

la révolution de la génomique : les nouvelles méthodes de séquençage et leurs applications

publié le 08.01.21 | par [erwin van dijk](#), [claire thermes](#)

des technologies qui permettent de déterminer la succession des nucléotides, c'est-à-dire la séquence d'une molécule d'acide nucléique (adn ou arn) existent depuis les années 70. la méthode sanger, la plus utilisée, a permis le séquençage de divers génomes dont celui de l'être humain, et a ainsi révolutionné la génomique et la biologie de manière générale. toutefois, le séquençage du génome humain avec la technologie sanger a été un immense effort qui a pris plus de dix ans et a coûté environ trois milliards de dollars [1]. il est donc évident que cette méthode n'est pas adaptée au séquençage de grands génomes, et c'est pour cette raison que de nouvelles technologies ont été développées. dans cet article nous discutons de ces méthodes dites « séquençage à haut débit » ou « de nouvelle génération » et qui ont révolutionné la génomique en permettant de séquencer un génome humain en quelques jours pour moins de 1000 dollars. nous présenterons également les technologies de troisième génération, encore plus récentes, qui, depuis quelques années, constituent une nouvelle révolution.

1. la méthode sanger

la technologie sanger a été décrite dans les articles [le séquençage d'un adn](#) et [le séquençage des génomes](#). ici nous résumerons de façon succincte le principe de cette méthode. le séquençage sanger d'un fragment d'adn nécessite d'abord de le cloner dans un plasmide, qui est ensuite introduit dans une cellule hôte, en général une bactérie ou une levure (figure 1). en se multipliant, cette cellule hôte produit un grand nombre de copies de chaque fragment d'adn d'origine. après purification, cet adn peut être séquencé en utilisant une polymérase qui synthétise un brin complémentaire à partir d'un brin matrice du fragment d'adn d'intérêt ; quand la polymérase incorpore un des quatre « didésoxynucléotides » (ddatp, ddctp, ddgtp, ou ddttp présents séparément dans quatre réactions individuelles), la synthèse s'arrête. cela génère un mélange de molécules qui se terminent à chaque position où se trouve un a, un c, un g, ou un t (selon le type de didésoxynucléotide présent). les fragments dans ce mélange sont séparés selon leur taille par électrophorèse sur gel. la connaissance du didésoxynucléotide qui a été incorporé dans chaque réaction permet ainsi de déduire la séquence du fragment d'adn d'intérêt.

2. la naissance des méthodes de séquençage de nouvelle génération

après la complétion du projet de séquençage du génome humain (*human genome project*, hgp), l'ambition était de séquencer un grand nombre de génomes afin d'étudier la variation génétique et de réaliser des études d'association pangénomique (*genome-wide association studies - gwas*), dans lesquelles on essaye d'identifier des liens entre des maladies génétiques et des profils génétiques spécifiques. dans ce but, des méthodes de séquençage de nouvelle génération (*next-generation sequencing*, ngs) ont été développées. ces méthodes partagent trois améliorations majeures par rapport au séquençage sanger :

1. au lieu d'un clonage moléculaire des fragments d'adn suivi par l'introduction dans des cellules hôtes et l'isolement de chaque clone individuellement, une banque qui contient l'ensemble des fragments est faite directement dans un tube (*in vitro*). la figure 1 montre de façon schématique la procédure ; des fragments d'adn, générés par coupures aléatoires (enzymatiques ou mécaniques) de l'adn génomique, sont reliés à des petites molécules d'adn de séquences connues appelés « adaptateurs ». les coupures aléatoires génèrent des fragments d'une grande diversité de tailles (entre environ 50 nt et 2000 nt). une sélection de taille est généralement effectuée pour deux raisons : (1) éliminer les fragments plus courts que la longueur de séquençage, et (2) éliminer les fragments trop longs (d'une taille supérieure à environ 1000 nt). cette dernière étape est importante pour les techniques de ngs qui nécessitent ensuite une amplification par pcr (*polymerase chain reaction*), qui est moins efficace sur de longs fragments. notez bien que tandis que dans la figure 1, une seule molécule d'adn génomique est montrée, en réalité une banque est faite à partir d'un grand nombre de copies des molécules d'adn génomique à séquencer^[1].
2. alors que les machines développées pour le projet de séquençage du génome humain n'étaient capables d'effectuer que quelques centaines de réactions de séquençage sanger en parallèle, les séquenceurs ngs peuvent faire des millions voire des milliards de réactions de séquençage en parallèle.
3. les technologies ngs ne nécessitent pas de séparation des fragments par électrophorèse ; la détection des nucléotides incorporés par la polymérase est faite directement après chaque cycle d'incorporation.

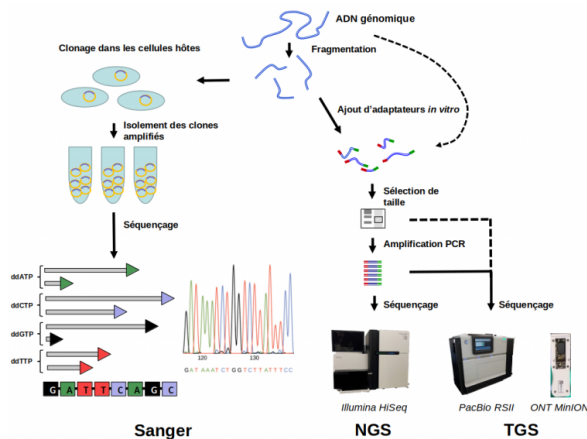


figure 1 - comparaison des méthodes de séquençage sanger, de nouvelle génération (ngs, next-generation sequencing) et de troisième génération (tgs, third-generation sequencing)
 les préparations des banques ngs et tgs suivent globalement les mêmes procédures, sauf que pour le tgs les étapes de fragmentation de l'adn génomique et d'amplification par pcr sont optionnelles (indiqué par flèches interrompues). les adaptateurs sont indiqués en rouge et vert.

crédits images : séquençage sanger : [openstax](#), cc by ; illumina hiseq 2500 : konrad förstner, cc0, [wikimedia](#) ; pacbio rsii : konrad förstner, cc by, [flickr](#) ; ont minion : cirosantilli2, cc by-sa, [wikimedia](#).

auteur(s)/autrice(s) : erwin van dijk, d'après différentes sources licence : [cc-by-sa](#)

la première technologie ngs était la méthode « 454 », lancée en 2005 par la société 454 life sciences. leur séquenceur 454 genome sequencer produisait environ 200 000 lectures d'une longueur de 110 paires de bases par run, c'est-à-dire par cycle de fonctionnement. un an plus tard, la technologie de solexa (maintenant illumina) apparaissait, suivie par solid (sequencing by oligo ligation detection), commercialisée par life technologies en 2007. les premiers séquenceurs illumina et solid produisaient plusieurs dizaines de millions de lectures, donc beaucoup plus que la machine 454. en revanche, ces lectures étaient beaucoup plus courtes, environ 35 pb seulement. une comparaison plus complète de ces trois technologies, qui prend également en compte les temps de runs et leurs coûts a été présenté par mardis (2008) [2]. trois ans plus tard, en 2010, ion torrent apparaissait, une technologie à base de semi-conducteurs, sans détection optique des nucléotides fluorescents. cette technologie était moins chère et plus rapide que les autres. la première machine ion torrent produisait environ 3 millions de séquences d'une longueur de 100 nucléotides maximum.

au début, ces différentes technologies étaient en concurrence, mais grâce à une évolution plus rapide de la qualité des données, du débit et de la longueur de lectures que les autres technologies, illumina a progressivement renforcé sa position sur le marché et a aujourd'hui un quasi-monopole.

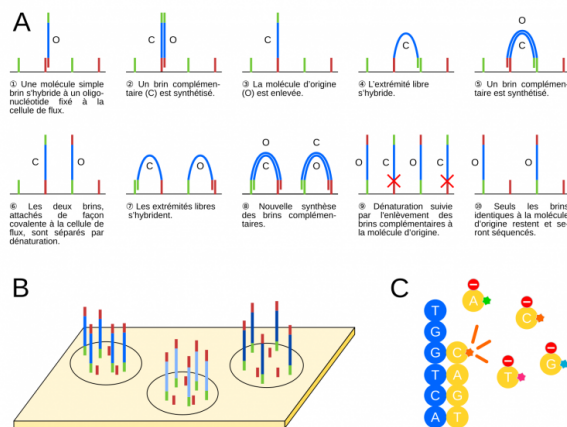
2.1. comment fonctionne la technologie illumina ?

illumina est aujourd'hui la technologie la plus largement utilisée ; nous discuterons donc seulement de cette technologie, sans rentrer dans les détails des autres méthodes. pour une description détaillée des diverses technologies, voir [3]. comme mentionné ci-dessus, les molécules d'une banque ngs contiennent des adaptateurs aux extrémités, qui ont deux fonctions principales : (1) l'accrochage et l'amplification sur une cellule de flux (*flow cell*), qui est une plaque en verre où les réactions de séquençage et la visualisation des nucléotides incorporés ont lieu, et (2) l'hybridation d'une amorce qui sert de point de départ pour le séquençage. la figure 2a montre schématiquement l'amplification des molécules de la banque sur une cellule de flux illumina. les fragments d'adn à séquencer sont d'abord dénaturés. les molécules simple brin obtenues vont alors s'hybrider à différents endroits de la cellule de flux, grâce à la complémentarité de leurs adaptateurs avec les oligonucléotides attachés de façon covalente à la cellule de flux (①). suivons le devenir de l'une de ces molécules d'adn. un brin complémentaire au brin matrice est synthétisé par une polymérase (②). ensuite, une étape de dénaturation sépare le brin matrice d'origine et le brin complémentaire néosynthétisé. après une étape de lavage, qui permet d'éliminer le brin matrice d'origine (③), l'extrémité libre du brin complémentaire s'hybride à un oligonucléotide complémentaire fixé à proximité sur la cellule de flux (④). cela permet une nouvelle synthèse d'un brin complémentaire (et donc identique au brin matrice d'origine, ⑤), suivie par une nouvelle étape de dénaturation afin de séparer les deux brins. on a alors les deux brins côte à côte sur la cellule de flux (⑥). le même processus se répète un certain nombre de fois (⑦,⑧), ce qui aboutit à la formation d'un « cluster », un groupe d'un millier de molécules identiques au brin matrice d'origine. ce processus est connu sous le nom d'amplification en pont (*bridge amplification*) dans la littérature. à d'autres endroits de la cellule de flux, se trouvent d'autres clusters correspondant à l'amplification d'autres fragments de la banque d'adn à séquencer (figure 2b). cette étape d'amplification est nécessaire à l'obtention d'un signal suffisamment important au moment du séquençage. à l'issue de cette étape, des amorces de séquençage s'hybrident à tous les brins de tous les clusters, et servent de point de démarrage du séquençage. un seul nucléotide marqué est incorporé à chaque cycle, suivi par une étape d'imagerie afin de déterminer quel nucléotide a été incorporé dans chaque cluster (figure 2c).

au fil des années, illumina a agrandi sa gamme de séquenceurs et propose aujourd'hui des machines qui peuvent générer de 4 millions à 20 milliards de séquences par run, avec une longueur maximum de 150 à 300 pb[2].

figure 2 - le principe de la technologie illumina

(a) représentation schématique du processus d'amplification en pont (*bridge amplification*) des clusters sur la cellule de flux (*flow cell*, trait noir). les adaptateurs et les oligonucléotides complémentaires présents sur la cellule de flux sont indiqués en rouge et vert, les fragments d'adn à séquencer en bleu. la molécule s'étant initialement fixée à la cellule de flux, et les brins dont la séquence est identique, sont notés o ; les brins complémentaires sont notés c. à l'étape 9, les brins complémentaires à la molécule d'origine sont éliminés par clivage chimique d'un nucléotide modifié dans l'oligonucléotide « rouge » fixé à la cellule de flux (croix rouges).



(b) différentes molécules d'adn sont amplifiées puis séquencées en même temps sur une même cellule de flux.

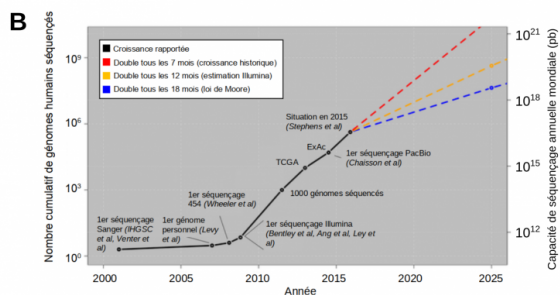
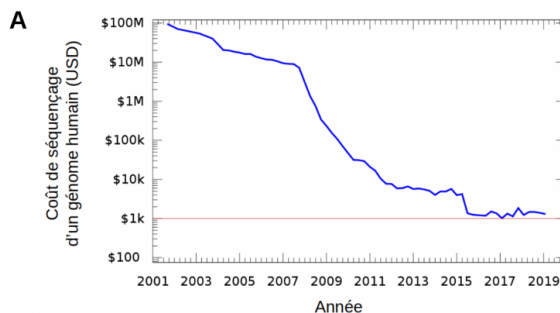
(c) réaction de séquençage. une polymérase synthétise un brin complémentaire (jaune) au brin situé sur la cellule de flux (bleu) en incorporant des nucléotides qui portent des groupes fluorescents (bleu, rose, orange, ou vert) et un groupe « stop » qui bloque la polymérisation. s'ensuit une étape d'imagerie afin d'identifier le nucléotide incorporé. ensuite, les groupes fluorescents et « stop » sont enlevés et le processus se répète.

auteur(s)/autrice(s) : pascal combemorel et erwin van dijk
licence : [cc-by-sa](#)

2.2. les applications du ngs

l'apparition des technologies ngs a permis le séquençage des génomes à une vitesse sans précédent et à un coût qui a chuté de façon spectaculaire par rapport à la technologie sanger (figure 3). le human genome project, achevé en 2003 après 13 ans d'effort, a coûté environ 3 milliards de dollars. seulement six ans plus tard, en 2009, le coût du séquençage d'un génome humain a été réduit à 200 000 dollars et en 2015, la barre symbolique de 1000 dollars a été franchie [4, 5]. ceci a mis le séquençage des génomes à la portée des laboratoires modestes, et le ngs a ainsi entraîné une révolution dans le domaine de la génomique. en effet, comme illustré dans la figure 3, le nombre de génomes séquencés a explosé depuis l'arrivée des technologies ngs.

figure 3 - évolution du coût et des capacités de séquençage
(a) évolution du coût du séquençage d'un génome humain. source : ben moore, réécrit en gnuplot par grendel|khan, cc0, [wikimedia](#).



(b) l'explosion du nombre de génomes séquencés depuis l'apparition des technologies ngs. le nombre total de génomes humains séquencés est indiqué à gauche et la capacité de séquençage annuelle mondiale est à droite. la croissance rapportée d'après les publications scientifiques est illustrée par la ligne noire avec des événements marquants indiqués (grands projets de séquençage tcga = *the cancer genome atlas*, exac = *exome aggregation consortium*). les lignes en pointillés représentent trois extrapolations différentes de la croissance de la capacité de séquençage ; en rouge, la croissance historique, en jaune la croissance estimée par illumina en 2014, et en bleue la croissance selon la loi de moore, formulée pour la puissance informatique, qui croît exponentiellement. source : adapté de stephens et coll., 2015, cc by, [plos biology](#).

auteur(s)/autrice(s) :
 traduction par erwin van dijk
 de différentes sources licence :
[cc-by-sa](#)

les méthodes de séquençage de nouvelle génération, utilisées pour l'étude des génomes, ont été mises à profit pour l'étude des transcriptomes. c'est ainsi que les méthodes de séquençage des arn (*rna-seq*) ont rapidement remplacé les technologies précédentes, qui étaient basées sur des micro-réseaux (*microarrays*). ces derniers utilisent des sondes avec des séquences prédéfinies auxquelles les arn s'hybrident et sont ainsi détectés. le rna-seq a plusieurs avantages par rapport à ces micro-réseaux. premièrement, il permet de détecter et de quantifier des transcrits sans connaissance a

priori de leur séquence. deuxièmement, il autorise la quantification des transcrits sans effet de saturation. en effet, les micro-réseaux étaient limités par le nombre des sondes ; la détection d'un arn extrêmement abondant pouvait arriver à un plateau suite à une saturation de celles-ci. avec le rna-seq au contraire, il n'y a pas de limite de détection quantitative. troisièmement, le rna-seq présente moins de bruit de fond que les micro-réseaux. le ngs est également devenu la technique de référence pour la détection des interactions protéine-adn par immunoprécipitation de la chromatine (chip), car le ngs est plus quantitatif et permet une meilleure résolution que les méthodes qui existaient jusqu'ici.

dans les années suivantes, le ngs a été utilisé pour des applications de plus en plus diverses. voici un résumé des applications les plus courantes.

2.2.1. séquençage des génomes à l'échelle de la population

bien évidemment, le séquençage des génomes constitue une application majeure du ngs qui a été mentionnée ci-dessus. mais les technologies ngs ont été développées initialement dans un but encore plus ambitieux : le séquençage d'un très grand nombre de génomes de la population humaine. les données collectées permettent alors d'identifier les bases génétiques des variations phénotypiques. ces connaissances sont ensuite utilisées pour des études d'association entre variations génétiques et maladies (*genome-wide association studies*, gwas). ces variations peuvent être de deux types. d'une part, il existe des variations rares qui ont un effet grave sur des caractéristiques simples (par exemple la mucoviscidose, la maladie de huntington). d'autre part, certaines variations plus fréquentes ont des effets plus légers et pourraient être impliquées dans des phénotypes complexes (par exemple le diabète, les maladies cardiovasculaires). un des premiers projets était le *1000 genomes project*, initié en 2008, dans lequel un millier de génomes humains ont été séquencés afin d'établir un catalogue des variantes génétiques ayant une fréquence d'au moins 1 % dans les populations étudiées [6]. ce projet a été suivi par des initiatives encore plus ambitieuses, par exemple le *100,000 genomes project* [7].

2.2.2. rna-seq

une autre application majeure est le séquençage d'arn (*rna-sequencing*, rna-seq) qui concerne le séquençage de l'ensemble des transcrits (transcriptome) d'une cellule, d'un tissu ou d'un organisme. le rna-seq est majoritairement utilisé pour l'analyse de l'expression différentielle des gènes entre des cellules soumises à des conditions différentes.

2.2.3. détection des interactions et/ou des modifications des acides nucléiques

le ngs est aussi un outil très puissant pour la détection d'une grande diversité d'interactions entre moléculr d'acides nucléiques, entre acides nucléiques et protéines, ainsi que des modifications « épigénétiques » des acides nucléiques[3]. une technique couramment utilisée est le « *chromatin immunoprecipitation-sequencing* » ou « chip-seq ». cette méthode est basée sur la reconnaissance par un anticorps d'un partenaire (souvent une protéine) ou d'une modification épigénétique. après fragmentation de la chromatine, cet anticorps permet de précipiter spécifiquement les fragments en interaction avec le partenaire d'intérêt, ou qui contiennent la modification épigénétique recherchée. cela permet de constituer une banque de fragments d'adn d'intérêt qui sont ensuite séquencés.

une autre technique très utilisée est celle du « *chromosome conformation capture* » qui permet l'étude des interactions entre différentes régions du génome. il existe plusieurs variantes de cette technique, qui permettent soit de détecter des interactions entre deux régions connues (3c), soit des interactions entre une région connue et une ou plusieurs régions inconnues (4c), ou de cartographier toutes les interactions entre des régions *a priori* inconnues (hic).

2.2.4. séquençage ciblé - séquençage des amplicons

pour un grand nombre d'applications il n'est pas nécessaire de séquencer la totalité du génome ou du transcriptome. plusieurs techniques ont été développées qui permettent de séquencer uniquement la partie du génome ou du transcriptome qui est la plus intéressante pour répondre à la question posée. par exemple, une technique couramment utilisée pour étudier des maladies génétiques rares est le séquençage de l'exome (*whole exome sequencing*, wes), c'est-à-dire de l'ensemble des exons. sachant que chez l'espèce humaine l'exome constitue uniquement environ 1 % du génome, cette technique permet de réduire énormément la profondeur de séquençage requise et donc de séquencer un très grand nombre d'échantillons avec relativement peu de lectures.

2.2.5. séquençage des cellules uniques

le progrès des techniques de séquençage à haut-débit a permis l'apparition des méthodes de séquençage de cellules uniques, avec le rna-seq comme application majeure. en effet, les méthodes classiques de rna-seq ne fournissent qu'une moyenne des profils transcriptomiques de toutes les cellules dans une population. à l'inverse, le séquençage arn sur cellule unique (*single cell rna-seq*, scrna-seq) a pour avantage de montrer la diversité des profils transcriptomiques entre cellules. de plus, le scrna-seq permet une bien meilleure résolution des profils d'expression des gènes. par exemple, cette technique peut déterminer si des gènes coexprimés dans une population sont coexprimés dans la même cellule ou dans des cellules différentes.

une autre possibilité intéressante est le séquençage des génomes de cellules uniques. ceci permet par exemple d'étudier la diversité génétique entre cellules dans une tumeur et d'étudier son évolution, c'est-à-dire l'accumulation des mutations et réarrangements génétiques au cours de la progression tumorale. une autre application importante est le séquençage des génomes des microorganismes qui ne peuvent être cultivés en laboratoire.

3. la troisième révolution : l'apparition des technologies de séquençage *long-read*

même si les technologies ngs sont extrêmement puissantes, elles présentent aussi quelques faiblesses, par exemple la faible longueur des fragments séquencés. des génomes contiennent souvent de nombreuses séquences répétées dont la longueur excède celle des fragments lus par ngs, ce qui mène à des difficultés dans l'assemblage du génome. par conséquent, un grand nombre de génomes ne sont pas publiés sous la forme d'une séquence continue mais sont disponibles sous une forme très fragmentée, en centaines ou milliers de contigs[4]. pour la même raison, les variantes structurales (réarrangements génomiques de plus de 50 paires de bases) entre les génomes de différents individus sont souvent difficiles à détecter. ceci est un problème important,

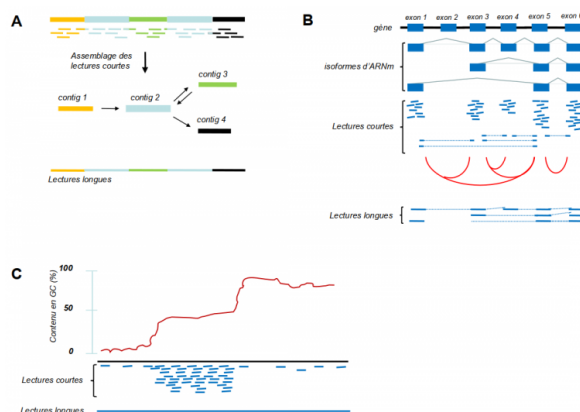
vu que les variantes structurales sont souvent liées à des maladies génétiques. un autre problème est que les méthodes ngs ne permettent pas toujours de caractériser des isoformes de transcrits générés par épissage alternatif. de plus, du fait que les méthodes ngs nécessitent une étape d'amplification par pcr, les régions génomiques avec un pourcentage élevé de gc sont souvent sous-représentées, car elles sont moins bien amplifiées par pcr.

rapidement après l'apparition des méthodes ngs, des technologies dites de troisième génération (tgs, *third generation sequencing*) ou de lectures longues (*long read*) sont apparues. elles se caractérisent par le séquençage, *en temps réel* [5] de molécules *uniques* [8]. ces technologies permettent de générer des lectures d'une longueur de plusieurs kilobases voire même des centaines de kilobases. la figure 4 résume les avantages des technologies *long read* par rapport au ngs.

figure 4 - comparaison des performances des lectures courtes ngs et des lectures longues tgs

(a) exemple qui montre le problème des régions répétées (en bleu clair) dans un génome, qui ne peuvent être positionnées à des endroits uniques avec des lectures courtes (ngs). des lectures qui se situent à l'intérieur de ces régions seront assemblées en un seul contig. la région entre les répétitions (vert) peut être placée en amont ou en aval de ce contig bleu (illustré par les flèches en double-sens), ce qui génère une ambiguïté. par contre, des lectures longues (*long reads*, tgs) qui traversent ces régions répétées ne laissent aucune ambiguïté. (b) plusieurs transcrits (isoformes) peuvent être générés à partir d'un seul gène par épissage alternatif. des lectures courtes de ces isoformes donneront des lectures à l'intérieur des exons présents dans le mélange (lignes continues) ainsi que des lectures qui chevauchent les jonctions entre les exons (lignes interrompues). ainsi, les événements d'épissage alternatif seront détectés ; les jonctions exon-exon détectées dans le mélange sont indiquées par des lignes rouges. toutefois, les informations concernant les combinaisons des jonctions exon-exon dans les transcrits individuels manquent. des lectures longues qui couvrent les transcrits entiers fournissent ces informations. (c) les méthodes ngs dépendent de

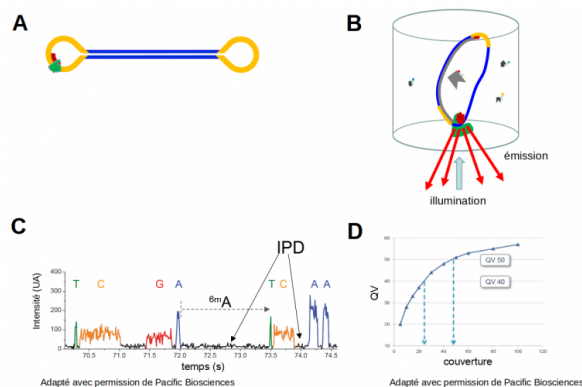
l'amplification par pcr, ce qui introduit des biais dans des régions pauvres en gc (à gauche de la courbe) ou très riche en gc (à droite). par conséquent, ces régions seront peu couvertes. les technologies de troisième génération pacbio et nanopore n'ont pas besoin de l'amplification pcr et présentent donc beaucoup



en 2011, pacific biosciences lançait sa technologie *single molecule real time sequencing* (*smrt*), basée sur une polymérase attachée au fond d'un micro-puits. cette polymérase synthétise un brin complémentaire à partir de molécules d'adn simple brin circulaires (figure 5). elle utilise des nucléotides marqués, comme dans la technologie illumina, mais avec pour différence importante que la synthèse ne s'arrête pas après l'incorporation de chaque nucléotide. la synthèse continue jusqu'au moment où la polymérase s'arrête spontanément (après 10 000 à 100 000 nucléotides incorporés), et l'incorporation des nucléotides est suivie en temps réel dans un film (*movie*). un aspect important à mentionner est que cette méthode permet de détecter directement des nucléotides méthylés, car le temps d'incorporation d'un nucléotide modifié est un peu différent de celui d'un nucléotide non modifié. le taux d'erreurs de cette méthode est assez élevé pour une seule lecture (environ 10 %), mais du fait de la circularité de la molécule d'adn à séquencer, la polymérase fait le tour plusieurs fois, ce qui augmente énormément la fiabilité jusqu'à des valeurs comparables ou même supérieures à celles de la technologie illumina.

figure 5 - la technologie de séquençage single molecule real time (smrt) de pacific biosciences (pacbio)

(a) pendant la préparation de la banque, des adaptateurs en tige boucle (jaune) sont attachés à des molécules d'adn double-brin (bleu), ce qui donne des molécules circulaires (*smrt bells*). ensuite, une amorce (rouge) et une polymérase (vert) se fixent sur l'adaptateur. (b) représentation schématique d'un puits dans lequel le séquençage a lieu ; la polymérase, qui se fixe sur le fond du puits, incorporera des nucléotides fluorescents, et à chaque évènement d'incorporation un signal fluorescent (« émission ») sera généré après illumination par un laser. ces signaux sont enregistrés par une caméra en temps réel, ce qui génère un film (*movie*). (c) exemple de movie. les différents nucléotides sont identifiés par leur couleur ; le temps entre les incorporations successives est enregistré également (*interpulse duration - ipd*, en noir). on remarque que dans cette technologie de séquençage, la vitesse de polymérisation (2 à 4 nucléotides par seconde) est bien plus faible qu'*in vivo*. la présence d'une modification épigénétique, comme la 6-méthyladénosine (6ma) est responsable d'un ipd plus long, ce qui permet son identification. (d) la qualité des données de séquençage (*quality value - qv*) augmente fortement quand la polymérase fait plusieurs fois le tour des molécules circulaires. par exemple, quand la polymérase fait 25 fois le tour, la précision atteint 99,999 % (qv40) et pour 50 tours, la précision atteint 99,9999 % (qv50), ce qui correspond à un taux d'erreur inférieur à celui du séquençage illumina.



en 2014, la méthode d'oxford nanopore technologies (ont) est apparue sur le marché. la technologie nanopore est unique car elle n'utilise pas de polymérase pour la synthèse d'une copie de la molécule à séquencer. la séquence d'une molécule d'adn ou d'arn est directement déduite grâce à la perturbation d'un courant électrique traversant un « nanopore » incorporé dans une membrane qui sépare deux compartiments contenant des solutions ioniques (figure 6). la perturbation du signal est une fonction de la composition des nucléotides présents dans le pore à chaque instant, et c'est en suivant en temps réel les successions de ces perturbations qu'on peut en déduire la séquence. comme pour la technologie smrt de pacific biosciences, on peut détecter directement les modifications des nucléotides car celles-ci peuvent perturber le signal électrique de façon spécifique.

la technologie nanopore peut générer des lectures extrêmement longues, qui peuvent aller jusqu'à 1 mb ou même plus. *a priori* le seul facteur limitant serait la longueur des fragments d'adn que l'on réussit à obtenir lors de l'extraction de cette molécule. en effet, les molécules d'adn génomique sont généralement fragmentées pendant cette étape. un vrai défi est donc de trouver des méthodes d'extraction qui laissent l'adn le plus intact possible.

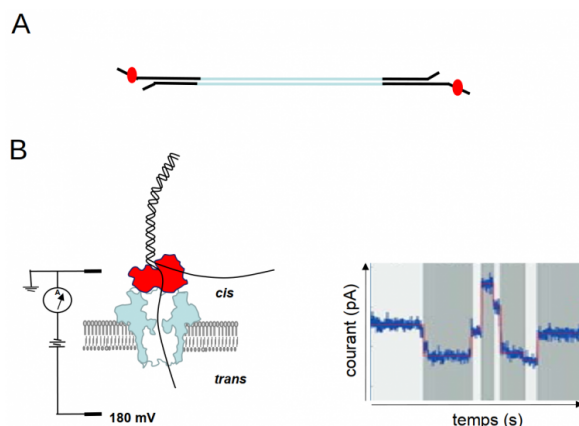


figure 6 - la technologie nanopore

(a) dans une banque nanopore, des adaptateurs (en noir) sont attachés aux molécules d'adn double-brin à séquencer (en bleu), comme pour les banques illumina. sur chaque adaptateur, une protéine motrice (rouge) est fixée. (b) dans une cellule de flux d'oxford nanopore technologies (ont), deux compartiments (*cis* et *trans*) contenant des solutions ioniques sont séparés par une membrane qui contient des pores ménagés par des protéines. un acide nucléique (noir) traverse le pore (en bleu) de façon contrôlée grâce à la présence de la protéine motrice (rouge). avant de traverser le pore, l'acide nucléique est dénaturé par la protéine motrice et ainsi un seul des deux brins traverse le pore. quand une molécule d'adn ou d'arn traverse le pore, le courant électrique est modifié ; celui-ci est enregistré en temps réel et graphiquement représenté par un *squiggle plot* (à droite).

auteur(s)/autrice(s) : erwin van dijk licence : [cc-by-sa](#)

4. traitement des données de séquençage : le rôle de la bio-informatique.

les technologies ngs et tgs ont nécessité le développement d'outils de bio-informatique spécifiques, capables d'analyser les données générées par les différentes méthodes. un grand nombre d'algorithmes capables de gérer des lectures courtes ngs ainsi que des algorithmes permettant l'assemblage de génomes *de novo*, la détection des polymorphismes mononucléotidiques (snp), l'analyse chip-seq et rna-seq ont été développés [9, 10]. des programmes corrigeant les biais introduits pendant la préparation des banques ont également été mis au point. l'ensemble des améliorations expérimentales et algorithmiques ont fortement renforcé l'utilité des technologies ngs.

l'apparition des technologies de troisième génération a également nécessité le développement d'outils adaptés permettant de gérer des lectures (très) longues et des taux d'erreurs élevés [11].

5. conclusions et perspectives

il y a un peu plus de quarante ans, l'apparition de la technologie de séquençage sanger a permis pour la première fois de déchiffrer des gènes et des génomes, ce qui a révolutionné le domaine de la génomique et de la biologie en général. dans les années 2000, le développement des méthodes ngs, beaucoup plus puissantes que le séquençage sanger, a fait chuter les prix du séquençage de façon considérable et a mis le séquençage des génomes à la portée des laboratoires modestes. plus récemment encore, le développement et la mise au point des technologies de troisième génération, ou *long read*, a ouvert des possibilités nouvelles permettant le séquençage des génomes avec une qualité sans précédent.

nous disposons donc aujourd'hui d'une boîte à outils extrêmement puissante pour l'exploration des génomes et les années récentes ont vu de nombreuses découvertes excitantes qui vont sans doute continuer dans les années à venir.

malgré les progrès impressionnants accomplis ces dernières années, il reste des défis majeurs pour l'avenir. au niveau de l'analyse bio-informatique des quantités massives de données de séquençage, un objectif important sera de développer des outils plus performants et rapides, notamment dans le cadre des utilisations cliniques, où chaque jour compte. il y a aussi des questions importantes concernant la gestion et l'utilisation des données de séquençage : comment stocker et protéger ces données, quelles informations faut-il partager et comment, etc.

afin d'obtenir une compréhension de la cellule plus profonde et complète, un défi majeur pour l'avenir sera de réussir à mettre en relation les résultats de la génomique avec les données épigénomiques, transcriptomiques et protéomiques. alors qu'il existe aujourd'hui des outils puissants et rapides pour l'étude des génomes, épigénomes et transcriptomes, l'exploration des protéomes reste techniquement très délicate. il est intéressant dans ce cadre de mentionner que la technologie nanopore permettrait, *a priori*, de séquencer des protéines [12], mais à ce jour cette application n'est pas encore opérationnelle. la mise au point d'une méthode rapide et économique pour séquencer l'ensemble des protéines d'une cellule, tissu ou organisme marquera sans doute la prochaine révolution dans l'« omique ».

6. références

- [1] [about the human genome project](#), human genome project information archive
- [2] elaine r. mardis. the impact of next-generation sequencing technology on genetic trends in genetics, 2008, vol 24, n° 3 <https://doi.org/10.1016/j.tig.2007.12.007>
- [3] sara goodwin, john d. mcpherson and w. richard mcombie. coming of age: ten years of next generation sequencing technologies, nat. rev. genet. 2016, 17:333-351
- [4] [the \\$1,000 genome](#), infographie d'illumina
- [5] [genome sequencing stocks on the rise](#), ken berman, forbes, 21 février 2019
- [6] [the international genome sample resource](#)
- [7] [the 100,000 genomes project](#), genomics england
- [8] erwin l. van dijk, yan jaszczyszyn, delphine naquin, and claude thermes. the third

revolution in sequencing technology. trends in genetics, 2018, vol. 34, n° 9
<https://doi.org/10.1016/j.tig.2018.05.008>

[9] hatem, a. et al. benchmarking short sequence mapping tools. bmc bioinformatics, 2013, 14, 184

[10] rucha m. wadapurkar, renu vyas. computational analysis of next generation sequencing data and its applications in clinical oncology (2018).
<https://doi.org/10.1016/j.imu.2018.05.003>

[11] shanika l. amarasingh, shian su, xueyi dong, luke zappia, matthew e. ritchie and quentin gouil. opportunities and challenges in long-read sequencing data analysis. genome biology, 2020, 21:30 <https://doi.org/10.1186/s13059-020-1935-5>

[12] kolmogorov m, kennedy e, dong z, timp g, pevnzner pa. single-molecule protein identification by sub-nanopore sensors. *plos comput biol.* 2017;13(5):e1005356. doi:10.1371/journal.pcbi.1005356

CRÉDITS

AUTEUR(S)/AUTRICE(S)

[erwin van dijk](#)

ingénieur de recherche en génétique moléculaire. il travaille sur la plateforme de séquençage à haut débit de l'institut de biologie intégrative de la cellule (i2bc) à gif-sur-yvette.

[claire thermes](#)

directeur de recherche émérite responsable, depuis 2010, de la plateforme de séquençage de l'institut de biologie intégrative de la cellule (gif-sur-yvette).

MISE EN LIGNE

[pascal combemorel](#)

agréé de svt, il est le responsable éditorial du site planet-vie depuis septembre 2016.

LICENCE DU TEXTE DE L'ARTICLE



creative commons - attribution - pas d'utilisation commerciale

notes

1

en effet, pour couvrir l'ensemble du génome, il est nécessaire d'en séquencer plusieurs copies (on parle de profondeur de lecture). pour plus d'informations sur le taux de couverture et la profondeur de lecture, voir l'article [le séquençage des génomes](#).

2

pour un séquençage des fragments dans « un seul sens », tel que présenté figure 2. en plus de cette approche *single end*, une approche *paired end* existe, qui consiste à séquencer, en plus du brin matrice, le brin complémentaire. cela permet de détecter les insertions, délétions et autre réarrangements chromosomiques lors de l'assemblage d'un génome et de sa comparaison à un génome de référence.

3

pour une présentation des modifications épigénétiques, voir l'article [épigénétique et cancer](#).

4

un contig est une séquence génomique continue obtenue par l'assemblage des fragments chevauchants d'une banque d'adn (voir figure 4).

5

contrairement au ngs pour lequel le séquençage est mis en pause après l'incorporation de chaque base.