

Les mémoires du FUTUR

Comment stocker et traiter le déluge d'informations produites au quotidien ? Les mémoires numériques actuelles se rapprochant de leurs limites, les scientifiques développent de nouveaux dispositifs.

Luca Parniola

En 2012, l'industrie a fabriqué plus de puces électroniques de mémoire que l'agriculture n'a produit de grains de riz ! Une telle production est indispensable pour stocker et traiter les énormes quantités d'informations créées tous les jours. Dans le monde, $2,72 \times 10^{21}$ octets de données ont ainsi été produits en 2012. Les dispositifs actuels sont-ils capables de soutenir la croissance rapide du volume de données ? Cela implique des capacités de stockage importantes et des dispositifs qui permettent d'accéder rapidement aux données. Or les techniques utilisées aujourd'hui semblent proches de leurs limites. Aussi les laboratoires développent-ils une nouvelle génération de mémoires numériques, plus performantes et moins coûteuses.

La nécessité de conserver des informations a débuté très tôt dans l'histoire de l'humanité. Les premiers supports étaient analogiques – parmi les nombreux dispositifs, on trouve les tablettes d'argile, les parchemins ou les livres. Ils avaient pour principal inconvénient d'être encombrants. L'émergence du

L'ESSENTIEL

- Les appareils numériques utilisent des mémoires rapides d'accès pour interagir avec le microprocesseur et des mémoires à grande capacité de stockage, mais plus lentes. Chaque type a ses avantages et ses limites.
- Les mémoires universelles, rapides et capables de stocker de grands volumes de données, n'existent pas.
- Les scientifiques conçoivent de nouveaux dispositifs, fondés sur des effets physiques différents, qui seront rapides d'accès, fiables, économes en énergie...

numérique a bouleversé cet état de fait. Toute donnée – musique, texte, image – est convertie en une séquence de bits (des « 0 » et des « 1 ») qui permet, entre autres avantages, un gain de place important, un traitement automatisé, la consultation à distance, etc.

Le premier exemple de mémoire numérique est celui des cartes perforées utilisées sur les métiers à tisser Jacquard à partir du début du XIX^e siècle. Ces cartes présentaient une suite de trous (un trou correspondant, par exemple, à 0 et son absence à 1). Elles contrôlaient le mouvement des crochets du métier à tisser et ainsi automatisaient la fabrication des textiles. Ces cartes perforées se limitaient à indiquer une succession d'instructions.

En 1834, le mathématicien anglais Charles Babbage a eu l'idée d'utiliser les cartes perforées avec sa machine à calculer. Les cartes contenaient les instructions et les données pour les opérations à effectuer. Cette innovation a conduit à l'invention des ordinateurs et au développement des supports de mémoire numériques. Ceux-



© Vasya Kobelew/shutterstock

ci ont connu un essor important au cours des dernières décennies. Ils représentent aujourd'hui l'essentiel des supports de mémoire. Durant ces 20 dernières années, les performances de ces dispositifs ont aussi encouragé la production massive de données (*voir la figure 1*). La mise en mémoire étant devenue simple et peu chère, il n'y avait plus besoin de limiter le volume de données à stocker.

Les mémoires universelles : utopiques ?

Si l'accroissement du volume d'informations à stocker se poursuit au même rythme, les mémoires d'aujourd'hui seront-elles à la hauteur, en capacité de stockage comme en rapidité d'accès ? Il est probable que non. Les ingénieurs rêvent de mémoires dites universelles qui seraient notamment rapides d'accès, capables de stocker de grands volumes de données, fiables et économes en énergie. Malheureusement, un tel dispositif relève plus de l'idéal vers lequel tendre que d'une réalité concrète.

Les systèmes numériques d'aujourd'hui associent toujours plusieurs dispositifs : des mémoires dites de travail (ou « mémoires vives »), qui conservent de petits volumes de données mais qui ont des temps d'accès très courts, et des mémoires de stockage, qui peuvent enregistrer d'énormes quantités d'information, mais au prix d'un temps d'accès plus long. En associant ces deux types de mémoires, il est possible de construire un système qui peut exécuter des opérations en peu de temps et conserver un grand volume de données.

De nombreux concepts sont à l'étude et certains pourraient être opérationnels dans les années à venir. Pour comprendre l'intérêt de ces dispositifs, il faut commencer par s'intéresser aux mémoires utilisées aujourd'hui, avec leurs qualités et leurs limites. Je présenterai ensuite quelques solutions techniques qui pourraient succéder aux dispositifs actuels d'ici cinq à dix ans.

Comme je l'ai évoqué, nous ne savons pas fabriquer de mémoires universelles. La nature ne semble pas faire mieux dans ce domaine. Il existe, par exemple, différents modèles théoriques de la mémoire du cer-

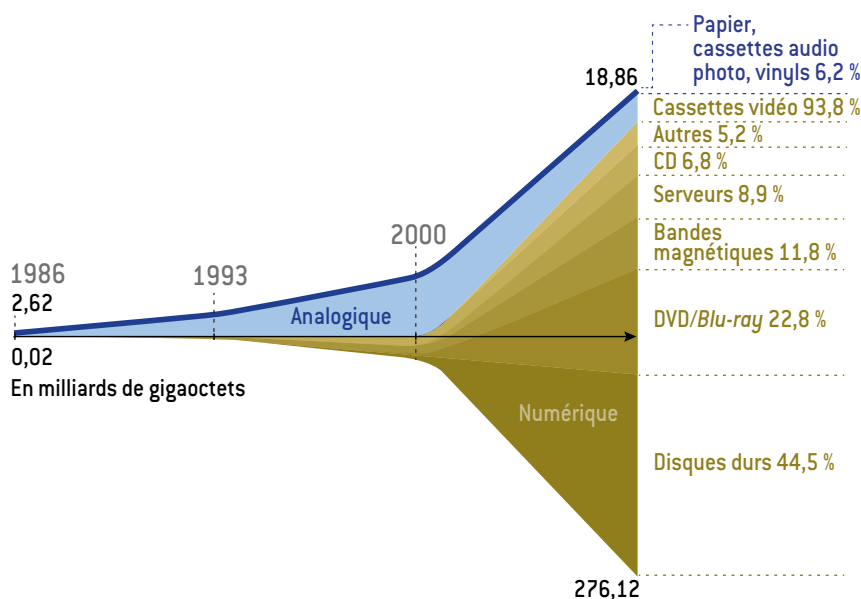
veau humain. Dans tous les cas, la mémoire est subdivisée en sous-systèmes spécialisés en différentes tâches. On observe ainsi de nombreuses analogies entre la mémoire biologique et les mémoires artificielles. Dans un modèle classique des années 1960, on distingue la mémoire sensorielle qui retient une grande quantité d'informations provenant des différents sens pendant un temps très bref et sans faire d'interprétation. Certaines informations sont ensuite transmises à la mémoire à court terme, dont la capacité est limitée et où les informations sont disponibles pendant plusieurs dizaines de secondes. Cette mémoire permet, entre autres, à un lecteur de ne pas oublier le début de la phrase pour pouvoir en comprendre le sens.

Enfin, une mémoire à long terme permet de conserver des informations longtemps et en quantités importantes. L'évolution a probablement conduit à cette structure qui permet par exemple de traiter rapidement une information perçue, comme quand nous sommes confrontés à un prédateur ou à un autre danger.

Depuis les années 1960, ce modèle de la mémoire humaine a subi des modifications et s'est enrichi de sous-catégories (*voir l'article de Franck Chaillan dans ce numéro*). Cependant, la conclusion est que le cerveau n'est pas une mémoire universelle et que la combinaison de mémoires spécialisées semble présenter des avantages.

L'industrie suit un schéma similaire, puisqu'elle associe différents types de mémoires. Si, à plusieurs reprises, de nouveaux dispositifs ont été présentés comme universels, des études approfondies les ont toujours mis en défaut. Quoi qu'il en soit, les fabricants se concentrent sur l'amélioration de leurs caractéristiques : la durée de vie, la fiabilité, etc. Une propriété importante est le temps d'accès, la durée nécessaire pour la lecture ou l'écriture d'une information sur une mémoire ; elle est comprise entre la nanoseconde et la milliseconde, suivant les dispositifs.

Le processeur d'un ordinateur, qui exécute des calculs, a besoin d'obtenir très rapidement les instructions et les données nécessaires à l'opération. Il doit donc être associé à des mémoires de travail à temps d'accès très court, pour ne pas ralentir les performances du processeur. De telles mémoires coûtent cher, et leur densité d'information stockée est faible : elles ne peuvent pas emmagasiner de grands



1. LES DONNÉES NUMÉRIQUES ont connu un essor spectaculaire. Le développement de supports numériques bon marché et simples d'utilisation marque un tournant au début des années 2000. Le numérique a supplanté les supports analogiques et le volume total de données enregistrées a augmenté rapidement. Cette tendance ne semble pas faiblir.

volumes de données. Pour archiver les informations, il faut donc ajouter au dispositif des mémoires de stockage.

Les deux types de mémoires sont complémentaires et cohabitent dans tous les dispositifs numériques (voir la figure 2) : ordinateurs, tablettes, téléphones portables, clés USB... On peut définir une hiérarchie pyramidale des mémoires. Au sommet se situent les mémoires de travail, aux temps d'accès très courts et aux capacités de stockage faibles. En descendant dans la pyramide, les volumes d'information qu'il est possible de stocker augmentent au détriment des temps d'accès. La base de la pyramide correspond aux mémoires de stockage.

Les mémoires de travail et de stockage se distinguent aussi par leur capacité à conserver l'information sans alimentation électrique. Les mémoires de travail sont dites volatiles : en l'absence d'une source d'énergie électrique, l'information qui y est enregistrée disparaît en quelques dizaines de millisecondes. Ainsi, quand un ordinateur effectue une opération et qu'une panne de courant survient, toutes les informations stockées dans la mémoire de travail sont perdues. Les mémoires de stockage sont, au contraire, non volatiles : l'information reste en mémoire. Si vous sauvegardez des fichiers sur une clé USB ou sur un disque dur et que vous éteignez l'ordinateur, l'information ne sera pas perdue et pourra être réutilisée ultérieurement.

Examinons le fonctionnement des différentes mémoires. Au plus près du processeur, on distingue deux types de mémoires de travail : les mémoires statiques (ou SRAM pour *Static Random Access Memory*) et les mémoires dynamiques (ou DRAM pour *Dynamic Random Access Memory*). Les premières occupent le sommet de la pyramide. Elles sont directement construites sur la puce de silicium du microprocesseur, car elles sont constituées de transistors, comme ce dernier. Leur temps d'accès est de quelques nanosecondes, mais le volume de stockage est très faible.

Les mémoires dynamiques, juste en dessous dans la pyramide des mémoires, ont une capacité de stockage plus importante ; elles sont en général placées sur des puces indépendantes, dont le temps d'accès est compris entre 10 et 100 nanosecondes. Ces mémoires de travail sont utilisées pour effectuer les opérations de calcul des microprocesseurs.

Des mémoires volatiles qui perdent les données

Dans une mémoire dynamique, le bit d'information est enregistré par un système composé d'un transistor et d'un condensateur. Ces composants sont gravés dans le substrat semi-conducteur en silicium de la puce de mémoire. Le transistor agit comme un interrupteur qui autorise, s'il est ouvert, la lecture du bit par un courant électrique. Physiquement, l'état du bit est défini par la charge électrique (l'accumulation d'électrons) du condensateur. Si le condensateur est chargé, le bit vaut 1, s'il est déchargé, le bit vaut 0.

Le problème des condensateurs est qu'ils retiennent mal les électrons, qui fuient dans le silicium environnant. Pour cette raison, la mémoire dynamique est volatile. Il est donc nécessaire de réécrire le bit d'information régulièrement, toutes les quelques dizaines de millisecondes, pour éviter que le condensateur ne se décharge trop et que le bit change d'état. Ainsi, sans alimentation électrique, l'information disparaît de la mémoire.

Dans les premières mémoires dynamiques des années 1960, les transistors et les condensateurs étaient accolés. La miniaturisation a conduit à réduire la taille des composants, mais aussi à changer leur placement relatif. Les condensateurs sont maintenant construits au-dessus des

L'AUTEUR



Luca PERNIOLA dirige le groupe des technologies de mémoires avancées au Laboratoire d'électronique et de technologie de l'information (LETI) au Commissariat à l'énergie atomique et aux énergies alternatives, à Grenoble.

transistors. De plus, la forme des condensateurs est aujourd'hui cylindrique et très allongée, ce qui réduit la consommation électrique. Avec un rapport de la longueur du condensateur sur la largeur de sa base qui peut atteindre 50, et une densité élevée de condensateurs, les fabricants se trouvent confrontés à plusieurs défis : la production de cavités aussi longues et étroites, et la gestion des problèmes de fuites. La limite de miniaturisation semble atteinte pour ces dispositifs, dont la densité de stockage (quatre gigaoctets) augmente difficilement.

Ainsi, les ingénieurs travaillent en priorité sur deux aspects relatifs aux mémoires dynamiques : l'augmentation de la densité d'information et l'économie d'énergie, cette dernière étant obtenue en supprimant la volatilité.

Les « mémoires flash » sont un type de mémoire non volatile. Mais elles ne peuvent pas remplacer les mémoires dynamiques. Situées plus bas dans la pyramide des mémoires, les mémoires flash ont un temps d'accès d'une dizaine de microsecondes – soit deux ordres de grandeur de plus que les mémoires dynamiques. Dans les mémoires flash, chaque bit est représenté par l'état ouvert ou fermé d'un transistor. Un réservoir au sein du transistor peut accumuler une certaine quantité d'électrons

qui créent un champ électrique. Celui-ci laisse passer ou bloque le courant qui traverse le transistor.

Les mémoires flash sont non volatiles, car les électrons restent bloqués dans le réservoir même en l'absence d'une source extérieure (qui, dans les mémoires dynamiques, rafraîchit les informations). Le blocage est dû au fait que le réservoir d'électrons est intercalé entre deux couches d'oxydes qui l'isolent du substrat environnant. Ainsi, une information archivée n'est pas altérée.

Depuis leur développement dans les années 1980, les mémoires flash ont été miniaturisées. Elles atteignent actuellement des capacités de stockage de 128 gigaoctets. Cependant, pour conserver les propriétés de non-volatilité, les couches d'oxydes ne peuvent pas être soumises à la même miniaturisation, qui affaiblirait le blocage des électrons. En outre, comme pour les mémoires dynamiques, la miniaturisation augmente le risque d'erreurs d'écriture et de lecture.

Aujourd'hui, pour augmenter la quantité d'information qu'il est possible d'enregistrer sur une plaque de silicium, les fabricants créent plusieurs couches de mémoire empilées les unes sur les autres. La capacité de mémoire est multipliée approximativement par le nombre de couches.

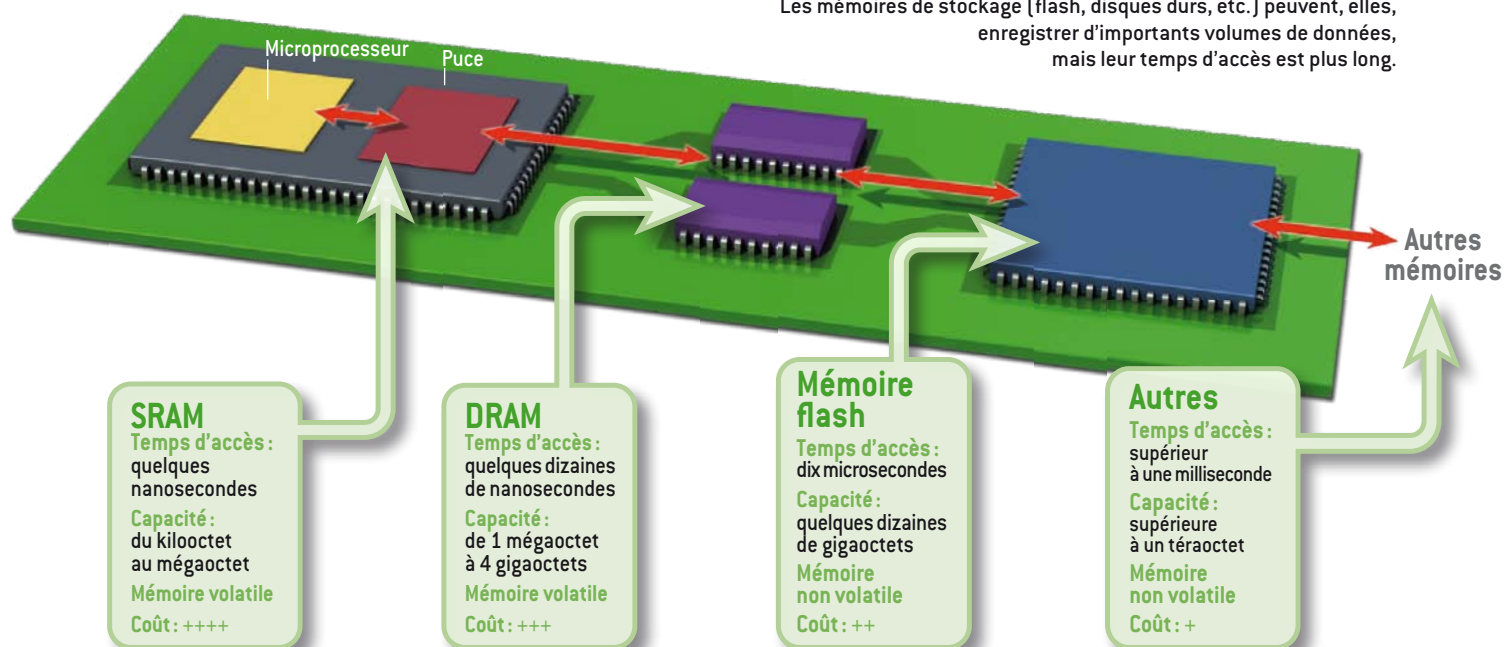
Même si les mémoires flash semblent aussi atteindre des limites de miniaturisation, elles pourraient remplacer une partie des « disques durs ». Ces supports numériques ont été développés dans les années 1950 par IBM. L'information est enregistrée sur des couches de matériau ferromagnétique déposés sur des disques rotatifs. Un élément, la tête de lecture, monté sur un bras mécanique, lit l'information sur le disque à partir de l'état magnétique local du bit. Il peut inscrire un bit grâce à des effets magnétiques réversibles.

Disques durs : volumineux, mais lents

Les disques durs occupent le bas de la pyramide des mémoires. Leur densité d'information est supérieure à celle des mémoires flash et leur prix par unité de mémoire stockée est plus faible. En revanche, le temps d'accès est de l'ordre de la milliseconde, car la tête de lecture doit d'abord être positionnée au bon endroit par un déplacement du bras mécanique et par une rotation du disque.

Dans les années à venir, les disques durs de certains appareils pourraient être remplacés par des dispositifs contenant des mémoires flash, les SSD, ou *Solid State*

2. UN SYSTÈME NUMÉRIQUE associe toujours des mémoires de travail et des mémoires de stockage. Les premières (mémoires statiques et mémoires dynamiques) sont placées au plus près du microprocesseur, chargé de faire les calculs. Elles échangent rapidement les informations avec le processeur, mais ne présentent que de petits volumes de stockage. Les mémoires de stockage (flash, disques durs, etc.) peuvent, elles, enregistrer d'importants volumes de données, mais leur temps d'accès est plus long.



Drive. Ces derniers ont plusieurs avantages. Les ordinateurs équipés de ce type de mémoires sont beaucoup plus rapides. Et l'absence d'éléments mécaniques les rend plus robustes et moins sujets aux pannes.

Du flash pour les nomades

Les habitudes nomades des utilisateurs de tablettes, d'ordinateurs ou de téléphones portables favorisent l'utilisation de systèmes résistants aux chocs. Toutefois, les disques durs ne disparaîtront pas dans un futur proche, leur coût restant dix fois inférieur à celui des mémoires flash.

Les mémoires dynamiques, les mémoires flash et les disques durs semblent atteindre leurs limites. Il est donc crucial de développer de nouveaux dispositifs

pour les remplacer. Mais les mémoires à venir ne pourront peut-être pas prétendre au statut de dispositifs universels. Les scientifiques continuent de développer des mémoires qui auront les caractéristiques soit de mémoire de travail, soit de mémoire de stockage, en améliorant leurs capacités de stockage, la vitesse de leur lecture ou de leur écriture, ou leur consommation d'énergie.

Ce dernier point est particulièrement étudié. Les mémoires non volatiles ont ici un réel avantage et devraient prendre une part de marché plus importante. Il est cependant difficile d'envisager un système uniquement composé de mémoires non volatiles. La mémoire la plus proche du microprocesseur est sollicitée en permanence et requiert des accès ultrarapides. La mémoire statique reste, dans ce cas, la

plus efficace et la plus économe en énergie. Il n'y a donc pas besoin de remplacer les mémoires statiques par des mémoires non volatiles (pour être exact, on distingue trois types de mémoires statiques, dont l'un fait l'objet de recherches pour le substituer par un système plus performant).

De nouveaux dispositifs sont pressentis pour remplacer les mémoires dynamiques et pour remplir l'espace de performance entre les mémoires dynamiques et les mémoires flash. Dans ces concepts, que nous allons présenter, le condensateur des mémoires dynamiques est remplacé par une résistance électrique qui peut varier, le bit correspondant valant 0 ou 1 selon que la résistance est faible ou élevée. Il n'y a pas de problème de fuite de charges électriques dans ces dispositifs ; dès lors, ces mémoires sont non volatiles.

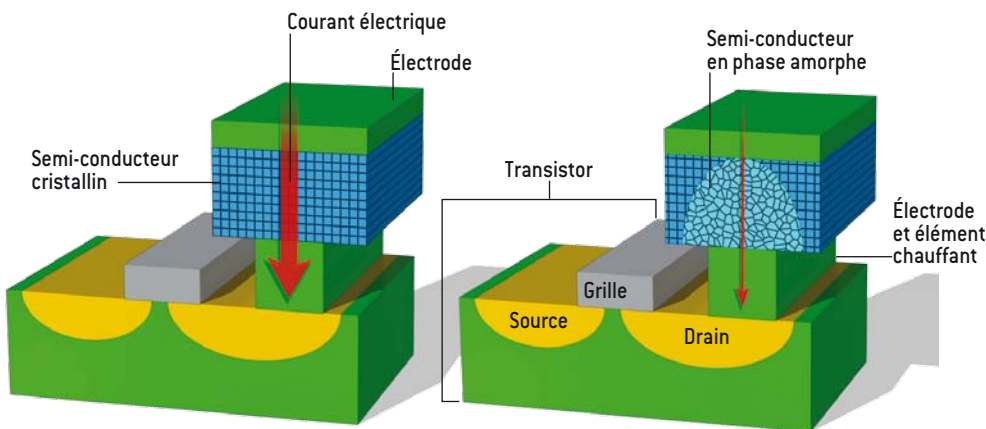
LES SUCCESSIONS DES

Pour gérer l'énorme flux d'informations à traiter et à archiver, de nouveaux dispositifs sont en cours de développement. Ils pourraient être bientôt sur le marché et remplacer les systèmes actuels. Les mémoires dynamiques sont des mémoires volatiles, qui nécessitent d'être réécrites en permanence. De ce fait, elles sont peu économes en énergie. Les mémoires à changement de phase et les mémoires magnétiques pourraient leur succéder en remplaçant les condensateurs, à l'origine de la volatilité, par des résistances.

Mémoires à changement de phase

Les mémoires à changement de phase (PCRAM ou *Phase Change Random Access Memory*) sont constituées d'un transistor et d'une résistance. Celle-ci se compose de deux électrodes séparées par un semi-conducteur, tel l'alliage germanium-antimoine-tellure, aux propriétés particulières. La structure de ce dernier peut être cristalline ou amorphe. Dans le premier cas, la circulation des électrons du courant électrique est aisée : la résistance est faible et le bit vaut 1 (*à gauche*). Avec une phase amorphe, la circulation des électrons est plus difficile, la résistance est supérieure (*à droite*). Cela correspond à un bit égal à 0.

Pour changer la valeur du bit, il faut changer la structure du semi-conducteur par chauffage. Un alliage résistif sert de source de chaleur. Il convertit un courant électrique en chaleur par effet Joule. Avec un chauffage bref, on obtient une structure amorphe. Si le chauffage est soutenu, une structure cristalline se met en place. Pour effectuer ces transitions réversibles,



des températures de 500 °C à 800 °C sont nécessaires, ce qui implique des courants assez importants dans les dispositifs. La différence de résistance des deux états est bien marquée (rapport supérieur à 100), ce qui permet d'avoir un temps de lecture assez court, car il n'y a pas d'ambiguïté entre les deux états.

Une telle amplitude de résistance permettrait d'utiliser un tel dispositif pour plusieurs bits. Dans ce cas, la plage de résistance n'est pas divisée en deux intervalles correspondant aux bits 0 et 1, mais en davantage d'intervalles, ce qui permet de coder deux bits (00, 01, 10 ou 11). Ce fonc-

tionnement est délicat à appliquer à cause de phénomènes de relaxation. Le matériau subit des chocs thermiques durant les étapes d'écriture et d'effacement. Entre ces étapes, une relaxation structurelle conduit à une variation stable de la résistance dans le temps, qui empêche une bonne utilisation des mémoires à changement de phase en mode multiniveaux.

En outre, ces mémoires ont une grande stabilité quand elles sont utilisées à haute température. Elles peuvent être utilisées dans certaines applications – automobiles, cartes à puces... – où les contraintes de température sont fortes.

Il existe toute une variété de mémoires à résistance variable, selon les principes physiques qui créent la résistance. Présentons ici deux successeurs possibles des mémoires dynamiques : les mémoires dites magnétiques et les mémoires dites à changement de phase. Ces mémoires consomment moins d'énergie et présentent, en principe, une densité d'information équivalant à celle des mémoires dynamiques.

Les mémoires magnétiques (MRAM pour *Magnetic Random Access Memory*) existent depuis les années 1950. Leur utilisation était cependant limitée à des niches, tel le domaine aérospatial. Ces mémoires résistent mieux aux effets des rayons cosmiques, qui peuvent induire des erreurs d'écriture des bits dans les autres dispositifs de stockage. Depuis dix ans, les mémoires magnétiques ont

connu de nombreuses innovations qui ont relancé leur intérêt.

MRAM et PCRAM, des pistes sérieuses

Il existe aujourd'hui des mémoires magnétiques ayant une capacité de stockage de 64 mégaoctets. La résistance qui code le bit est contrôlée par l'aimantation des éléments dont est formé le composant (voir l'encadré ci-dessous). Le temps d'écriture est de l'ordre de la nanoseconde, mais le temps nécessaire pour lire une information est plus long. En effet, la différence de résistance entre les états « 0 » et « 1 » n'étant pas très marquée (le rapport des résistances est inférieur à dix), il faut plus de temps pour lire correctement le bit d'information.

Le deuxième dispositif pour remplacer la mémoire dynamique est la mémoire à changement de phase (PCRAM ou *Phase Change Random Access Memory*). Il a été développé à la fin des années 1960 par l'inventeur américain Stanford Ovshinsky (un autodidacte qui a déposé plus de 400 brevets en 50 ans et dont un grand nombre sont exploités industriellement). Dans cette mémoire, la résistance est déterminée par la structure atomique d'un semi-conducteur. Ce dernier peut être cristallin ou amorphe, et devient ainsi conducteur ou isolant (voir l'encadré ci-dessous). Un élément chauffant permet de changer de façon réversible la structure du semi-conducteur. Ce type de mémoire demande donc plus d'énergie pour écrire ou effacer un bit que dans le cas de la mémoire magnétique.

MÉMOIRES DYNAMIQUES

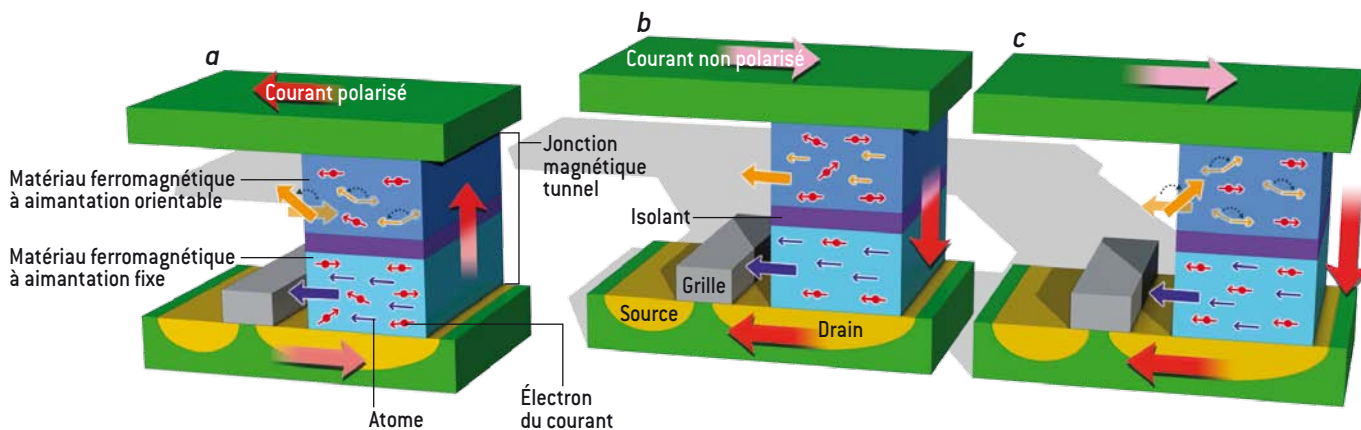
Mémoires magnétiques

Pour remplacer les mémoires dynamiques et dépasser leurs limites, une des pistes les plus intéressantes est celle des mémoires magnétiques (MRAM pour *Magnetic Random Access Memory*). Dans l'exemple de la mémoire magnétique à transfert de spin, le bit de mémoire est constitué d'un transistor et d'une résistance. Cette dernière est constituée d'une superposition de deux films minces de matériau ferromagnétique séparés par une barrière isolante qui ne peut être franchie par des électrons que par effet tunnel – on parle de jonction magnétique tunnel. L'information est portée par la résistance de la jonction : on a « 0 » si la résistance est importante, « 1 » si la résistance est faible. L'aimantation des deux couches ferromagnétiques contrôle cette ré-

sistance. L'aimantation de la couche inférieure est fixée. Dans l'autre, il est possible de l'orienter grâce à un courant dont le spin des électrons – leur moment cinétique intrinsèque – est polarisé : tous les spins sont dirigés dans le même sens. Ce spin permet d'orienter l'aimantation du matériau. Le schéma (a) indique le processus d'écriture et le passage d'un bit 0 à un bit 1. Un flux d'électrons (*en rouge*) non polarisé arrive par le bas. Lors du passage de la barrière de potentiel, le transfert des électrons de spin dirigé vers la gauche est facilité. Ainsi, dans le matériau ferromagnétique, le flux d'électrons est polarisé. Le spin des électrons interagit avec les spins des atomes du matériau (*en jaune*) et les réoriente de la droite vers la gauche. L'aimantation des deux

matériaux ferromagnétiques (*violet et jaune*) est alors dans le même sens (parallèle). La résistance est alors faible. Un courant plus faible peut être injecté dans le dispositif pour mesurer cette résistance et signaler que l'état correspond à un bit 1.

Pour faire passer l'état d'un bit de 1 à 0, on inverse le sens du courant de polarisation. En arrivant par le haut (b), il n'est pas polarisé, mais seuls les électrons dirigés vers la gauche traversent la barrière. Il ne reste que des électrons de spin dirigé vers la droite dans le milieu ferromagnétique du haut (c), ce qui entraîne un basculement de l'aimantation de la gauche vers la droite. Les aimantations sont antiparallèles. La résistance est grande, ce qui correspond à un bit 0.



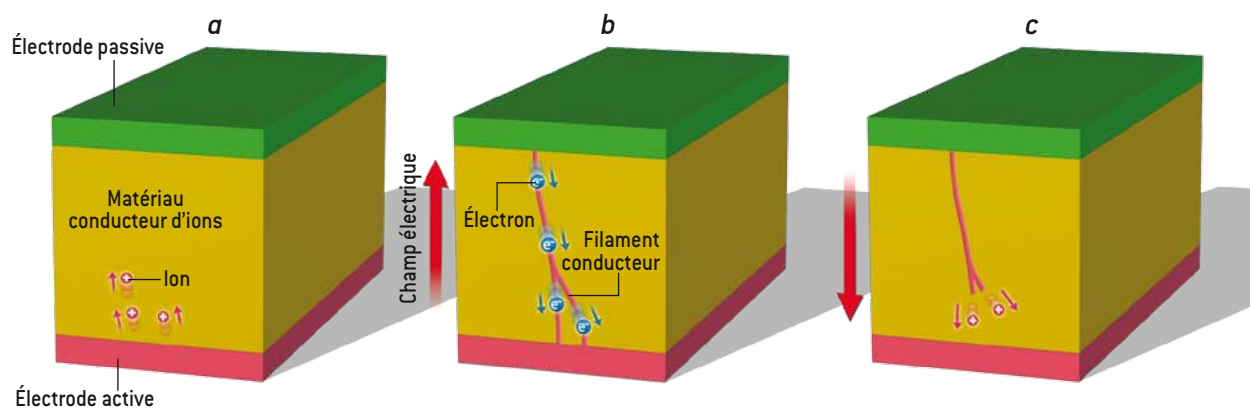
Les mémoires flash pourraient être remplacées par des dispositifs fondés sur les mémoires à oxydes (OxRAM) et les mémoires à pont conducteur (CBRAM). Dans les deux dispositifs, le bit d'information est porté par un transistor couplé à une résistance. Cette dernière est composée de deux électrodes séparées par un oxyde isolant ou un solide perméable à la conduction des ions. L'information est codée par la résistance du dispositif, qui dépend de la présence ou non d'un filament conducteur entre les électrodes. Ce filament est créé ou détruit par un stimulus électrique. Le phénomène électrique de claquage de l'oxyde (correspondant à la formation du filament conduc-

teur) a été étudié depuis les années 1960, mais l'industrie n'a commencé à s'y intéresser qu'à partir de 2004. En général, pour les mémoires à oxydes, le substrat est « vierge » lors de la première utilisation : il faut appliquer une tension importante pour amorcer la formation du filament. La tension d'utilisation par la suite sera beaucoup plus faible, car seul un petit élément du filament est construit ou détruit.

La physique des mémoires à pont conducteur est mieux établie que celle des mémoires à oxydes. La nature du pont dépend de la composition de l'électrode active (cuivre, argent...). Le filament se forme à partir de la dissolution de l'électrode active, qui émet des cations se propageant vers

l'électrode inerte grâce au champ électrique appliqué (a). Les cations sont réduits et forment un filament d'atomes neutres qui relie les électrodes (b). La résistance du dispositif chute (et fait passer le bit de l'état 0 à 1). Le processus est réversible : il faut inverser la polarité du champ électrique pour que les cations retournent à leur électrode initiale (c).

De nombreuses études cherchent à optimiser les performances de ce dispositif en associant divers matériaux pour les électrodes, aussi bien dans les mémoires à oxydes que dans celles à pont conducteur. De plus, la compréhension et la description microscopique des effets impliqués pourraient nécessiter plusieurs années.



Les mémoires à changement de phase atteignent aujourd'hui une capacité de un gigaoctet et les temps d'écriture sont de l'ordre de plusieurs dizaines de nanosecondes. De ce point de vue, ces dispositifs sont moins performants que les mémoires magnétiques. Mais la différence de résistance des deux états étant plus marquée (d'un facteur supérieur à 100), le temps de lecture est plus court. Par ailleurs, ces mémoires sont robustes et fonctionnent à haute température.

Le paysage des mémoires de stockage pourrait aussi changer d'ici cinq à dix ans. Comme nous l'avons vu, les disques durs sont les mémoires les moins chères, mais ils sont trop fragiles pour des utilisations nomades. Bien qu'ils soient déjà en partie remplacés par des dispositifs contenant des mémoires flash, de nouvelles mémoires de stockage se profilent. Si leurs caractéristiques deviennent intéressantes, elles pourront remplacer les mémoires

flash dans quelques années. Deux pistes sont prometteuses parmi les nombreuses solutions étudiées par les industriels : les mémoires à oxydes (OxRAM ou *Oxide-based Random Access Memory*) et les mémoires à pont conducteur (CBRAM ou *Conductive Bridge Random Access Memory*, voir l'encadré ci-dessus). Ces dispositifs sont récents, mais les recherches sont très actives et le nombre de prototypes augmente rapidement. Il n'est pas impossible que ces concepts de mémoires remplacent aussi les mémoires dynamiques. En 2013, une mémoire à oxydes de 32 gigaoctets a été annoncée par les Sociétés *Sandisk* et *Toshiba*.

Comme pour les mémoires magnétiques et à changement de phase, ces dispositifs sont constitués d'un transistor et d'une résistance variable. La valeur de la résistance détermine si le bit vaut 0 ou 1. Elle est contrôlée par un filament conducteur qui s'établit entre deux électrodes. Celui-ci peut être créé ou détruit par un champ et un courant électriques. Par exemple, dans une mémoire à pont conducteur, le filament métallique se

SI LES DISQUES DURS COMMENCENT déjà à disparaître dans les appareils nomades, les mémoires dynamiques pourraient être remplacées d'ici trois à cinq ans.

forme à partir d'ions positifs qui sont extraits d'une des électrodes. Quand la chaîne métallique relie les deux électrodes, le courant d'électrons peut facilement circuler et la résistance est faible. Le bit d'information vaut alors 1. En inversant le champ électrique, les atomes redeviennent des ions et migrent vers l'électrode : la chaîne se brise, la résistance augmente et le bit d'information passe de 1 à 0. Le processus est réversible.

Les progrès sont si rapides dans ce domaine qu'il est difficile de donner les caractéristiques de ces mémoires. Le temps d'écriture varie de la nanoseconde à la milliseconde suivant les dispositifs. Par ailleurs, les techniques de fabrication, surtout pour les mémoires à oxydes, sont relativement courantes, et leur coût de production est très compétitif. En revanche, les phénomènes microscopiques mis en jeu restent à mieux comprendre.

Tous les dispositifs présentés n'en sont pas au même point de maturité et leur présence sur le marché peut prendre

encore quelques années. Si les disques durs commencent déjà à disparaître dans les appareils nomades, les mémoires dynamiques pourraient n'être remplacées que d'ici trois à cinq ans.

Un horizon encore flou

La demande de mémoires toujours plus efficaces va continuer de croître. Cela encouragera les universités, les centres de recherche et les constructeurs à explorer les pistes évoquées, ainsi que de nombreuses autres. La concurrence étant forte et favorable à l'inventivité, l'évolution peut être rapide, et il est difficile de voir clairement quels composants s'imposeront dans cinq, dix ou vingt ans. La tendance qui se dégage est l'abandon de l'idée d'une technologie prédominante et universelle qui remplacera toutes les mémoires existantes. La solution optimale est de faire cohabiter plusieurs techniques et d'améliorer les performances de chacune. Jusqu'où ? On l'ignore. ■

■ BIBLIOGRAPHIE

H.-S. Wong *et al.*, **Metal oxide RRAM**, *Proceedings of the IEEE*, vol. 100, n° 6, pp. 1951-1979, 2012.

T. Kawahara *et al.*, **Spin-transfer torque RAM technology : Review and prospects**, *Microelectronics Reliability*, vol. 52, pp. 613-627, 2012.

M. Hilbert et P. Lopez, **The world's technological capacity to store, communicate and compute information**, *Science*, vol. 332, pp. 60-65, 2011.

G. W. Burr *et al.*, **Phase-change memory technology**, *J. Vac. Sci. Technol. B*, vol. 28, n° 2, pp. 223-262, 2010.



france culture

LA MARCHE DES SCIENCES

Découvertes, inventions, aventures savantes au fil de l'Histoire

Aurélie Luneau
14h-15h / chaque jeudi

franceculture.fr

SCIENCE

Stocker de l'information dans l'ADN

Christophe Dessimoz

Grâce à l'avancée des techniques de la génomique, il devient possible d'utiliser l'ADN comme support numérique pour archiver des données sur le long terme.

Deux piscines olympiques: tel est le volume qu'il faudrait pour contenir l'ensemble des données produites dans le monde en 2012 (environ 3×10^{21} octets), si on les stockait dans les toutes récentes clefs USB de un téraoctet (10^{12} octets). Au fil des ans, ce volume croît de façon exponentielle et, avec lui, le besoin de supports numériques fiables pour archiver l'ensemble des informations sur le long terme. De plus, la durée de vie des supports actuels étant limitée (pas plus de quelques dizaines d'années), la pérennisation des données nécessite une maintenance constante et lourde, fondée sur des copies multiples et régulières sur divers supports (voir l'article de Francis Daumas et Marion Massol dans ce numéro).

Pourtant, il existe depuis plus de trois milliards d'années un support numérique universel, compact et stable: l'ADN, la molécule porteuse de l'information génétique des organismes vivants. Codé sous la forme d'une succession de quatre molécules – les nucléotides A, T, G, C –, le programme de la vie est numérique: les nucléotides codent de l'information, au même titre que les bits (0 ou 1) de l'informatique. L'ADN est de plus très stable – on a retrouvé, dans des carottes de glace, des gènes datant de 450 000 à 800 000 ans – et dense: si l'on codait sur de l'ADN toute l'information produite en 2012, elle tiendrait dans le coffre d'une voiture...

En outre, le fait que l'ADN soit universel limite les risques d'obsolescence. Tant

qu'il y aura de la vie sur Terre et un intérêt pour la comprendre, la technique pour lire l'ADN existera. Aussi les progrès récents de la génomique nous ont-ils conduits, à l'Institut européen de bio-informatique, à étudier la possibilité d'utiliser l'ADN comme support d'archives numériques.

Tous les sonnets de Shakespeare sur ADN

Peu après la caractérisation de l'ADN, publiée en 1953 par James Watson et Francis Crick, l'idée d'utiliser cette molécule comme support d'information a été envisagée. En 1964, un physicien et officier radio russe, Mikhail Neïman, voit dans les systèmes et processus d'information biophysiques un moyen de miniaturiser les dispositifs de stockage et de traitement de l'information. L'idée a été mise en œuvre à la fin des années 1990, avec le développement du génie génétique. Plusieurs équipes ont su encoder de l'information sur l'ADN, mais les informations étaient courtes (un message secret annonçant le débarquement du 6 juin 1945, une comptine anglaise avec paroles, musique et petite image, la Déclaration des droits de l'homme...). Surtout, les manipulations génétiques étaient laborieuses et *ad hoc*.

Ces dernières années, cependant, les progrès des techniques de séquençage ont entraîné un tournant important. En 2010, le biologiste américain Craig Venter, de

L'AUTEUR



Christophe DESSIMOZ est professeur de bio-informatique au Collège universitaire de Londres. Il est aussi affilié à l'Institut européen de bio-informatique (EMBL-EBI). Article écrit en collaboration avec Marie-Neige Cordonnier.

BIBLIOGRAPHIE

N. Goldman *et al.*, Towards practical, high-capacity, low-maintenance information storage in synthesized DNA, *Nature*, vol. 494, pp. 77-80, 2013.

G. M. Church *et al.*, Next-generation digital information storage in DNA, *Science*, vol. 337, pp. 1628-1629, 2012.

Entretien avec Ch. Dessimoz à écouter sur: bit.ly/pls_433_dessimoz