

Eigenvalue Problem

One of the most widely used procedures to estimate eigenvalues of a matrix A is the power method. Our basic assumptions are that A is a real $n \times n$ matrix, its eigen values $\lambda_1, \lambda_2, \dots, \lambda_n$ are numbered so that

$$|\lambda_1| > |\lambda_2| \geq |\lambda_3| \geq \dots \geq |\lambda_n| \quad \dots (1)$$

and that A has n linearly ^{independent} eigenvectors $\{\underline{u}_i\}_{i=1}^n$, where $A\underline{u}_i = \lambda_i \underline{u}_i$, $1 \leq i \leq n$. When the power method works, it gives us an estimate of the dominant eigenvalue, λ_1 .

We first need the fact that any $n \times 1$ vector can be expressed as a linear combination of the eigenvectors $\underline{u}_1, \underline{u}_2, \dots, \underline{u}_n$. To see this, let B be ~~the~~ $(n \times n)$ matrix defined such that the ~~k~~ ^{k} th column vector of B is $\underline{u}_k$, $1 \leq k \leq n$, that is $B = [\underline{u}_1, \underline{u}_2, \dots, \underline{u}_n]$.

Since $\underline{u}_i$'s are linearly independent, B is invertible.

Now, let $\underline{x}_0$ be an arbitrary $(n \times 1)$ vector, and consider the linear system of equations $B\underline{x} = \underline{x}_0$. Since B is invertible, this system has a unique solution, $\underline{x} = (\alpha_1, \alpha_2, \dots, \alpha_n)^T$, and so $\underline{x}_0$ can be uniquely expressed as

$$\underline{x}_0 = \alpha_1 \underline{u}_1 + \alpha_2 \underline{u}_2 + \dots + \alpha_n \underline{u}_n \quad \dots (2)$$

Next, observe that if $A\underline{u}_i = \lambda_i \underline{u}_i$ then $A(A\underline{u}_i) = \lambda_i A\underline{u}_i$

$\sim A^2 \underline{u}_i = \lambda_i^2 \underline{u}_i$. ~~Contd~~ Continuing, it follows that

$A^3 \underline{u}_i = \lambda_i^3 \underline{u}_i$, or in general

$$A^k \underline{u}_i = \lambda_i^k \underline{u}_i \quad \text{--- (3)}$$

Thus if we form the sequence

$$\underline{x}_1 = A \underline{x}_0$$

$$\underline{x}_2 = A \underline{x}_1, (\underline{x}_2 = A^2 \underline{x}_0)$$

$$\underline{x}_3 = A \underline{x}_2, (\underline{x}_3 = A^3 \underline{x}_0)$$

$$\underline{x}_k = A \underline{x}_{k-1}, (\underline{x}_k = A^k \underline{x}_0)$$

$$\begin{aligned} \therefore \underline{x}_k &= A^k (\alpha_1 \underline{u}_1 + \alpha_2 \underline{u}_2 + \dots + \alpha_n \underline{u}_n) \\ &= \alpha_1 (A^k \underline{u}_1) + \alpha_2 (A^k \underline{u}_2) + \dots + \alpha_n (A^k \underline{u}_n) \\ &= \alpha_1 \lambda_1^k \underline{u}_1 + \alpha_2 \lambda_2^k \underline{u}_2 + \dots + \alpha_n \lambda_n^k \underline{u}_n \quad (\text{by (3)}) \\ &\quad \text{--- (4)} \end{aligned}$$

Recalling (1), it is natural to write $\underline{x}_k$ as

$$\underline{x}_k = \lambda_1^k \left\{ \alpha_1 \underline{u}_1 + \alpha_2 \left(\frac{\lambda_2}{\lambda_1} \right)^k \underline{u}_2 + \dots + \alpha_n \left(\frac{\lambda_n}{\lambda_1} \right)^k \underline{u}_n \right\} \quad \text{--- (5)}$$

In (5), we see for large k that the approximation

$$\underline{x}_k \approx \lambda_1^k \alpha_1 \underline{u}_1 \text{ should be good since } |\lambda_1| > |\lambda_2| \geq |\lambda_3| \geq \dots \geq |\lambda_n|.$$

In order to estimate λ_1 , we replace k by $(k+1)$ in (5) to obtain $\underline{x}_{k+1} \approx \lambda_1^{k+1} \alpha_1 \underline{u}_1$, or $\underline{x}_{k+1} \approx \lambda_1 \underline{x}_k$. This suggests a number of ways to approximate λ_1 . For example, $\underline{x}_k^T \underline{x}_{k+1} \approx \lambda_1 \underline{x}_k^T \underline{x}_k$ & thus $\lambda_1 \approx \underline{x}_k^T \underline{x}_{k+1} / \underline{x}_k^T \underline{x}_k$. Another common procedure is to divide the largest component of $\underline{x}_k$ into the corresponding component of $\underline{x}_{k+1}$. If the largest component is the j^{th} component, then this amounts to $\lambda_1 \approx \underline{e}_j^T \underline{x}_{k+1} / \underline{e}_j^T \underline{x}_k$.

Since $\underline{x}_k \approx \lambda_1^k \alpha_1 \underline{u}_1$ i.e. $\underline{x}_k$ differs from the eigen vector $\underline{u}_1$ solely by a numerical factor and, consequently, is also an eigen vector corresponding to the same eigenvalue λ_1 .

$$[\text{Proof}] \quad A \underline{x}_k = A \cdot \lambda_1^k \alpha_1 \underline{u}_1 = \lambda_1^k \alpha_1 A \underline{u}_1 = \lambda_1^k \alpha_1 (\lambda_1 \underline{u}_1)$$

$\therefore \underline{x}_k$ is an eigen vector corresponding to λ_1

We wish to be fairly careful about the description of the power method, for this description has two parts, namely the 'practical' part which shows how the method is actually carried out, and the 'theoretical' part which shows why the method works, what can go wrong, and how to interpret the output. First of all, powers of A are never actually computed in the implementation of the power method, nor is it necessary to know the λ_i 's and the u_i 's in (2).

We merely guess a vector $\underline{x}_0$ & form the sequence $\{\underline{x}_k\}$ from $\underline{x}_k = A \underline{x}_{k-1}$; $k=1, 2, 3, \dots$. The equation $\underline{x}_k = A^k \underline{x}_0$ is the theoretical basis for the power method, but all we really need is the vector $\underline{x}_k$. The estimates for λ can be found by generating the sequence of scalars $\{\beta_k\}$ from the equation

$$\beta_k = \underline{v}^T \underline{x}_{k+1} / \underline{v}^T \underline{x}_k \quad - - (6)$$

where $\underline{v}$ is usually $\underline{x}_k$ or the r th unit vector $\underline{e}_r$ (where the r th component of $\underline{x}_k$ is the dominant one).

When programming the power method, it is a little easier to generate the estimate β_k from $\beta_k = \underline{x}_k^T \underline{x}_{k+1} / \underline{x}_k^T \underline{x}_k$ than to include a test at each step to determine the maximum component of $\underline{x}_k$. The advantage in finding the maximum component of $\underline{x}_k$ is that round-off errors will tend to be reduced. Finally as k increases, the vectors $\underline{x}_k$ are usually growing quite large if $|\lambda| > 1$ (or quite small if $|\lambda| < 1$), as can be seen from (4). So it is desirable to 'scale' or 'normalize' $\underline{x}_k$ to keep it within the computational limits. This can be done by dividing $\underline{x}_k$ by $\|\underline{x}_k\|$ at each step. We now show that this has no effect on the convergence.

Let $\underline{x}_0$ be given by (2), where $\|\underline{x}_0\| = 1$. Let $\underline{y}_i = A \underline{x}_{i-1}$, $\beta_i = \underline{v}^T \underline{y}_i / \underline{v}^T \underline{x}_{i-1}$, $\underline{x}_i = \underline{y}_i / \|\underline{y}_i\|$, for $i=1, 2, \dots$

Put $\underline{w}_0 = \underline{x}_0$ & $\underline{w}_i = A \underline{w}_{i-1}$, $i=1, 2, \dots$

$$\therefore \underline{y}_1 = A \underline{x}_0 = \underline{w}_1$$

$$\underline{w}_2 = A \underline{w}_1 = A^2 \underline{x}_0 = A^2 \underline{w}_0$$

$$\underline{w}_3 = A \underline{w}_2 = A^3 \underline{w}_0 = A^3 \underline{x}_0$$

& in general, $\underline{w}_i = A^i \underline{x}_0 = A^i \underline{w}_0$

$$\text{Also } \underline{y}_2 = A \underline{x}_1 = A \frac{\underline{y}_1}{\|\underline{y}_1\|} = \frac{A^2 \underline{w}_0}{\|\underline{y}_1\|}$$

$$\underline{y}_3 = A \underline{x}_2 = A \frac{\underline{y}_2}{\|\underline{y}_2\|} = \frac{A^3 \underline{w}_0}{\|\underline{y}_1\| \|\underline{y}_2\|}$$

& in general, $\underline{y}_i = \frac{A^i \underline{w}_0}{\|\underline{y}_1\| \|\underline{y}_2\| \dots \|\underline{y}_{i-1}\|}$

$$\text{Now } \frac{\underline{v}^T \underline{y}_i}{\underline{v}^T \underline{x}_{i-1}} = \frac{\underline{v}^T \frac{A^i \underline{w}_0}{\|\underline{y}_1\| \|\underline{y}_2\| \dots \|\underline{y}_{i-1}\|}}{\underline{v}^T \frac{\underline{y}_{i-1}}{\|\underline{y}_{i-1}\|}}$$

$$= \frac{\underline{v}^T \frac{A^i \underline{w}_0}{\|\underline{y}_1\| \|\underline{y}_2\| \dots \|\underline{y}_{i-1}\|}}{\underline{v}^T \frac{A^{i-1} \underline{w}_0}{\|\underline{y}_1\| \|\underline{y}_2\| \dots \|\underline{y}_{i-2}\|}}$$

$$= \underline{v}^T A^i \underline{w}_0 / \underline{v}^T A^{i-1} \underline{w}_0$$

$$= \underline{v}^T \underline{w}_i / \underline{v}^T \underline{w}_{i-1}$$

which is similar to $\beta_k = \underline{v}^T \underline{x}_{k+1} / \underline{v}^T \underline{x}_k$

Thus, normalizing $\underline{x}_i$ in each step does not effect the convergence.

We like to mention here that $\underline{x}_{k+1} (= A \underline{x}_k)$ should be a finite vector if $\underline{x}_k$ is normalized.

In order to study the theoretical convergence of the power method, it is most convenient to consider a variation of the method in which β_k are calculated in a slightly different fashion from the two ways described above. In this variation, we assume $\underline{v}$ is a fixed vector

such that $\underline{v}^T \underline{u}_1 \neq 0$. We then define β_k to be given by $\beta_k = \underline{v}^T \underline{x}_{k+1} / \underline{v}^T \underline{x}_k$. If we set $v_j = \underline{v}^T \underline{u}_j$, equation (5) shows that

$$\beta_k = \lambda_1 \frac{\alpha_1 v_1 + \alpha_2 v_2 \left(\frac{\lambda_2}{\lambda_1}\right)^{k+1} + \dots + \alpha_n v_n \left(\frac{\lambda_n}{\lambda_1}\right)^{k+1}}{\alpha_1 v_1 + \alpha_2 v_2 \left(\frac{\lambda_2}{\lambda_1}\right)^k + \dots + \alpha_n v_n \left(\frac{\lambda_n}{\lambda_1}\right)^k} \quad (7).$$

The theoretical result (7) reveals that the sequence $\{\beta_k\}$ converges to λ_1 at essentially the same rate as $\left(\frac{\lambda_2}{\lambda_1}\right)^k$ converges to zero.

If we generate the sequence $\{\beta_k\}$ by $\beta_k = \underline{x}_k^T \underline{x}_{k+1} / \underline{x}_k^T \underline{x}_k$ where $\underline{x}_k$ has been normalized at each step, then also $v_j = \underline{x}_k^T \underline{u}_j$ are finite quantities and also converges to λ_1 , since the analysis is similar to that in (7).

When β_k is calculated by the formula

$$\beta_k = \underline{e}_j^T \underline{x}_{k+1} / \underline{e}_j^T \underline{x}_k$$

the convergence of $\{\beta_k\}$ to λ_1 can be shown in the following manner.

We have by (4)

$$\begin{aligned} \underline{x}_k &= \alpha_1 \lambda_1^k \underline{u}_1 + \alpha_2 \lambda_2^k \underline{u}_2 + \dots + \alpha_n \lambda_n^k \underline{u}_n \\ &= \sum_{j=1}^n \alpha_j \lambda_j^k \underline{u}_j \quad (4'). \end{aligned}$$

Let $\underline{x}_k = (x_1^k, x_2^k, \dots, x_n^k)^T$ and choose a basis $\{\underline{e}_1', \underline{e}_2', \dots, \underline{e}_n'\}$ in E_n .

Decomposing the eigen vectors $\underline{u}_j$ into the vectors of the basis, we have

$$\underline{u}_j = \sum_{i=1}^n u_i^j \underline{e}_i', \text{ where } \underline{u}_j = (u_1^j, u_2^j, \dots, u_n^j). \quad (8).$$

Substituting (8) into (4'), we obtain

$$\underline{x}_k = \sum_{j=1}^n \alpha_j \lambda_j^k \sum_{i=1}^n u_i^j \underline{e}_i$$

or changing the order of summation

$$\underline{x}_k = \sum_{i=1}^n \underline{e}_i \sum_{j=1}^n \alpha_j u_i^j \lambda_j^k \quad \dots (9)$$

Let the i^{th} component of $\underline{x}_k$ is the largest component.

Equating the i^{th} coordinate of $\underline{x}_k$ from (9), we have

$$x_i^k = \sum_{j=1}^n \alpha_j u_i^j \lambda_j^k \quad \dots (9')$$

Similarly, $x_i^{k+1} = \sum_{j=1}^n \alpha_j u_i^j \lambda_j^{k+1} \quad \dots (9'')$

Dividing the second sum by the first, we get

$$\rho_k = \frac{x_i^{k+1}}{x_i^k} = \frac{\alpha_1 u_i^1 \lambda_1^{k+1} + \alpha_2 u_i^2 \lambda_2^{k+1} + \dots + \alpha_n u_i^n \lambda_n^{k+1}}{\alpha_1 u_i^1 \lambda_1^k + \alpha_2 u_i^2 \lambda_2^k + \dots + \alpha_n u_i^n \lambda_n^k} \quad \dots (10)$$

Suppose that $\alpha_1 \neq 0$ & $u_i^1 \neq 0$. This can be achieved by appropriate choice of the initial vector $\underline{x}_0$.

Transform expression (10) as follows:

$$\rho_k = \frac{x_i^{k+1}}{x_i^k} = \lambda_1 \frac{1 + \frac{\alpha_2 u_i^2}{\alpha_1 u_i^1} \left(\frac{\lambda_2}{\lambda_1}\right)^{k+1} + \dots + \frac{\alpha_n u_i^n}{\alpha_1 u_i^1} \left(\frac{\lambda_n}{\lambda_1}\right)^{k+1}}{1 + \frac{\alpha_2 u_i^2}{\alpha_1 u_i^1} \left(\frac{\lambda_2}{\lambda_1}\right)^k + \dots + \frac{\alpha_n u_i^n}{\alpha_1 u_i^1} \left(\frac{\lambda_n}{\lambda_1}\right)^k} \quad \dots (11)$$

The theoretical result (11) reveals that the sequence $\{\rho_k\}$ also converges to λ_1 .