

Organizzazione corticale

Questo bellissimo pattern è una sezione di una piccola area del cervello. Le aree nere simili ad alberi sono spazi d'aria, mentre quelle bianche e arancioni sono tessuto cerebrale. Le aree bianche, denominate materia bianca, sono principalmente fibre nervose; le aree arancioni, denominate materia grigia, sono per la maggior parte i corpi cellulari dei neuroni. La materia grigia assieme a quella bianca, contengono dei meccanismi non solo per la percezione, ma per tutto ciò che il cervello controlla e crea. In questo capitolo consideriamo come le diverse funzioni percettive sono organizzate nel cervello.

La ricerca sulla fisiologia della visione è dominata da due maggiori tematiche: (1) descrivere i tipi di stimoli che causano la risposta di neuroni a diversi livelli del sistema visivo; (2) e descrivere come i neuroni sono organizzati nel sistema visivo. Abbiamo considerato il primo tema nel Capitolo 3 quando abbiamo mostrato come i neuroni ad alti livelli del sistema visivo rispondono a stimoli sempre più complessi. Questo capitolo considera il problema dell'organizzazione. Come vedremo, questo problema è stato approcciato in diversi modi, ognuno di questi importante per capire come il sistema visivo processa l'informazione.

Il sistema visivo organizzato

L'organizzazione è importante. Noi abbiamo bisogno di essere organizzati. Le aziende hanno un diagramma organizzativo. L'esercito ha una catena di comando.

Organizzare le informazioni in uno schedario o nel tuo computer rende più semplice accedere all'informazione quando ne hai bisogno.

La necessità dell'organizzazione è particolarmente importante nel sistema visivo a causa dei compiti che deve svolgere. Uno di questi compiti è processare le informazioni riguardo varie caratteristiche o peculiarità degli oggetti, quali la forma, la dimensione, l'orientamento, il colore, il movimento, la collocazione spaziale. Vedremo che ogni caratteristica è elaborata da differenti meccanismi localizzati in diverse aree del cervello. Una volta che le informazioni sono processate, come si combinano tutte le caratteristiche di un oggetto? Noi non vediamo separatamente "rosso", "camion", "lungo", "si sta muovendo verso sinistra". Noi vediamo un macchinario rosso fuoco che corre lungo la strada. L'organizzazione gioca un ruolo centrale nel riuscire a processare

specifiche informazioni e a combinare l'informazione per creare percezioni coerenti. Il sistema visivo raggiunge questa organizzazione in diversi modi. Inizieremo considerando *l'organizzazione spaziale*, ovvero come aree diverse dell'ambiente e sulla retina sono rappresentate nel cervello.

Un'esplorazione dell'organizzazione spaziale

L'*organizzazione spaziale* si riferisce al modo in cui gli stimoli in specifiche aree dell'ambiente sono rappresentate come attività in specifiche zone del sistema nervoso. Per esempio, quando guardiamo una scena, le cose sono organizzate attraverso il nostro campo visivo. Ci sono oggetti a destra e a sinistra, in alto e in basso. Quest'organizzazione nello spazio visivo si trasforma successivamente in organizzazione nell'occhio, quando un'immagine della scena si crea sulla retina. È facile apprezzare l'organizzazione spaziale a livello di immagine retinica perché quest'immagine è essenzialmente una foto della scena. Ma quando la foto si trasforma in segnali elettrici, si verifica un nuovo tipo di organizzazione sotto forma di mappa elettronica della retina in strutture localizzate più in alto nel sistema.

La mappa elettronica su V1

Per cominciare, chiediamoci come i punti nell'immagine retinica sono rappresentati spazialmente nella corteccia striata. Determineremo questo stimolando diversi luoghi sulla retina e determinando dove i neuroni si attivano nella corteccia. La **figura 4.1** raffigura un uomo

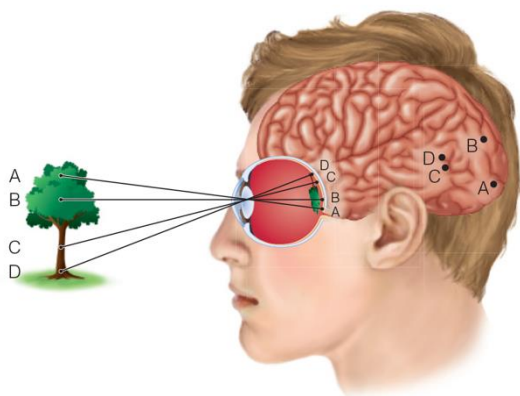
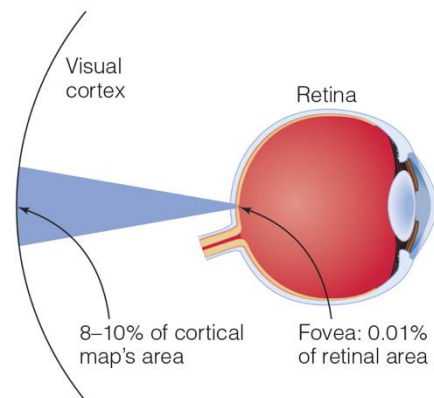


Figura 4.1: Una persona che sta guardando un albero, mostrando come i punti A, B, C, D sono rappresentati sulla retina e dove le attivazioni retiniche causano attività cerebrale. Sebbene le distanze tra A e B e tra C e D sono simili sulla retina, la distanza tra A e B è più grande sulla corteccia. Questo è un esempio di ingrandimento corticale, per cui più spazio è assegnato alle aree della retina vicino alla fovea.

che guarda la cima di un albero in modo che i punti A, B, C e D sull'albero stimolino i punti A, B, C, D sulla retina. Muovendoci sulla corteccia, l'immagine al punto A sulla retina causa l'attivazione dei neuroni al punto A nella corteccia. L'immagine al punto B causa l'attivazione dei neuroni al punto B e così via. Questo esempio dimostra come i punti sull'immagine retinica causano l'attività nella corteccia. Ma possiamo anche invertire il processo registrando da un neurone nella corteccia e determinando la localizzazione del suo campo ricettivo sulla retina. Quindi, se noi registriamo da un neurone al punto A nella corteccia, il suo campo ricettivo sarà localizzato al punto A sulla retina; se noi cominciamo dal punto B, il campo ricettivo è al punto B; e così via. Questi esempi dimostrano che la localizzazione sulla corteccia corrisponde alla localizzazione sulla retina. Questa mappa elettronica della retina sulla corteccia è denominata **mappa retinotopica**. Questa mappa spaziale organizzata indica che due punti che sono vicini su un oggetto e sulla retina attivano neuroni che sono vicini nel cervello (Silver e Kastner, 2009).

Ma guardiamo questa mappa retinotopica più da vicino, perché possiede un'interessante proprietà che è rilevante per la percezione. Sebbene i punti A, B, C e D sulla corteccia corrispondono ai punti A, B, C e D sulla retina, potresti notare qualcosa riguardo la *distanza* di queste aree. Considerando la retina, notiamo che l'uomo sta guardando alla cima dell'albero, quindi i punti A e B sono entrambi vicini alla fovea e le immagini dei punti C e D alla fine del tronco sono nella retina periferica. Sebbene la distanza tra A e B e tra C e D sia la stessa sulla retina, la distanza non è la stessa sulla corteccia. A e B occupano più spazio sulla

corteccia che C e D. Questo significa che la mappa sulla corteccia è distorta, con più spazio riservato alle aree più vicine alla fovea rispetto a quelle nella retina periferica. Sebbene la fovea conti solo per lo 0,01% dell'area della retina, i segnali dalla fovea contano dall'8 al 10 per cento della mappa retinotopica sulla corteccia (Van Essen e Anderson, 1995). Questa ripartizione alla piccola fovea di un'area grande sulla corteccia è denominata **ingrandimento corticale** (Figura



4.2).

Figura 4.2: Il fattore dell'ingrandimento nel sistema visivo. La piccola area della fovea è rappresentata da una grande area sulla corteccia visiva.

Il fattore dell'ingrandimento corticale è stato determinato nella corteccia umana usando una tecnica denominata *brain imaging*, che rende possibile creare immagini dell'attività cerebrale (**Figura 4.3**). Descriveremo la procedura del brain imaging e come questa è stata usata per misurare il fattore dell'ingrandimento corticale negli umani.



Figura 4.3: una persona in un apparecchio di scanning cerebrale.

METODO

Brain imaging

Il **Brain imaging** si riferisce a una serie di tecniche che si traducono in immagini che mostrano quali aree del cervello sono attive. Una di queste tecniche, la **tomografia a emissione di positroni (PET)**, è stata introdotta nella metà del 1970 (Hofmann et al., 1976; Ter-Pogossian et al., 1975). Nella procedura della PET, si inietta alla persona una bassa dose di tracciante radioattivo che non è pericolosa. Il tracciante entra nel flusso sanguigno e indica il volume del sangue che scorre. Il principio base della PET è che le parti del cervello che sono attive richiedono più apporto di sangue rispetto alle altre, parti meno attive del cervello. Quindi, cambiamenti nell'attività del cervello sono accompagnati da cambiamenti nel flusso sanguigno, così monitorare la radioattività del tracciante iniettato fornisce la misura dell'attività cerebrale.

Un'altra tecnica di neuroimaging è la **risonanza magnetica funzionale (fMRI)**. Come la PET, la fMRI è basata sulla misura del flusso sanguigno. Poiché l'emoglobina, che trasporta ossigeno nel sangue, contiene una molecola di ferro e quindi ha proprietà magnetiche, la presenza di un campo magnetico al cervello fa sì che le molecole di emoglobina si allineino come minuscoli magneti. La risonanza magnetica funzionale indica la presenza di attività cerebrale perché nelle aree di alta attività cerebrale le molecole di emoglobina perdono parte dell'ossigeno che trasportano. Questo rende l'emoglobina più magnetica, quindi queste molecole rispondono in maniera più forte al campo magnetico. L'apparecchio per la fMRI determina la relativa attività di varie aree del cervello tracciando cambiamenti nella risposta magnetica dell'emoglobina che si presentano quando una persona percepisce uno stimolo o intraprende uno specifico comportamento. Dato che la fMRI non richiede un tracciante radioattivo e dato che è più precisa, questa tecnica è diventata il metodo principale per localizzare l'attività cerebrale negli umani.

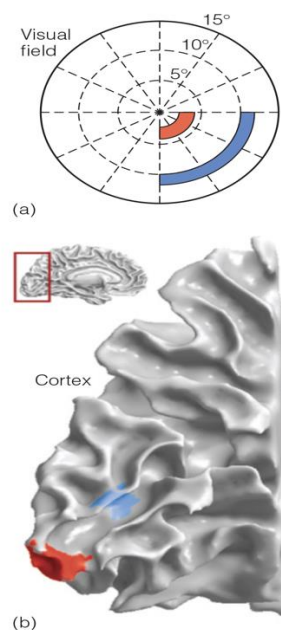


Figura 4.4: (a) Le aree blu e rosse rappresentano l'entità degli stimoli che erano presentati mentre la persona era in uno scanner per la fMRI. (b) Le aree blu e rosse indicano aree del cervello attivate dalla stimolazione in (a).

Robert Dougherty e altri collaboratori (2003) usarono il brain imaging per determinare il fattore dell'ingrandimento nella corteccia visiva umana. La **figura 4.4a** mostra il display degli stimoli visto dall'osservatore, che era posto in uno scanner fMRI. L'osservatore guardava dritto al centro dello schermo, quindi il punto al centro cadeva sulla fovea. Durante l'esperimento, lo stimolo luminoso era presentato in due luoghi: vicino al centro (area rossa), che illuminava una piccola area vicino alla fovea; e distante dal centro (area blu), che illuminava un'area della retina periferica. Le aree del cervello attivate da questi due stimoli sono indicate nella **figura 4.4b**. Questa attivazione illustra che il fattore di ingrandimento, grazie alla stimolazione di piccole aree vicino alla fovea, attiva un'area grande sulla corteccia (rosso) rispetto alla stimolazione di un'area più larga nella periferia (blu). (Vedi anche Wandell, 2011). La grande rappresentazione della fovea nella corteccia è anche illustrata nella **figura 4.5**, che mostra lo spazio che sarebbe assegnato alle parole su una pagina (Wandell et al, 2009). È da notare che la lettera "a", che è vicino a dove la persona sta guardando (freccia rossa), è rappresentata da un'area più grande nella corteccia rispetto alle lettere che sono più lontane dal punto di osservazione della persona.

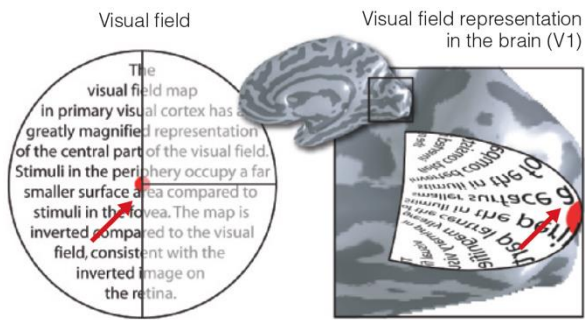


Figura 4.5 Dimostrazione del fattore d'ingrandimento. Una persona guarda il punto rosso nella figura a sinistra, l'area del cervello attivata da ogni lettera del testo viene mostrata nella figura a destra. La freccia indica la lettera a nel testo a sinistra e a destra si può vedere l'area cerebrale attivata dalla lettera a.

Lo spazio extra corticale designato alle lettere e alle parole che la persona sta guardando fornisce il lavoro neurale ulteriore ad eseguire compiti come la lettura, i quali richiedono una grande acutezza visiva (Azzopardi & Cowey, 1993).

Ciò che il fattore di ingrandimento significa quando osservi una scena è che la parte che stai guardando occupa un'area maggiore della corteccia rispetto alle zone che non stai guardando. Un altro modo per comprendere il fattore di ingrandimento è la seguente dimostrazione.

DIMOSTRAZIONE

Ingrandimento corticale del tuo dito

Allunga il braccio sinistro alzando l'indice e, mentre lo guardi, allunga anche il braccio destro a circa un piede di distanza in modo da avere il dorso della mano diretto verso il tuo viso. A questo punto il tuo indice sinistro attiva un'area della corteccia grande quanto la tua intera mano destra.

Una delle cose interessanti riguardo alla dimostrazione sopra è che quando il tuo dito è impresso sulla fovea occupa una porzione di corteccia pari alla tua mano impressa sulla tua retina periferica ma non percepisci il dito grande quanto la mano. D'altro canto osservi i dettagli del tuo dito con molta più definizione di quelli della mano. Il fatto che ad una maggiore area corticale occupata corrisponda una visione più dettagliata è un esempio di come ciò che percepiamo non è esattamente ciò che il nostro cervello elabora. Torneremo presto a quest'idea.

La corteccia è organizzata in colonne

Determinare la mappa retinica e il fattore di ingrandimento ci ha tenuti vicino alla superficie della corteccia. Ora considereremo ciò che succede sotto la superficie osservando i risultati degli esperimenti nei quali un elettrodo misurava l'attività corticale.

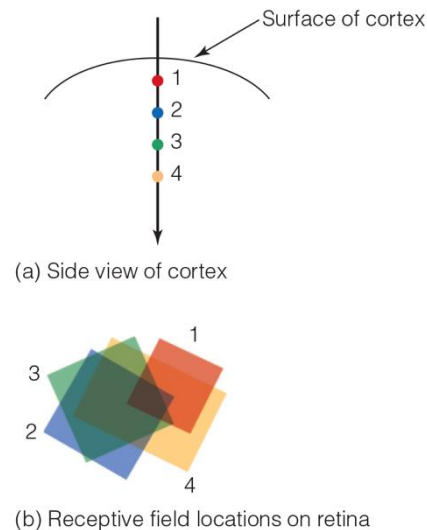


Figura 4.6 Colonne di posizione. Quando un elettrodo penetra perpendicolarmente la corteccia cerebrale i campi recettivi dei neuroni lungo il percorso si sovrappongono. Il campo recettivo registrato ad ogni posizione numerata del percorso è indicato dal quadrato enumerato corrispondente.

COLONNE DI POSIZIONE E DI

ORIENTAMENTO Hubel e Wiesel (1965) fecero

una serie di esperimenti nei quali registrarono l'attività dei neuroni corticali sottostanti agli elettrodi. Posizionando un elettrodo perpendicolarmente rispetto alla superficie della corteccia cerebrale di un gatto notarono che tutti i neuroni incontrati avevano il proprio campo recettivo più o meno nella stessa posizione retinica. I loro risultati sono mostrati nella **figura 4.6a**, la quale mostra quattro neuroni lungo il percorso dell'elettrodo, mentre nella **figura 4.6b** viene mostrata la parziale sovrapposizione dei campi recettivi dei quattro neuroni nella retina. Da questi risultati Hubel e Wiesel conclusero che la corteccia striata è organizzata in **colonne di posizione**, perpendicolari alla superficie corticale in modo che tutti i neuroni attorno ad una colonna di posizione abbiano i loro campi percettivi nella stessa posizione retinica.

Hubel e Wiesel notarono anche che i neuroni di una colonna di posizione preferivano stimoli con lo stesso orientamento. Perciò tutte le cellule incontrate lungo l'elettrodo A della **figura 4.7** si attivavano osservando linee orizzontali mentre quelle lungo l'elettrodo B osservando linee a 45°. Basandosi su questi risultati Hubel e Wiesel conclusero che la corteccia cerebrale è organizzata anche in **colonne di orientamento**;

ogni colonna contiene cellule che rispondono meglio a stimoli con particolari orientazioni. Hubel e Wiesel mostrarono anche che colonne di orientamento adiacenti hanno cellule con leggere differenze tra le preferenze di orientamento. Quando muovevano un elettrodo obliquamente attraverso la corteccia (ma non perpendicolarmente alla superficie) in modo da attraversare le colonne di orientamento trovarono che gli orientamenti favoriti dei neuroni cambiavano con un certo ordine, nello specifico una colonna di cellule che favorivano i 90° si trovava vicino ad una colonna che rispondeva agli 85° (**Figura 4.8**). Hubel e Wiesel scoprirono anche che se muovevano di 1 mm l'elettrodo passavano attraverso tutto il range di orientamenti delle colonne di orientamento, da ciò compresero che 1 mm era il diametro di una colonna di posizione.

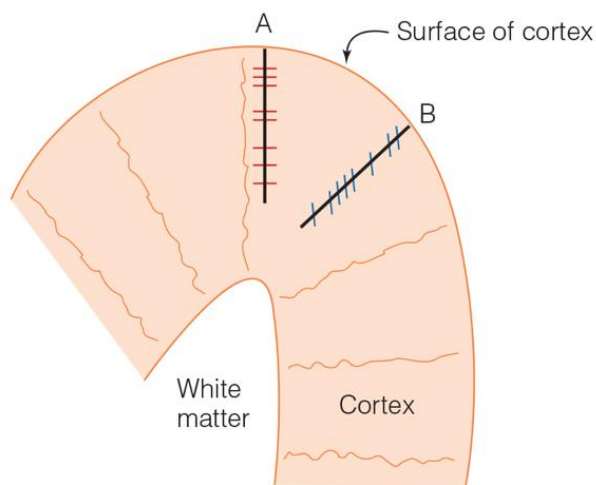


Figura 4.7 Colonne di orientamento. Tutti i neuroni corticali lungo il percorso A rispondono allo stimolo visivo delle barre orizzontali (indicate dalle linee rosse) mentre quelli lungo il percorso B rispondono alle barre a 45° .

UNA COLONNA DI POSIZIONE: MOLTE COLONNE DI ORIENTAMENTO La dimensione di 1mm per le colonne di posizione indica che una colonna di posizione è sufficientemente larga da contenere colonne di orientamento che comprendano tutte le possibili orientazioni. Perciò la colonna di posizione nella **Figura 4.9** è connessa ad una porzione della retina (in quanto tutti i neuroni nella colonna hanno i campi recettivi attorno allo stesso punto) e contiene neuroni che rispondono a tutti i possibili orientamenti di una figura.

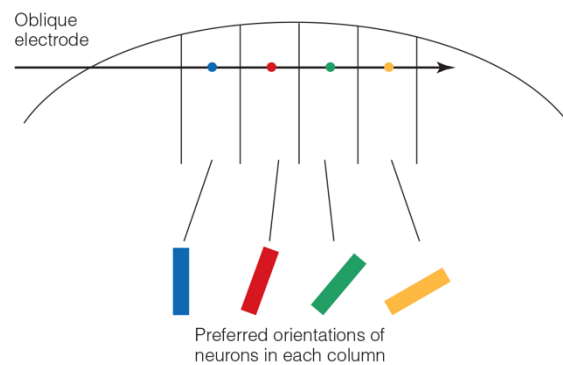


Figura 4.8 Se un elettrodo viene inserito obliquamente nella corteccia attraverserà una serie di colonne di orientamento. L'orientamento caratteristico di ogni neurone, indicato dalle barre colorate, cambia ordinatamente con il passaggio dell'elettrodo. La distanza percorsa in figura è ovviamente esagerata.

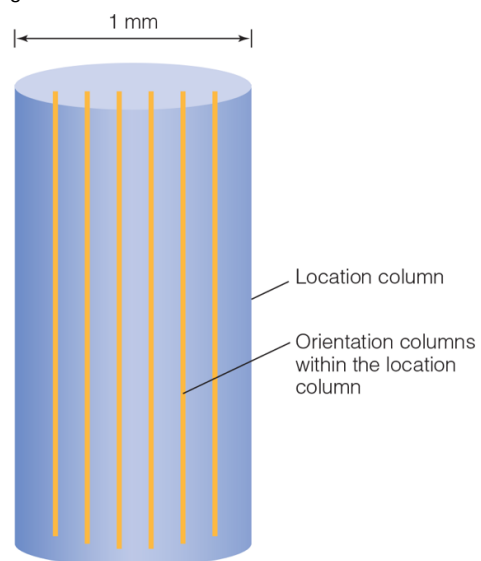


Figura 4.9 Una colonna di posizione contiene tutto un range di colonne di orientamento. Una colonna come questa veniva chiamata *ipercolonna* da Hubel e Wiesel e riceve informazioni su tutte gli orientamenti visivi in una piccola area retinica.

Pensa a cosa significa tutto ciò. I neuroni in quella colonna di posizione ricevono segnali da una particolare zona nella retina, la quale corrisponde ad una piccola porzione del campo visivo. Poiché questa colonna di posizione contiene alcuni neuroni che rispondono ai vari orientamenti, ogni oggetto orientato osservato dall'area retinica della colonna di posizione causerà una messa a fuoco attuata dai neuroni contenuti in essa. Una colonna di posizione con tutte le colonne di orientamento al suo interno, chiamata *ipercolonna* da Hubel e Wiesel, riceve informazioni riguardo a tutte le possibili orientazioni osservate da una piccola area retinica ed è ben programmata per elaborare le informazioni che arrivano da una porzione del campo visivo¹.

¹ Oltre alle colonne di posizione e di orientamento Hubel e Wiesel identificarono le **colonne di dominanza**

oculare. La maggior parte dei neuroni rispondono prevalentemente agli stimoli di un occhio rispetto a

Come risponde ad una scena visiva il rilevamento delle caratteristiche?

Determinare come i milioni di neuroni corticali rispondono ad una scena come quella presente nella **Figura 4.10a** è un progetto ambizioso. Semplificheremo il compito concentrandoci su una piccola parte della scena, il tronco dell'albero nella **Figura 4.10b**, in particolare nella parte tra i tre cerchi A, B e C.

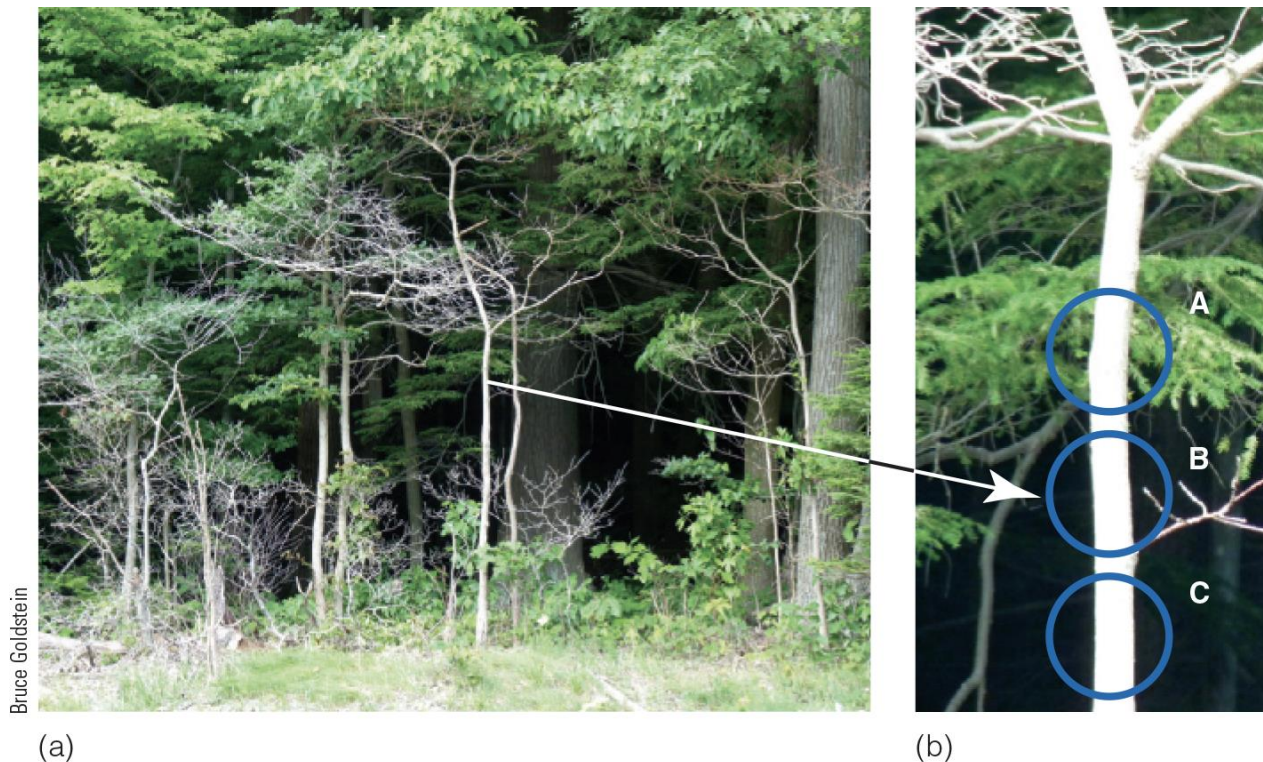


Figura 4.10 (a) una foto del Pennsylvania Woods. (b) le parti A, B e C del tronco rappresentano zone percepite dai campi recettivi di aree retiniche.

La **Figura 4.11a** mostra come l'immagine di quest'albero viene percepita dalla retina. Ognuno dei cerchi rappresenta l'area sottoposta ad una colonna di posizione. La **Figura 4.11b** mostra invece le colonne di posizione nella corteccia. Ricorda che ognuna delle colonne di posizione contiene un set completo di quelle di orientamento (Figura 4.9). Ciò significa che il tronco d'albero attiverà le colonne di orientamento a 90° in ogni colonna di posizione, come indicato dall'area arancione nelle colonne. Il tronco viene rappresentato dalla scarica dei neuroni delle colonne presenti nella corteccia. Nonostante possa sorprendere il fatto che venga rappresentato da colonne separate, ciò conferma una proprietà del sistema percettivo nominata poc'anzi: la rappresentazione corticale degli

stimoli non deve *assomigliare* allo stimolo ma contiene le informazioni che permettono di *rappresentarlo*. La rappresentazione dell'albero nella corteccia visiva viene scaricata dai neuroni nelle diverse colonne corticali e in certi punti della corteccia le informazioni separate verranno combinate per creare la percezione dell'albero.

quelli dell'altro secondo la loro **dominanza oculare**. I neuroni con la stessa dominanza sono organizzati in colonne di dominanza nella corteccia cerebrale. Ciò significa che ogni neuroni incontrato lungo il passaggio perpendicolare di un elettrodo risponde meglio

all'occhio destro o a quello sinistro. Ogni ipercolonna ha due colonne di dominanza oculare al suo interno, una per ogni occhio.

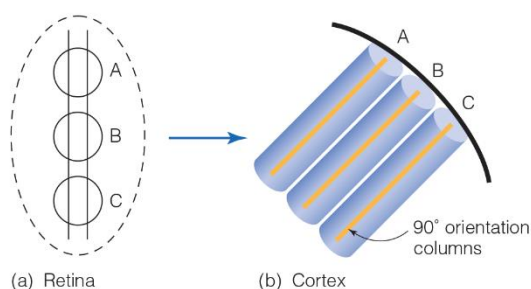


Figura 4.11 (a) campi recettivi delle tre sezioni del tronco della figura 4.10b; i neuroni associati ai campi si trovano in diverse colonne di posizione. (b) tre colonne di posizione corticali. I neuroni connessi all'orientamento del tronco sono nelle aree arancioni.

Prima di descrivere come gli oggetti vengono rappresentati dai rilevatori di caratteristiche torniamo alla nostra scena (**Figura 4.12**). Ogni cerchio o ellisse rappresenta un'area che invia informazioni da una colonna di posizione; queste colonne, lavorando assieme, coprono l'intero campo visivo con un effetto detto **tiling**. Proprio come una parete può essere ricoperta di piastrelle adiacenti, il campo visivo è composto da colonne di posizione vicine e spesso sovrapponibili (Nassi & Calway, 2009). (Ti suona familiare? Ricorda l'analogia sul calcio per il campo recettivo delle fibre nervose visive nel Capitolo 3, nella quale ogni spettatore osservava una piccola area del campo ed assieme lo mappavano nella sua interezza)

L'idea che ogni scena viene rappresentata attivamente da varie colonne di posizione significa che viene registrata dalla corteccia striata da una scarica neuronale incredibilmente complessa. Immagina il processo appena descritto per le tre piccole aree di un tronco moltiplicato centinaia di migliaia di volte. Questa rappresentazione è ovviamente il primo passo. Come vedremo ora, i segnali dalla corteccia striata viaggiano altri luoghi corticali per ulteriori processi.



Bruce Goldstein

Figura 4.12 I cerchi e le ellissi in sovrapposizione rappresentano ciascuno un'area che invia informazioni ad una colonna di posizione.

Vie dell'informazione visiva: il cosa, il dove e il come

Fino ad adesso, mentre osservavano nella corteccia i tipi di neuroni e come la corteccia è spazialmente organizzata in mappe e colonne, abbiamo descritto ricerche principalmente degli anni sessanta e settanta.

La maggior parte delle ricerche durante questo periodo erano concentrate sulla corteccia striata o sulle aree vicine alla stessa. Anche se alcuni pionieri hanno visto le funzioni visive al di fuori della corteccia striata (Gross e altri, 1969, 1972; vedi Capitolo 3, pagina 69), fu solo dopo gli anni ottanta che un gran numero di ricercatori cominciarono a investigare come la stimolazione visiva causi attività nell'area molto lontano dalla corteccia striata.

Una delle più influenti idee che uscirono da queste ricerche è che ci sono dei sentieri, o "flussi", che trasferiscono informazioni dalla corteccia striata ad altre aree del cervello. Questa idea fu introdotta nel 1982, quando Leslie Ungerleider e Mortimer Mishkin descrissero esperimenti che distinsero due flussi che servono differenti funzioni.

vie dell'informazione del cosa e dove

Ungerleider e Mishkin (1982) usarono una tecnica chiamata *Ablazione* (anche chiamata *lesioning*). "Ablazione" si riferisce alla distruzione o alla rimozione di tessuti nel sistema nervoso.

METODO

Ablazione del cervello

L'obiettivo degli esperimenti sull'ablazione del cervello è di determinare le funzioni di particolari aree del cervello. Per primo la capacità degli animali è determinata da test comportamentali. Molti esperimenti di ablazione hanno avuto come soggetto le scimmie per via della somiglianza del loro sistema visivo con quello dell'uomo e poiché le scimmie possono essere allenate in modo da permettere ai ricercatori di determinare capacità percettive come accuratezza, visione dei colori, percezione della profondità e degli oggetti.

Una volta che la percezione degli animali è stata misurata, una area particolare del cervello ha ricevuto l'ablazione (rimozione o distruzione), o per chirurgia o per iniezione chimica che distruggono i tessuti vicini al luogo dove viene

iniettata. Nel caso ideale una particolare area è stata rimossa e il resto del cervello rimane intatto. Dopo l'ablazione, la scimmia è riallenata per determinare quali capacità percettive rimangono e quali invece sono affette dalla ablazione.

Ungerleider e Mishkin presentarono scimmie con due compiti: (1) un problema di discriminazione degli oggetti e (2) un problema di discriminazione di punti di riferimento. Nel **problema di discriminazione degli oggetti**, ad una scimmia viene mostrato un oggetto, come un solido rettangolare, ed era dopo presentato con un compito a doppia scelta come l'esempio mostrato nella figura 4.13a, che comprendeva l'oggetto "target"

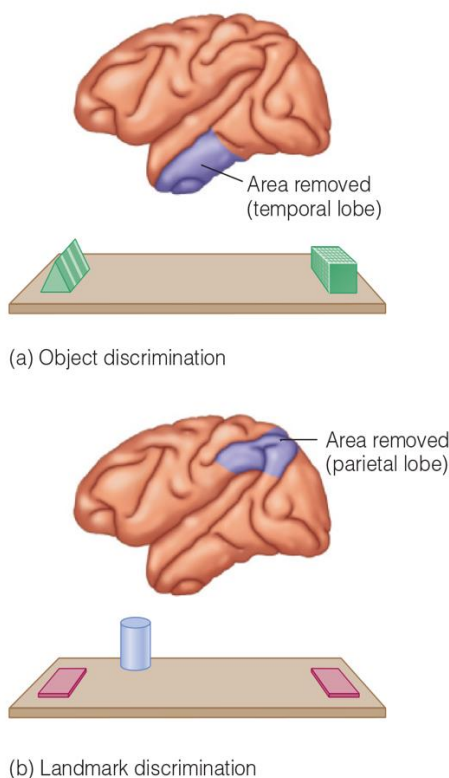


Figura 4.13 I due tipi di compiti di discriminazione usati da Ungerleider e Mishkin. (a) Discriminazione di oggetti: scegliere la figura corretta. Lesionare il lobo temporale (area offuscata) rende il compito difficile. (b) Discriminazione di punti di riferimento: prendere il cibo più vicino al cilindro. Lesionare il lobo parietale rende il compito difficile.

(il solido rettangolare) e un altro stimolo, come il solido triangolare. Se la scimmia spinge da parte l'oggetto target, riceve il cibo premio che era stato nascosto in una fessura sotto l'oggetto. Il **problema della discriminazione dei punti di riferimento** è mostrato nella **Figura 4.13b**. Qui, il compito della scimmia era di rimuovere il coperchio del contenitore del cibo che era più vicino al cilindro alto.

Nella parte dell'ablazione dell'esperimento, parte del lobo temporale è

rimosso in alcune scimmie. Dopo l'ablazione, dei test comportamentali mostrarono che il problema della discriminazione dell'oggetto era molto difficile per le scimmie con il loro lobo temporale rimosso. Questo risultato indica che il percorso che raggiunge il lobo temporale è responsabile per determinare l'identità dell'oggetto. Perciò Ungerleider e Mishkin chiamarono il percorso che collega la corteccia striata e il lobo temporale il **percorso del cosa** (**Figura 4.14**).

Altre scimmie con il lobo parietale rimosso, hanno difficoltà nel risolvere il problema della discriminazione dei punti di riferimento. Questo risultato indica che il percorso che porta al lobo parietale è responsabile per determinare la locazione dell'oggetto. Perciò Ungerleider e Mishkin chiamarono il percorso che unisce la corteccia striata e il lobo parietale il **percorso del dove** (**Figura 4.14**).

I percorsi del cosa e del dove sono anche chiamati il **percorso ventrale** (cosa) e il **percorso dorsale** (dove), perché la parte inferiore del cervello, dove il lobo temporale è situato, è la parte ventrale del cervello, e la parte superiore del cervello, dove il lobo parietale è localizzato, è la parte dorsale del cervello. Il termine dorsale si riferisce al retro e alla parte superiore della superficie di un organismo; così, la pinna dorsale di uno squalo o di un delfino è la pinna nella schiena che esce fuori dall'acqua. **Figura 4.15** mostra che per i bipedi, animali camminanti come l'uomo, la parte dorsale del cervello si trova nella sommità del cervello. (Immagina una persona con una pinna dorsale che esce dalla sommità del suo capo!) *Ventrale* è l'opposto del *dorsale*, anzi si riferisce alla parte inferiore del cervello.

La scoperta dei due percorsi nella corteccia- una per identificare oggetti (cosa) e una per locare oggetti (dove)-

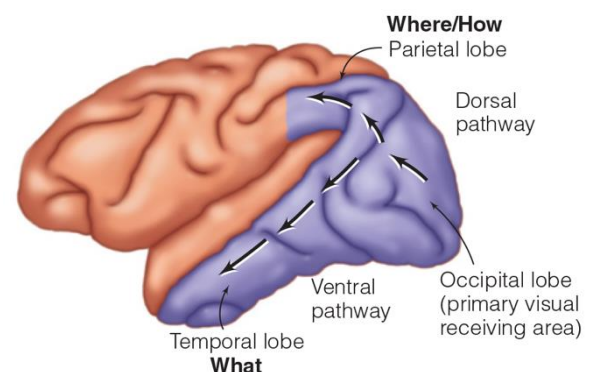


Figura 4.14 La corteccia delle scimmie, mostra il percorso del cosa, o ventrale, dal lobo occipitale al lobo temporale, e il percorso del dove, o dorsale, dal lobo occipitale al lobo parietale. Il percorso del dove è anche chiamato percorso del come.

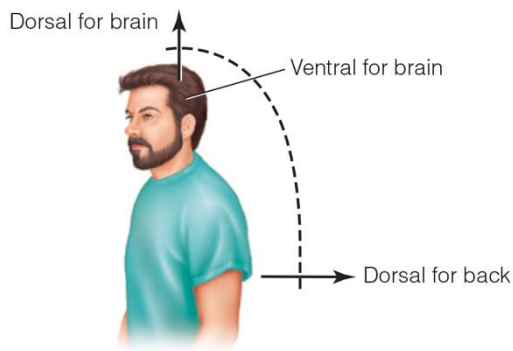


Figura 4.15 Dorsale si riferisce al retro di un organismo. In un animale eretto come l'uomo, dorsale si riferisce al retro del corpo e al retro della testa, come indicato dalle frecce e dalle curve tratteggiate. Ventrale è l'opposto del dorsale.

ha portato qualche ricercatore a rivedere la retina e il corpo genicolato laterale (LGN). Usando entrambe tecniche di registrazione di neuroni e di ablazione, hanno scoperto che le proprietà di sequenze ventrali e dorsali sono stabilite da due differenti tipi di cellule del ganglio nella retina, che trasmette segnali a differenti strati del LGN. Così, i flussi corticali ventrali e dorsali possono infatti essere fatti risalire alla retina e al LGN.

Anzitutto c'è una buona evidenza che il percorso ventrale e dorsale serve differenti funzioni, è importante notare che (1) i percorsi non sono totalmente separati, ma hanno connessioni tra di loro; e (2) il flusso di segnali non solo "sale" attraverso i lobi parietali e temporali, ma "scende" anche (Merigan & Maunsell, 1993; Ungerleider & Haxby, 1994). Ha senso che sia una comunicazione tra i percorsi perché nel nostro comportamento di tutti i giorni abbiamo bisogno sia di identificare che localizzare oggetti, e coordiniamo automaticamente queste due attività ogni volta che identifichiamo qualcosa (per esempio, una penna) e compiamo un'azione con essa di qualsiasi genere (prendere su una penna e scrivere con essa). Quindi, ci sono due differenti percorsi, ma qualche informazione è condivisa tra di loro. Il flusso "al contrario" di informazioni, chiamato feedback (vedi pagina 64), fornisce informazioni dai centri più alti che può influenzare i segnali che scorrono nel sistema. Questo feedback è uno dei meccanismi dietro al processo top-down, introdotto nel Capitolo 1 (pagina 10).

Flussi di informazioni riguardo al cosa e al come

Anche se l'idea di flussi ventrali e dorsali è stato generalmente accettata, David Milner e Melvyn Goodale (1995; vedi anche Goodale & Humphrey, 1998, 2001) hanno suggerito che il flusso dorsale fa molto di più che indicare solo dove sia un

oggetto. Milner e Goodale propongono che il flusso dorsale serva a compiere azioni, come prendere un oggetto. Compiere questa azione implicherebbe conoscere la locazione dell'oggetto, compatibilmente con l'idea del *dove*, ma va oltre il *dove* implicando un'interazione fisica con l'oggetto. Quindi, arrivare a prendere la penna implica le informazioni riguardanti alla locazione della penna *in più* a quelle del movimento della mano verso la penna. Secondo questa idea, il flusso dorsale provvede informazioni riguardo al *come* dirigere l'azione riguardo allo stimolo.

Evidenze che supportano l'idea che il flusso dorsale è implicato nel come dirigere l'azione è provvenuto dalla scoperta di neuroni nella corteccia parietale che risponde (1) quando una scimmia guarda un oggetto e (2) quando tende verso un oggetto (Sakata e altri, 1992; vedi anche Taira e altri, 1990). Ma la più drammatica evidenza che supporta l'idea di un flusso dorsale del *come* o *azione* viene dalla **neuropsicologia** - lo studio degli effetti comportamentali a seguito di danni al cervello negli umani.

METODO

Doppie dissociazioni in neuropsicologia

Uno dei principi fondamentali della neuropsicologia è che possiamo comprendere gli effetti dei danni cerebrali determinando **doppie associazioni**, che coinvolgono due persone: in una persona, i danni ad un'area del cervello provocano l'assenza della funzione A mentre la funzione B è presente; nell'altra persona, i danni ad un'altra area cerebrale causano l'assenza della funzione B mentre è presente la funzione A. Le scimmie di Ungerleider e Mishkin forniscono un esempio di doppia dissociazione. La scimmia con danni al lobo temporale era incapace di discriminare gli oggetti (funzione A, ma possedeva la capacità di risolvere il compito di discriminazione di punti di riferimento (funzione B). La scimmia con danni al lobo parietale era incapace di risolvere il compito di discriminazione di punti di riferimento (funzione B) ma era capace di discriminare gli oggetti (funzione A). Queste due conclusioni, prese assieme, sono un esempio di una doppia dissociazione. Il fatto che la discriminazione di oggetti e il compito di discriminazione di punti di riferimento possono essere interrotti separatamente e in modi opposti significa che queste due funzioni operano indipendentemente l'una dall'altra. Un esempio di doppia dissociazione negli esseri umani è fornito da due ipotetici pazienti. Alice, che ha subito danni al suo lobo temporale, ha difficoltà a nominare gli oggetti, ma non ha difficoltà ad indicare dove si trovano (tabella 4.1a). Bert, che

ha danni al lobo parietale, ha il problema opposto – è in grado di identificare gli oggetti ma non può dire esattamente dove si trovano (tabella 4.1b). I casi di Alice e Bert, presi assieme, rappresentano una doppia dissociazione e ci permettono di concludere che riconoscere gli oggetti e localizzarli funzionano indipendentemente l'una dall'altra.

Il comportamento del paziente D.F. Il metodo di determinazione della dissociazione è stato utilizzato da Milner e Goodale (1995) per studiare D.F., una donna di 34 anni che ha subito danni al suo percorso ventrale per un avvelenamento da monossido di carbonio causato da una perdita di gas in casa sua. Un risultato del danno cerebrale era che D.F. non era capace di abbinare l'orientamento di una carta tenuta nella sua mano a diversi orientamenti di una fessura. Questo è mostrato nel cerchio a sinistra nella **figura 4.16a**. Ogni linea nel cerchio indica l'orientamento a cui D.F. ha regolato la carta. La prestazione di orientamento perfetta sarebbe indicata da una linea verticale per ogni prova, ma le risposte di D.F. sono in larga parte sparpagliate. Il cerchio a destra mostra le performance accurate dei controlli normali.

Poiché D.F. aveva difficoltà ad orientare una carta per abbinarla all'orientamento della fessura, sembrerebbe ragionevole ritenere che il paziente avesse anche difficoltà a *disporre* la carta attraverso la fessura, perché per fare questo avrebbe dovuto girare la carta in modo che fosse allineata alla fessura. Ma quando è stato chiesto a D.F. di "imbucare" attraverso la fessura, poteva farlo! Anche se D.F. non poteva girare la carta per abbinare l'orientamento della fessura, una volta che iniziava a muovere la carta verso la fessura, era in grado di ruotarla per appaiare l'orientamento di essa (**figura 4.16b**). Pertanto, D.F. ha eseguito malamente il compito di *static orientation-matching* ma lo ha eseguito bene appena l'*azione* è stata coinvolta (Murphy et al., 1996). Milner e Goodale interpretarono il comportamento di D.F. per mostrare che esiste un meccanismo per determinare l'orientamento e un altro per coordinare la visione e l'azione. Questi risultati per D.F. dimostrano una doppia dissociazione rispetto ad altri pazienti i cui sintomi sono l'opposto di quelli di D.F., e persone così, infatti, esistono. Queste persone possono stabilire un'orientazione visiva, ma non possono raggiungere il compito che combina la visione e l'azione. Come ci si aspetterebbe, mentre il flusso ventrale di D.F. è danneggiato, le altre persone hanno danneggiati i loro flussi dorsali. Basandosi su questi risultati, Milner e Goodale proposero che il percorso ventrale dovrebbe ancora essere chiamato il *percorso del cosa*,

come suggerirono Ungerleider e Mishkin, ma che una migliore descrizione del percorso dorsale sarebbe stata *il percorso del come o il percorso dell'azione*, perché stabilisce *come* una persona svolge un'*azione*. Come a volte accade nella scienza, non tutti utilizzano gli stessi termini. Pertanto, alcuni ricercatori chiamano il percorso dorsale il *percorso del dove*, e alcuni lo chiamano il *percorso del come o dell'azione*.

Tabella 4.1 Una doppia dissociazione

	NOMINAR E GLI OGGETTI	DETERMINAR E LA POSIZIONE DEGLI OGGETTI
(a) ALICE: Danno al lobo tempora le (flusso ventrale)	NO	Sì
(b) BERT: Danno al lobo parietale (flusso dorsale)	Sì	NO

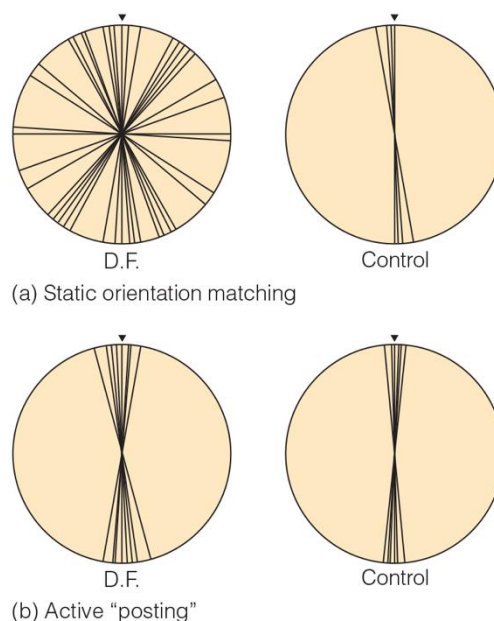


Figura 4.16 Prestazione di D.F. e di una persona senza danno cerebrale in due compiti: (a) stabilire l'orientamento di una fessura; e (b) posizionare una carta attraverso la fessura. Vedi il testo per i dettagli.

Il comportamento delle persone senza danni cerebrali Nel nostro comportamento normale di tutti i giorni, non siamo consapevoli di due flussi di elaborazione visiva, uno per il *cosa* e uno per il *come*, poiché essi lavorano senza

interruzioni mentre percepiamo gli oggetti e agiamo verso di essi. Casi come quello di D.F., in cui un flusso è danneggiato, rivelano l'esistenza di questi due flussi. Ma cosa dire riguardo alle persone senza il cervello danneggiato? Esperimenti psicofisici che misurano come le persone percepiscono e reagiscono alle illusioni visive hanno dimostrato la dissociazione tra percezione e azione che era evidente nel caso di D.F.

La **figura 4.17a** mostra lo stimolo usato da Tzvi Ganel e collaboratori (2008) in un esperimento progettato per dimostrare una separazione della percezione e dell'azione in soggetti non danneggiati al cervello. Lo stimolo crea un'illusione visiva: la linea 1 è effettivamente più lunga della linea 2 (vedi **figura 4.17b**), ma la linea 2 *sembra* più lunga.

Ganel e collaboratori presentarono i soggetti con due compiti: (1) un compito di *stima della lunghezza* in cui è stato chiesto di indicare come hanno percepito la lunghezza delle linee estendendo il loro dito pollice ed indice, come mostrato in **figura 4.17c**; e (2) un compito di

prensione in cui è stato chiesto di allungare il braccio verso le linee e afferrare ogni linea dalle sue estremità. I sensori sulle dita dei soggetti misuravano la distanza tra le dita mentre i soggetti afferravano le linee. Questi due compiti sono stati scelti poiché dipendono da diversi flussi di elaborazione. Il compito di stima della lunghezza coinvolge il flusso ventrale o *il flusso del cosa*. Il compito di prensione coinvolge il flusso dorsale o *flusso del come*.

I risultati di questo esperimento, mostrati nella **figura 4.17d**, indicano che nel compito di stima della lunghezza, i soggetti ritenevano che la linea 1 (la linea più lunga) sembrava più corta della linea 2, ma nel compito di prensione, distanziavano più lontano le loro dita per la linea 1. Pertanto, l'illusione funziona per la percezione (il compito di stima della lunghezza), ma non per l'azione (il compito di prensione). Questi risultati sostengono l'idea che la percezione e l'azione sono servite da meccanismi diversi. Un'idea che ha avuto origine con osservazioni di pazienti con danno cerebrale è supportata dalle prestazioni di osservatori senza danno cerebrale.

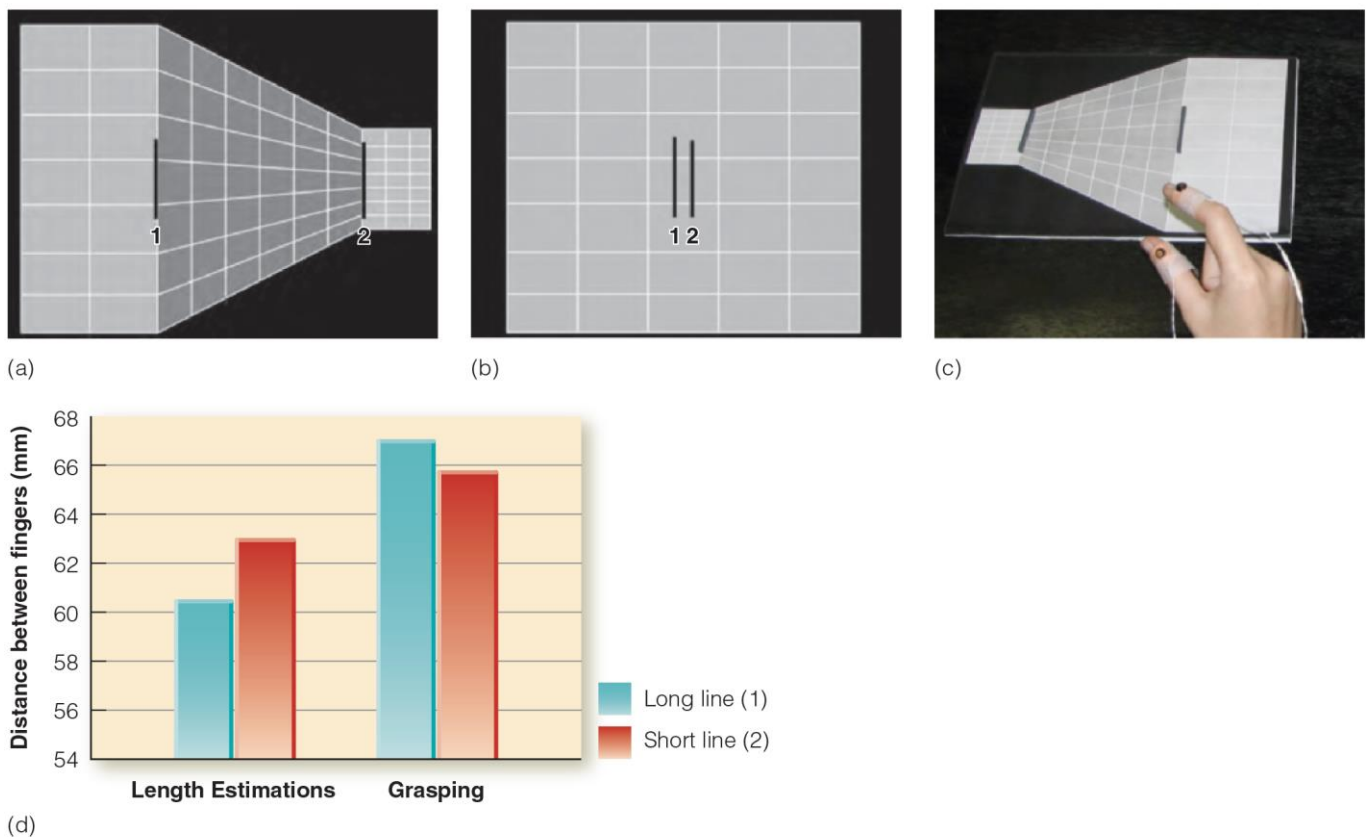


Figura 4.17 (a) L'illusione delle dimensioni utilizzata da Ganel e collaboratori (2008) in cui la linea 2 appare più lunga della linea 1. I numeri non erano presenti nel display visto dai soggetti. (b) Le due linee verticali da (a), mostrando che la linea 2 è effettivamente più corta della linea 1. (c) I soggetti nell'esperimento hanno regolato lo spazio tra le loro dita o per stimare la lunghezza delle linee (compito di stima della

lunghezza), o per allungare il braccio verso esse per afferrarle (compito di presa). La distanza tra le dita è misurata dai sensori sulle dita. (d) I risultati della stima della lunghezza e dei compiti di prensione nell'esperimento di Ganel et al. Il compito di stima della lunghezza indica l'illusione, perché la linea più corta (linea 2) erano ritenute essere più lunghe. Nel compito di prensione, i soggetti separavano di più le dita per la linea più lunga (linea 1), che era coerente con le lunghezze fisiche delle linee.

Modularità: strutture per facce, posti e corpi

Abbiamo visto come lo studio del sistema visivo è progredito dalla scoperta di Hubel e Wiesel dei neuroni nella corteccia striata e nelle aree vicine che rispondono a barre orientate ai primi pionieristici esperimenti di Charles Gross e collaboratori (1969, 1972), descritti nel capitolo 3 (vedi pagine 69), che mostrano che i neuroni della corteccia infero temporale delle scimmie rispondono a stimoli complessi come ritagli di mani e immagini di facce.

Sebbene c'era un ritardo tra la pubblicazione del lavoro di Gross e quando altri ricercatori iniziarono a trovare neuroni che rispondevano a stimoli complessi, una volta che incominciarono, le barriere si aprirono e uno studio di ricerca dopo l'altro descrissero neuroni che rispondevano a stimoli complessi. Per esempio, Keiji Tanaka e i suoi collaboratori (Ito et al., 1995; Kobatake & Tanaka, 1994; Tanaka, 1993; Tanaka et al., 1991) registrano da celle nella corteccia temporale che rispondevano al meglio a stimoli complessi, come mostra la figura 4.18a. Questa cella, che risponde al meglio al disco circolare con barra sottile,

risponde male alla barra da sola (figura 4.18b) o al disco da solo (figura 4.18c). La cella risponde alla forma quadrata con la barra (figura 4.18d), ma non come il cerchio e la barra. Oltre alla scoperta di neuroni che rispondono a stimoli complessi, i ricercatori fondarono anche una prova che i neuroni che rispondono a stimoli simili sono spesso raggruppati insieme in un'area cerebrale. La struttura che è specializzata nel processo di informazione rispetto a un tipo particolare di stimoli è chiamato modulo. Ci sono tante prove che dimostrano che aree specifiche del lobo temporale rispondono meglio a particolari tipi di stimoli.

Neuroni facciali nella corteccia infero-temporale della scimmia

Edmund Rolls e Martin Tovee (1995) misurarono la risposta dei neuroni nella corteccia infero temporale della scimmia (vedi Figura 3.35a). Quando presentavano immagini di stimoli facciali e di stimoli non facciali (soprattutto paesaggi e cibo), trovarono molti neuroni che rispondevano meglio alla facce. La **figura 4.19** mostra i risultati di un neurone che risponde rapidamente alle facce ma difficilmente a tutti gli altri tipi di stimoli.

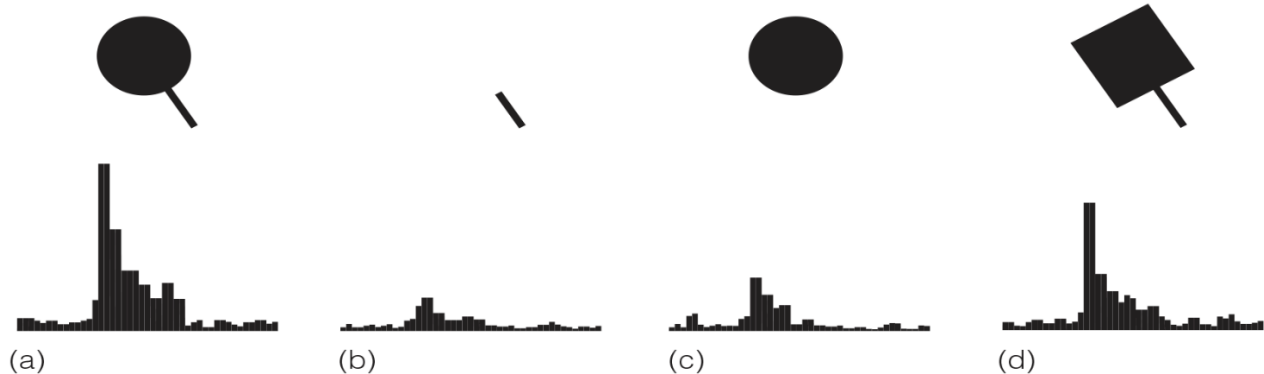


Figura 4.18 come un neurone nel lobo temporale di una scimmia risponde ad alcuni stimoli. Questo neurone risponde meglio al disco circolare con una barra fina (a).

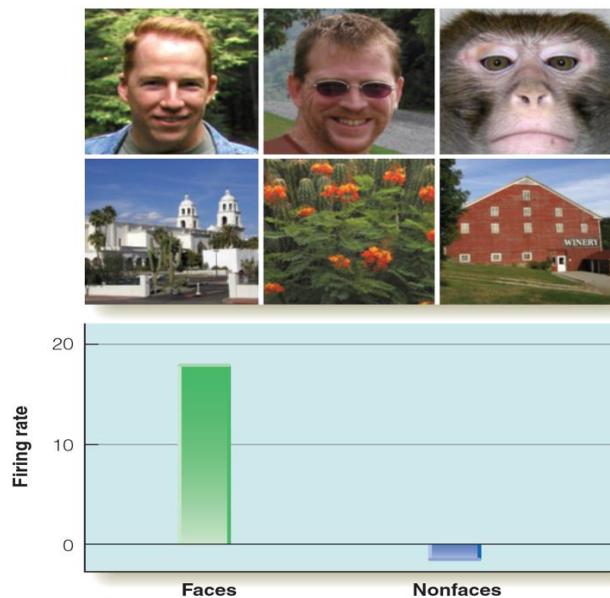


Figura 4.19 la portata della risposta di un neurone nella corteccia IT della scimmia che risponde a stimoli facciali ma non a stimoli non facciali. (Basato sui dati da Rolls & Tovee, 1995.)

La cosa particolarmente significativa dei “neuroni facciali” è che ci sono aree del lobo temporale della scimmia che sono particolarmente ricchi in questi neuroni. Doris Tsao e collaboratori (2006) presentarono 96 immagini di facce, corpi, frutti, arnesi, mani, e mischiarono i pattern a due scimmie mentre registravano i neuroni corticali all’interno di questa area facciale. Classificarono i neuroni come “selettivo di faccia” se rispondevano più intensamente alle facce nello stesso modo delle non facce. Usando questo criterio, hanno trovato che 97 per cento delle celle erano “selettive di faccia”. L’alto livello di selettività delle facce all’interno di questa area è illustrato nella **figura 4.20**, che mostra la risposta media per entrambe le scimmie per ognuno dei 96 oggetti. La risposta alle 16 facce, alla sinistra, è molto più grande della risposta di qualunque altro oggetto.

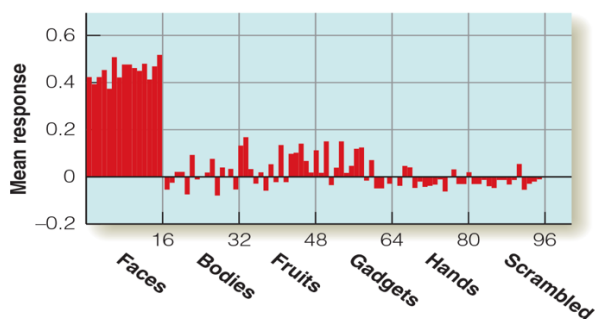


Figura 4.20 risultati dell’esperimento di Tsao et al. (2006) in cui l’attività dei neuroni nel lobo temporale di una scimmia era registrata in risposta a stimoli facciali, di altri oggetti e mescolati.

Potreste chiedervi come mai ci sono dei neuroni che rispondono meglio a stimoli complessi come ad esempio le facce. Abbiamo visto come il

processo neurale che implica il meccanismo della convergenza, dell’eccitazione e dell’inibizione possono creare neuroni che rispondono meglio a piccole chiazze di luce (vedi **figura 3.21**, pagina 61). Gli stessi meccanismi sono presumibilmente implicati nella creazione di neuroni che rispondono a stimoli più complessi. Naturalmente, i circuiti neurali implicati nella creazione di neuroni “rivelatori di facce” devono essere estremamente complessi. Comunque, il potenziale per questa complessità c’è, perché ci sono cento miliardi (10^{11}) neuroni e ogni neurone nella corteccia riceve input da una media di 1000 altri neuroni. Basandosi su questi numeri, è stato stimato che ci sono diverse centinaia di miliardi di connessioni sinaptiche tra i neuroni nel cervello (Marois & Ivanoff, 2005). Quando consideriamo la vasta complessità delle interconnessioni neurali che devono essere implicate nella creazione del neurone che risponde meglio alle facce, è semplice essere d’accordo con la descrizione del cervello fatta da William James, professore di psicologia alla Harvard e autore di uno dei primi libri di testo di psicologia (James, 1890/1981), come “la cosa più misteriosa dell’universo”.

Area cerebrali per il riconoscimento di facce, posti e corpi

Il brain imaging (vedi Metodi, pagina 79) è stato usato per identificare aree del cervello umano che contengono neuroni che rispondono meglio alle facce, e altri che rispondono meglio alle immagini di scene e corpi umani. In uno di questi esperimenti, Nancy Kanwisher e collaboratori (1997) usarono fMRI per determinare le attività cerebrali in risposta a immagini di facce e altri oggetti, come ad esempio facce mescolate/disordinate, oggetti familiari, case e mani. Quando tolsero la risposta agli altri oggetti dalla risposta alle facce, Kanwisher e collaboratori trovarono che l’attività rimaneva in un’area che chiamarono l’**area della faccia fusiforme (FFA)**, che è collocata nella circonvoluzione fusiforme nella parte inferiore del cervello direttamente al di sotto della corteccia IT (vedi figura 3.35b). Quest’area è all’incirca equivalente alle aree facciali nella corteccia temporale della scimmia. I risultati di Kanwisher, uniti a molti altri esperimenti, hanno mostrato che la FFA è specializzata a rispondere alle facce (Kanwisher, 2010). Un’altra evidenza dell’area specializzata per la percezione delle facce è che il danneggiamento il lobo temporale causa *prosopagnosia* – difficoltà a riconoscere le facce di persone familiari. Anche i volti molto familiari sono interessati, quindi

persone con prosopagnosia possono non essere in grado di riconoscere amici stretti o membri familiari – o anche il loro riflesso nello specchio – anche se possono facilmente identificare, per esempio, persone non appena li sentono parlare (Burton et al., 1991; Hecaen & Angelerques, 1962; Parkin, 1996).

In aggiunta alla FFA, che contiene neuroni che sono attivati dalle facce, due altre aree specializzate nella corteccia temporale sono state identificate. L'**area del para ippocampo (PPA)** è attivata da immagini raffiguranti scene all'interno e all'esterno come è mostrato nella **Figura 4.21a** (Aguirre et al., 1998; Epstein & Kanwisher, 1998). Apparentemente ciò che è importante per quest'area è l'informazione rispetto alla disposizione spaziale perché si verifica un'aumentata attivazione sia nelle stanze vuote sia nelle stanze che sono completamente arredate (Kanwisher, 2003). Un'altra area specializzata, l'**area del corpo extra striata (EBA)**, è attivata da immagini di corpi e parti di corpi (ma non da facce), come è mostrata dalla **Figura 4.21b** (Downing et al., 2001).

L'esistenza dei neuroni che sono specializzati a

rispondere alle facce, ai posti e ai corpi ci fa capire che siamo capaci di spiegare come la percezione sia basata sull'accensione/attivazione di neuroni. È probabile che la nostra percezione di facce, luoghi storici e di corpi di persone dipende specialmente da neuroni messi a punto in delle aree come per esempio la FFA, PPA e EBA.

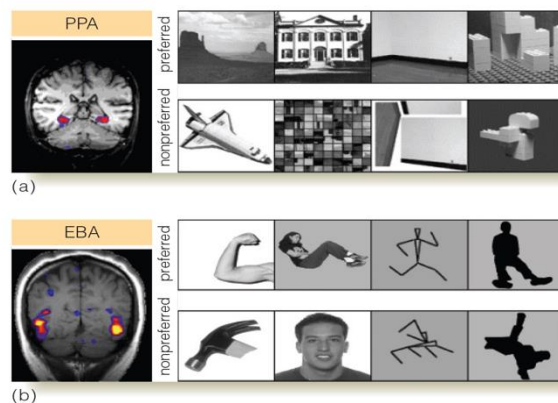


Figura 4.21 (a) l'area para ippocampale (PPA) è attivata da posti (riga superiore) ma non da altri stimoli (riga inferiore). (b) l'area del corpo extra striata (EBA) è attivata da corpi (riga superiore), ma non da altri stimoli (riga inferiore).

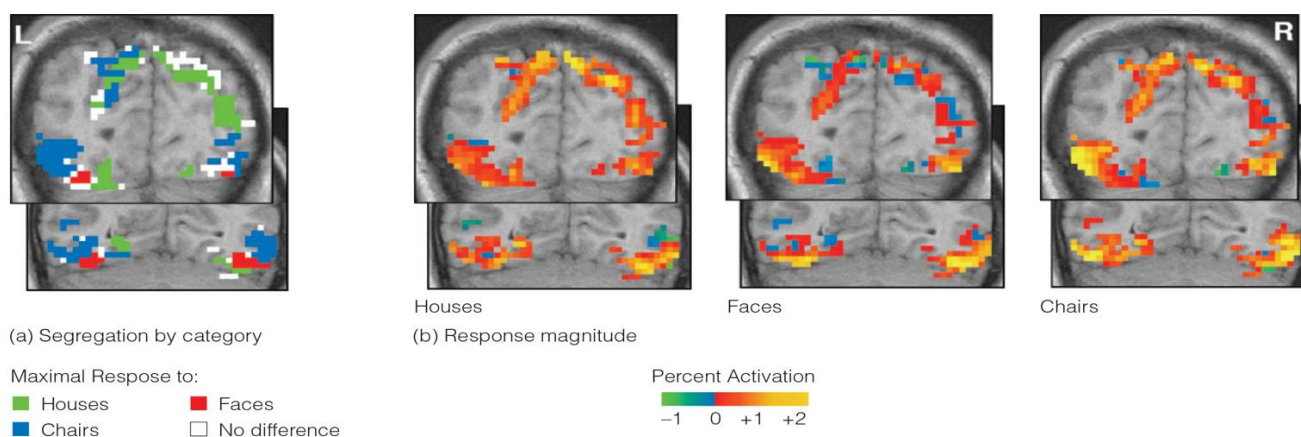


Figura 4.22 la fMRI risponde al cervello umano con vari tipi di stimoli: (a) le aree che erano più fortemente attivate da case, facce e sedie; (b) tutte le aree attivate da ogni tipo di stimoli.

Ma è anche importante ricordare che anche se stimoli come facce e costruzioni attivano specifiche aree del cervello, questi stimoli attivano anche altre aree del cervello nello stesso modo. Questo è illustrato nella **figura 4.22**, che mostra i risultati del esperimento di fMRI sugli uomini. La figura 4.22a mostra che immagini di case, facce e sedie causano un'attivazione estrema di tre aree separate nella corteccia IT. Comunque ogni tipo di stimolo causa anche un'attività sostanziale all'interno di altre aree, come è mostrato nelle tre figure limitate da solo queste aree (Ishai et al., 1999, 2000). Quindi oggetti come ad esempio facce dovrebbero causare un grande focus di

attività in un'area specializzata per le facce, come ad esempio la FFA, ma causano anche altre attività che sono distribuite sopra una vasta area della corteccia (Cohen & Tong, 2001; Riesenhuber & Poggio, 2000, 2002).

Possiamo riassumere quello che sappiamo riguardo l'organizzazione nel sistema visivo annotando che il sistema visivo è organizzato sia spazialmente sia funzionalmente. La mappa spaziale è retinotopica, cioè che punta alla LGN o alla corteccia corrispondente a specifici punti sulla retina o nella scena. Ma l'organizzazione spaziale diventa più debole quando ci muoviamo verso le aree corticali più alte, perché in aree come ad esempio la corteccia IT, i neuroni hanno campi recettoriali molto grandi che si estendono sulle grandi aree della retina e del campo visivo. Molti dei neuroni facciali rispondono quando la faccia è raffigurata nella fovea, che ha senso perché quando vogliamo identificare una faccia di solito la

guardiamo.

Il sistema visivo è organizzato *funzionalmente* con diversi flussi di dati per *cosa* e *dove/come* e con specifiche aree corticali che sono ricche di neuroni che rispondono a specifici tipi di stimoli come ad esempio facce, posti e corpi. Non è una coincidenza che stimoli che hanno delle specifiche aree nel cervello sono le uniche che vediamo ogni momento (facce e corpi) e che sono importanti per aiutarci a trovare una nostra via attraverso il contesto ambientale (neuroni dei posti).

Finora dalla nostra storia, dovrebbe essere allettante dire che la corteccia IT, con i suoi neuroni per le facce e gli altri oggetti complessi, è “la fine della linea” per l’elaborazione di informazioni riguardo oggetti. Come adesso vedremo, questa conclusione dovrebbe essere parzialmente corretta, ma dei segnali dalla corteccia IT continuano anche su altre strutture che dovrebbero essere implicate negli oggetti per il riconoscimento.

QUALCOSA DA CONSIDERARE:

Dove la visione incontra la memoria

Alcuni dei segnali che lasciano la corteccia inferotemporale raggiungono strutture nel lobo temporale mediale, come la corteccia paraippocampale, la corteccia entorinale, e l’ippocampo (**Figura 4.23a**). Queste strutture del lobo temporale mediale sono estremamente importanti per la memoria. La classica dimostrazione dell’importanza di una delle strutture nel lobo temporale mediale, l’ippocampo, è il caso di H.M., al quale è stato rimosso l’ippocampo da entrambi i lati del cervello nel tentativo di eliminare le crisi epilettiche che non avevano risposto ad altri trattamenti (Scoville & Milner, 1957).

L’operazione eliminò le crisi di H.M., ma eliminò inoltre la sua abilità di archiviare esperienze nella sua memoria. Perciò, quando H.M. faceva esperienza di qualcosa, come, ad esempio, la visita del suo dottore, non era in grado di ricordare quell’esperienza, quindi, alla successiva apparizione del dottore, H.M. non aveva memoria di averlo visto. La sfortunata situazione di H.M. si è verificata perché nel 1953 i chirurghi non avevano realizzato che l’ippocampo è fondamentale per la formazione di memoria a lungo termine. Una volta realizzati i devastanti effetti derivati dalla rimozione dell’ippocampo da entrambi i lati del cervello, l’operazione di H.M. non fu mai più ripetuta.

La connessione tra ippocampo e visione fu dimostrata nell’esperimento di R. Quian Quiroga e

collaboratori (2005, 2008) che abbiamo introdotto nel Capitolo 3 (Pagina 72). Questi esperimenti mostrano che sono i neuroni nell’ippocampo che rispondono a facce di specifiche persone, come Steve Carell (vedi Figura 3.20), e a specifiche strutture, come la Torre Eiffel o il Teatro dell’opera di Sydney. Vediamo ora questi esperimenti più nel dettaglio.

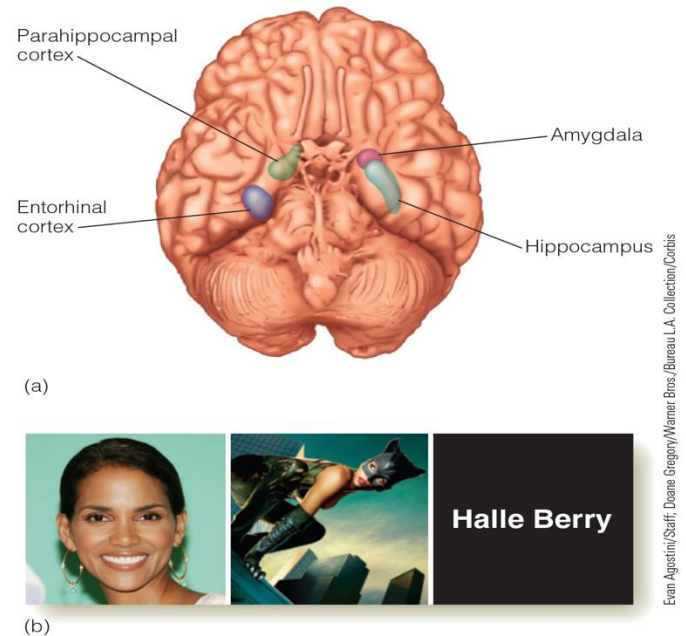


Figura 4.23 (a) Posizione dell’ippocampo e di altre strutture che sono state studiate da Quiroga e collaboratori (2005). (b) Alcuni degli stimoli che causano l’attivazione di un neurone dell’ippocampo.

Quiroga ha registrato otto pazienti affetti da epilessia a cui, in preparazione alla chirurgia, erano stati impiantati elettrodi nel loro ippocampo o in altre aree del lobo temporale mediale per localizzare con precisione dove erano originati i loro attacchi. I pazienti hanno visto un numero di visualizzazioni diverse di individui e oggetti specifici oltre a immagini di altre cose, come visi, edifici e animali. Non sorprendentemente, un certo numero di neuroni ha risposto ad alcuni di questi stimoli. Ciò che ha sorpreso, tuttavia, era che alcuni neuroni rispondevano a una serie di punti di vista diversi su una sola persona o edificio o su un numero di modi di rappresentare quella persona o edificio. Per esempio, un neurone rispondeva a tutte le foto dell’attrice Jennifer Aniston, ma non rispondeva a facce di altre persone famose e non famose, punti di riferimento, animali ed altri oggetti. Come abbiamo notato nel Capitolo 3, un altro neurone rispondeva alle immagini dell’attore Steve Carell. Ancora un altro neurone rispondeva alle fotografie di Halle Berry, ai suoi disegni, alle foto di lei vestita da Catwoman in “Batman”, e anche alle parole “Halle Berry” (**Figura 4.23b**).

Secondo Quiroga, il lobo temporale mediale – e specialmente l’ippocampo – non è responsabile

del riconoscimento degli oggetti. Il paziente H.M., per esempio, che non aveva l'ippocampo, poteva comunque riconoscere gli oggetti. Solo che non dopo li ricordava. Perciò, solo perché un neurone dell'ippocampo risponde a uno stimolo visivo non significa che sia responsabile del vedere. Ciò di cui è responsabile è il ricordare.

Il possibile ruolo di questi neuroni nella memoria è supportato dal modo in cui rispondono a diversificate visioni dello stimolo, a diversi modi di rappresentazione e persino a parole che concettualizzano lo stimolo. Questi neuroni non rispondono alle caratteristiche visive delle immagini, ma ai concetti - "Jennifer Aniston", "Halle Berry", "Teatro dell'opera di Sydney" - che rappresentano gli stimoli. Quindi, il fatto che il neurone che ha risposto a Jennifer Aniston abbia anche risposto a Lisa Kudrow non è stata una coincidenza, perché entrambe sono apparse nella serie TV di Friends. La risposta di questi neuroni del lobo temporale mediale agli stimoli visivi sembra dipendere, quindi, dalle esperienze passate di una particolare persona. Quindi, un appassionato di calcio potrebbe concepirne avere un neurone che risponde vedendo una foto di Tom Brady, e questo stesso neurone potrebbe anche rispondere ad Aaron Rogers.

Il legame tra questi neuroni del lobo temporale mediale, che rispondono a stimoli visivi, e la memoria ha ricevuto ulteriore supporto dai risultati di un esperimento di Hagan Gelbard-Sagiv e collaboratori (2008). Questi ricercatori hanno fatto visualizzare a pazienti con epilessia una serie di video-clip da 5 a 10 secondi, un certo numero di volte mentre venivano monitorati i neuroni nel lobo temporale mediale. Le clip mostravano

personaggi famosi, punti di riferimento e personaggi non famosi con animali impegnati in varie azioni. Mentre la persona stava visualizzando le clip, alcuni neuroni rispondevano meglio a determinate clip. Ad esempio, un neurone in uno dei pazienti ha risposto meglio ad una clip del programma TV "I Simpson".

L'attivazione di video-clip specifici è simile a quella che Quiroga ha trovato per la visualizzazione di immagini fisse. Tuttavia, questo esperimento ha fatto un ulteriore passo in avanti facendo in modo che gli sperimentatori chiedessero ai pazienti di ripensare a qualsiasi delle clip che avevano visto mentre continuavano a registrare le attività dei neuroni del lobo temporale mediale dei pazienti. Un risultato è mostrato nella **Figura 4.24**, che indica la risposta del neurone che ha risposto alla clip de "I Simpson". La descrizione del paziente di ciò che ricordava è mostrata in fondo alla suddetta figura. Prima la paziente ha ricordato "qualcosa su New York", poi "il segno di Hollywood". Il neurone risponde debolmente o per niente a quei due ricordi. Tuttavia, il ricordo de "I Simpson" provoca una grande risposta, che continua come la persona continua a ricordare l'episodio (indicato dalle risate). Risultati come questo supportano l'idea che i neuroni nel lobo temporale mediale che rispondono alla percezione di oggetti o eventi specifici possono anche essere coinvolti nel ricordare questi oggetti ed eventi. Moran Cerf e collaboratori (2010) hanno fornito un'altra dimostrazione di come i pensieri possano influenzare l'attivazione dei neuroni. Vai a "CourseMate" per visualizzare i video che descrivono questa ricerca.



Figura 4.24 Attività di un neurone del lobo temporale mediale in un paziente epilettico mentre ricordava le cose indicate sotto il registratore. Una risposta arriva mentre la persona ricordava il programma TV "I Simpson". Prima, questo neurone è stato visto rispondere alla visione di un videoclip de "I Simpson".

DIMENSIONI DI SVILUPPO: esperienza e risposte neurali

In questo capitolo abbiamo visto che il cervello è organizzato per funzioni, con neuroni specializzati a rispondere alle facce, ai posti e ai corpi posizionati in specifiche aree del cervello. Qui considereremo il ruolo dell'esperienza nella creazione di questi neuroni specializzati. C'è evidenza, che descriveremo più tardi in

Dimensioni di Sviluppo, che alcune capacità percettive, come l'abilità di percepire il movimento, il contrasto luce-buio, le facce, la profondità, i gusti, e gli odori, sono presenti o quasi alla nascita, nonostante non siano sviluppate ai livelli dell'adulto. Altre capacità, quali la profondità di percezione del colore che può

essere vista con un occhio e l'attenzione visiva, emergono leggermente dopo e neanche a livelli adulti. Con il tempo, queste capacità migliorano - alcune velocemente, come l'acuità visiva, che raggiunge un livello quasi adulto ai 9 mesi d'età, e altre più lentamente, come il riconoscimento di volti, che continua a svilupparsi nell'adolescenza (Grill-Spector et al., 2008; Sherf et al., 2007).

Cosa causa questo miglioramento nel tempo? La maturazione biologica è chiaramente implicata, come abbiamo visto quando abbiamo descritto la connessione tra il miglioramento dell'acuità visiva e lo sviluppo dei coni e dei bastoncelli. Su un lungo lasso di tempo, c'è evidenza che alcuni aspetti del riconoscimento dei volti dipendono dall'emergenza dell'area fusiforme del volto (fusiform face area "FFA"), che non è completamente sviluppata fino all'adolescenza.

In aggiunta alla maturazione biologica, anche l'esperienza nella percezione dell'ambiente gioca un ruolo nello sviluppo percettivo. Una serie di testimonianze che supporta il ruolo dell'esperienza è quella della ricerca sulla plasticità esperienza-dipendente di cui abbiamo parlato nel Capitolo 3. Gli esperimenti di Blakemore e Cooper, nei quali hanno cresciuto dei gattini in tubi a strisce, mostrano che i sistemi visuali di quei gatti erano adattati all'ambiente dove erano cresciuti, perciò i gatti cresciuti vedendo solamente strisce verticali avevano neuroni che rispondevano solo a orientamenti verticali o quasi verticali.

Gli umani solitamente non vengono cresciuti in ambienti deprivati, ma in un ambiente nel quale diverse caratteristiche si ripetono regolarmente, queste caratteristiche ripetute dell'ambiente possono influenzare come il nostro sistema visivo si sviluppa e, perciò, come percepiamo. Un esempio di questo, che abbiamo descritto nel Capitolo 1, è la scoperta che le persone percepiscono le linee orizzontali e verticali con più facilità degli altri orientamenti, l'*effetto obliquo*. C'è evidenza che l'effetto obliquo esiste perché c'è un maggior numero di neuroni corticali che risponde ad orientamenti orizzontali e verticali, e non è una coincidenza che nell'ambiente linee orizzontali e verticali siano più frequenti di linee inclinate.

Il fatto che l'esperienza con l'ambiente possa modellare il sistema nervoso è alla base dell'*ipotesi competenza* (expertise hypothesis), che propone che la nostra competenza nel percepire certe cose può essere spiegata dai cambiamenti del cervello causati da lunghe esposizioni, pratica o allenamento (Bukach et al., 2006; Gauthier et al., 1999). Isabel Gauthier e collaboratori (1999) dimostrarono un effetto competenza usando l'fMRI per determinare il livello di attività nell'area fusiforme del volto (FFA) in risposta a facce e oggetti chiamati Greebles - famiglie di "esseri" generati da computer che hanno tutti la stessa configurazione di base ma differiscono nelle forme delle loro parti (**Figura**

4.25a). Inizialmente, venivano mostrate agli osservatori sia facce umane sia Greebles. I risultati per questa parte dell'esperimento, mostrati dalle coppie di barre nella **Figura 4.25b**, indicano che i neuroni FFA hanno risposto poco ai Greebles ma molto alle facce.

I partecipanti erano poi allenati nel "riconoscimento Greebles" per 7 ore lungo un periodo di 4 giorni. Dopo le sessioni di allenamento, i partecipanti erano diventati "esperti Greeble," fatto indicato dalla loro abilità a identificare rapidamente molti diversi Greeble secondo i nomi che avevano imparato durante l'allenamento. La coppia di barre a destra nella Figura 4.25b mostra come diventare un esperto Greeble abbia influenzato la risposta neurale nell'FFA dei partecipanti. Dopo l'allenamento, i neuroni FFA rispondevano quasi allo stesso modo ai Greebles e alle facce.

Questo risultato mostra che l'area FFA della corteccia risponde non solo alle facce ma anche ad altri oggetti complessi, e che gli oggetti ai quali i neuroni rispondono possono essere modificati dall'esperienza con tali oggetti.

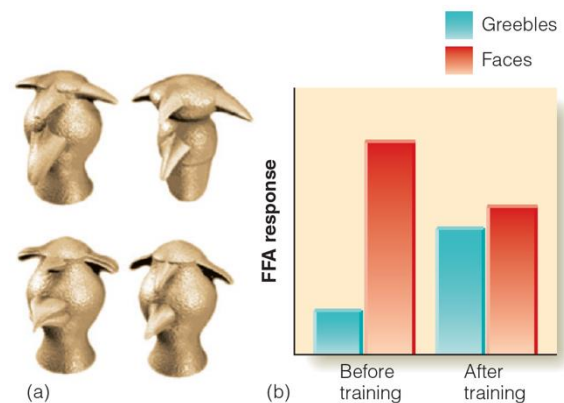


Figura 4.25 (a) Stimoli Greeble usati da Gauthier. I partecipanti erano stati allenati a chiamare per nome ciascun Greeble diverso. (b) Le risposte del cervello ai Greebles e alle facce prima e dopo l'allenamento Greeble.

Infatti, Gauthier ha anche mostrato che i neuroni nell'area FFA delle persone che sono esperti nel riconoscimento di macchine o uccelli rispondono bene non solo a volti umani ma anche alle macchine (per gli esperti di auto) e agli uccelli (per gli esperti di volatili; Gauthier et al., 2000). Recentemente, un altro studio ha mostrato che vedere la posizioni di pezzi degli scacchi su una scacchiera causa una maggiore attivazione dell'FFA negli esperti di scacchi piuttosto che nei non-esperti (Bilalic et al., 2011). I risultati come questo hanno portato molti ricercatori a suggerire che la ragione per cui l'FFA risponde bene alle facce è perché siamo "esperti di volti".

E' importante notare che nonostante ci sia buona evidenza che l'esperienza può influenzare i tipi di stimoli ai quali i neuroni rispondono, il ruolo dell'esperienza nello stabilire l'uso dell'FFA come

modulo per le facce è dibattuto. Alcuni ricercatori come Gauthier ritengono che l'esperienza è importante per la decisione di utilizzare l'FFA come modulo per le facce (Bukach et al., 2006); altri rispondono che il ruolo dell'FFA come area delle facce non dipende dall'esperienza (Kanwisher, 2010).

Qualsiasi sia il risultato di questo dibattito sull'FFA, non c'è dubbio sul fatto che le proprietà dei neuroni sono influenzate dalla nostra esperienza con gli stimoli nell'ambiente. Questa esperienza, che "regola" il nostro sistema percettivo per rispondere al meglio a ciò che è spesso presente nell'ambiente, è probabile che giochi un ruolo nella determinazione dei miglioramenti nella percezione che abbiamo tra l'infanzia e l'età adulta.

Ora che hai finito questo capitolo, hai le conoscenze necessarie per capire il materiale psicologico nei capitoli che seguono. Nei prossimi sei capitoli continueremo a discutere del sistema visivo, con ogni capitolo dedicato ad una specifica qualità visiva o processo. Il Capitolo 5 continua la nostra discussione su come percepiamo gli oggetti. Ci preoccupiamo ancora delle facce, ma il nostro focus principale saranno gli oggetti in generale, come anche come percepiamo diversi oggetti che sono organizzati per creare scene. Una cosa che noterete mentre leggete il prossimo capitolo è che non incontrerete la parola neurone fino ai due-terzi del capitolo. Uno dei messaggi del Capitolo 5 è che un grande numero delle ricerche sulla percezione se ne occupa a livello comportamentale, misurando la relazione tra gli stimoli e la percezione. Certamente, non ci libereremo mai dei neuroni, perché la fisiologia è parte della storia. Ma quando i neuroni riappariranno, sarete pronti per loro!