



**Informatique distribuée : la gestion de données
pair-à-pair par Serge Abiteboul. Recherche
collaborative, entretien avec Anne -Marie Kermarrec,
propos recueillis par Dominique Chouchan.**

Serge Abiteboul, Anne-Marie Kermarrec

► **To cite this version:**

Serge Abiteboul, Anne-Marie Kermarrec. Informatique distribuée : la gestion de données pair-à-pair par Serge Abiteboul. Recherche collaborative, entretien avec Anne -Marie Kermarrec, propos recueillis par Dominique Chouchan.. Les Cahiers de l'INRIA - La Recherche, 2008, Soleil et climat : la polémique, 425 décembre 2008. inria-00536672

HAL Id: inria-00536672

<https://inria.hal.science/inria-00536672>

Submitted on 16 Nov 2010

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

INFORMATIQUE DISTRIBUÉE

La gestion de données pair-à-pair

Le pair-à-pair, déjà adopté par des amateurs de musique, permet de partager des tâches ou des ressources entre de nombreux partenaires, à la fois clients et serveurs, au lieu d'utiliser un système centralisé.

En une quinzaine d'années à peine, le web a révolutionné nos modes de travail et notre vie sociale. Des milliards de pages sont aujourd'hui accessibles sur des millions de serveurs web, sur des milliards de téléphones portables et assistants personnels, sur des milliards d'objets communicants. Cette croissance considérable du volume de données et du nombre d'utilisateurs a toutefois conduit à augmenter la taille des serveurs ou à créer des plates-formes s'appuyant sur d'énormes grappes de machines comme les fermes de serveurs Google*. Cette hypercentralisation atteint cependant ses limites. Le web serait-il victime de son succès ? Une solution alternative se développe depuis peu, fondée sur un tout autre concept : la gestion des données dite pair-à-pair (*peer-to-peer*)⁽¹⁾.

Le web, c'est surtout des histoires de données. La gestion de données est au cœur de ses plus beaux succès : l'index du web de Google, les catalogues de vente d'Amazon, les réseaux sociaux de Facebook, les vidéos de Youtube, les photos de Flickr, les cartes géographiques de Google

© INRIA-GANG



Fig. 1 : Dans un réseau pair-à-pair, chaque ordinateur est à la fois client et serveur. Une fois connecté au réseau, via un logiciel commun, chacun peut participer aux tâches communes, par exemple échanger des fichiers de musique. Sur ce graphique, les points représentent des machines autonomes, et les lignes les communications entre ces machines.

Maps, l'encyclopédie Wikipedia, les profils de personnes d'un site de rencontre comme Meetic. Une très large majorité des sites proposant du contenu sur le web sont aujourd'hui articulés autour d'un système de gestion de bases de données (SGBD) relationnel, un des grands succès de l'informatique de la fin du siècle dernier (voir l'encadré).

Mais le monde a changé ! Il a changé d'abord par la nature de l'information que l'on manipule avec des structures plus riches (souvent des arbres comme en HTML* ou XML*, les standards pour les données du web) et surtout moins rigides et plus dynamiques. Il a changé aussi par le volume des données que l'on considère, le nombre de sources de données et leur distribution géographique. Il a changé

SYSTÈMES DE GESTION DE BASES DE DONNÉES EN BREF

Un SGBD relationnel est un logiciel qui permet de traiter de gros volumes de données, stockées dans des tableaux à deux dimensions, sur un serveur centralisé⁽³⁾. La technologie est fondée sur des bases mathématiques rigoureuses et sûres, le calcul des prédicats du premier ordre* en particulier. Elle est le fruit d'une recherche fondamentale⁽⁴⁾ qui s'est développée en lien étroit avec une industrie florissante (IBM, Oracle ou Microsoft). Ces logiciels sont devenus la norme pour le stockage de données.

enfin parce que tout un chacun peut maintenant publier des données sur le web, sans pour autant avoir les compétences d'un administrateur de SGBD. Si elle n'est pas optimale du point de vue des performances, la centralisation des données apparaît aussi de plus en plus discutable au plan socio-économique. Dans de nombreuses applications, on préférerait éviter de donner le contrôle de toute l'information à un seul partenaire.

Une solution alternative existe, la gestion de données en pair-à-pair (P2P). Un tel système est constitué d'un très grand nombre de machines autonomes qui coopèrent pour réaliser une tâche (fig. 1) : ce sont des pairs égaux en droits, tour à tour clients et serveurs pour d'autres pairs. Un pair peut être un ordinateur personnel, une boîte triple play*, un assistant personnel, un simple objet communicant. Les tâches P2P sont réalisées en parallèle avec les tâches « normales » des pairs, comme télécharger de la musique tout en travaillant.

Comment de tels systèmes se substitueraient-ils à des systèmes centralisés pour gérer des

informations (rapports, photos, courriels, contacts, catalogues) mis en commun par une communauté d'utilisateurs? D'un point de vue technique, le nœud du problème réside dans l'hétérogénéité des machines, des systèmes d'exploitation, des langages utilisés pour les applications, des modèles de données... Pour tout cela, la gestion de données distribuées a longtemps constitué un écueil difficile à contourner⁽²⁾. Les applications mettant en jeu de la distribution de données n'impliquaient qu'un nombre modeste de machines (deux, trois, une dizaine) et demandaient des investissements lourds. Les standards du web ont changé la donne. On dispose maintenant d'un modèle d'échange de données XML, de modèles plus ou moins acceptés pour décrire des connaissances, comme RDF* et Owl*, qui vont permettre aux systèmes de comprendre la nature des données qu'ils trouvent sur le web. On dispose surtout des services web, c'est-à-dire d'un protocole de communication entre machines indépendant de leurs particularités qui rend possible le calcul distribué et permet

* Le **moteur de recherche Google** utilise des centaines de milliers de machines, organisées en grappes de plusieurs milliers. Ces grappes sont appelées des « fermes ».

* Le **calcul des prédicats du premier ordre**, ou logique du premier ordre, est une formalisation du langage des mathématiques proposée au début du XX^e siècle, bien avant les débuts officiels de l'informatique.

* Le **langage HTML** est le langage de balises le plus utilisé dans les pages du web. Il permet de décrire des documents textuels sous forme arborescente et de spécifier les liens qui font passer d'une page du web à l'autre.

* Le **langage XML** est un langage fondé sur des balises qui permettent de structurer l'information. C'est le standard pour l'échange d'information sur le web.

* Une **boîte triple play** comme la Freebox permet d'avoir accès à trois types de services : internet, téléphonie et télévision.

* **RDF** (*Ressource Description Framework*) est un langage proposé pour décrire le contenu de sources d'information sur le web et **Owl** (*Web Ontology Language*) un langage de représentation des connaissances.

* **BitTorrent** est un protocole très populaire pour le partage de fichiers en P2P.

WebContent, une plate-forme pour la gestion de contenus

Le but du projet WebContent de l'Agence nationale de la recherche (ANR), qui a démarré en 2006, est de réaliser une plate-forme logicielle permettant de faciliter la gestion de contenus du web. Au sein de l'équipe GEMO, qui associe l'INRIA-Saclay-Île-de-France et l'université Paris-Sud, nous nous penchons plus particulièrement sur le développement de la gestion des connaissances de cette plate-forme dans un environnement pair-à-pair. Les applications ciblées concernent essentiellement la surveillance du web : veille économique dans le domaine aéronautique avec EADS, veille stratégique avec Thales, veille sismique avec le CEA ou surveillance du risque microbiologique et chimique dans le domaine alimentaire avec la Société de recherches et de développement alimentaire Bongrain (Soredab) et l'Institut national de la recherche agronomique (INRA). Ces organismes participent au développement de logiciels avec quelques startups (Exalead et New Phenix) et une dizaine d'équipes de recherche. Les verrous technologiques viennent notamment de la masse d'information disponible sur le web, du nombre de serveurs, de l'hétérogénéité syntaxique et sémantique des données et du multilinguisme.

Dans un tel cadre, le P2P permet d'obtenir d'importantes ressources à faible coût. Il offre aussi la possibilité à plusieurs sociétés de partager leurs données tout en gardant le contrôle de leurs propres informations. Les défis technologiques à relever exigent de combiner des techniques de divers domaines de l'informatique : bases de données (optimisation de requêtes), apprentissage (fouille de texte), bases de connaissance, linguistiques...

Il est prévu que la plate-forme WebContent, qui devra être achevée à la mi-2009, soit ouverte à d'autres établissements qu'aux organismes partenaires : son cœur sera distribué sous une licence de logiciel libre de telle sorte que les utilisateurs puissent y intégrer leurs propres outils. S. A.

Pour en savoir plus : <http://www.webcontent.fr/>

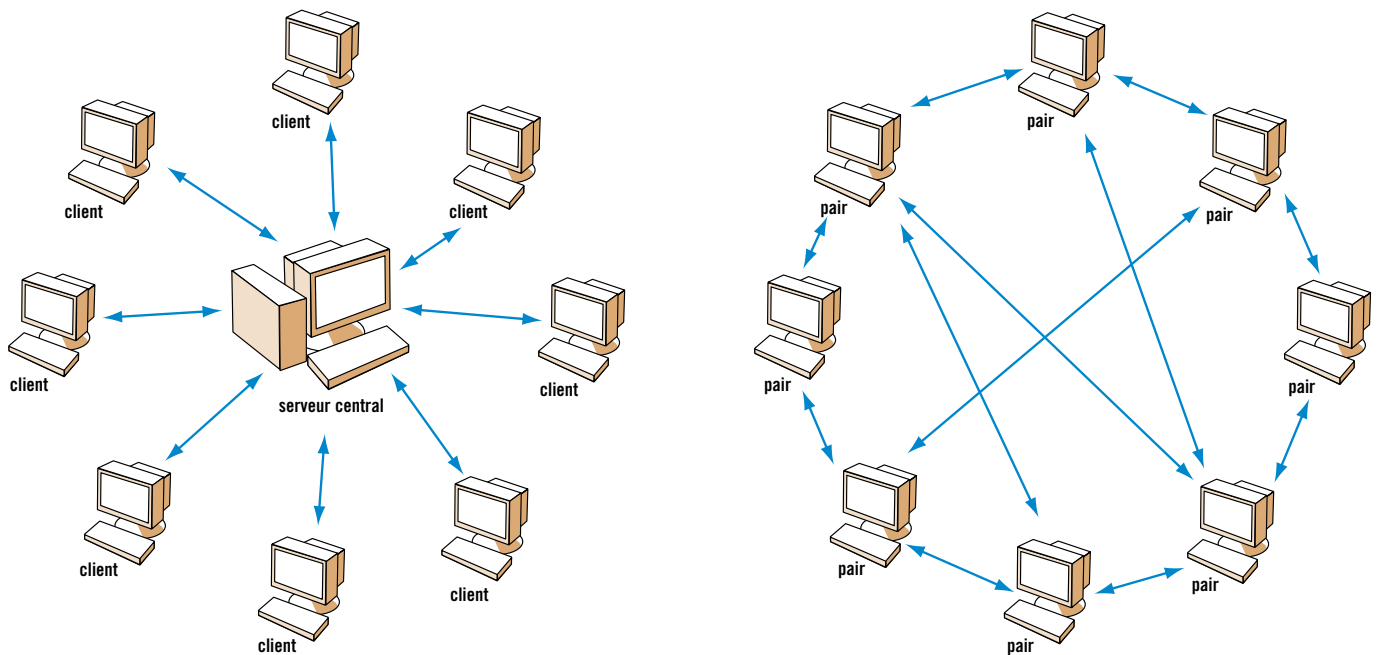


Fig. 2 : A gauche est schématisée une architecture client-serveur, à droite un réseau pair-à-pair.

de transformer des données locales en une ressource accessible partout sur le web.

Les avantages du P2P sont considérables. En disposant de nombreuses machines (de leurs processeurs, de leurs mémoires, de leurs disques), on peut offrir de meilleures performances et une meilleure disponibilité. Par exemple, dans l'approche centralisée, plus une information est populaire, moins elle est accessible, car le serveur est vite saturé, ce qui est pour le moins paradoxal. Avec un système P2P comme *BitTorrent*^{*}, c'est l'inverse. Un client ayant téléchargé l'information devient serveur à son tour (fig. 2). Donc un gros fichier (une vidéo) très populaire est vite disponible chez de nombreux pairs et, pour la télécharger, il suffit de récupérer en parallèle des fragments de plusieurs machines à la fois.

Les systèmes P2P sont surtout intéressants d'un point de vue socio-économique. Au lieu d'investir dans un gros entrepôt de données centralisé (avec des coûts importants en serveurs et en salaires d'ingénieurs systèmes), il est possible de mettre en place quasi gratuitement un système P2P. De plus, avec un tel système, on garde le contrôle sur sa propre information. Cet aspect peut être illustré par les réseaux sociaux. Dans des systèmes centralisés comme Facebook, on nourrit le système d'informations personnelles et on ignore ce qu'il en sera fait. On aimerait être en mesure de rester maître de

telles informations, et notamment savoir qui y a accès. Avec un système P2P, c'est possible !

Pour réaliser ce genre d'approche, un certain nombre de problèmes devra être résolu, comme celui de savoir localiser l'information recherchée ou d'arriver à garder cohérentes différentes versions d'une même donnée. Il faut aussi pouvoir garantir que seuls les utilisateurs habilités auront accès à une information et seront autorisés à la modifier, et qu'il leur sera donné les moyens de se protéger des écueils classiques du web (internauts ou systèmes indéliçats). Tous ces problèmes ont des solutions et on devrait voir arriver des systèmes P2P de plus en plus fonctionnels et complets (voir l'encadré sur WebContent). Le frein principal à leur développement reste que, pour les industriels, le « business model » du P2P n'est pas clair. Que peut-on vendre quand le but du jeu est d'utiliser des ressources déjà existantes et de préférence avec des logiciels libres ? De la publicité, peut-être. On note par exemple l'arrivée timide de systèmes de publicité comme Altnet. Pourtant, on préfère rêver à des systèmes libres d'où publicité et intentions commerciales seraient absentes.

Serge Abiteboul, directeur de recherche à l'INRIA-Saclay-Île-de-France, dirige une équipe sur la gestion de données et de connaissances distribuées. Il a été professeur à Stanford et, à temps partiel, à l'École polytechnique. Lauréat de l'ACM SIGMOD Innovation Award (1998) et du prix d'informatique de l'Académie des sciences (2007), il vient d'obtenir une *Advanced Grant* du Conseil européen de la recherche pour son projet Webdam sur les fondements de la gestion de données du web.

⁽¹⁾ S. Androutsellis-Theotokis et D. Spinellis, « A survey of peer-to-peer content distribution technologies », in *ACM Computing Surveys*, 36 (4), 2004

⁽²⁾ T. Ozsu et P. Valduriez, *Principles of distributed database systems*, Prentice Hall, 1999

⁽³⁾ G. Gardarin, *Base de données*, Eyrolles, 2003

⁽⁴⁾ S. Abiteboul, R. Hull et V. Vianu, *Foundations of databases*, Addison-Wesley, 1995



INRIA-C. LEBEDINSKY

« Si tous les gars du monde voulaient se donner la main... » : nul n'ignore ce poème de Paul Fort ! Le but de Gossple, un Google collaboratif, est certes plus prosaïque, mais c'est bien sur l'établissement de liens sociaux qu'il s'appuie.

Anne-Marie Kermarrec, directrice de recherche INRIA (Rennes), a créé l'équipe-projet ASAP centrée sur les systèmes distribués dynamiques et de très grande taille, après quatre ans chez Microsoft Research, à Cambridge (Grande-Bretagne).

Votre projet de moteur de recherche fondé sur des technologies pair-à-pair vient d'être primé par le Conseil européen de la recherche*. De quoi s'agit-il ?

Anne-Marie Kermarrec : Je suis partie du constat que Google est extrêmement bon pour indexer les documents qui existent sur le web : il serait absurde de le concurrencer

D'UNE FRONTIÈRE À L'AUTRE

Cela fait un peu penser à la manière dont certains établissements commerciaux recueillent des informations pour cibler leur publicité...

A.-M. K. : En effet, j'ai eu l'occasion de discuter avec un spécialiste d'économétrie et de statistique : il étudie lui aussi les affinités, par exemple sur la base des achats inscrits sur les cartes de fidélité de supermarché. Mais là où nous divergeons, c'est sur le caractère décentralisé : aucun serveur, nulle part, n'est là pour maintenir les connexions et centraliser les informations. Notre objectif est de réaliser un système totalement distribué, sur des millions de machines. Bien entendu, nous n'éviterons pas les comportements « déviants » (mensonges, rumeurs...) : la veille sur ce problème pose des questions techniques à résoudre, comme d'ailleurs le problème des virus ou celui d'utilisateurs cherchant à biaiser le système (raisons commerciales, publicitaires...).

Comment, dans cet immense désordre de terminaux interconnectés, créer un ordre au point d'en extraire de l'information pertinente ?

A.-M. K. : C'est exactement le défi du pair-à-pair. Nous mettons notamment en œuvre des protocoles dits épidémiques*. Typiquement, si l'on veut communiquer une information, une solution consiste à la faire passer à la radio à

Entretien avec Anne-Marie Kermarrec Recherche collaborative

sur ce terrain. Il l'est moins pour des requêtes plus spécifiques. Par exemple, si je cherche une baby-sitter de profil particulier (anglophone...) qui, typiquement, s'adresserait à des étudiants ou assistants, j'aurai du mal *via* Google, alors qu'il me suffirait de réunir les étudiants étrangers pour la trouver presque immédiatement. L'idée de Gossple est de favoriser la création de réseaux d'internautes (réseaux sociaux) ayant des affinités communes, dans le but de trouver des réponses à des requêtes assez pointues. Ces réseaux existent dans la réalité (associations, institutions, amis...) : comment faire exister l'équivalent sur le web, et ce de manière entièrement décentralisée ? La solution ne peut venir de Google car les techniques mises en œuvre ne sont pas centrées sur l'utilisateur. Il faut notamment « capturer » les affinités (personnelles, professionnelles...) d'une manière ou d'une autre, et ce de manière transparente. Cela revient à fusionner l'aspect recherche de données et l'aspect connexions sociales.

une heure de grande écoute. Autre solution, celle consistant à répartir le travail : chacun transmet l'information à cinq personnes, lesquelles font de même, et ainsi de suite. C'est, très schématiquement, l'une des choses que peut faire un protocole épidémique. Il s'avère que l'information arrive très vite chez tout le monde, même si la moitié des personnes contactées (des machines connectées) ne fait pas son boulot. Au plan scientifique, la mise au point d'un tel système va entre autres s'appuyer sur des résultats de théorie des graphes, d'algorithmique distribuée...

Qu'attendez-vous de la subvention du Conseil européen de la recherche ?

A.-M. K. : D'abord, c'est extraordinaire pour un chercheur d'être doté des moyens lui permettant d'explorer une idée en toute tranquillité. Ensuite, je suis convaincue qu'avec les algorithmes épidémiques, tout est possible...

Propos recueillis par Dominique Chouchan

* Instauré par le **Conseil européen de la recherche** (European Research Council, ou ERC), le Starting Independent Research Grant permet de financer le projet de recherche d'un(e) jeune chercheur. Sur 10 000 candidatures, 300 ont été sélectionnées dont celle d'Anne-Marie Kermarrec, qui dispose ainsi d'un million d'euros pour cinq ans.

* Un **protocole épidémique** permet de faire émerger des propriétés globales dans un système distribué de grande taille à partir d'interactions locales (pair-à-pair) : périodiquement chaque machine échange des données avec une autre. Ces protocoles sont simples, robustes et puissants.