

4

Propriedades gerais de matrizes de Toeplitz

Neste capítulo desenvolveremos propriedades elementares de matrizes de Toeplitz e usaremos as mesmas para construir precisamente a região de matrizes Landau-regulares $\mathcal{R}_k$. Para um vetor $v = (v_1, \dots, v_k) \in \mathbb{R}^k$, sua reversão é $v_{\text{rev}} = (v_k, \dots, v_1)$. Dizemos que v é *par* se $v = v_{\text{rev}}$ e v é *ímpar* se $v = -v_{\text{rev}}$.

Proposição 4.1 *Se v é um autovetor de uma matriz de Toeplitz M , então v_{rev} também é autovetor associado ao mesmo autovalor. Em particular, um autovalor simples tem autovetor correspondente par ou ímpar.*

Prova. Para qualquer vetor v , é fácil ver que $Mv_{\text{rev}} = (Mv)_{\text{rev}}$. Se v é um autovetor, $Mv_{\text{rev}} = (Mv)_{\text{rev}} = \lambda v_{\text{rev}}$. Se λ é simples, v_{rev} deve ser múltiplo de v . Mas como v e v_{rev} têm a mesma norma (euclidiana), $v_{\text{rev}} = \pm v$. ■

Para $v = (v_1, v_2, \dots, v_k)$, sua *extensão* é o vetor $v_{\text{ext}} = (v_1, \dots, v_k, 0)$. A operação inversa, a remoção de última coordenada nula de um vetor é um *truncamento*.

Proposição 4.2 *Seja $n_k(\lambda) > 0$ a multiplicidade de λ como autovalor de $M_k = M_k(t_1, t_2, \dots, t_k)$ e $E_k(\lambda)$ o autoespaço associado a λ . Então*

$$(a) \quad n_{k+1} - n_k = \pm 1.$$

(b) *Os seguintes itens são equivalentes.*

- (i) *para todo $v \in E_k(\lambda)$, $v_{\text{ext}} \in E_{k+1}(\lambda)$,*
- (ii) *$n_{k+1}(\lambda) = n_k(\lambda) + 1$,*
- (iii) *$E_{k+1}(\lambda)$ contém algum vetor com primeira coordenada não nula.*

(c) *A propriedade (b)(iii) implica que o mesmo vale para $E_k(\lambda)$.*

Prova. Começaremos provando (b). Escolha uma base $\{w^{(i)}\}$ para $E_{k+1}(\lambda)$. Se a última componente de algum desses vetores, digamos $w^{(1)}$, não se anula, substituímos $w^{(i)}$ por $w^{(i)} - \alpha_i w^{(1)}$, $i \geq 2$, com α_i escolhido de forma a

anular a última coordenada dessa combinação linear. Esses $n_{k+1} - 1$ vetores são linearmente independentes e truncando a última coordenada de cada vetor, obtemos $n_{k+1} - 1$ autovetores de M_k . Se todos os $w^{(i)}$ tivessem última coordenada nula, teríamos n_{k+1} autovetores para M_k . De qualquer forma, $n_k \geq n_{k+1} - 1$ e a igualdade só se verifica quando E_{k+1} contém algum vetor com última coordenada não nula. A implicação $(b)(ii) \Rightarrow (b)(iii)$ também está provada porque reverter um autovetor com última coordenada não nula obtém um autovetor com primeira coordenada não nula.

Analogamente, seja $\{v^{(i)}\}$ uma base para $E_k(\lambda)$. A extensão $v_{\text{ext}}^{(i)}$ é autovetor de M_{k+1} se, e somente se, a última coordenada de $M_{k+1}v_{\text{ext}}^{(i)}$ for nula. Caso isso não se verifique para todos os vetores da base, podemos supor que $M_{k+1}v_{\text{ext}}^{(1)}$ tem última coordenada diferente de zero. Substituímos $v^{(i)}$ por $v^{(i)} - \beta_i v^{(1)}$, para $i \geq 2$, anulando a última coordenada de $M_k(v^{(i)} - \beta_i v^{(1)})$. Construimos assim $n_k - 1$ autovetores de M_{k+1} associados a λ , mostrando que $n_{k+1} \geq n_k - 1$ e, em consequência,

$$n_k - 1 \leq n_{k+1} \leq n_k + 1. \quad (4-1)$$

Para verificar $(b)(i) \Rightarrow (b)(ii)$, escolha $v^{(1)}$ de modo que sua primeira coordenada não nula tenha o menor índice entre os vetores $v \in E_k$, digamos $v_j^{(1)} \neq 0$. Assim, tanto $v^{(i)}$ quanto $v_{\text{rev}}^{(i)}$ têm as $j - 1$ primeiras e últimas coordenadas nulas. Por hipótese, $v_{\text{ext}}^{(i)}$ é autovetor de M_{k+1} e $(v_{\text{ext}}^{(1)})_{\text{rev}}$ também é. A coordenada $k + 2 - j$ de $(v_{\text{ext}}^{(1)})_{\text{rev}}$ é diferente de zero e a mesma coordenada de $v_{\text{ext}}^{(i)}$ é nula. Portanto $n_{k+1} \geq n_k + 1$ e a igualdade $n_{k+1} = n_k + 1$ segue de (4-1).

Supondo que vale $(b)(iii)$, sejam $w \in E_{k+1}$ com $w_{k+1} \neq 0$, $v \in E_k$ e $\mu = (M_{k+1}v_{\text{ext}})_{k+1}$.

$$\lambda(w, v_{\text{ext}}) = (M_{k+1}w, v_{\text{ext}}) = (w, M_{k+1}v_{\text{ext}}) = \lambda(w, v_{\text{ext}}) + w_{k+1} \mu,$$

e então $\mu = 0$, estabelecendo $(b)(i)$.

Agora, provaremos (c). Vimos que E_{k+1} é gerado por $\{v_{\text{ext}}^{(i)}, (v_{\text{ext}}^{(1)})_{\text{rev}}\}$. Como os vetores $v_{\text{ext}}^{(i)}$ têm a última coordenada nula, $(b)(iii)$ implica que a última coordenada de $(v_{\text{ext}}^{(1)})_{\text{rev}}$ é diferente de zero. Mas a última coordenada de $(v_{\text{ext}}^{(1)})_{\text{rev}}$ é a primeira coordenada de $v^{(1)}$.

Finalmente, vamos provar (a). Suponha que não vale $(b)(ii)$. Então a última coordenada de qualquer vetor em E_{k+1} é nula. Todos os vetores de

*A equação é uma consequência do entrelaçamento dos autovalores de matrizes simétricas encaixadas, que será demonstrado no apêndice C, mas neste caso especial apresentamos uma demonstração mais elementar e direta.

E_{k+1} podem ser truncados para E_k e depois estendidos para E_{k+1} . Como não vale (b)(i), esses truncamentos não geram E_k e $n_k \geq n_{k+1} + 1$. Usando (4-1), $n_{k+1} = n_k - 1$. ■

Corolário 4.3 *Se $n_{k+1}(\lambda) < n_k(\lambda)$, então n_k decresce até zero.*

Prova. Pela proposição 4.2(b), todos os vetores em E_{k+1} têm primeira coordenada nula. Pela proposição 4.2(c), o mesmo vale para E_{k+2} , e a condição de decrescimento se propaga. ■

Proposição 4.4 *(Comportamento dos autovetores múltiplos)*

- (a) *Suponha que $n_{k-1}(\lambda) = 0$, $n_k(\lambda) = 1$ e $v = (v_1, \dots, v_k)$ é o autovetor de M_k associado a λ . Então $v_1 \neq 0$ e $n_{k+1} = 2$ se, e somente se,*

$$t_{k+1} = -v_1^{-1}(t_k v_2 + \dots + t_2 v_k).$$

Neste caso $E_{k+1}(\lambda)$ é gerado por $(v_1, \dots, v_k, 0)$ e $(0, v_1, \dots, v_k)$.

- (b) *Se $n_j(\lambda)$ cresce de $n_k(\lambda) = 1$ até $n_{k+m}(\lambda) = m + 1$, $m \geq 2$, cada entrada t_{k+q} , $1 \leq q \leq m$, é determinada pela convolução discreta*

$$t_{k+q} = -v_1^{-1}(t_{k+q-1}v_2 + t_{k+q-2}v_3 + \dots + t_{q+1}v_k), \quad (4-2)$$

com o autovetor v de M_k . O autoespaço correspondente E_{k+q} é gerado pelos vetores formados pelo bloco $(v_1, \dots, v_k)$ precedido por i zeros e seguido por $q - i$ zeros, com $0 \leq i \leq q$. Na convolução (4-2), o valor de t_{k+q} permanece inalterado se trocamos v por outro autovetor $w = (w_1, \dots, w_{k+l}) \in E_{k+l}(\lambda)$, com $l + 1 \leq q \leq m$, desde que $w_1 \neq 0$.

$$t_{k+q} = -w_1^{-1}(t_{k+q-1}w_2 + \dots + t_{q-l+1}w_{k+l})$$

- (c) *Se t_{k+m+1} difere do valor dado por 4-2, $n_{k+m+1}(\lambda) = m$ e $n_{k+m+j}(\lambda)$ decresce monotonamente até $n_{k+2m+1}(\lambda) = 0$. O autoespaço E_{k+m+j} é gerado por vetores formados pelo bloco $(v_1, \dots, v_k)$ precedido por i zeros e seguido por $m + j - i$ zeros, com $j \leq i \leq m$.*
- (d) *$E_{k+q}(\lambda)$ pode ser decomposto na soma dos seus subespaços par e ímpar (ie, formados pelos vetores pares ou ímpares), cujas dimensões diferem por no máximo um. Além disso, quando $n_{k+q}(\lambda)$ é ímpar, o subespaço de maior dimensão tem a paridade de v , o autovetor simples em $E_k(\lambda)$.*

(e) Se λ é o maior autovalor de M_k , então sua multiplicidade como autovalor de M_{k-1} é $n_{k-1}(\lambda) = n_k(\lambda) - 1$. E se $n_k(\lambda) \geq 2$, então λ também é o maior autovalor de M_{k-1} . O mesmo vale para o menor autovalor.

Prova. Começaremos por (a). Se $n_{k-1}(\lambda) = 0$, $n_k(\lambda) = 1$ e $v = (v_1, \dots, v_k)$ é o autovetor de M_k associado a λ , então $v_1 \neq 0$ pois de outra forma $(v_2, v_3, \dots, v_k)$ seria autovetor de M_{k-1} . Pela proposição 4.2(b)(i), $n_{k+1} = 2$ desde que v_{ext} seja autovetor de M_{k+1} . Mas isso é equivalente à última coordenada de $M_{k+1}(v_{\text{ext}})$ se anular, ou seja,

$$t_{k+1} = -v_1^{-1}(t_k v_2 + \dots + t_2 v_k).$$

Com t_{k+1} desta forma, o autoespaço E_{k+1} é gerado por $v_{\text{ext}} = (v_1, \dots, v_k, 0)$ e $(v_{\text{ext}})_{\text{rev}} = (0, v_k, \dots, v_1) = \pm(0, v_1, \dots, v_k)$.

Vamos a (b). Para manter o crescimento de n_k , precisamos que cada vetor da base de M_{k+j} se estenda para um autovetor de M_{k+j+1} pela operação $(\cdot)_{\text{ext}}$. Supondo por indução que E_{k+j} é gerado por $(0, \dots, 0, v_1, \dots, v_k), \dots, (v_1, \dots, v_k, 0, \dots, 0)$, é suficiente verificar que a extensão de um autovetor com primeira coordenada não nula está em E_{k+j+1} . E novamente, isso equivale a

$$t_{k+j+1} = -v_1^{-1}(t_{k+j} v_2 + t_{k+j-1} v_3 + \dots + t_{j+2} v_k)$$

e E_{k+j+1} é gerado por $(0, \dots, 0, v_1, \dots, v_k, 0), \dots, (v_1, \dots, v_k, 0, \dots, 0)$, extensões de vetores em E_{k+j} e por $(0, \dots, 0, v_k, \dots, v_1) = \pm(0, \dots, 0, v_1, \dots, v_k)$.

Pela proposição 4.2, podemos calcular t_{k+j+1} substituindo v por qualquer autovetor w de M_{k+j} com primeira coordenada não nula.

Agora, (c). Se t_{k+m+1} não é dado pela convolução 4-2, $n_{k+m+1} = n_{k+m} - 1 = m$ e, pelo corolário 4.3, a multiplicidade continua a decrescer até zero. Uma base para o autoespaço E_{k+m+j} é formada por uma base para $E_{k+m+j-1}$ de onde retiramos o vetor com menor número de componentes iniciais nulas e estendemos os demais vetores acrescentando um zero.

Provaremos (d). Considerando as bases descritas em (b) e (c) para E_{k+q} , temos que se u é um elemento da base, $\pm u_{\text{rev}}$ também é. Se $u \neq \pm u_{\text{rev}}$, podemos substituir u e $\pm u_{\text{rev}}$ por $u + u_{\text{rev}}$ e $u - u_{\text{rev}}$. Observe ainda que $u = \pm u_{\text{rev}}$ para no máximo um vetor da base e nesse caso u tem a mesma paridade que o autovetor simples de E_k . Ou seja, construímos uma base de vetores pares e ímpares para E_{k+q} , na qual a diferença entre o número de vetores pares e ímpares é no máximo 1 e a paridade dominante é a mesma de E_k .

Para (e), seja λ o maior autovalor de M_k . O entrelaçamento de autovalores (apêndice C) garante que a multiplicidade de λ em M_k não pode ser menor que a sua multiplicidade em M_{k-1} e que se λ for autovalor de M_{k-1} , deve ser o maior deles. ■

No Capítulo 2, observamos que o grau de $\sigma_{\text{norm}} : \mathcal{T}_{\text{norm}}^k \rightarrow L_k$ não está definido e sugerimos restringir σ_{norm} a alguma componente conexa limitada de $\mathcal{U}_{\text{norm}}$, o aberto de matrizes de espectro simples em $\mathcal{T}_{\text{norm}}^k$. Como observado por Delsarte e Genin ([DG]), nem toda componente conexa $\mathcal{C}$ de $\mathcal{U}_{\text{norm}}$ tem chance de ser mapeada sobrejetivamente por Λ no interior de L_k . De fato, os autovalores das matrizes em $\mathcal{C}$ podem ser rotulados p ou i , dependendo da paridade do autovetor associado. Mais até, a rotulação é a mesma para todas as matrizes em $\mathcal{C}$, pela continuidade dos autovalores e autovetores para matrizes de espectro simples (Apêndice A). Para exemplificar, considere uma componente $\mathcal{C}$ na qual os autovalores ordenados são rotulados por $piip$. Uma matriz com autovalores $\lambda_1 < \lambda_2 = \lambda_3 < \lambda_4$ em $\partial\mathcal{C}$ teria um autovalor duplo com um autoespaço de dimensão dois consistindo apenas de autovetores ímpares, por continuidade novamente, o que contradiz a proposição 4.4(d): autoespaços de matrizes de Toeplitz se decompõem na soma direta das partes par e ímpar, cujas dimensões não diferem por mais do que 1. Assim, é conveniente considerar uma componente conexa de $\mathcal{U}_{\text{norm}}$ na qual autovalores adjacentes tenham paridades opostas. O argumento motiva a definição a seguir.

Definição 4.5 *Uma matriz de Toeplitz $M_k(t_1, \dots, t_k)$ é Landau-regular se todas as submatrizes $M_j(t_1, \dots, t_j)$ têm autovalores simples com paridades alternadas, sendo o maior deles par. Seja $\mathcal{R}_k$ o conjunto das matrizes de Toeplitz Landau-regulares da forma $M_k(0, 1, t_3, \dots, t_k)$.*

No texto, identificamos o conjunto $\mathcal{R}_k$ com um subconjunto de $\mathbb{R}^{k-2}$ através da imersão $M_k(0, 1, t_3, \dots, t_k) \mapsto (t_3, \dots, t_k)$. Os valores de (t_3, t_4) para os quais $M_4(0, 1, t_3, t_4)$ tem espectro múltiplo recortam o plano em componentes conexas nas quais a paridade dos autovalores está indicada na figura 4.1. O conjunto de matrizes Landau-regulares $\mathcal{R}_4$ é exatamente a componente limitada.

Por continuidade dos autovalores ordenados, matrizes Landau-regulares formam um conjunto aberto. Um exemplo de matriz em $\mathcal{R}_k$ é $M_k(0, 1, 0, \dots, 0)$. A menos de um múltiplo aditivo da identidade e de uma constante multiplicativa, a matriz é uma discretização da segunda derivada para funções que se anulam nos extremos de um intervalo, digamos

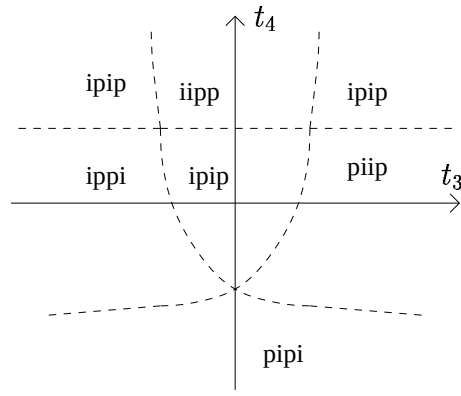


Figura 4.1: Paridades de autovalores nas componentes de $\mathcal{T}_{\text{norm}}^4$

$[0, \pi]$. E, de fato, seus autovetores são discretizações dos autovetores da segunda derivada que valem 0 nos extremos do intervalo, $\text{sen}(\ell x)$. Para que $v_\ell = (\text{sen}(\frac{\ell\pi}{k+1}), \text{sen}(\frac{2\ell\pi}{k+1}), \dots, \text{sen}(\frac{k\ell\pi}{k+1}))$ seja autovetor precisamos ter $\text{sen}((i-1)\theta) + \text{sen}((i+1)\theta) = \lambda_\ell \text{sen}(i\theta)$, para $1 \leq i \leq k$, onde $\theta = \frac{\ell\pi}{k+1}$. Para isso basta tomar $\lambda_\ell = 2\cos(\frac{k\ell\pi}{k+1})$. Observe ainda que $\lambda_k < \dots < \lambda_2 < \lambda_1$ e que v_1 é par, v_2 é ímpar, $\dots$, logo M_k tem a alternância certa.

Lema 4.6 (*Geometria de $\mathcal{R}_k$*)

- (a) Para $k > 2$, $\mathcal{R}_k \subset \mathbb{R}^{k-2}$ consiste de pontos $(t_3, \dots, t_k)$ tais que $(t_3, \dots, t_{k-1}) \in \mathcal{R}_{k-1} \subset \mathbb{R}^{k-3}$. A fronteira de $\mathcal{R}_k$ está contida no fecho da união das $(k-1)$ superfícies B_i , definidas como gráficos das funções $f_i: \mathcal{R}_{k-1} \rightarrow \mathbb{R}$ dadas por

$$t_k = f_i(0, 1, t_3, \dots, t_{k-1}) = -(v_1^{(i)})^{-1}(t_{k-1}v_2^{(i)} + \dots + t_2v_{k-1}^{(i)}),$$

onde $v^{(i)} = (v_1^{(i)}, v_2^{(i)}, \dots, v_{k-1}^{(i)})$ é o i -ésimo autovetor de $M_{k-1}(0, 1, t_3, \dots, t_{k-1})$, $1 \leq i \leq k-1$. As matrizes em B_i têm espectro duplo $\lambda_i = \lambda_{i+1}$.

- (b) O conjunto $\mathcal{R}_k$ está acima de B_i quando $v^{(i)}$ é par, e abaixo de B_i quando $v^{(i)}$ é ímpar.
- (c) Se M_k está na fronteira de $\mathcal{R}_k$, qualquer autovalor de M_k satisfaz $n_k(\lambda) = n_{k-1}(\lambda) + 1$.
- (d) $\overline{\mathcal{R}_k}$ é compacto.

Prova. A prova será feita por indução. A matriz $M_2(0, 1)$ tem autovetores $(1, 1)$ e $(1, -1)$ e autovalores ± 1 . Em dimensão 3, $M_3(0, 1, t_3)$ tem o

autovalor $-t_3$ associado ao autovetor ímpar $(1, 0, -1)$ e autovalores pares $(t_3 \pm \sqrt{t_3^2 + 8})/2$. Então $M_3(0, 1, t_3)$ é Landau-regular quando

$$t_3 - \sqrt{t_3^2 + 8} < -2t_3 < t_3 + \sqrt{t_3^2 + 8},$$

ou seja, para $-1 < t_3 < 1$. As afirmações sobre a fronteira de $\mathcal{R}_3$ são imediatas, provando o lema para $k = 3$.

Por definição, $\mathcal{R}_k \subset \mathcal{R}_{k-1} \times \mathbb{R}$. A fronteira de $\mathcal{R}_k$ é formada por matrizes de espectro múltiplo ou matrizes na parede cilíndrica $\partial\mathcal{R}_{k-1} \times \mathbb{R}$. Veremos adiante que a primeira possibilidade já contém a fronteira toda.

Dada uma matriz $M_k \in \mathcal{R}_k$, estudaremos o efeito da variação de t_k sobre os autovalores. O Apêndice A mostra que se $\lambda_i(t)$ é o i -ésimo autovalor de $M(t) = M_k + tM_k(0, \dots, 0, 1)$, então

$$\lambda'_i(t) = \langle M_k(0, \dots, 0, 1)u^{(i)}(t), u^{(i)}(t) \rangle = \pm \left(u_1^{(i)}\right)^2, \quad (4-3)$$

onde o sinal é positivo (resp. negativo) quando o autovetor $u^{(i)}$ é par (ímpar). Assim os autovalores pares de $M(t)$ crescem com t e os ímpares decrescem. O caminho $M(t)$ só atinge a fronteira de $\mathcal{R}_k$ quando dois autovalores colapsam em um autovalor duplo. Como os autovalores de M_{k-1} são simples, a proposição 4.4(a) fornece uma fórmula para esses pontos, estabelecendo (a) para esse trecho da fronteira.

Se a matriz $M_k(0, 1, \dots, t_k)$ está em B_i , ela tem autovalor duplo $\lambda_i = \lambda_{i+1}$ e $M(\epsilon)$ tem autovalores $\lambda_i(\epsilon)$ e $\lambda_{i+1}(\epsilon)$ simples pois a proposição 4.4 garante que no máximo $n - 1$ matrizes na reta $\{M_{k-1}\} \times \mathbb{R}$ têm autovalores múltiplos. Trataremos em separado os casos em que $v^{(i)}$, o autovetor simples de M_{k-1} , for par ou ímpar. No primeiro caso, comparando os autovalores de M_k e M_{k-1} , vemos que para $M_k(\epsilon)$ estar em $\mathcal{R}_k$, $\lambda_i(\epsilon)$ deve ser ímpar e $\lambda_{i+1}(\epsilon)$ deve ser par. Para tornar possível a coincidência $\lambda_i(0) = \lambda_{i+1}(0)$, ϵ deve ser positivo. Caso $v^{(i)}$ seja ímpar, para que $M_k(\epsilon) \in \mathcal{R}_k$, devemos ter $\lambda_i(\epsilon)$ par e $\lambda_{i+1}(\epsilon)$ ímpar e para isso ϵ precisa ser negativo. Uma vez que t_k cruza a fronteira de B_i , a alternância de autovalores necessária para que a matriz esteja em $\mathcal{R}_n$ só pode ser restabelecida através de uma nova mudança de paridade entre os autovalores λ_i e λ_{i+1} , mas já observamos que a coincidência $\lambda_i = \lambda_{i+1}$ acontece para apenas um valor de t_k . Concluimos que $\mathcal{R}_k$ se localiza inteiramente de um dos lados de B_i e o lado depende da paridade de $v^{(i)}$. Isso estabelece (b).

Para finalizar a prova de (a), passaremos a considerar matrizes M_k na fronteira de $\mathcal{R}_k$ para as quais $M_{k-1} \notin \mathcal{R}_{k-1}$. Dada uma seqüência M_k^n

de matrizes em $\mathcal{R}_k$ convergindo para M_k , M_{k-1}^n tende a $M_{k-1} \in \partial\mathcal{R}_{k-1}$. A compacidade da esfera unitária permite escolher uma subsequência de M_{k-1}^n com autovetores unitários convergentes.

Por indução, M_{k-1} tem autovalor múltiplo λ e $n_{k-1}(\lambda) > n_{k-2}(\lambda)$. A proposição 4.2(b)(iii) garante que o autoespaço $E_{k-1}(\lambda)$ tem um autovetor com primeira coordenada não nula e, usando a base obtida na demonstração da proposição 4.4(d), sabemos que os seus subespaços par e ímpar também têm vetores com a mesma propriedade. Os autovetores de M_{k-1}^n associados aos autovalores que convergem para λ são ortogonais entre si e convergem para alguma base ortogonal de E_{k-1} .

Pelo menos uma das seqüências de autovetores pares converge para um autovetor de M_{k-1} com primeira coordenada não nula, digamos $v_n^{(i)} \xrightarrow{n \rightarrow \infty} v_+$, caso contrário teríamos uma base para o subespaço par de E_{k-1} formada por vetores com primeira coordenada nula. A superfície B_i converge ao longo de M_{k-1}^n e sua altura em M_{k-1} é dada pela fórmula em (a) aplicada ao vetor v_+ , devido à continuidade da fórmula para vetores com primeira coordenada não nula. Se t_i^n é a altura de B_i no ponto M_{k-1}^n , $t_i^n < t_k^n$ pois M_k^n é Landau-regular e por isso está acima de B_i . O mesmo argumento aplicado a autovetores ímpares obtém uma seqüência de vetores ímpares $v_n^{(j)}$ convergindo para um autovetor ímpar v_- com primeira coordenada diferente de zero. A superfície B_j converge ao longo de M_{k-1}^n para o valor dado pela mesma fórmula em (a), aplicada agora ao vetor v_- . Se t_j^n é a altura de B_j em M_{k-1}^n , como $v_n^{(j)}$ é ímpar, $t_j^n > t_k^n$. Tomando o limite de $t_i^n < t_k^n < t_j^n$, temos $h \leq t_k \leq h$. Observe ainda que as superfícies B_i e B_j são contínuas em M_{k-1} pois a escolha de t_k depende apenas de E_{k-1} : t_k é o único valor que aumenta a multiplicidade de λ para $n_k = n_{k-1} + 1$. Então $M_k \in \overline{B_i} \cap \overline{B_j}$.

Seja λ um autovalor de $M_k \in \partial\mathcal{R}_k$. Se $n_{k-1}(\lambda) = 0$, (c) é automaticamente verdadeiro. Se $n_{k-1}(\lambda) = 1$, $n_k(\lambda) = 0$ ou 2. Como assumimos λ autovalor de M_k , $n_k(\lambda)$ deve ser 2, em concordância com (c). Se $n_{k-1} \geq 2$, M_{k-1} está na fronteira de $\mathcal{R}_{k-1}$, e o argumento acima mostra que t_k é determinado pela fórmula de crescimento da multiplicidade, provando (c).

Finalmente, como $\overline{\mathcal{R}_{k-1}}$ é compacto e t_k é localmente limitado por funções contínuas, $\overline{\mathcal{R}_k}$ se mantém compacto. ■