intelligence » et de « Big Data Analytics » (Source Squid)



MYTHE N° 5 :

LE MARCHÉ DU BIG DATA, QUI SE CHIFFRE GLOBALEMENT A 300 MILLIARDS DE DOLLARS, EST A PRENDRE TOUT DE SUITE

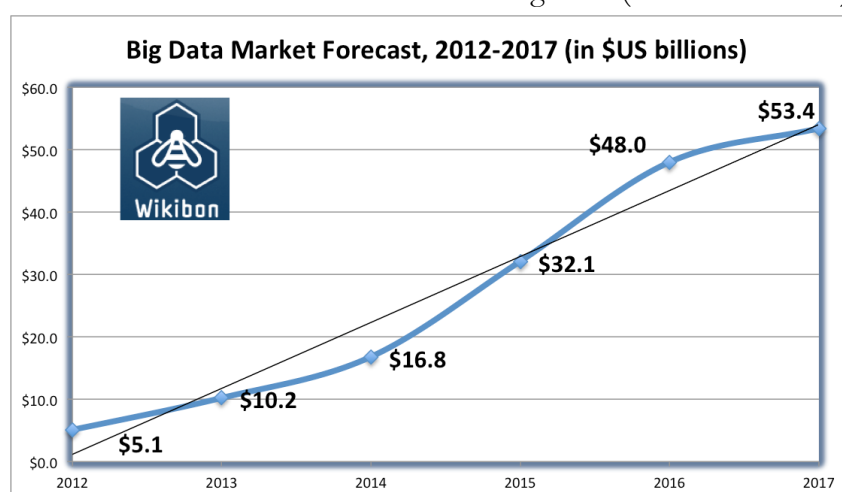
Le marché potentiel est énorme, si on commence à imaginer toutes les applications possibles du Big Data, si l'on observe les premières initiatives et leurs résultats (exemple de la UK Biobank qui regroupe des informations médicales sur plus de 500 000 citoyens anglais), cela peut vite faire tourner la tête...

Le marché existe bel et bien c'est indéniable ! Mais va-t-il réaliser la croissance annoncée ?

Le marché du Big Data représentait 5,1 milliards de \$ en 2012 et devrait réaliser plus de 53 milliards de \$ d'ici 2017. Le marché global est quant à lui estimé à terme à 300 milliards de \$!

Il est aujourd'hui dominé par des grands groupes IBM, Intel, HP. Ce marché se répartit entre 3 pôles : le logiciel 25%, le service 44% et le matériel 31%. Sur le marché des logiciels Big Data (environ 1,3 M\$ en 2012), les purs players du Big Data représentent environ 468 m\$, les autres entreprises ont été intégrées ou font partie de grands groupes. Rien ne semble arrêter le Big Data...

Fig 5 – Prévision du chiffre d'affaires du marché du Big Data (source : Wikibon)



Pour saisir ce marché, il faut investir aujourd'hui pour répondre aux besoins de demain, ou en est-on aujourd'hui ?

Concernant les investissements effectués (quasiment exclusivement aux US).

Plus de 260 m\$ ont été investis dans les start-ups basées sur des technologies Hadoop. Les start-ups basées sur des technologies NoSQL ont levé plus de 90 m\$. L'entreprise Splunk a été introduite au Nasdaq pour un montant de 229 m\$. Plus de 40 acquisitions technologiques significatives ont été effectuées dans le domaine du Big Data par les grands groupes informatiques (Google, Amazon, IBM, EMC, SAP...). À la vue de cette dynamique, le Big Data est un des sujets technologiques d'avenir les plus importants.

Qu'en est-il en France ?

En France, dans le cadre du grand emprunt, 25 M€ ont été spécialement débloqués pour financer des projets innovants sur le sujet du Big Data.

Quels sont les freins potentiels à cette croissance ?

Les principaux freins identifiés à la croissance du marché du Big Data sont les suivants : le recrutement des data scientists, la formation des équipes, la complexité des traitements temps réel, le coût des licences et l'intégration.

Mais le plus grand frein reste sans contexte les usages. S'il n'y a pas de cas d'utilisation, le marché ne sera pas au rendez-vous. Le potentiel est donc indéniable, mais il n'est pas à saisir tout de suite et dépend principalement des usages.

CONCLUSION

Le Big Data existe bel bien, ce n'est pas un mythe et il a un énorme potentiel. Inutile de s'emballer, prudence et expertise sont les maîtres-mots. La technique est là et les coûts restent raisonnables, tout repose donc sur le cas d'usage. L'étude précise du cas d'usage est nécessaire et c'est là-dessus qu'une entreprise/institution doit être ambitieuse et prendre des risques. Il ne sert à rien de reprendre des cas d'usage business intelligence sur base de données relationnelles avec plus de données, il faut réinventer l'exploitation des données de l'institution/entreprise, aller en chercher ailleurs ...

La donnée est la clé, mais attention sans traitement, elle ne devient pas information et sans interprétation elle ne devient pas savoir ou aide à la décision ! La révolution Big Data n'a pas eu lieu, elle se prépare.

Fig 6 – Pour conclure, un peu d'humour ! (Source Dilbert)



MYTHES ET LEGENDES DE LA 4G

Jean-Marc Do Livramento, Consultant Télécom Fixe et Mobile

INTRODUCTION

Depuis fin 2012, des offres 4G sont disponibles en France. La 4G, ou 4^{ème} génération de télécommunication mobile, regroupe un certain nombre de technologies permettant de transférer des données sur un réseau cellulaire en respectant des critères de performance définis par le groupement Radiocommunication de l'Union Internationale des Télécommunications (UIT-R ou ITU-R en anglais), un organisme dépendant des Nations Unies en charge de coordonner au niveau mondial les télécommunications. (Achevé de rédiger en janvier 2013).

Les générations de télécommunication mobile, c'est d'abord une affaire de réglementation.

A ce jour, quatre générations de télécommunication mobile ont été développées dans le monde. Pour chacune d'elles, plusieurs technologies peuvent être utilisées selon les pays voire les opérateurs.

La première génération (1G) regroupait des technologies analogiques. En France, la plus connue d'entre elles était Radiocom 2000.

La deuxième génération (2G) a introduit les technologies numériques. Elle est toujours en usage aujourd'hui. En France, il s'agit du GSM et de ses évolutions permettant le transport de données sur le protocole IP, le GPRS et l'EDGE. D'autres technologies comme le PDC au Japon ou le CDMAOne aux Etats Unis font également partie de la 2G.

La troisième génération (3G) a vu une amélioration des performances notamment au niveau des débits lors des transferts de données et l'introduction de nouveaux services comme la visiophonie. En France, il s'agit de l'UMTS et son évolution le HSPA. D'autres technologies 3G existent par ailleurs comme le CDMA2000 1xEV aux Etats Unis ou le TD-SCDMA en Chine.

L'UIT-R définit les critères que doivent respecter les technologies de télécommunication mobile afin d'appartenir à l'une ou l'autre de ces générations. La quatrième génération de télécommunication mobile (4G) est définie au sein d'une spécification appelée *IMT évoluées* (ou *IMT-Advanced* en anglais). Les technologies voulant être classées dans cette catégorie doivent respecter un certain nombre de critères dont le plus emblématique est le débit pic du réseau vers le terminal : 100 Mbps en déplacement rapide et 1 Gbps en déplacement lent voire à l'arrêt.

Les générations de télécommunication mobile, c'est ensuite une affaire de technique.

En France, comme ailleurs dans le monde, chaque génération de télécommunication mobile a vu évoluer ses caractéristiques techniques.

La 1G était une technologie analogique sur la bande de fréquences 400 MHz ne permettant que de téléphoner.

La 2G (le GSM, le GPRS et l'EDGE) comprend des technologies numériques utilisant les bandes de fréquences 900 et 1800 MHz, une technique de partage des ressources radio TDMA/FDMA (répartition dans le temps et en fréquences) et les modulations GMSK et

8PSK (codage des données dans un intervalle de temps). En 2G, on peut téléphoner et transférer des données en mobilité à un débit qui a été d'abord de quelques dizaines puis d'un peu plus d'une centaine de Kbps.

La 3G (l'UMTS et le HSPA) a vu arriver une nouvelle bande de fréquences dans le patrimoine des opérateurs mobiles avec le 2100 MHz avant que ceux-ci soient autorisés par le régulateur à utiliser également la bande 900 MHz (ainsi que la bande 1800 MHz s'ils le souhaitent). La technique de partage des ressources radio est de type WCDMA (distinction des communications grâce à un code) et les modulations sont de type QPSK et 16QAM pour les premières versions de ces technologies. En 3G, on peut téléphoner en vocal ou en visio et transférer des données à un débit de plusieurs Mbps.

Bien que l'on parle maintenant de 4G, la 3G continue pour autant à évoluer afin d'offrir plus de débits et une meilleure qualité de services. Le HSPA utilise des bandes de fréquences de plus en plus larges, des modulations de plus en plus élevées et une multiplication des antennes en émission et en réception (MIMO).

En parallèle du HSPA, une nouvelle technologie voit le jour : le LTE. Celle-ci s'est vue attribuer de nouvelles bandes de fréquences, le 800 MHz et le 2600 MHz ; une nouvelle technique d'utilisation des ressources radio, l'OFDM ; une modulation encore plus performante, le 64QAM ainsi que la multiplication des antennes en émission et en réception (MIMO). L'architecture du réseau est totalement refondue puisque le mode circuit disparaît et tous les services sont sur IP y compris la téléphonie qui devient alors une application VoIP. Les débits promis sont de plusieurs dizaines de Mbps ce qui constitue une évolution importante... mais pas suffisante pour respecter les critères 4G définis par l'ITU-R. Le LTE a été classé d'emblée dans la 3G et c'est son évolution, le *LTE-Advanced*, qui a décroché le label 4G grâce à l'utilisation de bandes de fréquences encore plus importantes et un plus grand nombre d'antennes en émission et en réception.

Les générations de télécommunication mobile, c'est enfin une affaire de... marketing.

Avant de pouvoir commercialiser le LTE, les opérateurs doivent investir massivement et refondre une partie de leur architecture réseau :

- Acquérir pour plusieurs centaines de M€ l'autorisation d'utiliser les bandes 800 et 2600 MHz.
- Acquérir et déployer un nouveau réseau :
 - De nouvelles antennes sur les sites radio (les antennes actuelles ne couvrent pas ces nouvelles bandes de fréquences).
 - De nouveaux équipements radio (appelés eNodeB).
 - Un réseau de transport IP vers chaque site radio (l'ATM déployé lors de l'arrivée de la 3G n'est pas réutilisé par le LTE).
 - Un nouveau cœur de réseau (appelé EPC) et un nouveau système de gestion des communications (appelé IMS).

Au regard des montants investis dans cette nouvelle technologie, de l'importance de la refonte des infrastructures réseau et de l'amélioration des services apportés aux clients, un certain nombre d'acteurs télécoms ont cherché à augmenter la différenciation entre le LTE et la 3G et sollicité l'ITU-R afin de requalifier cette technologie pour l'intégrer à la 4G. C'est ainsi que depuis décembre 2010 le LTE, tel que déployé par les opérateurs aujourd'hui, est devenu une

technologie 4G. Statut qui a également été accordé aux évolutions du HSPA (dénommé HSPA+ à partir de la Release 7).

Plusieurs versions du LTE existent actuellement en normalisation : la version initiale appelée « *LTE* » dont les premières ouvertures commerciales ont eu lieu fin 2009 et son évolution appelée « *LTE-Advanced* ». La version initiale du LTE a par ailleurs été normalisée selon deux modes : un mode dit FDD pour lequel les communications en sens montant (du terminal vers le réseau) n'utilisent pas les mêmes porteuses (bandes de fréquences) que les communications en sens descendant) et un mode dit TDD pour lequel les communications en sens montant et descendant se partagent les mêmes porteuses en communiquant alors de manière alternée. Le LTE-TDD fait l'objet d'une attention plus particulière de la part des acteurs chinois dans la mesure où ces opérateurs mobiles utilisaient déjà ce mode pour leur 3G (technologie TD-SCDMA). Seule la version FDD, ouverte commercialement en France depuis fin 2012, sera abordée dans ce *Mythes et légendes de la 4G*.

MYTHE N° 1 :

LA 4G, C'EST MIEUX QUE LA 3G !

C'est vrai !

Dès sa première version, le LTE a été conçu pour optimiser les performances dans le transfert de données en mobilité :

- Une gestion des ressources radio (OFDM du réseau vers les terminaux) réduisant les risques d'interférences lors de la réception des données : les débits sont plus élevés car il y a moins de données à retransférer.
- Une modulation du signal radio transmettant plus d'informations dans un même intervalle de temps (jusqu'à 64QAM) : pour une même largeur de bande de fréquences, les débits sont plus élevés.
- Une canalisation plus grande (jusqu'à 20 MHz) : les débits sont plus élevés parce qu'un transfert de données peut utiliser plus de ressources radio simultanément. Le LTE reste toutefois une technologie agile puisqu'il est possible de l'utiliser également sur des canalisations plus petites (1,4 MHz, 3 MHz, 5 MHz, 10 MHz et 15 MHz) au prix tout de même d'une réduction du débit.
- Une transmission simultanée sur, au moins, deux antennes en émission et en réception (MIMO 2x2) : le débit est plus élevé puisque, en bonnes conditions radio, ce multiplexage répartit les données à transférer sur deux antennes. Même en mauvaises conditions radio, le débit est plus élevé grâce à la *diversité* du MIMO : les mêmes données sont transmises simultanément sur les deux antennes afin que la recombinaison des flux radio donne un meilleur signal.
- Une architecture de réseau simplifiée avec moins d'équipements à traverser lors des transferts de données : c'est une meilleure réactivité dans les échanges car il y a moins de latence lors de la traversée du réseau.
- Une architecture de réseau conçue entièrement en IP : c'est également une meilleure réactivité dans les échanges ainsi qu'une meilleure interopérabilité des services parce qu'il y a moins de transcodages de protocoles.

- Des ressources réseau entièrement dédiées au transport de données (la téléphonie est désormais une application VoIP) : c'est encore une meilleure réactivité dans les échanges car il n'est plus nécessaire d'avoir une coordination entre des ressources gérées en mode circuit (téléphonie, SMS, ...) et des ressources gérées en mode paquet (transfert de données).
- Une accélération des mécanismes de connexion au réseau :
 - Une transition accélérée (100 ms) entre l'état de repos du terminal et son état actif (temps après le 1^{er} clic sur un lien HTML par exemple).
 - Une latence améliorée (20 ms) en état actif (temps après les clics suivants).
 - ➔ C'est une meilleure réactivité pour les applications exigeantes en QoS (jeux en réseau, streaming, ...). On a enfin l'impression d'être connecté en permanence.
- Une optimisation de la mobilité entre les réseaux de générations différentes (fonction appelée *Packet Switched HandOver*) : le temps d'interruption est imperceptible lors du passage d'un réseau 4G vers un réseau 3G ou 2G (ce qui se produit par exemple sur perte de couverture 4G le temps que celle-ci soit équivalente à celles des générations précédentes). Cette fonctionnalité est importante notamment pour ne pas ressentir d'interruption lors d'un appel voix ou visio.
- Une meilleure couverture réseau grâce à l'utilisation de bandes de fréquences moins élevées que pour la 3G : la bande 800 MHz permettra de couvrir les zones rurales et l'intérieur des bâtiments (au lieu de la bande 900 MHz en 3G) et la réutilisation de la bande 1800 MHz permettra de couvrir les zones urbaines (au lieu de la bande 2100 MHz en 3G).

La 4G, c'est donc :

- **Plus de débit du réseau vers les terminaux afin d'échanger encore plus vite des gros fichiers ou voir des vidéos en très haute définition.**
- **Mais aussi plus de débit du terminal vers le réseau afin de développer de nouveaux usages (peer to peer mobile, cloud personnel, ...).**
- **Une meilleure réactivité du réseau pour surfer encore plus vite ou faire des jeux en streaming.**
- **Un réseau de plus grande capacité afin de palier à la saturation des réseaux 3G.**

Avec la 4G, l'expérience client du réseau mobile rejoint l'expérience du réseau fixe. Les services peuvent être conçus d'emblée fixe/mobile. L'utilisateur peut être connecté en permanence à ses services et contenus sans se soucier du réseau qui lui donne l'accès.

MYTHE N° 2 :

LA 4G, C'EST LA CONTINUITÉ DE LA 3G !

Concernant l'expérience client, il y a effectivement une continuité dans la progression des technologies : la qualité sonore des appels voix progresse grâce à l'intégration de CODECs de plus en plus performants (normes AMR-WideBand et AMR WideBand+ qui échantillonne la voix jusqu'à 48 KHz) ; Il en va de même pour l'Internet mobile qui s'améliore constamment

avec des débits de plus en plus élevés et un réseau de plus en plus réactif. Par contre, pour les opérateurs qui doivent la déployer, la 4G introduit un certain nombre de ruptures technologiques.

De nouvelles bandes de fréquences dans le patrimoine des opérateurs mobiles.

Pour proposer des services mobiles, les opérateurs disposent de plusieurs bandes de fréquences :

- **La bande 2100 MHz** : elle a été attribuée afin de pouvoir être utilisée par des infrastructures 3G. Elle est découpée pour cela en bandes (« porteuses ») de 5 MHz. Pour chacun de leurs sites, les opérateurs peuvent utiliser une ou plusieurs porteuses de 5 MHz selon l'importance du trafic à écouler. Dans la pratique, avec la croissance des usages liés aux smartphones et tablettes, il n'est pas rare que des sites radio 3G utilisent la totalité des porteuses présentes dans le patrimoine des opérateurs. Elle est donc considérée comme saturée et ne permettra pas, avant longtemps, de proposer des services 4G.
- **La bande 900 MHz** : elle a été attribuée à l'origine pour y déployer des infrastructures 2G. Depuis, sa réglementation a évolué afin de pouvoir être utilisée par des infrastructures 3G. Du fait d'une meilleure propagation du signal radio que pour la bande 2100 MHz, elle peut être utilisée pour déployer des cellules de grande taille en zone rurale et couvrir efficacement l'intérieur des bâtiments (« *deep indoor* ») en zones urbaines. Afin d'accueillir la 3G, il est donc nécessaire de libérer 5 MHz dans cette bande de fréquences. Cela est rendu possible grâce à la généralisation des terminaux 3G qui privilégient le mode 3G 2100 MHz par rapport au mode 2G 900 MHz lorsque la double couverture existe. Pour autant, la bande 900 MHz continue d'être utilisée. Il y a encore un parc de terminaux exclusivement 2G toujours en service, des objets communicants hors de portée de la couverture 2100 MHz (ex : intérieur profond des bâtiments), etc. Pouvoir libérer 5 MHz n'est pas toujours facile notamment dans les zones les plus denses en population (le centre des grandes agglomérations) et plus particulièrement pour les opérateurs dont la base client est plus importante que d'autres.

De surcroît, les trois opérateurs « historiques » ont été tenus de rétrocéder une partie de leur spectre 900 MHz afin d'inclure 5 MHz de cette bande dans la 4^{ème} autorisation 3G. Bien que, pour chacun de ces opérateurs, la rétrocession ne se fasse pas sur la totalité du territoire, cela augmente à nouveau la difficulté pour réserver une porteuse de 5 MHz sur la bande 900 MHz.

Cette bande ne permettra donc pas non plus de proposer des services 4G avant longtemps.

- **La nouvelle bande 2600 MHz pour la 4G** : cette bande de 2 x 70 MHz (70 MHz pour le sens montant du terminal vers le réseau et 70 MHz pour le sens descendant) a été découpée en France en deux blocs de 2 x 15 MHz et deux blocs de 2 x 20 MHz. Ceux-ci ont été attribués aux opérateurs fin 2011 après enchères. La largeur de la bande de fréquences utilisée conditionnant le débit du LTE, « *toutes choses égales par ailleurs* », les services 4G proposés sur des blocs de 2 x 20 MHz auront un débit supérieur à ceux proposés sur des blocs de 2 x 15 MHz. Toutefois, les versions ultérieures du LTE donneront la possibilité d'agréger des bandes de fréquences entre elles. Des blocs de 2 x 20 MHz voire plus pourront être constitués en associant la bande 2600 MHz avec une ou plusieurs autres bandes de fréquences.

Parmi les contraintes qui pèsent sur cette bande de fréquences, on notera :

- Sa situation plus élevée dans le spectre de fréquences qui induit une portée plus faible de quelques pourcents par rapport à celle de la bande 2100 MHz : comme pour la 2G et la 3G, il est nécessaire de l'associer à une bande basse pour couvrir les zones rurales et l'intérieur des bâtiments en zones urbaines.
- Sa disponibilité progressive d'ici mars 2014 selon les régions françaises.
- Son indisponibilité sur certaines zones de grandes villes en raison d'un risque de brouillage lié à la proximité de radars fonctionnant dans la bande 2700-2900 MHz (aviation civile, météo nationale, défense nationale). La mise à jour progressive de ces radars est prévue jusqu'en 2015.

Dans l'optique d'un déploiement rapide de la 4G, ces contraintes de brouillage redonnent de l'intérêt à la bande 1800 MHz qui présente par ailleurs un gain conséquent en terme de couverture.

- **Une autre nouvelle bande pour la 4G, le 800 MHz** : cette bande de fréquences est issue du dividende numérique qui est le spectre libéré par l'extinction de la télévision analogique. Une partie de ce dividende numérique a été réaffectée pour des services mobiles. Pour la France, une bande de 2 x 30 MHz a été découpée en deux blocs de 2 x 10 MHz et deux blocs de 2 x 5 MHz suite au remplacement de la télévision analogique par la télévision numérique de terre moins consommatrice en fréquences. Ces blocs ont également été attribués après enchères début 2012. Trois des quatre opérateurs français ont reçu chacun un lot de 2 x 10 MHz sachant que le quatrième opérateur pourra demander l'itinérance à l'un des trois détenteurs.

Sa situation relativement basse dans le spectre des fréquences rend cette bande attractive pour couvrir de vastes zones rurales à faible densité de population ainsi que l'intérieur des bâtiments dans les zones urbaines. Par contre, l'étroitesse de chaque lot (2 x 10 MHz) fait qu'il ne sera pas possible de proposer les débits les plus élevés de la 4G ; sauf si les opérateurs s'associent pour mettre en commun leurs fréquences au sein d'un partage d'infrastructures réseau.

Comme la bande 2600 MHz, la bande 800 MHz doit composer avec la proximité d'autres technologies. En l'occurrence, il s'agit ici de la TNT qui occupe le spectre de fréquences juste en dessous. Concrètement, il est à craindre que les émissions sur la bande 800 MHz brouillent la réception TNT. La solution passe alors par l'insertion d'un filtre derrière l'antenne TV dont les modalités de financement et de déploiement font actuellement débat.

- **La bande 1800 MHz** : comme la bande 900 MHz, cette bande a été attribuée à l'origine pour la 2G avant d'être autorisée pour la 3G. Elle aussi est progressivement libérée au fur et à mesure de la généralisation des terminaux 3G qui privilégient le mode 3G 2100 MHz par rapport au mode 2G 1800 MHz lorsque la double couverture existe.

Bien que des terminaux exclusivement 2G soient toujours en service, dans la mesure où elle n'a pas été utilisée par la 3G, il est donc intéressant de chercher à libérer au moins 2 x 10 MHz pour y déployer des services 4G. D'une part, cette bande de fréquences permet une meilleure couverture que la bande 2600 MHz. D'autre part, les antennes compatibles 1800 MHz ont déjà été déployées à l'occasion de la 2G. Enfin, la rénovation des réseaux 2G a permis de déployer des stations de base multistandard pour lesquelles l'adjonction de la 4G se réduit à un simple ajout de cartes électroniques. Cette stratégie nécessite alors de diminuer la capacité allouée au réseau 2G 1800 MHz. Comme pour la

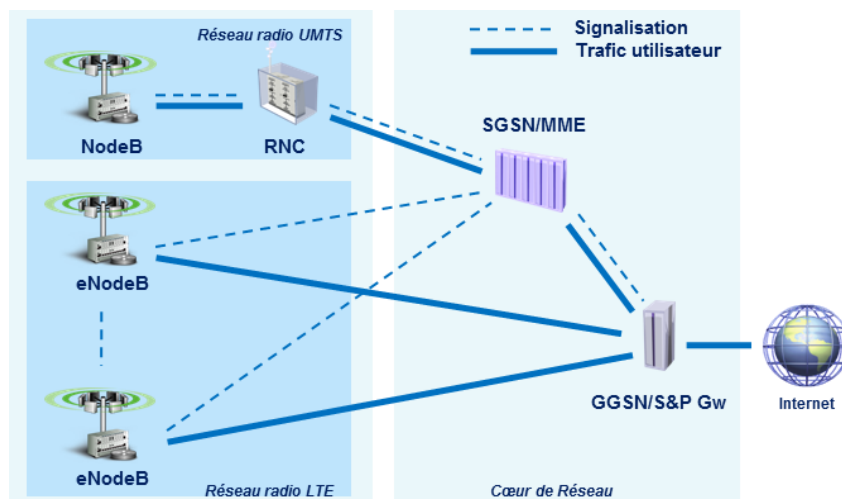
bande 900 MHz, cela est plus difficile dans les zones les plus denses en population pour les opérateurs ayant une base client importante. Réallouer 10 MHz à la 4G pourrait conduire à saturer leur réseau 2G sur cette bande de fréquences. Il faut par ailleurs réviser l'autorisation d'utiliser cette bande pour permettre d'accueillir des infrastructures 4G. Ceci se fait à la suite d'un réexamen de l'usage de la bande qui peut conduire à un « *refarming* » consistant à réallouer équitablement le spectre entre tous les opérateurs.

Déployer un réseau sur deux nouvelles bandes de fréquences (2600 MHz et 800 MHz) nécessitera de déployer sur chaque site radio des nouvelles antennes dans la mesure où celles utilisées jusqu'à présent ne couvrent pas ce spectre. Par ailleurs, une coordination technique doit être réalisée site par site afin de gérer, avec tous les acteurs concernés, les risques de brouillage avec les technologies (radar et TNT) déjà déployées sur les bandes de fréquences adjacentes. C'est dans ce contexte que des solutions alternatives sur 1800 MHz sont étudiées actuellement bien que n'étant pas répliquables tout de suite par tous les opérateurs.

Une architecture de réseau simplifiée

Un réseau de télécommunication mobile se décompose en deux parties principales :

- Un « réseau d'accès » comprenant l'ensemble des sites radio sur lesquels sont déployés les stations de base ainsi que la transmission vers les sites de gestion des ressources radio.
- Un « cœur de réseau » comprenant l'ensemble des équipements de gestion des communications, du suivi de la mobilité des utilisateurs, du transport des trafics voix et data et de l'interconnexion avec les autres opérateurs.



En 2G et en 3G, les stations de base (BTS en 2G et NodeB en 3G) sont reliées à un contrôleur de stations de base (BSC/PCU en 2G et RNC en 3G) chargé de répartir les ressources radio entre les différents terminaux présents dans les cellules dont il a la charge. En 4G, les stations de base (eNodeB) sont suffisamment évoluées pour gérer entre elles la répartition des ressources radio : ***moins d'équipements traversés c'est moins de latence donc un réseau plus rapide notamment pour le trafic vers Internet.***

En 2G et en 3G, le SGSN est un équipement qui gère l'utilisateur lors de ses échanges de données : ses droits d'accès au réseau, la QoS appropriée pour le service demandé et sa mobilité. Par contre, il n'intervient pas sur le trafic Internet mobile de l'utilisateur et ne fait que réémettre vers le GGSN les données reçues du réseau d'accès radio. En 4G, les fonctions du SGSN sont reprises en grande partie par le MME tandis que les fonctions du GGSN sont reprises en grande partie par le S&P GW. Cette fois, le trafic Internet mobile de l'utilisateur ne passe plus par le MME. Les données sont transférées directement du eNodeB vers le S&P GW : ***dans le cœur de réseau aussi, moins d'équipements traversés c'est moins de latence donc un réseau plus rapide.***

En 2G et en 3G, les appels voix utilisent le mode connecté dans lequel un circuit est établi jusqu'au destinataire de l'appel. Un autre équipement du cœur de réseau est utilisé, le MSC. En 4G, le cœur de réseau circuit disparaît et les communications téléphoniques sont réalisées par une application utilisant le protocole IP au même titre que l'accès à Internet : ***un réseau dans lequel les ressources ne doivent plus être partagées entre un cœur circuit pour la téléphonie et un cœur paquet pour l'Internet mobile c'est un réseau qui peut optimiser la qualité de service fournie à l'utilisateur.***

Grâce aux performances des réseaux 3G-HSPA et à l'ergonomie apportée par les smartphones (grand écran tactile, processeur multicœur, browser web, ...) la consommation d'Internet mobile a fortement augmenté avec des coûts à la clé : pour l'opérateur, le réseau sature et il doit investir pour en augmenter la capacité. L'utilisateur, lui, atteint plus vite son *fair use* et doit migrer vers l'offre supérieure pour conserver son confort d'usage. Pour ces raisons, la généralisation du WiFi est une double opportunité : l'opérateur décharge une partie du trafic de son réseau (on parle alors d'*offload WiFi*) et l'utilisateur peut consommer de l'Internet mobile hors forfait. En 2G et en 3G, la mobilité vers le WiFi en cours de communication ne peut être assurée qu'au prix d'une ingénierie laborieuse. De fait, elle n'est pas mise en œuvre. En 4G, le cœur de réseau est dit « *access agnostic* » c'est-à-dire capable de piloter tous les types de réseaux d'accès : 2G, 3G, 4G, WiFi voire même CDMA2000 pour les opérateurs (US, Asie, Europe de l'Est, ...) qui utilisaient cette norme et souhaitent migrer vers le vaste écosystème LTE. Ainsi, la mobilité devient totalement transparente. On peut changer de cellule, de réseau (2G/3G vers 4G) ou de technologie (4G vers WiFi) en restant en communication et sans avoir besoin de se réauthentifier. ***A tout instant, on surfe sans le savoir sur le meilleur réseau disponible : le WiFi de la box à la maison, la 4G dans la rue, le WiFi communautaire chez des amis, la 3G dans les zones non encore couvertes en 4G, le hotspot WiFi à l'hôtel, etc.***

L'avènement du protocole IP dans les réseaux mobiles

L'utilisation du protocole IP a été très progressive dans les réseaux mobiles :

- **Fin des années 90, le GSM data** : à l'origine, le GSM ne servait qu'à téléphoner et envoyer des SMS. Par la suite, la technologie CSD (Circuit Switched Data) a été utilisée pour transporter des données via le protocole IP à un débit de 9,6 Kbps. Cela a permis d'accéder, par exemple, à des portails WAP pour consulter des services de type météo, infos, résultats sportifs, etc.
- **Début des années 2000, le GPRS puis l'EDGE** : avec le GPRS, un cœur de réseau paquet est déployé par les opérateurs. Les appels voix et les SMS continuent d'être gérés

par un cœur circuit utilisant un réseau TDM tandis que les communications de données sont gérées par un cœur paquet utilisant un réseau IP.

- **Quelques années après, le NGN** : la première génération d'équipements du cœur circuit (les MSC) arrive en fin de vie. La nouvelle génération de MSC suit l'architecture NGN qui distingue clairement les fonctions de transport, de contrôle et de service. Les nouveaux MSC sont décomposés en deux équipements : la media gateway (MGW) qui transporte les flux voix encapsulés dans le protocole IP et le MSC Server (MSC-S) qui gère les communications grâce à des nouveaux protocoles de signalisation transportés eux-aussi sur IP. Cette nouvelle architecture a permis d'optimiser le dimensionnement des équipements en distinguant, par exemple, la charge liée à la répartition géographique des clients de celle liée à leurs usages.
- **Fin des années 2000, le protocole IP gagne l'interconnexion entre le cœur de réseau et le réseau d'accès radio** : un autre avantage du protocole IP est de pouvoir mettre les MSC et SGSN *en pool*. Cela signifie que le contrôleur de station de base (BSC/PCU en 2G, RNC en 3G) peut être connecté à plusieurs MGW/MSC-S et plusieurs SGSN au lieu d'un de chaque auparavant. L'avantage est de pouvoir répartir plus facilement la charge en fonction des évolutions de trafic et d'avoir une redondance automatisée en cas de panne d'un équipement du cœur de réseau.
- **Au même moment, le protocole IP gagne aussi l'interconnexion entre les opérateurs** : progressivement, tous les opérateurs ont déployé l'architecture NGN pour gérer les communications voix. Cependant, l'interconnexion entre opérateurs étant restée en TDM, un double transcodage de la voix était réalisé lors d'un appel téléphonique entre deux interlocuteurs appartenant à deux opérateurs différents : un transcodage IP-NGN vers TDM par l'opérateur A puis un transcodage TDM vers IP-NGN par l'opérateur B. Le transcodage entraînant de la latence et un risque de dégradation de la QoS voix, une interconnexion IP était donc plus pertinente.
- **Début des années 2010, le protocole IP va jusqu'aux stations de base** : avec l'arrivée du HSDPA et l'augmentation des usages, le débit que doit écouler chaque site radio augmente fortement. Dès lors, il n'est plus économiquement pertinent de relier les sites par une agrégation de liens à 2 Mbps. Les stations de base se dotent alors d'interfaces fast-ethernet transportant les flux encapsulés dans le protocole IP. Le réseau de transmission reliant les sites radio au cœur de réseau venait, lui, d'évoluer en ATM avec l'arrivée de la 3G. Mais l'ATM n'est plus pérenne car non retenu par la normalisation LTE. Le réseau de transmission doit donc à nouveau évoluer vers un ethernet « carrier-grade » transportant les flux IP grâce à de nouveaux équipements dits PTN (Packet Transport Network). Toute l'infrastructure réseau est désormais en IP et seuls les communications téléphoniques et les SMS continuent à utiliser le mode circuit.
- **Avec l'arrivée de la 4G, c'est la fin d'une évolution qui aura pris près de 15 ans. L'ensemble du réseau et des services utilise maintenant le protocole IP** : avec le LTE, le réseau devient maintenant 100% IP. Téléphoner consiste à utiliser une application IP sur son smartphone. Le cœur circuit disparaît et tous les services sont contrôlés par le cœur paquet.

De nouveaux terminaux

Bien entendu, accéder à un réseau 4G nécessitera de changer son terminal. Toutefois, l'attrait des nouvelles fonctionnalités et les effets de mode font que, depuis bien longtemps, le changement pour le dernier modèle de terminal n'attend plus les ruptures technologiques.

Début février 2013, l'association GSA (Global mobile Suppliers Association) dénombrait plus de 660 terminaux LTE dans le monde (tous ne fonctionnent pas sur les bandes de fréquences disponibles en France). Des smartphones, des tablettes, des clés USB mais aussi des routeurs qui se connectent à un réseau 4G et redistribuent les flux en WiFi permettant ainsi à plusieurs PC portables d'accéder à Internet grâce à un hotspot mobile.

Pour les opérateurs, la 4G introduit un certain nombre de ruptures technologiques :

- Des nouvelles bandes de fréquences impliquent de déployer des nouvelles antennes sur chaque site radio.
- Des nouvelles stations de base doivent être déployées sur chaque site radio mais elles peuvent maintenant être mutualisées avec des stations de base 2G et 3G (*baies multistandards* parfois déjà déployées dans le cadre d'un programme de gestion de l'obsolescence).
- La généralisation de l'IP dans tout le réseau nécessite la refonte totale du réseau de transmission entre les sites radio et le cœur de réseau. Par ailleurs, le protocole IP étant beaucoup plus répandu que ses prédécesseurs, cela entraîne également une refonte de la politique de sécurité télécom.
- Le mode circuit disparaît. La téléphonie devient une application IP comme les autres. Cette rupture existait déjà dans le fixe mais cela sera plus sensible pour le mobile où la QoS est très dépendante des conditions radio...

MYTHE N° 3 :

LA 4G, C'EST LA FIN DE LA 3G !

Faux ! Il se passera de longues années avant que les réseaux 3G soient démontés !

De même que la 3G n'a pas mis fin à la 2G, la 4G coexistera avec les générations précédentes de télécommunications mobiles.

Il faut plusieurs années pour couvrir la totalité du territoire avec une nouvelle technologie.

La couverture actuelle de plus de 99% de la population *outdoor* (extérieur des bâtiments) en 2G et/ou 3G a nécessité de déployer entre 15 et 20000 sites radio. N'étant pas concevable que des utilisateurs ne soient soudainement plus couverts, l'extinction de la 2G et de la 3G ne pourra se faire que si la 4G dispose d'une couverture au moins équivalente. Etant donné les opérations techniques à réaliser sur chaque site radio (installer de nouvelles stations de base, déployer de nouvelles antennes lorsque les bandes 800 ou 2600 MHz sont utilisées, trouver un site de remplacement si le site ne peut évoluer, etc.), on peut estimer qu'ouvrir commercialement la 4G sur plus de 2000 sites radio par an est une belle performance. A ce rythme, la couverture de 99% de la population prendra tout de même au moins 8 ans pendant lesquels la 3G et la 4G coexisteront. D'ailleurs, les obligations de couverture associées à l'autorisation 4G sur la bande 2600 MHz confirment de tels délais : 4 ans maximum pour couvrir 25% de la population (une estimation d'environ 2500 sites radio), 8 ans maximum pour couvrir 60% de la population

(une estimation d'environ 6500 sites radio) et 12 ans maximum pour couvrir 75% de la population (une estimation d'environ 8500 sites radio). La couverture restant toutefois un facteur indispensable dans l'acte d'achat, il va de soi que l'intensité concurrentielle contribuera à réduire plus ou moins rapidement ces délais.

Il faut plusieurs années pour que les terminaux de nouvelle technologie se généralisent dans un parc client.

La première étape dans la diffusion des terminaux est la disponibilité des chipsets 4G en provenance de fournisseurs de circuits intégrés télécom (Qualcomm, ST Ericsson, ...). Ceux-ci intègrent progressivement les différentes bandes de fréquences qui seront utilisées par la 4G ainsi que les fonctionnalités au fur et à mesure de leur normalisation (ex : LTE 100 Mbps puis 150 Mbps, support de la voix, agrégation de porteuses, etc.).

Dans un second temps, les fabricants de terminaux diffusent progressivement leurs produits. Les premiers à apparaître sont généralement les modems sous forme de clés USB et les routeurs Cellulaires/WiFi. Etant destinés à être connectés à un ordinateur, ils sont plus simples à concevoir (pas d'écran, pas de gestion de la voix, moins de fonctionnalités que dans un téléphone donc miniaturisation moins poussée, ...). Les tout premiers ne permettent de se connecter qu'à un réseau 4G puis les suivants intègrent aussi la 2G et la 3G afin que l'utilisateur soit connecté en permanence même dans les zones où la couverture 4G n'est pas encore déployée. La 4G arrive alors sur les smartphones et tablettes haut de gamme avec un surcôt pouvant être de plusieurs dizaines d'euros par rapport aux modèles uniquement 2G/3G. Progressivement, elle se diffuse ensuite dans les terminaux moyen puis entrée de gamme. De quelques références au début, la technologie finit par se généraliser sur tous les nouveaux modèles au bout de 3 à 4 ans.

Pour que les utilisateurs ne choisissent plus des produits exclusivement 2G/3G et adoptent massivement la 4G, celle-ci doit être disponible sur une gamme suffisamment large de produits attractifs et répondant aux besoins et au budget de chacun : des tablettes, des modems USB, des routeurs cellulaire/WiFi, des smartphones *high tech*, des terminaux simples d'utilisation sans oublier aussi des produits *low cost*.

L'expérience montre qu'il faut près de 3 à 4 ans après l'introduction d'une nouvelle technologie pour que les ventes de terminaux de nouvelle génération dépassent celles de l'ancienne génération. En d'autres termes, sans incitation commerciale particulière, pendant encore 3 à 4 ans, il se vendra plus de terminaux 3G que de terminaux 4G. Pour autant, au sein du parc total de clients, les terminaux exclusivement 2G/3G resteront majoritaires. Le remplacement total du parc par des terminaux 4G prendra plusieurs années supplémentaires. D'ailleurs, il y a toujours des clients aujourd'hui qui ne souhaitent que téléphoner et qui ne changeront leur terminal exclusivement 2G que lorsque la batterie sera hors d'usage... Tout cela sans compter la grande quantité de modules exclusivement 2G embarqués dans des objets communicants (ex : terminaux de paiement électronique) qui ne disparaîtront qu'avec l'objet en question.

Dans un premier temps, les communications voix vont continuer à emprunter les réseaux 2G et 3G.

L'organisme 3GPP, qui regroupe entre autres des opérateurs et constructeurs télécoms, a normalisé le LTE comme une technologie de transport de données. Dès lors, elle n'intègre pas nativement le support des services circuit comme les communications téléphoniques et les SMS. Passer un appel téléphonique sur un réseau 4G consiste alors à utiliser une des

applications VoIP déjà largement répandues sur Internet. Il faut cependant assurer l'interopérabilité de cette application VoIP avec les réseaux de téléphonie déjà déployés par les opérateurs (2G, 3G, fixe, ...). C'est cette fois la GSM Association, l'association mondiale des opérateurs mobiles, qui a publié des recommandations visant à standardiser le service VoLTE (Voice over LTE).

VoLTE s'appuie sur un réseau de transport IP. Les communications sont gérées par un protocole de signalisation lui aussi sur IP : SIP. Les appels voix peuvent être enrichis avec d'autres services comme des indicateurs de présence liés à un agenda en ligne (disponible pour un appel, disponible uniquement par messagerie instantanée, en rendez-vous, ...), le partage d'écran ou de documents, l'enchaînement avec d'autres applications, etc. Pour cela, il faut alors déployer dans le cœur de réseau un système de gestion IMS. Ceux qui ont déjà été déployés dans le monde du fixe peuvent être réutilisés mais doivent être mis à jour pour intégrer les fonctionnalités spécifiques du monde mobile.

Mais il ne suffit pas d'avoir une application VoIP sur son smartphone et un cœur de réseau IMS pour que tout fonctionne parfaitement. Un appel téléphonique ne s'accommode pas du mode *best effort* utilisé largement sur Internet. L'ensemble du réseau (radio, transmission, cœur) doit donc intégrer des fonctionnalités de qualité de service notamment lorsque, pendant la communication, on sort d'une couverture 4G pour passer sous une couverture 2G/3G. Le temps que ces fonctionnalités soient déployées, un autre mécanisme appelé *CSFB* (pour *Circuit Switched FallBack*) est utilisé. Celui-ci consiste à basculer la communication sur le réseau 2G ou 3G lors de l'appel. A la fin de la communication, le terminal retourne ensuite automatiquement sur le réseau 4G.

Même si elle ne la rattrapera pas, la 3G continue d'évoluer en bénéficiant des progrès techniques de la 4G.

En parallèle du développement de la 4G, le 3GPP continue à normaliser les évolutions de la 3G en réutilisant les innovations conçues initialement pour la 4G :

- **De nouvelles modulations** : jusqu'à la version 6, le HSPA culmine à 14,4 Mbps dans le sens descendant (du réseau vers le terminal) et 5,7 Mbps dans le sens montant. A partir de la version 7, la modulation 64 QAM remplace la modulation 16 QAM pour les communications descendantes et la modulation 16 QAM remplace la modulation QPSK pour les communications montantes. Le débit maximum théorique de la 3G évolue donc en conséquence jusqu'à 21 Mbps dans le sens descendant et 12 Mbps dans le sens montant.
- **La multiplication des antennes en émission et en réception (MIMO 2x2)** : à partir de la version 7, l'utilisation du MIMO 2x2 dans un terminal à modulation 16 QAM devait permettre de doubler les débits avec un maximum théorique de 28 Mbps. Cette évolution est dans la pratique peu considérée en raison de son coût par rapport aux installations déjà déployées (nécessité d'ajouter sur chaque station de base un deuxième amplificateur de puissance par cellule) et des doutes sur ses performances dès que les conditions radio se dégradent.
- **L'agrégation de porteuses** : jusqu'à la version 7, le HSDPA n'utilise qu'une seule porteuse de 5 MHz. Le débit atteint un maximum théorique de 21 Mbps dans le sens descendant comme nous venons de le voir. A partir de la version 8, le HSDPA intègre la fonctionnalité Dual Carrier qui permet à un terminal d'utiliser simultanément deux porteuses de 5 MHz. Le débit dans le sens descendant s'en trouve doublé avec un

maximum théorique pouvant aller jusqu'à 42 Mbps. Sur le plan opérationnel, dès qu'une cellule à une porteuse arrive à saturation, on « allume » une deuxième porteuse et on active simultanément la fonctionnalité Dual Carrier. Si, à un instant t , un utilisateur est seul à communiquer dans la cellule, il utilisera les deux porteuses (sous réserve qu'il dispose d'un terminal Dual Carrier). Par contre, si deux utilisateurs communiquent simultanément, les ressources radio restent réparties entre eux.

- **L'agrégation des porteuses provenant de bandes de fréquences différentes (Dual Band) :** à partir de la version 9, il est possible d'utiliser la fonctionnalité Dual Carrier sur deux bandes de fréquences différentes (en l'occurrence le 900 et le 2100 MHz). Cela permet également de doubler les débits mais, bien entendu, uniquement sur le périmètre où il y a la double couverture (c'est-à-dire uniquement sur le périmètre de la bande de fréquences 2100 MHz).
- **Des terminaux plus performants dotés de système de réception évolué (normalisé sous l'appellation « type 3i ») :**
 - ✓ Deux antennes internes afin de combiner les signaux reçus et améliorer les performances en réception (principe appelé *diversité*).
 - ✓ Un système d'égalisation amélioré G-RAKE2 (élément permettant de retrouver le signal émis à partir du signal reçu).
 - ✓ Un système de suppression des interférences provenant des cellules voisines (*interference cancellation*)
- **Un réseau de transmission 3G qui utilisait l'ATM et migre désormais vers de l'IP sur de l'ethernet carrier-grade :** la 3G a nécessité le déploiement d'un réseau de transmission ATM. La 4G, elle, impose de raccorder les sites radio en IP. Deux réseaux de transmission doivent donc coexister en parallèle dont un qui n'est plus pérenne. Dans la mesure où la croissance actuelle des usages 3G nécessite une augmentation de capacité du réseau, il est plus pertinent d'investir sur le réseau IP plutôt que sur le réseau ATM. Dès lors, au fur et à mesure des besoins de capacité, les stations de base 3G se dotent, elles-aussi, d'interfaces IP sur ethernet carrier-grade.

Comme on le voit, malgré l'arrivée de la 4G, les investissements 3G vont se poursuivre le temps que la couverture 4G soit suffisante et que la majorité des utilisateurs adopte les nouveaux terminaux. Ce n'est qu'à ce moment que le trafic migrera massivement de la 3G vers la 4G et que les investissements en capacité 3G ne seront plus nécessaires. De surcroît, le démontage d'un réseau devra prendre en compte les impacts comptables (ex : accélération d'amortissement) des derniers investissements réalisés dans la mesure où la durée d'amortissement des équipements est relativement longue (de l'ordre de 8 ans pour les équipements radio et de l'ordre de 5 ans pour les équipements de cœur de réseau). Une des solutions proposées par les fournisseurs consiste à investir dans des équipements multistandards (à la fois 2G, 3G et 4G) qui permettent d'éteindre à moindre coût une ancienne génération en réutilisant les équipements pour la nouvelle génération.

MYTHE N° 4 : APRES LA 4G, IL Y A LA 5G !

L'avenir nous dira ce que recouvrera le terme 5G ! Aujourd'hui, la 5G n'est pas mentionnée par le groupement Radiocommunication de l'Union Internationale des Télécommunications qui s'emploie plutôt à améliorer les technologies 3G (spécifications IMT-2000) et 4G (spécifications IMT-Advanced) existantes. Ces améliorations concerneront à la fois les performances vues des utilisateurs (plus de débit, plus de réactivité du réseau, ...) mais aussi l'exploitabilité du réseau pour les opérateurs (réseaux adaptatifs : SON pour Self Organized Network, réseaux de cellules hétérogènes : HetNet, etc.).

Concernant les technologies HSPA et LTE, c'est l'organisme 3GPP qui coordonne la normalisation de ces évolutions techniques. Celles-ci se font de manière progressive au sein de différentes versions technologiques appelées *releases*. La plupart des réseaux ouverts commercialement aujourd'hui suivent les spécifications de la release 8.

Les évolutions du HSPA+, une manière de maintenir la valorisation des spectres 2100 et 900 MHz

La release 8 : la version actuelle du HSPA+ permet d'avoir des débits descendants jusqu'à un maximum théorique de 42 Mbps grâce à la modulation 64 QAM et en utilisant deux porteuses de 5 MHz simultanément. La norme prévoit par ailleurs d'atteindre le même débit en gardant une seule porteuse de 5 MHz mais en doublant les antennes en émission et en réception (MIMO 2x2).

La release 9 : elle envisage de doubler les débits descendants jusqu'à un maximum théorique de 84 Mbps en utilisant conjointement la modulation 64 QAM, l'agrégation de deux porteuses de 5 MHz et le doublement des antennes en émission et en réception (MIMO 2x2).

La release 10 : elle propose de doubler à nouveau les débits descendants jusqu'à un maximum théorique de 168 Mbps en agrégeant non plus deux mais quatre porteuses de 5 MHz. Ces porteuses peuvent être prises dans la même bande de fréquences (2100 MHz) ou dans des bandes de fréquences différentes (2100 + 900 MHz par exemple) pour les opérateurs ne disposant pas de 20 MHz dans la bande 2100 MHz).

La release 11 : elle double une fois de plus les débits descendants jusqu'à un maximum théorique de 336 Mbps en agrégeant cette fois 8 porteuses de 5 MHz au lieu de 4. Une autre option consiste à proposer de doubler à nouveau les antennes en émission et en réception (MIMO 4x4). Une autre fonctionnalité, appelée *Multiflow HSPA*, permet cette fois d'augmenter le débit en limite de cellules, là où les conditions radio moins bonnes provoquent le plus de pertes de paquets. Le principe consiste à tirer parti du fait que les cellules radio se recouvrent afin de permettre au terminal de basculer de l'une à l'autre (notion de *handover*). Sur cette zone de recouvrement, les données sont réparties sur les deux stations de base et envoyées simultanément au terminal qui voit ainsi son débit doubler. L'opérateur tire également parti de cette fonctionnalité puisque cela lui permet de répartir le trafic sur un plus grand nombre de cellules.

Bien qu'étant une technologie 3G à l'origine, le HSPA va peu à peu intégrer des fonctionnalités lui permettant de remplir les critères de performance 4G définis par l'UIT-R.

Les évolutions du LTE :

La release 8 : sur une canalisation de 20 MHz, la version actuelle du LTE permet d'avoir des débits descendants jusqu'à un maximum théorique de 172 Mbps avec un doublement des antennes en émission et en réception (MIMO 2x2) et 326 Mbps avec un quadruplement de ces mêmes antennes (MIMO 4x4).

La release 9 : cette version introduit la fonctionnalité MBMS (Multimedia Broadcast Multicast Service) permettant de diffuser des contenus en multicast ou en broadcast. Déjà existante en 3G, cette fonctionnalité n'a jusqu'à présent pas réussi à consolider un écosystème autour d'elle. A l'instar du DVB-H d'ailleurs...

La release 10 : c'est avec cette version qu'apparaît le LTE Advanced qui, dès l'origine, devait répondre aux critères 4G définis par l'UIT-R. Les débits descendants peuvent aller jusqu'à un maximum théorique de 1 Gbps grâce à l'utilisation, cette fois, de 8 antennes en émission et en réception (MIMO 8x8) et une canalisation de 40 MHz (mais au sein d'une même bande de fréquences). Par ailleurs, la release 10 introduit la notion de *réseau hétérogène* (HetNet) à savoir la possibilité de gérer des besoins locaux de capacité en déployant des micro-cellules utilisant les mêmes fréquences que les cellules macro.

La release 11 : elle apporte quelques améliorations très utiles comme l'agrégation de porteuses mais provenant cette fois de bandes de fréquences distinctes. C'est cette version qui permettra réellement d'obtenir une canalisation de 40 MHz dans la mesure où les blocs de fréquences alloués aux opérateurs en France ne dépassent pas 20 MHz (par sens de transmission). Une autre fonctionnalité appelée CoMP (Coordination MultiPoint transmit/receive) permet d'améliorer les débits en limite de cellule grâce à différentes méthodes de coordination des stations de base visant à limiter les interférences entre elles. Cette technique nécessite toutefois de séparer en deux les stations de base en centralisant la partie contrôle (étage *bande de base*) et en gardant sur le site radio uniquement la partie transmission/réception (étage *radiofréquence*). Les deux parties de la station de base sont alors reliées par une fibre optique.

Progressivement, les technologies radio atteignent des débits et embarquent des fonctionnalités qui vont imposer de fibrer chaque site radio. Le déploiement de la fibre optique sur tout le territoire devient alors un enjeu qui dépasse le seul contexte du remplacement de la paire de cuivre dans les foyers.

Quelques-uns des enjeux des prochaines générations de télécommunications mobiles :

Plus de fréquences : pour répondre aux besoins de débit et de capacité engendrés par la croissance des usages notamment la vidéo délinéarisée, il sera nécessaire de dégager de plus en plus de spectre pour les télécommunications mobiles. D'ores et déjà, un deuxième dividende numérique a été identifié dans la bande 694 – 790 MHz lors de la dernière Conférence Mondiale des Radiocommunications organisée par l'UIT-R en février 2012. D'autres bandes

de fréquences sont en cours d'étude et seront à l'ordre du jour de la prochaine Conférence en 2015. L'identification de nouvelles bandes de fréquences est un processus long qui doit donc être pris en compte très en amont des besoins. Il a fallu 8 ans pour parvenir à l'attribution aux enchères du premier dividende numérique. Le second dividende numérique ne sera probablement disponible que vers la fin de la décennie. D'ici là, la TNT devra être replanifiée sur un spectre plus restreint et adopter conjointement un certain nombre d'évolutions technologiques (codec HEVC, diffusion DVB-T2, ...).

Des terminaux qui vont devenir de vrais challenges technologiques pour les constructeurs : il ne s'agit pas de considérer uniquement la puissance de traitement du processeur d'un terminal. Après tout, on trouve maintenant couramment des terminaux quad-core... Il s'agit, ici, de résoudre simplement quelques problèmes physiques comme la possibilité d'installer 8 antennes sur un smartphone (pour l'implémentation du MIMO 8x8) ou trouver une technologie de batterie faisant face à de telles consommations d'énergie sans nécessiter un rechargement quotidien.

De la fibre optique sur tout le territoire : avec l'évolution des débits, le raccordement des sites radio par faisceaux hertziens va atteindre ses limites car eux-mêmes auront besoin d'une bande de fréquences de plus en plus large. La fibre optique devra donc traverser tout le territoire y compris jusqu'aux sites radio les plus ruraux.

MYTHES ET LEGENDES DU PEER-TO-PEER

Francesca Musiani, CSI, MINES ParisTech

INTRODUCTION

Le *peer-to-peer* (P2P, « pair-à-pair » en français) est devenu l'un des termes les plus largement discutés dans le domaine des technologies de l'information et de la communication. Il se réfère à la notion que dans un réseau d'égaux ou de pairs, à l'aide de systèmes de communication et d'échanges appropriés, deux ou plusieurs individus sont en mesure de collaborer spontanément, sans nécessairement avoir besoin de coordination centrale.

Depuis la fin des années 90, les technologies P2P ont fait l'objet d'une évolution très rapide. Le grand succès dont les applications de cette technologie ont bénéficié a certes été un catalyseur important de la créativité de leurs développeurs, et des perfectionnements de ces outils en termes d'efficacité, mais les évolutions du secteur ont largement été influencées par des contraintes politiques, économiques et juridiques, notamment les menaces de procès mises sur la table par certains des grands acteurs de l'industrie du contenu numérique. Trois générations technologiques se sont ainsi succédées, tandis que le modèle P2P commençait à être appliqué non plus seulement au partage de fichiers mais à une variété d'usages et d'applications, dévoilant la complexité de l'objet et la multiplicité de ses mobilisations. Ce chapitre rend compte de comment, autour et au moyen de ces mobilisations, des discours partiels ou réducteurs ont pris forme avec le P2P – discours qui cachent trop souvent les expérimentations socio-économiques à l'œuvre de par ou avec le P2P, et qui empêchent ou entravent un renouvellement du débat politique autour de ces systèmes. (*Achévé de rédiger en janvier 2012*).

MYTHE N° 1 :

LE PEER-TO-PEER, C'EST DU PIRATAGE

Depuis que, en 1999, la naissance de *Napster* leur a donné visibilité et diffusion auprès du grand public, les réseaux informatiques *P2P* ont été considérés presque exclusivement comme une menace pour l'industrie des contenus numériques. L'usage principal de ces réseaux par le public étant le partage non autorisé de fichiers musicaux ou vidéo, le problème du droit de propriété intellectuelle, du droit d'auteur notamment, s'est imposé en tant que cadrage médiatique et politique prédominant des réseaux P2P et leurs usages. Cependant, l'argument qui consiste, essentiellement, à assimiler P2P et piratage, présente plusieurs limites.

Avec des millions d'utilisateurs à l'échelle mondiale (le pionnier *Napster* comptait, l'année même de sa création, 60 millions d'utilisateurs « partageurs »), les réseaux P2P facilitent la distribution massive de copies parfaites et gratuites de contenus numériques. Il serait difficile de nier que cette capacité soit la raison principale à la base du succès universel de ces dispositifs ; pourtant, comme soulignent quelques auteurs interdisciplinaires entre le droit et l'informatique (par exemple Niva Elkin-Koren et Barbara van Schewick) leur signification politique et technique serait à chercher ailleurs, dans un ensemble de propriétés qui tiennent à la fois du technique, du social, de l'économie et du légal.

La capacité de ces systèmes à tirer avantage de leur architecture décentralisée peut donner lieu, en effet, à une meilleure efficacité économique, une plus grande liberté et à l'émergence de

nouveaux principes organisationnels et légaux, rendus possibles par l'échange direct de contenus entre les différents nœuds du réseau. Une architecture décentralisée peut augmenter le niveau de liberté personnelle car il devient plus facilement possible pour les utilisateurs de rester anonymes et de protéger leur *privacy* : il n'est plus nécessaire de s'enregistrer en tant qu'utilisateur d'un serveur particulier, mais il est possible de se « déplacer » entre réseaux *ad-hoc*. Par ailleurs, une plus grande protection de l'anonymat peut dans certains contextes être « libératrice », en ouvrant plus de possibilités de développement de différents aspects de son identité, rendant plus facile l'expression de préférences authentiques, et facilitant par conséquent la formation d'un environnement plus participatif pour tester de nouvelles idées. Outre que dans la nature et l'étendue des libertés personnelles, la décentralisation des réseaux informatiques peut faciliter les processus de décision alternatifs, en augmentant notamment la capacité à favoriser l'exclusion des intermédiaires, le tout dans des contextes qui incluent, certes, mais qui dépassent amplement, le partage de contenus numériques protégés : par exemple, la démocratie directe, le débat public et la recherche de consensus.

MYTHE N° 2 :

LE PEER-TO-PEER, C'EST DU PARTAGE DE FICHIERS

Tout comme le peer-to-peer est devenu la « technologie des pirates », il est souvent considéré comme la technologie « du partage de fichiers », le plus souvent protégés par le droit d'auteur. Cependant, cette technologie de réseau ne sert pas seulement au partage de fichiers : elle a certes été, au cours de ses premiers pas, reléguée à ce seul domaine, ce qui constitue l'option technique la plus facile et nécessitant un minimum de ressources humaines et techniques pour sa réalisation – mais le P2P est aussi exploité, et ce de plus en plus, pour des applications « alternatives » et « légales », qui peuvent servir plus d'une nécessité des usagers/consommateurs/citoyens d'aujourd'hui, et qui se proposent en tant qu'alternatives décentralisées à des services et instruments aujourd'hui fondamentaux de notre vie quotidienne : moteurs de recherche, réseaux sociaux, stockage de fichiers en ligne. Cela se doit non seulement aux évolutions technologiques à large échelle (qualité des connexions Internet, espace disque à disposition sur chaque ordinateur), mais aussi à la prise de conscience (soit par les chercheurs, soit par le public) de l'existence d'une « *écologie Internet* » de plus en plus délicate et articulée.

Avec *Google*, *Facebook* ou encore *Picasa*, à chaque fois qu'un usager exécute une recherche, échange un message avec quelqu'un ou met un album photo en ligne pour le montrer à ses amis, des données sont envoyées et téléchargées à des serveurs avant de rejoindre leur destinataire prévu, contribuant à constituer le scénario de « concentration » de contenus dont on a parlé ci-dessus. En revanche, mettant à profit le potentiel décentralisateur du P2P, ces autres applications récentes visent à répondre aux mêmes exigences du point de vue de l'utilisateur final (qui continuera donc à rechercher des mots, à former des réseaux d'amis et à partager des photos), mais en se basant sur une architecture technique différente. Ce qui a des implications à plusieurs niveaux : meilleures performances techniques, certes, mais aussi possibilité de reconsidérer des concepts tels que la sécurité et la *privacy*, en reconfigurant les emplacements des données et des échanges, les frontières entre l'utilisateur et le réseau, la prise en compte des outils qu'on a à disposition : en somme, l'attribution, reconnaissance et modification de droits entre utilisateurs et fournisseurs des services.

Parmi les exemples les plus intéressants de ces applications pionnières, on retrouve bien sûr les services de voix sur IP qui ont chamboulé le marché traditionnel de la téléphonie ; mais aussi des applications moins connues, de stockage et accès distribué de fichiers privés ; de

moteur de recherche P2P, qui se fonde sur la détection des préférences personnelles et des affinités entre les usagers ; de streaming vidéo ; de messagerie ; de réseautage social.

MYTHE N° 3 :

LE TRAFIC EN PEER-TO-PEER EST EN BAISSSE PUISQUE LE TRAFIC EN STREAMING AUGMENTE

Nombre d'études constatent dans les dernières deux ou trois années « une baisse des échanges peer-to-peer face à la montée du streaming », notamment vers les sites de streaming et de téléchargements direct, comme RapidShare ou Megaupload. Si cela est certes indicatif d'une tendance de certaines pratiques de consommation numérique à se déplacer vers d'autres arènes – notamment à cause de mesures juridiques visant des technologies plutôt que des usages – mettre en corrélation directe la baisse du P2P et la hausse « du streaming », comporte, encore une fois, des imprécisions et une confusion entre les usages et les infrastructures qui les supportent. En effet, si le streaming vidéo correspond dans l'imaginaire d'une très grande majorité d'utilisateurs à des solutions centralisées proposées par des grandes plateformes, de YouTube et DailyMotion à nombre de sites au statut légal plus douteux, le streaming vidéo en P2P est déjà largement utilisé, et plusieurs projets de recherche étudient actuellement les moyens d'améliorer sa qualité de service. Ce système est plus particulièrement à l'œuvre dans le domaine du P2PTV.

Les applications P2PTV sont destinées à redistribuer des flux (streams) vidéo en temps réel sur un réseau P2P ; les flux vidéo distribués sont généralement des chaînes de télévision, mais peuvent aussi provenir d'autres sources. Le potentiel de ces applications est qu'elles peuvent rendre toute chaîne de télévision disponible au niveau mondial par toute personne alimentant ce flux au sein du réseau ; chaque pair qui rejoint le réseau pour voir la vidéo contribue au visionnage des autres pairs/télespectateurs, permettant un passage à l'échelle du système au fur et à mesure que le public augmente, sans coût supplémentaire pour la source du flux vidéo.

Dans un système de P2PTV, chaque utilisateur, quand il télécharge un flux vidéo, est simultanément en train de permettre aux autres utilisateurs de télécharger ce flux, contribuant ainsi à la bande passante totale disponible. Les flux qui arrivent sont typiquement en retard de quelques minutes par rapport aux sources originales. La qualité vidéo des canaux dépend généralement du nombre d'utilisateurs qui sont en train de les regarder ; la qualité vidéo est meilleure s'il y a plus d'utilisateurs. Un système de diffusion en P2PTV est généralement beaucoup moins cher que les alternatives, et peut être initié au niveau individuel. En revanche, il pose des problèmes de qualité de service si comparé à l'unicastré (l'architecture client-serveur généralement utilisée dans le streaming) puisque personne ne peut garantir une source fiable, chaque utilisateur étant aussi un réémetteur.

MYTHE N° 4 :

LE PEER-TO-PEER, C'EST DU LOGICIEL LIBRE ET/OU DE L'OPEN SOURCE

Comme les autres « mythes » présentés dans ce chapitre, l'idée que le peer-to-peer et le logiciel libre ou open source coïncident dérive d'un ensemble de conceptions et de facteurs qui tiennent à la fois du « politique », de l'économique, du social, mêlant des arguments établis à des idées reçues. Beaucoup d'outils en peer-to-peer, en particulier les premiers grands systèmes de partage de fichiers, sont effectivement nés au sein des communautés de logiciel libre et de l'open source et en ont, à leur tour, facilité le développement et l'organisation, dans une démarche éthique commune de partage de ressources, de gestion consensuelle et sans centre, d'attribution d'importance au choix et à la liberté de l'utilisateur. Pourtant, un nombre

important d'applications, sous-tendant une technologie P2P et servant des usages variés, sont à ce jour partiellement ou complètement propriétaires.

Comme a souligné en 2000 le développeur Dave Winer, « *The P in P2P is People* » : c'est-à-dire, ce qui est important dans les réseaux peer-to-peer, ce sont les gens. Ce commentaire souligne en quoi la connexion entre le développement d'applications peer-to-peer et le mouvement open source est significatif : les projets open source s'organisent autour de groupes de travail décentralisés, qui s'autogèrent et sont eux-mêmes rendus possibles par des technologies Internet en peer-to-peer. Si la P dans P2P est celle de « People » - note Tim O'Reilly - soit les technologies permettant aux gens de créer des communautés qui s'auto-organisent, soit les cadres organisationnels développés afin de gérer ces communautés, donnent d'importantes leçons pour ceux qui veulent travailler dans l'espace P2P.

L'open source n'est pas tout simplement déterminé par un ensemble de licences pour la distribution des logiciels, mais, à un niveau plus profond, par un ensemble de techniques pour un développement de logiciels collaboratif et global. C'est là que, en effet, que la boucle entre l'open source et le peer-to-peer se voit bouclée, comme avait déjà montré un des moteurs de la première communauté open source, Usenet : un système qui, sans contrôle central, copie des fichiers entre ordinateurs, et dans lequel chaque site participant au réseau sauvegarde ses copies des messages postés dans les forums, et se synchronise périodiquement avec ses pairs. Les « labels » open source et peer-to-peer indiquent donc tous les deux, généralement, des technologies ou des communautés permettant aux gens de s'associer librement, de manière directe, et sont souvent parmi les incubateurs d'innovation les plus prometteurs.

Ces similitudes, pourtant, ne devraient pas emmener à traiter le peer-to-peer et l'open source comme étant tout à fait coïncidents, ce que montrent plusieurs cas récents. Un désormais célèbre service de voix sur IP implémente une architecture P2P basée sur un protocole propriétaire. À côté de nombreux avantages, liés à la connexion directe et au partage de ressources de bande passante entre utilisateurs, le « mariage » entre P2P et propriétaire à l'œuvre dans ce logiciel a cependant donné lieu à une importante controverse quant au manque d'interopérabilité du système avec d'autres systèmes, P2P et non.

Une start-up proposant un moteur de recherche P2P qui vise une « distribution totale » de la recherche en ligne ne publie pas, quant à elle, son produit comme open source. Elle considère que le modèle open source est parfait lorsqu'on compète au moyen d'un avantage de coût avec un produit commercial au même niveau technologique, comme c'est le cas avec les systèmes d'exploitation Linux ou le logiciel de traitement de texte OpenOffice, mais que ce modèle n'est pas une bonne idée quand on possède un avantage technologique par rapport à un monopole, comme ce serait le cas pour un moteur de recherche P2P par rapport au « géant » Google, et on doit se confronter à un service fourni de façon gratuite, soutenu par un « pouvoir de marque » très puissant.

Une entreprise développant une application pour le stockage et l'accès distribué de fichiers privés, reposant à la fois sur une plateforme de serveurs et sur une approche d'architecture distribuée et décentralisée, où le fonctionnement du dispositif repose sur la mise à disposition de ressources *hardware* de la part des usagers, revendique quant à elle un statut de dispositif hybride P2P. Cependant, celui-ci est aussi un logiciel propriétaire, car son code source est fermé, même si les projets universitaires sur lesquels le logiciel se fonde sont partiellement, à l'origine, de l'open source, et leur contribution reconnue sur le site de l'entreprise.

MYTHE N° 5 :

LE PEER-TO-PEER PRIVE, C'EST DU « DARKNET », DE L'« INTERNET ILLEGAL »

A la fin de 2002, quatre ingénieurs faisant partie du groupe de recherche sur la sécurité de Microsoft créèrent le terme « darknet », dans un papier très influent pour la suite, appelé « *The Darknet and the Future of Content Distribution* », pour se référer à l'Internet « souterrain ». On était alors dans l'environnement d'après-Napster et d'avant-Gnutella. Les ingénieurs définirent dans ce papier le darknet en tant que « collection de réseaux et de technologies utilisées pour partager du contenu numérique ». Suivant le papier, le mot a infiltré les médias généralistes, et a été utilisé pour se référer à une variété d'activités et technologies « clandestines » sur l'Internet. Entre 2003 et 2005, le terme « darknet » a été utilisé comme étiquette d'une quantité d'activités à l'allure menaçante et incertaine, de « cyberclubs » privés, à bases de données en ligne mais puissamment sécurisées et non-traçables avec les moteurs de recherche grand public, ou encore, au monde du cybercrime et du spam, et les autres « endroits obscurs » de l'Internet utilisés pour échapper à la loi.

En même temps, le mot a été utilisé pour distinguer les réseaux distribués anonymes et privés de leurs prédécesseurs « publics ». L'élément de *privacy* a été introduit pour la première fois dans un travail juridique en 2004 ; le darknet y est défini comme la collection de réseaux et autres technologies qui permettent aux personnes de partager matériaux numériques « sans peur d'être découverts ». J. D. Lasica décrit ailleurs le darknet comme un réseau de personnes qui utilisent des espaces fermés – des ports francs, à la fois virtuels et réels, où il n'y a que peu ou pas de possibilité d'être découverts – pour partager du matériel numérique avec des autres, afin d'éviter les restrictions sur les médias numériques imposés par les compagnies du monde du divertissement. Il en résulte une superposition des darknet et des réseaux P2P privés dans une vision de supermarché de médias numériques avec une mentalité de « *wild west* », qui pourrait rivaliser les produits et services fournis par les grandes industries du contenu avec les armes de la *privacy*, de l'invisibilité même. L'ambiguïté entre le peer-to-peer privé et le réseautage Internet illégal et souterrain s'est donc vue renforcée.

Pourtant, le peer-to-peer « privé » ne se limite pas aux darknet. Au delà de la connotation d'illégalité qu'ils peuvent avoir assumé, la caractéristique principale de ces réseaux est leur statut de systèmes de connexion « friend-to-friend », à signifier que des connexions directes sont établies seulement entre des amis reconnus. Plus généralement, tout réseau de partage de fichiers privé peut être défini comme un réseau, ou un ensemble de réseaux, distribué et décentralisé (sans index central) qui inclut des fonctions de privacy, sécurité (encryptage), et anonymat de l'utilisateur, qui a le but primaire de partager de l'information avec des membres certifiés du réseau.

Une des entreprises proposant un service de P2P privé souligne que ce système permet d'apporter une « solution nouvelle à un problème classique », mieux que d'autres solutions existantes : les services de stockage en ligne sont limités au niveau de l'espace disponible et nécessitent la recopie des fichiers chez un tiers ; les services d'envoi de fichier ne conviennent pas pour partager des dossiers complets ; le FTP demande des connaissances techniques pointues ; les plateformes de streaming ne sont pas adaptées pour les échanges privés et les solutions de P2P existantes ne sont pas assez sécurisées.

Comme explique Fabrice Le Fessant, dans des connexions peer-to-peer privées, ami-à-ami, chaque usager héberge sa page personnelle sur son ordinateur, avec, en particulier, toute information ou donnée qu'il considère personnelle. Il autorise ses amis, un par un, à accéder à son ordinateur. A cette fin, il envoie à chaque ami une clé secrète, qui sera utilisée par l'ami lors de sa première connexion au moyen de l'application P2P. Au moment de cette première

connexion, un nouveau « secret » est échangé, et sera utilisé pour toutes les connexions suivantes. La première clé secrète n'est pas réutilisable, en évitant ainsi toute interception par une troisième personne ou entité.

Ce mécanisme d'identification et de distribution de contenus est actuellement utilisé dans nombre d'applications pour le partage de données et d'informations personnelles, en permettant un accès plus rapide aux contenus (puisque ceux-ci sont directement disponibles sur l'ordinateur de l'utilisateur, sans besoin de le recopier sur un site), tout en les rendant accessibles seulement à des amis utilisant le même programme.

MYTHE N° 6 :

LA DIFFUSION DU CLOUD COMPUTING SIGNIFIE LA MORT DU PEER-TO-PEER

Bien que la définition même de *Cloud* soit actuellement l'objet de vives controverses (une revue spécialisée a récemment réuni plus de vingt définitions différentes du concept), ce modèle indique généralement que le vendeur fournit l'infrastructure physique et le produit logiciel, abritant ainsi à la fois les applications et les données dans un lieu inconnu de l'utilisateur (le fameux « nuage », *Cloud* en anglais) et interagit avec ce dernier grâce à une interface client. On s'achemine dans ce cas vers un modèle de déportation et de virtualisation sur des serveurs distants de traitements informatiques traditionnellement localisés sur le poste utilisateur. Eben Moglen, professeur à Columbia University et inspirateur du réseau social décentralisé Diaspora, a récemment affirmé que, dans un paysage de services internet dominé par le paradigme client-serveur, ce qui est actuellement rangé sous l'étiquette de la tendance *Cloud Computing* n'est rien d'autre que « des serveurs qui ont gagné [davantage de] liberté. Liberté de bouger. Liberté de louvoyer ; de combiner et de diviser, de ré-agréger et d'utiliser toute sorte d'astuces. Les serveurs ont gagné en liberté. Les clients n'ont rien gagné ».

Dans ces conditions - alors qu'un modèle économique et technique dans lequel l'utilisateur final sollicite de puissants centres de serveurs, qui stockent l'information et gèrent le trafic sur le réseau - certains soutiennent que un « P2P turn », tel qu'on l'a décrit dans les sections précédentes, pourrait ne plus être possible. Certes, il s'agit là d'une tendance inverse à celle proposée avec le modèle P2P, qui vise à (re-)placer l'utilisateur et sa machine au centre de la création, du partage, de la gestion de contenus numériques. Toutefois, le nuage décentralisé ou P2P est aussi envisageable. En fait, les premières expérimentations avec le nuage décentralisé sont déjà à l'œuvre, et seraient conçues pour répartir la puissance de calcul et les ressources du nuage entre les terminaux de tous les utilisateurs/contributeurs, avec l'idée que la liberté de l'utilisateur au sein du 'nuage' et la possibilité pour lui de contrôler entièrement, et par ses propres moyens, ses données personnelles, ne sont pas des buts incompatibles entre eux.

MYTHES ET LEGENDES DE LA MOBILITE EN ENTREPRISE

Jean-Denis Garo, Directeur Communication et Marketing Support Astra

L'augmentation du nombre de situations de mobilité en entreprise n'est plus à prouver. L'entreprise moderne se doit d'être flexible, mobile et connectée.

Ces situations sont nombreuses : la mobilité dans le bâtiment, entre des sites, pendant les trajets ou même au domicile, mais pas nouvelles. La technologie, les usages ont pourtant modifié ces pratiques.

MYTHE N° 1 :

DES NOUVEAUX TERMINAUX, DONC DES NOUVEAUX USAGES

Les solutions proposées aux utilisateurs mobiles des entreprises sont diverses et dépendent essentiellement des usages. Et même si le Smartphone ou la tablette semblent plébiscités, le coût de chaque terminal ne permet pas l'équipement de l'ensemble des salariés. Ces terminaux sont réservés à des catégories bien définies, souvent direction, équipes commerciales, nomades...

Par rapport au Smartphone ou au terminal DECT, la tablette présente un écran plus grand, permettant de bénéficier lors des communications de l'apport de la vidéo ; l'expérience utilisateur sera ainsi nettement améliorée. L'arrivée de la mobilité à très haut débit (4G) va permettre d'apporter sur les équipements mobiles (Smartphone et tablette) d'avantage de fonctionnalités quelques soit la localisation géographique. La Vidéo, disponible en Wi-Fi actuellement, fait partie des nouveaux usages qui vont se démocratiser dans un avenir très proche sur les équipements mobiles. Pourtant le phénomène BYOD (Bring Your Own Device : utiliser son propre terminal dans son environnement professionnel) ne renverse pas la tendance, et au-delà de l'engouement médiatique ne semble pas recueillir aujourd'hui l'adoption des DSI.

En effet l'entreprise semble plus sensible aux solutions de type Convergence Fixe Mobile, qui permet d'offrir les facilités de communications de l'entreprise aux utilisateurs nomades à l'extérieur de l'entreprise. La notion de numéro unique ou même de terminal unique, que ce soit pour la réception ou l'émission d'un appel permet à l'utilisateur de gérer ses communications d'entreprise quel que soit sa localisation géographique. Cumulée avec les fonctions de Least Cost Routing, la Convergence Fixe Mobile permet à l'administrateur de maîtriser les coûts de sa flotte mobile pour les nomades à l'international.

Bien évidemment tous les utilisateurs ne vont pas avoir besoin de disposer de l'ensemble de ces équipements, par exemple un magasinier qui souvent se déplace sur une zone limitée à l'intérieur de l'entreprise n'aura souvent pas besoin d'un autre terminal que son terminal DECT, un commercial a contrario qui se déplace à l'extérieur de l'entreprise pourra utiliser son Smartphone couplé avec son PC et/ou sa tablette, un technicien se déplaçant sur site privilégiera plutôt le Smartphone et/ou la tablette.

MYTHE N° 2 : LE DECT SERAIT REMPLACÉ PAR LE WI-FI

Si les terminaux Wi-Fi ont encore du mal à se faire une place en entreprise, et en particulier dans le secteur hospitalier, le DECT lui est préféré. Cette technologie qui existe depuis pourtant longtemps a bénéficié d'innovations technologiques fortes. Le capital robustesse et fiabilité des solutions DECT n'a jamais été démenti, il est resté aussi, plus abordable financièrement que les terminaux Wi-Fi ou les Smartphones.

Au milieu des années 2000 une vague de rajeunissement, suivant l'impulsion du marché privé, a durablement inscrit ces terminaux dans la modernité : plus léger, plus compact, avec écran couleur. Mais l'innovation majeure reste l'évolution vers IP, suivant le développement d'équipement des serveurs de communication, le terminal DECT IP profite des atouts de ce protocole, pour devenir encore plus facile à déployer (sans nécessiter de câblage spécifique, dès lors qu'un réseau LAN/WAN est disponible. Les bornes prennent en compte la qualité de service sur le réseau IP).

L'intégration du PTI au terminal DECT est une autre innovation marquante, jusqu'alors réservée à des terminaux spécifiques couteux. (**PTI – Protection du Travailleur Isolé : norme visant à apporter des moyens d'alertes et de localisation pour le travailleur isolé**). Et plus récemment, citons l'arrivée du CAT-iq un standard de mobilité, visant à apporter encore plus de services innovants dans le terminal DECT (qualité audio, vidéo, services internet, etc...)

Enfin, les bornes SIP-DECT en intégrant un port USB font aussi leur révolution : par exemple pour raccorder des caméras vidéo USB aux bornes DECT, la composante vidéo est apportée au sein de l'application de géolocalisation. L'efficacité et la réactivité des superviseurs s'en trouvent renforcées.

Pour le même besoin de communications internes, la technologie de Voix sur réseau sans fil WLAN (Wi-Fi) est aussi utilisée, essentiellement lorsque l'entreprise souhaite mutualiser le déploiement d'un réseau WLAN pour des besoins de téléphonie en plus des besoins informatiques.

CONCLUSION

De nouveaux usages sont en train de se développer très rapidement et sont de plus en plus demandés, y compris par des usagers mobiles : les communications vidéo, les messages instantanés, le partage d'informations par les réseaux sociaux, etc. Des solutions sont d'ores et déjà proposées par certains acteurs précurseurs dans le domaine des communications unifiées. La mobilité en entreprise s'adaptera et se réinventera à nouveau.

MYTHES ET LEGENDES DE LA VIDEOCONFERENCE

Jean-Denis Garo, Directeur Communication et Marketing Support Astra

La vidéo redevient un sujet d'actualité pour les DSI et les DG. Il est vrai que l'offre s'est diversifiée : des solutions gratuites sur PC, aux solutions intégrées dans les solutions de communications unifiées, aux nouveaux terminaux dédiés, sans oublier les salles de téléprésence. Toutefois subsistent autour de ces solutions un certain nombre de mythes.

MYTHE N° 1 :

LA VIDEOCONFERENCE EST RESERVEE AUX GRANDES ENTREPRISES

L'usage de la vidéo au sein des entreprises est plus que jamais une réalité. Il persiste néanmoins un décalage entre les attentes des utilisateurs, les messages marketing des acteurs spécialisés, et la mise en œuvre réelle de ces solutions. Réservées auparavant à certaines catégories d'utilisateurs dans l'entreprise les solutions tendent aujourd'hui à se généraliser à la plupart des utilisateurs de l'entreprise. La vidéo est de plus en plus considérée comme une extension du système de communications et d'informations (SI), et donc comme un média complémentaire à la voix et la data.

La vidéoconférence se démocratise, elle répond à de nouvelles attentes, et prend de nouvelles formes. En effet les solutions historiques de vidéoconférence (salles dédiées) répondent généralement à des besoins de réunions longues, programmées, privilégiant la parole (à l'écrit), dans des salles dédiées, offrant parfois l'apparence d'une réunion virtuelle (téléprésence, co-présence physique).

Désormais plus facilement utilisables, les solutions de vidéoconférence peuvent aussi être initiées depuis un PC portable (favorisant le nomadisme ou le télétravail), et bientôt à partir d'un téléviseur, d'une tablette numérique, ou même d'un Smartphone. Les solutions de vidéoconférence sur PC sont, elles, plus utilisées pour des réunions impromptues, où le partage de document prendra rapidement le pas sur la fenêtre vidéo. Souvent utilisées pour un suivi de projet (follow up) elles sont, du fait de leurs coûts réduits, plus accessibles aux PME. L'appropriation de ces nouveaux modes de communication a profité de l'usage banalisé des applications Skype ou MSN par le grand public. Ils sont aujourd'hui également relayés par d'autres solutions de vidéoconférence comme les Webconférences.

L'émergence de terminaux dédiés à la vidéo et à certaines applications offre une troisième voie. Celle du confort et de la convivialité d'un terminal vidéo HD, utilisable directement sur le poste de travail individuel ou en salle de réunion, permettant de poursuivre le partage de documents sur le PC, mais surtout offrant une simplicité d'usage telle qu'elle révolutionne les comportements utilisateurs. Il n'est ainsi pas interdit de penser que demain les appels voix dans l'entreprise seront naturellement remplacés par des appels vidéos, de la même manière que les messages vocaux ou les messages d'attentes, d'accueil deviendront des messages vidéos.

...

Le besoin crée donc toujours la solution, l'ensemble des solutions de vidéoconférence s'interconnectent / inter-opèrent pour répondre aux nouveaux usages de la mobilité et aussi aux budgets des entreprises.

MYTHE N° 2 :

LA VIDEOCONFERENCE NECESSITE DES INVESTISSEMENTS COUTEUX.

Depuis 2005 et la version beta de Skype qui introduisit la vidéo sur les ordinateurs familiaux, il n'est plus possible de tenir cette position. On nous opposera que ce type de solutions est plutôt réservé à un usage grand public, notamment du fait de failles de sécurité révélées ces dernières années. Reste que cette solution est répandue dans nombre de PME . L'affirmation des coûts ne touche plus les solutions dédiées aux entreprises. Les grands acteurs historiques des salles de conférences ont en effet vu croître un certain nombre d'acteurs proposant des solutions logicielles sur PC, et offrant des garanties en terme de sécurité, mais sans être capables pour autant d'offrir la qualité HD souvent attendue par les entreprises pour les relations externes. Pour combler cette lacune, d'autres acteurs proposent depuis le début 2001 de porter des solutions de vidéoconférence sur des équipements dédiés HD, qu'il s'agisse de terminaux multimédia ou de tablettes. Le positionnement de ces offres se situe le plus souvent un peu au-dessus du prix d'un ordinateur portable haut de gamme.

Concernant le déploiement, l'émergence de protocoles comme SIP (Session Initiation protocol) facilite et simplifie tant l'administration que le temps dévolu à l'installation, engendrant des économies immédiates. En outre les interfaces tactiles et intuitives remettent les nombreuses télécommandes aux solutions du passé, et permettent de réduire significativement les coûts de formation et d'utilisation généralement induits par ce type d'outil.

MYTHE N° 3 :

LA VIDEOCONFERENCE EST RESERVEE AUX EXECUTIFS

Au-delà des impacts techniques et financiers amenés par la mise en place d'une solution vidéo IP dans une entreprise, on constate que la vidéo ne concerne souvent que certaines catégories de personnels et de groupes de travail. Le secteur d'activité de l'entreprise est aussi un élément très segmentant. Ce résultat n'est pas une surprise, dans la mesure où les entreprises recherchent de moins en moins à mettre en œuvre des solutions ponctuelles mais plutôt à faire de la vidéo une application clé de leur système d'information. Les entreprises sont au milieu du gué sur ce type de projet, partagées entre leur volonté de généralisation de ce type de projets et les difficultés technologiques et économiques qui président à intégrer complètement la vidéo dans le système d'information

Le coût n'étant plus un frein à l'utilisation, la démocratisation est en marche et de nouveaux usages apparaissent. Des entreprises commencent d'ailleurs à se servir des Terminals Multimédias HD comme elles utilisent un vidéoprojecteur. Elles le réservent pour une réunion dans une salle qui n'en serait pas encore équipée. Les télétravailleurs sont aussi équipés....

MYTHE N° 4 :

LA VIDEO C'EST JUSTE POUR LES SALLES DE CONFERENCE

La vidéo aujourd'hui est omniprésente. Sur son téléphone mobile, sa console de jeux, sur une tablette, sur son PC, sur un terminal multimédia, dans une salle de conférence, sur une borne interactive, sur le portier de son immeuble, et prochainement sur la télévision, ou dans les lunettes. La vidéo devient un média simple, facile à déployer, intuitif et attendu : là où nous prenions des photos, nous réalisons des minis films que l'on s'empresse de partager quelque soit le support. Dans les années à venir, la pression des usages grand public et le consumérisme des solutions technologiques feront qu'il deviendra rapidement indispensable de se voir aussi bien dans l'environnement entreprise, que dans l'environnement privé. Et ce besoin sera

récurrent, quel que soit l'environnement existant : à son bureau ou en salle de réunion, en déplacement professionnel ou en situation de télétravail. A ce titre, la salle de conférence ne deviendra qu'un des environnements possibles.

MYTHE N° 5 :

LA VIDEO EST UNE SOLUTION TOTALEMENT SEPARÉE

Une solution vidéo est considérée comme une extension du système de communication, et donc comme un média complémentaire à la voix et la data. Les entreprises restent encore attentives et prudentes à considérer la vidéo comme une application bureautique généralisée, sachant qu'elle nécessite une infrastructure système et réseau beaucoup plus rigoureuse. La tentation de déployer des solutions autonomes et indépendantes du SI de l'entreprise est donc grande.

Cependant, le simple fait que la vidéo soit de plus en plus considérée comme une extension du système de communication démontre sa place importante au sein même des processus métiers de l'entreprise. En conséquence, il devient aujourd'hui impensable que les solutions traditionnelles isolées (solutions traditionnelles de vidéoconférence reposant sur la technologie TDM) ne soient pas progressivement remplacées par des solutions vidéo capables de s'imbriquer presque nativement dans le SI et donc de contribuer à la productivité des directions métiers de l'entreprise.

MYTHE N° 6 :

LA VIDEO EST IMPERSONNELLE ET RETIRE LA VALEUR ATTACHEE A UN VRAI FACE A FACE

Rien ne remplace une rencontre, mais pour la suite, la qualité des solutions (voix, image) apportent le non corporel nécessaire, et enrichissent particulièrement la communication. En ajoutant l'image à la voix, la vidéo donne à la communication une dimension complémentaire dans laquelle la gestuelle et les comportements viennent enrichir le contenu oral : la communication vidéo est une bonne illustration d'un fait bien connu des professionnels de la communication, pour lesquels la signification d'un message passe autant par son contenant que par son contenu.

CONCLUSION

La vidéo est en passe de devenir un média aussi important et naturel que le mail ou le téléphone.

Dans une étude récente, les utilisateurs considèrent que la vidéo ne doit pas être cantonnée aux salles de conférences dédiées ou à une simple utilisation depuis un PC : si 30% des personnes interrogées estiment que la vidéo doit être utilisée uniquement pour les réunions, 45% considèrent aussi la vidéo comme un moyen naturel pour enrichir les communications interpersonnelles au sein de l'entreprise. Vraisemblablement vers un usage « à la volée ».

Plusieurs facteurs expliquent ce regain d'intérêt et ce nouveau regard sur la vidéo :

D'une part, les crises économiques de ces dernières années, conjuguées aux différents risques de catastrophes naturelles et pandémiques, ont contraint les entreprises à trouver des solutions pour préserver leur compétitivité (réduction des coûts et d'optimisation des dépenses) et la productivité de leurs employés ; et cela tout en renforçant leurs démarches de développement durable.

D'autre part, la vidéo a atteint une maturité technologique qui la place comme un média incontournable pour apporter de la valeur et améliorer la collaboration entre les équipes et les salariés. Ces derniers, éduqués pas les solutions grand public, ont trouvé dans les solutions professionnelles les garanties de sécurité et de qualité nécessaires à leur activité professionnelle.

Au final, c'est à la fin du paradigme que nous assistons : celui selon lequel les terminaux dédiés à la téléphonie et la vidéo seront totalement remplacés par les PC et les logiciels de web et vidéoconférence. Les utilisateurs se sont vite vus confrontés aux problèmes de performances des PC : disponibilité, ergonomie, etc.... Le choix d'un terminal n'est plus tant dicté par le contenu qu'il pourra relayer, que par l'environnement dans lequel il sera utilisé.

2° PARTIE : ASPECTS SECURITE ET SURETE



MYTHES ET LÉGENDES DES CHIFFREMENTS HOMOMORPHES

Hervé Lebning, ARCSI

MYTHE N° 1 :

LA NOTION D'HOMOMORPHIE EST UNE IDÉE NOUVELLE

Bien qu'elle soit courante, cette idée est fausse. En mathématiques, la notion d'homomorphie remonte à Camille Jordan (1838 – 1922), et elle est sous-jacente dans les travaux d'Évariste Galois (1811 – 1832). Elle concerne les structures algébriques, c'est-à-dire les ensembles comme celui des nombres entiers naturels 0, 1, 2, 3, 4, 5, *etc.* munis d'opérations comme l'addition ou la multiplication. Une transformation est homomorphe pour l'addition si la somme des transformés est égale au transformé de la somme, de même pour la multiplication en remplaçant le mot « somme » par le mot « produit ».

Dans ce cadre, la transformation qui consiste à multiplier par un nombre fixe, par exemple 35, est homomorphe pour l'addition car $35 \times (56 + 75) = 35 \times 56 + 35 \times 75$, ceci restant vrai si on remplace 56 et 75 par n'importe quels nombres.

De même, la transformation qui consiste à élever à une certaine puissance, comme 7 par exemple, est homomorphe pour la multiplication car $(25 \times 32)^7 = 25^7 \times 32^7$, ceci restant vrai si on remplace 25 et 32 par n'importe quels nombres.

À l'ère numérique, tout message étant un nombre et son chiffrement se faisant par une transformation en un autre nombre, la notion d'homomorphie s'applique à celle de chiffrement. On peut d'ailleurs considérer la multiplication par 35 et l'élevation à la puissance 7 comme des chiffrements, même s'ils manquent singulièrement de sécurité.

Messages	Chiffrement	Chiffrés
1101111000111010010	→	111100110000111110110110
1011001110011000011	→	110001000110111010101001
11001000111010010101	→	1101101110111111001011111

Principe d'un chiffrement homomorphe pour l'addition. Les messages comme les chiffrés sont des suites de bits (donc des nombres exprimés en binaire). Dans cet exemple, le chiffrement consiste à multiplier les messages par 35 (100011 en binaire). Le chiffré de la somme des deux messages est égal à la somme de leurs chiffrés. Pour déchiffrer, il suffit de diviser par 35.

La nouveauté n'est donc pas la notion d'homomorphie mais l'idée de l'appliquer à la cryptographie. Quel intérêt ? La raison est que si un chiffrement est homomorphe pour

l'addition, la somme des chiffrés est le chiffré de la somme. Autrement dit, les calculs peuvent être faits sur les chiffrés, donc en gardant le secret sur les données. L'intérêt de la cryptographie homomorphe est là, comme nous le verrons sur des exemples concrets plus loin.

MYTHE N° 2 :

LES CHIFFREMENTS HOMOMORPHES NE SONT QU'EN COURS DE DEVELOPPEMENT

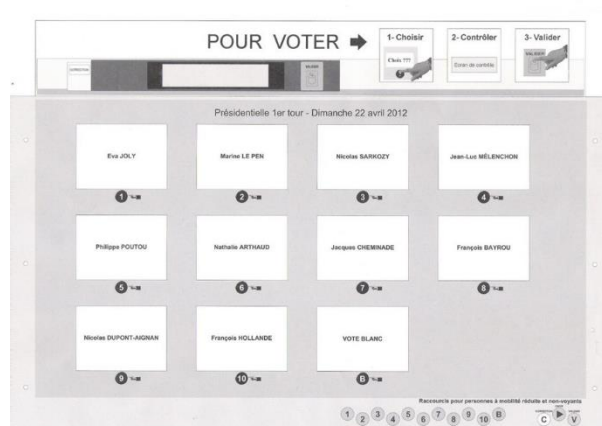
Cette idée vient d'une confusion entre chiffrement homomorphe et chiffrement doublement homomorphe, c'est-à-dire pour les deux opérations (addition et multiplication). L'élévation à une puissance est au cœur du système de chiffrement asymétrique RSA, inventé par Ronald Rivest, Adi Shamir et Leonard Adleman, en 1978, qui est donc homomorphe pour une seule opération, la multiplication. En s'inspirant de l'idée de base de ce chiffrement, en 1999, Pascal Paillier a inventé un chiffrement homomorphe pour l'addition. En ajoutant les chiffrés de deux nombres, on obtient donc le chiffré de leur somme, comme nous l'avons vu plus haut.

L'existence d'un chiffrement doublement homomorphe sûr a été conjecturée dès 1978 par Rivest, Adleman et Dertouzos. Un cap théorique a été franchi en 2009 quand Craig Gentry, doctorant à l'Université de Stanford, inventa le premier chiffrement doublement homomorphe et sûr. Malheureusement, l'avancée n'est que théorique car la clef de son chiffre est très longue. Plus de deux gigabits qui le rendent actuellement impraticable !

MYTHE N° 3 :

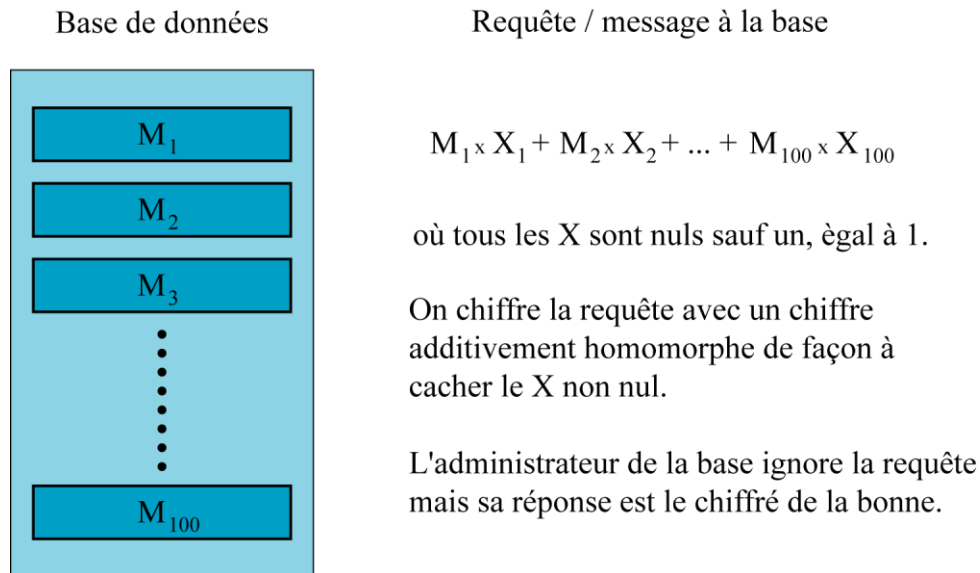
LES CHIFFREMENTS HOMOMORPHES N'ONT ENCORE AUCUNE APPLICATION

Un chiffrement homomorphe additif peut être utile dans un système de vote électronique, ou sur Internet. Chaque vote y est masqué au moyen d'un chiffrement. Le décompte est effectué sur ces chiffrés. Le déchiffrement de la somme des chiffrés donne le résultat sans que l'on déchiffre pour autant chaque vote. La méthode est en particulier utilisée par le système *Helios* sur Internet. Nous ne pouvons dire qui utilise ce principe car les logiciels utilisés par les machines à voter sont tenus secrets.



Une machine à voter utilisée lors du premier tour de l'élection présidentielle de 2012. Son logiciel est tenu secret, ce qui tranche singulièrement sur la transparence des urnes classiques et contredit les principes de la cryptographie moderne énoncés par Auguste Kerckhoffs (1835 – 1903), selon lesquels un système cryptographique doit être protégé par sa clef et non pas par le secret de l'algorithme.

De façon sans doute plus étonnante, et mathématiquement plus subtile, un chiffrement homomorphe additif permet de consulter des bases de données sans révéler la requête faite à l'administrateur de la base. En particulier, cette propriété est intéressante pour un investisseur qui ne veut pas que l'on sache à quelles sociétés il s'intéresse, ainsi que dans le domaine médical.



Principe de la requête cachée à une base de données, avec un chiffrement homomorphe additif. L'administrateur de la base ne peut savoir la requête effectuée car elle correspond au seul X non nul, qui est chiffré.

Un chiffrement doublement homomorphe est encore plus utile puisque tout calcul impliquant seulement des additions et des multiplications peut alors être fait sur les chiffrés. Cela inclut, en particulier, tous les calculs comptables, financiers ou épidémiologiques dont les modèles mathématiques n'utilisent que les deux opérations (addition et multiplication). La double homomorphie est donc particulièrement utile si on veut externaliser des calculs dans des nuages ou *clouds*. Il est alors possible de demander à un prestataire d'effectuer certains calculs sur les données chiffrées qu'on lui fournit et déchiffrer ses résultats pour obtenir les informations voulues.

	Chiffrement	
Données		Chiffrés
145	→	1001111010011
65	→	100011100011
42	→	10110111110
10	→	101011110
	Déchiffrement	
537	←	100100101101011

Calcul de deux fois la première donnée, plus trois fois la seconde, plus la troisième et la quatrième : il est fait sur les chiffrés des données (en binaire ici). En déchiffrant le résultat obtenu sur les chiffrés, on obtient le résultat sur les données initiales, que le prestataire de service ignore donc.

MYTHES ET LEGENDES DES APT

Nicolas RUFF – Google

Pour commencer, précisons immédiatement qu'APT signifie "*Advanced Persistent Threats*", soit "attaques complexes et récurrentes" dans une traduction approximative. Il s'agit d'un *buzzword* qui (si ma mémoire est bonne) a été inventé autour de 2008 par la société MANDIANT¹⁷, en réponse aux incidents de sécurité de plus en plus nombreux et de plus en plus graves rapportés par la presse à cette époque.

Il est difficile de définir plus précisément le terme APT. Pour paraphraser l'expert en sécurité Cesar Cerrudo: "Lorsqu'une petite entreprise est piratée, elle est victime de ses lacunes en sécurité. Lorsqu'une grande entreprise est piratée, elle est victime d'une APT".

La confusion est également alimentée par les vendeurs de produits de sécurité, qui prennent le train en marche et rajoutent une couche de marketing-fiction avec des noms futuristes tels que: opération "Titan Rain", opération "Aurora", opération "Ghost Net", opération "Night Dragon", opération "Shady RAT", et autres "*Advanced Evasion Threats*".

On notera toutefois que les entreprises piratées étaient quasiment toutes équipées avec la plupart des solutions de sécurité du marché. Et que les vendeurs de ces solutions apparaissent également sur la liste des victimes ... (*Achevé de rédiger en janvier 2012*).

MYTHE N° 1 :

LES APT NE ME CONCERNENT PAS

C'est l'idée fausse la plus répandue: seules les institutions politiques, les systèmes militaires et les centres de recherche des grandes multinationales occidentales seraient visés.

Or l'actualité nous démontre tout le contraire:

Une petite entreprise peut être ciblée, car elle dispose d'un savoir-faire ou d'une technologie unique en son genre.

Une entreprise commerciale quelconque peut être visée, car elle est sous-traitante d'une cible plus intéressante. C'est le cas de la société RSA, qui fournit des solutions de sécurité à Lockheed Martin¹⁸ (et de nombreux autres industriels). La pratique anarchique de l'externalisation multiplie le nombre de vecteurs d'entrée.

Une cible intéressante, mais difficile à pénétrer "en direct", peut être compromise par le biais d'une filiale de moindre importance. En effet, le niveau d'isolation et de surveillance entre réseaux internes est souvent bien moindre qu'avec Internet. Google suspecte fortement sa filiale chinoise d'être à l'origine de la compromission détectée en janvier 2010¹⁹.

¹⁷ http://www.mandiant.com/services/advanced_persistent_threat/

¹⁸ http://www.nytimes.com/2011/05/30/business/30back.html?_r=1

¹⁹ http://fr.wikipedia.org/wiki/Op%C3%A9ration_Aurora

Même une organisation à but non lucratif peut être victime. Ce fût le cas par exemple de plusieurs agences anti-dopage²⁰, ou d'association de défense des droits de l'homme au Tibet²¹.

MYTHE N° 2 :

LES APT FONT APPEL A DES COMPETENCES TECHNIQUES HORS DU COMMUN

C'est aussi une idée fausse, bien souvent propagée par les victimes pour se défaire de leurs responsabilités.

Mon analyse de plusieurs dizaines de cas réels démontre que:

- Les attaquants recyclent du code écrit par d'autres, parfois même sans le comprendre.
- Les attaquants utilisent des outils librement téléchargeables sur Internet (tels que la *backdoor* Poison Ivy²² ou la suite d'outils Pass-The-Hash²³).
- Les codes d'attaque sont souvent bogués.
- Les attaquants se propagent sur les réseaux internes en utilisant des techniques simples, tels que des mots de passe par défaut ou triviaux.

Patrick Pailloux, directeur général de l'ANSSI, a rappelé récemment²⁴ que la plupart des cas d'intrusions dans lesquels ses services sont intervenus découlaient d'une hygiène informatique déplorable, malgré une apparente maîtrise (telles que les certifications ISO 2700x, ITIL, CMMI, et consorts).

Il existe quelques points durs dans les intrusions, comme la découverte et l'exploitation de failles "0day"²⁵. Mais pour 1 attaquant qui dispose d'une telle faille, il en existe 10 (ou 100) qui copient son code sans même le comprendre. Les codes les plus populaires sont rapidement disponibles dans le projet libre et gratuit Metasploit²⁶, ce qui facilite leur dissémination.

Il est évident que toute stratégie de défense contre les intrusions doit prendre en compte la présence de failles dans les logiciels: entre les failles restant à découvrir, celles qu'on peut déjà acheter sur le marché gris, et celles pour lesquelles le correctif n'a pas été déployé sur l'intégralité du parc informatique ...

MYTHE N° 3 :

ON PEUT SE PROTEGER EFFICACEMENT CONTRE LES APT :

Bien entendu, chaque vendeur de solution de sécurité va vous promettre (mais pas vous garantir) que son produit vous protège contre "100% des attaques connues et inconnues".

²⁰ <http://www.mcafee.com/us/resources/white-papers/wp-operation-shady-rat.pdf>
<http://www.franceinfo.fr/faits-divers-justice/floyd-landis-condamne-pour-piratage-informatique-d-un-laboratoire-antidopage-442981-2011-11-10>

²¹ <http://fr.wikipedia.org/wiki/GhostNet>

²² <http://www.poisonivy-rat.com/>

²³ <http://oss.coresecurity.com/projects/pshtoolkit.htm>

²⁴ <http://www.ssi.gouv.fr/fr/anssi/publications/communiqués-de-presse/cloture-les-assises-de-la-securite-et-des-systemes-d-information-2011.html>

²⁵ Faille logicielle inconnue de l'éditeur, et a fortiori pour laquelle aucun correctif n'est disponible.

²⁶ <http://metasploit.com/>

Mais bien entendu, les attaquants achètent les mêmes produits et savent comment les mettre en défaut.

Compte-tenu du flot ininterrompu d'attaques diverses et variées (dont les APT) qui arrivent chaque jour dans vos boîtes aux lettres et dans vos navigateurs, il est illusoire de croire qu'aucun poste de travail ne sera jamais compromis chez vous. StuxNet (mais avant lui Conficker) ont démontré que même les réseaux déconnectés d'Internet pouvaient être attaqués ...

La plupart des entreprises disposent aujourd'hui d'une protection "à l'état de l'art": antivirus, pare-feu, proxy filtrant, utilisateurs non administrateurs de leurs postes de travail, politique de mots de passe robuste ...

A contrario, rares sont les entreprises qui disposent d'une politique de *détection* et de *réaction* adaptée à l'état de la menace ... Ce sont aujourd'hui deux axes d'amélioration majeurs sur lesquels il faut investir.

Il n'est pas question de stéganographie ou de canaux cachés improbables. La plupart des APT se soldent par l'envoi de fichiers de plusieurs gigaoctets, le samedi soir, vers des serveurs situés à l'étranger ... Comment une telle anomalie peut-elle échapper à la vigilance de tous les produits de sécurité, et des humains chargés de les faire fonctionner ?

Malheureusement, la plupart des intrusions sont détectées "par hasard", suite à un problème d'apparence banale: saturation de l'espace disque sur un serveur, lenteur de la connexion Internet ...

Et lorsque le problème est détecté, reste à savoir quoi faire ... Qui contacter ? Faut-il débrancher le serveur ou le laisser actif ? Quel est l'étendue des dégâts ? Comment repartir dans un état sain ? Les premières réactions sont souvent improvisées et catastrophiques: par exemple effacer toutes les traces de l'attaquant ...

MYTHE N° 4 :

L'ATTRIBUTION DES ATTAQUES EST IMPOSSIBLE SUR INTERNET

Il est vrai que la nature pervasive d'Internet permet de masquer efficacement ses traces, pour peu qu'on s'en donne la peine. Il est parfois possible de remonter jusqu'à une adresse IP ou une machine. Mais il reste impossible de savoir qui était au clavier lors de l'attaque ...

Toutefois cette assertion doit être relativisée dans le cas d'une APT. En effet, une attaque longue et récurrente va forcément laisser beaucoup plus de traces qu'une défiguration unitaire de site Web ou l'envoi d'un spam. Parmi tous les indices qu'on peut collecter, on peut citer:

1. Les motifs de connexion discernables.

En analysant les horaires de connexion des attaquants sur une période suffisamment longue, il est possible d'identifier plusieurs informations utiles, telles que le fuseau horaire ou les fêtes nationales. Bien entendu cette information peut être falsifiée si tous les attaquants se mettent d'accord pour agir les jours fériés à 3h du matin ... mais les informations obtenues actuellement par cette méthode sont cohérentes avec les autres sources d'informations.

2. Les sources des outils utilisés.

Il existe souvent de nombreux outils capables de réaliser la même tâche: par exemple il existe un nombre considérable de *backdoors* librement disponibles sur Internet. L'attaquant va privilégier les outils dont la documentation est rédigée en anglais ... ou dans sa langue maternelle.

3. Les traces involontaires.

L'outil de compilation d'un programme ou de génération d'un document va laisser de nombreuses traces à l'insu de l'utilisateur.

Dans le cas de StuxNet, la présence de la chaîne "myrtus" a fait gloser de nombreux observateurs. S'agit-il d'une référence biblique, ou faut-il lire "My RTUs" (*Real Time Unit*) ? La question reste ouverte.

Mais dans d'autres cas, les traces involontaires sont beaucoup plus faciles à interpréter: par exemple la langue par défaut de l'interface graphique, ou le nom d'utilisateur de l'attaquant ...

4. Le mobile de l'intrusion.

Indépendamment des aspects techniques, on peut noter que le nombre de personnes susceptibles d'être intéressées par des informations sur les négociations du G20 ou la construction de réacteurs nucléaires se compte sur les doigts d'une main. Contrairement à des numéros de CB volés, de telles informations sont difficilement exploitables sur le marché gris ... Ce qui fait dire à certains que les APT sont le fait d'états et non de criminels.

Il existe quelques cas où il a été possible d'approcher de très près la source de l'attaque. Lors de l'opération "Aurora", Google a nommément désigné une source chinoise²⁷. Il semble que les équipes sécurité de Google aient contre-attaqué et remonté le fil rouge jusqu'à la source de l'intrusion, bien que cela n'apparaisse qu'en filigrane dans leur communiqué officiel. Il est également arrivé qu'un serveur de rebond sous le contrôle direct de l'attaquant ait été saisi (cas de l'opération "Shady RAT").

MYTHE N° 5 :

LA SECURITE INFORMATIQUE A 100% N'EXISTE PAS

Ceci est l'un des mythes les plus destructeurs qu'on puisse entendre.

Il est vrai que le risque zéro n'existe pas. Mais cette assertion sert bien souvent à justifier des arbitrages totalement absurdes tels que: "... donc je garde mon iPhone et je fais suivre tout mon mail professionnel sur GMail" !

"La sécurité à 100% n'existe pas" sert bien souvent de prétexte pour faire une sécurité à 10%. L'information est comme un fluide: si votre plomberie est à 99% étanche ... alors vous avez déjà un sérieux problème de fuite !

Il est impossible d'empêcher les attaques d'arriver, ni d'appliquer les correctifs de sécurité sur l'intégralité d'un parc informatique hétérogène et étendu. Mais il est possible de faire beaucoup mieux qu'actuellement, grâce aux leviers de la détection et de la réaction évoqués précédemment:

- Une attaque détectée et éradiquée en 1 heure n'a aucun impact sérieux.
- Une attaque détectée et éradiquée en 1 journée nécessitera une analyse post-mortem pour déterminer l'étendue de la compromission.
- Une attaque détectée et éradiquée en 1 semaine laissera le temps à l'attaquant de collecter suffisamment de mots de passe et de poser suffisamment de *backdoors* pour pouvoir revenir à volonté ...

²⁷ http://www.theregister.co.uk/2010/02/19/aurora_china_probe_latest/

P.S. Les attaques les plus longues documentées dans la nature ont officiellement duré ... 3 ans.

MYTHES ET LEGENDES DE L'ANALYSE DE RISQUE

Louis Derathé, Thales

MYTHE N° 1 :

LES CRITERES TRADITIONNELS DE CONFIDENTIALITE, INTEGRITE ET DISPONIBILITE DES INFORMATIONS SONT INSUFFISANTS !

Depuis de très nombreuses années, la sécurité d'une information est appréciée à l'aune de trois critères²⁸ : la Confidentialité, l'Intégrité et la Disponibilité (la fameuse triade C I D).

Longtemps ces critères ont paru suffisants pour conduire à l'élaboration de politiques de sécurité techniques et organisationnelles dans les systèmes informations, permettant de protéger ces informations de toutes les attaques ... enfin, autant que cela était possible !

Mais, les Nouvelles Technologies de l'Information et de la Communication et leur cortège de cyber attaques ont bousculé les certitudes et entaché de doute cette caractérisation ; de nombreux critères sont apparus comme pouvant compléter cette caractérisation « sécurité » des informations et mieux répondre à la menace ambiante : de nombreux gourous ont donc complété la liste proposant en guise de critère, aussi bien des fonctions de sécurité que de sûreté, arguant que tout est dans tout et la sécurité partout.

Certes, la communauté de la SSI su faire la part des choses et rejeter l'ivraie, mais une suggestion récurrente engage encore aujourd'hui un vrai questionnement : **Est-ce que l'authenticité d'une information peut être un critère de sécurité ?**

En effet, une information authentique, c'est-à-dire vraie, ne peut pas nuire à la sécurité d'un système d'information ... et donc, cette qualité devrait être, à l'évidence, protégée ! La confiance, la véracité, la conviction ou la preuve sont des valeurs partagées sur lesquelles une sécurité devrait se construire ; voilà l'évidence qui nous aveugle soudain !

Et voilà surtout l'analyste sécurité plongé dans les affres du doute et de la crainte de l'incomplétude ... de la faille logique de sa démarche d'analyse !

Si l'on se réfère aux définitions couramment offertes par nos dictionnaires, « authenticité » se décline selon deux approches complémentaires : *ce qui est conforme à la vérité, ce dont l'exactitude ne peut être contestée* (Larousse).

Poursuivons notre analyse ! Deux concepts constitutifs de cette notion d'authenticité apparaissent donc : la vérité et sa preuve.

Nous conviendrons sans hésiter que la vérité, en soi, n'est pas une qualité relevant de la sécurité puisqu'elle procède plus de la foi ou de la confiance (du point de vue plus centré sur les systèmes d'information, elle relève des utilisateurs du SI ; c'est une qualité relevant de la source de l'information comme par exemple pour les systèmes de Renseignement et ses logiques de cotation).

Mais l'exactitude, la preuve ! Voilà qui fleure bon la sécurité, non ?

Récapitulons, ce ne serait donc pas l'authenticité en soi, mais bien sa composante preuve/démonstration qui serait alors un nouveau critère sécurité de l'information.

²⁸ Cf. la méthode EB IOS

Or, quelques mots d'un certain Jacques Stern²⁹, trouvés sur le web, éclairent notre réflexion :

« Un service d'authenticité garantit l'identité d'une entité donnée ou l'origine d'une communication ou d'un fichier. Lorsqu'il s'agit d'un fichier, et que l'entité qui l'a créé est la seule à avoir pu apporter la garantie d'authenticité, on parle de « non-répudiation ». Ce service de non-répudiation est réalisé par une signature numérique, qui a une valeur juridique depuis la loi du 20 mars 2000 »

.... « Signature numérique » ! Le mot est dit ... car, qu'est-ce qu'une signature, sinon effectivement un élément permettant à un émetteur de garantir l'intégrité et la provenance d'une information ?

L'authenticité ne serait donc finalement, du point de vue sécurité, que l'intégrité d'une information associée à l'identité de son émetteur ... ce que l'on traduit plus communément dans le langage technique : preuve d'origine ! Et nous voilà à nouveau revenus aux trois critères immémoriaux.

Vite un autre critère ! Un autre gourou !

MYTHE N° 2 :

MAÎTRISER LE RISQUE INFORMATIONNEL AU NIVEAU D'UN ORGANISME (HYPERVISION), C'EST POUR DEMAIN !

Même si les méthodes comme EBIOS 2010 et les travaux actuels regroupés sous le terme « d'hypervision³⁰ » visent à intégrer dans l'analyse du risque, les aspects « métiers » de l'entreprise (en particulier au travers des modalités d'expression des événements redoutés pour la première), la déclinaison d'événements redoutés en risques, traduits par des scénarios de menaces sur des biens supports, valorisés par des facteurs de gravité ou vraisemblance, ne se convertit pas facilement en états gradués d'occurrence des événements redoutés : nous restons à un stade d'appréciation de la probabilité/vraisemblance d'occurrence des risques sur le système informatique et non à celui du degré d'imminence de ces mêmes risques dans une appréhension globale de l'organisme.

Quelles sont les qualités nécessaires d'une véritable hypervision, une maîtrise du risque au niveau de l'organisme ?

TOUT D'ABORD, UNE ANALYSE DE RISQUE FONDÉE SUR LES MISSIONS DE L'ORGANISME

Lorsqu'on parle d'analyse de risque, on pense naturellement à la méthode préconisée par l'ANSSI : EBIOS. Dans le cadre d'une homologation de système d'information, cette méthode, au même titre que PILAR pour l'OTAN, vise à guider l'analyste dans l'expression des besoins de sécurité du système d'information.

Longtemps utilisée dans une optique d'analyse systématique de toutes les vulnérabilités d'un système d'information, la méthode EBIOS dans sa nouvelle version 2010 introduit une nouveauté dans le cadre de l'analyse de risque : elle vise à cerner les principaux enjeux de cette sécurité au travers de l'expression « d'Événements Redoutés » au niveau fonctionnel de

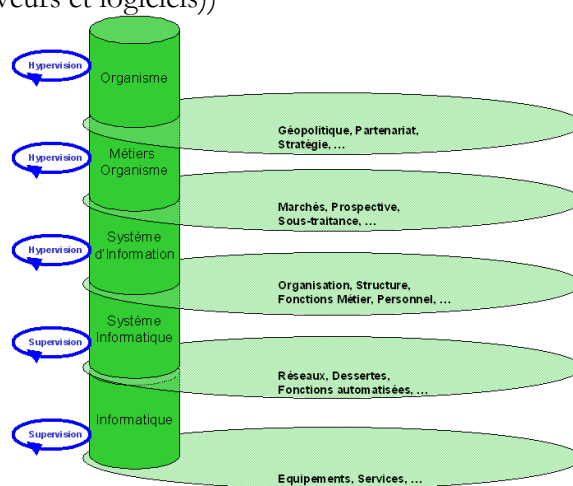
²⁹ Directeur du Laboratoire d'informatique de l'École normale supérieure

³⁰ Ces travaux visent à mettre en œuvre une supervision dynamique étendue à plusieurs systèmes d'information et intégrant des analyses selon les couches classiques : système d'information - système d'information – organisme ; la prise en compte de la dimension « métier » du risque est sa caractéristique actuelle ; plus qu'une méta-supervision, c'est une supervision au niveau de l'entreprise qui est recherchée.

l'organisme ; l'analyse de risque qui en découle est donc fortement orientée par les missions de l'Organisme.

Bien entendu, il n'est pas envisageable de traduire directement des événements redoutés relatifs aux missions d'un organisme en événements informatiques à surveiller, il faut bien décliner ces besoins primordiaux selon les 5 couches du système informationnel de l'organisme qui supportent ces missions :

- Une couche basse composée des architectures techniques informatiques (équipements, services)
- Support de celle supérieure du système informatique au sens de l'outil de traitement (réseaux, dessertes, grandes fonctions (serveurs et logiciels))
- Lui-même, composant principal du système d'information (fonctions supportées essentiellement par l'informatique, mais englobant aussi toute l'organisation de ses métiers comme les approvisionnements, l'organisation et la structure, la sélection du personnel, etc.)
- Au service des métiers de l'organisme (vision opérationnelle par grande fonction de l'organisme)
- Caractérisant l'Organisme dans toutes ses dépendances.



Ainsi, faut-il à l'évidence décliner l'analyse de risque, telle une poupée russe, selon ces couches dont les objectifs de sécurité sont de nature et d'expression différentes.

Le niveau supérieur exprime donc les événements redoutés de l'organisme, imagine les scénarii qui peuvent conduire à ces événements, les apprécie et exprime ainsi les risques encourus. Puis, en cascade, les niveaux subalternes prennent alors en tant qu'événements redoutés de leur analyse, les risques du niveau supérieur et ceci jusqu'aux composants informatiques.

Cette analyse de risque en cascade présente deux caractéristiques très particulières rarement formalisées : si les événements traduisant l'occurrence de risques sont uniquement techniques et de valeur souvent binaire (ON/OFF) aux plus bas niveaux, plus on monte dans les couches, plus ceux-ci quittent la technique pour englober des domaines plus larges comme la politique, le personnel, la stratégie, l'environnement, les pressions externes, etc. De plus, plus on monte dans ces couches, plus les événements à relever et les capteurs associés sont externes à l'organisme.

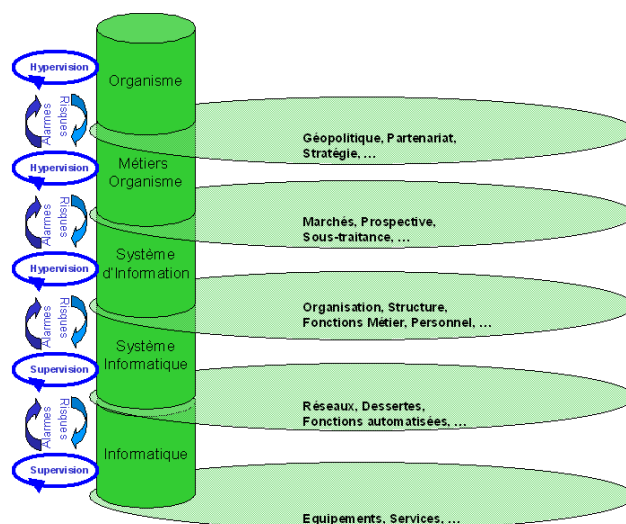
ENSUITE UN FLUX TRANSPOSE DE LA SUPERVISION VERS L'HYPERVISION

Si l'on considère que l'objet de l'hypermémoire est bien de détecter, mesurer et répondre aux atteintes subies par l'organisme, la surveillance devrait naturellement s'appliquer à ce que redoutent les acteurs de l'organisme ; les éléments surveillés devraient donc permettre de détecter l'occurrence des risques conduisant à l'ensemble des événements redoutés qui ont été déterminés lors de l'analyse présentée ci-dessus.

Cette surveillance devrait donc s'effectuer par niveau, et chaque alarme (risque échu) devrait remonter au niveau supérieur en tant qu'événement où elle serait prise en compte et corrélée

avec les événements de ce niveau ; les événements devraient donc remonter par bulles successives de la supervision telle qu'on la connaît jusqu'à l'hypervision :

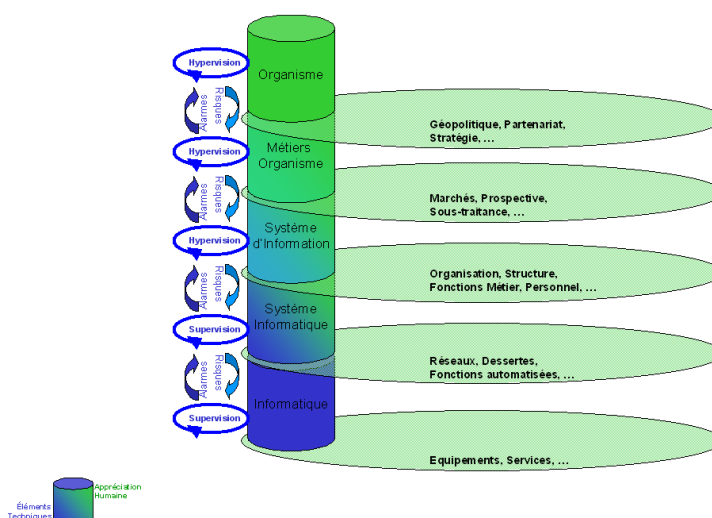
- Les événements informatiques vers le niveau système informatique lorsqu'ils constituent un risque pour ce dernier
- Les risques du niveau système informatique vers le niveau système d'information lorsque corrélés avec des événements relatifs aux architectures, dessertes, personnel informaticien, méthodes du service ou Politique de sécurité informatique
- Puis vers le niveau métier lorsque corrélés avec les aspects de distribution du système, de partage d'information ou de dépendance informationnelle
- Vers le niveau organismes lorsque confortés avec des événements relatifs aux marchés, la concurrence, à l'évolution des pratiques ou de l'état de l'art (intervention de l'Intelligence économique)
- Où enfin pourraient être analysées les conséquences de ces risques pour l'organisme en rapport avec la géopolitique, les groupes de pression, l'image ou l'espionnage



PASSER DE L'OBSERVATION D'ETATS A L'APPRECIATION D'UNE SITUATION

Si l'on veut atteindre un niveau d'appréciation de la sécurité d'un organisme, il faut obtenir une représentation de l'état de cette sécurité en référence aux événements redoutés de plus haut niveau et donc représenter non seulement des faits (panne d'un composant, dysfonctionnement d'une fonction) mais des renseignements, c'est à dire une analyse de ces faits ou d'états potentiels traduits en connaissance d'un état futur.

Ainsi, si au niveau informatique, les événements peuvent se traduire en état O/N, puis au niveau système informatique en probabilité selon des arbres d'attaque, autant dès le niveau Système d'information l'appréciation d'un événement relève d'une analyse plus complexe et généralement humaine (est-ce ou non un facteur de risque ? Est-ce que l'impact est direct ? Est-ce que l'on peut y répondre autrement ? Doit-on réorienter nos capteurs ?) et sa représentation (alarme/alerte) doit s'exprimer autrement : plus intuitive, plus proche d'une traduction en logique floue de type important, critiques, etc.



QUITTER L'INSTANTANE POUR LE CONTINU

Avoir la capacité d'apprécier une situation à un instant précis est indéniablement intéressant, mais s'attendre à ce que le système de surveillance et d'analyse, une fois défini et mis en œuvre, présente de lui-même un état toujours actualisé serait une erreur. Deux impératifs conduisent à faire évoluer ce système :

- L'évolution naturelle des menaces, en écho aux évolutions dans le monde de l'information (capacités, vulnérabilités) qui conduisent à une nouvelle expression des scénarios de risque, voire à un changement des priorités de l'organisme
- L'avancement d'une atteinte à l'organisme, conduisant à des adaptations du périmètre de surveillance (limitation des connexions, cloisonnement des sites) ou un renforcement des réponses avec de nouveaux indicateurs (ex : mise en œuvre d'une journalisation certifiées ou extension des éléments observés).

Ainsi, non seulement les capacités du système d'observation doivent être dynamiques, mais son mode de fonctionnement aussi.

Si l'on en revient à l'expression d'événements redoutés, dont on dit qu'elle devrait conduire la LID et ici l'hypervision, on peut affirmer de même qu'elle est influencée en retour par cette hypervision ; en effet, cette expression doit traduire :

- les nouvelles priorités de l'organisme,
- L'évolution, l'apparition ou disparition de métiers
L'inclusion de SI ou les externalisations
Les failles découvertes
Les attaques réalisées

C'est donc bien un cycle en boucle continue « objectifs de sécurité - hypervision » qui doit être mis en action.

Et si l'on projette ce cycle sur les différents niveaux cités plus haut c'est en réalité une accumulation de cycles imbriqués qu'il faut animer :

- Cycle d'évolution/maintenance du système informatique au travers de l'état de l'art
- Cycle de correction/évolution du système d'information au travers du MCS
- Cycle d'orientation de la LID pour ce qui concerne les fonctions/missions de l'organisme
- Cycle de définition d'un plan de maîtrise de l'information (en référence avec les concepts de guerre de l'information) pour une réelle hypervision.

AVOIR UNE VISION PROSPECTIVE

Cet état de fait amène les constatations et déductions complémentaires suivantes :

- Si l'analyse automatisée du risque ne peut viser actuellement à prédire/prévenir l'occurrence d'un événement redouté mais tente seulement de représenter la possibilité que cet événement advienne, les réponses du niveau hypervision ne peuvent en conséquence être automatisées ;

- Si l'on vise à composer une équipe d'hypervision, celle-ci devrait être composée de tous les acteurs de l'organisme pour traduire la gravité des impacts³¹ (l'arrêt d'une fonction peut être considéré comme *négligeable* par la direction de l'organisme du point de vue de l'inertie de son fonctionnement et des impacts sur la production, mais *grave* pour les personnes en charge de réaliser cette fonction, du point de vue des emplois, de la survie de la sous-traitance, etc.) et selon le modèle de la gestion de crise (Décision, conduite, communication, conseil) pour répondre efficacement et durablement à l'avènement des risques ;
- Si l'on envisage une hypervision, la surveillance ainsi que démontré plus avant, doit prendre en compte les événements de niveau « système » d'un organisme (grèves, changements de politique de l'entreprise, concurrence, géopolitique, etc.), et inclure les interactions et dépendances extérieures de l'organisme ; bien évidemment, une telle surveillance implique dès le niveau système d'information, de nouveaux capteurs (ex : cellule I.E.) et indicateurs puisqu'il faut prendre en compte des sources externes (réseaux personnels, analyse par cabinets).

CONCLUSION

En conclusion, si l'on veut une véritable « hypervision », il faut largement dépasser le paradigme actuel de l'informatique et concevoir cette supervision améliorée, étendue, dans une appréhension systémique de l'organisme, couvrant toutes ses facettes et dépendances internes et externes, et selon l'éclairage du concept de guerre de l'information : son avènement n'est pas a priori pour demain.

³¹ Voir à ce sujet de l'évaluation du risque, l'excellent article de Rey Senjen et Steffen Foss Hansen « Towards a nanorisk appraisal framework » présenté dans les « Comptes-rendus de l'Académie des Sciences » de septembre 2011 (ISSN 1631-0705)

MYTHES ET LEGENDES DE LA GUERRE DANS LE CYBERESPACE

Gérard Peliks, Forum ATENA

D'abord définissons ce qu'on entend par guerre dans le cyber espace ou cyberguerre. Si on sépare les acteurs entre organismes d'état et organismes privés et si on sépare les domaines d'action entre sécurité nationale et économie, la cyberguerre est menée par les organismes d'état, et leurs intermédiaires, et inquiète la sécurité nationale du pays attaqué.

A l'autre bout du spectre, quand des organismes privés, ou des individus autonomes, s'en prennent à l'économie d'un pays ou aux avoirs d'un particulier, on parle de cybercriminalité. Les deux autres espaces, cyber espionnage (état, économie) et cyber terrorisme (organismes privés, sécurité nationale) sont évidemment très proches de la cyberguerre et présentent avec elle, une surface de recouvrement non négligeable.

On entend dire que la troisième guerre mondiale sera une cyberguerre avec une quasi certitude qu'elle se produira dans les années qui viennent. On dit que les bits et les électrons vont remplacer les missiles et les arbalètes. Détrompez-vous ! La cyberguerre n'a pas obligation d'être mondiale et elle a déjà commencé depuis plusieurs années.

Que ce soit, en citant des faits qui ont été particulièrement médiatisés :

- Estonie en 2007 où les botnets ont constitué une arme de perturbation massive ;
- guerre entre la Russie et la Géorgie en 2008 qui a commencé par l'attaque des réseaux de communication ;
- attaque cybernétique de l'usine d'enrichissement d'uranium de Natanz, en Iran, en 2010 où le ver Stuxnet était dans les centrifugeuses ;
- cyber guéguerre entre Marocains et Algériens en 2011 où c'est à qui défigurera le plus de sites Web officiels de l'autre pays en y injectant des messages politiques ;

la cyberguerre n'est pas, loin de là, qu'une vue de l'esprit dans les pensées de quelques experts de la sécurité de l'information. Elle se concrétise dans la réalité, elle sévit autour de vous.

Des pays s'y préparent. Citons Gordon Brown qui fut premier ministre britannique : *"Tout comme au 19eme siècle nous avons eu à sécuriser les mers pour la défense de notre pays et sa prospérité, qu'au 20eme siècle ce furent les cieux qu'il fallut rendre plus sûrs, au 21eme siècle nous prenons place désormais dans le cyber espace"*.

Voilà le décor planté, mais que de mythes et légendes accompagnent déjà aujourd'hui la cyberguerre

MYTHE N° 1 :

UNE CYBERGUERRE NE FAIT QUE DES CYBER MORTS

Allez demander aux Géorgiens, durant l'attaque des Russes en Ossétie du Sud en 2008, s'ils n'ont pas eu de victimes humaines bien réelles, pas des morts d'avatars ou autres constructions cybernétiques ! Et ces victimes auraient sans doute pu être évitées si les systèmes de télécommunication, de transferts d'information et de commandements de la Géorgie avaient fonctionné correctement.

Bien sûr, on n'occupe pas un terrain avec des data qui transitent par les réseaux et s'introduisent dans un système d'information. Mais une guerre moderne commencera par paralyser un pays ; et plus le pays sera dépendant de ses systèmes d'information et de ses réseaux de données, plus vite le chaos s'installera et la panique gagnera la population attaquée. Et le pays tombera sous le joug de l'assaillant comme un fruit mur.

Du chaos et de la panique qui résulteraient (je pourrais sans doute tout aussi bien écrire "qui résulteront" de la cyber attaque préalable, on pourra aussi compter nombre de victimes (et là je n'écris pas de "cyber victimes").

MYTHE N° 2 :

MON ORGANISATION N'EST PAS CONNECTEE A L'INTERNET, ELLE N'A RIEN A CRAINDRE

Comment existe-t-elle alors, votre organisation, dans cette quatrième dimension que constitue le cybermonde, là où de plus en plus d'administrés se rencontrent et où se nouent des relations privilégiées entre les entreprises, leurs partenaires, leurs fournisseurs et leurs clients ?

Mais la question n'est pas là, vous pensez être protégés parce que vous n'êtes pas connectés ? L'usine d'enrichissement d'uranium de Natanz en Iran, site sensible par excellence car le programme nucléaire iranien dépend de sa bonne marche, vous vous en doutez, n'était bien entendu pas connectée à l'Internet. Et pourtant elle a fait l'objet, en 2010 d'une attaque perpétrée par un maliciel très sophistiqué venu d'ailleurs, le ver Stuxnet.

Tout est nominal sur les écrans de la salle de contrôle des centrifugeuses de l'usine de Natanz. Elles tournent à vitesse constante. Un ver, appelé depuis Stuxnet, est introduit au moyen d'une clé USB infectée. Rapidement l'infection se propage sur les ordinateurs sous Windows connectés au réseau interne de l'usine, puis se répand sur les PLC (automates programmables) de Siemens que ces ordinateurs contrôlent. Ces composants Siemens assurent des vitesses de rotation des centrifugeuses nominales et constantes. Ce sont précisément ces contrôleurs que cherche et trouve le ver Stuxnet avant de libérer sa charge létale. Les vitesses de rotation des centrifugeuses, désormais sous contrôle du ver Stuxnet, deviennent alors hors de contrôle de la centrale. Les centrifugeuses accélèrent, décélèrent, leurs axes de rotation vibrent, l'ensemble chauffe. Pendant ce temps, les écrans de la salle de contrôle disent que "tout va bien", car le ver envoie également des informations rassurantes vers les capteurs. La situation réelle est devenue catastrophique. L'attaque n'est pas venue de l'Internet mais le résultat final a été le même.

Connectés ou pas connectés, si vous avez des informations numériques, elles sont en danger. Si vous gérez des infrastructures sensibles de type SCADA, même non connectées, elles peuvent être infectées par un maliciel venant d'une clé USB, comme ce fut le cas pour Stuxnet, ou par un transfert direct par disque dur comme ce fut le cas pour le Pentagone.

MYTHE N° 3 :

JE TRAVAILLE DANS LE DOMAINE DE LA SANTE, JE NE SUIS DONC PAS CONCERNE

Noble secteur que celui de la santé, mais ultra sensible !

Que veulent les cyber agresseurs quand ils attaquent à des fins terroristes ? : Provoquer le chaos en causant une panique généralisée. Quoi de mieux que de s'en prendre aux hôpitaux, aux ambulances ? Les hôpitaux sont équipés de nombreux systèmes commandés par l'informatique qui sont autant de cibles intéressantes. On prend le contrôle des respirateurs,

de la climatisation : panique assurée, cela en parallèle bien sûr avec l'arrêt de la distribution d'électricité, l'empoisonnement de l'eau, la paralysie des transports.

Non, dans une cyberguerre, le secteur de la santé ne sera pas un havre de paix, bien au contraire.

MYTHE N° 4 :

SI ON M'ATTAQUE, JE RIPOSTE PLUS FORT ET L'ATTAQUE S'ARRETE

Si c'est de la légitime défense et si votre riposte n'est pas exagérée par rapport à l'attaque, l'opinion publique ne vous le reprochera pas. Mais au juste, contre qui allez-vous riposter ? Contre une adresse IP qui s'avèrera être celle d'un adolescent d'Europe de l'est ? Contre un pays dont les ordinateurs ont transmis l'attaque à l'insu de leur plein grés, parce que contaminés par des bots, et qui ne savent pas du reste que leurs ordinateurs sont devenus des "zombies" ? Contre la Chine parce que dans l'écosystème de l'insécurité, ce sont toujours les Chinois qui sont les méchants ?

Après avoir déplacé une statue érigée à la gloire du soldat soviétique durant la deuxième guerre mondiale, l'Estonie a été attaquée en déni de service distribué, simultanément par une cinquantaine de pays, peut être aussi par votre ordinateur. L'agresseur a utilisé un réseau de botnet pour déclencher une tempête numérique qui a bloqué les serveurs de plusieurs administrations estoniennes, pendant plusieurs jours.

Contre qui l'Estonie aurait du riposter ? Contre les cinquante pays d'où sont venues les attaques, alors qu'elles ne faisaient que transiter par les ordinateurs infectés de ces pays ? Contre la Russie parce que visiblement, c'était à elle que profitait le crime ? Et si c'était une attaque juste initialisée par un groupe de hackers russes échappant à tout contrôle ?

Oui, décidément il n'est pas facile de reconnaître qui est son adversaire dans le cybermonde, et si on se trompe, non seulement on n'arrête pas l'attaque mais on accumule ses ennemis.

MYTHE N° 5 :

DANS UNE CYBERGUERRE, NOUS SORTIRONS VAINQUEURS

Parce que nous sommes les plus forts ?

Si Sun Tzu, auteur, dans la Chine antique, de l'art de la guerre avait vécu à notre époque, il aurait sans doute écrit que la seule façon de gagner une cyberguerre était de l'éviter. Plus un pays est connecté, donc à priori plus il est fort, plus il est vulnérable. Dans une cyberguerre opposant un pays tel que les Etats-Unis et un pays tel que le Rwanda, à votre avis, qui pourrait faire le plus de mal aux systèmes d'information de l'autre ?

La guerre dans le cyberspace est assurément une guerre asymétrique. Il convient de s'y préparer et de connaître ses adversaires potentiels.

MYTHES ET LEGENDES SUR LA CONFIANCE NUMERIQUE

Jean-Baptiste FAHY

La confiance est un élément essentiel des activités numériques. Elle est composée d'éléments, des maillons que l'on peut comparer à ceux d'une chaîne.

Parmi ces éléments, le certificat a un rôle important.

Un certificat accorde un droit à son titulaire : par exemple, l'accès à des données, à un système informatique ; ou le droit de signer un document, un message ; ou de chiffrer et déchiffrer des informations, des documents, des messages. Un certificat a une période de validité (dates de début et fin), et peut être révoqué. Il est (un peu) comparable à un visa apposé sur un passeport, qui donne droit d'entrer dans un pays, pendant une période donnée.

Un certificat est toujours lié un à bi-clé (une clé publique et la clé privée correspondante). La clé publique est partie intégrante du certificat, qui est lui-même un élément public, c'est à dire publiable sans inconvénient pour son titulaire. La clé privée doit être conservée et gardée secrète par le titulaire, qui en est par principe le propriétaire. Il peut pour cela utiliser une carte à puce ou tout élément physique, ou utiliser un élément logiciel tel qu'un ordinateur, un smartphone ou une tablette.

Un certificat est signé par son émetteur (l'Administration, une banque, un commerçant, ...) au moyen de la clé privée associée à un certificat signataire, appelé parfois improprement certificat racine.

La signature d'un certificat permet de vérifier son authenticité, en déchiffrant cette signature à l'aide de la clé publique du certificat signataire.

Lorsqu'un titulaire veut utiliser son certificat, par exemple pour s'authentifier, il le présente à l'entité voulue (site web, système d'information, ...). Celle-ci contrôle son authenticité et sa validité, puis émet un aléa, qu'elle chiffre avec la clé publique de ce certificat. L'aléa chiffré est envoyé vers le titulaire, qui le déchiffre avec la clé privée liée au certificat, et retourne le résultat à l'entité accédée. Enfin, celle-ci compare la donnée reçue avec l'aléa. S'il y a égalité, c'est la preuve que la personne qui tente de se connecter dispose bien de la clé privée liée au certificat.

Inversement, un utilisateur de service informatique, qui accède par exemple à un site web, peut vérifier que le certificat du site web est bien valide (bonnes dates et non révoqué) et authentique (signature vérifiable avec le certificat signataire du site Web connu par le navigateur de l'utilisateur).

MYTHE N° 1 :

LE CERTIFICAT INSCRIT SUR MA CARTE A PUCE NE PEUT ETRE UTILISE QUE SI JE FOURNIS MON CODE PIN

Vrai et faux

La carte à puce que j'utilise m'a été fournie par un organisme : mon entreprise, ma banque, un tiers, ... peu importe, et mon certificat a été inscrit sur cette carte.

Deux cas sont possibles :

1. La clé privée liée au certificat a été générée par la carte ;

2. elle a été générée centralement par l'organisme et inscrite sur la carte.

Dans le cas 1, on a la certitude que si le certificat est utilisé, le code PIN a été fourni à la carte. Encore faut-il que le certificat ait été utilisé APRÈS que la carte a été remise à son titulaire. D'où la nécessité d'horodater les transactions.

Dans le cas 2, il est techniquement possible de dupliquer le certificat et sa clé privée, et de l'inscrire sur tout support, matériel ou logiciel. Et donc de l'utiliser.

En conclusion, pour être sûr que, si le certificat inscrit sur ma carte à puce n'a pu être utilisé qu'après que j'ai fourni le code PIN, il faut :

- a) que je fasse confiance au système d'horodatage (qui permettra de savoir à quel moment exactement la transaction a eu lieu)
- b) que je fasse confiance à l'organisme pour prouver à quel moment la carte m'a été remise
- c) que je fasse confiance à l'organisme pour que ce soit la carte à puce qui génère la clé privée du certificat
- d) que je fasse confiance à la carte à puce, qui ne doit pas être capable d'exporter une clé privée

On le voit, la confiance dans l'utilisation de certificat sur carte à puce est le résultat de plusieurs éléments.

Note : certaines cartes à puce peuvent indiquer si la clé privée d'un certificat a été générée par la puce, ou extérieurement. Si l'indicateur « générée par la puce » est positionné, je peux avoir confiance dans le point c)... si je fais confiance à la carte à puce.

MYTHE N° 2 :

JE PEUX AVOIR CONFIANCE EN UNE TRANSACTION SIGNÉE AVEC UN CERTIFICAT QUI EMANE D'UNE SOCIÉTÉ AYANT PIGNON SUR RUE (VERISIGN, CERTINOMIS, ...)

Vrai... sous réserve

Sous réserve que l'usage de la clé privée liée au certificat n'est pas usurpé au moment de la transaction, et que le certificat est bien valide. Donc sous réserve que j'ai bien vérifié qu'il n'est pas révoqué. Et encore...

- Il faut qu'en cas de vol, le titulaire l'ait signalé au service compétent ;
- que le service compétent ait mis à jour la liste des certificats révoqués ;
- que la liste à jour des certificats révoqués soit disponible au moment de la transaction.

Ces actions doivent être décrites dans des procédures qui résultent de la politique de sécurité de l'organisme qui gère le certificat, et ces procédures doivent être appliquées.

Conclusion : si le certificat est volé, mais que le vol n'est pas déclaré, ou bien s'il y a du retard dans la mise à jour de la liste des certificats révoqués, ou si la liste mise à jour n'est pas encore disponible, votre confiance peut être mal placée. C'est peu probable, mais c'est possible.

C'est pour cela que certains organismes prennent des assurances sur les transactions : en cas de vol, la transaction n'est pas révoquée tant que le vol n'a pas été déclaré, le titulaire reste responsable.

MYTHE N° 3 :

JE PEUX FAIRE CONFIANCE AU SITE WEB QUE JE CONSULTE AVEC MON NAVIGATEUR PARCE QU'UN CADENAS FERME APPARAÎT DANS LA BARRE D'ADRESSE

Faux

Outre le fait qu'une connexion en SSL, matérialisée par l'icône du cadenas, ne garantit pas de sécurité (cf Tome 1, Mythes et légendes de la navigation sur internet), la confiance envers le site consulté ne va pas de soi.

La seule implication de la présence de ce cadenas est que le site consulté utilise un certificat dit « certificat d'authentification de serveur », délivré par un organisme.

Si l'organisme en question a pignon sur rue (Verisign, Certinomis, ...), il figure dans la liste des autorités de confiance de votre navigateur : avec Firefox, voir dans Outils, Options, Avancé, onglet Chiffrement, bouton Afficher les certificats, onglet Autorités. On y trouve entre autres 2 certificats signataires de Certinomis, et plusieurs certificats signataires de Verisign.

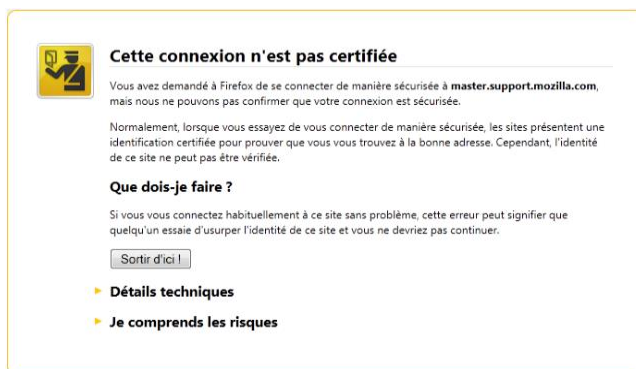
Lorsque vous vous connectez à un site web qui utilise une connexion SSL, donc un certificat de serveur web, votre navigateur vérifie que ce certificat a été signé par un certificat signataire présent dans sa liste des autorités de confiance. C'est généralement le cas (voir le cas inverse plus bas). Il vérifie ensuite la validité du certificat (dates et non révocation), et enfin que le site web possède la clé privée liée au certificat de site web. Si tout va bien, il y a affichage du cadenas, et vous pouvez visiter le site web à votre guise.

La seule chose dont vous êtes sûr, c'est que le certificat a bien été délivré par une autorité de confiance connue de votre navigateur, et que ce certificat est valide d'après cette autorité.

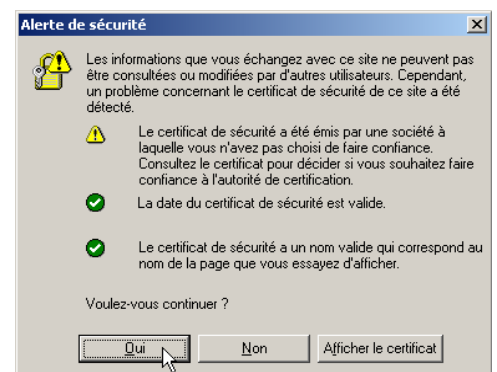
La confiance existe si :

- Vous faites confiance à l'autorité qui a émis le certificat de site web : elle a dû vérifier les dires de la société qui a demandé ce certificat, elle ne délivre pas de certificat à n'importe qui, elle traite les révocations en temps et en heure ;
- vous faites confiance au site web de l'organisme qui a acquis le certificat de serveur web.

Voyons maintenant le cas où le certificat a été signé par un certificat qui ne figure pas dans la liste des autorités de confiance : lorsque vous vous connectez à un tel site, votre navigateur émet un message d'avertissement, tel que



ou bien



À vous de décider si vous pouvez faire confiance à ce site web.

Si vous acceptez, vous lui faites momentanément confiance. Lors d'une prochaine visite, le même message d'avertissement apparaîtra.

S'il est possible de télécharger le certificat signataire du certificat de site web, vous pouvez l'incorporer dans la liste des autorités de confiance. Mais attention, ce n'est pas un acte anodin : le message d'avertissement évoqué ci-dessus ne sera plus présenté lorsque vous visiterez n'importe quel site ayant un certificat signé par ce certificat signataire, ce qui peut entraîner une confiance implicite.

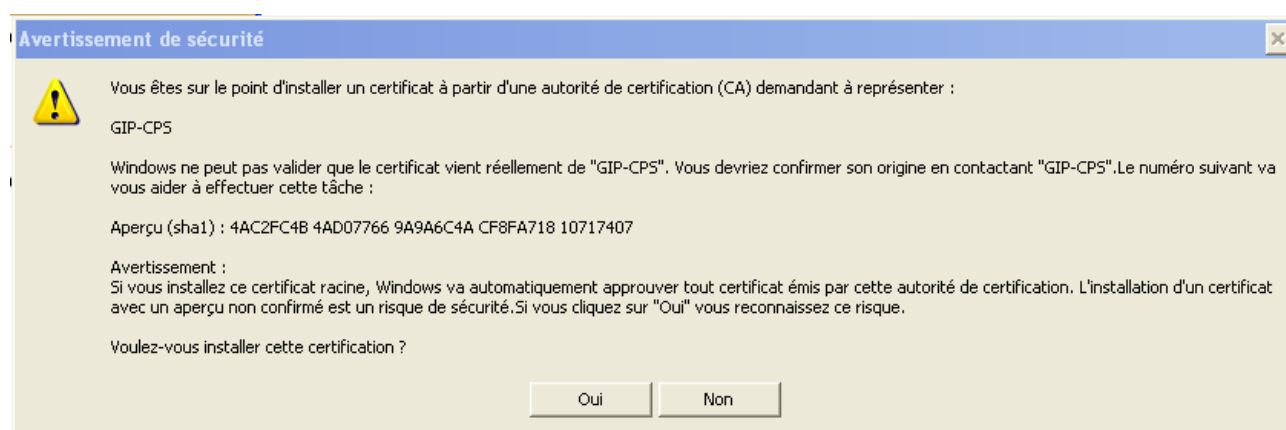
En conclusion, lors de la visite de sites web en connexion sécurisée par certificat, et qu'il n'y a pas de message d'alerte, les seules choses qu'on puisse affirmer sont :

le site web utilise un certificat valide

ce certificat a été émis par une autorité reconnue par votre navigateur

Note :

lorsqu'on ajoute un certificat signataire de certificat de site web à la liste des autorités de confiance d'internet explorer, il s'affiche un message tel celui-ci :



MYTHE N° 4 :

MA SOCIÉTÉ EXPLOITE SA PROPRE PKI, JE PEUX DONC AVOIR ENTIÈRE CONFIANCE DANS LES CERTIFICATS QU'ELLE GÉNÈRE

Vrai... sous conditions

La délivrance de certificats se fait grâce à de nombreuses procédures qui découlent de la politique de sécurité.

Une demande de certificat est présentée à un service, le service d'enregistrement, qui est chargé de vérifier d'une part l'identité du futur titulaire, et d'autre part ses droits. En effet, un certificat donne un certain nombre de droits : accès à un site web ou à un système informatique, signature de documents, ...

Ce service d'enregistrement effectue ensuite une demande de certificat auprès de la PKI, qui signe le certificat. Le certificat peut alors être remis à son titulaire.

Toutes les procédures qui permettent la délivrance de certificats doivent permettre une sécurité optimale. La majorité du travail est le fait de personnes : opérateurs et administrateurs, ... Seulement une petite partie du travail est effectuée par des programmes informatiques : génération de clés, chiffrement, signatures, ...

De plus, lors de la création de la PKI, un bi-clé a été généré, le bi-clé correspondant au certificat racine de la PKI. Sa clé privée est l'élément le plus critique : si elle est divulguée, la confiance dans la PKI tombe instantanément.

En conclusion, pour avoir confiance dans les certificats générés par la PKI de sa société, il faut :

- Que la clé privée correspondant à son certificat racine ait été générée de manière sûre (par exemple lors d'une « cérémonie des clés », ou la clé privée générée est découpée en plusieurs morceaux, répartis entre les participants à la cérémonie ; pour reconstituer cette clé privée, les participants devront réunir leurs informations) ;
- que les participants à la cérémonie des clés ne divulguent pas la clé privée ;
- que le système de gestion de la PKI ne permette pas de divulguer la clé privée, même aux administrateurs ;
- que les administrateurs et opérateurs de la PKI soient honnêtes (par exemple, ils doivent suivre à la lettre les procédures issues de la politique de sécurité) ;
- que les procédures soient solides.

Il faut également gérer la révocation des certificats suite à vol, perte ou utilisation frauduleuse de leurs clés privées.

MYTHES ET LEGENDES DE LA CONFIANCE DANS LA PROTECTION DES DONNEES

Jean-Baptiste FAHY

Les données qui nous concernent, que ce soit en tant qu'individu ou en tant que personne morale, sont de deux sortes :

- celles gérées par des tiers : les banques, fournisseurs, clients, administrations, ...
- celles gérées par nous : agenda, courriers, stock, comptabilité, et surtout les documents stratégiques tels que les plans, les tarifs, ...

Ces données sont-elles bien gardées ? Sont-elles à l'abri de mes concurrents et autres ennemis ?

MYTHE N° 1 :

LES DONNEES QUI ME CONCERNENT ET QUI SONT GEREES PAR DES TIERS SONT EN SECURITE, PUISQU'IL FAUT MONTRER PATTE BLANCHE POUR Y ACCEDER.

On utilise en effet des VPN ou des accès web sécurisés pour accéder à nos données bancaires, nos commandes chez nos fournisseurs, nos données administratives, ...

Mais la seule garantie des VPN et des accès "en HTTPS", c'est que vos données sont bien à l'abri... pendant le transport entre le site tiers et vous. et encore, il faut que vos crédentils (mots de passe, cartes à puce, ...) ne soient pas utilisés frauduleusement.

Et de fait, les données sont en clair à chaque bout de la liaison VPN ou https.

Voyons alors la situation à chaque extrémité :

- de votre côté, à vous de protéger les données que vous consultez ou que vous rapatriez dans un fichier. Méfiez-vous de ceux qui pourraient regarder pas dessus votre épaule, ou qui auraient la possibilité de fouiller dans vos fichiers temporaires ou archivés sur votre ordinateur ou sur un disque du réseau local. Ici, le problème est simple, la solution ne l'est pas forcément, mais c'est votre problème et vous devez pouvoir le gérer.
- du côté du tiers, la situation est bien différente. Vous êtes bien obligé de lui faire confiance. Mais un administrateur ne peut-il pas consulter les bases de données ? ou plus simplement, n'importe quel conseiller bancaire ne peut-il pas accéder à votre relevé de compte et en faire un usage frauduleux ? En fait, vous ne pouvez être sûr de rien.

Pour que la confiance soit établie, il faut premièrement que le tiers qui gère vos données mette en place les procédures qui permettront d'éviter l'accès indu aux données, deuxièmement qu'il mette en place les procédures qui permettront de constater les manquements, et enfin la transparence, qui permet à chacun de savoir s'il y a eu accès à ses données.

Notons qu'une série de procédures semblables peut s'appliquer dans votre propre entreprise, donc pour votre côté du problème.

Et on peut rêver de la « Société Idéale », où ces procédures existeraient, seraient connues de tous, et où chacun pourrait vérifier a posteriori qu'elles sont appliquées, et dans la négative, les contrevenants risqueraient des sanctions.

MYTHE N° 2 :

LES DONNÉES QUE JE GERE SONT EN SECURITE MEME QUAND ELLES SONT HEBERGEES SUR DES SERVEURS DISTANTS PARCE QUE PERSONNE NE PEUT LES INTERCEPTER.

On suppose que la sécurité de vos informations est assurée chez vous, que ce soit votre domicile ou votre entreprise.

Pour la même raison que plus haut, vos informations sont à l'abri lorsqu'elles circulent dans un VPN ou en connexion https. Mais à l'autre bout, elles sont de nouveau en clair, ne fut-ce qu'un instant, avant stockage. Une personne mal intentionnée peut donc les intercepter au moment où elles sortent du VPN ou de la connexion https. Si en plus, elles ne sont pas stockées chiffrées sur le site distant, elles sont visibles de tout administrateur de ce site.

L'idéal serait de chiffrer les données avant la sortie de votre machine, et de les laisser ainsi sur les sites distants. Malheureusement, si cette méthode ne pose pas de difficulté pour les fichiers plats, elle est aujourd'hui très complexe pour les éléments mettant à jour une base de données distante. On nous annonce des techniques de chiffrement qui devraient bientôt résoudre ce problème (le chiffrement homomorphique), mais ce n'est pas pour aujourd'hui.

Vous êtes donc contraint à faire confiance à votre prestataire. Mais comme il s'agit de votre prestataire, vous pouvez lui imposer des procédures et la transparence.

Sauf bien sûr s'il s'agit de prestataires étatsuniens : le « USA patriot act » les oblige à fournir toute donnée sur simple requête de l'administration fédérale étatsunienne.

Un exemple : si vous utilisez un smartphone et/ou une tablette sous Android, il y a de fortes chances que votre agenda soit lié à votre compte Google. Dès que vous créez un événement, il est recopié sur le serveur de Google, puis sur vos éventuels autres dispositifs mobiles. Ces informations sont protégées d'un accès de l'extérieur par un mot de passe. Mais elles sont accessibles par le FBI.

Que les utilisateurs d'iPhones ne s'inquiètent pas : c'est la même chose avec iCloud, le service fourni par Apple. Quant à ceux qui préfèrent un Blackberry, il y a de fortes chances que leurs données soient gérées sur des serveurs étatsuniens, même si leur fabricant, RIM, est canadien.

Et bien sûr, ça ne s'arrête pas aux événements de votre agenda, cela concerne également vos contacts, éventuellement vos photos ou tout document que vous aurez confiés à Google ou Apple.

Il existe une solution, toute théorique aujourd'hui, pour éviter l'accès à vos données par des tiers : chiffrez vos données avant de les envoyer vers Google ou Apple.

Pour conclure, le seul moyen de protéger les données que vous gérez quand elles sont hébergées sur des serveurs, c'est de les chiffrer avant qu'elles sortent de l'environnement que vous contrôlez. Mais ce n'est pas toujours simple.

3° PARTIE : ASPECTS PHYSIQUES



MYTHES ET LEGENDES DES DANGERS SANITAIRES DE LA TELEPHONIE MOBILE

Viken TORAMANIAN, Orange

Les sujets de la santé publique liés à la téléphonie mobile sont une préoccupation croissante pour l'ensemble des populations. Tout le monde sait que les antennes relais ainsi que les téléphones portables produisent un champ électromagnétique couramment appelé une onde. Ces ondes émises et reçues par les terminaux mobiles et les antennes sont le support de l'information qui permet ainsi une communication non liée par un élément physique. Quelles sont les conséquences des ondes sur l'organisme humain, voire sur l'environnement des vivants ? Le sujet anime beaucoup de débats et rentre dans les considérations globales de l'homme « moderne ».

Aujourd'hui les antennes qui permettent l'usage de la téléphonie mobile sont très fortement déployées et les conséquences de « troubles » liés à ces dispositifs pourraient être catastrophiques s'ils existaient. Il faut s'attendre à une réorganisation de la société pour la préserver, si les nuisances sanitaires sont avérées. Dans le monde professionnel, public et privé, qui peut actuellement se passer des téléphones portables ? Qui organiserait le démantèlement des équipements radios installés dans le monde entier ? Quels moyens de remplacement disposerions-nous ? Comment communiquer cette information aux populations afin qu'elles changent leurs usages de l'utilisation du téléphone portable ? Qui paierait le prix fort de ce changement ? Nous avons le script idéal pour lancer un film catastrophe !

Les études sanitaires sont menées sur les dangers possibles de l'utilisation du téléphone portable et les systèmes qui le font fonctionner. Nous allons comprendre quelles sont les conséquences de l'usage de la téléphonie mobile. *(Achévé de rédiger en février 2012).*

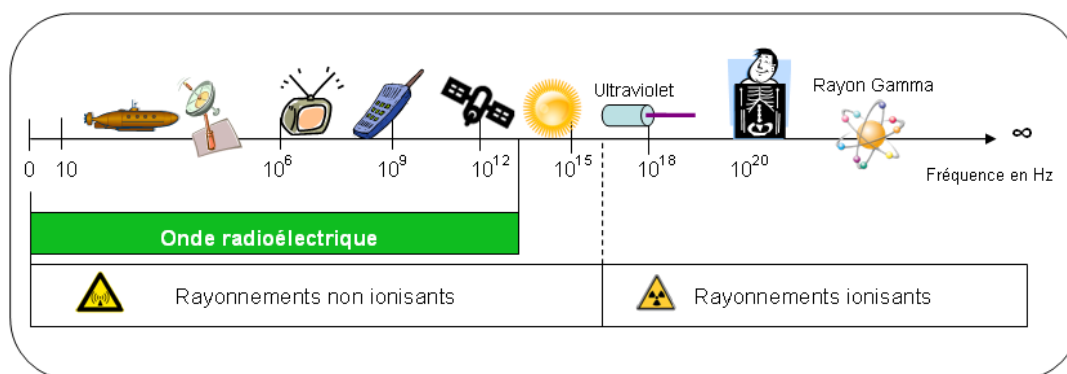
MYTHE N° 1 :

LES CHAMPS ELECTROMAGNETIQUES PRODUITS PAR LES MOBILES ET LES ANTENNES RELAIS SONT NOCIFS POUR L'HOMME :

Pour répondre à cette question nous nous appuyons sur des organismes objectifs dont le but est de préserver la santé de tous. Nous tirons nos conclusions de l'Organisation Mondiale de la Santé, de l'Académie de médecine ou encore de l'ANSES (Agence Nationale de Sécurité Sanitaire de l'alimentation, de l'environnement et du travail). Nous ne parlerons pas des effets sociaux d'isolement ou d'hyperactivité liés aux nouvelles technologies, dont la téléphonie mobile, et qui ont fait évoluer les comportements sociaux. Ainsi la réponse doit être faite en deux temps.

Premièrement pour une utilisation à court terme, il est faux d'affirmer qu'il y a un danger pour l'homme. Cette conclusion est vraie pour une utilisation en contact ou non d'un équipement de téléphonie mobile en fonctionnement ou non. La notion de « court terme » est variable selon les études : il peut s'agir d'une utilisation de quelques minutes sur un jour à quelques heures pour une semaine. Plusieurs expériences d'organismes indépendants ont démontré qu'il n'y avait aucun risque direct chez l'homme. Un des travaux de l'ANSES a été de regrouper ces expériences et de synthétiser les résultats des analyses.

Le champ électrique d'un mobile émis vers une cellule vivante est non destructeur : l'onde émise dans les télécommunications sans fils est dans la catégorie des ondes non-ionisantes. Contrairement aux rayonnements ionisants, elles ne détruisent pas les cellules vivantes.



Les expériences montrent que les effets des ondes existent sur les cellules vivantes, mais elles sont bénignes pour l'homme. Les conséquences connues actuellement sont l'augmentation de un à deux degrés des tissus pour une utilisation intense, il s'agit d'une propriété des ondes non-ionisantes sur les cellules. L'effet micro-onde, également une onde non-ionisante et qui agite les molécules d'eau, ne peut pas avoir de conséquences sur le corps humain, car les puissances émises par les mobiles (inférieures à 2W pour les vieux mobiles et inférieures à 1W pour la majorité des mobiles actuels) sont 1000 fois inférieures en puissance pour permettre à la molécule d'eau de s'agiter. Le risque sanitaire par cette piste est à exclure.

Par contre, il ne faut pas oublier qu'une partie de la population perçoit des effets lorsqu'elle est exposée aux champs électromagnétiques de nos portables. Il s'agit d'une intolérance environnementale idiopathique, communément appelée hypersensibilité électromagnétique (HSE). Les patients ressentent des malaises, des démangeaisons ou encore des vertiges lorsqu'ils sont proches de sources radio. Les causes de ces symptômes ne sont pas encore clairement définies et les recherches, qui sont difficiles à mener du fait des multitudes d'interactions qui peuvent être subjectives, sont en cours. Il s'agit là de cas rares mais reconnus.

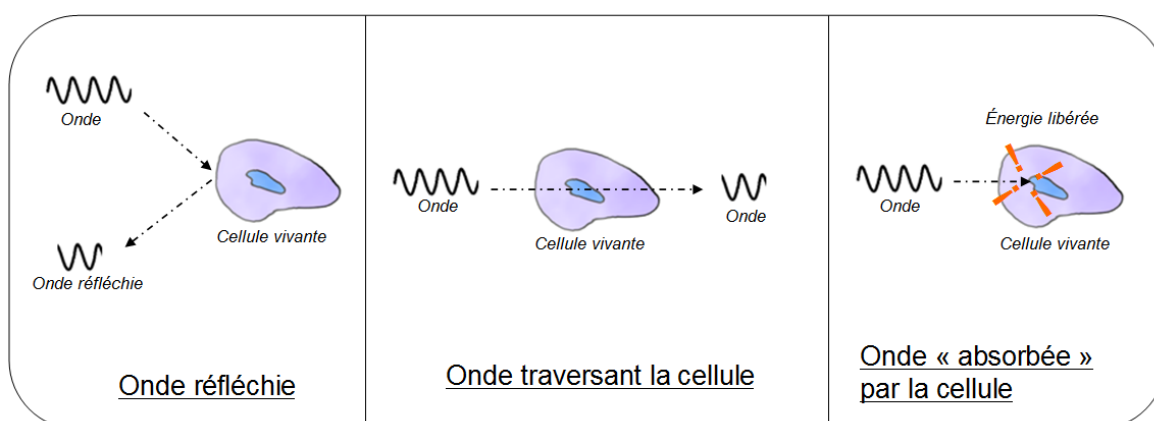
Deuxièmement sur le long terme (10 ans, voire 20 ans), les expériences entreprises actuellement n'ont rien prouvé. Il est aussi faux de dire que l'usage du mobile est dangereux que de dire qu'il n'y a aucun risque sur l'usage d'un mobile sur le long terme. Il est extrêmement coûteux, long et difficile de mener des études sur les dangers des réseaux mobiles. D'une part, les observations doivent être effectuées sur plusieurs années. Les cas de cancers doivent être recensés, suivis et les statistiques doivent déboucher sur des hypothèses valides. D'autre part, les technologies évoluent très rapidement et rendent difficiles l'observation des conséquences sur le long terme. Le passage de l'utilisation des réseaux 2G aux réseaux 2G-3G, qui n'utilisent pas les mêmes fréquences, ni les mêmes protocoles, perturbent le suivi des expériences en cours. Les expériences menées jusqu'à aujourd'hui ne nous permettent pas de conclure sur la nuisance ou non sur l'usage de la téléphonie mobile.

MYTHE N° 2 :

LES ETRES VIVANTS RESSENTENT TOUS DE FAÇON EGALITAIRE LES CHAMPS ELECTROMAGNETIQUES

Faux : les ondes électromagnétiques sont des déplacements d'énergie dans l'espace. Il ne faut pas les confondre avec les ondes mécaniques, qui sont la vibration de particules physiques et que nous percevons à certaines fréquences de vibration : c'est le son (20Hz à 20 000Hz).

Les cellules vivantes réagissent aux ondes. Il a été prouvé que certains animaux se guident grâce aux champs magnétiques terrestres (ex : oiseaux migrateurs, certaines tortues...). Nous-mêmes nous percevons les couleurs, qui sont également une onde électromagnétique très haute fréquence (350 à 750 THz). Les ondes dites non-ionisantes ne détruisent pas les cellules vivantes, mais elles peuvent y pénétrer. Trois phénomènes se produisent. Premièrement une partie de l'énergie de l'onde est absorbée, ce qui atténue la puissance de l'onde et sa propagation. Deuxièmement, l'énergie restante de l'onde continue sa propagation et traverse la cellule. Enfin, une partie de l'onde émise sur la cellule ne la traverse pas, mais est directement réfléchiée dans une autre direction.



Une partie des ondes des téléphones mobiles nous traverse, ou du moins essaie de nous traverser. La peau des êtres vivants protège le corps des ondes. Une partie de l'énergie de l'onde qui percute une cellule est dissipée. Par l'effet Joule, cette dissipation se transforme en chaleur. Mais, il n'y a aucune inquiétude à avoir, cette chaleur est négligeable. Des tests montrent que l'élévation température est inférieure à 1°C pour plusieurs heures d'utilisation de son mobile. Il n'y a pas de quoi faire cuire un œuf ! L'effet du soleil sur l'homme a une conséquence thermique plus grande que celle des ondes électromagnétiques.

Aussi pour revenir à notre question, nous savons que le tissu de protection naturelle qu'est notre peau, est différent selon les parties du corps. De même entre plusieurs individus, les constitutions corporelles et cellulaires peuvent varier. Cette variation pourrait expliquer l'hypersensibilité de certains individus face aux ondes électromagnétiques.

MYTHE N° 3 :

IL N'EXISTE AUCUNE LEGISLATION, NI NORME DE PRECAUTION CONTRE L'USAGE DES RESEAUX MOBILES.

Faux : à plusieurs niveaux, locaux, régionaux et même internationaux, des précautions et des normes sont mises en place pour limiter la puissance d'émission des champs électromagnétiques. Ces règles ont d'abord été mises en place pour empêcher les interactions

de brouillage qui peuvent exister entre machines émettrices et réceptrices. Par précaution sanitaire, ces normes limitent les puissances émises comme nous allons le voir. Le respect de ces normes est obligatoire pour mettre sur le marché des nouveaux produits ou utiliser un service dans les lieux publics. Voici différentes obligations appliquées dans le monde :

Au niveau international :

L'IEEE (IEEE Std C95.1 2005) et l'organisation non gouvernementale CIPRNI (Commission internationale de protection contre les rayonnements non ionisants) prévoient des limitations quasiment équivalentes ; limitation à $4,5\text{W/m}^2$ pour les fréquences 900MHz et 9W/m^2 pour les fréquences à 1800MHz.

Au niveau régional :

Dans l'article 1999/519/CE, l'Europe conseille le respect de la limitation du DAS à un maximum de 2W/Kg . Il demande également d'avoir une puissance du champ électromagnétique inférieure à 41,25 volts/mètre, 58,33 volts/mètre et 61 volts/mètre pour respectivement le GSM (900MHz), le DCS (1800 MHz) et l'UMTS (2100MHz). Enfin l'Europe préconise d'éviter les expositions des champs l'électromagnétique aux individus fragiles.

Dans une grande partie des pays de l'ex-URSS la norme prévoit le respect à une limitation de $0,1\text{ W/m}^2$.

Au niveau national :

Chaque pays a mis en place des règles. Si la plupart des pays demande que les recommandations de l'IEEE, de l'Europe, de la CIPRNI... soient respectées, certains pays ont leur particularité. Par exemple nous retiendrons que la Chine ou la Suisse prévoient le respect d'une limitation du champ d'émission inférieure à $0,1\text{ W/m}^2$. Quant à la France, elle impose pour les équipements de radio-télécommunication un DAS local inférieur à 2W/Kg .

Au niveau local :

Plusieurs villes imposent aux opérateurs des limites particulières. À Paris, les opérateurs se sont engagés à ne pas dépasser les 2 V/m (moyenne sur 24h). Dans la commune de Varades (Loire Atlantique), les opérateurs s'engagent à ne pas dépasser la valeur de $0,6\text{ V/m}$. A Bordeaux, les opérateurs ont interdiction de déployer une antenne à moins de 100 m d'un établissement scolaire (enfant de moins de 12 ans)

Nous verrons plus loin la signification des différentes unités utilisées ci-dessus.

MYTHE N° 4 :

IL EXISTE UN MOYEN DE MESURER LE RISQUE DE DANGEROUSITE DE L'UTILISATION D'UN EQUIPEMENT RADIO :

Vrai : les champs électromagnétiques ont sur l'homme des effets en fonction de leur puissance et de leur fréquence. Pour définir des seuils, les organismes utilisent différentes unités liées aux champs électromagnétiques.

L'indice le plus utilisé est le DAS (Débit d'Absorption Spécifique) qui est la quantité d'énergie émise par un mobile - ou équipement électromagnétique – vers un utilisateur. L'unité est le Watt par kg. Le SAR (Specific Absorption Rate) est la traduction en anglais du DAS. Tous les téléphones mobiles indiquent l'indice de DAS. La valeur seuil appliquée en France est de 2 W émis par kilogramme et moyennée sur 10 cm^2 de tissu humain.

Comme il est difficile de mesurer le DAS sur une personne face à un équipement mobile, une autre unité est couramment utilisée : le volt par mètre ou V/m. Un simple spectromètre permet de mesurer sa valeur.

L'unité volt par mètre (V/m) est la force électrique mesurée en un point de l'espace.

L'unité ampère par mètre (A/m) est la force magnétique mesurée en un point de l'espace.

Le Watt par mètre carré (W/m²) est la densité surfacique de puissance. Il est le résultat du produit du champ électrique par le champ magnétique.

Par hypothèse, plus les forces électromagnétiques sont élevées et plus les risques sur l'homme augmentent.

MYTHE N° 5 :

AUCUNE ETUDE N'A ETE MENEES A CE JOUR SUR LES DANGERS DE L'UTILISATION INTENSIVE DES MOBILES :

Faux, plusieurs études existent sur le sujet et ont permis de donner les conclusions que nous avons citées : pas de danger pour une utilisation sur le cours terme et aucune conclusion sur le long terme ne peuvent être faites sur l'utilisation du mobile.

Sur les différentes études sérieuses que nous pouvons retenir, les plus abouties sont celles du CIRC (projet Interphone), de la Commission Internationale de Protection contre les Rayonnements non Ionisants (ICNIRP) ou celle de l'Affet. Prenons quelques exemples d'études.

Tout d'abord l'étude Interphone est une étude multinationale qui s'est déroulée sur 10 ans (2000-2010) et qui a réuni treize pays : Australie, Canada, Danemark, Finlande, France, Allemagne, Israël, Italie, Japon, Nouvelle Zélande, Norvège, Suède, et Royaume-Uni. Les études ont consisté à auditer des personnes atteintes de cancer et à relever leurs habitudes en termes d'utilisation de la téléphonie mobile. Aucun lien n'a été défini avec certitude entre le cancer des sujets et l'utilisation du téléphone.

Dans les mêmes cas, l'étude Schuz (2006) a essayé de trouver un rapport entre les expositions d'une population aux antennes relais et l'augmentation du risque de cancer. L'étude conclut qu'il n'y a pas de lien entre l'augmentation de risque d'avoir un cancer et habiter en face d'une station de base.

L'Affet - Agence française de sécurité sanitaire de l'environnement et du travail - a également mené une analyse bibliographique sur les travaux scientifiques menés sur les possibles dangers « électromagnétiques » des mobiles sur l'homme (2008-2009). Plus de 3500 références et travaux sur les conséquences des ondes sur les cellules vivantes ont été sélectionnés pour cette étude. L'analyse des expérimentations ne permet pas de conclure sur un risque de cancer plus important lors de l'utilisation du téléphone portable.

MYTHE N° 6 :

IL N'EXISTE AUCUN MOYEN DE LIMITER SON EXPOSITION AUX CHAMPS ELECTROMAGNETIQUES ET A SES POSSIBLES EFFETS NEGATIFS SUR LA SANTE :

Faux : chacun est libre de se protéger si telle est sa crainte et des moyens techniques existent. Les techniques de limitations aux champs électriques sont multiples et recommandées par principe de précaution par une multitude de sociétés privées et publiques (l'OMS, l'Affet, l'ICNIRP, opérateurs mobiles...). Voici les principaux conseils :

- Limiter le nombre d'appels et la durée pour chaque appel.

- Utiliser un kit main libre, afin d'éloigner de son cerveau les ondes émises par son mobile.
- Eviter d'utiliser son mobile lorsque ce dernier ne capte pas correctement (une ou deux barres sur l'indicateur) : les mobiles émettent plus fortement lorsqu'ils sont en limite de couverture.
- Couper sa borne Wi-Fi lorsque celle-ci n'est pas utilisée. De même couper le mode Wi-Fi et Bluetooth de son téléphone mobile, il économisera vos batteries.

Si vous souhaitez plus d'information, le site de <http://www.lesondesmobiles.fr/> est peut être le plus exhaustif sur le sujet.

MYTHE N° 7 :

L'UTILISATION DU TELEPHONE MOBILE N'A AUCUN EFFET SANITAIRE SUR LA SOCIETE ACTUELLE

Faux : rien ne prouve qu'il n'y a pas d'effet sanitaire sur l'être humain. Les principes de précaution sont là pour nous le rappeler. Les recherches continuent et s'orientent toujours vers la possibilité d'avoir un risque de cancer plus important pour l'homme. Pour rappeler que rien n'est prouvé, l'OMS a classé les ondes électromagnétiques dans la catégorie des cancérogènes possibles pour l'homme (groupe 2B).

De plus les téléphones portables ont apporté à notre société une évolution socioculturelle qui peut nuire à la santé de l'être humain : dépendance, isolement de l'individu, violence, hyperactivité... Ces effets ne sont pas liés aux ondes électromagnétiques, mais bien à la « mutation » de la société - ou à son comportement - face à l'évolution technologique.

Enfin, l'OMS nous rappelle que le téléphone tue : l'usage du mobile en voiture multiplie le risque d'avoir un accident par trois ou quatre. Il faut donc savoir se contrôler et s'interdire toute utilisation de son téléphone au volant, ajouté aux risques dans les stations-services et autres milieux explosifs : micro étincelles dégagées par la batterie essentiellement pendant les phases de sonnerie.

MYTHE N° 8 :

LES REACTIONS DES ONDES ELECTROMAGNETIQUES NE CONCERNENT PAS L'HOMME

Faux : les équipements entre eux fonctionnent sans qu'ils s'auto-perturbent. Aussi tous les équipements émetteurs et récepteurs doivent être certifiés avant leur commercialisation. Cette certification consiste à passer un ensemble de tests électromagnétiques : émission de champs à différentes puissances et fréquences et vérification de la non-réactivité de l'équipement en cours de certification. De même ces tests vont vérifier que les champs émis par ces équipements ne dépassent pas des valeurs seuils réglementaires.

Par principe de précaution l'utilisation d'un mobile dans les lieux sensibles est interdite. Vous verrez souvent des signalisations vous interdisant l'utilisation de votre mobile dans les aéroports, les avions, les hôpitaux ou encore dans les stations service. L'interdiction de l'utilisation de son mobile se justifie par rapport à la sensibilité – sensibilité vis-à-vis de la fonction principale de l'équipement - des équipements présents dans ces endroits. Il n'y a ici pas d'impact direct sur la santé, mais les interférences que pourraient créer ces mobiles peuvent être la cause de dérèglement des équipements présents et cela même si la compatibilité électromagnétique avait été validée précédemment. Le principe de précaution est cette fois encore appliqué.

4° PARTIE : ASPECTS METIERS



MYTHES ET LEGENDES DE LA RESPONSABILITÉ SOCIÉTALE DE L'ENTREPRISE INFORMATIQUE QUI GÉNÈRE DE LA VALEUR

Jean-Marie CORRIERE, manager expert en transition des organisations

INTRODUCTION

Si être informaticien c'est être exclusivement tourné vers le système d'information, alors je ne suis pas cet informaticien. Les technologies de l'information et de la communication sont un incroyable vecteur d'innovation et de création de valeur lorsqu'elles s'étendent au monde qui les entoure. Dans ce monde, être informaticien c'est être au carrefour de toutes les pratiques, au carrefour de toutes les sciences, où l'apprentissage ne peut plus être que permanent.

Les sciences "du matériel" nous emmènent parfois très loin parfois très près de nous. Les sciences "de l'immatériel" nous aident à retrouver notre chemin dans ces allers et ces retours. Les sciences humaines, économiques, sociales, devraient être incontournables dans une démarche de rationalisation de nos organisations et pour soutenir nos innovations, respecter notre environnement.

De l'approche systémique des TIC, à la Responsabilité Sociétale de l'Entreprise, l'espace se réduit et c'est une bonne chose. A mi-chemin entre industrie et service, les TIC sont le point de jonction d'où la création de valeur pourrait être la plus forte autant pour l'organisation entrepreneuriale que pour la communauté, qu'elle soit d'usage ou territoriale.

Pour vous y convaincre, je vous propose de balayer quelques idées qui méritent d'être dépoussiérées. *(Achevé de rédiger en janvier 2012).*

MYTHE N° 1 :

EN MATIÈRE DE RSE, L'ENTREPRISE INFORMATIQUE S'INTERESSE SURTOUT À LA COMPENSATION CARBONE

Non, être une entreprise informatique responsable ne se limite pas à tenter de compenser les émissions de carbone. Sur le site du Ministère du Développement Durable, on peut lire la définition suivante en préambule des actions menées dans le cadre de la politique RSE :

“La responsabilité sociétale des entreprises (RSE) est la contribution des entreprises aux enjeux du développement durable. La démarche consiste pour les entreprises à prendre en compte les impacts sociaux et environnementaux de leur activité pour adopter les meilleures pratiques possibles et contribuer ainsi à l'amélioration de la société et à la protection de l'environnement. La RSE permet d'associer logique économique, responsabilité sociale et écoresponsabilité.”

Il s'agit de donc de “contribuer” aux “enjeux du développement durable”. Si nous nous limitons à cette première phrase, avec nos actions de tri sélectif et de lutte contre l'effet de serre, nous n'irions pas très loin. Les objectifs se précisent avec les explications complémentaires sur la démarche que doit engager l'entreprise : “prendre en compte les impacts sociaux et environnementaux”, “meilleures pratiques possibles”, “amélioration de la société”, “protection de l'environnement”.

L'objectif est aussi d'améliorer la société et ce n'est pas en plantant un arbre à Montpellier à chaque fois que nous ferions construire un firewall en Malaisie pour le vendre à Lille que nous pourrions espérer avoir rempli nos devoirs d'entrepreneur responsable.

La responsabilité sociétale est un véritable challenge quotidien qui demande au prestataire informatique de réfléchir à son propre modèle et surtout de faire preuve d'intelligence économique. Comme tout autre, il doit écouter, analyser, partager, définir et construire tout en cherchant à s'améliorer. Plus que tout autre, il contribue à soutenir les services qui feront la qualité de vie future.

L'entreprise informatique est justement placée au cœur de ces préoccupations puisqu'elle participe déjà par ses domaines et par ses métiers à nous faire évoluer d'un monde industriel concentré sur ses biens matériels, vers un monde serviciel qui respecte son patrimoine matériel et immatériel.

Au delà de la compensation carbone, le prestataire informatique a les compétences nécessaires pour jouer un rôle majeur dans la mise en œuvre de politiques RSE. L'entreprise informatique et les membres qui la composent peuvent contribuer au ciment qui va lier « logique économique », « responsabilité sociale » et « ecoresponsabilité ».

MYTHE N° 2 :

LA MUTUALISATION DES RESSOURCES INFORMATIQUES IMPLIQUE NECESSAIREMENT L'OPTIMISATION DE NOTRE IMPACT ENVIRONNEMENTAL

Non, mutualiser les ressources informatiques ne permet pas systématiquement d'optimiser nos impacts environnementaux.

Parlons du Cloud, par exemple. Un des premiers arguments avancés est celui de la maîtrise des coûts de production et de l'optimisation de la consommation. Effectivement, optimiser son usage et ne payer que ce que nous utilisons réellement constitue une nette amélioration des conditions de consommation dans un monde où nous sommes particulièrement consommateurs de biens et services qui ne nous intéressent pas forcément. Combien de fois nous plaignons nous du fait qu'un fournisseur nous contraint à acheter un ensemble de produits ou de services dont nous n'allons consommer que la moitié, voire le tiers, ou même aucun ?

Si vous en doutez, intéressez vous aux assurances liées à vos cartes bancaires.

Donc, effectivement, cesser d'acheter des serveurs, cesser de construire des infrastructures locales et cesser de renforcer des équipements dédiés à une informatique géographiquement restreinte, pour passer à une consommation de services internationaux dont nous allons payer l'usage et non la propriété, peut paraître aller dans le bon sens, a priori.

Pourtant, justement, quand je loue une ressource en mode SAAS (Software As A Service) ou IAAS (Infrastructure As A Service), qui me dit que je paie un service dont j'ai besoin en totalité ? Construire et fournir un service automatisé à l'international implique de définir des niveaux de services, des standards, et de procéder à une industrialisation maximale de sa production. Or, optimiser, c'est tendre vers l'ensemble plutôt que de considérer les groupes ou les unités.

Nous savons qu'en nous éloignant des besoins spécifiques de nos utilisateurs nous nous éloignons aussi de leur satisfaction et nous en payons les conséquences lorsque nous l'oublions : standardisé au maximum, le service ne répond plus toujours aux attentes. Le monde du libre l'a bien compris par sa gestion communautaire limitée aux usages du plus grand nombre mais ouverte aux développements spécifiques. Le lien direct entre créateur et consommateur de

valeurs a du sens et gagne à être préservé, voire gagnerait à être partagé par de nouvelles formes de coopération.

Dans un projet de mutualisation, se limiter à ne considérer qu'une partie des coûts de production nous empêche d'identifier l'ensemble des impacts de la mise en commun des ressources et nous limite considérablement dans la prise en compte des impacts environnementaux. Considérer le travail local, la prise en compte des territoires et des collectivités, réinventer les communautés et intégrer de nouvelles coopérations, etc... peut largement contribuer à agir pour préserver notre environnement.

Dans cette démarche, l'entreprise informatique est une ressource fondamentale génératrice d'une valeur dont les limites restent encore à découvrir.

MYTHE N° 3 :

L'INDUSTRIALISATION DE L'INFORMATIQUE GENERE AUTOMATIQUEMENT DE LA VALEUR

Non, l'industrialisation de l'informatique n'implique pas systématiquement la génération de valeur.

Comme dans toute entreprise, l'augmentation de la valeur est un objectif en soi. C'est même un objectif en tête du hit-parade des plans stratégiques. A tel point qu'il devient banal de l'évoquer. C'est directement sous cet objectif que sont défendus les axes de réduction de coûts et donc d'optimisation de production des biens et services.

Industrialiser, c'est standardiser, procédurer, automatiser. C'est faire en sorte que d'un prototype, d'un pilote, d'un POC (Proof of Concept) on puisse sortir un produit ou un service de façon automatique et à un niveau de qualité suffisamment stable pour être à même de le proposer à la vente. Si standardiser un bien matériel peut avoir ses limites, industrialiser le service peut aussi nous éloigner de notre objectif commercial en générant une perte d'adéquation entre le service et le consommateur. La différence ici, c'est que le service est produit en même temps qu'il est consommé et qu'il est impossible de le retenir en bout de chaîne. S'éloigner de la demande génère un impact immédiat sur l'image du prestataire de service (informatique ou non).

En production matérielle on a appris qu'en industrialisant l'ensemble des actes, on peut optimiser le processus et garantir un niveau de qualité relativement stable. Pourtant, on sait aussi qu'appliquer la procédure « à la lettre » mène tout droit à l'arrêt de la chaîne de production. De la même façon, respecter les normes et procédures ne nous empêche pas de vivre parfois des catastrophes (AZF, Fukushima...). C'est probablement aussi ça qu'on appelle « l'épreuve du réel ».

Ce qui reste, une fois qu'on a tout décrit, tout analysé, tout traduit dans des procédures enregistrées et certifiées, s'appelle le travail. Le travail c'est l'intelligence que met l'être humain dans son expérience du réel. Autrement dit, c'est ce que convoquera le technicien informatique pour résoudre un incident, un problème, ou même pour apporter une réponse à un besoin et qui ne sera pas décrit dans une procédure ou même identifié dans son champ de compétence.

Si industrialiser les éléments techniques et matériels paraît immédiatement générateur de valeur, appliquer des méthodes et procédures systématiques et restrictives aux services et aux individus limite la valeur de l'ensemble. Cela peut même aller jusqu'à générer de la souffrance et de l'incompréhension et finalement faire baisser considérablement la valeur intrinsèque de l'entreprise et des services qu'elle produit.

S'agissant d'informatique, veiller à ne pas dépasser les limites de bon sens d'une optimisation raisonnable peut avoir des impacts positifs à l'intérieur et aussi à l'extérieur de la structure et donc générer de la valeur même pour son environnement et cela dans une démarche responsable, sociétale, éclairée.

MYTHE N° 4 :

L'ENTREPRISE INFORMATIQUE N'EST PAS LEGITIME A INNOVER EN MATIERE DE RSE

Au contraire, vu son domaine de compétence et sa connaissance de la gestion de projets, le prestataire informatique est un acteur de choix pour la mise en œuvre d'une véritable politique RSE dans son entreprise mais également pour l'environnement institutionnel et social de cette entreprise.

L'informatique n'est pas seulement une affaire d'experts techniques et de sciences dures. Il y a quelques décennies il n'y avait pas beaucoup de formation en informatique et les équipes se constituaient également de biologistes, d'enseignants, qui ont contribué à nous apporter rigueur et méthode autant qu'écoute et analyse.

Dans la série des particularités, les sciences du bâtiment semblent même nous avoir offert leur vocabulaire spécifique (Maîtrise d'œuvre, Maîtrise d'ouvrage, Cahier des charges, pour exemple) et nous avons développé en retour des outils de gestion de projet qui paraissent souvent obscurs à ceux qui n'en ont jamais eu l'utilité mais qui même s'ils ne sont pas parfaits nous permettent aujourd'hui de gérer un portefeuille de projets avec professionnalisme.

Aujourd'hui, les formations informatiques sont pléthores, et les profils le restent aussi. Être informaticien c'est faire partie d'un groupe social plutôt indéfinissable aux compétences certaines et extrêmement variées. Pour la petite histoire, on y trouve également des ergonomes, des psychologues, et des artistes que l'on paie pour l'être (pas seulement en dehors de leurs heures de travail).

Alors face à la nécessité d'inventer de nouveaux modèles et afin de répondre aux besoins d'un environnement sociétal et économique qui a tendance à se tendre, s'appuyer sur ces compétences transverses c'est tout simplement à nouveau faire preuve de bon sens, de sens commun.

Innover demande autant de créativité qui supporte mal les contraintes, que de rigueur qui deviendra nécessaire pour passer de la volonté au résultat. Le prestataire informatique est un animateur et un traducteur interprète qui a la science des moyens sur lesquels reposent les services délivrés aujourd'hui et donc la science des ressources nécessaires à leur évolution. Il y a nécessité à réintroduire des échanges entre les différents acteurs pour être capable de prendre en compte l'expertise de chacun.

En matière d'innovation RSE, l'enjeu consiste notamment à créer les nouveaux services qui feront votre bien-être, en tant que client, usager, ou patient. La France innove relativement naturellement dans les sciences dures, mais peine à lancer de grands projets d'innovation sociale qui pourraient par exemple intégrer des travaux de recherche en psychodynamique du travail, en sociologie, avec des objectifs clairement assumés de perspectives RSE.

L'entreprise informatique, déjà sollicitée sur des projets d'étude et d'innovation autour des concepts de « ville intelligente » et de services aux personnes, est un porteur/vecteur de choix pour ce type de projets.

MYTHE N° 5 :

LA POLITIQUE RSE DE L'ENTREPRISE INFORMATIQUE SE LIMITE A DEMANDER AU DSI DE FAIRE DU « GREEN-IT ».

Bien entendu, c'est faux et voici pourquoi :

Depuis mes débuts dans l'informatique, la Direction des Systèmes d'information a toujours été une équipe, un pôle, un département, une Direction, bref un « élément » distinct et transverse par lequel les décisions concernant les technologies de l'information transitaient, souvent mais pas toujours.

Suivant les différentes périodes et les différents sujets, le Directeur des Systèmes d'Information est tour à tour :

- en charge de décider des moyens à mettre en œuvre, quel que soit le périmètre du système concerné et du moment qu'il s'agit de ce grand sac qu'on appelle « informatique ». Parce que finalement, « ces **« d'informaticiens »** !! d'utilisateurs ne peuvent de toute façon pas savoir ce qu'ils veulent, ni ce qui est bien pour eux. ».
- pas du tout consulté par les équipes « métiers » qui ont décidé pour le coup qu'il suffit et que ce site web « on va se l'acheter soi même au Québec, on sera tranquille sans ces **« d'informaticiens »** !! d'informaticiens qui ne nous donnent jamais ce qu'on veut. ».

Alors de la même façon qu'on s'y est longtemps pris avec les solutions informatiques (et qui ne fonctionnent plus vraiment) on peut choisir de tayloriser les tâches induites par la politique RSE, tenter de les répartir en fonction des domaines de compétences qu'on aurait a priori identifiés dans l'entreprise, et ce faisant demander au DSI de faire du « Green IT », si c'est tout ce que l'on a compris et qu'on attend de lui.

On peut aussi considérer que grâce à une politique RSE réfléchie, l'entreprise informatique peut innover, faire évoluer son positionnement sociétal et son modèle économique.

Il serait opportun d'inclure sa politique RSE au plus haut niveau de la stratégie d'entreprise et de la construire en collaboration avec l'ensemble des acteurs.

S'agissant d'informatique, les décisions que nous prenons aujourd'hui feront les innovations servicielles de demain. Une entreprise qui s'oriente vers des partenariats nationaux ou internationaux avec une démarche RSE ouverte et proactive, développe des projets d'Open Innovation, sublimant les usages d'une relation client-fournisseur.

En informatique plus qu'ailleurs, la politique RSE nous pousse à réinventer nos modèles et à reconquérir nos responsabilités.

Soyons créatifs !

MYTHES ET LEGENDES DU SOCIAL SHOPPING

Agnès TEISSIER, Diginove, Agence Conseil E-Commerce

INTRODUCTION

Facebook et surtout ses boutons J'aime, Partager, Commenter ont révolutionné les comportements et engendré le social shopping en multipliant les opportunités d'achats chez les e-commerçants et les marques. Ces nouveaux outils sont considérés comme un engagement émotionnel du client qui place la marque au même niveau que ses amis. En un clic, le lien du site marchand ou d'une fiche produit apparaît sur le mur Facebook de l'internaute, permettant de diffuser rapidement l'offre vers l'ensemble de la communauté d'amis virtuels d'un consommateur, qui partage les mêmes centres d'intérêt que lui. Les boutons servent de relais efficaces de l'information.

Dès à présent, les améliorations apportées par Facebook au-travers de ses nouvelles fonctionnalités, la Timeline, le Ticker, le nouveau fil d'actualité, le bouton « S'abonner », les Smart listes, et l'Open Graph 2.0, enrichissent les possibilités d'étendre sa communauté, de segmenter les informations selon ses propres centres d'intérêt, ses goûts, ses achats et intentions d'achats, et d'identifier les profils aux caractéristiques similaires.

Le Social Shopping, né de la combinaison de la communauté et de l'e-commerce, prend de l'ampleur en revêtant des formes diverses, notamment le F-commerce, l'achat groupé ou les deals. Dans tous les cas, la recommandation de ses amis et les prix imbattables restent les leviers principaux qui vont motiver l'achat et représentent plus qu'une tendance de consommation durable, un véritable mode de consommation. *(Achevé de rédiger en janvier 2012).*

MYTHE N° 1 :

LE SOCIAL SHOPPING OU L'ACHAT GROUPE NE PROFITE PAS AUX COMMERÇANTS

Le modèle économique des deals sous-tend des offres promotionnelles fortes de plus de 50% de remise sur le prix de base, qui vont permettre de recruter et de générer du trafic au point de vente. Les différents acteurs (comme Groupon, eBuyClub, KGBDeals, Lookingo, etc) se rémunèrent par une commission. Les commerçants vont réaliser des prestations de service ou vendre des produits avec une marge réduite. L'effet de volume généré par l'opération et le trafic permettra d'augmenter sensiblement la notoriété et la visibilité, mais il appartient ensuite au commerçant de fidéliser les clients.

Le Social Shopping profite aux commerçants à plusieurs niveaux, pas seulement pour recruter et fidéliser, mais aussi pour valoriser sa communauté et générer des ventes. Il existe plusieurs mécaniques pour mettre en valeur une communauté :

- Vendre directement depuis sa page Facebook
- Générer des ventes online en proposant des bons de réduction
- Générer du trafic qualifié sur un site web
- Rediriger du trafic qualifié vers un magasin physique grâce à des bons de réduction
- Ecouter ses consommateurs
- Fidéliser ses fans en leur proposant du contenu à valeur ajoutée, ou des avantages exclusifs

- Organiser un concours pour collecter des données opt-in
- Générer des recommandations entre amis
- Faire participer ses fans à la création de l'offre produits
- Créer un évènement

MYTHE N° 2 :

LE SOCIAL SHOPPING NE FAIT PAS VÉRITABLEMENT ÉCONOMISER MAIS POUSSE À DÉPENSER TOUS AZIMUTS

Les deals dynamisent les achats impulsifs, en proposant des offres fortes et en incitant à découvrir d'autres produits et services, qu'on n'aurait pas forcément achetés autrement, et font ainsi dépenser dans d'autres sphères de dépenses. Les deals compensent l'hypo-consommation, due à la crise, où les consommateurs ne trouvant plus autant de plaisir dans la consommation excessive se limitent à leur liste de courses. Grâce aux deals, ils peuvent consommer autrement, de manière plus avertie, en comparant et sont certains d'acheter les meilleurs produits au meilleur prix. Ils peuvent partager leurs bons plans avec leurs amis et vice-versa, profiter des bons plans de leurs amis. En économisant centimes après centimes, ils sont plus, ensuite, en mesure de dépenser pour de nouveaux produits et de nouvelles expériences.

MYTHE N° 3 :

LE SOCIAL SHOPPING EST PERÇU COMME DU SPAM, CAR LES DEALS NE SONT PAS CIBLES NI PERSONNALISÉS

Les deals ne sont pas un symptôme lié à la crise, mais sont une tendance de consommation durable, un nouveau mode de consommation. S'ils ne sont pas ciblés ni personnalisés, les deals peuvent être rapidement perçus comme intrusifs. Cependant, les consommateurs sont de plus en plus conditionnés à recevoir des deals pour tout, sans être lassés des offres, car les deals entretiennent la recherche de deals ciblés. Les deals sont désormais recherchés quel que soit le produit ou le service en fonction du besoin du moment. En ce sens, les internautes acceptent de faire le tri, de comparer et achètent le meilleur deal du marché qui correspond à leurs attentes.

Ainsi, entre janvier et juin 2011, Groupon a consacré 432 millions de dollars à la chasse aux adresses email et autres dépenses marketing, soit plus de 60% de son chiffre d'affaires. Ses membres se montrant peu fidèles, il faut sans cesse les solliciter.

MYTHE N° 4 :

LE SOCIAL SHOPPING PERMET DE CONSTRUIRE LA CONFIANCE ET DE DÉVELOPPER SES LIENS SOCIAUX

Le fan a tendance à faire naturellement confiance à une marque, surtout s'il la connaît déjà, il clique facilement sur le bouton J'aime de Facebook, mais il peut être rapidement infidèle si les informations communiquées ne correspondent plus aux promesses initiales, pour l'attirer, comme les bons plans ou offres exclusives.

La marque devra tisser un contact réel et très apprécié avec ses fans. Le lien social est important pour instaurer une relation de confiance. Les fans sont des clients et des clients potentiels, qui peuvent devenir des ambassadeurs de la marque, sans vraiment s'en rendre compte.

MYTHE N° 5 :

LE SOCIAL SHOPPING PERMET D'ACCROITRE LA FIDELITE A UNE MARQUE, MEME SI LE DEAL EST PLUS CHER QUE POUR UN PRODUIT SIMILAIRE D'UNE MARQUE INCONNUE DE NOUS

Aujourd'hui, le consommateur est de moins en moins fidèle aux marques et aux enseignes. L'internaute « aime » de nombreuses marques, il le fait pour montrer son attachement, parce qu'il est client mais aussi parce qu'il attend des actions concrètes comme des offres spéciales, promotions ou exclusivités. La marque a la responsabilité de le satisfaire et de répondre à ses attentes. En créant un lien privilégié, en engageant le dialogue avec sa communauté pour la fidéliser et en l'impliquant, les clients deviendront de véritables ambassadeurs de la marque. L'objectif est de connaître l'opinion de ses fans sur un sujet, par exemple sur un produit et de les impliquer, en leur demandant si la marque leur plaît, s'ils achètent le produit, s'ils ont envie d'autre chose, etc.

MYTHE N° 6 :

LE SOCIAL SHOPPING PERMET DE PARTAGER DES BONS PLANS AVEC SES AMIS POUR MAINTENIR L'EQUITE SOCIALE ET RENDRE LES FAVEURS QUE NOUS RECEVONS

On ne le dira jamais assez, mais la règle « Il faut donner pour recevoir » fonctionne aussi sur les réseaux sociaux. Désormais, au lieu de se cacher, il est accepté et de bon ton de faire de bonnes affaires. Cela peut être même une source de reconnaissance et de statut. Chez MisterGooddeal, les ventes des produits les mieux notés ont augmenté de 40 % après la mise en ligne d'avis de consommateurs.

Les réseaux sont la galerie marchande du XXI^e siècle : on y va faire son shopping entre amis et on doit y être. Ils constituent un levier incontournable sur le web marchand.

Cependant, le Social Shopping permet aussi de se rassurer sur ses achats et en ce sens, a une fonction de développement de réassurance personnelle quant à sa capacité de dépenser judicieusement son argent. Les recommandations de ses amis, la consultation des sites d'avis consommateurs, des comparateurs de prix, et la recherche de bons plans font partie intégrante du processus d'intention d'achat.

ACRONYMES

16QAM : 16 Quadrature Amplitude Modulation
64QAM : 64 Quadrature Amplitude Modulation
8PSK : 8 Phase Shift Keying
Affet : Agence française de sécurité sanitaire de l'environnement et du travail
AMR-WB : Adaptive MultiRate WideBand
ANSES : Agence Nationale de Sécurité Sanitaire de l'alimentation, de l'environnement et du travail
ATM : Asynchronous Transfer Mode
BSC : Base Station Controller
BTS : Base Transceiver Station
CDMA2000 1xEV : Code Division Multiple Access 2000 1xEVolution
CDMAOne : Code Division Multiple Access One
CIPRNI : Commission internationale de protection contre les rayonnements non ionisants.
En anglais : ICNIRP International Commission on Non-Ionizing Radiation Protection
CIRC : Centre international de Recherche sur le Cancer
CODEC : COder DECoder
CSD : Circuit Switched Data
DAS : Débit d'Absorption Spécifique. En anglais : SAR Specific Absorption Rate
DCS : Digital Communication System
EDGE : Enhanced Data rates for Gsm Evolution
eNodeB : evolved NodeB
EPC : Evolved Packet Core
FDD : Frequency Division Duplexing
FDMA : Frequency Division Multiple Access
GGSN : Gateway GPRS Support Node
GMSK : Gaussian Minimum Shift Keying
GPRS : General Packet Radio Service
GSA : Global mobile Supplier Association
GSM : Global System for Mobile Communications
HDFS : Hadoop Distributed File System, Système de fichier distribué adapté pour les traitements big data
HSE : hypersensibilité électromagnétique
HSPA : High Speed Packet Access
IEEE : Institute of Electrical and Electronics Engineers
IMS : Ip Multimedia Subsystem
IMT-Advanced : International Mobile Telecommunications – Advanced
IP : Internet Protocol
LTE : Long Term Evolution
MGW : Media GateWay
MIMO : Multiple Input Multiple Output
MME : Mobility Management Entity
MSC : Mobile Switching Center
MSC-S : Mobile Switching Center - Server
NGN : Next Generation Network
NoSQL : NoSQL signifie aujourd'hui Not Only SQL, que l'on pourrait littéralement traduire par « pas seulement SQL »
OFDM : Orthogonal Frequency Division Multiplexing

OMS : Organisation Mondiale de la Santé
PCU : Packet Control Unit
PDC : Personal Digital Cellular
PLC: Programmable Logic Controller
PTN : Packet Transport Network
QoS : Quality of Service
QPSK : Quadrature Phase Shift Keying
RNC : Radio Network Controller
SCADA : Supervisory Control and Data Acquisition
S&P Gw : Serving & Packet data network Gateway
SGSN : Serving GPRS Support Node
SMS : Short Message Service
SQL : SQL ou Structured Query Language, langage informatique effectuant des opérations sur des bases de données relationnelles)
TDD : Time Division Duplexing
TDMA : Time Division Multiple Access
TDM : Time Division Multiplexing
TD-SCDMA : Time Division Synchronous Code Division Multiple Access
TNT : Télévision Numérique de Terre
QPSK : Quadrature Phase Shift Keying
UMTS: Universal Mobile Telecommunications System
USB : Universal Serial Bus
VoIP : Voice over Internet Protocol
WAP : Wireless Application Protocol
WCDMA : Wideband Code Division Multiple Access

GLOSSAIRE

Certificat

c'est équivalent au visa d'un passeport : il vous donne personnellement un droit donné pendant une durée définie.

Clé publique, clé privée, bi-clé

le bi-clé est composé d'une clé publique et d'une clé privée, liées mathématiquement.

Un processus mathématique permet de générer des bi-clés. La clé publique peut être diffusée librement, le titulaire du bi-clé conservant cachée la clé privée. Ce qui est chiffré avec une clé ne peut être déchiffré qu'avec l'autre clé.

Si la génération a été bien faite, on ne connaît pas aujourd'hui de méthode (autre que la force brute) pour trouver la clé privée correspondant à une clé publique.

Les clés utilisées couramment aujourd'hui (clés de 2048 bits) correspondent à des nombres décimaux d'environ 1200 chiffres. Pour essayer toutes les combinaisons de clés de 2048 bits, il faut faire 2^{2048} essais !

Cloudera

Cloudera est une distribution d'Hadoop orientée vers les entreprises

Crédentiels

ce sont les éléments qui permettent de se faire reconnaître et de faire valoir ses droits : code utilisateur associé à son mot de passe, certificat et sa clé privée, token et son code PIN, ...

C'est l'équivalent d'un passeport et de ses visas, d'une carte d'identité, d'une lettre de créances, ...

Etatsunien

nom commun : habitant des États Unis d'Amérique. Sous forme d'adjectif : relatif aux États Unis d'Amérique.

Il est curieux que ce pays soit le seul dont les habitants n'ont pas de nom (comme le fait remarquer le cinéaste Jean-Luc Godard)

Hadoop

Hadoop est un kit de composants logiciels open source et écrit en Java, il permet la création d'applications distribuées et extensibles. Il permet à ces applications de travailler avec des milliers de nœuds distribués et des pétaoctets de données.

Hbase

Hbase est une base de données orientée colonne s'inscrivant dans la mouvance NoSQL dont le design s'inspire de BigTable (Google). Elle supporte jusqu'à plusieurs millions de lignes et fonctionne au-dessus de HDFS.

Hcatalog

HCatalog est un système de gestion de table et de stockage pour Hadoop. Il permet d'utiliser les outils de l'environnement Hadoop comme Pig, MapReduce et Hive.

Hive

Hive est un data warehouse libre implémentant un langage de requête orienté SQL (HiveSQL) dont la mise en œuvre se traduit par l'exécution de jobs Map/Reduce orchestrés par Hadoop.

HTTPS

HyperText Transfer Protocol Secure

c'est un protocole de transport d'information sécurisé utilisé pour consulter des sites web, qui permet la confidentialité des informations échangées, de contrôler l'authenticité du site consulté, et éventuellement de contrôler l'identité de celui qui consulte ce site.

La confidentialité est assurée par une couche SSL ou TLS

Langage R

le R est un langage de programmation comprenant des méthodes statistiques et des facilités graphiques importantes.

MapReduce

MapReduce est un environnement de développement permettant de traiter de manière distribuée des traitements sur des larges volumes de données. Issue de Google il est considéré comme le noyau Hadoop.

Mahout

Mahout est un logiciel d'apprentissage de traitement automatique de données et librairie de data mining

Oozie

Oozie est un moteur de workflow qui permet d'orchestrer des programmes Map/Reduce et des scripts Pig.

Pig

Pig est un langage de requête pour analyser un vaste ensemble de données, né chez Yahoo. C'est une abstraction comme peut l'être SQL pour écrire des requêtes qui sont alors traduites en job Map/Reduce et exécutées de façon distribuée sur le cluster. Cela peut simplifier considérablement l'écriture d'une requête.

Sqoop

Sqoop va, lui, vous permettre d'importer vos données stockées dans une base SQL dans HDFS. Le nom vient d'ailleurs de la contraction SQL-to-Hadoop. Il peut également importer les données dans Hive.

SSL, TLS

Secure Socket layer, Transport Layer Security

SSL est l'ancienne appellation de TLS. Il s'agit d'un protocole de chiffrement de données, permettant la confidentialité lors d'échange de données.

USA Patriot Act

Uniting and Strengthening America by Providing Appropriate Tools Required to Intercept and Obstruct Terrorism Act (oh, quel bel acronyme!)

Loi étatsunienne votée juste après les attentats du 11 septembre 2001, permettant entre autre l'accès par le FBI à toute information se trouvant sur le territoire des USA, ou traitée par une société étatsunienne ou une filiale d'une société étatsunienne.

VPN

Virtual Private Network

Réseau privé virtuel, permettant de transporter de manière confidentielle des informations d'un point A à un point B, même s'il faut traverser plusieurs réseaux publics pour relier physiquement ces deux points.

XML

L'Extensible Markup Language est un langage de balisage générique qui dérive du SGML. Cette syntaxe permet de définir différents espaces de noms, c'est-à-dire des langages avec chacun leur vocabulaire et leur grammaire, comme XHTML, XSLT, RSS, SVG.

Zookeeper

Zookeeper est un moteur de workflow qui permet d'orchestrer vos programmes Map/Reduce ou vos scripts Pig.

POUR ALLER PLUS LOIN DANS LA CONNAISSANCE DES TIC

Les contributeurs de cet ouvrage collectif ont également écrit, ou participé à l'écriture de livres dans leurs domaines d'expertises. Voici, sans être objectif, loin s'en faut, quelques-uns de ces livres qui peuvent vous permettre d'aller plus loin dans la connaissance des TIC.

Aurélie Banck

Auteure de :

- A paraître 2018 Vademecum de la protection des données personnelles pour le secteur bancaire et financier, Les essentiels de la Banque et de la Finance
- Analyse prospective pour le secteur financier : Les impacts du RGPD sur le profil, Banque & Stratégie, n°360, juillet 2017
- Neuf idées préconçues sur le règlement européen relatif à la protection des données Revue Banque n° 810, juin 2017
- GDPR et sous-traitance : un nouveau devoir de conseil? Dalloz IP/IT n°36, janvier 2017
- Panorama des obligations de notification des violations de données et d'incidents de sécurité, Revue Banque n°803, décembre 2016
- Données personnelles : l'impact du GDPR sur les éditeurs de logiciel, Expertises des systèmes d'information n°417, octobre 2016

Laurent Bloch

Auteur ou co-auteur de :

- Bloch, Laurent. Initiation à la programmation avec Scheme. Paris: Technip, 2011.
« La maîtrise d'Internet : des enjeux politiques, économiques et culturels ». Questions internationales, n° 47 (janvier 2011).
- Les systèmes d'exploitation des ordinateurs – Histoire, fonctionnement, enjeux. Paris: Laurent Bloch, 2017.
- L'Internet, vecteur de puissance des États-Unis? — Géopolitique du cyberspace, nouvel espace stratégique. Paris: Éditions du Diploweb, 2017.
<https://www.diploweb.com/Laurent-BLOCH.html>.
- Révolution cyberindustrielle en France. Paris: Economica, 2015.
<http://www.laurentbloch.net/MySpip3/Revolution-cyberindustrielle-en-307>.
- « Révolution dans la nouvelle cité électronique - Cyberspace, démocratie et logiciel libre ». Site WWW de Laurent Bloch, 2012.
<https://laurentbloch.net/MySpip3/Revolution-dans-la-nouvelle-cite-electronique>.
- Systèmes d'information, obstacles et succès – La pensée aux prises avec l'informatique. Paris: Vuibert, 2005.
<http://www.laurentbloch.net/MySpip3/-Systemes-d-information->

- Bloch, Laurent, Christophe Wolfhugel, Ary Kokos, G r me Billois, Arnaud Soulli , Alexandre Anzala-Yamajako, et Thomas Debize. S curit  informatique – Pour les DSI, RSSI et administrateurs. Paris: Eyrolles, 2016.
<https://laurentbloch.net/MySip3/Securite-informatique>.
- Eudes, Yves, et Laurent Bloch. « La France est en train de rater la troisi me r volution industrielle ». Le Monde, 7 janvier 2014.
https://www.lemonde.fr/economie/article/2014/07/01/laurent-bloch-la-france-est-en-train-de-rater-la-troisieme-revolution-industrielle_4448233_3234.html.
- Rochet, Claude, Michel Volle, Jean-Pierre Corniou, Laurent Bloch, Jacques Printz, Francis Jacq, et Pierre-Jean Benghozi. L'intelligence  conomique – Les nouveaux mod les d'affaires de la troisi me r volution industrielle. De Boeck, 2015

Alessandro Fiorentino

Auteur de :

- Quel sch ma organisationnel de gouvernance Informatique et Libert s ? Etude du Cabinet Infhotep, date de publication 18 mai 2017
- LE REVEIL DE LA VIE PRIVEE. Premi re publication de l'association PRIVACY TECH, date de publication 06 avril 2017
- Protection des donn es personnelles ; le compte   rebours a commenc . Etude du Cabinet Infhotep, date de publication 24 novembre 2016
- Economie des donn es personnelles » : Les enjeux d'un business  thique. Etude du CIGREF, date de publication 11 octobre 2015
- Privacy by Design   la lumi re du r glement europ en. Th se professionnelle ISEP soutenue le 2 d cembre 2013

Jean-Denis Garo

Auteur de :

- « Mon papa travaille dans l'Informatique et les T l coms » - 2007
- « Anita & B atrix – Le sens cach  du vocabulaire des IT » - 2010

Co-Auteur des livres collectifs :

- « S curit  des Syst mes d'Information » Les Guides ECOTER, Edition Mission Ecoter – 2002
- « Guide TIC des petites et moyennes collectivit s », Edition Ficome – 2004
- « La s curit    l'usage des d cideurs ». Edition etna France- 2005
- « La s curit    l'usage des PME et des TPE », Edition T nor – 2006
- « La S curit    l'usage des collectivit s locales et territoriales », Edition Forum ATENA-2009
- « Lexique des TIC », Edition Forum ATENA – 2010
- « L'Internet   l'usagede l' cosyst me num rique de demain », Edition Forum ATENA - 2011

Responsable  ditorial de :

- « L'Off-Shore et les centres de contacts Cap sur l'île Maurice », Edition 1Angle2Vues - 2007

Agnès Teissier

Auteur de plusieurs articles sur la distribution digitale, le Cash Back et Social Shopping, publiés notamment sur le Journal du Net.

- Evolution de la distribution digitale
<http://issuu.com/casualconnect/docs/2009winter>
- JDN (2011) : Cash Back, achat groupé, deals, la nouvelle donne du social shopping
<http://www.journaldunet.com/ebusiness/expert/50117/cash-back--achat-groupe--deals---la-nouvelle-donne-du-social-shopping.shtml>
- JDN (2008) : Distribution digitale
<http://www.journaldunet.com/solutions/expert/30364/distribution-dematerialisee--l-avenir-du-logiciel-et-des-contenus-numeriques.shtml>

WEBOGRAPHIE :

MYTHES ET LEGENDES DU BIG DATA

- <http://shop.oreilly.com/product/06369200266.do>
- http://wikibon.org/wiki/v/Big_Data_Market_Size_and_Vendor_Revenues
- <http://hadoop.apache.org/>
- http://www.mckinsey.com/insights/mgi/research/technology_and_innovation/big_data_the_next_frontier_for_innovation
- <http://www.forbes.com/sites/davefeinleib/>
- <http://strata.oreilly.com/2012/02/what-is-apache-hadoop.html>
- <http://hypedrivendev.wordpress.com/2011/12/>

MYTHES ET LEGENDES DES RISQUES SANITAIRES DE LA TELEPHONIE MOBILE

- Conclusion du ministre américain de la santé « Mobile Phones, Brain Tumours and the Interphone Study: Where Are We Now? » National Institutes of Health » U.S. Department
- Etude Interphone
<http://ije.oxfordjournals.org/content/39/3/675.abstract?sid=feb4993a-9e4a-4e5b-a6b9-c02b930fb63c>
http://www.iarc.fr/en/media-centre/iarcnews/2009/interphone_status.php
- “Base Stations and Wireless Networks: Exposures and Health Consequences” (OMS - http://www.who.int/peh-emf/meetings/base_stations_june05/en/index.html)
- Health Effects of Exposure to EMF” Scientific Committee on Emerging and Newly Identified Health Risks (SCENIHR)
http://ec.europa.eu/health/ph_risk/committees/04_scenihhr/docs/scenihhr_o_022.pdf
- European Bioelectromagnetics Association (EBEA)
- Société française de radio protection (SFRP)
 - Sensibilise le public aux effets de la radio sur l'être vivant (séminaire, publication)
- Union radio scientifique internationale (URSI)
- COMARE (UK)
- ISRN (France)
- « Champs électromagnétiques, environnement et santé » de Anne Perrin et Martine Souques
- Ligue suisse contre le cancer :
http://www.mobile-research.ethz.ch/var/Commentaire_Interphone_update01.pdf
- Organisation Mondiale de la santé : <http://www.who.int/fr/>
- ICNIRP : <http://www.icnirp.net/>
- Rapport de l'académie de médecine en France :
<http://www.academie-medecine.fr/detailPublication.cfm?idRub=26&idLigne=1752>

- Article de wikipédia : pollution électromagnétique :
http://fr.wikipedia.org/wiki/Pollution_%C3%A9lectromagn%C3%A9tique4
- Université d'Ottawa (Canada) :
<http://www.rfcom.ca/welcome/indexfr.shtml>
http://www.iarc.fr/en/media-centre/pr/2010/pdfs/pr200_E.pdf
- Règlementation mondiale en terme d'émission :
<http://www.who.int/docstore/peh-emf/EMFStandards/who-0102/Worldmap5.htm>
- Organisations indépendantes :
<http://www.teslabel.be>
<http://www.iemfa.org>
- Dossier santé et télécom de l'Anses :
<http://www.anses.fr/ET/PPNBB69.htm>
- Portail français sur les effets des mobiles sur l'environnement :
<http://www.radiofrequences.gouv.fr/>
- Prévention mobile :
<http://www.lesondesmobiles.fr/>
- Exposition aux champs électromagnétiques en France :
www.cartoradio.fr
- La Fondation Santé et Radiofréquences :
<http://www.sante-radiofrequences.org/>

A PROPOS DES AUTEURS

Les bio sont celles des dates où les auteurs ont remis leurs contributions.

Par ordre alphabétique :



Aurélie Banck est juriste en droit de la protection des données depuis 2006, elle a notamment occupé un poste de juriste à la CNIL. Elle est également responsable pédagogique du Diplôme d'Université Data Protection Officer de l'Université de Paris Nanterre. Certifiée CIPP/E, CIPM et FIP par l'IAPP et co-auteur du Vademecum de la protection des données personnelles pour le secteur bancaire et financier (Revue Banque) à paraître.



Laurent BLOCH a été précédemment responsable de l'informatique scientifique de l'Institut Pasteur, responsable de la sécurité des systèmes d'information de l'Inserm, puis Directeur du Système d'Information de l'Université Paris-Dauphine. Laurent Bloch est auteur de plusieurs ouvrages sur les systèmes d'information et leur sécurité. Il se consacre à la recherche en cyberstratégie.

Il a ainsi participé au démarrage de l'Internet en France. *lb (at) laurentbloch.org*



Jean-Marie CORRIERE est un manager expert en transition des organisations. Analyste de l'organisation en "Centres de Services", il accompagne des groupes internationaux des domaines tertiaires et industriels et des Sociétés de Services (IT, Logiciel Libre) dans la définition de leur stratégie et dans sa mise en œuvre opérationnelle. Créateur de services IT à très forte valeur ajoutée et passionné par la coopération et la co-crédation, il contribue et soutient des projets de recherches liés à l'économie de l'immatériel. *jmcorriere (at) up2it.fr*



Alessandro FIORENTINO est consultant en charge de l'offre de Protection des données à caractère personnel au sein du Cabinet Infhotep et Vice-Président de l'association PRIVACY TECH. Nommé "Ambassadeur du Privacy by Design", le 22 mai 2013 par l'Office du Commissaire à l'information et à la protection de la vie privée de l'Ontario, Canada. Il assure l'unité d'enseignement Privacy by Design pour le Mastère Spécialisé en Management et Protection des données à caractère personnel à l'Institut Supérieur d'Electronique de Paris et l'unité d'enseignement Méthodologies du DPO au sein du Mastère spécialisé en Management et Protection des données à caractère personnel de l'Institut Mines-Telcom Business School.



Ivan de LASTOURS travaille à la direction de l'innovation de l'Institut Mines-Télécom. Il participe au financement et au suivi de projets de recherches nationaux ou européens (Grand emprunt, FUI, FP7...). Sur l'axe Big Data, l'Institut Mines-Télécom travaille avec le GENES (Ensa, Ensa, Crest, Cepa) sur la mise en place d'une plateforme Big Data au service des projets de recherche et d'expérimentation nationaux (projet BADAP. Ivan contribue également à la promotion de Skolkovo (pôle technologique de la région moscovite) auprès des entreprises françaises. Il a auparavant été chargé d'affaires en fusions acquisitions, plus particulièrement sur des opérations dans le secteur de la défense.



Louis DERATHE présente une expérience de plus de 15 ans en SSI, successivement informaticien, responsable de l'organisation d'une Cie d'Assurance, consultant en Système d'Information, officier en charge de SSI et maintenant expert SSI à THALES. Il est à ce sujet l'auteur d'un roman sur la guerre de l'information « Opération ACCAPARE » Ed l'Harmattan (et bientôt d'un autre sur la Métacommunication) *Louis.DERATHE (at) thalesgroup.com*



Jean-Marc Do LIVRAMENTO est titulaire d'un DEA de Micro Electronique de l'Université Pierre et Marie Curie (PARIS VI). Il a été pendant 9 ans consultant en réseaux d'entreprises puis en réseaux d'opérateurs. Il a passé treize années chez un opérateur télécom où il a été en charge successivement de veille technologique, d'expertise télécom, d'études technico-économiques, d'instruction de dossiers réglementaires et enfin de plan stratégique réseau. Il est aujourd'hui consultant et conseille de grandes entreprises et opérateurs sur leur stratégie télécom.

jmdolivramento [at] laposte.net



Jean-Baptiste FAHY, est consultant en sécurité informatique, spécialisé dans l'identité numérique. Il a été auparavant directeur du support chez Credentiel, spécialiste de la gestion de l'identité numérique mobile et de la signature numérique. Il assiste les entreprises et les administrations sur le plan organisationnel (analyse du positionnement en terme d'identité numérique, élaboration de la politique de sécurité, plan d'action...) et sur le plan technique (formation, maîtrise des PKI et des certificats, rédaction des procédures, ...) *jb.fahy (at) free.fr*



Jean-Denis GARO, Directeur Communication et Marketing Support d'Aastra, est Titulaire d'un DEA Science et Technologie du CSTS, complétant un DUE à la Faculté de droit de Bordeaux et une école de Commerce dans la même ville.

Il a effectué sa carrière dans les IT, Matra Communication, Nortel Networks, EADS Telecom et Aastra. Administrateur du Forum ATENA, il est auteur de plusieurs ouvrages spécialisés. Il intervient dans les domaines touchant à l'évolution des usages dans les centres de contacts, les communications unifiées et collaboratives, la téléphonie IP, les solutions de vidéoconférence et les réseaux sociaux d'entreprises. *jgaro (at) aastra.com*



Jean-Yves GRESSER, X62, ENST 67, MSEE MIT ('68) a un passé de chercheur puis de directeur de recherches au CNET, de responsable informatique dans les télécom. et la finance (banque, banque centrale, assurance) puis de « dircom » et de mercatique sur la toile du premier groupe mondial d'assurance crédit (1998-2002).

Depuis, il continue de parrainer des jeunes innovateurs en Europe et aux Etats-Unis, dans le numérique et le cinéma, et d'inventer. Il est vice président du Black Forest Group Inc. (NY), membre fondateur de la Société française de terminologie, membre de plusieurs commissions spécialisées de terminologie et de néologie du dispositif d'enrichissement de la langue française et d'autres associations dont le Stéréo-Club de France, fondé en 1903. *jgresser (at) numericable.com*



Hervé LEHNING a une formation classique (école normale supérieure, agrégation de mathématiques). Parallèlement à une carrière de professeur en école d'ingénieurs (école nationale d'ingénieurs de Tunis, école centrale de Paris) puis en classes préparatoires (lycée Janson de Sailly), il a effectué des missions de conseils en informatique, actuariat et sécurité du numérique.

Il est rédacteur en chef du magazine de vulgarisation mathématique *Tangente*. Il contribue également aux magazines *Pour La Science* et *La Recherche*. Il est l'auteur de nombreux ouvrages de vulgarisation dont *l'Univers des codes secrets de l'Antiquité à*

internet paru chez Ixelles en 2012. *hervelehning (at) orange.fr*

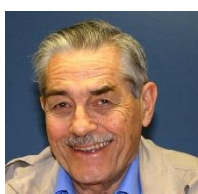


Francesca MUSIANI est ingénieure de recherche ARMINES, attachée de recherche et doctorante au Centre de sociologie de l'innovation (CSI) de MINES ParisTech (UMR 7185 CNRS) et enseigne à l'Université Pierre et Marie Curie. Diplômée en communication des organisations (Université de Padoue, Italie) et en droit international (Université pour la Paix des Nations Unies), elle participe actuellement au projet ANR ADAM et rédige sa thèse sur la technologie P2P appliquée aux services Internet.

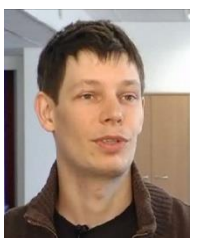
Francesca est l'auteur de plusieurs articles sur les pratiques "alternatives" du P2P, publiés dans Terminal, Observatorio et tripleC, et de Cyberhandshakes: How the Internet Challenges Dispute Resolution (...And Simplifies It), publié en 2009 par EuroEditions grâce à une bourse de publication de la European Foundation for the Information Society. *francesca.musiani (at) mines-paristech.fr*



Gérard PELIKS préside l'atelier sécurité de l'association Forum ATENA, et anime les Lundi de l'IE, pour les aspects sécurité de l'information dans le cadre du Cercle d'Intelligence Économique du Medef Hauts-de-Seine. Il préside l'association CyberEdu. Il est membre de l'ARCSI et de la Réserve Citoyenne de Cyberdéfense dans la Gendarmerie nationale. Gérard Peliks est chargé de cours dans des écoles d'Ingénieurs et universités, sur différentes facettes de la sécurité. *gerard.peliks (at) forumatena.org*



Louis POUZIN est entré dans l'informatique avant que le mot existe. Il a participé au MIT (Boston) à la construction du premier grand système de temps partagé, et inventé le premier SHELL (interpréteur de script). Plus tard il a construit le réseau Cyclades et inventé la technique de datagrammes, reprise dans TCP-IP. Il a fondé une association EUROLINC pour la promotion des langues natives dans l'internet, ainsi qu'une société Open-Root pour permettre aux PME d'acquérir des noms de domaine de 1er niveau.



Nicolas RUFF est chercheur au sein de la société EADS.

Il est l'auteur de nombreuses publications sur la sécurité des technologies Microsoft dans des revues spécialisées telles que MISC. Il dispense régulièrement des formations sur le sujet et participe à des conférences telles que SSTIC, les Microsoft TechDays ou la JSSI de l'OSSIR.

nicolas.ruff (at) eads.net



Agnès TEISSIER est Consultante, Fondatrice de l'Agence Conseil E-Commerce Diginнове, fruit de son expérience de plus de 20 ans du Marketing et de la Communication. Titulaire d'un DESS en Marketing, Agnès a occupé durant sa carrière des fonctions managériales dans les IT et les nouvelles technologies. Directrice de la Communication chez Nexway, leader de la distribution digitale de contenus logiciels et jeux vidéo, elle a développé la notoriété de l'entreprise dans les médias et l'industrie des logiciels et de l'entertainment. Plus récemment, chez Franfinance (Groupe Société Générale), elle a lancé un nouveau service web de Cash Back. Elle dispense aussi régulièrement des formations sur le sujet et les réseaux sociaux. *agnes.teissier (at) diginnove.com*



Viken TORAMANIAN, ingénieur de l'Institut Supérieur d'Electronique de Paris (ISEP) et titulaire d'un MS-ESSEC, a acquis une expérience en tant que consultant dans les Télécoms. Après plusieurs missions respectivement au sein d'Orange, SFR et Ericsson, il travaille actuellement à Orange Guinée en tant que responsable déploiement de réseaux 2G et 3G.
viken_toramanian (at) hotmail.com

Les idées émises dans ce livre n'engagent que la responsabilité de leurs auteurs et pas celle de Forum ATENA.

La reproduction et/ou la représentation sur tous supports de cet ouvrage, intégralement ou partiellement est autorisée à la condition d'en citer la source comme suit :

© Forum ATENA 2011 – Mythes et légendes des TIC

Licence Creative Commons

- Paternité
- Pas d'utilisation commerciale
- Pas de modifications



L'utilisation à but lucratif ou commercial, la traduction et l'adaptation sous quelque support que ce soit sont interdites sans la permission écrite de Forum ATENA.