

9 Localization

Localization is a phenomenon where the eigenvectors of a random matrix concentrate on certain favored directions, as opposed to how they are found in rotationally invariant ensembles: equally likely to be anywhere. The ultimate localized matrix is a matrix with random diagonal entries and zero everywhere else, or

$$D = \begin{bmatrix} x_1 & 0 & \cdots & 0 \\ 0 & x_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & x_N \end{bmatrix} \quad (1)$$

for $x_i \sim \mathcal{D}$ for some probability distribution $\mathcal{D}$. The eigenvalues of this matrix are precisely the x_i , while the eigenvectors are simply the unit vectors $\mathbf{e}_i = (0, \dots, 0, 1, 0, \dots, 0)$, which are completely localized.

The standard method to measure localization is to use the *inverse participation ratio* (IPR), defined for a vector $\mathbf{v}$ by

$$\mathcal{I} = \frac{\sum_{i=1}^N v_i^4}{\left(\sum_{i=1}^N v_i^2\right)^2} = \frac{\sum_{i=1}^N v_i^4}{\|\mathbf{v}\|^2} = \sum_{i=1}^N v_i^4 \quad (2)$$

where in the last step we assume the eigenvector is normalized. This is called the inverse participation ratio because it is a way of measuring the (inverse) of the number of components of $\mathbf{v}$ that take substantial nonzero values. For instance, if the values of $\mathbf{v}$ are all equal, i.e., they all ‘participate’, then they will each have value $v_i = \frac{1}{\sqrt{N}}$ and

$$\mathcal{I} = \sum_{i=1}^N \left(\frac{1}{\sqrt{N}}\right)^4 = \sum_{i=1}^N \frac{1}{N^2} = \frac{1}{N} \quad (3)$$

Because the participation is large, the inverse participation is small. On the other hand, if only one component of $\mathbf{v}$ participates, with $v_1 = 1$ and $v_i = 0$ for $i > 1$, then

$$\mathcal{I} = 1^4 = 1 \quad (4)$$

Because the participation is small, the inverse participation is large. We see here the two extremes: the smallest $\mathcal{I}$ can be is N^{-1} while the largest it can be is 1. If $\mathcal{I}$ is of order 1, then we say the vector is localized; if it is of order N^{-1} , then we say the vector is delocalized. If it instead behaves like $N^{-\gamma}$ for some $0 < \gamma < 1$, then we have a partially delocalized or fractal eigenvector.

The inverse participation ratio can be extracted from the resolvent in the following way. First, note that if a matrix is nondegenerate, it can always be written as a sum over rank-one matrices associated with each of its eigenvectors, or

$$H = \sum_{j=1}^N x_j \mathbf{v}^{(j)} (\mathbf{v}^{(j)})^T \quad (5)$$

where $\mathbf{v}^{(j)}$ is the normalized eigenvector associated with eigenvalue x_j . The inverse matrix shares the same eigenvectors and only inverts the eigenvalue, or

$$H^{-1} = \sum_{j=1}^N \frac{\mathbf{v}^{(j)}(\mathbf{v}^{(j)})^T}{x_j} \quad (6)$$

Taking now the matrix involved in the resolvent, we see that along the diagonal,

$$(z - H)_{ii}^{-1} = \sum_{j=1}^N \frac{(v_i^{(j)})^2}{z - x_j} \quad (7)$$

To get the inverse participation ratio, we should want the eigenvector components to the fourth power, so we might write

$$\begin{aligned} |(z - H)_{ii}^{-1}|^2 &= \left(\sum_{j=1}^N \frac{(v_i^{(j)})^2}{z - x_j} \right)^* \left(\sum_{k=1}^N \frac{(v_i^{(k)})^2}{z - x_k} \right) = \left(\sum_{j=1}^N \frac{(v_i^{(j)})^2}{z^* - x_j} \right) \left(\sum_{k=1}^N \frac{(v_i^{(k)})^2}{z - x_k} \right) \\ &= \sum_{j=1}^N \frac{(v_i^{(j)})^4}{(\text{Re } z - x_j)^2 + (\text{Im } z)^2} + \sum_{j \neq k}^N \frac{(v_i^{(j)})^2}{z^* - x_j} \frac{(v_i^{(k)})^2}{z - x_k} \end{aligned} \quad (8)$$

Take now $z = x - i\epsilon$. We have

$$|(x - i\epsilon - H)_{ii}^{-1}|^2 = \sum_{j=1}^N \frac{(v_i^{(j)})^4}{(x - x_j)^2 + \epsilon^2} + \sum_{j \neq k}^N \frac{(v_i^{(j)})^2}{x - x_j + i\epsilon} \frac{(v_i^{(k)})^2}{x - x_k - i\epsilon} \quad (9)$$

We now take

$$\lim_{\epsilon \rightarrow 0} \epsilon |(x - i\epsilon - H)_{ii}^{-1}|^2 \quad (10)$$

The first term contains

$$\lim_{\epsilon \rightarrow 0} \frac{\epsilon}{(x - x_j)^2 + \epsilon^2} = \pi \delta(x - x_j) \quad (11)$$

a common resolution of the Dirac δ -function, which can be confirmed by integrating the left-hand side over some interval containing x_j or not before take the limit. If the matrix is nondegenerate, the second term is completely regular in the limit. It is regular when $x \neq x_j$ for any j because nothing is divergent, but it is also regular when $x = x_l$

for some $j = l$. Under that condition,

$$\begin{aligned}
& \epsilon \sum_{j \neq k}^N \frac{(v_i^{(j)})^2}{x - x_j + i\epsilon} \frac{(v_i^{(k)})^2}{x - x_k - i\epsilon} \\
&= \epsilon \sum_{j \neq k \neq l}^N \frac{(v_i^{(j)})^2}{x - x_j + i\epsilon} \frac{(v_i^{(k)})^2}{x - x_k - i\epsilon} \\
&\quad + \epsilon \sum_{k \neq l} \frac{(v_i^{(l)})^2}{x - x_l + i\epsilon} \frac{(v_i^{(k)})^2}{x - x_k - i\epsilon} + \epsilon \sum_{j \neq l} \frac{(v_i^{(j)})^2}{x - x_j + i\epsilon} \frac{(v_i^{(l)})^2}{x - x_l - i\epsilon} \\
&= \epsilon \sum_{j \neq k \neq l}^N \frac{(v_i^{(j)})^2}{x - x_j + i\epsilon} \frac{(v_i^{(k)})^2}{x - x_k - i\epsilon} \\
&\quad - \frac{1}{i} \sum_{k \neq l} (v_i^{(l)})^2 \frac{(v_i^{(k)})^2}{x - x_k - i\epsilon} + \frac{1}{i} \sum_{j \neq l} \frac{(v_i^{(j)})^2}{x - x_j + i\epsilon} (v_i^{(l)})^2 \\
&= \epsilon \sum_{j \neq k \neq l}^N \frac{(v_i^{(j)})^2}{x - x_j + i\epsilon} \frac{(v_i^{(k)})^2}{x - x_k - i\epsilon}
\end{aligned} \tag{12}$$

where the two term that were potentially nonzero in the zero ϵ limit cancel each other. We therefore have

$$\lim_{\epsilon \rightarrow 0} \epsilon |(x - i\epsilon - H)_{ii}^{-1}|^2 = \pi \sum_{j=1}^N (v_i^{(j)})^4 \delta(x - x_j) \tag{13}$$

Summing now over i , we have

$$\begin{aligned}
\lim_{\epsilon \rightarrow 0} \epsilon |G(x - i\epsilon)|^2 &= \lim_{\epsilon \rightarrow 0} \epsilon \frac{1}{N} \sum_{i=1}^N |(x - i\epsilon - H)_{ii}^{-1}|^2 \\
&= \pi \frac{1}{N} \sum_{j=1}^N \delta(x - x_j) \sum_{i=1}^N (v_i^{(j)})^4 = \pi \frac{1}{N} \sum_{j=1}^N \delta(x - x_j) \mathcal{I}_j \\
&= \pi \rho(x) \mathcal{I}(x)
\end{aligned}$$

and therefore the expected IPR for an eigenvector with eigenvalue x is

$$\mathcal{I}(x) = \frac{1}{\pi \rho(x)} \lim_{\epsilon \rightarrow 0} \epsilon |G(x - i\epsilon)|^2 \tag{14}$$

We can check the results we already know: that the GOE is delocalized and that a diagonal matrix with random entries is localized. In the former case,

$$G(z) = z \pm \sqrt{2 - z^2} \tag{15}$$

and

$$|G(x - i\epsilon)|^2 = 2 \pm 2x\sqrt{2 - x^2} + O(\epsilon) \tag{16}$$

Because this is finite at zero ϵ for all values of x , there is no localization. On the other hand, for a random diagonal matrix,

$$\begin{aligned}\lim_{\epsilon \rightarrow 0} \epsilon |G(x - i\epsilon)|^2 &= \lim_{\epsilon \rightarrow 0} \int dx' \frac{\epsilon p(x')}{(x - i\epsilon - x')(x + i\epsilon - x')} \\ &= \lim_{\epsilon \rightarrow 0} \int dx' \frac{\epsilon p(x')}{(x - x')^2 + \epsilon^2} \\ &= \pi \int dx' p(x') \delta(x - x') = \pi p(x)\end{aligned}$$

Since the spectral density $\rho(x) = p(x)$, we have precisely that $I(x) = 1$ for all x where $p(x) \neq 0$.

Another signature of localization is in the statistics of neighboring eigenvalues. We know that in delocalized systems like the GOE, the probability distribution of the gap between neighboring eigenvalues follows nearly the Wigner surmise, which for real-valued symmetric matrices is

$$p(s) = \frac{\pi s}{2} e^{-\pi s^2/4} \propto s + O(s^2) \quad (17)$$

The characteristic of the delocalized system is a *pseudogap* in the gap distribution. However, localized matrices behave differently.

To compute the gap statistics of D , we first want to probability that two eigenvalues sit a distance Δx apart given that one is located at x and that no other eigenvalues lie between x and $x + \Delta x$. The probability that an eigenvalue is located at $x + \Delta x$ given that a particular one is located at x without worrying about the other $N - 2$ eigenvalues is $p(x + \Delta x)$, the probability that the second eigenvalue is where it should be. The probability that one eigenvalue is *not* in that interval is

$$p(x' < x \text{ \& } x' > x + \Delta x) = \int_{-\infty}^x dx' p(x') + \int_{x+\Delta x}^{\infty} dx' p(x') \quad (18)$$

Therefore, the conditional probability that there is an eigenvalue at $x + \Delta x$ given that there is one at x and with none between them is

$$p_N(\Delta x \mid x_i = x) = p(x + \Delta x) \left(\int_{-\infty}^x dx' p(x') + \int_{x+\Delta x}^{\infty} dx' p(x') \right)^{N-2} \quad (19)$$

The probability that this is true given that *any* eigenvalue has value x is then

$$p_N(\Delta x \mid x) = \sum_{i=1}^N p_N(\Delta x \mid x_i = x) p(x_i) = N p_N(\Delta x \mid x_i = x) p(x) \quad (20)$$

Finally, the probability that such a gap exists anywhere is

$$\begin{aligned}p_N(\Delta x) &= \int_{-\infty}^{\infty} dx p_N(\Delta x \mid x) \\ &= N \int_{-\infty}^{\infty} dx p(x) p(x + \Delta x) \left(\int_{-\infty}^x dx' p(x') + \int_{x+\Delta x}^{\infty} dx' p(x') \right)^{N-2}\end{aligned} \quad (21)$$

We are analyzing this at large N , where the probability of a gap becomes shrinks if the gap is held constant. So, we rescale the gap with N , and instead analyze s with $\Delta x = s/Np(x)$. This gives

$$p_N(s) = \int_{-\infty}^{\infty} dx p(x + \frac{s}{Np(x)}) \left(\int_{-\infty}^x dx' p(x') + \int_{x + \frac{s}{Np(x)}}^{\infty} dx' p(x') \right)^{N-2} \quad (22)$$

Expanding in large N , this gives

$$\begin{aligned} p_N(s) &= \int_{-\infty}^{\infty} dx (p(x) + O(N^{-1})) \left(\int_{-\infty}^x dx' p(x') + \int_x^{\infty} dx' p(x') - \frac{s}{N} + O(N^{-2}) \right)^{N-2} \\ &= \int_{-\infty}^{\infty} dx (p(x) + O(N^{-1})) \left(\int_{-\infty}^{\infty} dx' p(x') - \frac{s}{N} + O(N^{-2}) \right)^{N-2} \\ &= \int_{-\infty}^{\infty} dx (p(x) + O(N^{-1})) \left(1 - \frac{s}{N} + O(N^{-2}) \right)^{N-2} \\ &= \int_{-\infty}^{\infty} dx p(x) e^{-s} + O(N^{-1}) = e^{-s} + O(N^{-1}) \end{aligned}$$

The resulting asymptotic distribution is the *Poisson* distribution, and we say the eigenvalue gaps have *Poisson statistics*.

A simple model for localization is the Rosenzweig–Porter model, which consists of adding a GOE matrix H to a random diagonal matrix D , or

$$M = D + aH \quad (23)$$

where a is a constant that scales the relative weight of D in the combination. We can understand the resolvent of the sum by use of the R transform. Since the R transform of a sum of free matrices is the sum of their R transforms, and because GOE matrices are free of any other independent matrix, we therefore have

$$R_M(t) = R_D(t) + R_{aH}(t) = R_D(t) + R_H(at) \quad (24)$$

where we have used the scaling property of R transforms as well. This implies that the Blue function of M is

$$B_M(t) = R_M(t) + \frac{1}{t} = R_D(t) + R_H(at) + \frac{1}{t} = B_D(t) + R_H(at) \quad (25)$$

where we have absorbed the $\frac{1}{t}$ into the Blue function of D . Writing $t = G_M(z)$, we have

$$\underbrace{B_M(G_M(z))}_z = B_D(G_M(z)) + R_H(aG_M(z)) \quad (26)$$

Finally, isolating $B_D(G_M(z))$ and applying G_D to both sides gives

$$\underbrace{G_D(B_D(G_M(z)))}_{G_M(z)} = G_D(z - R_H(aG_M(z))) \quad (27)$$

This is therefore a self-consistency equation for G_M , relying on G_D and R_H . Up to here this has been entirely general: the formula is exact for any free matrices H and D . In our case, we take H to be GOE, so $R_H(t) = \frac{1}{2}t$, and we have

$$G_M(z) = G_D \left(z - \frac{1}{2}aG_M(z) \right) \quad (28)$$

To have an analytic solution for this formula requires that p_D be sufficiently simple that the resolvent can be computed and that the resulting formula for G_M is not transcendental. Surprisingly, Gaussian p_D is not such a simple case. It is much simpler in fact to use Cauchy-distributed p_D , with

$$p_D(x) = \frac{1}{\pi\omega} \frac{\omega^2}{(x - \mu)^2 + \omega^2} \quad (29)$$

where ω is the width and μ is the mean. Then

$$G_D(z) = \int dx p(x) \frac{1}{z - x} = \int dx \frac{1}{\pi\omega} \frac{\omega^2}{(x - \mu)^2 + \omega^2} \frac{1}{z - x} \quad (30)$$

This integral can be treated as a contour integral, with the only poles at $x = z$ and $x = \mu \pm i\omega$. Take the contour running along the real line and then with an arc in the upper half-plane. The residue due to the pole at $\mu + i\omega$ is $\frac{1}{2\pi i} \frac{1}{z - \mu - i\omega}$. If $\text{Im } z < 0$, then this is the only in the contour, otherwise we also have the residue at $x = z$ with $-\frac{1}{2\pi i} \frac{2i\omega}{(z - \mu)^2 + \omega^2}$. In the first case the resolvent is

$$G_D(z) = \frac{1}{z - \mu - i\omega} \quad (31)$$

while in the second it is

$$G_D(z) = \frac{1}{z - \mu - i\omega} - \frac{2i\omega}{(z - \mu)^2 + \omega^2} = \frac{1}{z - \mu + i\omega} \quad (32)$$

In either case this is a simple formula, and

$$G_M(z) = \frac{1}{z - \frac{1}{2}aG_M(z) - \mu \pm i\omega} \quad (33)$$

which is quadratic and solved by

$$G_M(z) = \frac{1}{a} \left(z - \mu \pm i\omega \pm \sqrt{(z - \mu \pm i\omega)^2 - 2a} \right) \quad (34)$$

As a sanity check, we see that when z is scaled by $\sqrt{a}$ and a is very large, we have

$$\sqrt{a}G_M(\sqrt{a}z) = z - \frac{1}{\sqrt{a}}(\mu \mp i\omega) \pm \sqrt{(z - \frac{1}{\sqrt{a}}(\mu \mp i\omega))^2 - 2} = G_H(z) + O(a^{-\frac{1}{2}}) \quad (35)$$

and when a is taken to zero (with the correct branch of the square root chosen so that the term in parentheses goes to zero as $a \rightarrow 0$), we have

$$G_M(z) = G_D(z) + O(a) \quad (36)$$

What about the localization properties of M ? Well,

$$\lim_{\epsilon \rightarrow 0} |G(x - i\epsilon)|^2 = \frac{1}{a^2} \left| x - \mu - i\omega \pm \sqrt{(x - \mu - i\omega)^2 - 2a} \right|^2 \quad (37)$$

for any finite a , so there is not localization for any a . This is a bit disappointing for those who want a model with a localization transition. However, there are some things to note. First, we are working in a limit where N has already been taken to infinity. If a is allowed to scale with N , different outcomes are possible: there is localization for any scaling where $a \propto N^{-\gamma}$ for $\gamma > 0$. However, the localization is not a trivial one, since there is a range of positive γ where the IPR is nonzero but the spectrum is still Wigner-like. Work on such models is a field of active research, since techniques for sums of matrices that scale with different powers of N are still being developed.

On the other hand, localization transitions *do* occur if one replaces the fully-connected H with a sparse H whose nonzero entries correspond to the adjacency matrix of a tree-like graph. We will see later in the course when we treat sparse random matrices how to analyze such situations.