

constituerait leur mise à l'épreuve. Si l'hypothèse de l'allocation de lien était venue assez naturellement, il n'en était pas de même de sa validation. Ce test devait attendre le moment propice.

La situation a changé quand Denise Cai et Justin Shobe, tous deux alors au laboratoire, se sont joints au projet. Denise Cai proposa une expérience astucieuse: Justin Shobe et elle ont commencé à placer des souris dans deux cages au cours de la même journée, à cinq heures d'intervalle, dans l'espoir que les souvenirs des deux cages se trouveraient connectés l'un à l'autre. Plus tard, ils leur administrèrent un léger choc électrique alors qu'elles étaient dans la seconde cage. Comme prévu, quand, ensuite, les souris ont à nouveau été placées dans la cage où elles avaient reçu le choc électrique, elles se sont immobilisées, probablement parce qu'elles se rappelaient avoir reçu un tel choc en ce lieu. L'immobilisation est une réaction naturelle chez les souris apeurées, parce que la plupart des prédateurs les détectent alors moins facilement.

Le résultat clé est apparu lorsque Denise Cai et Justin Shobe ont placé les souris dans la première cage, où elles n'avaient jamais reçu de choc électrique. Nous pensions que si les souvenirs des deux cages étaient liés, les souris placées dans la première se rappelleraient du choc électrique reçu dans la seconde et, par anticipation, s'immobiliseraient – et c'est exactement ce qui s'est produit.

Nous avons aussi prédit que les deux souvenirs auraient moins de chances d'être liés s'ils étaient séparés de sept jours. Et effectivement, les souris replacées dans la première cage ne semblaient pas se souvenir de la deuxième lorsque l'intervalle de temps entre les deux premières expositions dépassait une semaine. En général, au-delà d'une journée d'intervalle, les souvenirs ne s'associaient pas.

UN MICROSCOPE PORTABLE POUR SOURIS

Aussi excitantes que fussent ces découvertes comportementales, elles ne suffisaient pas à confirmer une prédiction essentielle découlant de notre hypothèse: l'idée que les souvenirs individuels formés à peu de temps d'intervalle soient stockés dans une même zone du cerveau, au sein de populations neuronales se recouvrant partiellement. Ce recouvrement physique devait lier les souvenirs, de sorte que le rappel de l'un entraîne celui de l'autre.

Pour tester l'hypothèse d'allocation de lien, il fallait idéalement être en mesure de voir les souvenirs dans le cerveau au moment où ils se forment. Des techniques d'imagerie des neurones *in vivo* existaient déjà, mais elles nécessitaient de maintenir la tête de l'animal immobile, fixée à de gros microscopes, un

montage incompatible avec les expériences de comportement que nous devions mener.

Par chance, la bonne technique est apparue juste au moment où nous en avions le plus besoin. Lors d'un séminaire à l'UCLA auquel j'assistais, Mark Schnitzer, de l'université Stanford, décrit un tout petit microscope que son laboratoire venait d'inventer, capable de visualiser l'activité de neurones dans le cerveau de souris libres de se déplacer. Ce microscope de 2 à 3 grammes était porté par l'animal comme un chapeau. Cet instrument s'est révélé idéal pour pister les neurones activés par un souvenir donné. Il nous a permis de savoir si ces mêmes neurones deviennent actifs quelques heures après, lorsqu'un autre souvenir est créé – c'est-à-dire qu'il a permis de tester notre prédiction.

Nous étions si excités par cette invention, que nous avons décidé d'en créer notre propre version. Nous nous sommes associés avec les laboratoires de Peyman Golshani et Baljit Khakh, tous deux établis à UCLA, et avons embauché un chercheur postdoctoral du nom de Daniel Aharoni, qui s'attela à la mise au point de ce que nous appellerions plus tard les miniscopes de l'UCLA. Équipé d'une lentille insérée à proximité immédiate des cellules nerveuses, l'appareil est fixé au crâne de la souris et reste stable tout au long des tests qu'elle subit.

Deux souvenirs formés à peu de temps d'intervalle ont plus de chances d'être liés

Pour visualiser uniquement les neurones activés par un souvenir, nous avons utilisé une autre technique dite d'indicateur de calcium codé génétiquement. Il s'agit d'injecter à l'animal une molécule qui devient fluorescente en présence de calcium: or un neurone qui s'active absorbe le calcium présent dans le milieu extracellulaire... Et devient donc visible au microscope.

Nous avons décidé de nous concentrer sur la région CA1 de l'hippocampe en raison de son rôle dans l'apprentissage et la mémorisation des lieux – par exemple, les cages que ➤