



Sommaire :

- 1 Récapitulatif
 - 1 Méthodes de détection des menaces internes
 - 8 Méthodes de détection des menaces externes
 - 16 Opérationnalisation des risques internes et externes
 - 17 À propos des solutions IBM SPSS d'analyse prédictive
 - 19 À propos d'IBM Business Analytics
-

Les données, une arme à exploiter pour la détection des menaces internes et externes

Récapitulatif

L'analyse de l'information est la pierre angulaire d'une politique de défense du territoire efficace. L'objectif des activités de veille et de lutte contre le terrorisme est de détecter à temps les menaces pour la sécurité, afin de pouvoir prendre les mesures appropriées. Or les modèles qui permettent d'identifier ces menaces sont souvent noyés dans une masse de données. Pour y parvenir, il existe une nouvelle forme d'analyse de l'information particulièrement utile : l'analyse prédictive. Les solutions d'analyse prédictive appliquent des techniques élaborées de statistiques, d'exploitation de données et d'apprentissage automatique aux informations historiques afin d'identifier des modèles et des tendances cachés, même dans des ensembles de données étendus et complexes. Contrairement aux méthodes de détection et d'analyse à base de règles, l'analyse prédictive permet d'identifier des comportements assez inhabituels, même ceux présentant des différences subtiles qui souvent ne sont pas identifiés par d'autres méthodes.

Les techniques d'analyse prédictive consistent à explorer et tirer les enseignements de toutes les dimensions de données, permettant ainsi aux analystes d'allier connaissances humaines, première expérience et intuition pour guider leur application des techniques d'analyse. Les solutions d'analyse prédictive les plus efficaces peuvent analyser non seulement les données tabulaires, mais aussi les données textuelles. La capacité de l'analyse prédictive à combiner en permanence une grande variété de dimensions, types et sources de données permet de détecter rapidement et de manière fiable les signatures involontaires de pirates informatiques, de criminels ou de terroristes lorsqu'ils créent des cyber-dialogues ou tentent de nouvelles tactiques pour accéder sans autorisation à des informations sensibles.

Ce livre blanc tente de définir différents défis d'évaluation des risques en termes de sécurité et d'aide décisionnelle auxquels les techniques de l'analyse prédictive peuvent apporter une réponse partielle ou complète. Il comprend des exemples de détection des menaces internes et externes, et présente des méthodes spécifiques de manipulation et de modélisation des données. Les exemples et les graphiques sont tirés de données échantillonnées.

Méthodes de détection des menaces internes

Les coûts liés à l'attaque d'une personne malveillante, qui dispose d'un accès ou d'informations privilégiées réservés aux initiés, auront souvent un impact plus important et de plus longue durée sur une entreprise qu'une menace externe. Les conséquences d'actes malveillants commis par des utilisateurs agissant de l'intérieur ont souvent été désastreuses : violation de la confidentialité, atteinte à l'intégrité de l'information, influence



néfaste sur la politique du gouvernement, révélation des sources et des méthodes, et compromission des opérations de terrain, avec une issue parfois fatale pour certains agents. ¹

A l'occasion d'une enquête en ligne réalisée en 2004 par la revue CSO Magazine, en coopération avec les Services secrets américains et le Centre de coordination CERT, 59 % des responsables de la sécurité et de la mise en application de la loi interrogés ont signalé des intrusions causées de l'intérieur par des initiés qui avaient impacté négativement leurs organisations. ²

Les cybercrimes d'initié sont souvent très difficiles à détecter car l'auteur du crime a souvent un motif légitime pour consulter, modifier et manipuler des données critiques et/ou sensibles. Toutefois, malgré ces difficultés, la plupart des organisations disposent d'un volume de données substantiel qui peut servir à caractériser et potentiellement limiter les effets d'une attaque par un initié malveillant. Il peut s'agir de données démographiques, d'évaluations des performances, de missions existantes ou passées, de communications électroniques internes et externes et de fichiers logs.

Modélisation des comportements antérieurs pour prévoir les comportements futurs

Une méthode de détection des menaces d'initié par l'analyse prédictive consiste à prendre des cas connus de comportement malveillant et à caractériser les différences entre ces cas précis et les cas « normaux » connus. Bien qu'il s'agisse là d'une approche idéale, dans le sens où les algorithmes de data mining permettent de reconnaître rapidement et facilement un comportement antérieur, il existe des difficultés inhérentes à l'utilisation exclusive de cette approche. D'une manière générale, les opérations d'initié malveillantes sont très rares. Les données historiques permettant de modéliser un comportement futur manquent souvent de cas d'illustration pour prédire précisément les cas similaires mais pas exactement identiques aux précédents cas connus d'opérations malveillantes.

Lorsqu'il s'agit de menace interne ou de détection de fraude, on se rend compte qu'une personne malveillante peut afficher des modèles comportementaux normaux, dynamiques et complexes. Dans ce cas, le délit peut être très difficile à détecter car le comportement de cette personne peut continuer à sembler légitime, et ne laisser transparaître que quelques légères évolutions au fil du temps. Il est donc important de déterminer non seulement le type de comportement affiché par les individus, mais aussi les individus dont le comportement diffère légèrement depuis peu de celui de leurs homologues.

Si une entreprise ou une agence est attentive à l'accès des initiés aux données électroniques sensibles, elle utilisera très souvent un ensemble de règles codées en dur pour identifier les comportements potentiellement anormaux. Par exemple, elle demandera à ce qu'une personne qui travaille normalement sur des dossiers de ressources humaines soit contrôlée si elle tente à plusieurs reprises d'accéder à des fichiers contenant des données sensibles du service d'ingénierie. Les opérations potentiellement malveillantes sont cependant souvent beaucoup plus subtiles et difficiles à détecter.

Le groupement, un moyen d'identifier les groupes importants

Une technique qui peut se révéler très efficace pour identifier les changements de comportements au fil du temps consiste à enregistrer un instantané d'un groupe à intervalles réguliers. Les modèles de groupement visent à identifier des groupes d'enregistrements similaires et à étiqueter les enregistrements en fonction du groupe auquel ils appartiennent. Cette

opération est effectuée sans connaissance préalable des groupes et de leurs caractéristiques. En fait, le nombre exact de groupes à rechercher peut même ne pas être connu.

C'est ce qui distingue les modèles de groupe des autres techniques d'apprentissage automatique : le modèle ne peut prévoir aucune zone de sortie ou de zone cible prédéfinie. Ces modèles sont souvent désignés comme des modèles d'apprentissage sans supervision, car il n'existe pas de norme externe en fonction de laquelle juger les performances de classification du modèle. Il n'existe pas de bonne ou de mauvaise réponse pour ces modèles. Leur valeur est déterminée par leur capacité à capturer des groupements intéressants dans les données et à fournir des descriptions utiles de ces groupements.

Les changements de comportement sont toujours attendus lorsque des groupes de personnes sont affectés à différents projets et tâches au fil du temps. Toutefois, lorsque les modifications effectuées par un individu diffèrent de celles de ses homologues, il peut s'agir d'un problème potentiel justifiant une enquête.

Dans la Figure 1, ci-dessous, une agence souhaitant analyser les changements de comportement au cours du temps peut d'abord cumuler toutes les variables intéressantes correspondant à un mois et une année spécifiques. Ces comptages ou agrégations ont des chances de comprendre des variables du type suivant : nombre de fichiers consultés, examens antérieurs des performances, emplacement et responsabilités actuels de la personne susceptibles de l'amener à accéder à des enregistrements particuliers. Toutes les données pertinentes peuvent alors être traitées via un algorithme de cluster pour créer un instantané des comportements du mois. Chaque personne est automatiquement affectée à un groupe spécifique par l'algorithme.

Table (59 fields, 75 records) #3				
	Project Status	Sum of Access to Classified Files	Kohonen Cluster - November	Kohonen Cluster - December
1	1.000	1 30	20	
2	1.000	1 30	20	
3	1.000	1 20	10	
4	1.000	1 30	20	
5	1.000	1 10	00	
6	1.000	1 10	00	
7	1.000	1 30	20	
8	1.000	1 30	20	
9	1.000	1 20	10	
10	1.000	1 30	20	
11	1.000	1 30	20	
12	1.000	1 30	20	
13	1.000	1 30	20	
14	1.000	1 10	00	
15	1.000	1 30	20	
16	1.000	1 30	20	
17	1.000	1 10	00	
Table	Annotations			

Figure 1: Un algorithme de groupement peut être utilisé à des intervalles définis (dans ce cas, mensuels) pour affecter des salariés à un groupe sur la base de leur comportement en réseau. La comparaison d'appartenance au groupe sur ces intervalles permet d'identifier les salariés dont le comportement a changé et/ou diffère de ceux de leurs homologues.

Le mois suivant, des données concernant le même groupe de personnes sont collectées. Les données du nouveau mois sont agrégées et notées par le même algorithme créé au cours du mois précédent. L'appartenance au groupe du mois initial est comparée à l'appartenance du deuxième mois. Un rapport d'exception est généré pour les personnes dont le comportement a suffisamment changé pour qu'elles sortent d'un groupe et entrent dans un autre. Une enquête plus poussée peut alors être menée sur ces personnes.

L'avantage du groupement sur la durée est que, contrairement à un modèle de série temporelle, toute variable, catégorique ou continue, permet de déterminer l'appartenance au groupe. Dans un modèle de série temporelle, toutes les zones doivent être numériques. Des facteurs de prévision pouvant être importants, par exemple le projet en cours, l'emplacement et le niveau de sécurité, ne peuvent pas être utilisés dans un modèle de série temporelle, mais peuvent être incorporés dans un modèle de groupement qui suit les changements sur la durée. Avec un modèle de groupement, il n'est pas nécessaire que l'enquêteur ou l'analyste se concentre sur une variable de sortie spécifique, par exemple le « nombre de fichiers consulté, » pour identifier un comportement inhabituel.

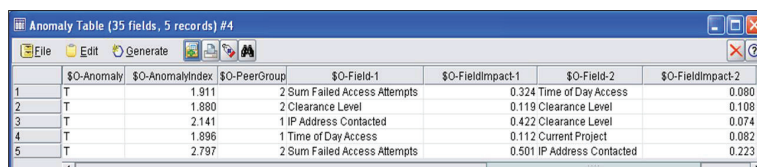
Les résultats de l'appartenance au groupe générés par un modèle de groupement sont souvent utilisés comme données en entrée dans les modèles créés dans les analyses suivantes. Leur utilité pour l'analyse exploratoire en fait un élément important de toute analyse dans laquelle des événements à haut risque doivent être isolés d'une grande quantité d'autres données.

Détection d'anomalies permettant de repérer les cas inhabituels

Il existe de nombreuses méthodes de groupement et de détection des anomalies. L'exemple ci-dessous met en évidence une méthode dans laquelle un algorithme regroupe d'abord les données, puis détecte des comportements atypiques dans ces groupes. Cela est particulièrement utile pour les travaux de veille car le processus peut être automatisé et permettre ainsi aux analystes de passer au peigne fin des millions d'enregistrements afin d'identifier les comportements atypiques ou les anomalies au sein de sous-groupes précis de personnes. Les méthodes automatisées de détection des anomalies sont disponibles dans certains outils de data mining et font partie de la famille des algorithmes de groupement.

Alors que les méthodes traditionnelles de recherche des comportements atypiques examinent en général uniquement une ou deux variables à la fois, un algorithme automatique de détection des anomalies peut étudier un nombre important de zones pour identifier des groupes ou des groupes d'homologues dans lequel on retrouve des enregistrement similaires. Chaque enregistrement peut alors être comparé avec d'autres enregistrements de son groupe d'homologues pour identifier les anomalies éventuelles. Plus un cas est éloigné du centre normal, plus il a des chances d'être inhabituel. Par exemple, l'algorithme peut regrouper des enregistrements dans trois groupes distincts et baliser ceux qui aboutissent loin du centre de l'un de ces groupes.

Dans l'exemple de la Figure 2, cinq cas ont été identifiés comme anormaux par rapport à leurs homologues. Un indice d'anomalies identifie l'écart des zones du cas par rapport à ce qui constitue la norme du groupe d'homologues de ce cas. En outre, des informations sont fournies concernant les variables ayant eu le plus grand impact du fait de l'éloignement de chaque cas du centre de son groupe d'homologues.



	\$O-Anomaly	\$O-AnomalyIndex	\$O-PeerGroup	\$O-Field-1	\$O-FieldImpact-1	\$O-Field-2	\$O-FieldImpact-2
1	T	1.911	2	Sum Failed Access Attempts	0.324	Time of Day Access	0.080
2	T	1.880	2	Clearance Level	0.119	Clearance Level	0.108
3	T	2.141	1	IP Address Contacted	0.422	Clearance Level	0.074
4	T	1.896	1	Time of Day Access	0.112	Current Project	0.082
5	T	2.797	2	Sum Failed Access Attempts	0.501	IP Address Contacted	0.223

Figure 2 : Dans ces cinq cas d'anomalies, l'index d'anomalie identifie l'écart des zones du cas par rapport à la norme du groupe d'homologues de ce cas. Les zones 1 et 2 décrivent les variables qui ont eu le plus gros impact sur la détermination des comportements anormaux.

Une autre forme de détection des anomalies peut être exécutée à l'aide du text mining associé à l'analyse des correspondances afin de réaliser une analyse du réseau social. Lorsqu'il est important de déterminer la distance relative entre les catégories d'intérêt, l'analyse des correspondances permet d'obtenir cette perspective supplémentaire au sein d'un réseau. L'un des objectifs de l'analyse des correspondances consiste à décrire les relations entre deux variables nominales dans une table de correspondances dans un espace à faible dimension, tout en décrivant en même temps les relations entre les catégories de chaque variable.

Pour chaque variable, les distances entre les points de catégories d'un tracé correspondent à la force de la relation entre les catégories, les catégories similaires étant tracées plus proches l'une de l'autre.

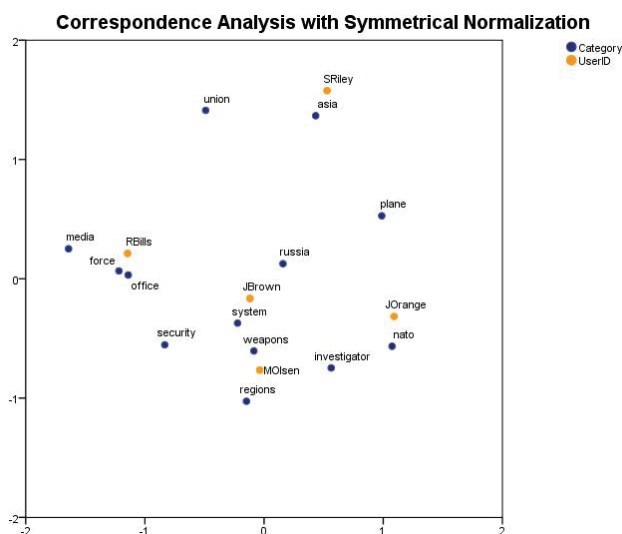


Figure 3: Mappage vectoriel fournissant un affichage visuel des résultats de l'analyse de texte et permettant de comprendre les domaines des sujets auxquels les salariés ont accédé.

Pour déterminer la distance entre les catégories, l'analyse des correspondances prend en compte les fréquences de cellules individuelles ainsi que plusieurs autres facteurs. Le calcul des fréquences de cellule est similaire à l'analyse des tabulations croisées. L'analyse des correspondances crée également plusieurs statistiques intermédiaires qui, entre autres, mesurent l'influence, la variance, et la distance entre un objet et un autre.

L'analyse des correspondances aide les analystes à comprendre les différences entre les catégories d'une variable ainsi que les différences entre des variables.

Le tracé de la Figure 3, ci-dessus, est le résultat de l'analyse des correspondances effectuée sur les salariés participant à un projet particulier et montre les catégories de sujet extraites de l'analyse de texte. Les points de ligne et de colonne qui sont les plus proches les uns des autres indiquent une correspondance ou une association plus étroite que celles des points plus distants les uns des autres. Dans ce tracé, le salarié SRiley est étroitement associé aux documents concernant l'« Asie, » tandis que le reste de ses homologues semblent accéder à d'autres rubriques (c'est -à-dire qu'ils y sont plus étroitement associés).

L'avantage du text mining pour ce type d'analyse est qu'il permet aux analystes de « lire » et de trier littéralement plusieurs milliers de documents afin de définir qui accède à quel contenu, en fonction du sujet, par rapport à leur groupe d'homologues.

Analyse de série temporelle permettant de visualiser les comportements futurs

Un avantage de l'analyse des menaces venues des initiés par rapport aux menaces externe est que dans certaines zones de données, la quantité des données disponibles est en général beaucoup plus complète et que des points de données peuvent en général être attribués à une personne, une date/heure et un événement spécifiques. Cette possibilité est utile car elle permet de réaliser une analyse de série temporelle. En dépit d'erreurs courantes de compréhension, l'analyse de série temporelle peut être utilisée non seulement pour tracer des données historiques, mais aussi afin de prévoir des points de données futurs, et de prendre en compte des facteurs tels que les événements saisonniers, les occurrences ponctuelles et les interventions ou les modifications des attentes. Pour en faire la démonstration, dans l'exemple suivant, l'utilisation de la bande passante pour chaque salarié a été « journalisée », puis a été utilisée dans un algorithme de série temporelle pour analyser les tendances d'utilisation.

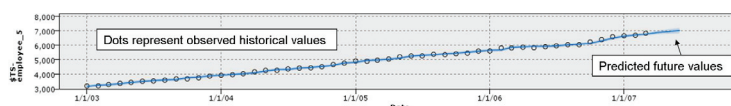


Figure 4 : Les modèles ARIMA (Auto-Regressive Integrated Moving Average) permettent non seulement de tracer les valeurs historiques et prévues par rapport aux valeurs réelles, mais aussi de prévoir les valeurs futures. Lorsque la valeur observée dépasse la valeur prévue, le score de risque augmente.

Il en résulte non seulement un agrégat d'utilisation historique du réseau, mais aussi un composant permettant les prévisions. Dans cet exemple, un modèle ARIMA (Auto-Regressive Integrated Moving Average) prend en compte plusieurs facteurs, tels que les modèles d'utilisation saisonniers et les comportements atypiques, pour prévoir l'utilisation de la bande passante pour plusieurs périodes. Si les valeurs observées sur les périodes futures dépassent les valeurs prévues, une augmentation d'un score de risque peut être déclenchée ou un auditeur peut être averti.

Résumé des méthodes analytiques pour la détection des menaces d'initiés

Très souvent, la meilleure approche de l'évaluation des risques en interne regroupe plusieurs méthodes idéalement adaptées aux objectifs spécifiques de l'agence et aux données disponibles. Comme pour tout type d'application adapté à l'analyse statistique et à l'exploitation des données, il est important de commencer par déterminer les objectifs et les problèmes potentiels découlant de l'analyse. Si l'on évalue et réduit la menace potentielle d'une attaque d'initié, ces objectifs peuvent être les suivants :

- déterminer une prédiction, un niveau de confiance et le score de propension au risque d'une attaque venue de l'intérieur,
- calculer le coût ou l'impact prévu de la violation de l'information,
- identifier l'ampleur et le périmètre d'une attaque dans les cas où la perte ou le coût de l'information n'est pas quantifiable,
- prioriser les audits des anomalies ou des menaces détectées en fonction du niveau de risque et des ressources disponibles pour mener l'enquête.

Méthodes de détection des menaces externes

Les techniques d'analyse des menaces internes sont également souvent utiles pour l'analyse des menaces externes. La principale différence entre l'analyse des menaces internes et l'analyse des menaces externes concerne la disponibilité des données. Les attaques provenant de sources externes fournissent rarement le type de données démographiques disponibles pour l'analyse des menaces internes. Les champs de données, tels que l'âge, l'appartenance à un groupe, la situation géographique et les modèles comportementaux historiques qui peuvent être attribués à un individu ou à un groupe sont beaucoup plus difficiles à obtenir lorsque vous analysez des menaces externes.

Analyse des réseaux sociaux pour compléter les données manquantes

Lorsque les menaces externes ne fournissent pas suffisamment d'informations sur l'individu ou le groupe responsable d'une attaque ou d'une attaque potentielle, l'utilisation des techniques d'analyse des réseaux sociaux (SNA - Social Network Analysis) peut aider les enquêteurs à mieux évaluer le risque d'une menace externe spécifique en faisant des associations avec d'autres individus, groupes ou précédentes tentatives d'attaque connus.

Un réseau social est une structure sociale composée de noeuds (qui sont en général des individus ou des organisations) liés entre eux par un ou plusieurs types spécifiques d'interdépendance (par exemple valeurs, visions, idées, échanges financiers, amitié, parenté, ou conflits).³

Les fonctions analytiques dans SNA comprennent la théorie des diagrammes ainsi qu'une analyse des liens. D'autres approches de SNA sont les règles d'association, l'analyse des correspondances, l'analyse de régression et l'analyse des séries temporelles.

A l'aide d'un atelier d'analyse prédictive généralisé, nous pouvons exécuter certains types d'analyses SNA via la combinaison d'algorithmes de classification et de régression, d'algorithmes d'association et d'une analyse des correspondances. Les algorithmes d'association, tels que les algorithmes apriori, sont utiles pour prédire des résultats multiples : par exemple, des personnes présentant un ensemble donné de caractéristiques ont des chances d'être associées à une personne, un lieu géographique ou une organisation donné(e).

L'avantage des algorithmes des règles d'association par rapport aux algorithmes classiques d'arbre de décisions est que des associations peuvent exister entre n'importe lesquels des attributs. En d'autres termes, un algorithme d'arbre de décisions créera des règles ayant une seule conclusion, tandis que les algorithmes d'association tentent de trouver de nombreuses règles, dont chacune peut avoir une conclusion différente. En outre, contrairement aux algorithmes d'analyse de lien simple, les algorithmes d'association permettent l'analyse de liens de un à plusieurs et de plusieurs à plusieurs.

Consequent	Antecedent	Support %	Confidence %
Associate1 = Ali Atwa	Associated Mosque = Flint Islamic Center Orig Country = Syria	1.961	100.0
	Associate2 = Khalfan Associated Mosque = Flint Islamic Center Location = Washington	1.961	100.0
Associate1 = Ali Atwa	Associate2 = Khalfan Orig Country = Syria Location = Washington	1.961	100.0
Associate1 = Ali Atwa	Associate3 = Abdul Wasay Aghajan Associated Mosque = Sheepside mosque Orig Country = Egypt	3.922	100.0
Associate1 = Ali Atwa	Associate3 = Abdul Wasay Aghajan Associated Mosque = Sheepside mosque Location = Washington	3.922	100.0
Associate1 = Ali Atwa	Associate3 = Abdul Wasay Aghajan Orig Country = Egypt	3.922	100.0

Figure 5 : Utilisée avec les techniques SNA, l'analyse de texte permet de détecter les relations entre les personnes concernées, les groupes, et les lieux géographiques, afin de prévoir les associations probables.

A la Figure 5, un organisme de renseignement souhaitait identifier des associations entre des individus et un ensemble spécifique de caractéristiques et de comportements. Après examen des enregistrements d'un regroupeur de distribution de nouvelles, un algorithme apriori a été appliqué à toutes les zones considérées comme intéressantes par un analyste de l'organisme.

Le résultat mis en évidence peut être interprété comme suit : concernant les enregistrements dans lequel l'individu concerné est associé à « Khalfan » et dont le pays d'origine est la « Syrie » et le lieu de résidence actuel est « Washington », « Ali Atwa » est alors un associé probable. Dans cet exemple, lorsque tous les antécédents étaient présents, le conséquent « Ali Atwa » était également constaté dans 100 % des cas. On peut donc en conclure que si l'individu concerné correspond à tous les antécédents de la règle de l'algorithme, il a aussi des chances d'être associé à l'individu indiqué dans le conséquent.

Une difficulté inhérente de SNA est l'impossibilité de créer efficacement des associations à partir d'une quantité illimitée de données disponibles. Comme dans la plupart des fonctions d'analyse des données, l'outil est surtout efficace lorsqu'il est utilisé par un expert de domaine qui peut concentrer l'analyse sur un ensemble particulier de données ou de zones.

A titre d'exemple, l'une des observations les plus courantes de SNA est que des individus similaires sur le plan démographique ont le plus de chance de former des liens sociaux. ⁴ Ce type d'observation peut être évident pour un analyste ayant une expérience de SNA, et c'est aussi le type de connaissance qui peut améliorer fortement la probabilité d'obtenir des résultats concrets de SNA dans des bases de données complexes de grande taille. Les algorithmes d'association peuvent tirer parti de cette expérience en générant un ensemble d'interactions et de mesures plus riche qu'une simple analyse de lien de un à un.

Par exemple, dans l'analyse de personnes similaires sur le plan démographique, la liste ou le graphique de chaque personne liée par une nationalité commune peut rapidement submerger l'analyste de données parasites. Par contraste, un algorithme d'association peut expliquer une affiliation à un groupe de nationalité, un groupe social ou un groupe professionnel, de même que la ville de résidence peut fournir un ensemble de caractéristiques concis qui a depuis toujours fait ressortir une association au sein d'un groupe d'individus.

Comme dans l'exemple précédent d'analyse des correspondances appliqué à une menace d'initié, la Figure 6 représente la même technique appliquée aux menaces externes. Les documents texte de format libre collectés auprès des distributions de nouvelles, des blogs et des forums de discussion Internet sont transmis via un algorithme d'analyse de texte (traitement automatique du langage naturel) afin de déterminer le ou les sujets abordés dans ces documents. Un algorithme dédié d'analyse de lien de texte est alors appliqué afin de connecter la mention d'une sécurité spécifique aux lieux concernés (dans ce cas, les villes). Les points de données obtenus sont transmis à un algorithme d'analyse de correspondance afin de mieux déterminer les sujets ayant une corrélation étroite avec les villes concernées.

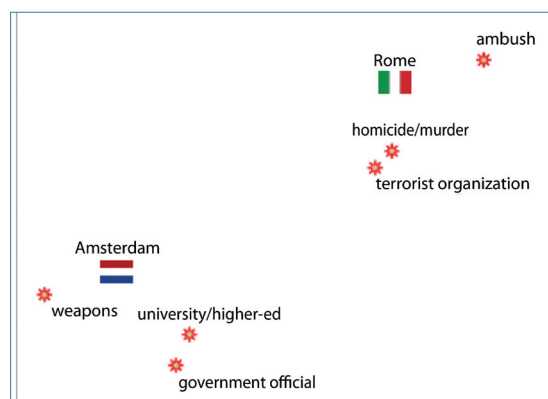
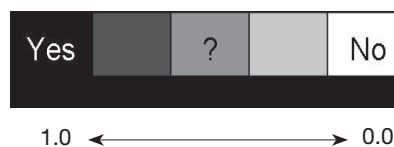


Figure 6 : L'analyse des correspondances des distributions de nouvelles, des blogs, des forums de discussion sur Internet, et des autres sources de texte peut être tracée afin de montrer les domaines de sujet qui sont couramment liés aux villes concernées.

Modèles de scoring mesurant l'impact potentiel

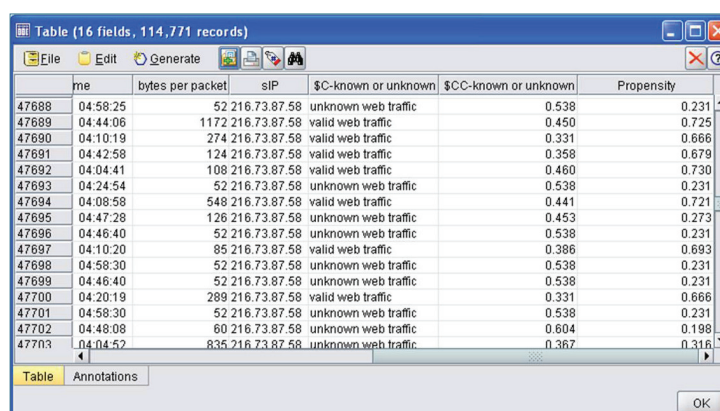
Lorsqu'un modèle a été créé, l'étape suivante consiste à mesurer la probabilité et l'impact potentiel d'un événement identifié ou prévu par le modèle.

Lors de la modélisation des événements de sécurité, les caractéristiques d'un cas sont rarement assez claires pour permettre de déterminer de façon absolue si le cas est positif ou négatif, ou de définir si un événement s'est produit ou non. Comme la plupart des cas relèvent d'une décision nuancée et non pas sans ambiguïté, il est souvent utile de convertir un événement de type non ou oui en une échelle de propension de 0 à 1, à l'aide des chiffres de fiabilité fournis par l'algorithme prédictif. Les scores de propension indiquent la probabilité d'un résultat ou d'une réponse spécifiques. Cela fournit une mesure de certitude et non une classification absolue pour un cas donné. Les décisions proches des extrêmes (1.0 ou 0.0) sont claires et celles du milieu sont plus incertaines.



Dans cet exemple, un « pot de miel », c'est-à-dire un piège réseau monté pour détecter et contrer les accès non autorisés à des informations sensibles, a été créé par un organisme qui cherchait à modéliser les comportements normaux/malveillants. Les « pots de miel » qui sont en général un ordinateur, un jeu de données, ou un site réseau comprenant des informations ou des ressources falsifiées susceptibles d'avoir de la valeur pour les agresseurs n'ont en règle générale aucune valeur réelle et ne doivent par conséquent pas faire l'objet d'opérations ou d'activités légitimes. Les informations de trafic qu'ils captent peuvent donc être considérées comme malveillantes ou non autorisées. Dans l'exemple ci-dessous, toutes les données collectées auprès du « pot de mie » ont été étiquetées comme malveillantes. Un jeu de données contenant du trafic normal connu a ensuite été ajouté au jeu de données du pot de miel, afin d'établir la différence entre un accès autorisé et un accès incorrect. Un algorithme de classification a ensuite été utilisé pour créer des profils de trafic normal et de trafic malveillant. En outre, l'algorithme fournissait un niveau de fiabilité pour chaque prévision (classification).

Un score de fiabilité est dérivé des chiffres de fiabilité de façon à avoir les propriétés d'un score, c'est-à-dire proche de 0 lorsque le trafic a des chances d'être normal et proche de 1 lorsque le trafic a des chances d'être malveillant.



	me	bytes per packet	sIP	\$C-known or unknown	\$CC-known or unknown	Propensity
47688	04:58:25	52	216.73.87.58	unknown web traffic	0.538	0.231
47689	04:44:06	1172	216.73.87.58	valid web traffic	0.450	0.725
47690	04:10:19	274	216.73.87.58	valid web traffic	0.331	0.666
47691	04:42:58	124	216.73.87.58	valid web traffic	0.358	0.679
47692	04:04:41	108	216.73.87.58	valid web traffic	0.460	0.730
47693	04:24:54	52	216.73.87.58	unknown web traffic	0.538	0.231
47694	04:08:58	548	216.73.87.58	valid web traffic	0.441	0.721
47695	04:47:28	126	216.73.87.58	unknown web traffic	0.453	0.273
47696	04:46:40	52	216.73.87.58	unknown web traffic	0.538	0.231
47697	04:10:20	85	216.73.87.58	valid web traffic	0.386	0.693
47698	04:58:30	52	216.73.87.58	unknown web traffic	0.538	0.231
47699	04:46:40	52	216.73.87.58	unknown web traffic	0.538	0.231
47700	04:20:19	289	216.73.87.58	valid web traffic	0.331	0.666
47701	04:58:30	52	216.73.87.58	unknown web traffic	0.538	0.231
47702	04:48:08	60	216.73.87.58	unknown web traffic	0.604	0.198
47703	04:04:52	835	216.73.87.58	unknown web traffic	0.367	0.316

Figure 7 : Les scores de propension brute calculés à partir des niveaux de fiabilité permettent de déterminer la probabilité d'une réponse ou d'un résultat donné, mais ils ne sont pas fiables tant qu'ils n'ont pas été testés par rapport à des données supplémentaires et ajustés en fonction de ces dernières.

Les scores de propension bruts fournis par un outil d'analyse sont basés purement sur des estimations données par le modèle et peuvent par conséquent être exagérés, et entraîner des estimations de propension inexactes. Les propensions ajustées tentent de compenser cela en examinant le comportement du modèle sur le test ou sur des partitions de validations, et en ajustant les propensions afin d'obtenir une meilleure estimation.

Matrices de risque pour la classification de la gravité des menaces

Jusqu'à présent, nous avons fait principalement porter l'analyse sur l'identification des anomalies, la prévision des événements et l'identification d'associations entre des individus et des groupes. La plupart de ces techniques fournissent en général une mesure de distance des anomalies, un niveau de fiabilité des prévisions et des classifications (comme dans l'exemple ci-dessus) et des mesures de proximité et de force des relations pour les associations et l'analyse du réseau.

Afin de déterminer le niveau de menace général d'un événement, nous pouvons utiliser ces mesures pour remplir une matrice de risques. En général, les matrices de risques permettent de déterminer la gravité du risque de voir se produire un événement. Le risque d'un événement donné peut être défini comme sa probabilité multipliée par sa conséquence (son impact).

Le graphique ci-dessous montre un exemple d'une matrice de risque 3x3. Si nous souhaitions établir plus de distinctions entre les niveaux de risque, la matrice de risque pourrait facilement être développée en une matrice de 4x4, 5x5, ou plus grande.

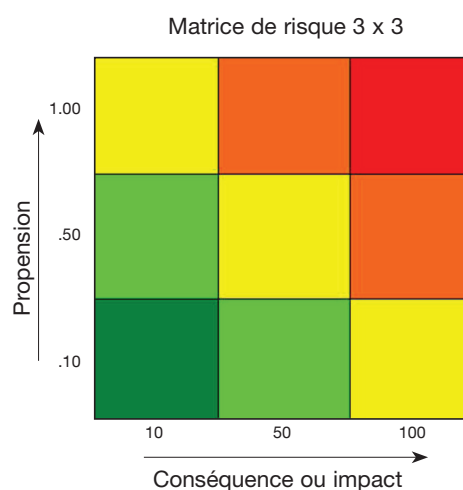


Figure 8 : Pour obtenir un score de menace général, la valeur numérique de la probabilité de l'événement peut être multipliée par la valeur numérique de l'impact si l'événement se produit. Le tracé de ce résultat sur une matrice de risque codée par couleur aide à démontrer visuellement la gravité d'une menace. ⁵

Si nous voulons exprimer sous forme numérique les valeurs d'une matrice de risque, il existe plusieurs formules mathématiques et calculateurs d'évaluation du risque pour certains secteurs d'activité et applications. Pour conserver un exemple relativement simple, nous allons utiliser une formule simple qui mesure les valeurs d'impact sur une échelle de 1 à 100 et les valeurs de propension (probabilité) sur une échelle de 0 à 1.00. Pour obtenir un score de menace général, nous pouvons multiplier la valeur numérique de la probabilité de l'événement par la valeur numérique de l'impact si l'événement se produit.

Par exemple, l'impact d'un initié divulguant des informations à haute sensibilité depuis un organisme de renseignements peut être considéré comme « élevé » dans une matrice 3x3. Toutefois, la probabilité de cet événement peut être considérée comme « faible ». Dans ce cas, le niveau de menace global exprimé sous forme numérique est dérivé en multipliant 100 (impact élevé) par 0,10 (faible probabilité), ce qui donne un niveau de menace global de 10.

L'exemple de la Figure 9 ci-dessous montre un modèle d'arborescence de classification et de régression qui est utilisé pour déterminer la probabilité qu'une demande réseau émane d'une source connue et sûre (valide). Le score de propension est alors utilisé pour créer trois catégories (faible, moyen et élevé) sur la base d'une matrice de risque 3x3.

	he	bytes per packet	bytes	flags	Classification	Confidence Level	Propensity	Propensity Category
1	04:25:37	52	52 A		unknown web traffic	0.574	0.426	Medium
2	04:17:51	365	1094	FSPA valid web traffic		0.692	0.692	High
3	04:25:37	277	830	FSPA valid web traffic		0.778	0.778	High
4	04:18:36	874	6117	FSPA valid web traffic		0.681	0.681	High
5	04:27:26	245	734	FSPA valid web traffic		0.778	0.778	High
6	04:28:58	97	292	FSPA unknown web traffic		0.671	0.329	Low
7	04:30:18	1207	8452	PA	unknown web traffic	0.836	0.164	Low
8	04:28:06	1140	148	FSPA unknown web traffic		0.530	0.470	Medium
9	04:20:07	1088	7619	FSPA valid web traffic		0.617	0.617	Medium
10	04:17:35	195	586	FSPA valid web traffic		0.655	0.655	Medium
11	04:17:23	413	1651	FSPA valid web traffic		0.560	0.560	Medium
12	04:17:25	1037	8293	FSPA valid web traffic		0.617	0.617	Medium
13	04:20:15	327	980	FSPA valid web traffic		0.516	0.516	Medium
14	04:26:11	1061	169	FSPA unknown web traffic		0.509	0.491	Medium
15	04:28:06	386	1159	FSPA valid web traffic		0.774	0.774	High
16	04:31:21	1177	364	FSPA unknown web traffic		0.594	0.406	Medium
17	04:23:03	355	2487	FSPA unknown web traffic		0.923	0.077	Low
18	04:17:23	997	5981	FSPA valid web traffic		0.681	0.681	High
19	04:25:37	52	52 A		unknown web traffic	0.574	0.426	Medium

Figure 9 : Sur la base d'une matrice de risque 3x3, les événements (tels qu'un trafic Web anormal) peuvent être classifiés en tant que risque faible, moyen ou élevé.

Pour cet exemple, nous supposons que l'impact global d'une demande réseau émanant d'une source inconnue ou non confirmée pose un risque « moyen » pour la sécurité du réseau. Une nouvelle zone classifiant l'impact global comme « Moyen » pour une matrice de 3x3 (sous forme numérique, un score d'impact de 50) est créée.

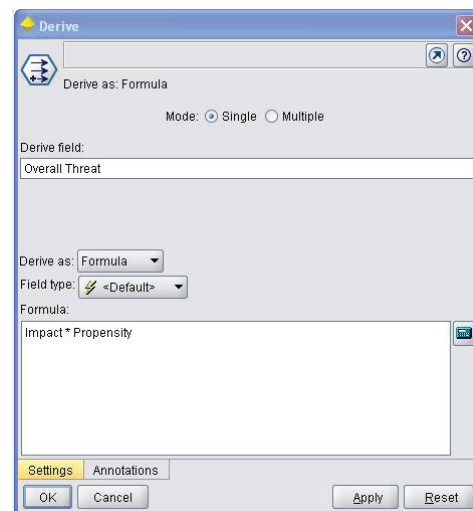


Figure 10 : Les algorithmes peuvent être configurés pour calculer le score de menace global, en prenant en compte le type de l'événement, sa gravité et sa probabilité au sein d'un seul score combiné.

Pour calculer la menace globale, nous créons une zone qui multiplie l'« Impact » (50 dans ce modèle) par la « Propension ». Le résultat est un score de menace global (Figure 10) qui peut être placé dans un environnement opérationnel avec d'autres modèles mesurant différents types d'événements de sécurité (éventuellement avec des niveaux d'impact différents). Lorsque tous les modèles sont combinés, nous pouvons trier et classer tous les événements entrants en fonction de son score de menace global. Cette valeur nous permet de prioriser une réponse à l'événement tout en expliquant le type de l'événement, sa gravité et sa probabilité au sein d'un seul score combiné.

Récapitulatif des méthodes analytiques pour la détection des menaces externes

Comme dans le cas de l'analyse des menaces internes, la meilleure approche de l'évaluation des risques externes consiste à se concentrer sur un objectif spécifique pouvant être réalisé, ou au moins amélioré, grâce à l'analyse des données. Ce n'est pas toujours une tâche facile. Lors de l'analyse des menaces externes, la quantité des données disponibles est parfois supérieure à celle qui peut effectivement être analysée. A d'autres moments, lors de l'analyse d'événements très rares, tels qu'une agression chimique, il existe peu de données historiques et une corrélation très médiocre, voire nulle, entre un événement et un autre.

Outre qu'elle facilite la détection des variations constantes des schémas d'attaques, l'opérationnalisation de ces processus comporte un avantage supplémentaire, qui est celui de rationaliser le déploiement des résultats. Pour préciser : lorsqu'un modèle est déployé dans les processus métier d'un organisme, il ne s'exécute pas simplement de son propre chef.

Au lieu de cela, l'atelier d'analyse prédictive collecte et distribue activement l'information de façon bidirectionnelle. En d'autres termes, lorsqu'un modèle est déployé en temps réel ou en mode par lots, il collecte automatiquement l'information d'autres sources de données (telles que les systèmes d'information décisionnelle et d'ERP), mais peut aussi renvoyer ces résultats vers ces systèmes pour mettre à jour les enregistrements ayant des scores de propension et/ou d'autres variables de classification. En ajoutant de nouvelles variables aux enregistrements dans les autres jeux de données, ce déploiement de résultats permet aux utilisateurs finaux d'accéder aux informations quand et là où ils le veulent, dans le format de leur choix.

À propos des solutions IBM SPSS d'analyse prédictive

Des perspectives et des prévisions de meilleure qualité

L'analyse prédictive apporte aux entreprises une vue plus claire des conditions existantes et une perspective plus approfondie des événements futurs. Grâce à IBM SPSS Modeler, notre atelier d'analyse prédictive leader sur le marché, votre organisme peut réaliser des analyses incorporant de nombreux types de données, et donnent une connaissance plus approfondie de chaque aspect de vos opérations, notamment une compréhension plus complète de vos données de veille.

IBM SPSS Modeler est une solution ouverte et normalisée. Il s'intègre aux systèmes d'information existant de votre entreprise, à la fois lors de l'accès aux données et lors du déploiement des résultats. Vous n'avez pas besoin de convertir les données dans ou vers un format propriétaire. Cela vous aide à conserver les ressources, fournit des résultats plus rapidement, et réduit les coûts d'infrastructure.

En outre, IBM SPSS Modeler est apprécié des analystes et des utilisateurs professionnels partout dans le monde, car il vous permet d'effectuer les opérations suivantes :

- Consulter, préparer et intégrer facilement des données structurées et aussi du texte, des données Web et des données d'enquête.
- Créer et valider rapidement des modèles, à l'aide des techniques statistiques et d'apprentissage automatique les plus sophistiquées qui existent.
- Déployer efficacement des perspectives et des modèles prédictifs sur une base planifiée, ou en temps réel, pour les personnes prenant des décisions et émettant des recommandations, et pour les systèmes qu'elles utilisent.

Tirer parti de toutes vos données pour obtenir des modèles améliorés

IBM SPSS Modeler est la seule solution qui vous permet d'accéder directement et facilement à des données texte, Web et d'enquêtes, et d'intégrer ces types de données traditionnels à vos modèles prédictifs. Les clients IBM SPSS se sont aperçus que le recours à des types supplémentaires de données augmente l'utilité ou la précision des modèles prédictifs, et aboutit à des recommandations plus utiles et à de meilleurs résultats. Avec le produit IBM SPSS Text Analytics entièrement intégré, vous pouvez extraire des concepts et des opinions de n'importe quel type de texte, par exemple des emails, un dialogue dans un « chat », des blogs, etc.

Automatisation sur les processus analytiques critiques

Les tâches associées au développement et au déploiement de modèle analytique et prédictif sont souvent répétées régulièrement. IBM SPSS Collaboration and Deployment Services vous aide à améliorer la productivité, à garantir la cohérence, et obtenir plus de précision dans ces processus, en fournissant un environnement puissant qui automatise les différentes étapes du processus analytique, telles que la préparation des données, les transformations, la création de modèle, l'évaluation et le scoring. Par conséquent, vos analystes peuvent se concentrer sur la résolution des problèmes métier, au lieu de créer et d'exécuter manuellement les processus de chaque nouveau projet.

En outre, vos analystes produisent une grande quantité de données en sortie de grande valeur : scores, règles, graphiques, rapports, autres types d'informations analytiques. Pour retirer une valeur optimale de vos solutions analytiques, vous devez fournir ou déployer les résultats pour vos décideurs dans toute l'organisation, d'une façon qui les aide à prendre de meilleures décisions.

Pour ce faire, IBM SPSS Collaboration and Deployment Services vous permet de fournir les éléments suivants :

- Des scores des enregistrements des personnes sur lesquelles vous enquêtez, susceptibles de faire apparaître la probabilité de leurs liens à une organisation terroriste connue
- Des jeux de règles ou de critères définissant un profil d'activités normales pour un segment précis de salariés
- Des graphiques comparant la précision de plusieurs modèles de risques.
- Des rapports démontrant la précision d'une prévision par rapport aux résultats réels

** Les anciens noms d'IBM SPSS Text Analytics et IBM SPSS Collaboration and Deployment Services sont PASW Text Analytics et PASW Collaboration and Deployment Services.

À propos d'IBM Business Analytics

Les logiciels IBM Business Analytics aident les entreprises à mesurer, comprendre et anticiper leur performance financière et opérationnelle en fournissant des informations exactes, cohérentes et complètes. Ce portefeuille complet de solutions décisionnelles, d'analyse prédictive, de gestion des performances financières et de stratégies et d'applications d'analyse fournit une vision claire, immédiate et concrète des performances actuelles et permet de prévoir les résultats futurs. Alliant de riches solutions sectorielles, des services professionnels et les meilleures pratiques, la solution permet aux entreprises de toutes tailles d'optimiser leur productivité, d'automatiser en toute confiance leurs décisions et de générer de meilleurs résultats.

Faisant partie de ce portefeuille, la solution logicielle IBM SPSS Predictive Analytics permet aux entreprises de prévoir les événements et d'agir proactivement en fonction des données obtenues pour obtenir de meilleurs résultats métier. Les clients des secteurs privé, public et universitaire du monde entier font appel à la technologie IBM SPSS, qui les dote d'un atout concurrentiel pour attirer, fidéliser et développer leur clientèle, tout en réduisant les fraudes et en atténuant les risques. En intégrant les logiciels IBM SPSS à leurs opérations quotidiennes, les entreprises deviennent proactives, capables de diriger et d'automatiser leurs décisions pour atteindre leurs objectifs métier et se doter d'un atout concurrentiel. Pour plus d'informations ou pour contacter un interlocuteur IBM, visitez le site ibm.com/spss/fr

Remarques :

1. Atelier Défi 2005 ARDA. « Insider Threat : Analysis and Detection of Malicious Insiders ».
2. Enquête de surveillance 2004 eCrime. Réalisée par la revue CSO Magazine en coopération avec les Services secrets américains & CERT® Centre de coordination : « Insider Threat Study : Computer System Sabotage in Critical Infrastructure Sectors ».
3. Peter Sheridan Dodds, Roby Mubamad, Duncan J. Watts. « An Experimental Study of Search in Global Social Networks ». (8 août 2003)
4. Aili E. Malm et al. « Social Network and Distance Correlation of Drug Production ».
5. Wikipédia. « Risk Matrix ». [http : //en.wikipedia.org/wiki/Risk_Matrix](http://en.wikipedia.org/wiki/Risk_Matrix)



Compagnie IBM France

17 Avenue de l'Europe
92 275 Bois-Colombes Cedex

L'adresse de la page d'accueil IBM est :

ibm.com

IBM, le logo IBM, ibm.com, WebSphere, InfoSphere et Cognos sont des marques d'International Business Machines aux Etats-Unis et/ou dans certains autres pays. Si ces marques et d'autres marques d'IBM sont accompagnées d'un symbole de marque (® ou ™), ces symboles indiquent qu'il s'agit de marques déposées aux Etats-Unis ou reconnues par la législation générale comme étant la propriété d'IBM au moment de la publication de ce document. Ces marques peuvent également exister et éventuellement avoir été enregistrées dans d'autres pays. La liste actualisée de toutes les marques d'IBM est disponible sur la page Web « Copyright and trademark information » à l'adresse suivante : ibm.com/legal/copytrade.shtml

SPSS est une marque de SPSS Inc., une société du groupe IBM, dans de nombreux pays.

Les raisons sociales, noms de produit et de service peuvent être des marques déposées de leurs propriétaires respectifs.

US Government Users Restricted Rights - Use, duplication of disclosure restricted by GSA ADP Schedule Contract with IBM Corp.

© Copyright IBM Corporation 2012



Merci de recycler ce document