

Un déluge de données

Serge Abiteboul et Pierre Senellart

Dans de nombreux domaines, scientifiques ou non, les données s'accumulent en masse.

Les gérer et les exploiter est le défi posé à l'informatique du Big Data.

Au CERN, près de Genève, le collisionneur LHC (*Large Hadron Collider*) est équipé d'énormes détecteurs capables d'enregistrer les traces de dizaines à centaines de millions de collisions proton-proton par seconde. Évaluons grossièrement le volume de données que cela représente, en faisant des hypothèses basses. L'information relative aux produits de chaque collision est représentée à l'aide de quatre octets (soit $(2^8)^4 = 256^4 = 4,3$ milliards de possibilités), l'accélérateur fonctionne dix heures par jour en moyenne, et 100 millions de collisions sont enregistrées chaque seconde. On calcule facilement qu'au bout d'un an, le LHC produit ainsi 5×10^{15} octets, soit cinq pétaoctets d'information (une évaluation plus précise fournit des valeurs supérieures, voir l'article de F. Malek dans ce numéro). Il faut 5 000 disques durs du commerce pour stocker une telle masse de données.

Cependant, ce n'est pas le stockage qui est la difficulté principale, mais l'exploitation : extraire de l'information utile – de la connaissance – de cette masse de données dépasse largement les capacités humaines.

Pour relever de tels défis, une nouvelle science est apparue, la science des données. Il s'agit d'extraire des contenus intéressants à partir de masses de données gigantesques, changeantes, hétérogènes, incertaines, et ce à l'aide d'algorithmes efficaces. L'objet de cet article est de décrire les principaux enjeux de cette science des données et les difficultés qu'elle doit surmonter.

L'exemple du LHC par lequel nous avons commencé est extrême en termes de volumes d'octets, mais de nombreux autres domaines scientifiques et industriels doivent faire face à une avalanche similaire de données. En recherche médicale, par exemple, l'un des problèmes a longtemps été l'absence de données en nombre suffisant ; aujourd'hui, c'est plutôt l'inverse : les chercheurs sont confrontés à la difficulté de produire des connaissances utiles à partir des myriades de données accumulées par les différents essais cliniques, les services hospitaliers, etc. Citons la base de données MIMIC-II (*Multiparameter Intelligent Monitoring in Intensive Care*), constituée aux États-Unis par une équipe de chercheurs du

ANALYSE DES DONNÉES : LES TÂCHES CLASSIQUES

- Rechercher des données et les acquérir.
- Homogénéiser les données provenant de plusieurs sources.
- Détecter et éliminer les doublons et les erreurs.
- Interagir avec des humains pour obtenir plus de données, résoudre les contradictions et combler les manques (« crowdsourcing »).
- Faciliter le travail des analystes en leur fournissant des outils de visualisation
- Réaliser des analyses statistiques automatiques des données.
- Développer des applications et de nouveaux services.

© D'après Matvienko Vaidini/Shutterstock.com

MIT, de Philips et du Centre médical Beth Israel Deaconess. Elle regroupe des données médicales portant sur 32 000 personnes hospitalisées en unités de soins intensifs, soit plus de 40 000 séjours. Le volume de données est nettement plus modeste que celui produit par le LHC, mais son analyse n'en nécessite pas moins des compétences spécifiques et des algorithmes complexes.

Les exemples de domaines scientifiques où l'on doit recueillir et gérer des données massives ne manquent pas. On peut citer la génomique (voir l'article de G. Perrière), les relevés astronomiques (voir l'article de C. Reylé), les recensements de la flore et de la faune, la recherche pharmacologique, les études démographiques, etc.

On le comprend, des algorithmes d'analyse de données massives sont indispensables à la recherche scientifique d'aujourd'hui. Mais il en faut également pour des applications plus quotidiennes. En voici un exemple.

Vous voulez choisir un film pour votre soirée ; vous vous connectez à Internet et allez sur un site détaillant les possibilités. Son système de recommandation vous

fera des propositions. Il peut procéder de manière très simpliste, par exemple en vous suggérant le film récent le plus regardé au cours de la semaine, ou celui ayant reçu les meilleures critiques. Mais la recommandation peut aussi résulter d'un algorithme complexe qui prend en compte les films que vous avez aimés, votre humeur, peut-être les goûts de celui ou celle qui partagera votre soirée. Pour vous aider à répondre à la simple question « Quel film pourrais-je regarder ? », l'algorithme tiendra peut-être compte d'un océan de données : des avis de centaines de millions de personnes sur des milliers de films.

Or les programmes qui permettent de réaliser des analyses statistiques, même très simples, sur de telles quantités de données sont d'une grande complexité. Dans cette masse, le logiciel de recommandation de films trouvera des personnes qui ont des goûts voisins des vôtres, révélera des proximités avec des internautes que vous ne connaissez pas. Il pourra découvrir vos goûts cinématographiques et vous proposer des films en moins d'une seconde !

Dans la jungle des chaînes de télévision de plus en plus nombreuses, des VoD (*Video on Demand*, vidéo à la demande) et SVoD (*Subscription Video on Demand*, vidéo à la demande avec abonnement), des innombrables films disponibles légalement ou pas sur la Toile, l'utilisateur est perdu ; le but des systèmes de recommandation est de l'aider à s'y retrouver.

L'ingrédient principal qui alimente de tels systèmes est bien sûr constitué par les données numériques, lesquelles tiennent une place de plus en plus importante dans le monde d'aujourd'hui. Depuis les années 1960, les logiciels de bases de données se sont imposés pour partager des données au sein d'une entreprise ou d'une organisation. D'abord isolées dans des centres de calcul, ces données sont devenues accessibles partout dans le monde avec l'arrivée d'Internet, le réseau des réseaux de machines, puis du Web, le réseau des contenus, et finalement du Web 2.0 avec la participation de chacun, les réseaux d'individus.

Et il n'y a pas que les données stockées, il y a aussi les données échangées. Nous sommes entourés de milliards d'objets communicants. En 2008, le Web comptait déjà plus de 1000 milliards de pages et, chaque mois, les internautes y réalisaient des dizaines de milliards de recherches. On estime que le monde numérique double en volume tous les 18 mois et le trafic sur Internet est déjà, chaque année, supérieur à tout ce que l'on pourrait stocker en utilisant tous les disques et autres supports disponibles.

Analyser les données pour les valoriser

L'ensemble de ces données disponibles sur le réseau mondial constitue un énorme gisement de connaissances à découvrir et à valoriser. L'analyse de données a été un domaine très actif presque depuis les débuts de l'informatique et a revêtu divers noms, tels que *fouille de données* ou *business intelligence*. En raison de l'accroissement des capacités des disques et des mémoires, ainsi que des puissances de calcul avec des grappes d'ordinateurs pouvant réunir jusqu'à des milliers de machines, en raison aussi de l'explosion du volume de données disponibles, l'analyse de données pour en extraire de la valeur est devenue une industrie florissante. Et c'est sous le nom

de *Big Data* (données massives) qu'elle se développe aujourd'hui.

Le point de départ est de valoriser les gisements massifs de données. Le *Big Data* inclut en général deux aspects. D'une part, il sous-entend l'idée de croiser des données très structurées, par exemple celles d'une entreprise, avec des masses de données moins structurées et plus défectueuses disponibles sur le Web. D'autre part, il nécessite la mise en œuvre de calculs massivement parallèles en utilisant des techniques logicielles telles que *Hadoop*, issues des moteurs de recherche du Web (voir l'encadré page ci-contre).

Des difficultés multiples

L'objectif étant de faire émerger de nouvelles connaissances à partir des données, les tâches sont celles, classiques, de l'analyse de données, qui vont de leur acquisition à leur exploitation (voir l'encadré page 33). Les difficultés sont nombreuses et tiennent en premier lieu des quatre « V » du *Big Data* : le volume des données (il se compte typiquement en téraoctets ou pétaoctets, soit 10^{12} ou 10^{15} octets), leur variété ou hétérogénéité (sur le plan de la structure, de la langue, du format, etc.), leur vélocité (le rythme auquel elles sont modifiées), leur véracité (erreurs, incomplétude, confiance, provenance, fraîcheur, etc.).

D'autres difficultés proviennent de la répartition des données dans l'espace, des protections éventuelles (droit d'accès, restriction sur leur usage), etc. Par ailleurs, la nature des traitements auxquels ces données sont soumises a une importance considérable. Par exemple, un algorithme dont le temps d'exécution est proportionnel au cube du nombre de données reste inutilisable sur une base contenant un milliard d'enregistrements, même avec des milliers de machines qui fonctionneraient pendant des centaines d'années. Quant aux logiciels tels que *Hadoop*, développés spécifiquement pour traiter des données massives, ils sont encore relativement jeunes et compliqués à utiliser, en particulier parce qu'ils font travailler ensemble un grand nombre de machines.

L'analyse de grands volumes de données intervient dans d'innombrables domaines, afin de faire des prédictions de plus en plus fines, d'anticiper des épidémies, de mieux comprendre l'évolution du climat, d'aider à soigner les cancers, etc. Au niveau du

LES AUTEURS



Serge ABITEBOUL est directeur de recherche à l'INRIA-Saclay et à l'École normale supérieure de Cachan. Il est membre du Conseil national du numérique.



Pierre SENELLART est maître de conférences en informatique à Télécom ParisTech.

BIBLIOGRAPHIE

S. Abiteboul, *Sciences des données : de la logique du premier ordre à la Toile*, Leçon inaugurale au Collège de France, Fayard, 2012 : www.college-de-france.fr/site/serge-abiteboul/

S. Abiteboul et al., *Web Data Management*, Cambridge University Press, 2011 : <http://webdam.inria.fr/Jorge>

D. Agrawal et al., *Big data and cloud computing : current state and future opportunities*, EDBT/ICDT 2011, www.edbt.org/Proceedings/2011-Uppsala/papers/edbt/a50-agrawal.pdf

C. Lynch, *Big data : How do your data grow ?*, *Nature*, vol. 455, pp. 28-29, 2008.

G. Linden et al., *Amazon.com recommendations : Item-to-item collaborative filtering*, *IEEE Internet Computing*, vol. 7(1), pp. 76-80, 2003.