
Universidade Estadual Paulista

Câmpus de São José do Rio Preto

Instituto de Biociências, Letras e Ciências Exatas

Intersecções Homoclínicas

Marcus Augusto Bronzi

Orientador: Prof. Dr. Vanderlei Minori Horita

Dissertação apresentada ao Instituto de Biologia, Letras e Ciências Exatas da Universidade Estadual Paulista, Câmpus São José do Rio Preto, como parte dos requisitos para a obtenção do título de Mestre em Matemática

Bronzi, Marcus Augusto.

Intersecções homoclínicas / Marcus Augusto Bronzi. - São José do
Rio Preto : [s.n], 2006
143 f. : il. ; 30 cm.

Orientador: Vanderlei Minori Horita

Dissertação (mestrado) – Universidade Estadual Paulista, Instituto de
Biociências, Letras e Ciências Exatas

1. Sistemas dinâmicos. 2. Dinâmica hiperbólica. 3. Ferradura de
Smale. 4. Intersecção homoclínica. 5. Bifurcação homoclínica. 6.
Tangência homoclínica. 7. Dinâmica simbólica. I. Horita, Vanderlei
Minori. II. Universidade Estadual Paulista, Instituto de Biociências,
Letras e Ciências Exatas. III. Título.

CDU – 517.93

COMISSÃO JULGADORA

Titulares

Vanderlei Minori Horita

Professor Doutor - IBILCE - UNESP

Orientador

Ali Tahzibi

Professor Doutor - ICMC - USP

1º Examinador

Paulo Ricardo Silva

Professor Doutor - IBILCE - UNESP

2º Examinador

Suplentes

Ali Messaoudi

Professor Doutor - IBILCE - UNESP

1º Suplente

Daniel Smania

Professor Doutor - ICMC - USP

2º Suplente

*Dedico à minha mãe Célia e
à minha companheira Elba.*

Agradecimentos

Agradeço profundamente:

- Primeiramente ao Professor Doutor Vanderlei por sua orientação, pelas oportunidades proporcionadas, por sua grande paciência e amizade sem a qual não teria aprendido a caminhar nesse universo que é a matemática.
- À Fapesp pelo fomento.
- À minha família, Fabiana, Carol e Joãozinho pela força, principalmente à minha mãe, Célia, por nunca me deixar desistir de meus objetivos. Lembro-me até hoje do dia em liguei para ela com os olhos já vermelhos querendo voltar para casa e ela, como faz a águia com seus filhotes, me empurrou para fora do ninho.
- Ao Joaquim por sua grande amizade, copanheirismo, momentos de distração e festança. Também ao João Gabriel pela “brodagem” que sempre tivemos um pelo outro. E à todos os que me apoiaram no tempo que estive na Moradia Estudantil da Unesp.
- À Elba por todas as formas de fraternidade e amor que uma pessoa pode ter pela outra. Sem ela eu perderia a coragem.
- À Osmir, Sônia e Nathan, minha segunda família.
- À todos meus velhos amigos de Ribeirão Preto que apesar da distância sempre me estimularam e elogiaram essa trilha por mim escolhida.

- Aos moradores da República Vacuidade: Júlio Bola de Fogo, Luis Koike, Durval Bói da Fêra. Especialmente ao Júlio, por nossa escalada e por nossa garra.
- À todos meus amigos de mestrado, em especial à minha turma de 2004. Nossa turma sempre unida para esgotar as listas de exercícios, na luta por uma sala de estudos com doze lugares e principalmente para vencer os desafios do editor de texto TeXnicCenter e do Beamer.
- À todos os professores do Departamento de Matemática do IBILCE, em particular ao Cláudio por sua atenção e dedicação, sempre me recebendo em sua sala para sanar quaisquer dúvidas.
- À Cidinha, dona da casa sede de nossa república, e sua família pelo apoio.
- À Deus e a todas as demais pessoas que me ajudaram.

"Os grandes navegadores só
o são por causa das grandes
tormentas", disse um sábio.

Sumário

Introdução	15
1 Órbitas Homoclínicas	20
1.1 Órbitas Homoclínicas em Sistemas Dinâmicos	20
1.2 Órbitas Homoclínicas em uma Aplicação Linear Deformada	22
2 A Ferradura de Smale	28
2.1 Conjuntos Recorrentes por Cadeias e a Ferradura de Smale	32
2.2 Análise das Variedades Estável e Instável da Ferradura de Smale	36
2.3 O <i>Shift</i> Bilateral e a Ferradura de Smale	40
2.4 Bifurcações Homoclínicas	51
3 Conseqüências Dinâmicas de uma Intersecção Homoclínica Trans-	53
versal	
3.1 Coordenadas Linearizantes e o Domínio R	54
3.2 Análise Topológica do Subconjunto Invariante Maximal de R	59
3.3 O Conjunto Λ , Hiperbolicidade e Folheações Invariantes	61
3.4 A Estrutura do Conjunto Λ	71
3.5 Pontos Homoclínicos de Órbitas Periódicas	79
4 Tangências Homoclínicas	81
4.1 Cascatas de Tangências Homoclínicas	82
4.2 Bifurcações de Pontos Fixos	87
4.2.1 Bifurcação Sela-nó	88

4.2.2	Bifurcação de Duplicação de Período	92
4.3	Formas Genéricas das Bifurcações Sela-nó e Flip	99
4.4	Cascatas de Bifurcações de Duplicação de Período e Poços	106
4.5	Renormalização de Tangências Homoclínicas	115
Referências Bibliográficas		132
A Conjunto recorrente por cadeias		134
B Folheações		139

Lista de Símbolos

		A	
$\alpha(x)$	conjunto α -limite, página 134.		
$\underline{a_i}$	é $\underline{0}$ ou $\underline{1}$, as componentes conexas de $R \cap \Psi(R)$, página 60.		
		C	
$C(r)$	campo de cones (contínuo, estável ou instável), página 63.		
C^r	conjunto das aplicações com a r -ésima derivada contínua, página 20.		
		D	
$\partial(X)$	fronteira (ou bordo) do conjunto X , página 70.		
$D\varphi_p$	derivada de φ no ponto p , página 21.		
		E	
$\mathbb{E}^s$	subespaço invariante estável, página 22.		
$\mathbb{E}^u$	subespaço invariante instável, página 22.		
e_i, e'_i	coordenadas linearizantes, página 63.		
		F	
$\mathcal{F}$	folheação (topológica), página 139.		
$\mathcal{F}^u$	folheação instável, página 69.		
		G	
$\Gamma_\mu^s, \Gamma_\mu^u$	pedaços de parábolas, página 84.		
		H	
H_j	faixa horizontal, página 28.		
V_j	faixa vertical, página 28.		
		I	

$\text{int}(X)$	interior do conjunto X , página 64.
L	
ℓ^s, ℓ^u	caminhos contidos nas variedades estável e instável, respectivamente, página 55.
$\ell^s(q, q')$	arcos que liga o ponto q ao q' em $\mathcal{F}^s$, página 72.
$\ell^u(q, q)$	arcos que liga o ponto q ao q em $\mathcal{F}^u$, página 72.
$\hat{\Lambda}$	conjunto φ -invariante correspondente a Λ , página 75.
Λ	ferradura geométrica ou de Smale, página 38.
Λ_A	um A -bloco, página 72.
M	
M_0	caminho de Möbius, página 111.
N	
$\mathcal{N}$	vizinhança em Σ_2 , página 43.
O	
$\mathcal{O}(q)$	órbita do ponto q , página 22.
$\mathcal{O}_r$	vetor zero de $T_r M$, página 64.
$\Omega(f)$	conjunto dos pontos não errantes da aplicação f , página 136.
$\omega(x)$	conjunto ω -limite, página 134.
P	
P	espaço dos pares (x, μ) com $x \in \text{Per}(\varphi_\mu)$, página 109.
$\text{Per}(f)$	conjunto dos pontos periódicos para a aplicação f , página 49.
$\text{Per}(X)$	conjunto dos pontos periódicos de uma aplicação no conjunto X , página 76.
Q	
Q	quadrado unitário do plano, página 28.
R	
$\mathbb{R}^n$	espaço euclidiano n -dimensional, página 21.
$\mathcal{R}(f)$	conjunto recorrente por cadeias para f , página 32.
R	domínio especial (retângulo) que é utilizado na construção da ferradura, página 58.
S	
$\mathbb{S}^2$	espaço $\mathbb{R}^2$ unido com o ponto no infinito p_∞ , página 29.

$\sigma(s)$ aplicação shift bilateral, página 41.

Σ_2 espaço das seqüências de dois símbolos, página 40.

T

$T_q(X)$ espaço tangente a X no ponto q , página 22.

$T_\Lambda M$ espaço (ou fibrado) tangente a Λ , página 68.

W

$W^c, W^c(p)$ variedade central, página 101.

$W^s(p)$ variedade estável de p , quando a função é conhecida, página 20.

$W^s(p, \varphi)$ variedade estável de p para φ , página 20.

$W^u(p)$ variedade instável de p , quando a função é conhecida, página 20.

$W^u(p, \varphi)$ variedade instável de p para φ , página 20.

Z

$\mathbb{Z}$ conjunto dos números inteiros, página 22.

Resumo

Estudamos intersecções homoclínicas de variedades estável e instável de pontos periódicos. Toda intersecção homoclínica produz um comportamento curioso na dinâmica. Nosso modelo de tal fenômeno é a famosa ferradura de Smale, a qual é um conjunto hiperbólico para um difeomorfismo. Além disso, estudamos dinâmica não hiperbólica cuja perda de hiperbolicidade é devido à tangências homoclínicas. Elas têm um papel central na teoria de sistemas dinâmicos. O desdobramento de uma tangência homoclínica produz dinâmicas muito interessantes. Neste trabalho estudamos a criação de cascatas de bifurcações de duplicação de período e um esquema de renormalização para uma tangência homoclínica.

Palavras chave: Intersecção Homoclínica, Ferradura de Smale, Desdobramento de uma Tangência Homoclínica, Bifurcação Sela-nó e Flip.

Abstract

We study homoclinic intersection of stable and unstable manifolds of periodic points. Every homoclinic intersection produce a intricate behavior of the dynamics. Our model of such phenomena is the so called Smale's horseshoe, which is a hyperbolic set for a diffeomorphism. We also study non hyperbolic dynamics whose lack of hyperbolicity is due to homoclinic tangencies. They play a central role in the theory of dynamical systems. The unfolding of a homoclinic tangency produce many interesting dynamics. In this work we study creation of cascade of period doubling bifurcations and a renormalization scheme for a homoclinic tangency.

Key words: Homoclinic Intersection, Horseshoe of Smale, Unfolding of a Homoclinic Tangency, Flip and Saddle-node Bifurcation.

Introdução

Métodos que envolvem o uso de equações diferenciais para descrever como sistemas físicos evoluem tem uma limitação: enquanto as equações diferenciais são suficientes para determinar o comportamento (no sentido que soluções das equações existem), freqüentemente é difícil encontrá-las, mesmo quando a expressão da equação é simples.

Quando soluções podem ser encontradas elas descrevem movimentos muito regulares. A partir da década de 60 um grande número de cientistas depararam-se com um novo tipo de movimento, atualmente chamado de caótico. Esta nova noção de movimento é errático (pelo menos do ponto de vista determinístico): não é apenas quasiperiódico com um grande número de períodos; e não é causado pelo grande número de iterações de elementos. Este tipo de comportamento pode ser detectado mesmo em sistemas muito simples.

Em termos gerais, a teoria de Sistemas Dinâmicos está interessada em descrever, para a maioria dos sistemas, como a maioria das órbitas se comportam, especialmente quando o tempo vai para o infinito. Mais ainda, quando e em que sentido este comportamento é robusto sob pequenas perturbações do sistema.

A formulação matemática de um sistema dinâmico inclui dois ingredientes principais: um conjunto M (geralmente uma variedade diferenciável de dimensão $d \geq 1$), o *espaço de estados*, cujos pontos representam os possíveis estados do sistema; e uma *lei de evolução*, descrevendo como o sistema evolui de um estado a outro. Esta lei de evolução pode ser *tempo discreto*: uma transformação $T : M \rightarrow M$ mapeando

cada estado $x_0 \in M$ a outro estado $x_1 = T(x_0)$ que o sistema estará uma unidade de tempo depois. A sequência $x_0, x_1, x_2, \dots$ definida por $x_{j+1} = T(x_j)$ é a *trajetória* de um *estado inicial* x_0 . Se a transformação T for inversível então temos também a trajetória passada $\dots, x_{-2}, x_{-1}, x_0$ definida por $x_{-j+1} = T(x_{-j})$. Outro modelo de evolução é o *tempo contínuo*, expresso por uma equação diferencial

$$\frac{dx}{dt} = F(x).$$

A trajetória de um estado inicial x_0 é a solução $x(t)$, $t \in \mathbb{R}$ da equação diferencial com $x(0) = x_0$. Assumimos que a solução está definida para todo tempo. De fato sempre consideraremos M uma variedade diferenciável (compacta) e o sistema diferenciável em relação ao tempo e ao espaço M .

Estamos interessados principalmente em:

1. descrever o comportamento da maioria das órbitas da maioria dos sistemas, especialmente quando o tempo vai para o infinito

A descrição de *todas* as órbitas ou sistemas não é, em geral, um objetivo realístico, pois existem muitas formas de comportamentos excepcionais.

2. entender se este comportamento é estável por pequenas modificações da lei de evolução.

A formulação matemática da lei de evolução é quase sempre uma simplificação do processo real. Por isso, é muito importante que as conclusões obtidas dela não seja muito específica: elas devem permanecer válidas para leis próximas.

Estudaremos difeomorfismos em superfícies que apresentam intersecções ou órbitas homoclínicas isto é, órbitas que no passado e no futuro possuem o mesmo conjunto limite. Em geral, órbitas homoclínicas transversais estão associadas à complexidade dinâmica de um sistema (uniformemente) hiperbólico, e, a presença de uma tangência homoclínica é uma obstrução para que um sistema seja hiperbólico.

Além disso, as bifurcações homoclínicas possuem um papel central na teoria de Sistemas Dinâmicos.

Um marco de referência fundamental na evolução das equações diferenciais que descrevem um sistema é o trabalho de Poincaré “Mémoire sur les courbes définies par une équation différentielle” (1881) no qual são lançadas as bases da Teoria Qualitativa das Equações Diferenciais. Tendo início com Poincaré em 1890, com seu ensaio sobre a estabilidade do sistema solar, as intersecções homoclínicas já eram associadas à complexidade do sistema. Em 1935, Birkhoff mostra que próximo de uma órbita homoclínica existe uma infinidade de órbitas periódicas com períodos muito grandes. Passando da teoria abstrata para a aplicada, Van der Pol, em 1920, introduziu uma classe de equações que descreviam o comportamento de circuitos eletrônicos que empregavam tubos de vácuo e que foram usados nos primeiros rádios. Ele mostrou que esses circuitos possuíam oscilações estáveis, que hoje são conhecidas como ciclos limites. Em 1927, com seu amigo Van der Mark, Van der Pol descobrem o caos determinístico. Em 1937 Andronov e Pontrjagin introduzem o conceito de estabilidade estrutural e Maurício Peixoto (1958-62) dá a caracterização, a abertura e densidade das equações diferenciais estruturalmente estáveis em superfícies, o que constitui um marco fundamental para o desenvolvimento contemporâneo das equações diferenciais. Até a década de 50 os pesquisadores já haviam encontrado soluções extremamente complicadas de equações como as de Van der Pol. Mas foi na década de 60 que *Smale*, com um simples exemplo geométrico, descreveu o comportamento complicado que uma órbita homoclínica pode gerar. Esse exemplo é conhecido como a Ferradura Geométrica de Smale. Nasce então a chamada teoria hiperbólica ou dinâmica hiperbólica. Ainda na década de 60, Newhouse combina as bifurcações homoclínicas com a complexidade da teoria hiperbólica para obter sistemas dinâmicos mais complicados que os hiperbólicos. Na década de 80, Levi reanalisa as equações de Van der Pol usando agora a teoria hiperbólica e prova que existem soluções para tais equações.

Começamos nosso trabalho descrevendo, no primeiro capítulo, o que é uma órbita homoclínica e damos um exemplo onde ocorre uma órbita homoclínica (transversal),

o qual denominamos aplicação linear deformada. Depois, no Capítulo 2, veremos a Ferradura de Smale e suas propriedades tais como a existência de um subconjunto invariante maximal, o qual denotamos por Λ , como é a configuração de suas variedades estável e instável. Ainda nesse capítulo, veremos também que os pontos periódicos da aplicação f que define a ferradura são densos em Λ e que existe uma *representação simbólica* para a aplicação f restrita ao conjunto Λ (uma *dinâmica simbólica*), ou seja, o comportamento das órbitas dos pontos em Λ são de certo modo equivalentes ao comportamento das órbitas de pontos no *espaço de símbolos* Σ_2 , que é o conjunto das seqüências bilaterais com entradas 0 e 1. Desta equivalência conseguimos obter mais propriedades do conjunto Λ , como por exemplo, que a aplicação f possui dois pontos fixos, p e $\tilde{p}$, em Λ e que o conjunto dos pontos homoclínicos ao ponto fixo p é denso em Λ . Finalizamos o capítulo definindo o conceito de *bifurcação homoclínica* e dando um exemplo geométrico.

No Capítulo 3, estudamos as conseqüências dinâmicas de uma bifurcação homoclínica transversal. A principal é a criação de uma ferradura. Além disso, estudaremos as propriedades que decorrem da existência dessa ferradura. Ela possui um subconjunto invariante maximal Λ que provaremos ser hiperbólico utilizando folheações invariantes. Uma intersecção homoclínica pode não ser transversal, nesse caso dizemos que é uma *tangência homoclínica*. No Capítulo 4, veremos que em famílias a um parâmetro de difeomorfismos $\{\varphi_\mu\}$ uma tangência homoclínica quadrática (sob certas condições) associada a um ponto fixo p_0 para aplicação φ_0 é acumulada por tangências homoclínicas dessa mesma família. Estudaremos bifurcações nessas famílias a um parâmetro de difeomorfismos tais como a sela-nó e o flip e suas formas genéricas, isto é, que ocorre na maioria dos casos. Depois veremos que o desdobramento de uma bifurcação homoclínica pode criar uma ferradura e que durante a criação de uma ferradura existem infinitos poços ou fontes e bifurcações de duplicação de período. Finalmente, vamos comparar a ferradura criada por uma tangência homoclínica com a família quadrática (funções reais que dependem de um parâmetro μ que têm a forma $y \mapsto y^2 + \mu$). A principal ferramenta utilizada será a renormalização. O Apêndice A complementa o Capítulo 2 e o Apêndice B

complementa o Capítulo 3.

Colocando em forma de tópicos, temos que os seguintes fenômenos ocorrem durante o desdobramento de uma tangência homoclínica:

- órbitas homoclínicas transversais, que estão associadas a ferraduras (conjuntos de Cantor invariantes hiperbólicos),
- cascatas de tangências homoclínicas, isto é, seqüências de parâmetros reais cujos correspondentes difeomorfismos apresentam uma tangência homoclínica e
- para famílias de difeomorfismos localmente dissipativas (isto é, onde o determinante do jacobiano no ponto homoclínico tem valor absoluto menor que 1), temos cascatas de bifurcações de duplicação de período de atratores periódicos (poços).

Além desses fenômenos, existem outros que não estudamos e que podem ser encontrados em [5]:

- cascatas de ciclos de selas-nó críticos,
- subconjuntos residuais de intervalos contidos na reta dos parâmetros cujos correspondentes difeomorfismos apresentam infinitos poços coexistindo um atrator estranho como o atrator de Hénon,
- subconjuntos com medida de Lebesgue positiva de intervalos contidos na reta dos parâmetros cujos correspondentes difeomorfismos apresentam um atrator estranho como o atrator de Hénon,
- prevalência da hiperbolicidade quando a dimensão fractal (de Hausdorff) do conjunto básico associado é menor do que 1 e
- a não prevalência da hiperbolicidade quando dimensão de fractal acima é menor do que 1.

Capítulo 1

Órbitas Homoclínicas

Nesse capítulo vamos rever as definições de *variedade estável* e *instável*, veremos também a definição de pontos *homoclínicos* e um exemplo no qual ocorre a intersecção da variedade estável com a instável transversalmente, ou seja, aparece um ponto *homoclínico transversal*.

1.1 Órbitas Homoclínicas em Sistemas Dinâmicos

Vamos agora recordar as definições de *ponto fixo hiperbólico* e de *variedades estável* e *instável*.

Definição 1.1.1. Seja $\varphi : M \longrightarrow M$ um difeomorfismo de classe C^r e seja p um ponto na variedade M . Dizemos que p é um *ponto fixo hiperbólico* se $\varphi(p) = p$ e se $D\varphi_p$ não possui autovalores de norma um. Para esse ponto fixo hiperbólico p , definimos as *variedades estável* e *instável* por

$$W^s(p, \varphi) = \{x \in M \mid \varphi^i(x) \longrightarrow p \text{ quando } i \rightarrow \infty\}$$

$$W^u(p, \varphi) = \{x \in M \mid \varphi^i(x) \longrightarrow p \text{ quando } i \rightarrow -\infty\}$$

Quando não existe dúvida a respeito da aplicação que define as variedades estável e instável denotamos apenas por $W^s(p)$ e $W^u(p)$.

Observação 1. No caso em que p é um ponto periódico de período k , para algum $k \in \mathbb{N}$, isto é, $\varphi^k(p) = p$ e k é o menor número natural que satisfaz essa condição, definimos a hiperbolicidade e as variedades estável e instável trocando φ por φ^k na definição anterior.

De acordo com o Teorema da Variedade Estável, as variedades $W^s(p)$ e $W^u(p)$ são subvariedades imersas injetivamente em M , têm a mesma classe de diferenciabilidade que φ e têm dimensões igual ao número de autovalores de $D\varphi_p$ com norma menor que um e maior que um, respectivamente.

Exemplo 1. Considere $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}^n$ uma aplicação linear sem autovalores de norma um, digamos com m ($0 \leq m \leq n$) autovalores de norma menor que um. Logo, a origem $\mathbf{0}$ é um ponto fixo hiperbólico (por linearidade) e $W^s(\mathbf{0})$ e $W^u(\mathbf{0})$ são subespaços lineares complementares

$$\mathbb{R}^n = W^s(\mathbf{0}) \oplus W^u(\mathbf{0}),$$

onde a dimensão de $W^s(\mathbf{0})$ é m e a dimensão de $W^u(\mathbf{0})$ é, logicamente, $n - m$.

Em se tratando de uma aplicação linear, o estudo da dinâmica é simples. Porém, a maioria das aplicações são não lineares e, nesses casos, podemos encontrar comportamentos “estranhos” na dinâmica como, por exemplo, as variedades estável e instável intersectarem-se formando o que chamamos de uma *intersecção homoclínica*.

Definição 1.1.2. Seja $\varphi : M \rightarrow M$ uma aplicação de classe C^1 e $p \in M$ um ponto fixo hiperbólico para φ . Dizemos que $q \in M$ é um *ponto homoclínico* para p se $q \neq p$ e

$$q \in W^s(p) \cap W^u(p),$$

isto é, $\lim_{i \rightarrow \pm\infty} \varphi^i(q) = p$.

Dizemos que q é um ponto *homoclínico transversal* se $W^s(p)$ e $W^u(p)$ interceptam-se transversalmente em q , isto é, se

$$T_q M = T_q(W^s(p)) \oplus T_q(W^u(p)).$$

Aqui, $T_q(W^\iota(p))$ quer dizer o espaço gerado (em $T_p M$) pelos vetores tangentes a $W^\iota(p)$ no ponto q , para $\iota = s, u$.

Observação 2. Seja φ um difeomorfismo, p um ponto fixo hiperbólico e q um ponto homoclínico para p . Logo, para todo $n \in \mathbb{Z}$, $\varphi^n(q)$ é também homoclínico para p , pois as variedades estável e instável são invariantes. Assim, dizemos que a órbita $\mathcal{O}(q)$ de q é uma *órbita homoclínica*.

Observe que os difeomorfismos lineares não possuem pontos homoclínicos. Suponha que $\varphi : E \rightarrow E$ é um difeomorfismo linear e que a origem $\mathbf{0}$ é seu (único) ponto fixo hiperbólico. Então, pela linearidade e pela bijetividade de φ , $W^s(\mathbf{0}) = \mathbb{E}^s$ e $W^u(\mathbf{0}) = \mathbb{E}^u$. Logo,

$$E = T_{\mathbf{0}}E = \mathbb{E}^s \oplus \mathbb{E}^u = W^s(\mathbf{0}) \oplus W^u(\mathbf{0}),$$

ou seja, $W^s(\mathbf{0}) \cap W^u(\mathbf{0}) = \{\mathbf{0}\}$, pela definição de soma direta. Portanto, não existem pontos na intersecção entre as variedades estável e instável, exceto a origem.

1.2 Órbitas Homoclínicas em uma Aplicação Linear Deformada

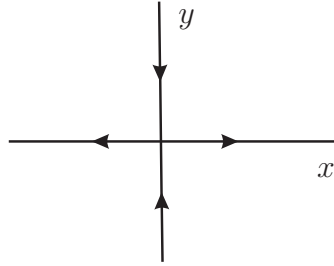
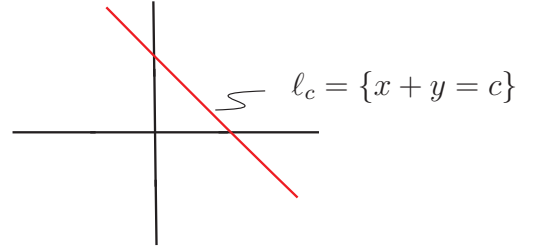
Nesta seção veremos um exemplo de aplicação na qual ocorre uma órbita homoclínica. Esta provém da composição de uma transformação linear com uma outra função que modifica a configuração da aplicação original fazendo com que suas variedades estável e instável se interceptem.

Exemplo 2. Seja $\varphi : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ uma aplicação linear dada por $\varphi(x, y) = (2x, \frac{y}{2})$. Claramente, $(0, 0)$ é ponto fixo hiperbólico de φ . Vamos ver como são as variedades estável e instável de $p = (0, 0)$ usando diretamente as definições

$$W^s(0, 0) = \{(x, y) \in \mathbb{R}^2 \mid \varphi^i(x, y) \rightarrow (0, 0) \text{ quando } i \rightarrow \infty\}$$

$$W^u(0, 0) = \{(x, y) \in \mathbb{R}^2 \mid \varphi^i(x, y) \rightarrow (0, 0) \text{ quando } i \rightarrow -\infty\}.$$

Note que $\varphi^i(x, y) = (2^i x, \frac{y}{2^i})$. Logo, a única maneira de fazer $2^i x \rightarrow 0$, quando $i \rightarrow \infty$, é exigindo que $x = 0$. Por outro lado, temos que $\frac{y}{2^i} \rightarrow 0$ para qualquer y em $\mathbb{R}$. Logo, $W^s(0, 0) = \{(x, y) \in \mathbb{R}^2 \mid x = 0\}$. De forma semelhante, como $\varphi^{-i}(x, y) = (\frac{x}{2^i}, 2^i y)$, a única maneira de fazer $2^i y \rightarrow 0$, quando $i \rightarrow \infty$, é exigindo que $y = 0$, enquanto que para todo $x \in \mathbb{R}$, temos $\frac{x}{2^i} \rightarrow 0$. Logo, temos que $W^u(0, 0) = \{(x, y) \in \mathbb{R}^2 \mid y = 0\}$. Portanto, a variedade estável é o eixo y e a variedade instável é o eixo x . Veja a Figura 1.1.

Figura 1.1: Eixos x e y .Figura 1.2: Reta ℓ_c .

Agora, considere a composição $\psi \circ \varphi$, com $\psi : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ um difeomorfismo da forma

$$\psi(x, y) = (x - f(x + y), y + f(x + y)),$$

onde f é alguma função diferenciável. Com isso, temos que ψ translada pontos sobre retas da forma $\ell_c = \{(x, y) \in \mathbb{R}^2 : x + y = c\}$ a uma distância que depende somente da constante real c . Na verdade, se $(x, y) \in \ell_c$, então

$$\|\psi(x, y) - (x, y)\| = \|(x - f(c), y + f(c)) - (x, y)\| = \sqrt{2}f(c).$$

Devemos colocar algumas hipóteses adicionais na aplicação diferenciável f . Suponhamos que $f \equiv 0$ em $(-\infty, 1]$ e que $f(2) > 2$ (no restante do domínio f é qualquer aplicação diferenciável, por isso temos uma infinidade de aplicações nesse exemplo). Assim, as variedades estável e instável $W^s((0, 0), \psi \circ \varphi)$ e $W^u((0, 0), \psi \circ \varphi)$ interceptam-se em um ponto diferente da origem. Vamos justificar esse fato com algumas afirmações.

Afirmção 1. As seguintes inclusões ocorrem

$$\{(x, y) \in \mathbb{R}^2 \mid x = 0, y \leq 2\} \subset W^s((0, 0), \psi \circ \varphi) \quad (1.1)$$

$$\{(x, y) \in \mathbb{R}^2 \mid x \leq 1, y = 0\} \subset W^u((0, 0), \psi \circ \varphi). \quad (1.2)$$

De fato: como $f((-\infty, 1]) = 0$, segue que $\psi = id$ restrita ao conjunto $\{(x, y) \in \mathbb{R}^2 \mid x + y \leq 1\}$. Logo, $\psi = id$ restrita aos conjuntos

$$W^s((0, 0), \varphi) \cap (\{0\} \times (-\infty, 1]) \quad \text{e} \quad W^u((0, 0), \varphi) \cap ((-\infty, 1] \times \{0\}),$$

o que prova (1.2). Para justificar (1.1), resta ver que $\{0\} \times (1, 2] \subset W^u((0, 0), \psi \circ \varphi)$.

Observe que

$$\psi \circ \varphi(0, y) = \left(-f\left(\frac{1}{2} \cdot y\right), \frac{1}{2} \cdot y + f\left(\frac{1}{2} \cdot y\right) \right).$$

Assim, se $y \in (1, 2]$, então $\frac{1}{2} \cdot y \in \left(\frac{1}{2}, 1\right]$ e $f\left(\frac{1}{2}y\right) = 0$. Com isso, obtemos

$$\psi \circ \varphi(0, y) = \left(0, \frac{y}{2} \right),$$

para todo $y \in (1, 2]$. Além disso, $(\psi \circ \varphi)^2(0, y) = \left(0, \frac{y}{4} \right)$ e por indução

$$(\psi \circ \varphi)^n(0, y) = \left(0, \frac{y}{2^n} \right) \xrightarrow{n \rightarrow \infty} (0, 0),$$

para todo $y \in (1, 2]$. O que prova (1.1).

Afirmção 2. $\psi(\{(x, y) \in \mathbb{R}^2 \mid 1 \leq x \leq 2, y = 0\}) \subset W^u((0, 0), \psi \circ \varphi)$.

De fato: note que ψ é inversível com inversa

$$\psi^{-1}(x, y) = (x + f(x + y), y - f(x + y)),$$

pois,

$$\begin{aligned}
 \psi \circ \psi^{-1}(x, y) &= \psi(x + f(x + y), y - f(x + y)) \\
 &= (x + f(x + y) - f(x + y + 0), y - f(x + y) + f(x + y + 0)) \\
 &= (x, y);
 \end{aligned}$$

analogamente, obtemos $\psi^{-1} \circ \psi = id$.

Seja $(a, b) \in \psi(\{(x, y) \in \mathbb{R}^2 \mid 1 \leq x \leq 2, y = 0\})$. Logo, $(a, b) = \psi(x, y)$ com $1 \leq x \leq 2$ e $y = 0$. Assim,

$$\begin{aligned}
 (\psi \circ \varphi)^{-1}(a, b) &= \varphi^{-1} \circ \psi^{-1} \circ \psi(x, 0) \\
 &= \varphi^{-1}(x, 0) \\
 &= \left(\frac{x}{2}, 0\right).
 \end{aligned}$$

Suponhamos por indução em n que $(\psi \circ \varphi)^{-n}(a, b) = \left(\frac{x}{2^n}, 0\right)$ e provemos que essa igualdade também é válida para $n + 1$. Como $x \in [1, 2]$, $f\left(\frac{x}{2^n}\right) = 0$, logo

$$\begin{aligned}
 (\psi \circ \varphi)^{-n+1}(a, b) &= (\psi \circ \varphi)^{-1} \circ (\psi \circ \varphi)^{-n}(a, b) \\
 &= \varphi^{-1}\left(\psi^{-1}\left(\frac{x}{2^n}, 0\right)\right) \\
 &= \varphi^{-1}\left(\frac{x}{2^n} + f\left(\frac{x}{2^n}\right), -f\left(\frac{x}{2^n}\right)\right) \\
 &= \varphi^{-1}\left(\frac{x}{2^n}, 0\right) \\
 &= \left(\frac{x}{2^{n+1}}, 0\right).
 \end{aligned}$$

Assim, obtemos que

$$(\psi \circ \varphi)^{-n}(a, b) = \left(\frac{x}{2^n}, 0\right) \xrightarrow{n \rightarrow \infty} (0, 0),$$

para todo $x \in [1, 2]$. Portanto, vale a Afirmação 2.

Agora, considere o caminho $\gamma(t) = (1+t, 0)$, com $t \in [0, 1]$, que liga os pontos $(0, 1) = \gamma(0)$ e $(2, 0) = \gamma(1)$ que, pela Afirmação 1, está contido em $W^s((0, 0), \psi \circ \varphi)$. Como ψ é contínua, segue que $\psi \circ \gamma$ é um caminho que liga os pontos

$$\begin{aligned}\psi \circ \gamma(0) &= \psi(1, 0) = (1, 0) \quad \text{e} \\ \psi \circ \gamma(1) &= \psi(2, 0) = (2 - f(2), f(2)),\end{aligned}$$

onde sabemos que $2 - f(2) < 0$ e $f(2) > 0$, pela definição de f . Logo, $\psi \circ \gamma(0)$ está à direita do eixo y e $\psi \circ \gamma(1)$ está à esquerda do eixo y . Veja a Figura 1.3. Portanto, pelo Teorema da Alfândega, existe $\alpha \in (0, 1)$ tal que $\psi \circ \gamma(\alpha)$ pertence ao eixo y . Aqui, pela Afirmação 2, sabemos que o caminho $\psi \circ \gamma$ está contido em $W^s((0, 0), \psi \circ \varphi)$. Além disso, temos que ψ preserva retas do tipo $\ell_c = \{x + y = c\}$, pois

$$\psi(\ell_c) = \{\psi(x, y) \mid x + y = c\} = \{(x - f(c), y + f(c)) \mid x + y = c\}$$

e chamando $u = x - f(c)$ e $v = y + f(c)$, obtemos que

$$\begin{aligned}\psi(\ell_c) &= \{(u, v) \mid u = x - f(c), v = y + f(c), x + y = c\} \\ &= \{(u, v) \mid u + v = x - f(c) + y + f(c) = x + y = c\} \\ &= \{(u, v) \mid u + v = c\} = \ell_c.\end{aligned}$$

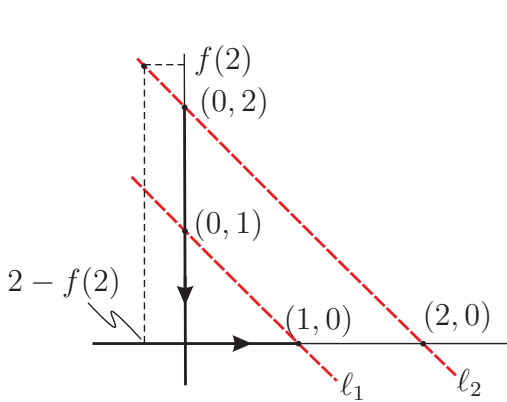


Figura 1.3: Ponto $(2 - f(2), f(2))$.

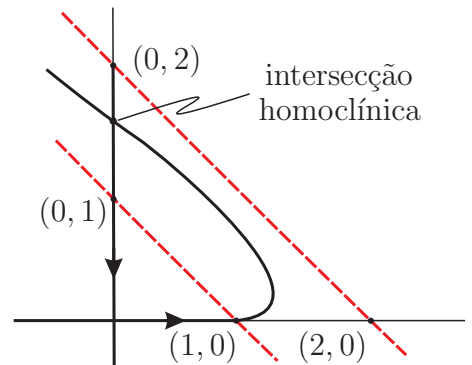


Figura 1.4: Intersecção homoclínica.

Como $(2, 0) \in \ell_2$, então $\psi(2, 0) \in \ell_2$. Pelo fato de ψ ser bijetiva (pois é inversível) e o caminho γ não interceptar ℓ_1 nem ℓ_2 em pontos diferentes de $(1, 0)$ e $(2, 0)$, respectivamente, concluimos que $\psi(\gamma(\alpha))$ está no eixo y entre os pontos $(0, 1)$ e $(0, 2)$, ou seja, que pertence a $W^u((0, 0), \psi \circ \varphi)$. Portanto, $\psi(\gamma(\alpha))$ é uma intersecção homoclínica. Veja a Figura 1.4 acima.

Capítulo 2

A Ferradura de Smale

Neste capítulo vamos estudar a chamada *ferradura de Smale* ou *ferradura geométrica*. Recebe este nome, pois foi introduzida por Smale em 1965. É o principal exemplo que apresenta dinâmica caótica sobre um conjunto invariante que é um conjunto de Cantor.

Exemplo 3. Seja $Q = [0, 1] \times [0, 1]$ o quadrado unitário. Seja H_0 e H_1 duas faixas horizontais na forma

$$H_j = \{(x, y) \in Q \mid 0 \leq x \leq 1, y_1^j \leq y \leq y_2^j\},$$

com $0 \leq y_1^0 < y_2^0 < y_1^1 < y_2^1 \leq 1$, para $j = 0, 1$. E sejam V_1 e V_2 faixas verticais dadas por

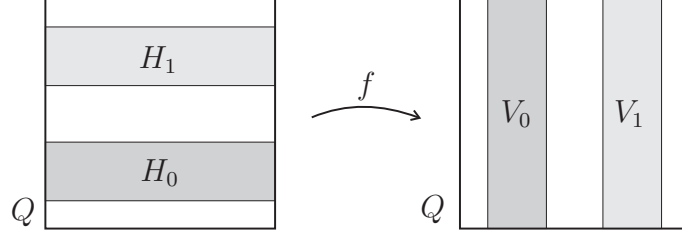
$$V_j = \{(x, y) \in Q \mid x_1^j \leq x \leq x_2^j, 0 \leq y \leq 1\},$$

com $0 \leq x_1^0 < x_2^0 < x_1^1 < x_2^1 \leq 1$, para $j = 0, 1$, como na Figura 2.1.

Vamos assumir inicialmente que f é um difeomorfismo entre subconjuntos do $\mathbb{R}^2$ satisfazendo

(i) $f(H_j) = V_j$, para $j = 0, 1$,

(ii) $Q \cap f^{-1}(Q) = H_0 \cup H_1$ e

Figura 2.1: Difeomorfismo f .

(iii) para $p \in H_0 \cup H_1$, tem-se

$$Df_p = \begin{pmatrix} a_p & 0 \\ 0 & b_p \end{pmatrix},$$

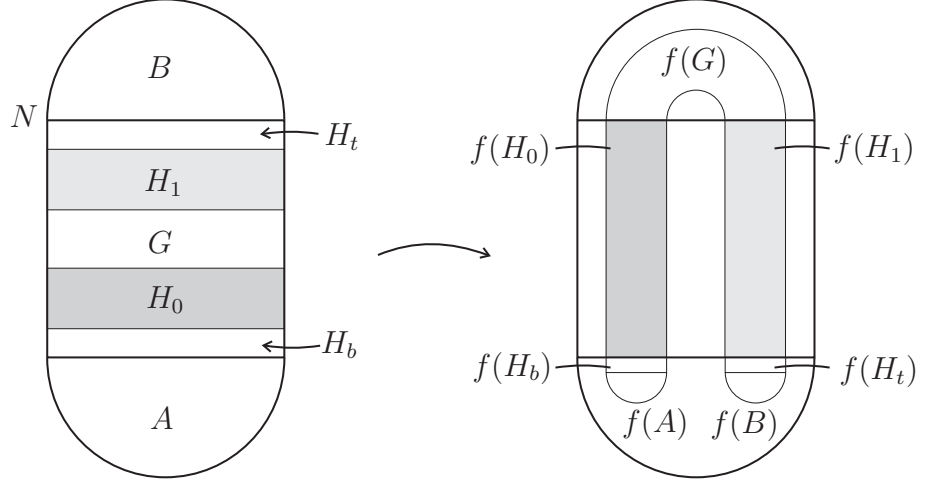
com $|a_p| = \mu < \frac{1}{2}$ e $|b_p| = \lambda > 2$. Veja a figura acima.

Essa aplicação f estende-se para todo o $\mathbb{S}^2$ do seguinte modo, onde $\mathbb{S}^2$ é o espaço $\mathbb{R}^2$ unido com o ponto no infinito. Sejam A o semi-disco inferior de raio $\frac{1}{2}$ localizado abaixo da base de Q , B o semi-disco superior de raio $\frac{1}{2}$ localizado acima do topo de Q e chame de $N = Q \cup A \cup B$ o disco topológico (isto é, um conjunto homeomorfo a um disco comum, por exemplo, a bola unitária centrada na origem) obtido pela união dessas três regiões. Seja G a lacuna horizontal entre H_0 e H_1 , H_t a faixa horizontal no topo de Q acima de H_1 e H_b a faixa horizontal na base de Q abaixo de H_1 . Como na Figura 2.2.

Vamos agora estender nossa aplicação f para uma função que, por abuso de linguagem, continuaremos a denotar por f , e que satisfaz $f(G) \subset B$ e a imagem das fronteiras de G com H_0 e H_1 são dois arcos que começam no topo de V_0 e terminam no topo de V_1 . Tomamos esta parte da aplicação satisfazendo

$$\left| \frac{\partial f}{\partial x}(q) \right| = \mu,$$

para $q \in G$. A função extensão, além disso, deve ter a imagem de $H_b \cup H_t$ contida em A , sendo que a primeira função coordenada de f , restrita à $H_b \cup H_t$, é uma

Figura 2.2: Extensão de f .

contração com fator de contração μ e a segunda função coordenada de f , restrita à $H_b \cup H_t$, muda de uma expansão por um fator λ na fronteira de $H_b \cup H_t$ com $H_0 \cup H_1$ para uma contração na fronteira com $A \cup B$ (isto é, essa função coordenada começa expandindo distâncias próximo dos conjuntos H_0 e H_1 vai diminuindo o fator de expansão até passar a contrair distâncias próximo da fronteira com o conjunto $A \cup B$). Tomamos a extensão tal que $f(A) \subset A$, logo, segue que f possui um ponto fixo p_0 em A , e tal que $f(B) \subset A$. Portanto, a nova função f está definida do disco topológico N em si mesmo.

Podemos estender novamente f para todo o plano $\mathbb{R}^2$, fazendo com que todos os pontos em $\mathbb{R}^2 \setminus N$ entrem no disco topológico N sob um número finito de iterados por f . Finalmente, estendemos f para $\mathbb{S}^2$ exigindo que a nova f deixe fixo o ponto no infinito, o qual denotamos por p_∞ , e que p_∞ seja uma fonte, isto é, que a órbita passada de qualquer ponto em $\mathbb{S}^2 \setminus (N \cup \{p\})$ converge para $\{p_\infty\}$.

Podemos imaginar esta aplicação (final) f , restrita à N , como a composição de duas outras aplicações. A primeira aplicação, que chamaremos de L , pega a região N estica sua altura mais que o dobro de seu tamanho original e comprime sua largura a menos que a metade de seu tamanho original, ou seja, é da forma $L(x, y) = (\mu x, \lambda y)$, onde μ e λ são as constantes definidas anteriormente. A segunda

aplicação, a qual denotaremos por g , pega este “retângulo” mais comprido e mais fino e o coloca sobre N de modo que atravesse Q duas vezes. A aplicação g deve ter as imagens das regiões próximas das extremidades desse retângulo de tal modo que $f = g \circ L$ seja uma contração, restrita ao conjunto A , e tal que $f(B) \subset A$. Veja a figura abaixo.

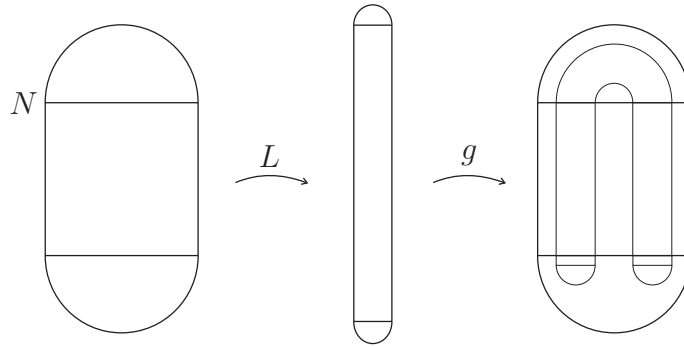


Figura 2.3: Outra construção.

A imagem $f(N)$ está representada, na Figura 2.4, junto com a região N . Como, da definição de f , $f(N) \subset N$, segue que $f^2(N) \subset f(N) \subset N$. Note que $L \circ f(N)$ é uma versão mais comprida e mais fina de $f(N)$ com duas faixas verticais. Logo,

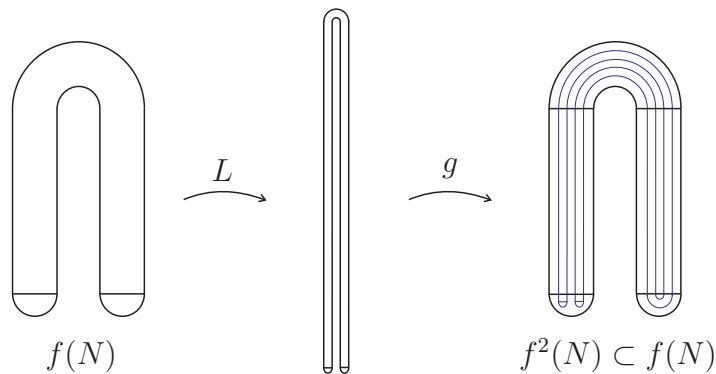


Figura 2.4: A segunda imagem do conjunto N .

$f^2(N) = g \circ L \circ f(N)$ é dobrada ao meio novamente e tem sua imagem dentro de $f(N)$. Portanto, a parte da imagem em Q , $f(N) \cap Q$, consiste de faixas verticais de largura μ^2 .

2.1 Conjuntos Recorrentes por Cadeias e a Ferradura de Smale

Agora, consideremos o *conjunto recorrente por cadeias para f , $\mathcal{R}(f)$* (veja o Apêndice A). Ainda usando as notações da seção anterior, temos que se $q \in A \cup B$, então $f^n(q) \rightarrow p_0$ (o ponto fixo de f em A), quando $n \rightarrow \infty$, pois como já vimos que $f|_{A \cup B}$ é uma contração forte, segue que $\omega(q) = \{p_0\}$. Com isso, podemos fazer a seguinte afirmação

Afirmção 1. $\mathcal{R}(f) \cap (A \cup B) = \{p_0\}$.

De fato: se $q \in B$, então $f(q) \in A$. Como $f(A) \subset A$, $f^n(q)$ permanecerá afastado de q para todo $n \geq 1$

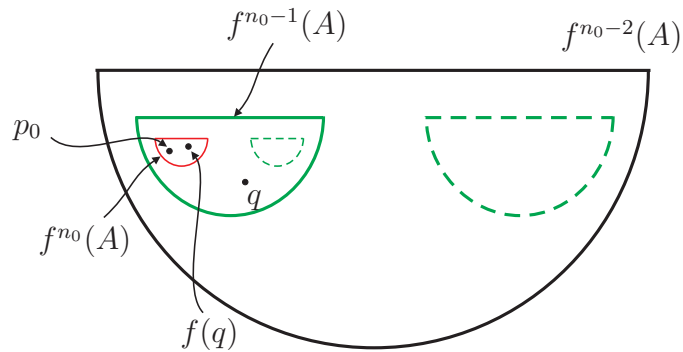


Figura 2.5: Imagens do conjunto A .

Agora, se $q \in A$ e $q = p_0$, então q é recorrente por cadeias pois p_0 é ponto fixo e, nesse caso, $\{q, q\}$ é ε -cadeia de q para q , para todo $\varepsilon > 0$. Suponha que $q \neq p_0$ e considere $r = \frac{d(p_0, q)}{2}$. Como $f|_A$ é contração, existe algum $n_0 \in \mathbb{N}$ tal que o iterado $f^{n_0}(A)$ não contém q e $f^{n_0-1}(A)$ contém q , a saber, basta tomarmos n_0 tal que $\text{diam}(f^{n_0}(A)) < \frac{r}{2}$. Agora, basta tomarmos uma vizinhança de $f(q)$ que não contenha q como, por exemplo, $B(f(q), \varepsilon)$ com ε pequeno o suficiente para que $B(f(q), \varepsilon) \subset \text{int}(f^{n_0}(A))$. Logo, não existem ε -cadeias que começam em q (veja Figura 2.5) e retornam ε -próximo de q . Portanto, $\mathcal{R}(f) \cap (A \cup B) = \{p_0\}$. O que

conclui a afirmação.

Se $q \in \mathbb{S}^2 \setminus N$, então $\lim_{n \rightarrow \infty} f^{-n}(q) = p_\infty$ e, assim, $\alpha(q) = \{p_\infty\}$. Estamos supondo que todos os pontos do conjunto $\mathbb{S}^2 \setminus N$, após um número finito de iterados de f , entram em N e convergem para p_0 , ou seja, que em algum momento entram em A .

Afirmação 2. Fixemos um natural $n > 0$ e $x \in \mathbb{S}^2 \setminus N$. Para todo $\varepsilon > 0$ dado, existe $\delta > 0$ tal que qualquer δ -cadeia $\{x, y_1, \dots, y_n\}$ satisfaz $d(f^n(x), y_n) < \varepsilon$.

De fato: se $n = 1$, para todo $\varepsilon > 0$, basta tomar $\delta < \varepsilon$ que toda δ -cadeia $\{x, y_1\}$ satisfaz $d(f(x), y_1) < \delta < \varepsilon$, por ser uma δ -cadeia.

Seja $n = 2$ e $\varepsilon > 0$ dado arbitrariamente. Como f é contínua, escolhemos $\delta_1 < \frac{\varepsilon}{2}$ tal que

$$d(f(x), z) < \delta_1 \implies d(f^2(x), f(z)) < \frac{\varepsilon}{2}.$$

Como temos válido o caso $n = 1$, para esse dado $\varepsilon > 0$, existe $\delta < \frac{\varepsilon}{2}$ tal que $d(f(x), y_1) < \delta_1$, para toda δ -cadeia $\{x, y_1\}$. Se y_2 é um ponto que estende a δ -cadeia $\{x, y_1\}$, isto é, a delta cadeia torna-se $\{x, y_1, y_2\}$ com

$$d(f(y_1), y_2) < \delta,$$

então temos

$$\begin{aligned} d(f^2(x), y_2) &\leq d(f^2(x), f(y_1)) + d(f(y_1), y_2) \\ &< \frac{\varepsilon}{2} + \delta < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon, \end{aligned}$$

pois $d(f(x), y_1) < \delta_1$ e por ser uma δ -cadeia. (Note que, como $\delta_1 < \varepsilon$, provamos também que $d(f^j(x), y_j) < \varepsilon, j = 1, 2$.)

Suponha por indução que o resultado seja válido para n e provemos que para $n + 1$ também é válido. Da continuidade de f , para o dado $\varepsilon > 0$, segue que existe

$\delta_2 < \frac{\varepsilon}{2}$ tal que

$$d(f^n(x), z) < \delta_2 \implies d(f^{n+1}(x), f(z)) < \frac{\varepsilon}{2}.$$

Por indução, escolha $\delta_0 < \frac{\varepsilon}{2}$ tal que para toda δ_0 -cadeia $\{x, y_1, \dots, y_n\}$ tenha-se

$$d(f^n(x), y_n) < \delta_2.$$

Seja y_{n+1} um ponto que estende a δ_0 -cadeia $\{x, y_1, \dots, y_n\}$, isto é, a nova δ_0 -cadeia é da forma $\{x, y_1, \dots, y_n, y_{n+1}\}$ e y_{n+1} satisfaz

$$d(f(y_n), y_{n+1}) < \delta_0.$$

Logo, dessas três últimas desigualdades segue que

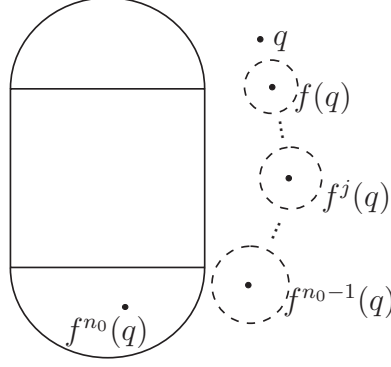
$$\begin{aligned} d(f^{n+1}(x), y_{n+1}) &\leq d(f^{n+1}(x), f(y_n)) + d(f(y_n), y_{n+1}) \\ &< \frac{\varepsilon}{2} + \delta_0 < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon, \end{aligned}$$

o que prova a Afirmação 2. (Observe que, como $\delta_0 < \varepsilon$, também provamos por indução que $d(f^j(x), y_j) < \varepsilon, j = 1, \dots, n+1$.)

Afirmação 3. $\mathcal{R}(f) \cap (\mathbb{S}^2 \setminus N) = \{p_\infty\}$.

De fato: seja $q \in \mathbb{S}^2 \setminus N$. Se $q = p_\infty$, então q é ponto fixo, portanto, $\{q, q\}$ é ε -cadeia de q para q , para todo $\varepsilon > 0$. Suponhamos que $q \neq p_\infty$, então, como foi visto anteriormente, existe $n_0 \in \mathbb{N}$ para o qual $f^{n_0}(q) \in N$. Suponha que n_0 é o primeiro iterado de f tal que $f^{n_0}(q) \in N$. Logo, pela Afirmação 2, para o n_0 tomado acima e para $\varepsilon = \frac{1}{3} \min\{dist(f^j(q), N), j = 1, \dots, n_0 - 1\}$, existe $\delta_0 > 0$ tal que toda δ_0 -cadeia $\{q = y_0, y_1, \dots, y_{n_0-1}\}$ satisfaz

$$d(f^j(q), y_i) < \varepsilon, j = 1, \dots, n_0 - 1.$$

Figura 2.6: Órbita do ponto q .

Logo, temos que

$$\begin{aligned}
 \text{dist}(N, y_j) &\geq \text{dist}(N, f^j(q)) - d(f^j(q), y_j) \\
 &\geq \text{dist}(N, f^j(q)) - \varepsilon \\
 &= \text{dist}(N, f^j(q)) - \frac{1}{3} \left(\min_{1 \leq j \leq n_0-1} \{ \text{dist}(f^j(q), N) \} \right) \\
 &\geq \frac{2}{3} \text{dist}(N, f^j(q)) > 0,
 \end{aligned}$$

para todo $j = 1, \dots, n_0 - 1$. Logo, tais δ_0 -cadeias não interceptam N .

Agora, considere $\eta > 0$ pequeno o suficiente para que $B(f^{n_0}(q), \eta) \subset N$. Como f é contínua, existe δ_1 , $0 < \delta_1 < \frac{\eta}{2}$, tal que

$$d(y, f^{n_0-1}(q)) < \delta_1 \implies d(f(y), f^{n_0}(q)) < \frac{\eta}{2}.$$

Seja $\delta = \min\{\delta_0, \delta_1\}$. Então, para toda δ -cadeia $\{q, y_1, \dots, y_{n_0}\}$, como $d(f^{n_0-1}(q), y_{n_0}) < \delta < \delta_1$, segue que

$$\begin{aligned}
 d(f^{n_0}(q), y_{n_0}) &\leq d(f(f^{n_0-1}(q)), f(y_{n_0-1})) + d(f(y_{n_0-1}), y_{n_0}) \\
 &< \frac{\eta}{2} + \delta < \frac{\eta}{2} + \frac{\eta}{2} = \eta.
 \end{aligned}$$

E como $f(N) \subset N$, segue que para esse $\delta > 0$ toda δ -cadeia começando em q não retorna próximo de q . Portanto, $\mathcal{R}(f) \cap (\mathbb{S}^2 \setminus N) = \{p_\infty\}$.

Finalmente, se $q \in Q \cap \mathcal{R}(f)$, então $f^j(q)$ deverá ficar em Q para todo $j \in \mathbb{Z}$. Caso contrário, isto é, se existe algum $j_0 \in \mathbb{Z}$ tal que $f^{j_0}(q) \in \mathbb{S}^2 \setminus Q$, então $f^n(q) \rightarrow p_0$ ou $f^{-n}(q) \rightarrow p_\infty$ e, portanto, q não é recorrente por cadeias. o que é absurdo.

Portanto,

$$\mathcal{R}(f) \subset \Lambda \cup \{p_0\} \cup \{p_\infty\},$$

onde $\Lambda = \bigcap_{j \in \mathbb{Z}} f^j(Q)$. Adiante veremos que vale mais do que a inclusão, usando dinâmica simbólica veremos que vale a igualdade

$$\mathcal{R}(f) = \Lambda \cup \{p_0\} \cup \{p_\infty\}.$$

2.2 Análise das Variedades Estável e Instável da Ferradura de Smale

Agora, voltemos à análise do conjunto Q . Definimos os conjuntos (de maneira semelhante à construção que se faz para conjuntos de Cantor da reta),

$$Q_m^n = \bigcap_{j=m}^n f^j(Q).$$

Com essa definição, temos que o conjunto $Q_0^1 = Q \cap f(Q)$ é a união de duas faixas verticais, que denotamos por V_0 e V_1 , de largura μ . Além disso, temos que

$$\begin{aligned} Q_0^n &= f(Q_0^{n-1}) \cap Q \\ &= [f(Q_0^{n-1}) \cap V_0] \cup [f(Q_0^{n-1}) \cap V_1] \\ &= f(Q_0^{n-1} \cap H_0) \cup f(Q_0^{n-1} \cap H_1). \end{aligned}$$

Em particular, para $n = 2$, temos

$$\begin{aligned} Q_0^2 &= [f(Q_0^1) \cap V_0] \cup [f(Q_0^1) \cap V_1] \\ &= f([V_0 \cup V_1] \cap H_0) \cup f([V_0 \cup V_1] \cap H_1). \end{aligned}$$

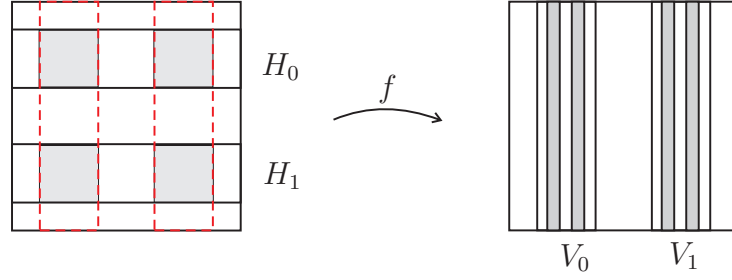


Figura 2.7: Faixas verticais.

Então, para $k = 0, 1$, $f(Q_0^1) \cap V_k = f([V_0 \cup V_1] \cap H_k)$ é a união de duas faixas verticais de largura μ^2 . Portanto, Q_0^2 é a união de 2^2 faixas verticais com largura μ^2 cada. Por indução, obtemos que

$$Q_0^n = f(Q_0^{n-1} \cap H_0) \cup f(Q_0^{n-1} \cap H_1)$$

é a união de 2^n faixas verticais de largura μ^n cada uma. Tomando-se a intersecção infinita, vem que

$$Q_0^\infty = \bigcap_{n=0}^{\infty} Q_0^n = C_1 \times [0, 1]$$

é um conjunto de Cantor de segmentos de reta verticais, onde estamos denotando por C_1 o conjunto de Cantor que obtemos na base (e também no topo) do quadrado Q . Para um ponto $q \in Q_0^\infty$ temos que $q \in f^j(Q)$ e $f^{-j}(q) \in Q$ para todo $j \in \mathbb{N}$. Portanto, Q_0^∞ é o conjunto dos pontos cujos iterados passados permanecem em Q .

Consideremos agora o conjunto Q_{-m}^0 . Logo, $Q_{-1}^0 = H_1 \cup H_2$ é a união de duas faixas horizontais de altura λ^{-1} cada uma. Assim, Q_{-2}^0 é a união de 2^2 faixas horizontais de alturas λ^{-2} .

Prosseguindo como antes, obtemos por indução que Q_{-m}^0 é a união de 2^m faixas

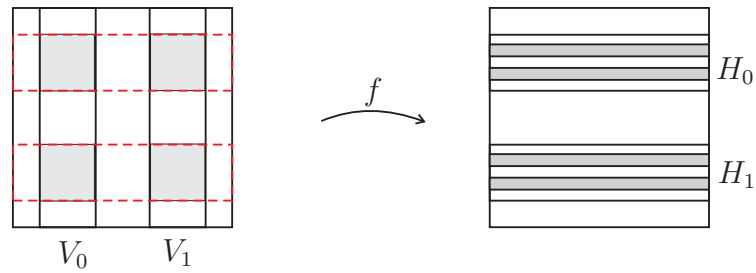


Figura 2.8: Faixas horizontais.

horizontais de alturas λ^{-m} e que

$$Q_{-\infty}^0 = \bigcap_{m=0}^{\infty} Q_{-m}^0 = [0, 1] \times C_2$$

é um conjunto de Cantor de segmentos de reta horizontais. Aqui, C_2 é o conjunto de Cantor obtido no lado esquerdo (e no lado direito) do quadrado Q . Se $q \in Q_{-\infty}^0$, então para todo $j \geq 0$, temos que $q \in f^{-j}(Q)$ e $f^j(q) \in Q$. Portanto, $Q_{-\infty}^0$ é o conjunto dos pontos cujos iterados (futuros) permanecem em Q .

A intersecção dos dois conjuntos Q_0^∞ e $Q_{-\infty}^0$ é o conjunto

$$\Lambda = Q_{-\infty}^\infty = Q_0^\infty \cap Q_{-\infty}^0 = C_1 \times C_2.$$

Assim, Λ é o produto cartesiano de dois conjuntos de Cantor e é caracterizado como o conjunto dos pontos cujas órbitas positivas e negativas permanecem em Q . Além disso, afirmamos o seguinte.

Afirmção 4. O conjunto Λ obtido acima tem as propriedades de um de Cantor da reta, isto é, o conjunto Λ é

- (i) fechado,
- (ii) perfeito e
- (iii) é totalmente desconexo.

De fato: (i) é válido pois Λ é a intersecção das imagens inversas, por funções

contínuas f^j com $j \in \mathbb{Z}$, do conjunto fechado $Q = [0, 1] \times [0, 1]$. Além disso, é compacto por estar contido no conjunto limitado Q . Como os conjuntos C_1 e C_2 são perfeitos, segue que seu produto cartesiano, o conjunto Λ , é perfeito. Logo vale (ii).

Para justificar (iii), basta ver que cada conjunto Q_{-n}^n é a união de 2^{2n} retângulos de dimensões μ^n por λ^{-n} . Desse modo, se Λ tivesse algum subconjunto conexo contendo mais que um ponto, digamos contendo x_1 e x_2 com $d(x_1, x_2) = r$, então, para algum $n \in \mathbb{N}$ grande o suficiente para que $\max\{\mu^n, \lambda^{-n}\} < r^2$, teríamos que x_1 e x_2 pertenceriam a diferentes retângulos de Q_{-n}^n .

Até agora, sabemos que o conjunto Λ assemelha-se com o conjunto de Cantor da reta. Vamos ver agora como são as variedades estável e instável de pontos de Λ . Considere o seguinte conjunto

$$Q_{-\infty}^0 = \bigcap_{m=0}^{\infty} Q_{-m}^0 = [0, 1] \times C_2.$$

Um ponto $p \in Q_{-\infty}^0$ se, e somente se, $q \in f^j(Q)$ para todo $j \leq 0$, ou seja, $f^j(q) \in Q$ para todo $j \geq 0$. Portanto, q pertence à variedade estável de Λ , $W^s(\Lambda)$ (isto é, à união das variedades estável de todos os pontos de Λ).

Afirmção 5. Se $q \in [0, 1] \times C_2$ e $p \in C_1 \times C_2 = \Lambda$ têm a mesma coordenada y , então

$$|f^j(q) - f^j(p)| \leq \mu^j, j \geq 0$$

e $q \in \text{comp}_p(W^s(p) \cap Q)$, a componente conexa de $W^s(p) \cap Q$ que contém p .

De fato: como p e q têm a mesma coordenada y , então ambos estão em um mesmo segmento de reta horizontal contido em $H_1 \cup H_2$. Logo, quando aplicamos f nesses pontos, contraímos distâncias na direção x com um fator de contração μ e expandimos distâncias na direção y com fator de expansão λ . (Logo, é de se esperar que $f^j(p)$ e $f^j(q)$ se aproximem à uma razão μ^j .) Considere o caminho que liga p a q da

forma $\gamma(t) = p + t(p - q)$, $t \in [0, 1]$. Então, pela Desigualdade do Valor Médio,

$$\begin{aligned}
 |f^j(q) - f^j(p)| &= |f^j \circ \gamma(1) - f^j \circ \gamma(0)| \\
 &\leq \sup_{t \in [0, 1]} \|D(f^j \circ \gamma)_t\| \cdot |1 - 0| \\
 &= \sup_{t \in [0, 1]} \|D(f^j)_{\gamma(t)} \gamma'(t)\| \\
 &= \sup_{t \in [0, 1]} \|D(f)_{\gamma(t)}^j (q - p)\| \\
 &= \sup_{t \in [0, 1]} \|a_p^j \pi_1(q - p)\| \leq \mu^j,
 \end{aligned}$$

onde π_1 é a projeção do $\mathbb{R}^2$ na primeira coordenada x . Assim, $q \in W^s(\Lambda)$. Além disso, segue que q está na mesma componente conexa que p em $W^s(\Lambda)$.

Portanto, os segmentos de reta horizontais são as variedades estáveis locais dos pontos de Λ . De forma análoga mostra-se que as variedades instáveis locais de pontos de Λ são os segmentos de reta verticais que passam por esses pontos.

2.3 O *Shift* Bilateral e a Ferradura de Smale

Nessa seção, vamos mostrar uma forma de representar a dinâmica da aplicação f restrita ao conjunto Λ . Para isso, vamos considerar o *shift bilateral* sobre o *espaço de símbolos* constituído por pontos que são seqüências bilaterais de 0's e 1's. Chamamos a aplicação *shift* de uma *representação simbólica* para a aplicação f , e costuma-se dizer que ela dá a *dinâmica simbólica* da aplicação. Vamos ver a definição desses novos objetos.

Definição 2.3.1. O *espaço das seqüências de dois símbolos* Σ_2 é definido por

$$\Sigma_2 = \{0, 1\}^{\mathbb{Z}} = \{s = (\dots, s_{-2}, s_{-1}, s_0, s_1, s_2, \dots) \mid s_j \in \{0, 1\}, j \in \mathbb{Z}\}.$$

Uma métrica em Σ_2 é definida por

$$d(s, t) = \sum_{j=-\infty}^{\infty} \frac{|s_j - t_j|}{2^{|j|}},$$

onde $s = (s_j)_{j \in \mathbb{Z}}$, $t = (t_j)_{j \in \mathbb{Z}}$ e $s_j, t_j \in \{0, 1\}$ para todo $j \in \mathbb{Z}$.

Um fato que segue direto da Definição 2.3.1 é que essa métrica faz de Σ_2 um espaço métrico.

Definição 2.3.2. A aplicação *shift* (ou *shift bilateral*) sobre Σ_2 é definida por $\sigma(s) = t$, onde $t_k = s_{k+1}$, para $t = (t_k)_{k \in \mathbb{Z}}$ e $s = (s_k)_{k \in \mathbb{Z}}$ em Σ_2 . Também denotamos por

$$\sigma(\dots, s_{-2}, s_{-1}, s_0, s_1, s_2, \dots) = (\dots, s_{-1}, s_0, s_1, s_2, s_3, \dots).$$

O espaço Σ_2 junto com a aplicação σ é chamado *shift bilateral*.

Uma consequência imediata da definição da aplicação *shift* é sua injetividade: para $(t_j)_{j \in \mathbb{Z}}$ e $(s_j)_{j \in \mathbb{Z}}$ em Σ_2 ,

$$\sigma((t_j)_{j \in \mathbb{Z}}) = \sigma((s_j)_{j \in \mathbb{Z}}) \Rightarrow (t_{j+1})_{j \in \mathbb{Z}} = (s_{j+1})_{j \in \mathbb{Z}} \Rightarrow (t_j)_{j \in \mathbb{Z}} = (s_j)_{j \in \mathbb{Z}}.$$

Agora vamos ver a prova de que a aplicação σ é contínua, para isso utilizaremos os seguintes lemas.

Lema 2.3.3. Dados $m \in \mathbb{N}$, $m > 1$, e $s = (s_i)_{i \in \mathbb{Z}}, t = (t_i)_{i \in \mathbb{Z}} \in \Sigma_2$ tais que $d(s, t) < \frac{1}{2^m}$, então $s_i = t_i$ para todo $-m + 1 \leq i \leq m - 1$.

Demonstração:

Provaremos por absurdo. Suponha que existe um inteiro j , $-m + 1 \leq j \leq m - 1$, tal que $s_j \neq t_j$. Então,

$$\begin{aligned} d(s, t) &= \sum_{i=-\infty}^{\infty} \frac{|s_i - t_i|}{2^{|i|}} \geq \frac{|s_j - t_j|}{2^{|j|}} \\ &\geq \frac{1}{2^{|j|}} \geq \frac{1}{2^{m-1}} > \frac{1}{2^m}, \end{aligned}$$

o que é um absurdo, pois $d(s, t) < \frac{1}{2^m}$. ■

Lema 2.3.4. *Sejam $s = (s_i)_{i \in \mathbb{Z}}, t = (t_i)_{i \in \mathbb{Z}} \in \Sigma_2$ e $m \in \mathbb{N}$ tais que $s_i = t_i$ para $-m \leq i \leq m$, então $d(s, t) \leq \frac{1}{2^{m-1}}$.*

Demonstração:

Para s e t como no enunciado, temos

$$\begin{aligned} d(s, t) &= \sum_{i=-\infty}^{\infty} \frac{|s_i - t_i|}{2^{|i|}} = \sum_{i=-\infty}^0 \frac{|s_i - t_i|}{2^{|i|}} + \sum_{i=1}^{\infty} \frac{|s_i - t_i|}{2^{|i|}} \\ &= \sum_{i=-\infty}^{-m-1} \frac{|s_i - t_i|}{2^{|i|}} + \sum_{i=m+1}^{\infty} \frac{|s_i - t_i|}{2^{|i|}} \leq \sum_{i=-\infty}^{-m-1} \frac{1}{2^{|i|}} + \sum_{i=m+1}^{\infty} \frac{1}{2^{|i|}} \\ &\leq 2 \sum_{j=m+1}^{\infty} \frac{1}{2^{|j|}} = 2 \cdot \frac{1}{2^{m+1}} \cdot \sum_{j=0}^{\infty} \frac{1}{2^{|j|}} = \frac{1}{2^{m-1}}. \end{aligned}$$

■

Proposição 2.3.5. *A aplicação shift $\sigma : \Sigma_2 \longrightarrow \Sigma_2$ é contínua.*

Demonstração:

Sejam $s = (s_i)_{i \in \mathbb{Z}} \in \Sigma_2$ e $\varepsilon > 0$ dado arbitrariamente. Queremos exibir um $\delta > 0$ tal que se $t \in \Sigma_2$ e $d(s, t) < \delta$, então $d(\sigma(s), \sigma(t)) < \varepsilon$. Tome $m \in \mathbb{N}$ tal que $\frac{1}{2^{m-3}} < \varepsilon$ e seja $\delta = \frac{1}{2^m}$. Logo, pelo Lema 2.3.3 temos

$$d(s, t) < \delta = \frac{1}{2^m} \Rightarrow s_i = t_i, \quad -m+1 \leq i \leq m-1. \quad (2.1)$$

Denotemos $\sigma(s) = (\dots, s_{-1}^*, s_0^*, s_1^*, \dots)$ e $\sigma(t) = (\dots, t_{-1}^*, t_0^*, t_1^*, \dots)$. Então, pelo Lema 2.3.4, temos que (2.1) acarreta que $s_i^* = t_i^*$ para todo $-m+2 \leq i \leq m-2$ e, portanto que

$$d(\sigma(s), \sigma(t)) \leq \frac{1}{2^{m-3}} < \varepsilon.$$

■

Mais do que provado anteriormente, temos que a aplicação shift é um homeomorfismo. Basta notar que sua inversa é a função que desloca uma unidade à direita

e, de maneira semelhante à Proposição 2.3.5, prova-se que essa inversa é contínua.

Definição 2.3.6. Dada uma aplicação $f : X \rightarrow X$, dizemos que a aplicação $g : Y \rightarrow Y$ é *topologicamente conjugada* a f se existe um homeomorfismo $h : X \rightarrow Y$ tal que $h \circ f = g \circ h$.

$$\begin{array}{ccc} X & \xrightarrow{f} & X \\ h \downarrow & & \downarrow h \\ Y & \xrightarrow{g} & Y \end{array}$$

Vamos agora mostrar que a aplicação $f|_{\Lambda} : \Lambda \rightarrow \Lambda$, definida na seção anterior, é topologicamente conjugada a $\sigma : \Sigma_2 \rightarrow \Sigma_2$. Mas antes vamos ver como é a topologia de Σ_2 . A base para a topologia é formada pelas seguintes vizinhanças: dado um ponto $s \in \Sigma_2$, uma vizinhança de s é dada por

$$\mathcal{N} = \{t \in \Sigma_2 \mid t_j = s_j \text{ para } -n_0 \leq j \leq n_0\}.$$

Teorema 1. *Seja $h : \Lambda \rightarrow \Sigma_2$ definida por $h(q) = s$ tal que $f^j(q) \in H_{s_j}$ para todo $j \in \mathbb{Z}$, onde $s = (s_j)_{j \in \mathbb{Z}}$ e H_{s_j} são as faixas horizontais definidas no Exemplo 3. Então, h é uma conjugação topológica de $f|_{\Lambda}$ para $\sigma : \Sigma_2 \rightarrow \Sigma_2$.*

Demonstração:

Mostremos inicialmente que h satisfaz a condição $\sigma \circ h = h \circ f$ em Λ . Sejam $q \in \Lambda$ e $s, t \in \Sigma_2$ tais que $h(q) = s$ e $h(f(q)) = t$. Então, segue que $f^{j+1}(q) \in H_{s_{j+1}}$ para todo $j \in \mathbb{Z}$ e, além disso, $f^{j+1}(q) = f^j \circ f(q) \in H_{t_j}$ para todo $j \in \mathbb{Z}$. Portanto, pela definição de h , temos a seguinte igualdade dos índices de H_i

$$s_{j+1} = t_j, \text{ para todo } j \in \mathbb{Z}.$$

Logo, obtemos que $\sigma(s) = t$, ou seja, que $\sigma(h(q)) = h(f(q))$, para todo $q \in \Lambda$.

Mostremos agora que h é contínua. Seja $q \in \Lambda$, $s \in \Sigma_2$ tal que $h(q) = s$ e considere a seguinte vizinhança de s

$$\mathcal{N} = \{t \in \Sigma_2 \mid t_j = s_j \text{ para } -n_0 \leq j \leq n_0\}.$$

Para $n_0 \in \mathbb{N}$ dado arbitrariamente, queremos exibir um $\delta > 0$ tal que para $p \in \Lambda$,

$$|p - q| < \delta \implies t = h(p) \in \mathcal{N},$$

ou seja, que $f^j(p)$ pertença a H_{s_j} para todo $-n_0 \leq j \leq n_0$. Seja $n_0 \in \mathbb{N}$ dado. Como f é contínua, para cada $-n_0 \leq j \leq n_0$ existe um $\delta_j > 0$ tal que se $p \in \Lambda$ e $|p - q| < \delta_j$, então $f^j(p) \in H_{s_j}$. Tome $\delta = \min\{\delta_j \mid -n_0 \leq j \leq n_0\}$ e temos que se $p \in \Lambda$ e $|p - q| < \delta$, então $f^j(p) \in H_{s_j}$ para todo $j \in \{-n_0, \dots, n_0\}$. Ou seja, $t = h(p)$ satisfaz $t_j = s_j$ para todo $-n_0 \leq j \leq n_0$. Portanto, $h(p) = t \in \mathcal{N}$. Logo, h é contínua.

Mostremos que h é injetora. Sejam $p, q \in \Lambda$ tais que $h(p) = h(q) = s$. Logo, para todo $j \in \mathbb{Z}$, vale $f^{-j}(p), f^{-j}(q) \in H_{s_{-j}}$, ou seja, $p, q \in f^j(H_{s_{-j}})$. Variando j de 1 até ∞ , obtemos que $p, q \in \bigcap_{j=1}^{\infty} f^j(H_{s_{-j}})$, ou seja, que p e q estão no mesmo segmento de

reta vertical. Por outro lado, variando j de $-\infty$ até 0, segue que $p, q \in \bigcap_{j=-\infty}^0 f^j(H_{s_{-j}})$. Assim, p e q estão em um mesmo segmento de reta horizontal. Portanto, variando j de $-\infty$ até ∞ , obtemos que p e q estão na intersecção de dois segmentos transversais, isto é, $p = q$. O que prova a injetividade de h .

Provemos que h é sobrejetiva. Dado $s = (s_i)_{i \in \mathbb{Z}} \in \Sigma_2$, queremos encontrar um $q \in \Lambda$ tal que $h(q) = s$. Pela construção da aplicação $f : \Lambda \rightarrow \Lambda$, temos que as imagens $f^n(Q) \cap Q$ são faixas verticais de largura μ^n para cada $n \geq 0$. O mesmo vale para as imagens $f^n(H_0 \cup H_1) \cap Q$ com $n \geq 1$. Cada uma dessas faixas é da forma $\bigcap_{k=1}^n f^k(H_{s_k})$, com $s_k = 0$ ou 1. Quando n vai para infinito, essa intersecção converge para um segmento de reta vertical. De forma semelhante, concluímos que $\bigcap_{k=-\infty}^n f^k(H_{s_k})$, com $s_k = 0$ ou 1, é uma faixa horizontal. Em particular, para nosso

ponto $s = (s_i)_{i \in \mathbb{Z}}$, temos que $\bigcap_{j=-\infty}^{\infty} f^j(H_{s_{-j}})$ é a intersecção de dois segmentos de reta e, portanto, é um ponto q . Esse ponto q satisfaz $f^{-j}(q) \in H_{s_{-j}}$, ou seja, $f^i(q) \in H_{s_i}$, para todo $i \in \mathbb{Z}$, o que implica que $h(q) = s = (s_i)_{i \in \mathbb{Z}}$ (pela definição de h). Então, h é sobrejetiva.

Logo, h é contínua e bijetiva. Como, além disso, seu domínio Λ é compacto e seu contra-domínio Σ_2 é Hausdorff (pois é espaço métrico), segue que h é um homeomorfismo (veja [3] página 167 Teorema 26.6). ■

Note que a aplicação shift $\sigma : \Sigma_2 \longrightarrow \Sigma_2$ possui apenas dois pontos fixos, a saber

$$s = (\dots, 0, 0, \dot{0}, 0, 0, \dots) \quad \text{e} \quad \tilde{s} = (\dots, 1, 1, \dot{1}, 1, 1, \dots).$$

Logo, o mesmo vale para $f|_{\Lambda}$, pois como h é bijeção, existe $p \in \Lambda$ tal que $h(p) = s$, logo $f(p) = h^{-1} \circ \sigma \circ h(s) = h^{-1}(\sigma(s)) = h^{-1}(s) = p$. Ou seja, p é um ponto fixo para f . O mesmo vale para $\tilde{s}$, existe $\tilde{p} \in \Lambda$ que é ponto fixo para f .

Olhando novamente para p , como temos $h(p) = s$, segue da definição de h que

$$f^j(p) \in H_{s_j}, \forall j \in \mathbb{Z} \implies f^j(p) \in H_0, \forall j \in \mathbb{Z}.$$

Em particular,

$$p \in f^0(H_0) \cap f^1(H_0) = H_0 \cap V_0 \quad \text{e}$$

$$p \in H_0 \cap f(H_0) \cap f^{-1}(H_0) \cap f^2(H_0) = (H_0 \cap f^{-1}(H_0)) \cap (V_0 \cap f(V_0)).$$

Logo, p encontra-se em um local próximo ao indicado na Figura 2.9.

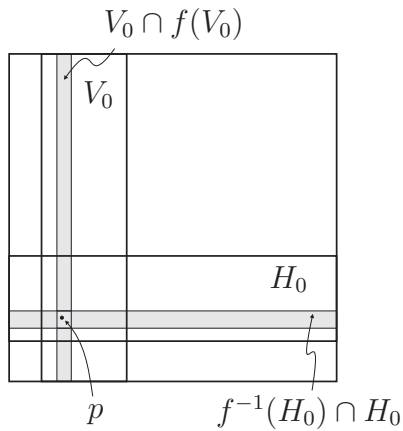


Figura 2.9: Posição de p .

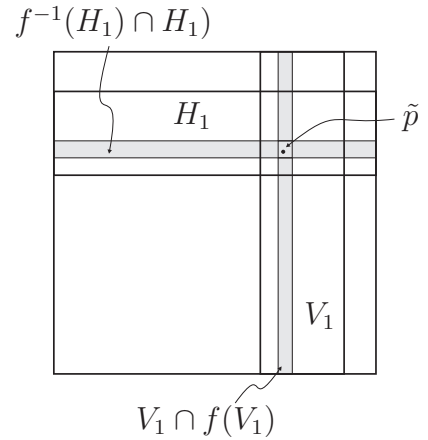


Figura 2.10: Posição de $\tilde{p}$.

Para $\tilde{p}$, temos $h(\tilde{p}) = \tilde{s} = (\dots, 1, \dot{1}, 1, \dots)$, logo da mesma forma que foi feito para p , obtemos

$$\begin{aligned}\tilde{p} &\in f^0(H_1) \cap f^1(H_1) = H_1 \cap V_1 \text{ e} \\ \tilde{p} &\in f^{-1}(H_1) \cap H_1 \cap f(H_1) \cap f^2(H_1) = (f^{-1}(H_1) \cap H_1) \cap (V_1 \cap f(V_1)).\end{aligned}$$

Logo, $\tilde{p}$ encontra-se em uma posição próxima à indicada na Figura 2.10.

Proposição 2.3.7. *O espaço métrico (Σ_2, d) é completo.*

Demonstração: Para provar que espaço métrico Σ_2 é completo, basta mostrar que toda seqüência de Cauchy em Σ_2 converge para um ponto de Σ_2 .

Como o espaço métrico Σ_2 é homeomorfo à Λ que é um espaço métrico compacto, segue que Σ_2 também é um espaço métrico compacto. Além disso, por Σ_2 ser um compacto em um espaço métrico, ele tem a seguinte caracterização: “toda seqüência de Cauchy em Σ_2 admite uma subsequência convergente”. Logo, dada uma seqüência de Cauchy $(s_k)_{k=1}^\infty$ em Σ_2 , esta admite uma subsequência convergente e, portanto, a própria seqüência $(s_k)_{k=1}^\infty$ é convergente em Σ_2 . De fato, suponha que $(s_{k_j})_{j \in \mathbb{N}}$ é a subsequência convergente, digamos para $s \in \Sigma_2$, então dado $\varepsilon > 0$ existe $j_0 \in \mathbb{N}$ tal que para todo $j \geq j_0$ temos

$$d(s_{k_j}, s) < \frac{\varepsilon}{2}. \quad (2.2)$$

Como (s_k) é de Cauchy, existe $k_0 \in \mathbb{N}$ tal que para todo $k, \ell \geq k_0$ temos

$$d(s_k, s_\ell) < \frac{\varepsilon}{2}. \quad (2.3)$$

Logo, tomando-se $k_1 = \max\{k_0, k_{j_0}\}$, segue de (2.2) e (2.3) que se $k \geq k_1$, então

$$d(s_k, s) \leq d(s_k, s_{k_j}) + d(s_{k_j}, s) < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

O que mostra que $s_k \longrightarrow s$, quando $k \rightarrow \infty$, satisfazendo assim a condição que pedimos no início da demonstração. Portanto, Σ_2 é completo. ■

Proposição 2.3.8. *Dados $s = (s_i)_{i \in \mathbb{Z}} \in \Sigma_2$, então $W^s(s, \sigma)$ consiste das seqüências $t = (t_i)_{i \in \mathbb{Z}}$ em Σ_2 tais que existe um $n \in \mathbb{N}$ para o qual $t_i = s_i$ para todo $i \geq n$.*

Demonstração:

Note que cada iterado de um ponto em Σ_2 desloca sua seqüência uma unidade à esquerda. Assim, todo ponto $t \in \Sigma_2$ tal que $t_i = s_i$ para todo $i \geq n$ pertence a $W^s(s, \sigma)$ pois, para $k \geq 2n + 1$, tem-se que $\sigma^k(t)$ e $\sigma^k(s)$ coincidem nas entradas i tais que $-n \leq i \leq n$. Logo, pelo Lema 2.3.4,

$$\lim_{k \rightarrow \infty} d(\sigma^k(t), \sigma^k(s)) \leq \frac{1}{2^{k-1}} = 0.$$

Agora, vamos mostrar que se um ponto t pertence a $W^s(s, \sigma)$, então ele é da forma $t = (\dots, t_{-2}, t_{-1}, t_0, t_1, t_2, \dots, t_{n-1}, s_n, s_{n+1}, s_{n+2}, \dots)$, para algum $n \in \mathbb{N}$. Suponha que existe $t \in W^s(s, \sigma)$ para o qual não existe $n \in \mathbb{N}$ tal que $t_i = s_i$, para todo $i \geq n$, ou seja, que para todo $n \in \mathbb{N}$ existe um $m_n \in \mathbb{N}$ tal que $s_{n+m_n} \neq t_{n+m_n}$. Logo

$$\begin{aligned} d(\sigma^{m_n}(t), \sigma^{m_n}(s)) &= \sum_{i=-\infty}^{\infty} \frac{|t_{i+m_n} - s_{i+m_n}|}{2^{|i|}} \\ &\geq \frac{|t_{0+m_n} - s_{0+m_n}|}{2^{|0|}} = 1, \end{aligned}$$

para todo $n \in \mathbb{N}$. Então, $\lim d(\sigma^n(t), \sigma^n(s)) = \lim(\sigma^{m_n}(t), \sigma^{m_n}(s)) \geq 1$. Portanto, $t \notin W^s(s, \sigma)$, o que é absurdo. ■

Proposição 2.3.9. *Dados $s = (s_i)_{i \in \mathbb{Z}} \in \Sigma_2$, então $W^u(s, \sigma)$ consiste das seqüências $t = (t_i)_{i \in \mathbb{Z}} \in \Sigma_2$ tais que existe um $n \in \mathbb{N}$ para o qual $t_i = s_i$ para todo $i \leq -n$.*

Demonstração:

É análoga à demonstração da Proposição anterior. ■

Uma propriedade para a aplicação $f : \Lambda \longrightarrow \Lambda$ que pode ser herdada via homeomorfismo do shift bilateral é a de ser topologicamente transitiva.

Definição 2.3.10. Dada uma função $f : X \longrightarrow X$, dizemos que f é *topologicamente transitiva* se para quaisquer dois conjuntos abertos $U, V \subset X$ existe um $k > 0$ tal

que $f^k(U) \cap V \neq \emptyset$.

Proposição 2.3.11. *A aplicação $\sigma : \Sigma_2 \longrightarrow \Sigma_2$ é topologicamente transitiva.*

Demonstração:

Dados dois conjuntos abertos U e V contidos em Σ_2 , considere uma bola de raio $\delta > 0$ e centro $s = (\dots, s_{-2}, s_{-1}, s_0, s_1, s_2, \dots)$ totalmente contida em U e outra de raio $\varepsilon > 0$ e centro $t = (\dots, t_{-2}, t_{-1}, t_0, t_1, t_2, \dots)$ totalmente contida em V . Sejam $m, k \in \mathbb{N}$ tais que $\frac{1}{2^{m-1}} < \delta$ e $\frac{1}{2^{k-1}} < \varepsilon$. Considere o ponto

$$p = (\dots, p_{-m-1}, s_{-m}, \dots, s_{-1}, \dot{s}_0, s_1, \dots, s_m, t_{-k}, t_{-k+1}, \dots, t_0, t_1, \dots, t_k, p_{m+k+1}, \dots)$$

que pertence a $B(s, \delta)$, pois como o ponto p coincide com o ponto s nas entradas $-m, -m+1, \dots, m-1, m$ segue da Proposição 2.3.4 que $d(s, p) \leq \frac{1}{2^{m-1}} < \delta$. De modo semelhante, a Proposição 2.3.4 também garante que o ponto $\sigma^{m+k+1}(p)$ na forma

$$(\dots, p_{-m-1}, s_{-m}, \dots, s_{-1}, s_0, s_1, \dots, s_m, t_{-k}, t_{-k+1}, \dots, \dot{t}_0, t_1, \dots, t_k, p_{m+k+1}, \dots)$$

pertence a $B(t, \varepsilon)$, o que implica que $p \in U \cap \sigma^{m+k+1}(V)$, ou seja, $U \cap \sigma^{m+k+1}(V) \neq \emptyset$. ■

Recordemos que um ponto q é homoclínico de um dado ponto periódico p se $q \neq p$ e $q \in W^s(p) \cap W^u(p)$. Com esse conceito em mente, temos o seguinte resultado.

Proposição 2.3.12. *O conjunto $\mathcal{H}$ dos pontos homoclínicos do ponto fixo $\mathbf{0} = (\dots, 0, \dot{0}, 0, \dots)$ é denso em Σ_2 .*

Demonstração:

Inicialmente, pelas Proposições 2.3.8 e 2.3.9, temos que os pontos homoclínicos têm a seguinte forma

$$s = (\dots, 0, 0, 0, s_{-n}, \dots, s_{-1}, s_0, s_1, \dots, s_n, 0, 0, 0, \dots).$$

Dado $t = (\dots, t_{-1}, t_0, t_1, \dots) \in \Sigma_2$ e $\varepsilon > 0$ vamos exibir um ponto $p \in \mathcal{H}$ tal que $d(t, p) < \varepsilon$. Tome $n \in \mathbb{N}$ grande o suficiente para que $\frac{1}{2^{n-1}} < \varepsilon$ e considere o seguinte ponto homoclínico de Σ_2

$$p = (\dots, 0, t_{-n}, \dots, t_{-1}, t_0, t_1, \dots, t_n, \dots).$$

Pela Proposição 2.3.4, segue que $d(t, p) \leq \frac{1}{2^{n-1}} < \varepsilon$, o que completa a demonstração. ■

Agora, vamos ver uma propriedade muito interessante com respeito aos pontos periódicos de σ em Σ_2 .

Proposição 2.3.13. *O conjunto dos pontos periódicos da aplicação shift é denso em Σ_2 .*

Demonstração:

Dado $s = (\dots, s_{-1}, s_0, s_1, \dots) \in \Sigma_2$ e $\varepsilon > 0$. Queremos encontrar um ponto periódico $t \in \Sigma_2$ tal que $d(s, t) < \varepsilon$. Para isso, escolha $m \in \mathbb{N}$ tal que $\frac{1}{2^{m-1}} < \varepsilon$. Tome o ponto $t = (\dots, s_{-m}, s_{-m+1}, \dots, s_{-1}, s_0, s_1, \dots, s_{m-1}, s_m, \dots)$ com $s_i = s_{i+2m+1}$ para todo $i \in \mathbb{Z}$, isto é, t é constituído de infinitos “blocos” $s_{-m}, \dots, s_m$ colocados um na frente do outro. Logo, t é periódico de período $2m+1$ e satisfaz $d(s, t) \leq \frac{1}{2^{m-1}} < \varepsilon$, como queríamos. ■

Como já sabemos que $f : \Lambda \rightarrow \Lambda$ é topologicamente conjugada ao shift sobre Σ_2 , a Proposição 2.3.13 é também válida para f sobre Λ , isto é, se denotarmos o conjunto dos pontos periódicos de f restrita a Λ por $Per(f)$, então

$$\overline{Per(f)} = \Lambda.$$

Observação 3. Como visto anteriormente, o segmento horizontal do conjunto de Cantor $[0, 1] \times C_2$ que contém p é a variedade estável local $W_{loc}^s(p)$. Basta ver que

para todo $q \in [0, 1] \times C_2$ que pertence a esse segmento, como $p \in C_1 \times C_2$ e p tem a mesma coordenada y que q , então

$$|p - f^j(q)| = |f^j(p) - f^j(q)| \leq \mu^n \xrightarrow{n \rightarrow \infty} 0.$$

De modo análogo, obtemos que $W_{loc}^u(p)$ é, localmente, o segmento de reta vertical de $C_1 \times [0, 1]$ que contém p . O mesmo vale para o outro ponto fixo $\tilde{p}$, porém vamos nos restringir apenas ao ponto p . Obtemos a Figura 2.11.

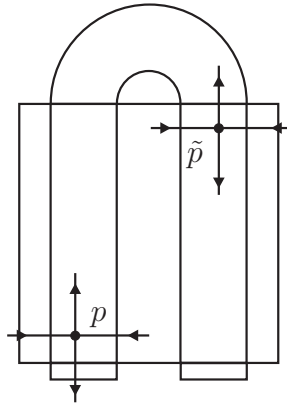


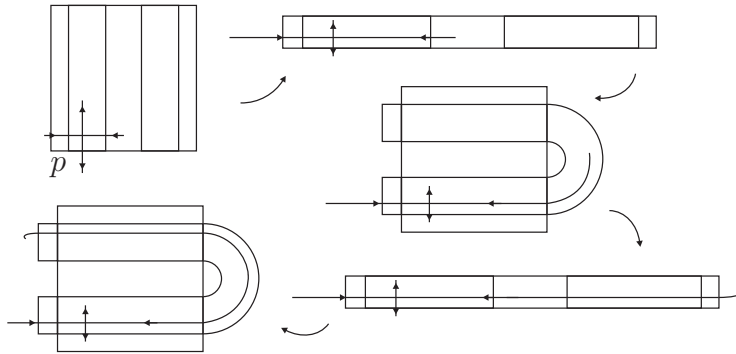
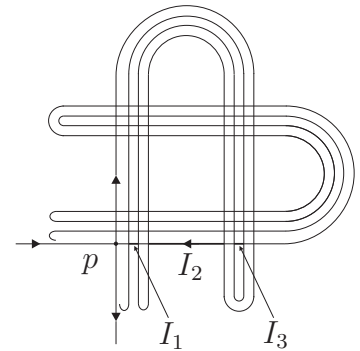
Figura 2.11: Posições de p e $\tilde{p}$.

Essa é uma representação local das variedades estável e instável. Para obter a continuação da variedade estável $W_{loc}^s(p)$, devemos olhar para os iterados por f^{-1} do desenho local feito acima e, para ver a continuação da variedade instável $W_{loc}^u(p)$, devemos olhar para os iterados por f do desenho local. Lembre-se que p é fixo, então, aplicando f^{-1} duas vezes em $W^s(p)$ obtemos a Figura 2.12.

Obtemos uma figura semelhante para $W^u(p)$. Quando olhamos para as duas variedades estável e instável em p , vemos a Figura 2.13.

Note que dentro do retângulo Q , ficam apenas os segmentos de reta verticais, para $W^u(p)$, e horizontais, para $W^s(p)$.

Uma observação final é que, depois de alguns iterados negativos, os intervalos I_1, I_2 e I_3 contidos em $W^s(p)$ representados na Figura 2.13 sempre caem nos arcos que ficam fora de Q e, portanto, nunca interceptam $W^u(p)$. Na verdade, existe uma

Figura 2.12: Como obter $W^s(p)$.Figura 2.13: Representação de $W^s(p)$ e $W^u(p)$.

quantidade infinita enumerável desses intervalos e o complementar desses intervalos é um conjunto de Cantor em $W^s(p)$. A intersecção $W^s(p) \cap W^u(p)$ consiste dos pontos limite desse conjunto de Cantor.

2.4 Bifurcações Homoclínicas

Vamos agora dar a definição e um breve exemplo onde ocorre uma bifurcação homoclínica.

Definição 2.4.1. Dada uma família a um parâmetro de difeomorfismos, temos uma *bifurcação homoclínica* se um par de intersecções homoclínicas colidem, formam uma tangência e então desaparecem, quando variamos o parâmetro, ou, no sentido contrário, quando um par de pontos homoclínicos é gerado depois de uma tangência.

Exemplo 4. Consideremos o Exemplo 3, a ferradura geométrica que é uma aplicação $f : \mathbb{R}^2 \longrightarrow \mathbb{R}^2$. Agora basta compô-la com uma aplicação

$$(x, y) \longmapsto (x, y - \mu)$$

que “translada” a imagem, em particular a do quadrado Q , para baixo. Na Figura 2.14 vemos o efeito dessa “translação” sobre a geometria das variedades estável e instável, $W^s(p_\mu)$ e $W^u(p_\mu)$, para valores decrescentes do parâmetro μ .

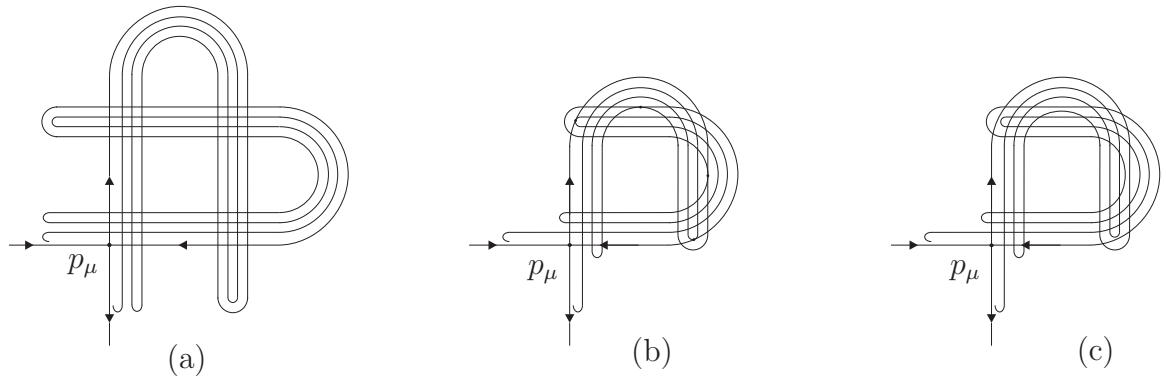


Figura 2.14: Variedades estável e instável da ferradura quando variamos o parâmetro μ .

Na figura (a) temos o Exemplo 3. Na Figura (b) vemos a primeira órbita homoclínica não transversal e temos 4 iterados indicados. Já na Figura (c) vemos que próximo de uma bifurcação existem muitas outras, mas isso estudaremos melhor nas próximas seções.

Capítulo 3

Conseqüências Dinâmicas de uma Intersecção Homoclínica Transversal

Nesta seção vamos estudar como uma intersecção homoclínica transversal pode gerar uma grande complexidade na dinâmica de um dado sistema. Para facilitar a escrita e, principalmente, o entendimento, vamos nos restringir ao caso bidimensional. Os resultados, feitas as devidas adaptações, estendem-se à dimensões arbitrárias e, em alguns casos, para espaços de Banach.

A principal característica de uma intersecção homoclínica transversal é a ocorrência de conjuntos invariantes hiperbólicos. Queremos estudar as propriedades geométricas que ocorrem próximo de uma órbita homoclínica transversal. Assumiremos de agora em diante que os difeomorfismos são de classe C^2 .

Veremos neste capítulo vários conceitos, resultados e propriedades que demonstram o seguinte teorema:

Teorema 2. *Sejam $\varphi : M \longrightarrow M$ um difeomorfismo de classe C^2 onde M é uma superfície bidimensional e $p \in M$ um ponto de sela fixo para φ . Se q é um ponto homoclínico primário para p , então existem coordenadas linearizantes em uma vizinhança que contém a região das variedades estável e instável de p compreendidas*

entre p e q , existe um retângulo R contido nesta vizinhança, um inteiro $N > 0$ e um subconjunto $\Lambda \subset R$ que é maximal, invariante pela aplicação $\Psi := \varphi^N|_R : R \rightarrow R$ e hiperbólico. Além disso, a aplicação $\Psi|_\Lambda : \Lambda \rightarrow \Lambda$ é topologicamente conjugada à aplicação shift $\sigma : \Sigma_2 \rightarrow \Sigma_2$.

Além disso, veremos que a partir deste conjunto Λ podemos obter um outro conjunto $\hat{\Lambda}$ que é φ -invariante, hiperbólico e mantém algumas propriedades do conjunto Λ .

3.1 Coordenadas Linearizantes e o Domínio R

Considere $\varphi : M \rightarrow M$ um difeomorfismo de classe C^2 onde M é uma superfície bidimensional e seja $p \in M$ um ponto de sela fixo, isto é $D\varphi_p$ tem dois autovalores λ e σ com $0 < |\lambda| < 1 < |\sigma|$ e $\varphi(p) = p$. Por simplicidade, assumimos que $0 < \lambda < 1 < \sigma$. Vamos agora enunciar um teorema importante, cuja demonstração pode ser encontrada em [2].

Teorema 3 (Hartman, 1964). *Seja φ um difeomorfismo de classe C^2 em uma superfície M e $p \in M$ um ponto de sela para φ . Então φ admite coordenadas linearizantes de classe C^1 numa vizinhança de p .*

Aqui, entendemos por *coordenadas linearizantes de classe C^1* uma mudança de coordenadas que é de classe C^1 e que torne a função linear, nas novas coordenadas. Assumiremos aqui que numa vizinhança U de p , x_1 e x_2 são coordenadas de classe C^1 tais que $p = (0, 0)$ e tal que

$$\tilde{\varphi}(x_1, x_2) = (\lambda \cdot x_1, \sigma \cdot x_2), \quad (3.1)$$

onde $\tilde{\varphi}$ é a função nas novas coordenadas, isto é, se h é a mudança de coordenadas, então $\tilde{\varphi} = h \circ \varphi \circ h^{-1}$. Porém, nas próximas seções omitiremos o sinal $\sim$ de φ por simplicidade de notação. Além disso, pelo Teorema da Variedade Estável, sabemos que as variedades estável e instável de p , $W^s(p)$ e $W^u(p)$ são de classe C^2 . Vamos

assumir que $W^s(p)$ e $W^u(p)$ possuem pontos de intersecção homoclínica transversal que são diferentes de p .

Definição 3.1.1. Um ponto homoclínico q é *primário* se os caminhos ℓ^u , que liga os pontos p e q em $W^u(p)$, e ℓ^s , que liga os pontos p e q em $W^s(p)$ formam um ponto duplo livre de curvas fechadas. Veja as Figuras 3.1 e 3.2.

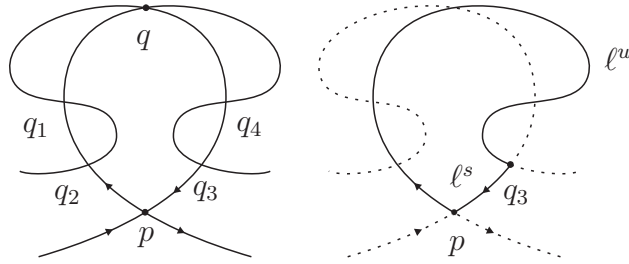


Figura 3.1: Os pontos q , q_1 , q_2 , q_3 e q_4 são primários. O ponto q_3 , por exemplo, é livre de curvas fechadas.

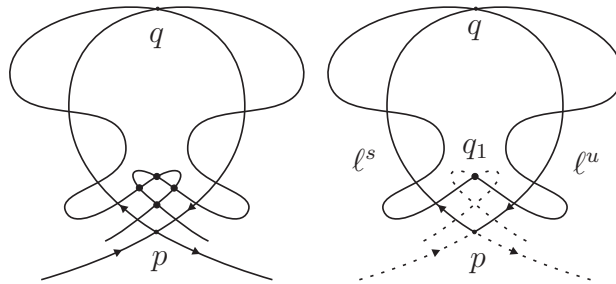


Figura 3.2: Os quatro pontos internos destacados não são primários. O ponto q_1 , por exemplo, está contido em uma curva fechada.

Observe que toda vez que o ponto p possui algum ponto homoclínico, ele também possuirá pontos homoclínicos primários. Além disso, se todas as intersecções entre $W^s(p)$ e $W^u(p)$ são transversais, então o número de órbitas homoclínicas primárias é finito. No exemplo ilustrado nas Figuras 3.1 e 3.2 temos apenas uma órbita homoclínica representada, mas poderíamos ter duas: a que já está desenhada e outra, que estaria abaixo do ponto p , seria a outra intersecção possível entre $W^s(p)$ e $W^u(p)$, no caso de estar contida no plano.

Continuando, vamos assumir que a imagem das coordenadas linearizantes x_1 e x_2 seja o quadrado $(-1, 1) \times (-1, 1) \subset \mathbb{R}^2$. Assim, a imagem desse quadrado por $\tilde{\varphi}$

é $\tilde{\varphi}((-1, 1) \times (-1, 1)) = (-\lambda, \lambda) \times (-\sigma, \sigma)$. Note que $h^{-1}([\lambda, 1) \times (-1, 1)) \subset U$ e

$$\tilde{\varphi}^{-1}([\lambda, 1) \times (-1, 1)) = [1, \lambda^{-1}) \times (-\sigma^{-1}, \sigma^{-1}).$$

Assim, para estender o domínio de definição das coordenadas linearizantes basta defini-las em $[1, \lambda^{-1}) \times (-\sigma^{-1}, \sigma^{-1})$. Uma maneira natural é definir as extensões $\bar{x}_1$ e $\bar{x}_2$ no conjunto $\varphi^{-1}(h^{-1}([\lambda, 1) \times (-1, 1)))$ por

$$\bar{x}_1 = \lambda^{-1} \cdot (x_1 \circ \varphi) \quad \text{e} \quad \bar{x}_2 = \sigma^{-1} \cdot (x_2 \circ \varphi),$$

quando $\varphi^{-1}(h^{-1}([\lambda, 1) \times (-1, 1))) \cap U = \emptyset$. Desse modo, obtemos que o ponto $\tilde{\varphi}(\bar{x}_1, \bar{x}_2) = (\lambda \cdot \bar{x}_1, \sigma \cdot \bar{x}_2) \in [\lambda, 1) \times (-1, 1)$ uma vez que $1 \leq \bar{x}_1 < \lambda^{-1}$ e $-\sigma^{-1} < \bar{x}_2 < \sigma^{-1}$. Veja a Figura 3.3.

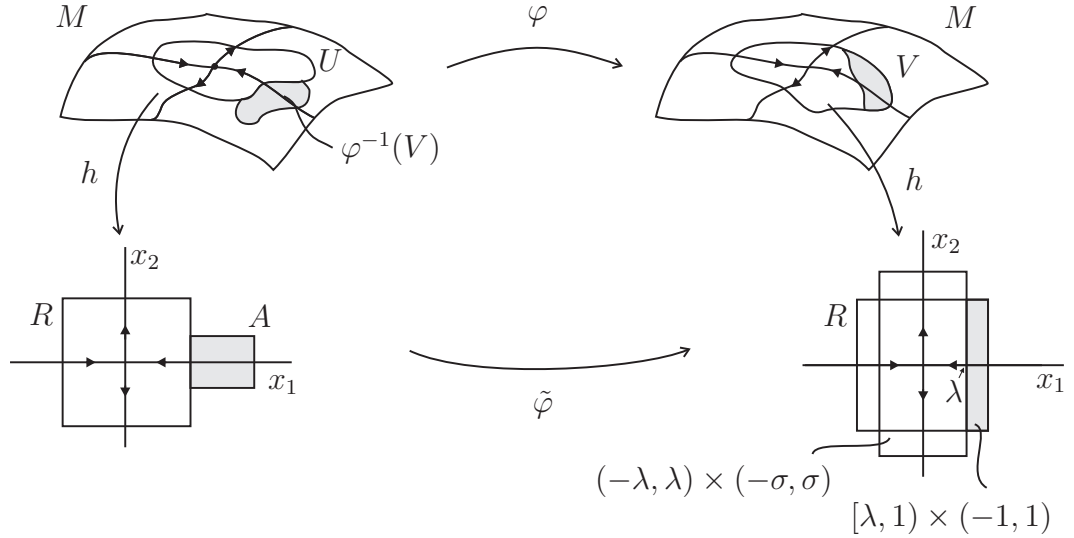


Figura 3.3: O conjunto V é a pré-imagem $h^{-1}([\lambda, 1) \times (-1, 1))$, o conjunto A é $h(\varphi^{-1}(V))$ e $R = (-1, 1) \times (-1, 1)$.

Repetindo esta construção, podemos estender as coordenadas linearizantes ao longo de qualquer segmento contido em $W^s(q)$ que comece no ponto p , pois $W^s(p)$ não possui auto-intersecções. Para isso, basta diminuirmos o domínio U suficientemente. De forma semelhante, podemos estender o domínio de nossas coordenadas linearizantes ao longo da variedade instável $W^u(p)$. Porém, as intersecções ho-

moclínicas formam uma obstrução para a extensão simultânea de tais coordenadas linearizantes ao longo das duas variedades estável e instável simultaneamente. No nosso caso, onde q é um ponto homoclínico primário e ℓ^s, ℓ^u são os caminhos que ligam os pontos p e q por $W^s(p)$ e $W^u(p)$, respectivamente, podemos estender as coordenadas linearizantes ao longo de ℓ^s e ℓ^u , mas em uma vizinhança de q esses sistemas de coordenadas podem não coincidir. Veja a Figura 3.4.

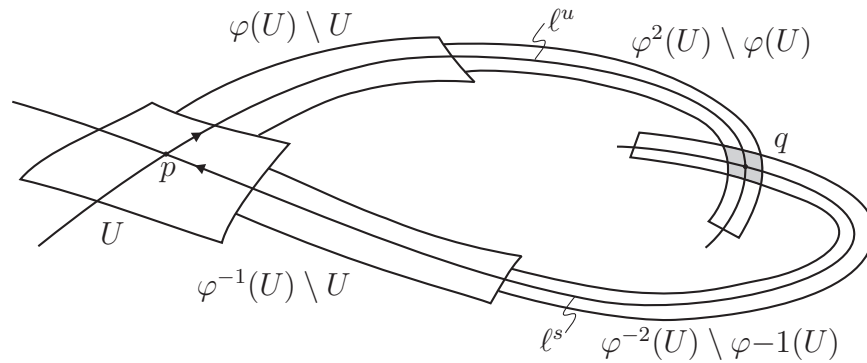


Figura 3.4: Extensão das coordenadas linearizantes. A área cinza é a vizinhança de q com dois sistemas de coordenadas.

Para ter uma visualização no $\mathbb{R}^2$, veja a Figura 3.5. A vizinhança de q na Figura 3.4 acima que apresenta problemas para a extensão está representada em cinza.

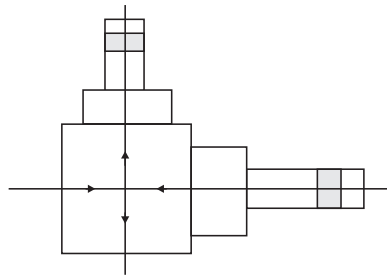


Figura 3.5: As áreas em cinza referem-se à vizinhança de q com com dois sistemas de coordenadas.

Agora estamos prontos para definir o domínio conveniente R . Consideremos o seguinte retângulo R contido no domínio das coordenadas linearizantes (já estendidas)

$$R = \{-a \leq x_1 \leq b, -\alpha \leq x_2 \leq \beta\},$$

onde as constantes a, b, α e β são positivas. Esse retângulo deve conter o caminho ℓ^s , que liga os pontos p e q em $W^s(p)$, e deve satisfazer as seguintes condições para algum $N \in \mathbb{N}$

- (i) $R \cap \varphi^n(R)$ é um retângulo que contém o ponto p para $0 \leq n < N$ e
- (ii) $R \cap \varphi^N(R)$ consiste de duas componentes conexas, sendo que uma delas contém o ponto p e a outra contém o ponto q , como mostra a Figura 3.6.

A condição (ii) acima, na verdade, exige que

$$\{x_1 = b, -\alpha \leq x_2 \leq \beta\} \cap \varphi^n(R) = \emptyset$$

$$\varphi^N(\{-a \leq x_1 \leq b, x_2 = \beta\}) \cap R = \emptyset.$$

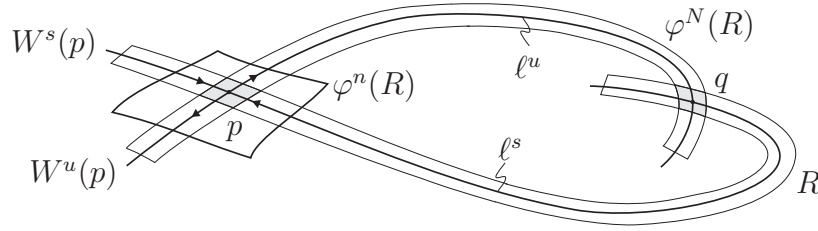


Figura 3.6: As áreas em cinza são as duas componentes conexas de $\varphi^N(R) \cap R$.

Podemos escolher R de tal modo que N fique arbitrariamente grande, basta tomar β arbitrariamente pequeno. Quanto menor for β , maior será N e, portanto, a região $\varphi^n(R)$ será uma faixa estreita ao longo de ℓ^u . Assim, pela transversalidade de $W^s(p)$ e $W^u(p)$ em q , temos que os lados de R e $\varphi^N(R)$ também são transversais.

3.2 Análise Topológica do Subconjunto Invariante Maximal de R

Nesta seção vamos estudar o *conjunto invariante maximal* do conjunto R obtido na seção anterior. Esse conjunto consiste dos pontos em R tais que qualquer múltiplo de sua N -ésima imagem por φ permanece em R . A partir de agora, vamos denotar a aplicação φ^N por Ψ e denotaremos o subconjunto Ψ - invariante maximal de R por

$$\Lambda = \{r \in R \mid \Psi^k(r) \in R \text{ para todo } k \in \mathbb{Z}\}.$$

Vamos “deformar” a Figura 3.6 para entendê-la melhor. Denotemos os vértices de R por $1, 2, 3$ e 4 e suas imagens em $\Psi(R)$ por $1', 2', 3'$ e $4'$ como na Figura 3.7 abaixo. Note que nesta figura, os lados de R e $\Psi(R)$ interceptam-se transversalmente e, além disso, estamos considerando a topologia e as posições de R , seus lados e seus vértices em relação às suas imagens por Ψ conforme mostra a Figura 3.7.

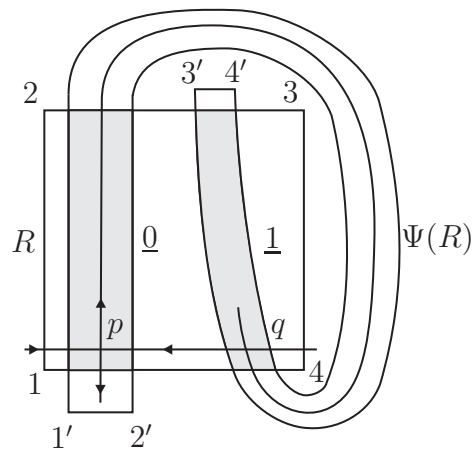


Figura 3.7: As regiões em cinza representam o conjunto $R \cap \Psi(R)$.

Definição 3.2.1. Dizemos que um subconjunto fechado $S \subset R$ é uma *faixa vertical* se S é limitado por duas curvas contínuas disjuntas ℓ_1 e ℓ_2 ligando o lado $(1, 2)$ com o lado $(3, 4)$, como mostra a Figura 3.8. Definimos uma *faixa horizontal* como o

subconjunto fechado $S \subset R$ que é limitado pelas curvas contínuas disjuntas ℓ_1 e ℓ_2 que ligam o lado $(1, 4)$ com o lado $(2, 3)$.

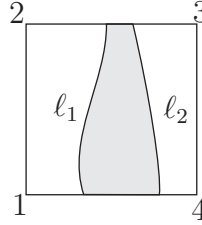


Figura 3.8: Faixa vertical.

Vamos denotar as componentes de $R \cap \Psi(R)$ por $\underline{0}$ a que contém o ponto p e por $\underline{1}$ a que contém o ponto q . Com essas considerações, podemos enunciar e demonstrar o seguinte resultado.

Proposição 3.2.2. *Para toda seqüência $\{a_i\}_{i \in \mathbb{Z}}$, onde $a_i = 0$ ou 1 , existe pelo menos um ponto $r \in \Lambda$ tal que $\Psi^i(r) \in \underline{a_i}$ para todo $i \in \mathbb{Z}$.*

Demonstração:

Seja S uma faixa vertical contida em R . Então, $\Psi(S) \cap R$ consiste de duas faixas verticais, uma contida em $\underline{0}$ e outra contida em $\underline{1}$. Consideremos a seqüência $\{a_i\}_{i \in \mathbb{Z}}$ como no enunciado. Construímos uma seqüência de faixas encaixadas definindo $S_0 = \underline{a_0}$, S_1 a faixa vertical $\Psi(\underline{a_{-1}}) \cap S_0$, S_2 a faixa vertical $\Psi^2(\underline{a_{-2}})$ e assim por diante e, finalmente, $S_\infty = \bigcap_{i=0}^{\infty} S_i$. Logo,

$$S_0 \supset S_1 \supset S_2 \supset \dots$$

Para cada ponto $r \in S_\infty$, temos que $\Psi^{-1}(r) \in \underline{a_{-i}}$ para todo $i \geq 0$. De forma semelhante, obtemos uma seqüência de faixas horizontais

$$T_1 \supset T_2 \supset T_3 \supset \dots$$

e o conjunto $T_\infty = \bigcap_{i=-\infty}^1 T_i$ tal que se $r \in T_\infty$, então $\Psi^i(r) \in a_i$ para todo $i \leq 1$. Afirmamos que $S_\infty \cap T_\infty \neq \emptyset$. De fato, suponhamos o contrário, isto é, que existe

um $i_0 \in \mathbb{Z}$ tal que $S_{i_0} \cap T_{i_0} = \emptyset$. Porém, S_{i_0} contém uma linha que vai do lado (1, 2) ao lado (3, 4) e T_{i_0} contém uma linha que vai do lado (1, 4) ao lado (2, 3), pois são seqüências encaixadas de conjuntos compactos. Logo, estas linhas contêm uma intersecção, o que é um absurdo.

Desse modo, para cada $r \in S_\infty \cap T_\infty$, segue que $\Psi^i(r) \in \underline{a_i}$ para todo $i \in \mathbb{Z}$. Além disso, como $S \subset R$, segue que $r \in \Lambda$, pela própria definição de Λ . ■

3.3 O Conjunto Λ , Hiperbolicidade e Folheações Invariantes

Neta seção, estamos interessados em mostrar que o conjunto invariante maximal Λ é um *conjunto hiperbólico*, como na seguinte definição (que pode ser encontrada em [7]).

Definição 3.3.1. Um conjunto fechado e invariante Λ tem uma *estrutura hiperbólica* para um difeomorfismo f sobre M se

- (i) para cada $p \in M$, o espaço tangente a M no ponto p é decomposto em soma direta $T_p M = \mathbb{E}_p^s \oplus \mathbb{E}_p^u$,
- (ii) para cada $p \in M$, essa soma direta é invariante pela derivada Df de f , isto é,

$$Df_p(\mathbb{E}_p^u) = \mathbb{E}_{f(p)}^u \quad \text{e} \quad Df_p(\mathbb{E}_p^s) = \mathbb{E}_{f(p)}^s,$$

- (iii) $\mathbb{E}_p^u$ e $\mathbb{E}_p^s$ variam continuamente com p , e

- (iv) existe $0 < \lambda < 1$ e $C \geq 1$ tais que para todo $n \geq 0$

$$|Df_p^n.v^s| \leq C\lambda^n|v^s| \quad \text{para } v^s \in \mathbb{E}_p^s \text{ e}$$

$$|Df_p^{-n}.v^u| \leq C\lambda^n|v^u| \quad \text{para } v^u \in \mathbb{E}_p^u.$$

Os espaços $\mathbb{E}_p^s$ e $\mathbb{E}_p^u$ que satisfazem os itens acima são chamados de subespaços *contrativo* e *expansor*, respectivamente. Se um conjunto invariante Λ tem uma estrutura hiperbólica para f , dizemos que Λ é um *conjunto hiperbólico*.

Observação 4. Se $m \in \mathbb{N}$ é tal que $\rho = C\lambda^m < 1$, então

$$\begin{aligned} |Df_p^m \cdot v^s| &\leq \rho |v^s|, v^s \in \mathbb{E}_p^s \\ |Df_p^{-m} \cdot v^u| &\leq \rho |v^u|, v^u \in \mathbb{E}_p^u. \end{aligned}$$

Logo $Df_p^m|_{\mathbb{E}_p^s}$ é uma contração e $Df_p^{-m}|_{\mathbb{E}_p^u}$ é uma expansão. A constante C , portanto, determina o número de iterados de f que são necessários para que os vetores comecem a ser contraídos em $\mathbb{E}_p^s$ e expandidos em $\mathbb{E}_p^u$.

Para mostrar que Λ é hiperbólico é necessário antes exigirmos algumas condições adicionais para a aplicação $\Psi = \varphi^N : \Lambda \rightarrow \Lambda$. Quando olhamos para a região R nas coordenadas linearizantes numa vizinhança pequena do caminho ℓ^s que liga o ponto p ao ponto q através de $W^s(p)$, temos a seguinte configuração

$$R = \{(x_1, x_2) \mid -a \leq x_1 \leq b, -\alpha \leq x_2 \leq \beta\}.$$

Devemos descrever o comportamento de Ψ no conjunto dos pontos em R que não são levados, via Ψ , em R , isto é, nos pontos em $\Psi^{-1}(R) \cap R$. Na componente conexa de $\Psi^{-1}(R) \cap R$ que contém os pontos p e q , Ψ é linear. Em particular, segue de (3.1) que

$$\Psi(x_1, x_2) = (\lambda^N \cdot x_1, \sigma^N \cdot x_2),$$

onde $0 < \lambda < 1 < \sigma$.

Por outro lado, a outra componente conexa de $\Psi^{-1}(R) \cap R$ é levada na componente conexa de $R \cap \Psi(R)$ que contém o ponto q . Veja a Figura 3.9. Esta componente de $R \cap \Psi(R)$ é justamente aquela região com dois sistemas de coordenadas (que não coincidem), construídos na Seção 3.1. Na verdade, nessa região, temos dois sistemas de coordenadas: o que acompanha $W^s(p)$ e que está representado na Figura 3.9

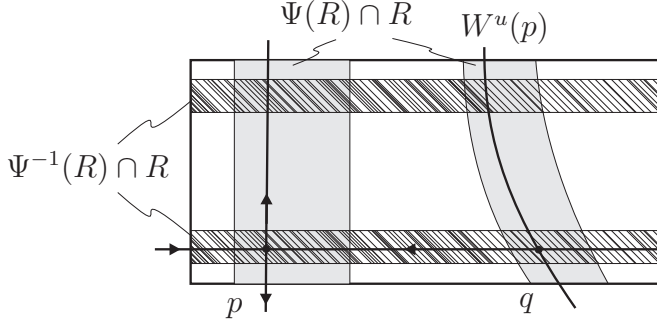


Figura 3.9: Componentes conexas da região $\Psi^{-1}(R) \cap R$.

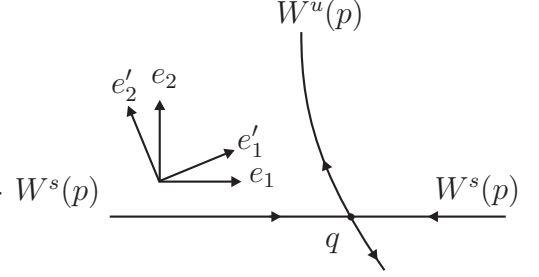


Figura 3.10: Dois sistemas de coordenadas.

pelas coordenadas cartesianas do plano e o que acompanha $W^u(p)$. Vamos, então, denotar por e_1 e e_2 o campo vetorial coordenado das coordenadas linearizantes que acompanham $W^s(p)$ e por e'_1 e e'_2 o campo vetorial coordenado das coordenadas linearizantes que acompanham $W^u(p)$, como na Figura 3.10.

Para r na componente de $R \cap \Psi^{-1}(R)$ que é levada em uma vizinhança de q , temos que

$$\begin{aligned} (D\Psi)(e_1(r)) &= D\Psi_r(e_1(r)) = \lambda^N \cdot e'_1(\Psi(r)) \quad \text{e} \\ (D\Psi)(e_2(r)) &= D\Psi_r(e_2(r)) = \sigma^N \cdot e'_2(\Psi(r)). \end{aligned}$$

Em razão da transversalidade de $W^u(p)$ e $W^s(p)$ e devido à espessura estreita de $\Psi(R)$, quando tomamos N bem grande, segue que e_1 e e'_2 são, inevitavelmente, linearmente independentes. Além disso, como escolhemos as regiões R e $\Psi(R)$ estreitas, podemos assumir que a matriz de transformação de $\{e_1, e_2\}$ para $\{e'_1, e'_2\}$ é praticamente constante. Vamos agora definir a principal ferramenta que será utilizada na prova de que Λ é hiperbólico. Faremos essa próxima definição sobre os conjuntos que estamos interessados, porém elas valem mais geralmente.

Definição 3.3.2. Seja M uma variedade diferenciável (no nosso caso de classe C^2).

- (a) Um *campo de cones contínuo* sobre $R \cap \Psi(R)$ é uma aplicação que associa cada $r \in R \cap \Psi(R)$ a um cone de dois lados $C(r)$ em $T_r M$, o qual é definido por

dois vetores linearmente independentes $w_1(r)$ e $w_2(r)$ do seguinte modo

$$C(r) = \{v \in T_r M \mid v = a_1 \cdot w_1(r) + a_2 \cdot w_2(r), \text{ para } a_1, a_2 > 0\}.$$

A continuidade de C significa que w_1 e w_2 dependem continuamente de r .

(b) Seja $\mathcal{O}_r$ o vetor zero de $T_r M$. Um *campo de cones instável* é um campo de cones contínuo sobre $R \cap \Psi(R)$ tal que

1. para cada $r \in R \cap \Psi(R) \cap \Psi^{-1}(R)$, temos

$$(D\Psi)(C(r)) \subset \text{int}(C(\Psi(r))) \cup \{\mathcal{O}_r\},$$

2. para cada $r \in R \cap \Psi(R) \cap \Psi^{-1}(R)$ e $v \in C(r)$, temos

$$\|D\Psi(v)\| \geq \eta \|v\|,$$

para algum $\eta > 1$, onde ambas as normas são tomadas com respeito à base e_1, e_2 .

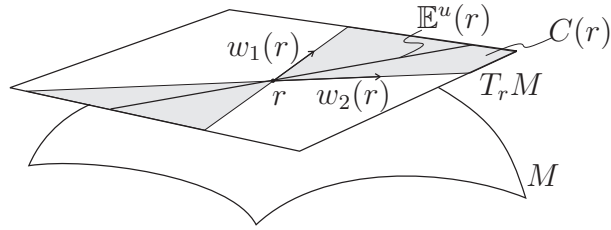


Figura 3.11: Campo de cones e direcional.

(c) Dado $r \in \bigcap_{i=0}^{\infty} \Psi^i(R)$, dizemos que $\mathbb{E}^u(r)$ é um *campo direcional contínuo* se

1. $\mathbb{E}^u(r) \subset C(r)$,
2. $D\Psi(\mathbb{E}^u(r)) \subset \mathbb{E}^u(\Psi(r))$ e

3. existe um $\eta > 1$ tal que todo $v \in \mathbb{E}^u(r)$, com $r \in \bigcap_{i=-1}^{\infty} \Psi^i(R)$, satisfaz

$$\|D\Psi(v)\| \geq \eta\|v\|.$$

Com essas informações estamos aptos a mostrar que Λ é hiperbólico.

Teorema 4. *Assuma que estamos na situação do parágrafo anterior. Então, para R suficientemente estreito e, portanto, N grande, o subconjunto invariante maximal*

$$\Lambda = \bigcap_{n \in \mathbb{Z}} \Psi^n(R)$$

de R é hiperbólico.

Demonstração:

Começaremos construindo um campo de cones sobre $R \cap \Psi(R)$. Na componente conexa de $R \cap \Psi(R)$ que contém p , vamos apenas tomar cones em torno de e_2 estendendo 45° para ambos os lados. Para a outra componente conexa de $R \cap \Psi(R)$, assumindo que R é suficientemente estreito, segue que existe um ângulo $\alpha < 90^\circ$ tal que todo r nesta componente conexa de $R \cap \Psi(R)$ satisfaz a seguinte condição: o cone que é construído em torno de e_2 estendendo-se um ângulo α de ambos os lados, contém e'_2 em seu interior. Essa condição é satisfeita devido ao fato de que $e_1(r)$ e $e'_2(r)$ são linearmente independentes. Veja a Figura 3.12.

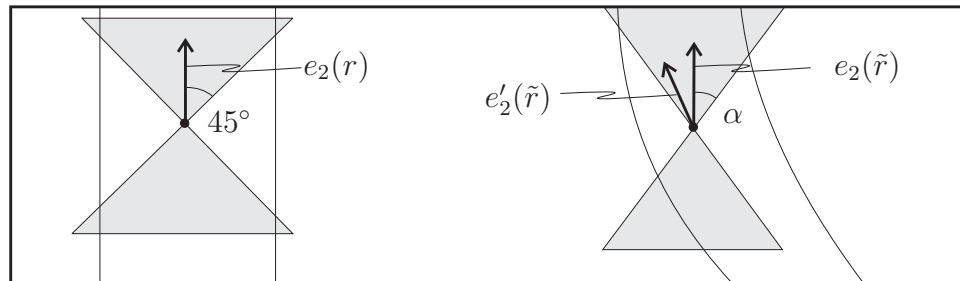


Figura 3.12: Cones em torno de $e_2(r)$ e $e_2(\tilde{r})$ com ângulos 45° e α , respectivamente.

Afirmção: O campo de cones instável C definido acima, isto é, centrado em e_2 e estendendo-se 45° e α para ambos os lados de e_2 dependendo da componente conexa de $R \cap \Psi(R)$, é um campo de cones instável.

De fato: considere as constantes A , B e B' obtidas como segue. Seja $v \in T_r M$ da forma $v = v_1 \cdot e_1(r) + v_2 \cdot e_2(r)$, com $r \in R \cap \Psi(R)$. Logo

(i) se $v \in C(r)$, então

- caso r esteja na componente conexa que contém p , temos

$$\frac{|v_1|}{|v_2|} \leq \tan 45^\circ = 1$$

- caso r esteja na outra componente conexa, temos

$$\frac{|v_1|}{|v_2|} \leq \tan \alpha.$$

Logo, tomando-se $A \geq \max\{\tan \alpha, 1\}$, obtemos que $|v_1| \leq A \cdot |v_2|$, para todo $v \in C(r)$, com $r \in R \cap \Psi(R)$.

(ii) Estamos interessados em exibir $B > 0$ tal que

$$|v_1| \leq B \cdot |v_2| \Rightarrow v \in C(r).$$

Tomemos $B = \min\{\tan \alpha, 1\}$ e temos que

- se r está na componente conexa que contém p , como $|v_1| \leq B \cdot |v_2| \leq |v_2|$, então $v \in C(r)$.
- se r está na outra componente conexa, como $|v_1| \leq B \cdot |v_2| \leq \tan \alpha \cdot |v_2|$, então $v \in C(r)$.

(iii) Seja $v = v'_1 \cdot e'_1(r) + v'_2 \cdot e'_2(r)$, queremos encontrar $B' > 0$ tal que se $|v'_1| \leq B' \cdot |v'_2|$, então $v \in C(r)$. Como as coordenadas linearizantes e'_1 e e'_2 estão definidas

somente na componente conexa de $R \cap \Psi(R)$ que contém q , então o ponto v pertence a esta componente conexa. Considere $\gamma = \angle(e'_2(r), e_2(r)) < \alpha$. Seja $\beta < \frac{1}{2} \min\{|\gamma|, |\alpha - \gamma|\}$, então o cone centrado em $e'_2(r)$ estendido por um ângulo β de ambos os lados, o qual denotaremos por $\tilde{C}(r)$, está contido em $C(r)$. Logo, se tomarmos $B' \leq \tan \beta$, temos que

$$\frac{|v'_1|}{|v'_2|} \leq \tan \beta \Rightarrow v \in \tilde{C}(r) \subset C(r),$$

como queríamos.

Agora, consideremos $N \in \mathbb{N}$ grande o suficiente para que $\left|\frac{\lambda}{\sigma}\right|^N \cdot A < \min\{B, B'\}$ (lembramos que estamos considerando $0 < \lambda < 1 < \sigma$). Então, para $v \in C(r)$, onde $r \in R \cap \Psi(R) \cap \Psi^{-1}(R)$, temos que

$$v = v_1 \cdot e_1(r) + v_2 \cdot e_2(r) \quad \text{e} \quad |v_1| \leq A \cdot |v_2|.$$

Logo, usando a linearidade da derivada, vem que

$$\begin{aligned} \frac{|\pi_1(D\Psi)(v)|}{|\pi_2(D\Psi)(v)|} &= \frac{|\pi_1(v_1 \cdot D\Psi(e_1(r)) + v_2 \cdot D\Psi(e_2(r)))|}{|\pi_2(v_1 \cdot D\Psi(e_1(r)) + v_2 \cdot D\Psi(e_2(r)))|} \\ &= \frac{|v_1| \cdot |\lambda^N| \cdot |e'_1(\Psi(r))|}{|v_2| \cdot |\sigma^N| \cdot |e'_2(\Psi(r))|} \\ &< A \cdot \left|\frac{\lambda}{\sigma}\right|^N < \min\{B, B'\}. \end{aligned} \tag{3.2}$$

Segue de (i) e (ii) que $(D\Psi)(v) \in C(\Psi(r))$ e, portanto, obtemos que $(D\Psi)(C(r)) \subset \text{int}(C(\Psi(r))) \cup \{\mathcal{O}\}_r$. (Observe que a hipótese de que r pertence a $R \cap \Psi(r) \cap \Psi^{-1}(R)$ garante que $\Psi(r)$ pertence a $R \cap \Psi(R) \cap \Psi^2(R) \subset R \cap \Psi(R)$.)

Agora, para cada $r \in R \cap \Psi(R) \cap \Psi^{-1}(R)$ e $v \in C(r)$, temos que $\frac{|v_1|}{|v_2|} \leq A$ e

$$|v| \leq |v_1| + |v_2| = |v_2| \left(\frac{|v_1|}{|v_2|} + 1 \right) \leq |v_2|(A + 1), \tag{3.3}$$

Logo, denotando $\zeta = \min\{B, B'\}$, de (3.2) e (3.3) temos

$$\begin{aligned}
|D\Psi(v)| &\geq |\pi_2 D\Psi(v)| - |\pi_1 D\Psi(v)| \\
&\geq |\pi_2 D\Psi(v)| - \zeta |\pi_2 D\Psi(v)| \\
&= |v_2| \cdot |\sigma|^N \cdot |e_2(\Psi(v))|(1 - \zeta) \\
&= |v_2|(1 + A) \cdot \left(\frac{1 - \zeta}{1 + A}\right) |\sigma|^N \\
&\geq \left(\frac{1 - \zeta}{1 + A}\right) |\sigma|^N \cdot |v|.
\end{aligned}$$

Denotemos $\eta = \frac{1 - \zeta}{1 + A} \cdot \sigma^N$ e obtemos que

$$|D\Psi(v)| \geq \eta \cdot |v|,$$

para cada $r \in R \cap \Psi(R) \cap \Psi^{-1}(R)$ e $v \in C(r)$. Note que para N suficientemente grande temos que $\eta > 1$ pois as constantes ζ e A não dependem de N . Portanto, o campo de cones é um campo de cones instável. Isto completa a prova da afirmação.

Continuando a demonstração, da existência desse campo de cones instável C , segue que existe um campo direcional contínuo $\mathbb{E}^u$. O campo $\mathbb{E}^u$, restrito à Λ , é obtido tomando-se a intersecção das imagens por $D\Psi$ do campo de cones instável C .

Substituindo Ψ por Ψ^{-1} , podemos construir, de modo análogo, um campo de cones estável e um campo direcional $\mathbb{E}^s$, o qual é $D\Psi$ -invariante e contrai distâncias por $D\Psi$. Portanto,

$$T_\Lambda M = \mathbb{E}_\Lambda^u \oplus \mathbb{E}_\Lambda^s$$

é a decomposição hiperbólica para o conjunto Λ desejada. ■

Observação 5. Note que o campo de cones C foi construído após a escolha de N , Ψ foi definida em termos de N ($\Psi = \varphi^N$) e o domínio do campo de cones, $R \cap \Psi(R)$, foi definido em termos de Ψ . No entanto, a maneira de aumentar N é fazer com que

R fique estreito. Com isso, diminuimos o domínio no qual o campo de cones está definido. O fato de Ψ mudar de φ^N para $\varphi^{N'}$, $N' > N$, não exerce influência sobre os argumentos, pois os campos vetoriais e_1, e_2, e'_1 e e'_2 não variam (sua definição é local). É por essa razão que podemos ajustar R e N depois.

Observação 6. Observe que provamos um pouco mais: existem campos vetoriais $e^u \in e_\Lambda^u$ e $e^s \in e_\Lambda^s$ e uma constante $\eta > 1$, tal que para $r \in \Lambda$ vale

$$\|D\Psi(e^u(r))\| \geq \eta \cdot \|e^u(r)\| \text{ e } \|D\Psi(e^s(r))\| \geq \eta^{-1} \cdot \|e^s(r)\|;$$

onde

$$e_\Lambda^\iota = \{e^\iota \mid e^\iota \text{ é campo direcional}\} \text{ e}$$

a aplicação e^ι é dada por $\Lambda \ni r \mapsto e^\iota(r) \in \mathbb{E}^\iota(r) \subset T_r M$, para $\iota = s, u$.

Passemos agora para o segundo objetivo desta seção, construir as folheações estável e instável. Descreveremos apenas a construção da folheação instável deixando alguns detalhes para a referência, uma vez que a construção da folheação estável é análoga. Inicialmente vejamos a seguinte definição.

Definição 3.3.3. Uma *folheação instável* para $\Lambda = \bigcap_{i \in \mathbb{Z}} \Psi^i(R)$ é uma *folheação* (ver a definição no Apêndice B) $\mathcal{F}^u$ de uma vizinhança de Λ (aqui consideraremos $R \cap \Psi(R)$) tal que

- (a) para cada $r \in \Lambda$, $\mathcal{F}^u(r)$, a folha de $\mathcal{F}^u$ que contém r , é tangente a $\mathbb{E}^u(r)$,
- (b) para cada r suficientemente próximo de Λ ,

$$\Psi(\mathcal{F}^u(r)) \supset \mathcal{F}^u(\Psi(r)) \text{ e}$$

- (c) as direções tangentes das folhas de $\mathcal{F}^u(r)$ variam continuamente.

Construção da Folheação Instável

Recordemos quais são as posições relativas de R , $\Psi(R)$ e $\Psi^{-1}(R)$ observando a Figura 3.13.

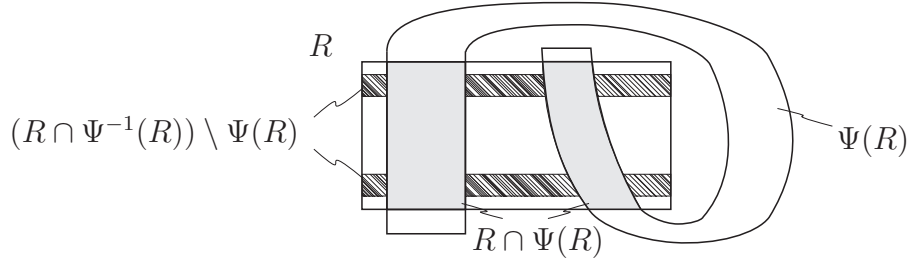


Figura 3.13: Posições relativas de R , $\Psi(R)$ e $\Psi^{-1}(R)$.

Considere uma folheação $\tilde{\mathcal{F}}^u$ de classe C^2 , que ainda não é a folheação instável, sobre o conjunto $(\Psi(R) \cup \Psi^{-1}(R)) \cap R$ tal que

- (a) em $R \cap \Psi(R)$, as direções tangentes às folhas estão contidas nos cones instáveis
- (b) as imagens por Ψ das folhas em $(R \cap \Psi^{-1}(R)) \setminus \Psi(R)$ têm as direções tangentes contidas nos cones instáveis (veja a Figura 3.14),

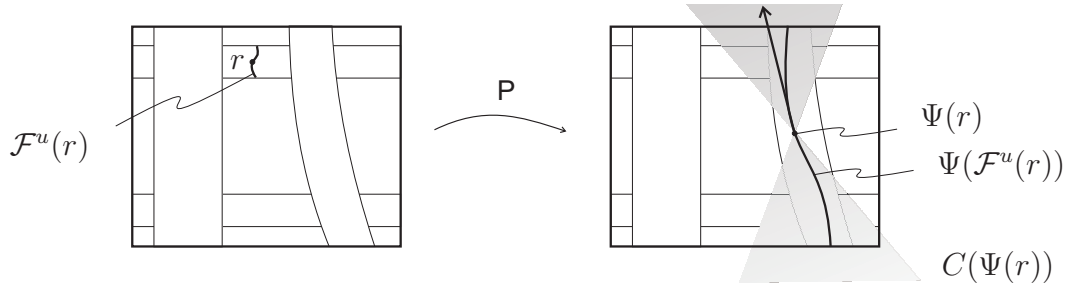


Figura 3.14: Folhas $\mathcal{F}^u(r)$ e $\Psi(\mathcal{F}^u(r))$

- (c) os quatro arcos de $\partial R \cap \Psi^{-1}(R)$ são folhas de $\tilde{\mathcal{F}}^u$, a união desses quatro arcos é denotada por $\mathbf{e}_0$,
- (d) os quatro arcos de $\partial(\Psi(R)) \cap R$ são folhas de $\tilde{\mathcal{F}}^u$, a união desses quatro arcos é denotada por $\mathbf{e}_1$
- (e) Ψ leva folhas de $\tilde{\mathcal{F}}^u$ próximas de $\mathbf{e}_0$ em folhas de $\tilde{\mathcal{F}}^u$ próximas de $\mathbf{e}_1$.

Como todos os cones do campo de cones instável estão centrados no campo vetorial e_2 e, onde estiverem definidos, esses cones contêm e'_2 , segue que existe folheações como a descrita acima. Para essas folheações definimos um operador Ψ_* como segue

- em $(R \cap \Psi^{-1}(R)) \setminus \Psi(R)$, as folhas de $\tilde{\mathcal{F}}^u$ e as de $\Psi_*(\tilde{\mathcal{F}}^u)$ coincidem e
- em $R \cap \Psi(R)$, as folhas de $\Psi_*(\tilde{\mathcal{F}}^u)$ são as componentes conexas das imagens, por Ψ , das folhas de $\tilde{\mathcal{F}}^u$ interceptadas com $(R \cap \Psi(R))$.

Devido às condições (c), (d) e (e) acima, temos que $\Psi_*(\tilde{\mathcal{F}}^u)$ é também de classe C^2 .

Vamos utilizar agora um fato que segue da Teoria das Variedades Invariantes, o qual pode ser encontrado um esboço da demonstração no Apêndice I de [5], o fato de que o seguinte limite existe

$$\lim_{i \rightarrow \infty} \Psi_*^i(\tilde{\mathcal{F}}^u) = \mathcal{F}^u$$

e que $\tilde{\mathcal{F}}^u$ é de classe C^1 . Tal limite depende da “folheação inicial” $\tilde{\mathcal{F}}^u$.

Podemos obter uma folheação estável $\mathcal{F}^s$ através de uma folheação instável para Ψ^{-1} . Observe que podemos estender nosso campo de vetores e^u e e^s em $\mathbb{E}^u$ e $\mathbb{E}^s$, respectivamente, para o campo dos vetores tangentes a $\mathcal{F}^u$ e $\mathcal{F}^s$ de tal modo que para uma constante $\tilde{\eta} > 1$ e para todo $r \in R \cap \Psi(R) \cap \Psi^{-1}(R)$ tenhamos

$$\|D\Psi(E^u(r))\| \geq \tilde{\eta} \cdot \|e^u(r)\|$$

e

$$\|D\Psi(E^s(r))\| \leq \tilde{\eta}^{-1} \cdot \|e^s(r)\|.$$

3.4 A Estrutura do Conjunto Λ

Nesta seção vamos ver que o conjunto Λ comporta-se localmente como um produto cartesiano, que existe uma conjugação topológica entre a aplicação $\Psi|_{\Lambda} : \Lambda \rightarrow \Lambda$ e

o shift $\sigma : \Sigma_2 \rightarrow \Sigma_2$ e, finalmente, veremos que propriedades da aplicação $\Psi = \varphi^N$ conseguimos recuperar para a aplicação φ sobre um conjunto correspondente ao já conhecido conjunto Λ .

Inicialmente, vamos dividir o conjunto invariante maximal $\Lambda \subset R$ em pequenos blocos.

Definição 3.4.1. Seja $k \in \mathbb{N}$ e $A = (a_{-k}, a_{-k+1}, \dots, a_{k-1}, a_k)$ uma $(2k + 1)$ -upla onde $a_i = 0$ ou 1 para $i = -k, \dots, k$. Definimos um A -bloco por

$$\Lambda_A = \{r \in \Lambda \mid \Psi^i(r) \in \underline{a_i}, \text{ para } i = -k, \dots, k\},$$

dizemos que k é o *raio* de A e o *diâmetro* de Λ_A é o supremo das distâncias em Λ_A , medido em $\mathbb{R}^2$.

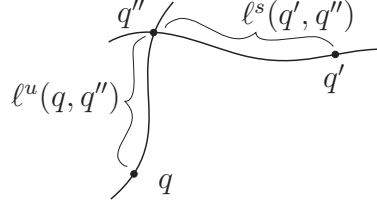
Vimos na seção anterior que os vetores ao longo das folheações estável e instável expandem e contraem por um fator $\tilde{\eta} > 1$ e $\tilde{\eta}^{-1}$, respectivamente. Veremos agora que podemos dar uma estimativa para o diâmetro de um A -bloco utilizando as folheações instável e estável.

Proposição 3.4.2. *Seja c o maior comprimento que uma componente de uma folha estável ou instável pode atingir em $R \cap \Psi(R)$ e tome $\tilde{\eta} > 1$ como definido anteriormente. Considere Λ_A um A -bloco e $k \in \mathbb{N}$ o raio de A . Então, o diâmetro de Λ_A é no máximo $2.c.\tilde{\eta}^{-k}$.*

Demonstração:

Sejam q e q' dois pontos na mesma componente conexa de $R \cap \Psi(R)$. Então existem únicos arcos $\ell^u(q, q'')$ nas folhas de $\mathcal{F}^u$ ligando q à q'' e $\ell^s(q', q'')$ nas folhas de $\mathcal{F}^s$ ligando q' à q'' , onde q'' é a intersecção da folha instável que passa por q com a folha estável que passa por q' .

No entanto, devemos considerar mais: os pontos q e q' devem estar em Λ_A . Em virtude disso, ao aplicarmos Ψ^i , $i = -k, \dots, k$, a configuração da Figura 3.15 permanecerá por completo em uma mesma componente conexa de $R \cap \Psi(R)$, pela própria definição de Λ_A . Como o comprimento maximal de folha instável ou estável

Figura 3.15: Arcos $\ell^u(q, q'')$ e $\ell^s(q', q'')$.

em $R \cap \Psi(R)$ é c e como Ψ^i expande o A -bloco por um fator $\tilde{\eta}^{|i|}$ em uma direção que depende do sinal de i , segue que o comprimento máximo da $\ell^u(p, p'')$ e $\ell^s(p', p'')$ é menor ou igual a $c \cdot \tilde{\eta}^{-k}$. Logo, a distância entre p e p' é menor que $2 \cdot c \cdot \tilde{\eta}^{-k}$. ■

Em posse desse resultado podemos demonstrar o seguinte teorema.

Teorema 5.

(a) *Seja A com raio k , como na Definição 3.4.1. O comprimento de um A -bloco Λ_A vai para zero quando o raio k vai para o infinito.*

(b) *Para cada seqüência*

$$(\dots, a_{-2}, a_{-1}, a_0, a_1, a_2, \dots)$$

com $a_i = 0$ ou 1 , existe exatamente um ponto $r \in \Lambda$ tal que $\Psi^i(r) \in \underline{a_i}$ para todo $i \in \mathbb{Z}$.

(c) *Existe um homeomorfismo $h : \Lambda \rightarrow \Sigma_2$ (com a topologia produto em Σ_2) definido, para $r \in \Lambda$, por*

$$h(r) = (\dots, a_{-2}, a_{-1}, a_0, a_1, a_2, \dots)$$

tal que $\Psi^i(r) \in \underline{a_i}$. Além disso, Ψ e o operador shift, $\sigma : \Sigma_2 \rightarrow \Sigma_2$, são topologicamente conjugados.

Demonstração:

A afirmação (a) segue da Proposição 3.4.2. A (b), decorre da Proposição 3.2.2.

E que h é uma conjugação topológica, (c), prova-se da mesma forma que foi feito para a Ferradura Geométrica de Smale no Teorema 1. ■

Observação 7. O conjunto dos pontos periódicos são densos em Λ . Esse fato segue diretamente da Proposição 2.3.13 com o fato de Ψ e o shift serem topologicamente conjugados. Em particular, todos os pontos periódicos são de sela, pois têm uma direção que expande e outra que contrai e seu comportamento é o de um ponto fixo para Ψ^ℓ , dado que ℓ é seu período.

Observação 8. O conjunto Λ tem uma estrutura local de produto, no seguinte sentido: se $r, r' \in \Lambda$ estão em uma mesma componente conexa de $R \cap \Psi(R)$, então a intersecção da folha estável $\mathcal{F}^s(r)$ que passa por r com a folha instável $\mathcal{F}^u(r')$ que passa por r' também pertence a Λ . De fato, se $h(r) = \{a_i\}_{i \in \mathbb{Z}}$ e $h(r') = \{a'_i\}_{i \in \mathbb{Z}}$, então $\mathcal{F}^s(r) \cap \Lambda$ corresponde à seqüência $\{\{b_i\}_{i \in \mathbb{Z}} \mid b_i = a_i \text{ para } i \geq 0\}$ e $\mathcal{F}^u(r') \cap \Lambda$ corresponde à seqüência

$$\{\{b'_i\}_{i \in \mathbb{Z}} \mid b'_i = a'_i \text{ para } i \leq 0\}, \quad (3.4)$$

uma vez que a construção da folheação $\mathcal{F}^u(r)$ (ou $\mathcal{F}^s(r)$) foi através do limite

$$\lim_{i \rightarrow \infty} \Psi^i(\tilde{\mathcal{F}}^u) = \mathcal{F}^u,$$

onde $\tilde{\mathcal{F}}^u$ é uma folheação de $R \cap \Psi(R)$ que satisfaz as condições de (a) a (e) da construção da folheação instável. Como r e r' estão na mesma componente conexa de $R \cap \Psi(R)$, segue que $a_0 = a'_0$ e, portanto, o ponto em $\mathcal{F}^s(r) \cap \mathcal{F}^u(r')$ corresponde à seqüência

$$(\dots, a'_{-2}, a'_{-1}, a'_0 = a_0, a_1, a_2, \dots),$$

que é um ponto de Λ .

Observação 9. Um ponto $r \in \Lambda$ tem uma *órbita caótica* se sua seqüência simbólica $h(r) = \{a_i\}_{i \in \mathbb{Z}}$ não é eventualmente periódica, isto é, $(a_n, a_{n+1}, \dots)$ não é periódica para nenhum $n \in \mathbb{N}$. Logo, as órbitas em Λ ou são assintoticamente periódicas ou

são caóticas. Em outras palavras, ou se aproxima de uma órbita periódica para n grande ou nunca tem um comportamento “quase”periódico.

Até agora estudamos as propriedades da aplicação $\Psi = \varphi^N$ e do subconjunto invariante maximal Λ . Vamos agora resgatar algumas dessas propriedades para a aplicação φ e estudar um conjunto correspondente ao conjunto Λ .

O conjunto Ψ -invariante maximal Λ está contido em $R \cap \Psi(R)$. Denotamos as componentes conexas de $R \cap \Psi(R)$ por $\underline{0}$ e $\underline{1}$. Definimos

$$\hat{\Lambda} = \bigcup_{i=0}^{N-1} \varphi^i(\Lambda).$$

Esse conjunto é o correspondente ao Λ , porém é φ -invariante, o que se mostra trivialmente. Este conjunto apresenta uma decomposição em subconjuntos disjuntos, é o que mostra a seguinte proposição.

Proposição 3.4.3. *O conjunto $\hat{\Lambda}$ definido acima é a união disjunta de $\{p\}$, $\Lambda \setminus \{p\}$, $\varphi(\Lambda \setminus \{p\})$, $\dots$, $\varphi^{N-1}(\Lambda \setminus \{p\})$.*

Demonstração: Começaremos mostrando que $\Lambda \cap \varphi^{-i}(\Lambda) = \{p\}$, para $0 < i < N$, e disso segue o resultado como veremos adiante. Seja $r \in \Lambda$ com $\varphi^i(r) \in \Lambda$, para algum $0 < i < N$. Então, $r \notin \underline{1}$ pois, na construção de R , tomamos N tal que $R \cap \varphi^n(R)$ consiste de um retângulo contendo p para $0 \leq n < N$ e, como $\Lambda \subset R$, segue que $\Lambda \cap \varphi^i(\Lambda) \subset R \cap \varphi^i(R)$ que está contido na componente conexa que contém p , que é a $\underline{0}$. Veja a Figura 3.16. Logo, para todo $k \in \mathbb{Z}$, $\varphi^{k.N}(r) = \Psi^k(r)$ tem as

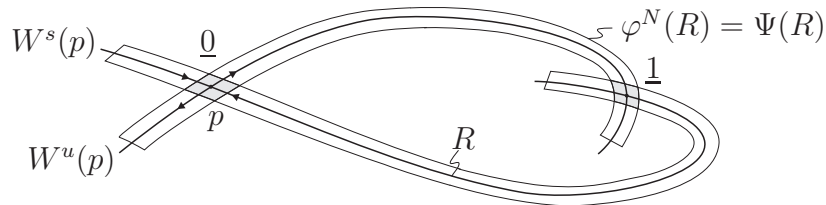


Figura 3.16: Regiões $\underline{0}$ e $\underline{1}$.

mesmas propriedades, isto é, $\Psi^k(r) \in \Lambda$ e $\varphi^i(\Psi^k(r)) \in \Lambda$, pois Λ é Ψ -invariante e

$$\varphi^i(r) \in \Lambda \implies \Psi^k(\varphi^i(r)) = \varphi^i(\Psi^k(r)) \in \Lambda.$$

Isto implica que $\Psi^k(r) \in \underline{0}$ para todo k e, portanto, que $h(r) = (\dots, 0, 0, 0, \dots) = h(p)$, ou seja, $r = p$. Com isso provamos que os conjuntos do enunciado são disjuntos dois a dois pois, para $i > j$ com $i, j \in \{1, \dots, N-1\}$,

$$\begin{aligned} \varphi^i(\Lambda \setminus \{p\}) \cap \varphi^j(\Lambda \setminus \{p\}) &= \emptyset \iff (\Lambda \setminus \{p\}) \cap \varphi^{j-i}(\Lambda \setminus \{p\}) = \emptyset \\ &\iff (\Lambda \cap \varphi^{j-i}(\Lambda)) \setminus \{p\} = \emptyset \\ &\iff \Lambda \cap \varphi^{j-i}(\Lambda) = \{p\}, \end{aligned}$$

ou seja, $\Lambda \cap \varphi^{-\ell}(\Lambda) = \{p\}$ para $k \in \{1, \dots, N-1\}$. O que prova a proposição. ■

Outras propriedades que conseguimos resgatar para $\hat{\Lambda}$ são a hiperbolicidade e densidade dos pontos periódicos que são denotados por $Per(\hat{\Lambda})$.

Proposição 3.4.4. *O conjunto $\hat{\Lambda}$ é hiperbólico e $\overline{Per(\hat{\Lambda})} = \hat{\Lambda}$.*

Demonstração:

Como Λ é hiperbólico segue que, para cada $x \in \Lambda$,

$$T_x M = \mathbb{E}_x^s \oplus \mathbb{E}_x^u,$$

e esta decomposição varia continuamente com x e

$$D\Psi_x(\mathbb{E}_x^s) = \mathbb{E}_{\Psi(x)}^s \quad \text{e} \quad D\Psi_x(\mathbb{E}_x^u) = \mathbb{E}_{\Psi(x)}^u.$$

Além disso, existem constantes $C > 0$ e $0 < \nu < 1$ tais que

$$\begin{aligned} \|D\Psi^n v^s\| &\leq C \cdot \nu^n \|v^s\| \quad \text{e} \\ \|D\Psi^n v^u\| &\geq C^{-1} \cdot \nu^{-n} \|v^u\|, \end{aligned}$$

para todo (v^s, v^u) pertencentes ao fibrado $\mathbb{E}^s \oplus \mathbb{E}^u$ e para todo $n \geq 1$. As direções naturais para as direções estáveis e instáveis são os iterados de $\mathbb{E}^s$ e $\mathbb{E}^u$ por $D\varphi$, isto é, $\mathbb{E}_{\varphi^i(x)}^\iota = D\varphi_x^i|_{\mathbb{E}_x^\iota}$, para $\iota = s, u$ e $i = 1 \dots, N-1$.

A direção (fibrado) $\mathbb{E}^s$ contrai pois dado $n \geq 1$, sejam $k \geq 0$ e $0 \leq r < N$ tais que $n = kN + r$, logo

$$\|D\varphi^n v\| \leq \|D\varphi^r\| \cdot \|D\Psi^k v\| \leq C_1 \nu_1^n \|v\|,$$

onde $\nu_1 = \nu^{1/N}$ e $C_1 = C \cdot \max\{\|D\varphi^j\|; 0 \leq j < N\}$. Analogamente, mostra-se que o fibrado $\mathbb{E}^u$ contrai no passado.

Mostremos agora que os pontos periódicos de $\hat{\Lambda}$ são densos. Como o conjunto $\hat{\Lambda} = \{p\} \cup (\Lambda \setminus \{p\}) \cup \varphi(\Lambda \setminus \{p\}) \cup \dots \cup \varphi^N(\Lambda \setminus \{p\})$, segue que se $q \in \hat{\Lambda}$ é um ponto periódico, isto é, $q \in \varphi^i(\Lambda) \setminus \{p\}$, então o inteiro k para o qual $\varphi^k(q) = q$ deverá ser um múltiplo de N , pois

$$\varphi^k(q) = \varphi^i(q) \iff k = \ell N \quad \text{para algum } \ell \in \mathbb{N}.$$

Seja então $k = \ell N$ o período de q . Logo, se $r \in \Lambda$ é tal que $q = \varphi^i(r)$, então $q \in \varphi^i(\text{Per}(\Lambda))$, pois

$$\Psi^\ell(r) = \varphi^{\ell N}(r) = \varphi^{\ell N}(\varphi^{-i}(q)) = \varphi^{-i}(\varphi^k(q)) = \varphi^{-i}(q) = r,$$

ou seja, $r \in \text{Per}(\Lambda)$. Portanto, $\text{Per}(\varphi^i(\Lambda)) \subset \varphi^i(\text{Per}(\Lambda))$.

Agora, se $q \in \varphi^i(\text{Per}(\Lambda))$ então existe $\ell \in \mathbb{N}$ tal que $q = \varphi^{i+N\ell}(r)$ para $r \in \Lambda$. Logo, $q \in \text{Per}(\varphi^i(\Lambda))$ porque existe $k = N\ell \in \mathbb{N}$ tal que $\varphi^k(q) = \varphi^{N\ell}(\varphi^{-i}(r)) = \varphi^{-i}(\varphi^{N\ell}(r)) = \varphi^{-i}(\Psi^\ell(r)) = \varphi^{-i}(r) = q$ e porque $q = \varphi^i(r)$ com $r \in \Lambda$. Portanto,

$$\varphi^i(\text{Per}(\Lambda)) = \text{Per}(\varphi^i(\Lambda)). \quad (3.5)$$

Logo,

$$\begin{aligned}
Per(\hat{\Lambda}) &= Per(\{p\} \cup (\Lambda \setminus \{p\}) \cup \dots \cup \varphi^{N-1}(\Lambda \setminus \{p\})) \\
&= \{p\} \cup Per(\Lambda \setminus \{p\}) \cup \dots \cup Per(\varphi^{N-1}(\Lambda \setminus \{p\})) \\
&= \{p\} \cup Per(\Lambda \setminus \{p\}) \cup \dots \cup \varphi^{N-1}(Per(\Lambda \setminus \{p\})).
\end{aligned}$$

Assim

$$\begin{aligned}
\overline{Per(\hat{\Lambda})} &= \{p\} \cup \overline{Per(\Lambda \setminus \{p\})} \cup \dots \cup \varphi^{N-1}(\overline{Per(\Lambda \setminus \{p\})}) \\
&= \{p\} \cup \Lambda \cup \dots \cup \varphi^{N-1}(\Lambda) = \hat{\Lambda},
\end{aligned}$$

como queríamos demonstrar. ■

Como visto anteriormente, a existência de uma órbita homoclínica no sistema gera uma grande complicação no desenho formado pelas variedades estável e instável. Em vista disso temos o seguinte resultado.

Proposição 3.4.5. *Na situação acima, $\hat{\Lambda}$ está contido no fecho de ambas variedades estável e instável do ponto p .*

Demonstração:

Como as órbitas periódicas são densas em $\hat{\Lambda}$ é suficiente provar que cada ponto periódico de $\hat{\Lambda}$ está contido em $\overline{W^\iota(p)}$, para $\iota = s, u$. Como $\overline{W^\iota(p)}$ é φ -invariante é suficiente provar que os pontos periódicos de Λ estão em $\overline{W^\iota(p)}$ pois, uma vez provado que $Per(\Lambda) \subset \overline{W^\iota(p)}$, segue de (3.5) que $Per(\varphi^i(\Lambda)) = \varphi^i(Per(\Lambda)) \subset \overline{W^\iota(p)}$ para $i = 1, \dots, N-1$. E assim obtemos que

$$Per(\hat{\Lambda}) = \bigcup_{i=0}^{N-1} Per(\varphi^i(\Lambda)) \subset \overline{W^\iota(p)}.$$

Provemos então que $Per(\Lambda) \subset \overline{W^\iota(p)}$, para $\iota = s, u$. Para um dado ponto periódico $r \in \Lambda$, a variedade instável $\overline{W^u(p)}$ contém a folha da folheação instável que passa por r e portanto, intercepta $\overline{W^s(p)}$ em um ponto que chamaremos de q . Então, ao iterarmos φ^{-1} , segue que esse ponto periódico está contido no fecho de $\overline{W^s(p)}$, pois a sequência $(\varphi^{-n}(q))_{n \in \mathbb{N}}$ de elementos de $\overline{W^s(p)}$ converge para r

por também estar contida em $\overline{W^u(p)}$. De modo semelhante provamos que o Λ está contido em $\overline{W^u(p)}$. ■

3.5 Pontos Homoclínicos de Órbitas Periódicas

Consideremos novamente o difeomorfismo $\varphi : M \longrightarrow M$ e seja $\{p_0, p_1, \dots, p_{k-1}\}$ uma órbita periódica de período k . Vamos assumir que essa órbita é do tipo sela e denotaremos as variedades instável e estável por $W^u(p_i)$ e $W^s(p_i)$. Temos assim dois tipos de órbitas homoclínicas:

- (i) órbitas homoclínicas de um ponto fixo hiperbólico para φ^k e
- (ii) intersecções de $W^u(p_i)$ e $W^s(p_j)$ para $i \neq j \pmod k$ (se esta condição não é satisfeita recaímos no caso anterior), veja a Figura 3.17.

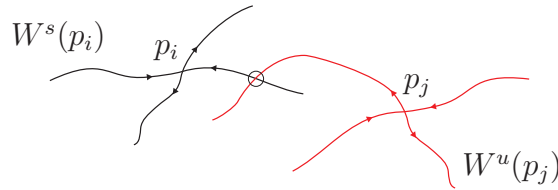


Figura 3.17: O ponto circulado é uma intersecção de $W^s(p_i)$ com $W^u(p_j)$.

No caso (ii), para $t = j - i$ temos também intersecções de $W^s(p_j)$ com $W^u(p_{j+t})$ e assim por diante. Por exemplo, se $k = 3$, $i = 1$ e $j = 2$, temos então que $t = j - i = 1$ e $j + t = 3$ e portanto existe uma intersecção entre $W^s(p_2)$ e $W^u(p_3)$. Veja a Figura 3.18. Assim obtemos um ciclo com período igual ao menor inteiro ℓ tal que k divide ℓt .

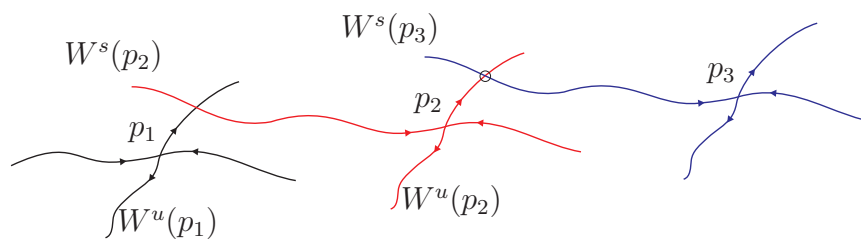


Figura 3.18: O ponto circulado é a nova intersecção de $W^u(p_2)$ com $W^s(p_3)$.

Em nossa ilustração (Figura 3.18), esse inteiro ℓ é 3 uma vez que o inteiro $k = 3$ divide o inteiro $1 \cdot 3$.

Capítulo 4

Tangências Homoclínicas

Nesse capítulo veremos resultados a respeito de desdobramentos de tangências homoclínicas em famílias a um parâmetro de difeomorfismos $\varphi_\mu : M \rightarrow M$, onde M é uma variedade de dimensão 2. Para isso vamos ver a definição de difeomorfismo *localmente dissipativo*.

Definição 4.0.1. Seja φ um difeomorfismo e p um ponto fixo para φ . Dizemos que φ é *dissipativo* (ou *localmente dissipativo*) se $|\det(D\varphi)_p| < 1$. De forma semelhante definimos para um ponto periódico de período k substituindo φ por φ^k .

Um fato que segue dessa definição é que se p , o ponto fixo para φ definido na Seção 3.1, é dissipativo, então para R suficientemente estreito todos os pontos periódicos no subconjunto invariante maximal $\Lambda \subset R$ serão dissipativos. De fato, se $q \in \Lambda$ é um ponto periódico de Ψ e, portanto, de φ e se R é suficientemente estreito, então q tem a “maioria” dos pontos de sua φ -órbita em uma pequena vizinhança de p . Na verdade, os iterados de q por φ passeiam pelas duas “alças” da Figura 3.16. Assim, tomando uma vizinhança U de p onde $|\det(D\Psi)_q| < 1$, para todo $q \in U$, e depois tomando R suficientemente pequeno (e estreito) de modo que as duas “alças” R e $\Psi(R)$ estejam em U , conseguimos fazer com que a maioria dos pontos da órbita de q fiquem em U . E como o $|\det(D\varphi)|$ é limitado (pois M é compacta e φ é contínua)

conseguimos obter que $|\det(D\varphi^k)_q| < 1$, uma vez que

$$\det(D\varphi^N)_q = \prod_{k=0}^{N-1} \det(D\varphi)_{\varphi^k(q)}.$$

Portanto, q é dissipativo.

Depois veremos a existência de cascatas de bifurcações de período e de poços e apresentaremos uma renormalização da família $\{\varphi_\mu\}$ no caso em que ela é dissipativa.

O desdobramento de uma tangência homoclínica rende um grande número de mudanças na dinâmica quando o parâmetro desenvolve, nesse caso, o surgimento de bifurcações. Em particular uma tangência homoclínica é um ponto de acumulação de outras tangências homoclínicas, veja o Teorema 6.

Vimos no capítulo anterior que que órbitas homoclínicas implicam na existência de órbitas caóticas. Quando desdobramos uma tangência homoclínica, esperamos dinâmicas caóticas e até mesmo atratores caóticos para diferentes aplicações. Nas seções seguintes descreveremos o fenômeno de bifurcação mencionado anteriormente que ocorre quando desdobramos uma tangência homoclínica quadrática. Veremos também uma ligação que existe entre este desdobramento e a *família quadrática* (veja [6]) de aplicações do intervalos da forma

$$f_\mu(y) = y^2 + \mu.$$

4.1 Cascatas de Tangências Homoclínicas

Seja $\varphi : M \times \mathbb{R} \longrightarrow M$ uma aplicação de classe C^3 tal que $\varphi_\mu(x) = \varphi(x, \mu)$ é um difeomorfismo sobre M para cada $\mu \in \mathbb{R}$. O motivo de tomarmos a família de classe C^3 é que adiante ao estudarmos o *flip* precisaremos dessa hipótese. Inicialmente estudaremos tangências homoclínicas e seus desdobramentos.

Definição 4.1.1. Seja $p = p_0$ um ponto fixo hiperbólico para φ_0 . Dizemos que o ponto q é uma *tangência homoclínica* associada a p se q é um ponto de tangência entre $W^s(p)$ e $W^u(p)$. Quando $W^s(p)$ e $W^u(p)$ têm contato quadrático (parabólico)

em q dizemos que q ou sua órbita $\mathcal{O}(q)$ é uma *tangência homoclínica quadrática*.

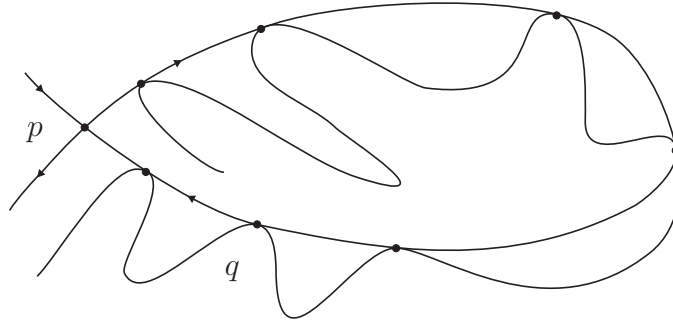


Figura 4.1: Tangência homoclínica quadrática.

Podemos escolher coordenadas locais (x_1, x_2) próximo de q de tal forma que as componentes locais de $W^s(p)$ e $W^u(p)$ que contém q tenham a seguinte forma

$$W^s(p) = \{(x_1, x_2) \mid x_2 = 0\}$$

$$W^u(p) = \{(x_1, x_2) \mid x_2 = ax_1^2\}$$

onde $a \neq 0$.

Antes de continuar, vamos definir quando uma propriedade ou condição é genérica.

Definição 4.1.2. Dado um conjunto topológico X , dizemos que uma propriedade é *genérica* se ela ocorre em um subconjunto de X que contém uma intersecção de conjuntos abertos e densos de X . Em particular, se tal propriedade ocorrer em um subconjunto aberto e denso ela também será genérica.

Genericamente, tangências homoclínicas são quadráticas (veja [4], Capítulo III, Seção 6), isto é, existem coordenadas locais, que dependem de μ , para as quais

$$W^s(p) = \{(x_1, x_2) \mid x_2 = 0\}$$

$$W^u(p) = \{(x_1, x_2) \mid ax_1^2 + bx_2 = \mu\}, \quad a \neq 0 \text{ e } b \neq 0.$$

Nesse caso, dizemos que a tangência homoclínica *desdobra genericamente*. Quando consideramos $a < 0$ e $b > 0$, obtemos a configuração dada na Figura 4.2. Note que

W^u sobe, tangencia W^s e depois a intercepta em dois pontos, quando μ varia de negativo para positivo.

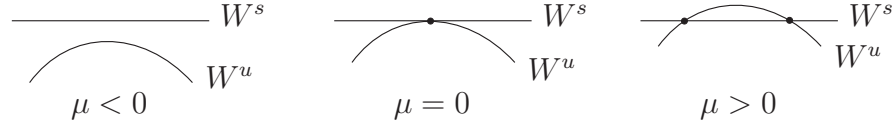


Figura 4.2: Posições das variedades W^s e W^u quando variamos o parâmetro μ .

Agora veremos um resultado que nos diz que, sobre certas condições, uma tangência homoclínica quadrática q para φ_0 pode ser aproximada por uma seqüência de tangências homoclínicas de aplicações da mesma família que φ_0 .

Teorema 6. *Seja $\{\varphi_\mu\}$ uma família a um parâmetro de difeomorfismos com uma tangência homoclínica quadrática q em $\mu = 0$ associada ao ponto hiperbólico fixo (ou periódico) p e suponha que ela desdobra genericamente. Então, existe uma seqüência $\mu_n \rightarrow 0$ tal que φ_{μ_n} têm tangências homoclínicas $q_{\mu_n} \rightarrow q$ associadas a $p_{\mu_n} \rightarrow p$.*

Demonstração:

Seja q como no enunciado e considere $r = \varphi_0^{-N}(q)$ para algum $N > 0$. Suponha que a tangência desdobra em pontos homoclínicos transversais para $\mu > 0$. Para $\mu > 0$ próximo de 0, existem pequenos pedaços de parábolas $\Gamma_\mu^u \subset W_\mu^u$ próximo de q e $\Gamma_\mu^s \subset W_\mu^s$ próximo de r que assumiremos estar com suas posições como na Figura 4.3.

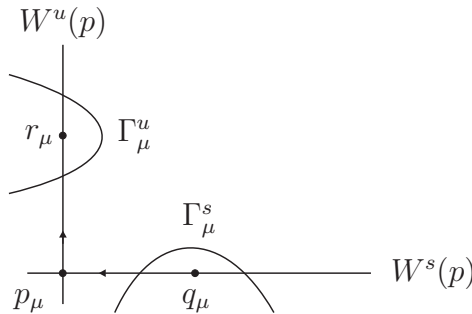


Figura 4.3: Pedaços de parábolas Γ_μ^s e Γ_μ^u .

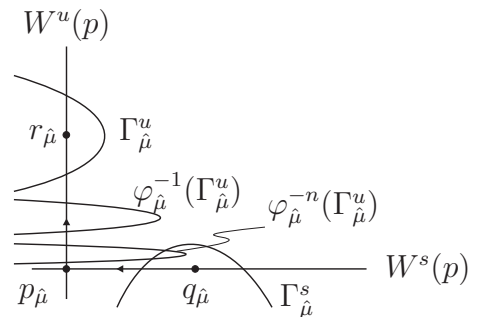


Figura 4.4: Intersecção entre as parábolas Γ_μ^u e $\varphi_\mu^{-n}(\Gamma_\mu^u)$.

Considere $\mu = \hat{\mu}$ arbitrariamente pequeno. Logo, para $n \in \mathbb{N}$ suficientemente grande, $\varphi_{\hat{\mu}}^{-n}(\Gamma_{\hat{\mu}}^s)$ intercepta $\Gamma_{\hat{\mu}}^u$, como na Figura 4.4. Tome agora $\mu > 0$ muito menor que $\hat{\mu}$ e temos que para o mesmo n vale que $\varphi_{\mu}^{-n}(\Gamma_{\mu}^s) \cap \Gamma_{\mu}^u = \emptyset$. Como $\varphi_{\mu}^{-n}(\Gamma_{\mu}^s)$ e Γ_{μ}^u dependem C^3 diferencialmente de μ , existe algum $0 < \mu_1 < \hat{\mu}$ para o qual $\varphi_{\mu_1}^{-n}(\Gamma_{\mu_1}^s)$ e $\Gamma_{\mu_1}^u$ são tangentes, digamos em $q_1 \in \Gamma_{\mu_1}^u$. Agora basta repetir o argumento para valores μ_n cada vez menores que $\hat{\mu}$ e construímos as seqüências $\mu_n \rightarrow 0$ e $q_n \rightarrow q$ desejadas. Dessa forma, provamos o caso indicado na Figura 4.3. Para os outros casos o argumento é análogo. A Figura 4.5 ilustra os outros casos. ■

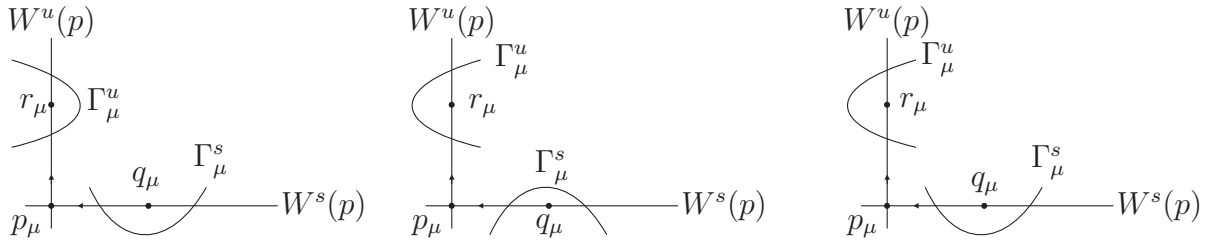


Figura 4.5: Outros casos.

Observação 10. Observe que o caso que demonstramos não ocorre no plano. Pode ocorrer por exemplo no bitoro.

Observação 11. Na demonstração do Teorema 6 podemos tomar as tangências homoclínicas q_{μ_n} com contato quadrático: devido à diferença das curvaturas, $\varphi_{\mu}^{-n}(\Gamma_{\mu}^q)$ e Γ_{μ}^u têm um contato quadrático em suas últimas tangências para valores decrescentes de μ . Pode-se ainda mostrar que essas tangências homoclínicas desdobram genericamente.

Observação 12. Pode-se ainda escolher os valores μ_n deste último teorema de tal modo que os ramos de $W^s(p_{\mu_n})$ e $W^u(p_{\mu_n})$, isto é, as componentes conexas de $W^s(p_{\mu_n}) \setminus \{p_{\mu_n}\}$ e $W^u(p_{\mu_n}) \setminus \{p_{\mu_n}\}$, as quais tem uma tangência homoclínica, também tenham intersecções homoclínicas transversais.

Vamos terminar esta seção com um fato que segue direto do Teorema da Função Implícita. Se um ponto fixo é hiperbólico então sua parte linear determina o tipo

de estabilidade. O seguinte resultado que diz que um ponto hiperbólico fixo persiste para pequenas mudanças na aplicação.

Teorema 7 ([6] p. 155). *Seja $\varphi_\mu(x)$ uma família a um parâmetro de aplicações diferenciáveis com $x \in \mathbb{R}^n$. Assuma que $\varphi_\mu(x)$ é de classe C^1 como função de ambas variáveis x e μ . Seja p_0 um ponto fixo hiperbólico para φ_{μ_0} . Suponha que $\varphi_{\mu_0}(x_0) = p_0$ e que 1 não seja autovalor de $D(\varphi_{\mu_0})_{p_0}$. Então existem*

- (i) *um conjunto aberto U que contém p_0 ,*
- (ii) *um intervalo aberto N que contém μ_0 e*
- (iii) *uma função $p : N \rightarrow U$ de classe C^1 tal que $p(\mu_0) = p_0$ e $\varphi_\mu(p(\mu)) = p(\mu)$.*

Além disso, para $\mu \in N$, a função φ_μ não possui outros pontos fixos em U exceto $p(\mu)$. E vale a seguinte derivação para a aplicação $p(\mu)$

$$p'(\mu) = -[D(\varphi_\mu)_{p(\mu)} - I]^{-1} \cdot \frac{\partial \varphi}{\partial \mu}(p(\mu), \mu).$$

Demonstração:

Queremos encontrar pontos x tais que $\varphi_\mu(x) = x$. Basta definir a função $G(x, \mu) = \varphi_\mu(x) - x$ e encontrar seus zeros. Note que temos as hipóteses $G(p_0, \mu_0) = 0$ e

$$\frac{\partial G}{\partial x}(p_0, \mu_0) = D(\varphi_{\mu_0})_{p_0} - I$$

é inversível, pois assumimos que $D(\varphi_{\mu_0})_{p_0}$ não tem autovalores iguais a 1. Logo, pelo Teorema da Função Implícita, existem abertos N que contém μ_0 , U que contém p_0 e uma aplicação $p(\mu) : N \rightarrow U$ de classe C^1 tal que $G(p(\mu), \mu) = 0$, para todo $\mu \in N$, e estes são os únicos zeros nessas vizinhanças. Além disso,

$$\begin{aligned} \varphi(p(\mu), \mu) = p(\mu) &\implies \frac{\partial \varphi}{\partial x}(p(\mu), \mu) \cdot p'(\mu) + \frac{\partial \varphi}{\partial \mu}(p(\mu), \mu) \cdot 1 = p'(\mu) \\ &\implies p'(\mu) = D(\varphi_\mu)_{p(\mu)} \cdot p'(\mu) + \frac{\partial \varphi}{\partial \mu}(p(\mu), \mu) \\ &\implies p'(\mu) = -[D(\varphi_\mu)_{p(\mu)} - I]^{-1} \cdot \frac{\partial \varphi}{\partial \mu}(p(\mu), \mu). \end{aligned}$$

O que conclui a demonstração. ■

A aplicação $p(\mu)$, que denotaremos por $\mu \mapsto p_\mu$, é chamada *continuação* do ponto fixo p . As componentes locais de $W^s(p)$ e $W^u(p)$ próximas de q dependem diferencialmente de μ , se μ está próximo de 0.

4.2 Bifurcações de Pontos Fixos

Nessa seção vamos considerar uma aplicação a um parâmetro $\varphi_\mu(x) = \varphi(x, \mu)$ onde $\mu \in \mathbb{R}$ e $x \in \mathbb{R}$. Já vimos no Teorema 7 que se $\varphi_{\mu_0}(x_0) = x_0$ é um ponto fixo hiperbólico, então o ponto fixo pode ser continuado para valores do parâmetro μ próximos de μ_0 . Esse não é um resultado de bifurcação. A principal ferramenta usada para mostrar que o ponto fixo pôde ser continuado foi o Teorema da Função Implícita, o qual utilizaremos no estudo das bifurcações consideradas nessa seção. O primeiro tipo de bifurcação, chamado *bifurcação sela-nó*, ocorre quando a hipótese acima sobre os autovalores é violada e 1 é um autovalor da derivada no ponto fixo. Sob certas condições sobre as derivadas superiores de $\varphi_\mu(x)$, obtemos que

- (i) para μ de um dos lados de μ_0 não ocorre pontos fixos próximos de x_0 ,
- (ii) e para μ do outro lado de μ_0 há dois pontos fixos. Desses, um é atrator e outro repulsor.

O outro tipo de bifurcação que vamos considerar são aquelas onde o ponto fixo persiste mas o tipo de estabilidade do ponto fixo muda quando μ passa por μ_0 . Esse segundo tipo, chamado de *bifurcação de duplicação de período* ou *bifurcação flip*, ocorre quando um autovalor é -1 . Nesse caso, um ponto atrator de período 1 torna-se repulsor em μ_0 e uma órbita estável de período 2 se ramifica a partir do ponto fixo.

Existe também um terceiro caso, chamado *bifurcação de Andronov-Hopf*, que ocorre quando temos um par de autovalores complexos, diferentes de ± 1 , com valor

absoluto 1. Porém como vamos exigir que a aplicação seja dissipativa e pelo fato de estarmos interessados no caso bidimensional, esse caso não ocorre.

4.2.1 Bifurcação Sela-nó

Nesta seção vamos estudar a bifurcação sela-nó, aquela que ocorre quando falhamos na hiperbolicidade admitindo que 1 é um dos autovalores da derivada. Exigiremos que $\varphi_{\mu_0}(x_0) = x_0$ e que $\varphi'_{\mu_0}(x_0) = 1$. Para que a tangência com o gráfico de φ_{μ_0} com a diagonal $\{(x, y) \mid y = x\}$ seja de maneira simples vamos exigir que $\varphi''_{\mu_0}(x_0) \neq 0$. E para que o gráfico de φ_μ translate de baixo para cima quando variar o parâmetro μ , pediremos que $\frac{\partial \varphi}{\partial \mu}(x_0, \mu_0) \neq 0$. Desse modo, para valores menores que μ_0 não se tem pontos fixos e para valores maiores que μ_0 surgem dois pontos fixos.

Exemplo 5. Seja $\varphi : \mathbb{R}^2 \rightarrow \mathbb{R}$ dada por $\varphi_\mu(x) = \mu + x - ax^2$, com $a > 0$. Queremos encontrar os pontos fixos de φ_μ , isto é, que pontos $x \in \mathbb{R}$ satisfazem $\varphi_\mu(x) - x = 0$, dependendo do parâmetro μ . Temos que

- se $\mu < 0$, então $\varphi_\mu(x) - x = \mu - ax^2 < 0$ para todo $x \in \mathbb{R}^2$, logo não existem pontos fixos para φ_μ nesse caso;
- se $\mu = 0$, então $\varphi_\mu(x) - x = ax^2 = 0$, logo temos um único ponto fixo para φ_0 , $x = 0$ e
- se $\mu > 0$, então $\varphi_\mu(x) - x = \mu - ax^2 = 0$. Logo, $x = \pm \sqrt{\frac{\mu}{a}}$ são os pontos fixos de φ_μ .

A Figura 4.6 ilustra esse exemplo.

Se plotarmos um gráfico dos pontos fixos por μ , obtemos a Figura 4.7. Note que temos dois gráficos de p_μ em função de μ , o que fica abaixo do ponto de bifurcação e o que fica acima.

Agora que já visualizamos a situação com um exemplo vamos ver o resultado geral.

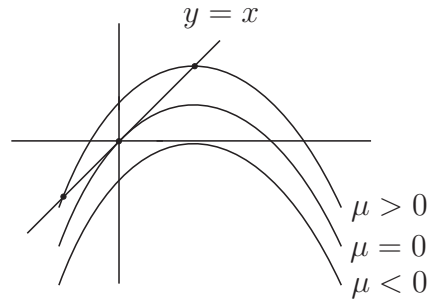


Figura 4.6: Os pontos fixos de φ_μ , nos casos $\mu < 0$, $\mu > 0$ e $\mu = 0$.

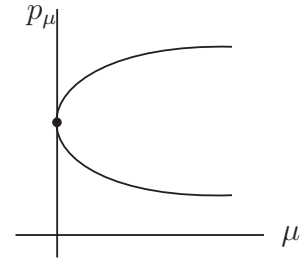


Figura 4.7: Gráficos de p_μ por μ .

Teorema 8 (Bifurcação Sela-nó). *Suponha que $\varphi : \mathbb{R}^2 \rightarrow \mathbb{R}$ é uma função de classe C^r em ambas as variáveis, onde $r \geq 2$. Suponha que*

- (1) $\varphi(x_0, \mu_0) = x_0$,
- (2) $\varphi'_{\mu_0}(x_0) = 1$,
- (3) $\varphi''_{\mu_0}(x_0) \neq 0$ e
- (4) $\frac{\partial \varphi}{\partial \mu}(x_0, \mu_0) \neq 0$.

Então, existem intervalos I que contém x_0 e N que contém μ_0 e uma função $m : I \rightarrow N$ de classe C^r tais que

- (i) $\varphi_{m(x)}(x) = x$ para todo $x \in I$,
- (ii) $m(x_0) = \mu_0$ e
- (iii) o gráfico de m nos fornece todos os pontos fixos em $I \times N$.

Além disso, $m'(x_0) = 0$ e

$$m''(x_0) = -\frac{\frac{\partial^2 \varphi}{\partial x^2}(x_0, \mu_0)}{\frac{\partial \varphi}{\partial \mu}(x_0, \mu_0)} \neq 0.$$

Estes pontos fixos são atratores de um lado e repulsores do outro.

Demonstração:

Para encontrar os pontos fixos de φ_μ , vamos considerar a aplicação $G(x, \mu) = \varphi(x, \mu) - x$. Então, os pontos fixos de φ correspondem aos zeros de G . Pela hipótese (1), temos que $G(x_0, \mu_0) = 0$. Além disso, observe que segue da hipótese (2)

$$\frac{\partial G}{\partial x}(x_0, \mu_0) = \varphi'_{\mu_0}(x_0) - 1 = 0,$$

logo não podemos encontrar x em função de y . No entanto, pela hipótese (4), temos

$$\frac{\partial G}{\partial \mu}(x_0, \mu_0) = \frac{\partial \varphi}{\partial \mu}(x_0, \mu_0) \neq 0,$$

logo podemos encontrar μ em função de x . Pelo Teorema da Função Implícita, existem intervalos abertos I que contém x_0 e N que contém μ_0 e uma função $m : I \rightarrow N$ de classe C^r tal que para todo $x \in I$, existe um único $m(x) \in N$ tal que $G(x, m(x)) \equiv 0$. Essa última igualdade nos traz que $\varphi_{m(x)}(x) = x$, o que prova (i). Como os zeros em $I \times N$ são únicos, $\varphi_{\mu_0}(x_0) = x_0$ e $\varphi_{m(x_0)}(x_0) = x_0$, segue que $m(x_0) = x_0$, o que prova (ii). E (iii) é direto.

Para calcular a deriva de $m(x)$ usaremos a derivação implícita. Por simplicidade de notação utilizaremos $G_x = \frac{\partial G}{\partial x}$ e semelhante para os outros casos. Diferenciamos $0 = G(x, m(x))$ em relação a x e obtemos

$$0 = G_x(x, m(x)) + G_\mu(x, m(x)) \cdot m'(x). \quad (4.1)$$

Avaliando em x_0 e utilizando o fato que $G_x(x_0, \mu_0) = 0$ e $G_\mu(x_0, \mu_0) = 0$, obtemos que

$$m'(x_0) = \frac{-G_x(x_0, \mu_0)}{G_\mu(x_0, \mu_0)} = 0.$$

Para obter a segunda derivada de m , diferenciamos (4.1) uma vez e obtemos

$$\begin{aligned} 0 &= G_{xx} + G_{x\mu} \cdot m' + (G_{\mu x} + G_{\mu\mu} \cdot m')m' + G_\mu \cdot m'' \\ &= G_{xx} + 2G_{\mu x} \cdot m' + G_{\mu\mu} (m')^2 + G_\mu \cdot m''. \end{aligned}$$

Avaliando a última expressão em x_0 e usando o fato que $m'(x_0) = 0$, obtemos que

$$m''(x_0) = \frac{-G_{xx}(x_0, m(x_0))}{G_\mu(x_0, m(x_0))} = \frac{-\frac{\partial^2 \varphi}{\partial x^2}(x_0, \mu_0)}{\frac{\partial \varphi}{\partial \mu}(x_0, \mu_0)},$$

como queríamos.

Para encontrar a estabilidade dos pontos fixos usamos a expansão em Série de Taylor de φ_x em torno de (x_0, μ_0) :

$$\begin{aligned} \frac{\partial \varphi}{\partial x}(x, \mu) &= \frac{\partial \varphi}{\partial x}(x_0, \mu_0) + D\left(\frac{\partial \varphi}{\partial x}\right)_{(x_0, \mu_0)} \cdot (x - x_0, \mu - \mu_0) + O(|(x - x_0, \mu - \mu_0)|^2) \\ &= 1 + \frac{\partial^2 \varphi}{\partial x^2}(x_0, \mu_0) \cdot (x - x_0) + \frac{\partial^2 \varphi}{\partial \mu \partial x}(x_0, \mu_0) \cdot (\mu - \mu_0) + \\ &\quad + O(|x - x_0|^2) + O(|x - x_0| \cdot |\mu - \mu_0|) + O(|\mu - \mu_0|^2). \end{aligned}$$

Como $m'(x_0) = 0$, temos que a expansão em Série de Taylor de m em torno de x_0 é

$$m(x) = m(x_0) + m'(x_0)(x - x_0) + O(|x - x_0|^2),$$

logo,

$$m(x) - \mu_0 = O(|x - x_0|^2).$$

Portanto,

$$\frac{\partial \varphi}{\partial x}(x, m(x)) = 1 + \frac{\partial^2 \varphi}{\partial x^2}(x_0, \mu_0) \cdot (x - x_0) + O(|x - x_0|^2)$$

Como $\frac{\partial^2 \varphi}{\partial x^2}(x_0, \mu_0) \neq 0$, então $\frac{\partial \varphi}{\partial x}(x, m(x)) - 1$ é maior que 1 de um lado de x_0 e menor que 1 do outro lado, dependendo do sinal de $\frac{\partial^2 \varphi}{\partial x^2}(x_0, \mu_0)$. Portanto, temos pontos fixos repulsores de um lado de x_0 e pontos fixos atratores do outro. ■

4.2.2 Bifurcação de Duplicação de Período

Nesta seção vamos estudar a bifurcação de duplicação de período ou flip, onde ocorre um ponto fixo que torna-se repulsor quando sua derivada é igual a -1 além de surgir uma órbita de período dois. Começamos vendo este exemplo no qual o ponto fixo $x = 0$ não varia com o parâmetro.

Exemplo 6. Seja $\varphi_\mu : \mathbb{R} \rightarrow \mathbb{R}$, $\mu \in \mathbb{R}$ dada por $\varphi_\mu(x) = -\mu x + ax^2 + bx^3$. Estamos com a hipótese que caracteriza o flip, ou seja, $\varphi'_1(0) = -1$. Queremos encontrar os pontos de período dois para parâmetros μ próximos de 1. Para isso vamos encontrar a cara de $\varphi_\mu^2(x)$ que é

$$\varphi_\mu^2(x) = \mu^2 x + x^2(-a\mu + a\mu^2) + x^3(-b\mu - 2a^2\mu - b\mu^3) + O(x^4).$$

Aqui $O(x^4)$ quer dizer termos de ordem maior ou igual a 4 que não descrevemos pois não os utilizaremos. Queremos encontrar os zeros de $\varphi_\mu^2(x) - x$. Como $\varphi_\mu(0) = 0$ obtemos que $\varphi_\mu^2(0) - 0 = 0$, já que x é um fator comum de $\varphi_\mu^2(x) - x$. Para encontrar os zeros de $\varphi_\mu^2(x) - x$ exceto $x = 0$, definimos a seguinte função

$$\begin{aligned} M(x, \mu) &= \frac{\varphi_\mu^2(x) - x}{x} \\ &= \mu^2 - 1 + x(-a\mu + a\mu^2) + x^2(-b\mu - 2a^2\mu - b\mu^3) + O(x^3). \end{aligned}$$

Esta função se anula no ponto $(x, \mu) = (0, 1)$, $M(0, 1) = 0$. Usando o fato que

$$\mu^2 - 1 = (\mu - 1)(\mu + 1) \approx 2(\mu - 1),$$

para $\mu \approx 1$, temos que os zeros de $M(x, \mu)$ são aproximadamente iguais aos zeros de

$$2(\mu - 1) - 2(b + a^2)x^2,$$

que são

$$\mu = 1 + 2(b + a^2)x^2 \quad \text{ou}$$

$$x = \pm \sqrt{\frac{\mu - 1}{b + a^2}}, \quad \text{para } \frac{\mu - 1}{b + a^2} > 0.$$

As derivadas parciais de M em $(0, 1)$ são

$$M_x(0, 1) = (-a\mu + a\mu)|_{\mu=1} = 0 \quad \text{e}$$

$$M_\mu(0, 1) = 2\mu|_{\mu=1} = 2 \neq 0.$$

Pelo Teorema da Função Implícita, podemos expressar μ em função de x para obtermos os zeros de M , onde $\mu = m(x)$ satisfaz $M(x, m(x)) = 0$, logo $\varphi_{m(x)}^2(x) - x = 0$. Dessa forma, podemos justificar as aproximações feitas acima calculando a derivada de $m(x)$ implicitamente e obter o coeficiente do termo de ordem 2 visto anteriormente. Vamos diferenciar duas vezes a expressão $0 = M(x, m(x))$ em relação a x para obtermos

$$\begin{aligned} 0 &= M_x(x, m(x)) + M_\mu(x, m(x)) \cdot m'(x) \quad \text{e} \\ 0 &= M_{xx}(x, m(x)) + M_{\mu x}(x, m(x)) \cdot m'(x) + M_{x\mu}(x, m(x)) \cdot m'(x) + \\ &\quad + M_{\mu\mu}(x, m(x)) \cdot (m'(x))^2 + M_\mu(x, m(x)) \cdot m''(x) \\ &= M_{xx}(x, m(x)) + 2M_{x\mu}(x, m(x)) \cdot m'(x) + M_{\mu\mu}(x, m(x)) \cdot (m'(x))^2 \\ &\quad + M_\mu(x, m(x)) \cdot m''(x). \end{aligned}$$

Como $M_x(0, 1) = 0$ e $M_\mu(0, 1) \neq 0$, segue da primeira expressão que $m'(0) = 0$. Já na segunda expressão, como $m'(0)$ anula todos os termos exceto M_{xx} , é este que devemos calcular para determinar $m''(0)$. Usando a expressão explícita obtida anteriormente, temos que

$$M_{xx}(0, 1) = -2(b\mu - 2a^2\mu - b\mu^2)|_{\mu=1} = -4(b + a^2).$$

Então,

$$m''(0) = \frac{-M_{xx}(0, 1)}{M_\mu(0, 1)} = 2(b + a^2).$$

Portanto, para obter uma expressão quadrática para os novos pontos de período dois devemos assumir que $b + a^2 \neq 0$. No resultado geral veremos que o sinal de $b + a^2$ também determina a estabilidade da órbita de período dois. Note que $-2(b + a^2)$ é o coeficiente do termo x^3 em φ_1^2 , onde 1 é o parâmetro de bifurcação.

Apesar deste exemplo ser bastante geral, nele o ponto fixo não varia com o parâmetro, logo $\frac{\partial \varphi}{\partial \mu}(0, 1) = 0$. No teorema que veremos adiante, assumiremos que $\frac{\partial \varphi}{\partial \mu}(0, 1) \neq 0$ e analisaremos o efeito deste termo. Além disso, vamos impor uma condição para que a derivada espacial (em relação a x) de φ varie ao longo da curva dos pontos fixos conforme varia o parâmetro.

Teorema 9 (Bifurcação de Duplicação de Período). *Assuma que $\varphi : \mathbb{R}^2 \rightarrow \mathbb{R}$ é uma aplicação de classe C^r , $r \geq 1$, e que φ satisfaz as seguintes condições*

- (1) *o ponto $x_0 \in \mathbb{R}$ é um ponto fixo para $\mu = \mu_0$, isto é, $\varphi(x_0, \mu_0) = x_0$;*
- (2) *$\varphi'_{\mu_0}(x_0) = -1$ (como esta derivada não é igual a 1, existe uma curva de pontos fixos $x(\mu)$ para μ próximo de μ_0);*
- (3) *a derivada de $\varphi'_\mu(x(\mu))$ em relação a μ é não nula (variando ao longo da curva dos pontos fixos):*

$$\alpha = \left[\frac{\partial^2 \varphi}{\partial \mu \partial x} + \frac{1}{2} \cdot \frac{\partial \varphi}{\partial \mu} \cdot \frac{\partial^2 \varphi}{\partial x^2} \right] \Big|_{(x_0, \mu_0)} \neq 0 \quad e$$

- (4) *o gráfico de $\varphi_{\mu_0}^2$ possui o termo cúbico não nulo em sua tangência com a diagonal:*

$$\beta = \left(\frac{1}{3!} \cdot \frac{\partial^3 \varphi}{\partial x^3}(x_0, \mu_0) \right) + \left(\frac{1}{2!} \cdot \frac{\partial^2 \varphi}{\partial x^2}(x_0, \mu_0) \right)^2 \neq 0.$$

Então, existe uma curva diferenciável de pontos fixos $x(\mu)$ que passa por x_0 em μ_0 , a estabilidade do ponto fixo muda em μ_0 e o lado de μ_0 que é atrator depende do sinal

de α . Existe uma curva diferenciável γ que passa por (x_0, μ_0) tal que $\gamma \setminus \{(x_0, \mu_0)\}$ é a união de órbitas hiperbólicas de período dois. A curva γ é tangente à reta $\mathbb{R} \times \{\mu_0\}$ em (x_0, μ_0) e é gráfico de uma função de x , $\mu = m(x)$, com $m'(x_0) = 0$ e $m''(x_0) = \frac{-2\beta}{\alpha}$. O tipo de estabilidade da órbita de período dois depende do sinal de β : se $\beta < 0$ então a órbita de período dois é atratora e se $\beta > 0$, a órbita esta é repulsora.

Demonstração:

Do Teorema 7 juntamente com a hipótese **(2)** segue que existe uma curva de pontos fixos parametrizada por μ , $x(\mu)$, e que sua derivada é da forma

$$\begin{aligned} x'(\mu) &= -\left(\frac{\partial \varphi}{\partial x}(x_0, \mu_0) - 1\right)^{-1} \cdot \frac{\partial \varphi}{\partial \mu}(x_0, \mu_0) \\ &= -(-1 - 1)^{-1} \cdot \frac{\partial \varphi}{\partial \mu}(x_0, \mu_0) = \frac{1}{2} \cdot \frac{\partial \varphi}{\partial \mu}(x_0, \mu_0). \end{aligned}$$

Queremos transladar as coordenadas de tal modo que 0 torne-se um ponto fixo para todo μ numa vizinhança, para isso, definimos

$$g(y, \mu) = \varphi_\mu(y + x(\mu)) - x(\mu).$$

Então, $g(0, \mu) = \varphi_\mu(x(\mu)) - x(\mu) \equiv 0$. Denotaremos $g(\cdot, \mu) \circ g(y, \mu)$ por $g^2(y, \mu)$. Note que as derivadas parciais de g em relação a y coincidem com as de φ em relação a x nos correspondentes pontos

$$\frac{\partial^j g}{\partial y^j}(0, \mu) = \frac{\partial^j \varphi}{\partial x^j}(x(\mu), \mu).$$

O valor da derivada espacial em relação a posição, $\frac{\partial g}{\partial y}(0, \mu)$, determina a estabilidade do ponto fixo, logo $\frac{\partial^2 g}{\partial \mu \partial y}(0, \mu)$ mede a variação ao longo da curva de pontos fixos e

$$\begin{aligned}
\frac{\partial^2 g}{\partial \mu \partial y}(0, \mu_0) &= \frac{\partial}{\partial \mu} \frac{\partial \varphi}{\partial x}(x(\mu, \mu)) \Big|_{\mu=\mu_0} \\
&= \frac{\partial^2 \varphi}{\partial \mu \partial x}(x(\mu_0), \mu_0) + \frac{\partial^2 \varphi}{\partial x^2}(x(\mu_0), \mu_0) \cdot x'(\mu_0) \\
&= \frac{\partial^2 \varphi}{\partial \mu \partial x}(x(\mu_0), \mu_0) + \frac{1}{2} \cdot \frac{\partial^2 \varphi}{\partial x^2}(x_0, \mu_0) \cdot \frac{\partial \varphi}{\partial \mu}(x_0, \mu_0) = \alpha \neq 0.
\end{aligned}$$

Este cálculo mostra que α mede o que foi descrito no enunciado do teorema.

Seja $a_j(\mu)$ o coeficiente de y^j na expansão em série de Taylor de g em torno de $y = 0$, logo temos

$$g(y, \mu) = a_1(\mu)y + a_2(\mu)y^2 + a_3(\mu)y^3 + O(y^4).$$

Para obter a expressão de $g^2(\mu, y)$ basta compor $g(y, \mu)$ consigo mesma e obtemos

$$g^2(y, \mu) = a_1^2 y + (a_1 a_2 + a_2 a_1^2) y^2 + (a_1 a_3 + 2a_1 a_2^2 + a_3 a_1^3) y^3 + O(y^4).$$

Note que omitimos a dependência de μ dos coeficientes a_j e como feito acima não descrevemos os termos em y^4 pois não nos interessam. Agora, queremos encontrar os zeros de $g^2(y, \mu) - y = 0$ que não sejam o próprio $y = 0$, pois este é sempre uma solução. Para isso, dividimos essa expressão por y e definimos

$$M(y, \mu) = \begin{cases} \frac{g^2(y, \mu) - y}{y}, & \text{para } y \neq 0; \\ \frac{\partial}{\partial y}(g^2(y, \mu)) \Big|_{y=0} - 1, & \text{para } y = 0. \end{cases}$$

Note que na definição para $y = 0$ apenas fizemos um limite de $y \rightarrow 0$. Usando a expressão de g^2 , obtemos

$$M(y, \mu) = (a^2 - 1) + (a_1 a_2 + a_2 a_1^2) y + (a_1 a_3 + 2a_1 a_2^2 + a_3 a_1^3) y^2 + O(y^3).$$

Para que possamos aplicar o Teorema da Função Implícita, observe que

$$a_1(\mu_0) = \frac{\partial g}{\partial y}(0, \mu_0) = \frac{\partial \varphi}{\partial x}(x(\mu_0), \mu_0) = \varphi'_{\mu_0}(x_0) = -1.$$

Logo,

$$M(0, \mu_0) = [(a_1^2 - 1) + 0] \Big|_{\mu_0} = (-1)^2 - 1 = 0$$

e as derivadas parciais são

$$\begin{aligned} M_y(0, \mu_0) &= a_1 a_2 + a_2 a_1^2 \Big|_{\mu_0} = -a_2(\mu_0) + a_2(\mu_0) = 0 \quad \text{e} \\ M_\mu(0, \mu_0) &= \frac{\partial}{\partial \mu} \frac{\partial g^2}{\partial y}(0, \mu) \Big|_{\mu=\mu_0} = \frac{\partial}{\partial \mu} (a(\mu))^2 \Big|_{\mu=\mu_0} = \frac{\partial}{\partial \mu} \left(\frac{\partial g}{\partial y}(0, \mu) \right)^2 \Big|_{\mu=\mu_0} \\ &= 2 \cdot \frac{\partial g}{\partial y}(0, \mu_0) \cdot \frac{\partial^2 g}{\partial \mu \partial y}(0, \mu_0) = 2 \cdot (-1) \cdot \alpha = -2\alpha \neq 0. \end{aligned}$$

Como $M_\mu(0, \mu_0) \neq 0$, pelo Teorema da Função Implícita, existe uma função diferenciável $m(y)$ tal que $M(y, m(y)) \equiv 0$. Assim, diferenciando implicitamente essa última expressão obtemos

$$\begin{aligned} 0 &= M_y(0, m(y)) + M_\mu(y, m(y)) \cdot m'(y) \Big|_{y=0} \\ &= M_y(0, \mu_0) + M_\mu(0, \mu_0) \cdot m'(0). \end{aligned}$$

Logo,

$$m'(0) = \frac{-M_y(0, \mu_0)}{M_\mu(0, \mu_0)} = 0.$$

Para obtermos a segunda derivada de $m(y)$, basta derivar duas vezes a expressão $0 = M(y, m(y))$,

$$\begin{aligned} 0 &= M_{yy} + M_{\mu y} \cdot m' + (M_{y\mu} + M_{\mu\mu} \cdot m') \cdot m' + M_y \cdot m'' \\ &= M_{yy} + 2 \cdot M_{\mu y} \cdot m' + M_{\mu\mu} \cdot (m')^2 + M_y \cdot m''. \end{aligned}$$

Avaliando em $y = 0$ e lembrando que $m'(0) = 0$ obtemos

$$m''(0) = \frac{-M_{yy}(0, \mu_0)}{M_y(0, \mu_0)}.$$

Logo, para obter $m''(0)$ basta calcular o numerador

$$\begin{aligned} M_{yy}(0, \mu_0) &= 2(a_1 a_3 + 2a_1 a_2^2 + a_3 a_1^3) \Big|_{\mu=\mu_0} = 2(-a_3 - 2a_2^2 - a_3) \Big|_{\mu=\mu_0} \\ &= -4(a_3(\mu_0) + (a_2(\mu_0))^2) = -4 \left[\frac{1}{3!} \cdot \frac{\partial^3 g}{\partial y^3}(0, \mu_0) + \left(\frac{1}{2!} \cdot \frac{\partial^2 g}{\partial y^2}(0, \mu_0) \right)^2 \right] \\ &= -4 \left[\frac{1}{3!} \cdot \frac{\partial^3 \varphi}{\partial x^3}(x_0, \mu_0) + \left(\frac{1}{2!} \cdot \frac{\partial^2 \varphi}{\partial x^2}(x_0, \mu_0) \right)^2 \right] = -4\beta \neq 0. \end{aligned}$$

Portanto,

$$m''(0) = \frac{-(-4\beta)}{-2\alpha} = \frac{-2\beta}{\alpha} \neq 0.$$

Além disso, podemos checar que

$$\begin{aligned} \frac{\partial^3 g^2}{\partial y^3}(0, \mu_0) &= 6(a_1 a_3 + 2a_1 a_2^2 + a_3 a_1^3) \Big|_{\mu=\mu_0} = 6(a_3(\mu_0) - 2a_2(\mu_0)^2 - a_3(\mu_0)) \\ &= -12(a_3(\mu_0) + a_2(\mu_0)^2) = 3M_{yy}(0, \mu_0) = -12\beta \neq 0. \end{aligned}$$

Agora resta apenas checar a estabilidade da órbita de período dois. Para isso, utilizaremos a expansão e série de Taylor de $\frac{\partial(g^2)}{\partial y^2}(y, m(y))$ em torno do ponto $(y, \mu) = (0, \mu_0)$:

$$\begin{aligned} \frac{\partial(g^2)}{\partial y}(y, m(y)) &= \frac{\partial(g^2)}{\partial y}(0, \mu_0) + \frac{\partial^2(g^2)}{\partial y^2}(0, \mu_0) \cdot y^2 + \frac{\partial^2(g^2)}{\partial \mu \partial y}(0, \mu_0) \cdot (\mu - \mu_0) + \\ &\quad + \frac{1}{2} \cdot \frac{\partial^3(g^2)}{\partial y^3}(0, \mu_0) \cdot y^2 + O(|y|^3 |\mu_0|^2). \end{aligned}$$

Já calculamos a maioria dos coeficientes

- $\frac{\partial(g^2)}{\partial y}(0, \mu_0) = (-1)^2$, pela regra da cadeia;
- $\frac{\partial^2(g^2)}{\partial y^2}(0, \mu_0) = 0$, pois da expressão explícita de $g^2(\mu, x)$ obtemos que $\frac{\partial^2(g^2)}{\partial y^2}(0, \mu_0) = M_y(0, \mu_0)$;
- $\frac{\partial(g^2)}{\partial \mu \partial y}(0, \mu_0) = M_\mu(0, \mu_0) = -2\alpha$.

Além disso, como $m(y) = m(0) + m'(0)y + \frac{1}{2}m''(0)y^2 + O(y^3)$, temos que

$$\begin{aligned} \frac{\partial^2(g^2)}{\partial\mu\partial y}(0, \mu_0) \cdot (m(y) - \mu_0) &= M_\mu(0, \mu_0) \cdot \frac{1}{2} \cdot m''(0) \cdot y^2 + O(y^3) \\ &+ M_\mu(0, \mu_0) \cdot \frac{1}{2} \cdot \frac{-M_{yy}(0, \mu_0)}{M_\mu(0, \mu_0)} \cdot y^2 + O(y^3) \\ &= 2\beta y^2 + O(y^3). \end{aligned}$$

E para finalizar, temos que

$$\frac{1}{2} \cdot \frac{\partial^3(g^2)}{\partial y^3}(0, \mu_0) = \frac{1}{2} \cdot 12\beta = -6\beta.$$

Combinando estes termos, obtemos a seguinte expressão

$$\begin{aligned} \frac{\partial^2(g^2)}{\partial y}(y, m(y)) &= 1 + 2\beta y^2 - 6\beta y^2 + O(y^3) \\ &= 1 - 4\beta y^2 + O(y^3). \end{aligned}$$

Portanto, considerando que o termo $O(y^3)$ não tem influência para y pequeno, segue que se $\beta > 0$ a órbita de período dois é atratora, uma vez que $\frac{\partial^2(g^2)}{\partial y^2}(y, m(y)) < 1$, e no caso em que $\beta < 0$ ela é repulsora. O que conclui a demonstração. ■

4.3 Formas Genéricas das Bifurcações Sela-nó e Flip

Como vimos no Capítulo 3, órbitas homoclínicas transversais implicam na existência de ferraduras. Logo, próximo de uma bifurcação homoclínica, esperamos o surgimento ou a destruição de ferraduras. Vimos na última seção exemplos de bifurcações genéricas de pontos fixos em famílias de difeomorfismos a um parâmetro. Recordando, para famílias genéricas temos três possíveis casos que dependem de quais são os autovalores ρ_1 e ρ_2 :

- (a) $\rho_1 = 1$ e $|\rho_2| < 1$ (ou $|\rho_2| > 1$);
- (b) $\rho_1 = -1$ e $|\rho_2| < 1$ (ou $|\rho_2| > 1$) e
- (c) $\rho_1 = e^{i\theta}$ e $\rho_2 = e^{-i\theta}$, para algum real $\theta \neq k\pi, k \in \mathbb{Z}$.

O caso (c) é o que não se enquadra pois estamos exigindo que nossas aplicações sejam dissipativas, isto é, que o determinante da derivada no ponto fixo tenha valor absoluto menor que 1, mas nesse caso o determinante é $|\rho_1 \cdot \rho_2| = |e^{i\theta} \cdot e^{-i\theta}| = 1$.

Para descrever os casos (a) e (b), vamos antes ver um resultado fundamental sobre variedades invariantes. Este resultado pode ser encontrado em [6] p.197. Vamos utilizar a notação $m(Df_p|_{\mathbb{E}^u}) = \inf_{\|v\| \neq 1} \frac{\|Df_p|_{\mathbb{E}^u} v\|}{\|v\|}$

Teorema 10 (Teorema da Variedade Central). *Seja $f : U \subset \mathbb{R}^n \rightarrow \mathbb{R}^n$ uma aplicação de classe C^k para $1 \leq k \leq \infty$ com $f(0) = 0$ e $A = Df_0$. Seja k' escolhido do seguinte modo*

- (i) k , se $k < \infty$
- (ii) algum inteiro satisfazendo $1 \leq k' < \infty$, se $k = \infty$.

Suponha que $0 < \sigma < 1 < \lambda$ e as normas sobre $\mathbb{E}^u$ e $\mathbb{E}^s$ são escolhidas de tal modo que $\|Df_p|_{\mathbb{E}^s}\| < \mu$ e $m(Df_p|_{\mathbb{E}^u}) > \lambda$. Seja $\varepsilon > 0$ pequeno o bastante para que $\|Df_p|_{\mathbb{E}^s}\| < \sigma - \varepsilon$ e $m(Df_p|_{\mathbb{E}^u}) > \lambda + \varepsilon$. Seja $r > 0$ e $\bar{f} : \mathbb{R}^n \rightarrow \mathbb{R}^n$ uma aplicação de classe $C^{k'}$ que satisfaz $\bar{f}|_{B(0,r)} = f|_{B(0,r)}$, $\|\bar{f} - A\|_{C^1} < \varepsilon$ e $\bar{f} = A$ fora de $B(0, 2r)$. Se $r > 0$ é pequeno o suficiente, então existe uma variedade centro-estável invariante $W^{cs}(0, \bar{f})$ de classe $C^{k'}$ que é um gráfico sobre $\mathbb{E}^c \oplus \mathbb{E}^s$, a qual é tangente a $\mathbb{E}^c \oplus \mathbb{E}^s$ em 0 e é caracterizada como segue

$$W^{cs}(0, \bar{f}) = \{q \in \mathbb{R}^n \mid d(\bar{f}^j(q), 0) \cdot \lambda^{-j} \rightarrow 0 \text{ quando } j \rightarrow \infty\}.$$

Isso significa que $d(\bar{f}^j(q), 0)$ cresce mais lentamente do que λ^j .

De forma semelhante, existe uma variedade central-instável invariante $W^{cu}(0, \bar{f})$ de classe $C^{k'}$ que é um gráfico sobre $\mathbb{E}^c \oplus \mathbb{E}^u$, é tangente a $\mathbb{E}^c \oplus \mathbb{E}^s$ em 0 e é

caracterizada como segue

$$W^{cu}(0, \bar{f}) = \{q \in \mathbb{R}^n \mid d(\bar{f}^{-j}(q), 0) \cdot \sigma^j \rightarrow 0 \text{ quando } j \rightarrow \infty\}.$$

Isso significa que $d(\bar{f}^{-j}(q), 0)$ cresce mais lentamente do que σ^{-j} quando $j \rightarrow \infty$.

Definição 4.3.1. Seja f como no teorema anterior. A variedade central da extensão $\bar{f}$ é definida por

$$W^c(0, \bar{f}) = W^{cs}(0, \bar{f}) \cap W^{cu}(0, \bar{f}),$$

que é de classe C^k e é tangente a $\mathbb{E}^c$. As variedades *centro-estável local*, *centro-instável local* e *central local* para f são definidas por

$$W^{cs}(0, B(0, r), f) = W^{cs}(0, \bar{f}) \cap B(0, r),$$

$$W^{cu}(0, B(0, r), f) = W^{cu}(0, \bar{f}) \cap B(0, r) \text{ e}$$

$$W^c(0, B(0, r), f) = W^c(0, \bar{f}) \cap B(0, r),$$

respectivamente. Estas variedades locais dependem da extensão, mas se $f^j(q)$ pertence a $B(0, r)$ para todo $j \in \mathbb{Z}$, então $q \in W^c(0, \bar{f})$ para toda extensão $\bar{f}$ e $q \in W^c(0, B(0, r), f)$. Quando for claro, omitiremos a bola $B(0, r)$ e a função das notações acima, por exemplo, $W^c(0)$ será o mesmo que $W^c(0, B(0, r), f)$.

Observação 13. A norma utilizada no enunciado para comparar $\bar{f}$ com A é dada por

$$\|\bar{f} - A\|_{C^1} = \sup\{|\bar{f}(x) - A(x)|, \|D\bar{f}_x - DA_x\|; x \in V\},$$

onde V é algum conjunto onde as aplicações $\bar{f}$ e A estão definidas.

Em posse deste resultado, para os casos **(a)** e **(b)**, segue que existe uma curva W_μ^c (ou $W^c(p_\mu)$) de classe C^3 , φ_μ -invariante, que é tangente ao subespaço associado à $\rho_1 = 1$ ou $\rho_1 = -1$, o qual denotamos por $\mathbb{E}_\mu^c$. Portanto, se denotarmos $f_\mu = \varphi_\mu|_{W_\mu^c}$ temos a seguinte expansão em série de Taylor para f_μ em torno de $(0, \mu_0)$

$$\begin{aligned}
f_\mu(x) &= f_{\mu_0}(0) + Df_{(0,\mu_0)}(x-0, \mu-\mu_0) + \frac{1}{2!} \cdot D^2 f_{(0,\mu_0)}(x-0, \mu-\mu_0)^2 + \\
&\quad + \frac{1}{3!} \cdot D^3 f_{(0,\mu_0)}(x-0, \mu-\mu_0)^3 + t.o.s. \\
&= f_{\mu_0}(0) \cdot x + \frac{\partial f}{\partial \mu}(0, \mu_0) \cdot (\mu - \mu_0) + \frac{1}{2} \cdot \frac{\partial^2 f}{\partial x^2}(0, \mu_0) \cdot x^2 + \\
&\quad + \frac{\partial^2 f}{\partial x \partial \mu}(0, \mu_0) \cdot x \cdot (\mu - \mu_0) + \frac{1}{2} \cdot \frac{\partial^2 f}{\partial \mu^2}(0, \mu_0) \cdot (\mu - \mu_0)^2 + \\
&\quad + \frac{1}{3!} \cdot \frac{\partial^3 f}{\partial x^3}(0, \mu_0) \cdot x^3 + t.o.s.
\end{aligned}$$

onde *t.o.s.* significa termos de ordem superior. Logo, para o caso **(a)** temos a seguinte expressão local

$$f_\mu(x) = x + ax^2 + b(\mu - \mu_0) + t.o.s. \quad (4.2)$$

onde $a = \frac{1}{2} \cdot \frac{\partial^2 f}{\partial x^2}(0, \mu_0) = \frac{1}{2} \cdot f''_{\mu_0}(0)$ e $b(\mu - \mu_0) = \frac{\partial f}{\partial \mu}(0, \mu_0)(\mu - \mu_0)$. Além disso, $b(0) = 0$.

Para o caso **(b)** precisaremos fazer uma mudança de coordenadas a fim de eliminar o termo em x^2 e o termo em $(\mu - \mu_0)$.

Afirmção. O caso **(b)** tem a seguinte expressão local

$$f_\mu(x) = -x + \kappa x^3 + \tau(\mu - \mu_0) + t.o.s. \quad (4.3)$$

onde $\tau(0) = 0$.

De fato: da expansão em série de Taylor de $f_\mu(x)$ em torno de $(0, \mu_0)$, obtemos a seguinte expressão

$$f_\mu(x) = -x + \alpha x^2 + \beta(\mu - \mu_0) + \gamma x^3 + \delta(\mu - \mu_0)x + t.o.s.$$

nesse caso, os *t.o.s.* são de ordens superiores ou iguais a x^4 ou $(\mu - \mu_0)x^2$, e onde $\alpha = \frac{\partial^2 f}{\partial x^2}(0, \mu_0)$, $\gamma = \frac{1}{3!} \cdot \frac{\partial f}{\partial \mu}(0, \mu_0)$ e $\delta = \frac{\partial^2 f}{\partial x \partial \mu}(0, \mu_0)$.

Queremos encontrar uma mudança de coordenadas que elimine o termos de x^2 e o de $(\mu - \mu_0)$. Para isso, vamos considerar a seguinte mudança de coordenadas

$$\bar{x} = h(x, \mu) = x + x^2 + b(\mu - \mu_0)$$

onde as constantes reais a e b ainda serão determinadas. Porém, precisamos determinar a expressão de h^{-1} . Ora, seja

$$h^{-1} = x + c^2 + d(\mu - \mu_0) + d(\mu - \mu_0) + ex(\mu - \mu_0) + t.o.s.$$

e encontremos as constantes c, d e e para as quais $h^{-1} \circ h = id$:

$$\begin{aligned} x &= h^{-1} \circ h(x) = h^{-1}(x + ax^2 + b(\mu - \mu_0)) \\ &= x + ax^2 + b(\mu - \mu_0) + c(x + ax^2 + b(\mu - \mu_0))^2 + d(\mu - \mu_0) + \\ &\quad + e(x + ax^2 + b(\mu - \mu_0))(\mu - \mu_0) + t.o.s. \\ &= x + (a + c)x^2 + (b + d)(\mu - \mu_0) + (2bc + e)(\mu - \mu_0)x + t.o.s.. \end{aligned}$$

Logo, basta exigir que

$$a + c = 0, \quad b + d = 0 \quad \text{e} \quad 2bc + e = 0,$$

ou seja, que $c = -a$, $d = -b$ e que $e = 2ab$. Assim, obtemos que

$$h^{-1}(x, \mu) = x - (ax^2 + b(\mu - \mu_0)) + 2ab(\mu - \mu_0)x + t.o.s..$$

Até agora não denotamos as variáveis de h^{-1} por $\bar{x}$ e $\bar{\mu}$ apenas para simplificar a notação. De agora em diante, consideremos

$$\begin{aligned}\bar{x} &= h(x, \mu) = x + ax^2 + b(\mu - \mu_0) \\ x &= h^{-1}(\bar{x}, \mu) = \bar{x} - a\bar{x}^2 - b(\mu - \mu_0) + 2ab(\mu - \mu_0)\bar{x} + t.o.s..\end{aligned}$$

Note que tomamos $\bar{\mu} = \mu$. Assim, com as novas coordenadas, temos que

$$\begin{aligned}\bar{f}_{\bar{\mu}}(\bar{x}) &= \bar{f}_{\bar{\mu}}(\bar{x}) = h \circ f_{\mu} \circ h^{-1}(\bar{x}, \bar{\mu}) \\ &= h \circ f_{\mu}(\bar{x} - a\bar{x}^2 + b(\mu - \mu_0) + 2ab(\mu - \mu_0)\bar{x} + t.o.s.) \\ &= h(-\bar{x} + a\bar{x}^2 + b(\mu - \mu_0) - 2ab(\mu - \mu_0)\bar{x} + \\ &\quad + \alpha(-\bar{x} + a\bar{x}^2 + b(\mu - \mu_0) - 2ab(\mu - \mu_0)\bar{x})^2 + \beta(\mu - \mu_0) + \\ &\quad + \gamma(-\bar{x} + a\bar{x}^2 + b(\mu - \mu_0) - 2ab(\mu - \mu_0)\bar{x})^3 + \\ &\quad + \delta(\mu - \mu_0)(-\bar{x} + a\bar{x}^2 + b(\mu - \mu_0) - 2ab(\mu - \mu_0)\bar{x}) + t.o.s.) \\ &= h(-\bar{x} + (a + \alpha)\bar{x}^2 + (b + \beta)(\mu - \mu_0) + (-2ab - \delta)(\mu - \mu_0)\bar{x} + t.o.s.) \\ &= -\bar{x} + (a + \alpha)\bar{x}^2 + (b + \beta)(\mu - \mu_0) + (-2ab - \delta)(\mu - \mu_0)\bar{x} + \\ &\quad + a(-\bar{x} + (a + \alpha)\bar{x}^2 + (b + \beta)(\mu - \mu_0) + (-2ab - \delta)(\mu - \mu_0)\bar{x})^2 + \\ &\quad + b(\mu - \mu_0) + t.o.s. \\ &= -\bar{x} + (2a + \alpha)\bar{x}^2 + (2b + \beta)(\mu - \mu_0) + (-4ab - \delta)(\mu - \mu_0)\bar{x} + t.o.s..\end{aligned}$$

Basta fazer $2a + \alpha = 0$ e $2b + \beta = 0$, ou seja, $a = \frac{-\alpha}{2}$ e $b = \frac{-\beta}{2}$, para obtermos que

$$\begin{aligned}\bar{x} &= h(x, \mu) = x - \frac{\alpha}{2} \cdot x^2 - \frac{\beta}{2} \cdot (\mu - \mu_0) \\ x &= h^{-1}(\bar{x}, \mu) = \bar{x} + \frac{\alpha}{2} \cdot \bar{x}^2 + \frac{\beta}{2} \cdot (\mu - \mu_0) + \frac{\alpha\beta}{2} \cdot (\mu - \mu_0) \cdot \bar{x} + t.o.s..\end{aligned}$$

E com essa mudança de coordenadas (lembremos que $\bar{\mu} = \mu$) a nova função $\bar{f}_{\mu}$ fica da seguinte forma

$$\bar{f}_{\mu}(\bar{x}) = -\bar{x} + \kappa\bar{x}^3 + \tau(\mu - \mu_0)\bar{x} + O(\bar{x}^4) + O(|\mu - \mu_0|\bar{x}^2),$$

onde κ é obtido, como acima, em função das constantes α, β, γ e δ , $\tau(0) = 0$ e $O(\iota)$ significa termos de ordens superiores ou iguais a ι , para $\iota = \bar{x}^4$ e $|\mu - \mu_0|\bar{x}^2$. Como

queríamos.

Em (4.2) tomamos $a \neq 0$ e a origem é uma sela-nó. Além disso, se tomarmos $b'(0) \neq 0$, dizemos que a sela-nó *desdobra genericamente*. Estas condições são claramente satisfeitas genericamente. Como vimos anteriormente na Seção 4.2.1, f_μ e, portanto, φ_μ têm dois pontos fixos hiperbólicos para $\mu < \mu_0$ e nenhum para $\mu > \mu_0$ ou vice-versa. Se considerarmos $a > 0$, $b > 0$ e $|\rho_2| < 1$ temos o seguinte desdobramento da sela-nó: um poço e uma sela colapsam e então desaparecem, como na Figura 4.8.

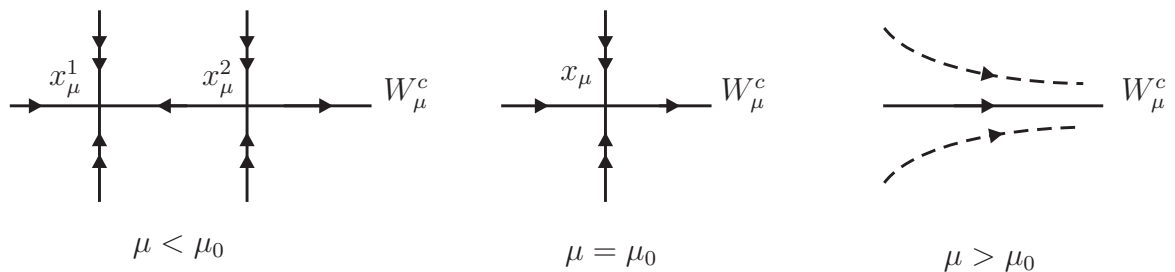


Figura 4.8: Desdobramento de uma sela-nó.

Na Figura 4.8, a flecha dupla significa que a contração na direção vertical é mais forte do que ao longo de W^c_μ . Se considerarmos as curvas $\mu \mapsto x^1_\mu$ e $\mu \mapsto x^2_\mu$ dos pontos fixos, para $\mu \leq \mu_0$, obtemos a Figura 4.9.

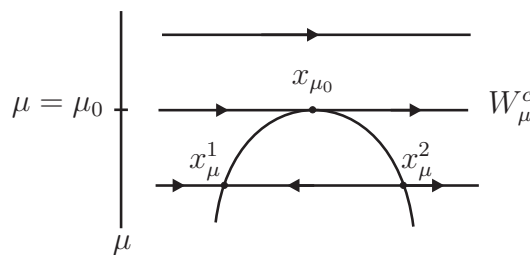


Figura 4.9: Curva dos pontos fixos de uma sela-nó.

Note que as duas curvas são diferenciáveis para μ próximo de μ_0 . Podemos orientar a curva da Figura 4.9 seguindo $\mu \mapsto x^1_\mu$ por valores crescentes de μ até $\mu = \mu_0$ e então retornar ao longo de $\mu \mapsto x^2_\mu$ por valores decrescentes de μ .

Vamos olhar para a expressão (4.3), que corresponde ao autovalor $\rho_1 = -1$. Tomemos $\kappa \neq 0$, que é uma condição genérica, e obtemos uma bifurcação do tipo flip. Dizemos que ela *desdobra genericamente* quando $\tau'(0) \neq 0$, o que é outra condição genérica. Quando tomamos $\kappa > 0$ e $\tau'(0) < 0$, obtemos pelo Teorema 9 que existe um único ponto fixo que é um poço para $\mu < \mu_0$ e uma sela para $\mu > \mu_0$ e, além disso, existe um poço de período dois (com autovalores positivos) para $\mu > \mu_0$.

Os resultados são semelhantes para bifurcações flip de órbitas periódicas de período k , basta considerar a família f_μ^k .

Para o flip considerado acima, se representarmos a curva dos poços e selas de período k e a curva dos poços de período $2k$ em um mesmo ambiente, obtemos a Figura 4.10.

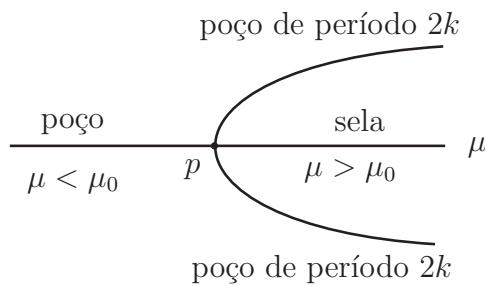


Figura 4.10: Curva dos pontos fixos e de período dois de um flip.

O poço à esquerda e a sela à direita têm o mesmo período e um correspondente autovalor negativo, -1 , para a derivada Df_μ^k .

4.4 Cascatas de Bifurcações de Duplicação de Período e Poços

Nesta seção nosso interesse é ver um resultado que mostra que durante a criação de uma ferradura existem infinitos poços ou fontes e bifurcações de duplicação de período. Para isso, antes introduziremos algumas definições e hipóteses que serão utilizadas neste resultado. Também veremos um exemplo que satisfaz o conjunto de condições que vamos considerar abaixo.

Condições Iniciais para a Criação de uma Ferradura

Seja $R \subset \mathbb{R}^2$ um retângulo e considere $\{\varphi_\mu\}$ uma família a um parâmetro de difeomorfismos de R em $\mathbb{R}^2$ e com $\mu \in [-1, 1]$ tais que

- (a) $\varphi_{-1}(R) \cap R = \emptyset$;
- (b) $\varphi_\mu|_R$ é dissipativa para $-1 \leq \mu \leq 1$, isto é, $|\det(D\varphi_\mu)| < 1$ sobre R ;
- (c) φ_1 possui pontos periódicos e são todos selas;
- (d) $\varphi_\mu(R) \cap S_1 = \emptyset$ e $\varphi_\mu(R) \cap S_2 = \emptyset$, para $-1 \leq \mu \leq 1$, onde S_1 e S_2 são os lados verticais no bordo do retângulo R ;
- (e) $\varphi_\mu(T) \cap R = \emptyset$ e $\varphi_\mu(B) \cap R = \emptyset$, para $-1 \leq \mu \leq 1$, onde T é o topo e B a base no bordo do retângulo R e
- (f) φ_μ tem no máximo uma órbita periódica não-hiperbólica para cada $-1 \leq \mu \leq 1$ e esta órbita deve corresponder ou a uma sela-nó ou a uma bifurcação de duplicação de período que desdobra genericamente. Esta é uma condição genérica sobre a família $\{\varphi_\mu\}$.

Observe que no item (f) não consideramos o caso da bifurcação de Andronov-Hopf pois, como já mencionamos, este caso não é dissipativo.

Embora não exigimos formalmente que φ_1 seja uma ferradura como no Exemplo 3 devemos ter exatamente esta situação em mente, veja a Figura 4.11. Neste caso costuma-se dizer que essa família é uma *família de área decrescente que cria uma ferradura*. Vamos agora ver um exemplo que satisfaz essas condições.

Exemplo 7. Seja $\{\varphi_\mu\}$ uma família a um parâmetro de difeomorfismos. Suponha que φ seja tal que

- (i) φ_0 possui um ponto de sela fixo p e que $|\det(D\varphi_0)_p| < 1$;
- (ii) existe uma tangência homoclínica q associada a p que é desdobrada genericamente.

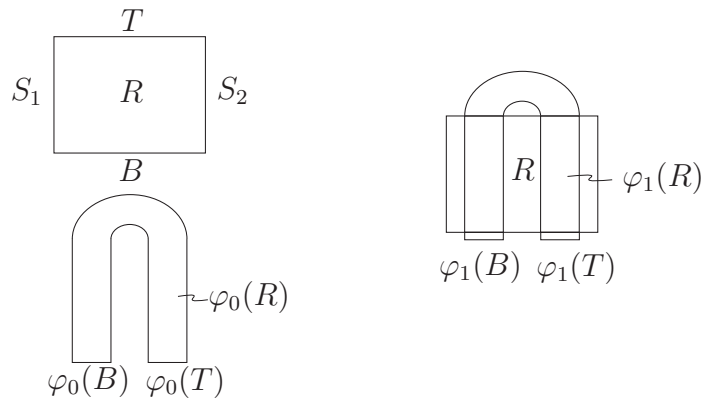


Figura 4.11: Situação descrita nas Condições Iniciais para a Criação de uma Ferradura.

Afirmção. Para cada vizinhança V de q existe um retângulo $R \subset V$, um número $\delta > 0$ e um inteiro $N > 0$ tal que $\varphi_\mu^N|_R$ cria uma ferradura para $-\delta < \mu < \delta$.

De fato: considere o retângulo R estreito, próximo de q e paralelo à componente local de W_μ^s como indicado na Figura 4.12.

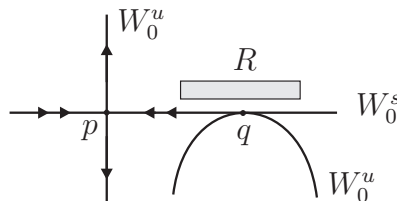


Figura 4.12: Posição do retângulo R paralelo à W_μ^s .

Na verdade, para μ pequeno, escolhemos coordenadas linearizantes de classe C^1 para cada φ_μ em uma vizinhança de p fixada que contenham um arco $\ell^s \subset W_0^s$ que liga os pontos p e q e que dependam continuamente de μ . Então, escolhemos R estreito suficientemente próximo de W_0^s de modo que sua projeção sobre W_0^u , paralela à W_0^s , contenha em seu interior o ponto $\varphi_0^{-N}(q)$, para algum N grande. Assim, $\varphi_0^N(R)$ é uma caixa curvada próxima de um arco W_0^u localizado perto de q . Logo, quando μ está bem perto de 0, temos a situação como na Figura 4.13.

Agora basta utilizarmos argumentos semelhantes aos da Seção 3 para mostrar que $\varphi_\mu^N|_R$ possui um conjunto invariante maximal hiperbólico com o subconjunto das

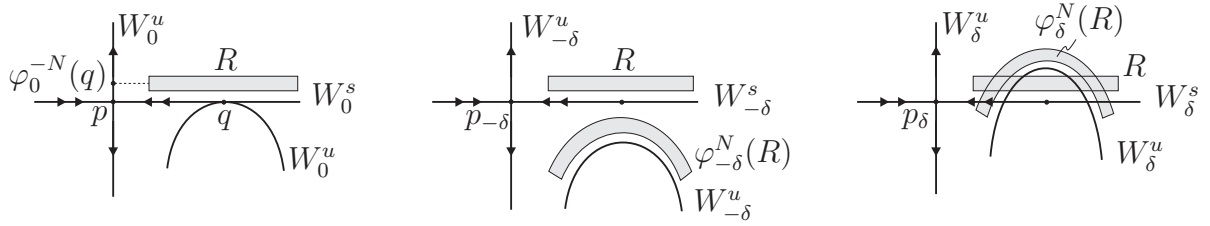


Figura 4.13: Criação de uma ferradura.

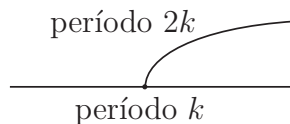
órbitas periódicas denso.

Observação 14. Note que a configuração de R e $\varphi_{\delta}^N(R)$ assemelha-se com a situação descrita na Seção 3, porém os retângulos em questão são bastante distintos, pois o da Seção 3 contém os pontos p e q , enquanto que o considerado no exemplo acima está contido em uma pequena vizinhança de q .

Continuando nossa discussão, vamos agora considerar um espaço mais simples para as órbitas de nossa família de aplicações $\varphi : R \subset \mathbb{R}^2 \longrightarrow \mathbb{R}^2$ que satisfaz as Condições Iniciais para a Criação de um Ferradura. Seja $Per(\varphi_{\mu})$ o conjunto das órbitas periódicas de φ_{μ} e considere

$$P = \{(x, \mu) \in R \times [-1, 1] \mid x \in Per(\varphi_{\mu})\}.$$

Definimos o espaço topológico $\tilde{P} = P / \sim$, onde a relação de equivalência “ $\sim$ ” identifica os pontos de uma mesma órbita. Desse modo, dado um ponto $(\mathcal{O}(x), \mu) \in \tilde{P}$, a componente conexa que passa por esse ponto é uma curva contínua exceto na bifurcação de duplicação de período (ou no colapso das curvas dos pontos de período k com a dos pontos de período $2k$) onde ela se ramifica em duas como na Figura 4.14.

Figura 4.14: Curva de pontos periódicos no espaço $\tilde{P}$.

Agora estamos em condições de enunciar um dos resultados principais desta seção.

Teorema 11 ([9], 1983). *Seja $\varphi_\mu : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ uma família de difeomorfismos que preserva orientação e satisfaz as Condições Iniciais para a Criação de uma Ferradura. Então, cada ponto $(\mathcal{O}(x), 1) \in \tilde{P}$ possui uma componente conexa que contém órbitas periódicas de período $2^n k$ para todo $n \geq 0$, onde k é o período do ponto x para a aplicação φ_1 .*

Demonstração:

Seja $(\mathcal{O}(x), 1) \in \tilde{P}$ e suponha que x tem período k . Temos duas possibilidades: ou $D(\varphi_1^k)_x$ tem autovalores positivos ou tem autovalores negativos. Começaremos com o caso dos autovalores positivos. Como x é ponto de sela, pela condição **(b)**, segue do Teorema 7 que existe um único caminho Γ em $\tilde{P}$ que passa pelo ponto $(\mathcal{O}(x), 1)$, o qual orientaremos seguindo valores decrescentes de μ .

Afirmção 1. Se prolongarmos Γ chegaremos em uma bifurcação.

De fato: Não podemos prolongar Γ até $R \times \{-1\}$ pois, pela condição **(a)**, $\varphi_{-1}(R) \cap R$ não possui pontos periódicos, logo não pode conter pontos de Γ .

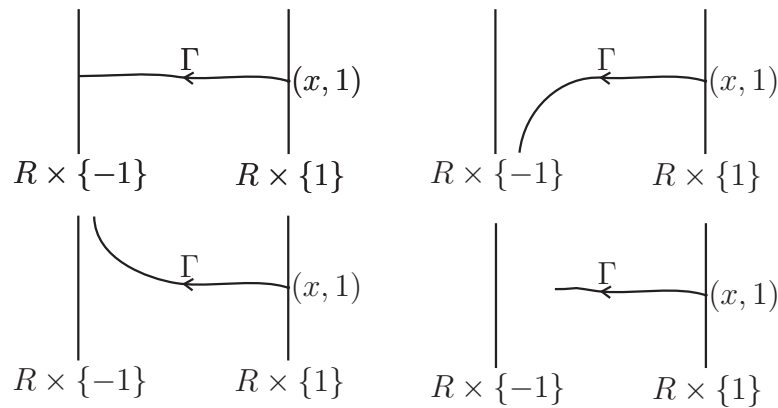


Figura 4.15: As situações que não ocorrem com Γ .

Pelas condições **(d)** e **(e)** segue que o conjunto invariante maximal de φ_μ em R é limitado pela fronteira de R , ∂R , o que impede os pontos periódicos de escapar por

∂R . Além disso, Γ não pode terminar em $R \times (-1, 1)$ pois Γ é um caminho de selas (é a continuação do ponto de sela $(\mathcal{O}(x), 1)$) e porque sempre podemos prolongar um caminho de selas. Veja a Figura 4.15.

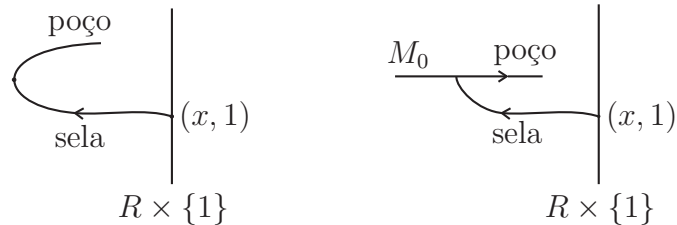


Figura 4.16: À esquerda temos uma sela-nó e à direita um flip.

Logo, a única situação que resta é Γ passar por um ponto de bifurcação que deve corresponder ou a uma sela-nó ou a um flip (no caso do flip o período é duplicado). Em ambos os casos, surge da órbita de bifurcação um caminho de poços mudando a direção de Γ em relação a μ , ou seja, orientada para valores crescentes de μ . Veja a Figura 4.16.

De agora em diante vamos sempre prolongar Γ para um flip evitando caminhos de selas com autovalores negativos para $D(\varphi_\mu^\ell)$, onde ℓ é o período de x para φ_μ . Tais caminhos chamaremos de *caminhos de Möbius* e denotaremos por M_0 . Convencionamos a seguinte orientação: se for um caminho de poços, orientamos positivamente, isto é, para valores crescentes de μ , e se for um caminho de selas, orientamos negativamente, isto é, para valores decrescentes de μ . Veja a Figura 4.17.

O primeiro caso, por exemplo, representa uma bifurcação onde x_μ é um poço que ao passar pela órbita de bifurcação torna-se um 2-poço (isto é, um poço com o dobro do período de x_μ que tem a órbita denotada por x_μ^1 e x_μ^2) e uma sela $\bar{x}_\mu$, como na Figura 4.18. Nesse caso evitamos o caminho de selas por que a derivada $D(\varphi_\mu^k)_{\bar{x}_\mu}$ tem autovalores negativos, já que temos -1 como um correspondente autovalor para a bifurcação flip, o que implica que a sela à direita e o poço à esquerda (que possuem o mesmo período) têm um correspondente autovalor negativo.

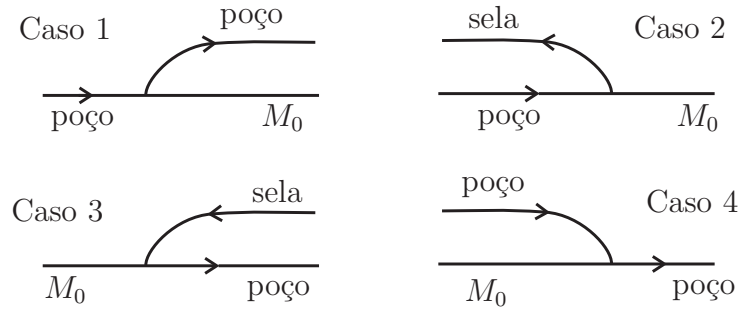
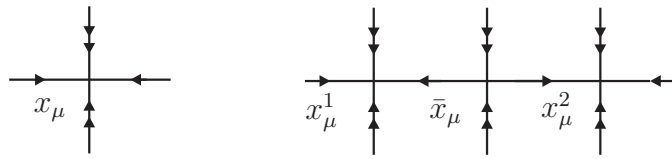
Figura 4.17: Orientação de Γ .

Figura 4.18: Exemplo de bifurcação.

Note que no caso considerado não poderíamos duplicar o período de nosso caminho de selas para um de poços pois estaríamos começando com um caminho de Möbius M_0 . Logo, a bifurcação seria colapsar uma sela em um poço com a metade do período, como por exemplo, na Figura 4.19 onde uma 2-sela x_μ^1 , x_μ^2 e um poço $\bar{x}_\mu$ se colapsam e se tornam uma sela x_μ .

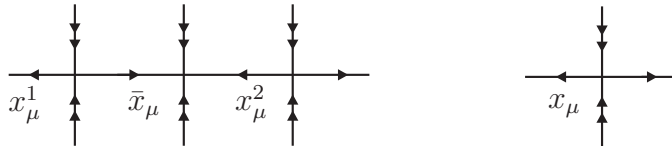


Figura 4.19: Representação do Caso 3.

Nas selas-nó os caminhos estão orientados seguindo a mesma convenção tomada acima. Continuemos a prolongar nossa curva Γ . Não podemos voltar em $R \times \{1\}$ pois φ_1 não possui poços, pela condição (c). Além disso, não podemos ter um ciclo, isto é, Γ não pode retornar em si mesma, pois em cada “ponto de bifurcação de retorno” existe um caminho de poços e um de selas e isso contraria a nossa convenção porque estaríamos admitindo um caminho de Möbius (deve-se analisar a nossa situação inicial, Figura 4.16, juntamente com todos os casos da Figura 4.17 para se

convencer melhor).

Afirmção 2. O caminho Γ não pode terminar em $R \times (-1, 1)$ se ele passar por um número finito de órbitas de bifurcação ou se ele passar por infinitas órbitas de bifurcação, porém com períodos limitados.

De fato: No primeiro caso, isto é, se fosse um número finito de bifurcações, o caminho terminaria com um poço ou com uma sela. Não pode ser poço porque senão, prolongando o caminho, chegaria a $R \times \{1\}$ com um poço. Não pode ser sela porque, prolongando o caminho, chegaria a $R \times \{-1\}$ com um ponto periódico. O caminho não pode sair pela lateral. Logo passamos por infinitas bifurcações. No segundo caso, como as órbitas periódicas são infinitas com períodos limitados, existe algum inteiro s que é o período de infinitas dessas órbitas, digamos (x_i, μ_i) . Logo, como estamos num conjunto compacto, a sequência (x_i, μ_i) admite um ponto limite que denotaremos por $(\tilde{x}, \tilde{\mu})$. Como as aplicações φ_μ são contínuas para todo $\mu \in [-1, 1]$, segue que

$$\varphi_{\tilde{\mu}}^s(\tilde{x}) = \varphi_{\tilde{\mu}}^s(\lim_{i \rightarrow \infty} x_i) = \lim_{i \rightarrow \infty} \varphi_{\tilde{\mu}}^s(x_i) = \lim_{i \rightarrow \infty} x_i = \tilde{x},$$

ou seja, $\tilde{x}$ é periódico de período s . Mas, pela condição **(f)**, $\tilde{x}$ deve ser sela-nó ou flip. Em ambos os casos, essas bifurcações atraem ou repelem pontos próximos, isto é, não podem existir pontos periódicos todos com períodos limitados numa vizinhança. O que justifica nossa afirmação.

Da Afirmção 2, segue o resultado, pois como temos infinitas bifurcações com períodos ilimitados, temos bifurcações com período $2^n k$, para todo $n \geq 0$.

Consideremos agora o segundo caso, onde os autovalores de $D(\varphi_1^k)_x$ são negativos e k é o período de x . Nesse caso, o caminho que passa por x é um caminho de Möbius. Para esse caso utilizaremos a seguinte orientação: se o caminho for de poços, orientamos para valores crescentes de μ , se for de selas (com autovalores positivos)

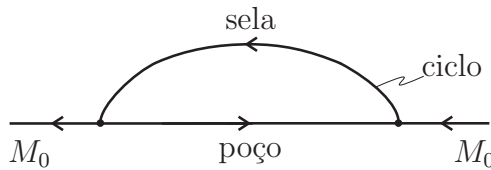


Figura 4.20: Um ciclo.

ou de Möbius, orientamos para valores decrescentes de μ . Seja Γ o caminho de Möbius iniciando no ponto $(\mathcal{O}(x), 1)$. Com o mesmo argumento do caso anterior, concluímos que Γ passa por órbitas de bifurcação e a primeira deve ser um flip (por ser de Möbius e ter um autovalor correspondente negativo). Então, seguimos com Γ pelo caminho de selas com o dobro do período que emana do caminho de selas no qual estávamos. Na próxima órbita de bifurcação repetimos o procedimento e prolongamos Γ ao longo do único caminho de Möbius que surge. E neste ponto já podemos ter um ciclo, isto é, um caminho de órbitas periódicas fechado e orientado que não contém curva de Möbius alguma. Veja a Figura 4.20 para visualizar a situação.

Se um desses ciclos não aparecer em Γ , então está pronto. Caso apareçam, resolvemos o problema do seguinte modo. Note que em cada ciclo, o número de períodos que duplicam é igual ao número de períodos que se colapsam. Em cada período de duplicação existe um nova ramificação de um caminho de Möbius e em cada período colapsante há uma perda de um caminho de Möbius. Resolvemos este problema considerando um novo espaço $\hat{P}$: identificamos em $\tilde{P}$, o espaço das órbitas periódicas em $R \times [-1, 1]$, cada ciclo com um ponto. Esses pontos têm agora o mesmo número (finito) de caminhos de Möbius criados e de caminhos de Möbius eliminados. Agora, no espaço $\hat{P}$, consideramos o caminho $\hat{\Gamma}$ que começa em $(\mathcal{O}(x), 1)$ e que é constituído completamente de caminhos de Möbius tais que

- a orientação da curva Γ coincide com a orientação dos caminhos de Möbius que ela segue e
- $\hat{\Gamma}$ passa por cada caminho de Möbius no máximo uma vez.

Esse tipo de caminho sempre pode ser prolongado toda vez que ele chegar em um

ciclo colapsado, pois o número de entradas de caminhos de Möbius é igual ao de saída de caminhos de Möbius. $\hat{\Gamma}$ não pode terminar em uma órbita periódica de $R \times \{1\}$ porque todos os caminhos de Möbius estão orientados no sentido decrescente dos valores de μ . Logo, existem duas possibilidades:

- ou $\hat{\Gamma}$ é finita, mas nesse caso ela termina em um ponto de duplicação de período que não pertence a um ciclo e então basta proceder como antes para concluir o resultado,
- ou $\hat{\Gamma}$ é infinita. Neste caso o caminho $\hat{\Gamma}$ induz um gráfico conexo Γ no espaço $\tilde{P}$, constituído dos “pedaços de arcos de $\hat{\Gamma}$ ” e dos ciclos por onde $\hat{\Gamma}$ passou (que em $\hat{P}$ eram pontos). Claramente Γ é infinita e, portanto, contém um número infinito de órbitas periódicas de bifurcação. Agora basta usar o argumento anterior para concluir a demonstração. ■

Observação 15. Para famílias $\{\varphi_\mu\}$ com orientação reversível basta considerar a família dos quadrados $\{(\varphi_\mu)^2\}$. Se esta nova família satisfaz as Condições Iniciais para a Criação de um Ferradura, então obtemos o mesmo resultado.

4.5 Renormalização de Tangências Homoclínicas

Nesta seção vamos ver como podemos comparar uma ferradura criada por uma tangência homoclínica com a família quadrática. A principal ferramenta utilizada será a renormalização.

Consideremos M uma variedade de dimensão 2 e $\varphi_\mu : M \longrightarrow M$ uma família de difeomorfismos a um parâmetro. Suponha que a família φ_μ possui uma tangência homoclínica em $\mu = 0$, relativa ao ponto de sela p_0 para φ_0 , e que essa tangência é genérica, isto é, tem contato parabólico, e desdobra genericamente.

Inicialmente daremos uma idéia sem muitos detalhes e justificativas da construção que faremos adiante na demonstração do resultado principal e algumas de suas conseqüências. Numa vizinhança do ponto de sela p_0 mencionado acima existem

coordenadas linearizantes x e y tais que $\varphi_0(x, y) = (\lambda \cdot x, \sigma \cdot y)$, onde $0 < |\lambda| < 1 < |\sigma|$. Assumiremos que λ e σ são positivos e que $\lambda \cdot \sigma < 1$ pois, caso eles fossem negativos, bastaria tomar φ_μ^2 no lugar de φ_μ e, caso $\lambda \cdot \sigma > 1$, bastaria tomar φ_μ^{-1} no lugar de φ_μ (se $\lambda \cdot \sigma = 1$, nossa construção não funciona). Considere como na Figura 4.21,

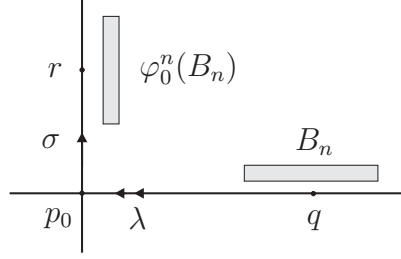


Figura 4.21: Posições dos pontos p e r e de B_n .

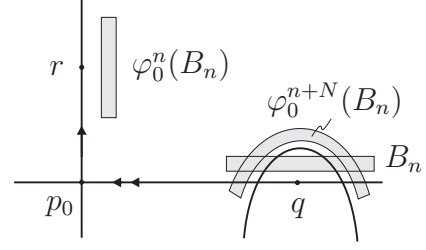


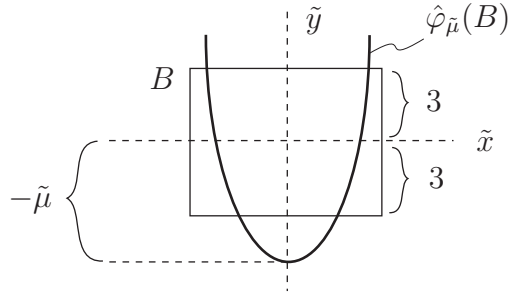
Figura 4.22: Posição de B_n e de suas imagens por φ_μ .

dois pontos p e r da órbita de tangência contidos no domínio das coordenadas linearizantes. Desse modo, existe um $N \in \mathbb{N}$ para o qual $r = \varphi_0^N(q)$. Considere, para cada n suficientemente grande, uma caixa B_n próxima de q de modo que $\varphi_\mu^n(B_n)$ fique próxima de r , como na Figura 4.21. Olhemos para a posição de $\varphi_\mu^{n+N}(B_n)$ em relação à posição de B_n . Podemos tomar B_n cuidadosamente de modo que, para n suficientemente grande, a caixa curvada $\varphi_\mu^{n+N}(B_n)$ atravesse B_n e crie uma ferradura, do mesmo modo que foi feito no Exemplo 7. Veja a Figura 4.22. Agora, aplicamos mudanças de coordenadas, que dependem de n , nas variáveis x e y e no parâmetro μ e denotamos as novas coordenadas por $\tilde{x}$, $\tilde{y}$ e $\tilde{\mu}$. Com essas novas coordenadas a aplicação φ_μ^{n+N} converge, quando $n \rightarrow \infty$, para a aplicação $\hat{\varphi}_{\tilde{\mu}}$ que tem a seguinte expressão $\hat{\varphi}_{\tilde{\mu}}(\tilde{x}, \tilde{y}) = (\tilde{y}, \tilde{y}^2 + \tilde{\mu})$. Para n suficientemente grande, tomemos a caixa B_n nas novas coordenadas n -dependentes igual a seguinte caixa

$$B = \{(\tilde{x}, \tilde{y}) \mid |\tilde{x}| \leq 3, |\tilde{y}| \leq 3\}.$$

Quando fazemos o parâmetro $\tilde{\mu}$ variar de -4 à 4 temos a formação de uma ferradura. Veja a Figura 4.23.

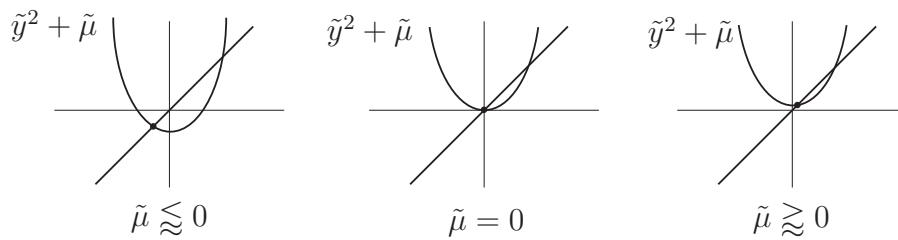
Como φ_μ contrai área em p_μ e, portanto, φ_μ^{n+N} tem essa propriedade, quando

Figura 4.23: Imagem de B pela aplicação $\hat{\varphi}_{\tilde{\mu}}$.

$n \rightarrow \infty$, a caixa curvada $\varphi_{\mu}^{n+N}(B_n)$ tende exatamente para uma curva, logo, a aplicação limite não é mais um difeomorfismo. Esta aplicação limite, restrita à variável $\tilde{y}$, tem a seguinte expressão

$$\tilde{y} \mapsto \tilde{y}^2 + \tilde{\mu}, \quad (4.4)$$

enquanto que a variável $\tilde{x}$ não tem importância para ela. A expressão (4.4) é a família quadrática à um parâmetro em dimensão 1 que pode ser encontrada em [6]. Logo, como a aplicação limite se aproxima da família quadrática, ela apresenta os principais complexidades dinâmicas dessa família como por exemplo conjuntos hiperbólicos e bifurcações de duplicação de período. Vamos ver uma dessas propriedades: para $\tilde{\mu}$ próximo de zero, a aplicação $\tilde{y} \mapsto \tilde{y}^2 + \tilde{\mu}$ possui um ponto fixo próximo de zero que é atrator, veja a Figura 4.24.

Figura 4.24: Pontos fixos da aplicação $\tilde{y} \mapsto \tilde{y}^2 + \tilde{\mu}$ para diferentes valores de $\tilde{\mu}$.

Temos que

$$\tilde{y}^2 + \tilde{\mu} = \tilde{y} \implies \tilde{y}^2 - \tilde{y} + \tilde{\mu} = 0 \implies \tilde{y} = \frac{1 \pm \sqrt{1 - 4\tilde{\mu}}}{2}$$

para $\tilde{\mu} < \frac{1}{4}$. Logo, o ponto fixo mais próximo de zero é $\frac{1 - \sqrt{1 - 4\tilde{\mu}}}{2}$. Além disso, a derivada do n -ésimo iterado da aplicação $\tilde{y} \mapsto \tilde{y}^2 + \tilde{\mu}$ é $2^n \tilde{y}$ tem módulo menor que 1 se $|\tilde{y}| < \frac{1}{2^n}$. Seja $\mu_n \rightarrow 0$ a seqüência dos valores μ correspondentes à $\tilde{\mu} = 0$ nas diferentes reparametrizações da variável μ . Logo, para n suficientemente grande e μ próximo de μ_n , a aplicação φ_μ^{n+N} tem um ponto fixo atrator. De fato, seja y_0 o ponto fixo da aplicação $\tilde{y} \mapsto \tilde{y}^2 + \tilde{\mu}$, que denotamos por $\hat{\varphi}_{\tilde{\mu}}$, que satisfaz a condição $y_0 < \frac{1}{2}$, ou seja,

$$y_0 = y_0^2 + \mu \quad \text{e} \quad 2y_0 < 1.$$

Logo, o ponto (y_0, y_0) é um ponto fixo da aplicação $(\tilde{x}, \tilde{y}) \mapsto (\tilde{y}, \tilde{y}^2 + \tilde{\mu})$. Além disso, esse ponto é hiperbólico pois

$$D\hat{\varphi}_{\tilde{\mu}}(y_0, y_0) = \begin{bmatrix} 0 & 1 \\ 0 & 2y_0 \end{bmatrix}$$

possui os autovalores $\lambda_1 = 0 < 1$ e $\lambda_2 = 2y_0 < 1$. Então, pelo Teorema 7, segue que o ponto fixo (y_0, y_0) tem uma continuação. Portanto, para n bem grande, e μ próximo de $\tilde{\mu}$ a aplicação φ_μ^{n+N} possui um ponto fixo que é atrator.

Agora que temos uma idéia heurística do resultado, vamos apresentá-lo formalmente. Mas antes vamos considerar algumas hipóteses a respeito da família φ_μ .

Hipóteses para a família φ_μ . (*)

Já assumimos que os autovalores λ e σ são positivos e satisfazem $\lambda \cdot \sigma < 1$ e que precisamos de coordenadas linearizantes x e y de classe C^2 para φ_μ em uma vizinhança de p_μ . Além dessas condições exigiremos também que $\varphi_\mu(x, y)$ seja de classe C^∞ como função de (μ, x, y) . A existência dessas coordenadas linearizantes de classe C^2 e que dependem de μ segue do fato que os autovalores λ e σ satisfazem as seguintes condições genéricas:

(i) $\lambda \neq \lambda^m \cdot \sigma^n$, para todo $m, n \in \mathbb{N} \setminus \{0\}$ tais que $n + m > 1$ e

(ii) $\sigma \neq \lambda^m \cdot \sigma^n$, para todo $m, n \in \mathbb{N} \setminus \{0\}$ tais que $n + m > 1$.

Além disso, elas dependem continuamente de μ na topologia C^2 .

De fato, para justificar a existência, em [8], temos o seguinte teorema.

Teorema 12. *Seja φ um difeomorfismo de classe C^ℓ de $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ definido em alguma vizinhança V da origem $(0,0)$, a qual supomos ser um ponto fixo. Sejam λ e σ os autovalores de $D\varphi_{(0,0)}$. Então, se σ e λ satisfazem (i) e (ii), existe uma vizinhança da origem e uma função $\xi = \xi(\lambda, \sigma; \ell)$ tal que existe uma mudança de coordenadas H de classe C^ξ definida em V que lineariza φ . Além disso, para todo λ e σ fixados, segue que $\xi(\lambda, \sigma; \ell) \rightarrow \infty$, quando $\ell \rightarrow \infty$. Em particular, para $\ell = \infty$, H pode ser tomada de classe C^∞ .*

As condições (i) e (ii) são chamadas *condições de não-ressonância*. Elas são genéricas porque basta impor um número finito de condições sobre os autovalores λ e σ para que sejam satisfeitas. A prova deste fato também pode ser encontrada em [8], Seção 7, nas páginas 629 e 630.

Teorema 13. *Para uma família a um parâmetro φ_μ satisfazendo as hipóteses (*), com q um ponto da órbita de tangência para o valor do parâmetro $\mu = 0$, existe uma constante inteira $N > 0$ e, para cada $n \in \mathbb{N}$, reparametrizações $\mu = M_n(\tilde{\mu})$ da variável μ e mudanças de coordenadas $\tilde{\mu}$ -dependentes*

$$(\tilde{x}, \tilde{y}) \longmapsto (x, y) = \Psi_{n, \tilde{\mu}}(\tilde{x}, \tilde{y})$$

tais que,

(a) *para cada conjunto compacto K no espaço das variáveis $(\tilde{\mu}, \tilde{x}, \tilde{y})$, as imagens de K pelas aplicações*

$$(\tilde{\mu}, \tilde{x}, \tilde{y}) \longmapsto (M_n(\tilde{\mu}), \Psi_{n, \tilde{\mu}}(\tilde{x}, \tilde{y}))$$

convergem, quando $n \rightarrow \infty$, para $(0, q)$ no espaço das variáveis $(\tilde{\mu}, \tilde{x}, \tilde{y})$;

(b) *os domínios das aplicações*

$$(\tilde{\mu}, \tilde{x}, \tilde{y}) \longmapsto (M_n(\tilde{\mu}), (\Psi_{n,\tilde{\mu}}^{-1} \circ \varphi_{M_n(\tilde{\mu})}^{n+N} \circ \Psi_{n,\tilde{\mu}})) \quad (4.5)$$

convergem, quando $n \rightarrow \infty$, para todo o $\mathbb{R}^3$ e as aplicações (4.5) convergem, na topologia C^2 , para a aplicação

$$(\tilde{\mu}, \tilde{x}, \tilde{y}) \longmapsto (\tilde{\mu}, \hat{\varphi}_{\tilde{\mu}}(\tilde{x}, \tilde{y}))$$

onde $\hat{\varphi}_{\tilde{\mu}}(\tilde{x}, \tilde{y}) = (\tilde{y}, \tilde{y}^2 + \tilde{\mu})$.

Demonstração:

Começamos, como de praxe, escolhendo coordenadas linearizantes μ -dependentes de classe C^2 (x, y) numa vizinhança de p_μ . Assim, para $\mu = 0$, $\varphi_0(x, y) = (\lambda \cdot x, \sigma \cdot y)$ com $0 < \lambda < 1 < \sigma$ e $\lambda \cdot \sigma < 1$. Considere q um ponto na órbita de tangência da variedade estável de p e r um ponto na variedade instável de p , ambos nessa vizinhança de p_μ onde as coordenadas linearizantes estão definidas. Para simplificar os argumentos, multipliquemos x e y por constantes de modo que $q = (1, 0)$ e $r = (0, 1)$, como na Figura 4.25.

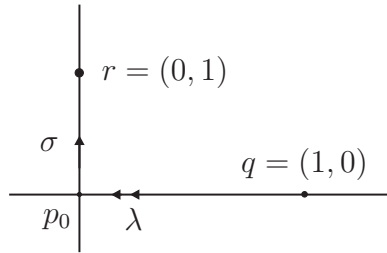


Figura 4.25: Posições de p e q nas variedades estável e instável locais.

Assim, como r e q pertencem à órbita de tangência, existe algum $N \in \mathbb{N}$ para o qual $q = \varphi_0^N(r)$. Vamos novamente fazer adaptações em nossas coordenadas. Para valores de μ próximo de zero adaptamos nossas coordenadas linearizantes tal que

(i) $\varphi_\mu(0, 1)$ é um máximo local da coordenada y quando restrita à $W^u(p_\mu)$;

(ii) a coordenada x de $\varphi_\mu^N(0, 1)$ é 1 e

(iii) reparametrizamos μ de tal modo que a coordenada y de $\varphi_\mu^N(0, 1)$ é μ .

Podemos fazer estas adaptações pois, quando $\mu = 0$, (i), (ii) e (iii) são satisfeitos e, para uma pequena variação de μ , as posições de $\varphi_\mu^N(0, 1)$ mudam ligeiramente. Veja a Figura 4.26.

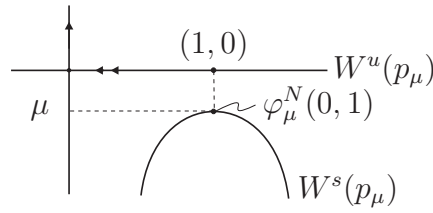


Figura 4.26: Adaptação das coordenadas x, y e do parâmetro μ .

Depois dessas adaptações faremos a seguinte afirmação.

Afirmção 1. Podemos escrever φ_μ^N , próximo de $(0, 1)$, da seguinte forma

$$(x, 1 + y) \longmapsto (1, 0) + (H_1(\mu, x, y), H_2(\mu, x, y))$$

com

$$\begin{aligned} H_1(\mu, x, y) &= \alpha \cdot y + \tilde{H}_1(\mu, x, y) \\ H_2(\mu, x, y) &= \beta \cdot y^2 + \mu + \gamma \cdot x + \tilde{H}_2(\mu, x, y), \end{aligned}$$

onde α, β e γ são constantes não nulas e onde, para $\mu = x = y = 0$, temos

$$\begin{cases} \tilde{H}_1 = \partial_y \tilde{H}_1 = \partial_\mu \tilde{H}_1 = 0 & \text{e} \\ \tilde{H}_2 = \partial_x \tilde{H}_2 = \partial_y \tilde{H}_2 = \partial_\mu \tilde{H}_2 = \partial_{yy} \tilde{H}_2 = \partial_{y\mu} \tilde{H}_2 = \partial_{\mu\mu} \tilde{H}_2 = 0. \end{cases} \quad (4.6)$$

De fato: consideremos a expansão em série de Taylor da aplicação φ_μ^N em torno do

ponto $(x, y) = (0, 1)$

$$\begin{aligned}\varphi_\mu^N(x, y+1) &= \varphi_\mu^N((0, 1) + (x, y)) \\ &= \varphi_\mu^N(0, 1) + D(\varphi_\mu^N)_{(0,1)} \cdot (x, y) + \frac{1}{2!} \cdot D^2(\varphi_\mu^N)_{(0,1)} \cdot (x, y)^2 + t.o.s..\end{aligned}$$

Consideremos $\varphi_\mu^N = (\varphi_{\mu,1}^N, \varphi_{\mu,2}^N)$, então temos que

$$\begin{aligned}D(\varphi_\mu^N)_{(0,1)} \cdot (x, y) &= \begin{bmatrix} \partial_x \varphi_{\mu,1}^N(0, 1) & \partial_y \varphi_{\mu,1}^N(0, 1) \\ \partial_x \varphi_{\mu,2}^N(0, 1) & \partial_y \varphi_{\mu,2}^N(0, 1) \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} \\ &= (\partial_x \varphi_{\mu,1}^N(0, 1) \cdot x + \partial_y \varphi_{\mu,1}^N(0, 1) \cdot y, \partial_x \varphi_{\mu,2}^N(0, 1) \cdot x + \partial_y \varphi_{\mu,2}^N(0, 1) \cdot y)\end{aligned}$$

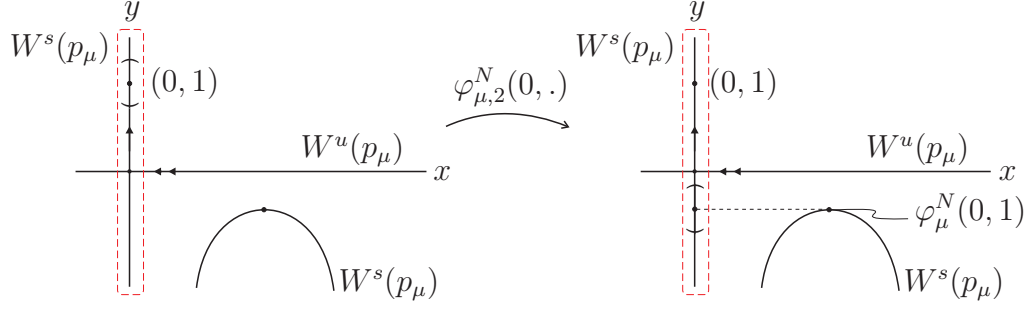
e, além disso, temos que $D^2(\varphi_\mu^N)_{(0,1)} \cdot (x, y)^2 = D^2(\varphi_\mu^N)_{(0,1)} \cdot ((x, y), (x, y))$ é dada por

$$\begin{bmatrix} \partial_{xx} \varphi_{\mu,1}^N(0, 1) \cdot x + \partial_{xy} \varphi_{\mu,1}^N(0, 1) \cdot y & \partial_{yx} \varphi_{\mu,1}^N(0, 1) \cdot x + \partial_{yy} \varphi_{\mu,1}^N(0, 1) \cdot y \\ \partial_{xx} \varphi_{\mu,2}^N(0, 1) \cdot x + \partial_{xy} \varphi_{\mu,2}^N(0, 1) \cdot y & \partial_{yx} \varphi_{\mu,2}^N(0, 1) \cdot x + \partial_{yy} \varphi_{\mu,2}^N(0, 1) \cdot y \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}.$$

Das hipóteses **(ii)** e **(iii)** obtemos que $\varphi_\mu^N(0, 1) = (1, \mu)$. Analisemos a hipótese **(i)**, isto é, que a função real $\varphi_{\mu,2}^N$ restrita à $W^u(p_\mu)$ tem 1 como máximo local. A função $\varphi_{\mu,2}^N$ vai de um aberto do $\mathbb{R}^2$ para $\{0\} \times \mathbb{R} \simeq \mathbb{R}$, ou seja, o eixo y . Quando restringimos nosso domínio à variedade instável local $W^u(p_\mu)$ estamos olhando para as duas componentes locais de $W^u(p_\mu)$. Porém, como queremos o máximo local no ponto $(0, 1)$ do domínio, basta olharmos para uma pequena vizinhança deste ponto em $W^u(p_\mu)$ (que estará contida somente no eixo y). Quando fazemos isso, estamos considerando, localmente, $x = 0$. Veja a Figura 4.27.

Feito isso, devemos olhar para as condições de maximalidade da função restrição $y \mapsto \varphi_{\mu,2}^N(0, y)$, que são

$$\partial_y \varphi_{\mu,2}^N(0, 1) = 0 \quad \text{e} \quad \beta := \partial_{yy} \varphi_{\mu,2}^N(0, 1) < 0.$$

Figura 4.27: Aplicação $\varphi_{\mu,2}^N(0, \cdot)$.

Assim, obtemos que

$$\det(D(\varphi_\mu^N)_{(0,1)}) = \partial_y \varphi_{\mu,1}^N(0,1) \cdot \partial_x \varphi_{\mu,2}^N(0,1).$$

Como φ_μ^N é um difeomorfismo, segue que $D(\varphi_\mu^N)_{(0,1)}$ é inversível, ou seja, que $\det(D(\varphi_\mu^N)_{(0,1)}) \neq 0$. Portanto, obtemos que

$$\alpha := \partial_y \varphi_{\mu,1}^N(0,1) \neq 0 \quad \text{e} \quad \gamma := \partial_x \varphi_{\mu,2}^N(0,1) \neq 0.$$

Assim, a expansão em série de Taylor fica na forma

$$\varphi_\mu^N(x, 1+y) = (0,1) + (\alpha \cdot y + \tilde{H}_1(\mu, x, y), \beta \cdot y^2 + \mu + \gamma \cdot x + \tilde{H}_2(\mu, x, y))$$

com

$$\begin{aligned} \tilde{H}_1(\mu, x, y) &= \partial_x \varphi_{\mu,1}^N(0,1) \cdot x + \frac{1}{2} \cdot \partial_{xx} \varphi_{\mu,1}^N(0,1) \cdot x^2 + \partial_{xy} \varphi_{\mu,1}^N(0,1) \cdot xy + \\ &\quad + \partial_{yy} \varphi_{\mu,1}^N(0,1) \cdot y^2 + O(3) \\ \tilde{H}_2(\mu, x, y) &= \frac{1}{2} \cdot \partial_{xx} \varphi_{\mu,2}^N(0,1) \cdot x^2 + \partial_{xy} \varphi_{\mu,2}^N(0,1) \cdot xy + O(3) \end{aligned} \quad (4.7)$$

satisfazendo, claramente, (4.6) por causa de sua expressão polinomial.

As funções H_i e $\tilde{H}_i$ são de classe C^2 porque φ_μ é de classe C^∞ e as coordenadas linearizantes x e y são de classe C^2 . O que justifica nossa afirmação.

Agora, definimos uma reparametrização n -dependente para parâmetro μ e uma mudança de coordenadas μ -dependente pelas seguintes fórmulas

$$\begin{aligned}\mu &= \sigma^{-2n} \cdot \bar{\mu} - \gamma \cdot \lambda^n + \sigma^{-n}, & \bar{\mu} &= \sigma^{2n} \cdot \mu + \gamma \lambda^n \cdot \sigma^{2n} - \sigma^n, \\ x &= 1 + \sigma^{-n} \cdot \bar{x}, & \bar{x} &= \sigma^n(x - 1), \\ y &= \sigma^{-n} + \sigma^{-2n} \cdot \bar{y}, & \bar{y} &= \sigma^{2n} \cdot y - \sigma^{-n}.\end{aligned}$$

Nas expressões acima λ e σ dependem de μ e, portanto, de $\bar{\mu}$ embora não esteja expresso nas fórmulas acima.

Afirmção. Fixe uma caixa B nas coordenadas $(\bar{x}, \bar{y})$. Se passarmos B para as coordenadas n -dependentes (x, y) , então B torna-se caixas n -dependentes B_n que convergem para $\{q\}$ quando $n \rightarrow \infty$.

De fato: seja $B = \{(\bar{x}, \bar{y}) \mid a \leq \bar{x} \leq b, c \leq \bar{y} \leq d\}$ uma caixa nas coordenadas $(\bar{x}, \bar{y})$. Então, para cada $n \in \mathbb{N}$,

$$\begin{aligned}B_n &= \{(x, y) \mid a \leq \sigma^n(x - 1) \leq b, c \leq \sigma^{2n}y - \sigma^n \leq d\} \\ &= \{(x, y) \mid a\sigma^{-n} + 1 \leq x \leq b\sigma^{-n} + 1, (c - \sigma^n)\sigma^{-2n} \leq y \leq (d - \sigma^n)\sigma^{-2n}\}.\end{aligned}$$

Quando fazemos $n \rightarrow \infty$ temos que

$$\begin{aligned}\lim_{n \rightarrow \infty} (a\sigma^{-n} + 1) &= 1 \leq x \leq 1 = \lim_{n \rightarrow \infty} (b\sigma^{-n} + 1) \\ \lim_{n \rightarrow \infty} (c\sigma^{-2n} - \sigma^{-n}) &= 0 \leq y \leq 0 = \lim_{n \rightarrow \infty} (d\sigma^{-2n} - \sigma^{-n}),\end{aligned}$$

ou seja, as caixas B_n convergem para $\{(1, 0)\} = \{q\}$, nas coordenadas (x, y) .

Agora vamos aos nossos cálculos principais: expressar φ_μ^{n+N} em termos de $\bar{\mu}, \bar{x}$ e $\bar{y}$. Seja $(\bar{\mu}, \bar{x}, \bar{y})$ um ponto. As coordenadas (μ, x, y) deste ponto são

$$\mu = \sigma^{-2n}\bar{\mu} - \gamma\lambda^n + \sigma^{-n}, \quad x = 1 + \sigma^{-n}\bar{x} \quad \text{e} \quad y = \sigma^{-n} + \sigma^{-2n}\bar{y}.$$

Aplicamos φ_μ neste ponto e obtemos

$$\begin{aligned}\varphi_\mu^n(x, y) &= (\lambda^n x, \sigma^n y) \\ &= (\lambda^n(1 + \sigma^{-n}\bar{x}), \sigma^n(\sigma^{-n} + \sigma^{-2n}\bar{y})) \\ &= (\lambda^n(1 + \sigma^{-n}\bar{x}), 1 + \sigma^{-n}\bar{y}).\end{aligned}$$

Portanto, $x = \lambda^n(1 + \sigma^{-n}\bar{x})$ e $y = 1 + \sigma^{-n}\bar{y}$. Agora, aplicamos φ_μ^N e encontramos

$$\begin{aligned}x &= 1 + \alpha\sigma^{-n}\bar{y} + \tilde{H}_1(\sigma^{-2n}\bar{\mu} - \gamma\lambda^n + \sigma^{-n}, \lambda^n(1 + \sigma^{-n}\bar{x}), \sigma^{-n}\bar{y}) \\ y &= \beta(\sigma^{-n}\bar{y})^2 + (\sigma^{-2n}\bar{\mu} - \gamma\lambda^n + \sigma^{-n}) + \gamma(\lambda^n(1 + \sigma^{-n}\bar{x})) + \\ &\quad + \tilde{H}_1(\sigma^{-2n}\bar{\mu} - \gamma\lambda^n + \sigma^{-n}, \lambda^n(1 + \sigma^{-n}\bar{x}), \sigma^{-n}\bar{y}).\end{aligned}$$

Transformando-as de volta às coordenadas $(\bar{x}, \bar{y})$ e, para não haver confusão, denotando-as por $\bar{\bar{x}}$ e $\bar{\bar{y}}$, obtemos

$$\begin{aligned}1 + \sigma^{-n}\bar{\bar{x}} &= 1 + \alpha\sigma^{-n}\bar{y} + \tilde{H}_1 \\ \sigma^{-n} + \sigma^{-2n}\bar{\bar{y}} &= \beta\sigma^{-2n}\bar{y}^2 + \sigma^{-2n}\bar{\mu} + \sigma^{-n} - \gamma\lambda^n\sigma^{-n}\bar{x} + \tilde{H}_2,\end{aligned}$$

ou seja,

$$\begin{aligned}\bar{x} &= \alpha\bar{y} + \sigma^n \tilde{H}_1(\sigma^{-2n}\bar{\mu} - \gamma\lambda^n + \sigma^{-n}, \lambda^n(1 + \sigma^{-n}\bar{x}), \sigma^{-n}\bar{y}) \\ \bar{y} &= \beta\bar{y}^2 + \bar{\mu} - \gamma\lambda^n\sigma^n\bar{x} + \tilde{H}_2(\sigma^{-2n}\bar{\mu} - \gamma\lambda^n + \sigma^{-n}, \lambda^n(1 + \sigma^{-n}\bar{x}), \sigma^{-n}\bar{y}).\end{aligned}$$

Mostremos que certos termos da expressão anterior convergem para zero, quando $n \rightarrow \infty$, na topologia C^2 (uniformemente sobre compactos nas coordenadas $(\bar{\mu}, \bar{x}, \bar{y})$). Na expressão de $\bar{y}$, o termo $\gamma\lambda^n\sigma^n\bar{x} = \gamma(\lambda\sigma)^n\bar{x}$ converge para zero porque temos $0 < \lambda\sigma < 1$. Já os termos que envolvem $\tilde{H}_i$ são mais complicados.

Observe que se as variáveis $(\bar{\mu}, \bar{x}, \bar{y})$ são limitadas, então os valores

$$\mu = \sigma^{-2n}\bar{\mu} - \gamma\lambda^n + \sigma^{-n}, \quad x = \lambda^n(1 + \sigma^{-n}\bar{x}) \quad \text{e} \quad y - 1 = \sigma^{-n}\bar{y}$$

que são substituídos em $\tilde{H}_i$ satisfazem

$$\mu = O(\sigma^{-n})$$

$$x = O(\lambda^n)$$

$$y = O(\sigma^{-n})$$

quando $n \rightarrow \infty$. De fato, como $\bar{\mu}$ é limitado e $0 < \lambda\sigma < 1$ temos que

$$\begin{aligned}\lim_{n \rightarrow \infty} [\sigma^{-2n}\bar{\mu} - \gamma\lambda^n + \sigma^{-n}] &= \lim_{n \rightarrow \infty} [(\sigma^{-n})(\sigma^{-n}\bar{\mu} - \gamma(\lambda\sigma)^n + 1)] \\ &= \left(\lim_{n \rightarrow \infty} \sigma^{-n} \right) \left(\lim_{n \rightarrow \infty} [\sigma^{-n}\bar{\mu} - \gamma(\lambda\sigma)^n + 1] \right) \\ &= \lim_{n \rightarrow \infty} \sigma^{-n},\end{aligned}$$

então para n suficientemente grande, μ é da ordem de σ^{-n} . O mesmo raciocínio se aplica em x e y tendo em vista que

$$\begin{aligned}\lim_{n \rightarrow \infty} [\lambda^n(1 + \sigma^{-n}\bar{x})] &= \left(\lim_{n \rightarrow \infty} \lambda^n \right) \left(\lim_{n \rightarrow \infty} [1 + \sigma^{-n}\bar{x}] \right) \\ &= \lim_{n \rightarrow \infty} \lambda^n,\end{aligned}$$

pois $\bar{x}$ é limitado, e que

$$\lim_{n \rightarrow \infty} [\sigma^{-n} \bar{y}] = \bar{y} \cdot \lim_{n \rightarrow \infty} \sigma^{-n},$$

pois $\bar{y}$ é limitado.

Definimos

$$\bar{H}_1(\bar{\mu}, \bar{x}, \bar{y}) = \sigma^n \cdot \tilde{H}_1(\sigma^{-2n} \bar{\mu} - \gamma \lambda^n + \sigma^{-n}, \lambda^n(1 + \sigma^{-n} \bar{x}), \sigma^{-n} \bar{y}).$$

Então, utilizando a expressão de $\tilde{H}_1$ dada em (4.7) e o fato que $0 < \lambda\sigma < 1$, vem que

$$\begin{aligned} \bar{H}_1(0, 0, 0) &= \sigma^n \cdot \tilde{H}_1(-\gamma \lambda^n + \sigma^{-n}, \lambda^n, 0) \\ &= \sigma^n \cdot (\partial_x \varphi_{(-\gamma \lambda^n + \sigma^{-n}), 1}^N(0, 1) \cdot \lambda^n + \partial_{xx} \varphi_{(-\gamma \lambda^n + \sigma^{-n}), 1}^N(0, 1) \cdot \lambda^{2n} + O(\lambda^{2n})) \\ &= \sigma^n \cdot (O(\lambda^n) + O(\lambda^{2n})) \\ &= O((\lambda\sigma)^n) + O((\lambda\sigma)^n \lambda^n) \end{aligned}$$

que converge para zero quando $n \rightarrow \infty$.

Além disso, as derivadas de primeira e segunda ordem de $\bar{H}_1(\bar{\mu}, \bar{x}, \bar{y})$ convergem para zero uniformemente sobre compactos. De fato, segue de (4.6), de $0 < \lambda\sigma < 1$ e por $\tilde{H}_1$ ser de classe C^2 , que

- $\partial_x \bar{H}_1(0, 0, 0) = 0$, pois, uma vez que

$$\begin{aligned} \partial_x \bar{H}_1(0, 0, 0) &= [\sigma^n \cdot \partial_x \tilde{H}_1 \cdot \lambda^n \sigma^{-n}]|_{(0,0,0)} \\ &= \lambda^n \cdot \partial_x \tilde{H}_1(-\gamma \lambda^n + \sigma^{2n}, \lambda^n, 0), \end{aligned}$$

quando $n \rightarrow \infty$, segue que

$$\partial_x \bar{H}_1(0, 0, 0) = \lim_{n \rightarrow \infty} \lambda^n \cdot \partial_x \tilde{H}_1(0, 0, 0) = 0;$$

- $\partial_y \bar{H}_1(0, 0, 0) = 0$, pois, uma vez que

$$\begin{aligned}\partial_y \bar{H}_1(0, 0, 0) &= [\sigma^n \cdot \partial_y \tilde{H}_1 \cdot \sigma^{-n}]|_{(0,0,0)} \\ &= \partial_y \tilde{H}_1(-\gamma\lambda^n + \sigma^{2n}, \lambda^n, 0),\end{aligned}$$

quando $n \rightarrow \infty$, segue que

$$\partial_y \bar{H}_1(0, 0, 0) = \lim_{n \rightarrow \infty} \partial_y \tilde{H}_1(0, 0, 0) = 0;$$

- $\partial_\mu \bar{H}_1(0, 0, 0) = 0$, pois, uma vez que

$$\begin{aligned}\partial_\mu \bar{H}_1(0, 0, 0) &= [\sigma^n \cdot \partial_\mu \tilde{H}_1 \cdot \sigma^{-2n}]|_{(0,0,0)} \\ &= \sigma^{-n} \cdot \partial_\mu \tilde{H}_1(-\gamma\lambda^n + \sigma^{2n}, \lambda^n, 0),\end{aligned}$$

quando $n \rightarrow \infty$, segue que

$$\partial_\mu \bar{H}_1(0, 0, 0) = \lim_{n \rightarrow \infty} \sigma^{-n} \cdot \partial_\mu \tilde{H}_1(0, 0, 0) = 0;$$

- $\partial_{xx} \bar{H}_1(0, 0, 0) = 0$, pois, uma vez que

$$\begin{aligned}\partial_{xx} \bar{H}_1(0, 0, 0) &= [\partial_x(\lambda^n \cdot \partial_x \tilde{H}_1)]|_{(0,0,0)} \\ &= \lambda^{2n} \sigma^{-n} \cdot \partial_{xx} \tilde{H}_1(-\gamma\lambda^n + \sigma^{2n}, \lambda^n, 0),\end{aligned}$$

quando $n \rightarrow \infty$, segue que

$$\partial_{xx} \bar{H}_1(0, 0, 0) = \lim_{n \rightarrow \infty} \lambda^{2n} \sigma^{-n} \cdot \partial_{xx} \tilde{H}_1(0, 0, 0) = 0;$$

- $\partial_{xy}\bar{H}_1(0,0,0) = 0$, pois, uma vez que

$$\begin{aligned}\partial_{xy}\bar{H}_1(0,0,0) &= [\partial_x(\partial_y\tilde{H}_1)]|_{(0,0,0)} \\ &= \lambda^n\sigma^{-n} \cdot \partial_{xy}\tilde{H}_1(-\gamma\lambda^n + \sigma^{2n}, \lambda^n, 0),\end{aligned}$$

quando $n \rightarrow \infty$, segue que

$$\partial_{xy}\bar{H}_1(0,0,0) = \lim_{n \rightarrow \infty} \lambda^n\sigma^{-n} \cdot \partial_{xy}\tilde{H}_1(0,0,0) = 0;$$

- $\partial_{yy}\bar{H}_1(0,0,0) = 0$, pois, uma vez que

$$\begin{aligned}\partial_{yy}\bar{H}_1(0,0,0) &= [\partial_y(\partial_y\tilde{H}_1)]|_{(0,0,0)} \\ &= \sigma^{-n} \cdot \partial_{yy}\tilde{H}_1(-\gamma\lambda^n + \sigma^{2n}, \lambda^n, 0),\end{aligned}$$

quando $n \rightarrow \infty$, segue que

$$\partial_{yy}\bar{H}_1(0,0,0) = \lim_{n \rightarrow \infty} \sigma^{-n} \cdot \partial_{yy}\tilde{H}_1(0,0,0) = 0;$$

- $\partial_{\mu x}\bar{H}_1(0,0,0) = 0$, pois, uma vez que

$$\begin{aligned}\partial_{\mu x}\bar{H}_1(0,0,0) &= [\partial_\mu(\lambda^n \cdot \partial_x\tilde{H}_1)]|_{(0,0,0)} \\ &= \lambda^n\sigma^{-2n} \cdot \partial_{\mu x}\tilde{H}_1(-\gamma\lambda^n + \sigma^{2n}, \lambda^n, 0),\end{aligned}$$

quando $n \rightarrow \infty$, segue que

$$\partial_{\mu x}\bar{H}_1(0,0,0) = \lim_{n \rightarrow \infty} \lambda^n\sigma^{-2n} \cdot \partial_{\mu x}\tilde{H}_1(0,0,0) = 0;$$

- $\partial_{\mu y} \bar{H}_1(0, 0, 0) = 0$, pois, uma vez que

$$\begin{aligned} \partial_{\mu y} \bar{H}_1(0, 0, 0) &= [\partial_\mu(\partial_y \tilde{H}_1)]|_{(0,0,0)} \\ &= \sigma^{-2n} \cdot \partial_{\mu y} \tilde{H}_1(-\gamma\lambda^n + \sigma^{2n}, \lambda^n, 0), \end{aligned}$$

quando $n \rightarrow \infty$, segue que

$$\partial_{\mu y} \bar{H}_1(0, 0, 0) = \lim_{n \rightarrow \infty} \sigma^{-2n} \cdot \partial_{\mu y} \tilde{H}_1(0, 0, 0) = 0;$$

- $\partial_{\mu\mu} \bar{H}_1(0, 0, 0) = 0$, pois, uma vez que

$$\begin{aligned} \partial_{\mu\mu} \bar{H}_1(0, 0, 0) &= [\partial_\mu(\sigma^{-n} \cdot \partial_\mu \tilde{H}_1)]|_{(0,0,0)} \\ &= \sigma^{-3n} \cdot \partial_{\mu\mu} \tilde{H}_1(-\gamma\lambda^n + \sigma^{2n}, \lambda^n, 0), \end{aligned}$$

quando $n \rightarrow \infty$, segue que

$$\partial_{\mu\mu} \bar{H}_1(0, 0, 0) = \lim_{n \rightarrow \infty} \sigma^{-3n} \cdot \partial_{\mu\mu} \tilde{H}_1(0, 0, 0) = 0;$$

- $\bar{H}_1(0, 0, 0) = 0$, pois, uma vez que

$$\bar{H}_1(0, 0, 0) = \sigma^n \cdot \tilde{H}_1(-\gamma\lambda^n + \sigma^{2n}, \lambda^n, 0),$$

e usando a expressão (4.7) para $\tilde{H}_1$, quando $n \rightarrow \infty$, segue que

$$\begin{aligned} \bar{H}_1(0, 0, 0) &= \lim_{n \rightarrow \infty} \sigma^n \cdot \tilde{H}_1(O(\sigma^{-n}), O(\lambda^n), 0) \\ &= \lim_{n \rightarrow \infty} \sigma^n \cdot (O(\lambda^n) + O(\lambda^{2n})) \\ &= \lim_{n \rightarrow \infty} [O((\lambda\sigma)^n) + O((\lambda\sigma)^n \lambda^n)] \\ &= 0. \end{aligned}$$

O mesmo procedimento funciona para a expressão de $\bar{y}$. Logo, quando fazemos

$n \rightarrow \infty$, as fórmulas de mudança de coordenadas convergem para

$$\begin{pmatrix} \bar{x} \\ \bar{y} \end{pmatrix} \mapsto \begin{pmatrix} \alpha \bar{y} \\ \beta \bar{y}^2 + \bar{\mu} \end{pmatrix},$$

a qual denotaremos por $T_{\bar{\mu}}$. Fazendo a substituição

$$\tilde{\mu} = \beta \bar{\mu}$$

$$\tilde{x} = \alpha^{-1} \beta \bar{x}$$

$$\tilde{y} = \beta \bar{y},$$

que denotaremos por $h(\bar{x}, \bar{x}, \bar{y}) = (\tilde{\mu}, \tilde{x}, \tilde{y})$, a transformação T_{μ} torna-se

$$\begin{pmatrix} \tilde{x} \\ \tilde{y} \end{pmatrix} \mapsto \begin{pmatrix} \tilde{y} \\ \tilde{y}^2 + \tilde{\mu} \end{pmatrix}.$$

De fato, temos que

$$\begin{aligned} \tilde{T}_{\tilde{\mu}}(\tilde{x}, \tilde{y}) &= h \circ T_{\bar{\mu}} \circ h^{-1}(\tilde{\mu}, \tilde{x}, \tilde{y}) \\ &= h \circ T_{\bar{\mu}}(\bar{\mu}, \bar{x}, \bar{y}) \\ &= h(\bar{\mu}, \alpha \bar{y}, \beta \bar{y}^2 + \bar{\mu}) \\ &= (\beta \cdot \bar{\mu}, \alpha^{-1} \beta \cdot (\alpha \bar{y}), \beta \cdot (\beta \bar{y}^2 + \bar{\mu})) \\ &= (\beta \bar{\mu}, \beta \bar{y}, (\beta \bar{y})^2 + \bar{\mu}) \\ &= (\tilde{\mu}, \tilde{y}, \tilde{y}^2 + \tilde{\mu}). \end{aligned}$$

Com isso, o teorema está provado: temos a mudança de coordenadas enunciada como limite de φ_{μ}^{n+N} composta com uma ligeira mudança de coordenadas e reparametrizações de μ . ■

Referências Bibliográficas

- [1] CAMACHO, C.; NETO, A.L. *Teoria geométrica das folheações*. Rio de Janeiro: Instituto Nacional de Matemática Pura e Aplicada, 1979. (Projeto Euclides)
- [2] HARTMAN, P. *Ordinary differential equations*. Boston: Birkhouser, 1982.
- [3] MUNKRES, J.R. *Topology, 2 ed.* United States of America: Prentice Hall, Upper Saddle River, NJ 07458, 2000.
- [4] NEWHOUSE, S.; PALIS, J.; TAKENS, F. Bifurcations and stability of families of diffeomorphisms, *Publications Mathématiques de l'I.H.E.S.*, **57** (1983), p. 5-71.
- [5] PALIS, J.; TAKENS, F. *Hyperbolicity and sensitive-chaotic dynamics at homoclinic bifurcations*. Cambridge: Cambridge University Press, 1993. (Cambridge Studies in Advanced Mathematics)
- [6] ROBINSON, C. *Introduction to the theory of dynamical systems*. Boca Raton: CRC Press, 1995. (Studies in Advances Mathematics)
- [7] SHUB, M. *Global stability of dynamical systems*. New York, Heidelberg, Berlin: Springer-Verlag, 1987.
- [8] STERNBERG, S. On the structure of local homeomorphisms of Euclidean n -space, II, *American Journal of Mathematics*. The University of Chicago, **80** (1958), 623-631.

- [9] YORKE, J.A.;ALLIGOOD, K.T. Cascades of period doubling bifurcations: a prerequisite for horseshoes, *Bulletin of American Society of Mathematics*. **9** (1983), 319-322.

Apêndice A

Conjunto recorrente por cadeias

Esta seção apresenta algumas definições e resultados sobre o que chamamos de *conjuntos recorrentes por cadeia* que, a grosso modo, nada mais é do que o conjunto formado por “falsas órbitas” de um dado ponto, isto é, pontos que estão muito próximo da órbita original e que não são, em geral, pontos dessa órbita.

Definição A.0.1. Seja $f : X \longrightarrow X$ uma aplicação sobre um espaço métrico completo M . Sempre que necessário, consideraremos f um homeomorfismo.

(a) Um conjunto S é um *conjunto minimal* para f se

- (i) S é fechado, não vazio e invariante ($f(S) = S$) e
- (ii) sempre que $B \subset S$ é fechado, não vazio e invariante, então $B = S$.

(b) Dizemos que um ponto $x \in X$ é ω -*recorrente* se $x \in \omega(x)$, onde

$$\omega(x) = \{y \in X : \exists n_k \rightarrow \infty \text{ tal que } \lim_{k \rightarrow \infty} d(f^{n_k}(x), y) = 0\},$$

e é α -*recorrente* se $x \in \alpha(x)$, onde

$$\alpha(x) = \{y \in X : \exists n_k \rightarrow \infty \text{ tal que } \lim_{k \rightarrow \infty} d(f^{-n_k}(x), y) = 0\}.$$

Chamamos $\alpha(x)$ o conjunto α -*limite* de x e $\omega(x)$ o conjunto ω -*limite* de x .

- (c) Uma ε -cadeia de comprimento n de x para y para a aplicação f é uma seqüência $\{x = x_0, x_1, \dots, x_n = y\}$ tal que para todo $j = 1, \dots, n$ temos satisfeito $d(f(x_{j-1}), x_j) < \varepsilon$.

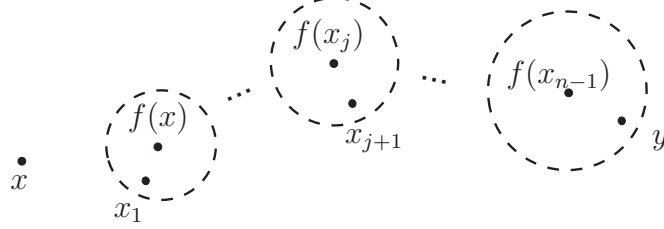


Figura A.1: ε -cadeia.

- (d) Seja $Y \subset X$. O conjunto ε -cadeia limite de Y para f é o conjunto

$$\Omega_{\varepsilon}^{+}(Y, f) = \{y \in X \mid \forall n \geq 1 \text{ existe } x \in Y \text{ e existe uma } \varepsilon\text{-cadeia de } x \text{ para } y \text{ de comprimento maior que } n\}.$$

Como isso, podemos definir o conjunto cadeia limite de Y como

$$\Omega^{+}(Y, f) = \bigcap_{\varepsilon > 0} \Omega_{\varepsilon}^{+}(Y, f).$$

No caso em que $Y = X$, escrevemos $\Omega^{+}(f)$.

De forma semelhante, definimos

$$\begin{aligned} \Omega_{\varepsilon}^{-}(Y, f) &= \{y \in X \mid \forall n \geq 1 \text{ existe } x \in Y \text{ e existe uma } \varepsilon\text{-cadeia de } x \text{ para } y \text{ de comprimento menor que } n\} \text{ e} \\ \Omega^{-}(Y, f) &= \bigcap_{\varepsilon > 0} \Omega_{\varepsilon}^{-}(Y, f). \end{aligned}$$

Finalmente, definimos o conjunto recorrente por cadeias de f por

$$\begin{aligned}\mathcal{R}(f) &= \{x \in X \mid \text{existe uma } \varepsilon\text{-cadeia de } x \text{ para } x, \text{ para todo } \varepsilon\} \\ &= \{x \in X \mid x \in \Omega^+(\{x\}, f)\} = \{x \in X \mid x \in \Omega^-(\{x\}, f)\}.\end{aligned}$$

(e) Um ponto $p \in X$ é chamado *não errante* quando, para toda vizinhança U de p , existe um inteiro $n > 0$ tal que $f^n(U) \cap U \neq \emptyset$. Ou, equivalentemente, quando existe um ponto $q \in U$ com $f^n(q) \in U$. O conjunto de todos pontos não errantes para f é chamado conjunto *não errante* e é denotado por $\Omega(f)$.

Proposição A.0.2. *Seja $f : X \longrightarrow X$ um homeomorfismo, onde X é um espaço métrico completo. Então*

- (i) $\overline{\bigcup_{x \in X} \omega(x, f)} \subset \Omega(f) \subset \mathcal{R}(f)$,
- (ii) $\overline{\bigcup_{x \in X} \alpha(x, f)} \subset \Omega(f) \subset \mathcal{R}(f)$ e
- (iii) $\mathcal{R}(f)$ é fechado.

Demonstração:

Vamos fazer apenas a demonstração de (i) pois (ii) é análogo. Começemos pela primeira inclusão

$$\overline{\bigcup_{x \in X} \omega(x, f)} \subset \Omega(f). \quad (\text{A.1})$$

Seja $y \in \omega(x, f)$, para um dado $x \in X$. Seja U uma vizinhança de y . Então, existem inteiros $n > m > 0$ tais que $f^m(x)$ e $f^n(x)$ pertencem a U . Logo, $U \cap f^{m-n}(U) \neq \emptyset$, uma vez que $f^n(x) \in U$ e $f^{m-n}(f^n(x)) = f^m(x) \in U$. Então, $y \in \Omega(f)$. Portanto, vale (A.1).

Afirmção 1. O conjunto $\Omega(f)$ é fechado.

De fato: seja (x_n) uma seqüência em $\Omega(f)$ tal que $x_n \longrightarrow p$. Mostremos que $p \in \Omega(f)$, isto é, que para toda bola aberta $B(p, \varepsilon)$,

$$\text{existe } k \in \mathbb{N} \text{ tal que } f^k(B(p, \varepsilon)) \cap B(p, \varepsilon) \neq \emptyset. \quad (\text{A.2})$$

Como $(x_n) \longrightarrow p$, para a bola $B(p, \varepsilon)$ acima, existe $n_0 \in \mathbb{N}$ tal que $x_n \in B(p, \varepsilon)$ para todo $n \geq n_0$. Logo, como $x_{n_0} \in \Omega(f)$ e $x_{n_0} \in B(p, \varepsilon)$, existe k_{n_0} tal que $f^{n_0}(B(p, \varepsilon)) \cap B(p, \varepsilon) \neq \emptyset$, como queríamos em (A.2). Logo, $p \in \Omega(f)$. Portanto, $\Omega(f)$ é fechado.

Da Afirmação 1, segue que

$$\overline{\bigcup_{x \in X} \omega(x, f)} \subset \overline{\Omega(f)} = \Omega(f).$$

Mostremos agora a segunda inclusão

$$\Omega(f) \subset \mathcal{R}(f). \quad (\text{A.3})$$

Seja $p \in \Omega(f)$. Mostremos que $p \in \mathcal{R}(f)$, isto é, que

$$\text{existe uma } \varepsilon\text{-cadeia de } x \text{ para } x, \text{ para todo } \varepsilon > 0. \quad (\text{A.4})$$

Seja $\varepsilon > 0$ dado arbitrariamente. Como f é contínua, existe $0 < \delta < \frac{\varepsilon}{2}$ tal que

$$d(p, y) < \delta \implies d(f(p), f(y)) < \frac{\varepsilon}{2}.$$

Seja U uma vizinhança de p contida em $B(p, \delta)$. Como $p \in \Omega(f)$, existe $n \in \mathbb{N}$ tal

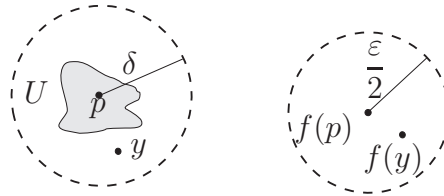


Figura A.2: Vizinhança do aberto U .

que $f^n(U) \cap U \neq \emptyset$. Se $n = 1$, então $\{p, p\}$ é a ε -cadeia procurada, pois existe $q \in U$ tal que $f(q) \in U$ e, com isso,

$$d(f(p), p) \leq d(f(p), f(q)) + d(f(q), p) < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

Se $n > 1$, então existe $q \in U$ tal que $f^n(q) \in U$. Assim, concluímos que $\{p, f(q), f^2(q), \dots, f^{n-1}(q), p\}$ é uma ε -cadeia. De fato, temos

- $d(f(p), f(q)) < \frac{\varepsilon}{2}$, pois $q \in U \subset B(p, \delta)$,
- $d(f(f^j(q)), f^{j+1}(q)) = 0 < \varepsilon$ para $j = 1, \dots, n-2$ e
- $d(f(f^{n-1}(q)), p) = d(f^n(q), p) < \delta < \frac{\varepsilon}{2}$, pois $f^n(q) \in U \subset B(p, \delta)$,

o que prova (A.4) e, portanto, (A.3).

Agora, seja (x_n) uma seqüência em $\mathcal{R}(f)$ tal que $x_n \rightarrow p$. Basta mostrar que $p \in \mathcal{R}(f)$ para concluir que $\mathcal{R}(f)$ é fechado. Mostremos que para todo $\varepsilon > 0$ existe uma ε -cadeia de p para p . Seja $\varepsilon > 0$ dado arbitrariamente. Como f é contínua, existe $0 < \delta < \frac{\varepsilon}{2}$ tal que

$$d(p, y) < \delta \implies d(f(p), f(y)) < \frac{\varepsilon}{2}.$$

Como $x_n \rightarrow p$, existe $n_0 \in \mathbb{N}$ tal que $d(x_n, p) < \delta$, para todo $n \geq n_0$. Logo, para x_{n_0} vale que

$$d(p, y) < \frac{\varepsilon}{2} \text{ e } d(f(x_{n_0}), f(p)) < \frac{\varepsilon}{2}.$$

E como $x_{n_0} \in \mathcal{R}(f)$, existe uma $\frac{\varepsilon}{2}$ -cadeia $\{x_{n_0}, y_1, \dots, y_k, x_{n_0}\}$. Logo, a cadeia $\{p, y_1, \dots, y_k, p\}$ é uma ε -cadeia. De fato, como $\{x_{n_0}, y_1, \dots, y_k, x_{n_0}\}$ é $\frac{\varepsilon}{2}$ -cadeia, temos

- $d(f(p), y) \leq d(f(p), f(x_{n_0})) + d(f(x_{n_0}), y_1) < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon$ e
- $d(f(y_j), f(y_{j+1})) < \frac{\varepsilon}{2}$.

E como, além disso, $d(f(x_{n_0}), p) < \delta < \frac{\varepsilon}{2}$, segue que

$$d(f(y_k), p) \leq d(f(y_k), x_{n_0}) + d(x_{n_0}, p) < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

Logo, $p \in \mathcal{R}(f)$ e, portanto, $\mathcal{R}(f)$ é fechado. O que prova (iii). ■

Apêndice B

Folheações

Neste apêndice daremos uma breve explicação do que é uma folheação e depois daremos uma definição mais formal. Retiramos essas definições e exemplos de [1].

Podemos imaginar que uma folheação de dimensão n para uma variedade M é uma decomposição em subvariedades conexas de dimensão n que são as folhas que se aglomeram localmente como os subconjuntos de $\mathbb{R}^m = \mathbb{R}^n \times \mathbb{R}^{m-n}$ com a segunda coordenada constante. Vejamos o seguinte exemplo.

Exemplo 8. Considere o espaço euclidiano $\mathbb{R}^m = \mathbb{R}^n \times \mathbb{R}^{m-n}$. Uma folheação para esse espaço (que também é uma variedade) é considerar como folhas os n -planos da forma $\mathbb{R}^n \times \{c\}$, com $c \in \mathbb{R}^{m-n}$. Se $m = 2$ e $n = 1$, temos o plano decomposto em retas verticais.

Veremos agora a definição formal.

Definição B.0.3. Seja M uma variedade de dimensão m e classe C^∞ . Uma *folheação* de classe C^r e dimensão n para M é um atlas (máximo) $\mathcal{F}$ de classe C^r em M com as seguintes propriedades:

- (a) Se $(U, \varphi) \in \mathcal{F}$, então $\varphi(U) = U_1 \times U_2 \subset \mathbb{R}^n \times \mathbb{R}^{m-n}$, onde U_1 e U_2 são discos abertos de $\mathbb{R}^n$ e de $\mathbb{R}^{m-n}$, respectivamente.
- (b) Se (U, φ) e (V, ψ) pertencentes a $\mathcal{F}$ são tais que $U \cap V \neq \emptyset$, então a mudança

de coordenadas $\psi \circ \varphi^{-1} : \varphi(U \cap V) \rightarrow \Psi(U \cap V)$ é da seguinte forma

$$h(x, y) = (h_1(x, y), h_2(y)), \quad (x, y) \in \mathbb{R}^n \times \mathbb{R}^{m-n},$$

isto é, $\psi \circ \varphi^{-1}(x, y) = (h_1(x, y), h_2(y))$. Dizemos também que M é *folheada* por $\mathcal{F}$.

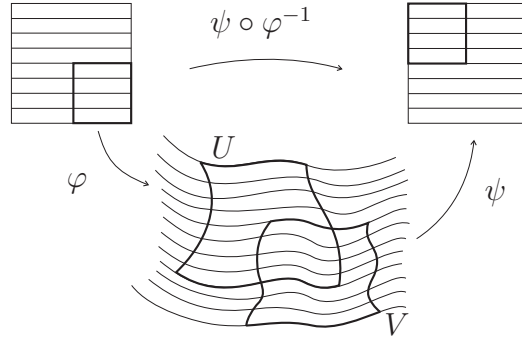


Figura B.1: Uma folheação vista localmente.

Veja a Figura B.1, onde ilustramos a aparência de uma variedade de dimensão 2 folheada por uma folheação de dimensão 1, vista localmente.

As cartas $(U, \phi) \in \mathcal{F}$ serão chamadas também de *cartas trivializadoras de $\mathcal{F}$* .

Seja $\mathcal{F}$ uma folheação de classe C^r e dimensão n , $0 < n < m$ de uma variedade M^n . Consideramos uma carta local (U, ϕ) de $\mathcal{F}$ tal que $\phi(U) = U_1 \times U_2 \subset \mathbb{R}^n \times \mathbb{R}^{m-n}$. Os conjuntos da forma $\phi^{-1}(U_1 \times \{c\})$, $c \in U_2$ são chamados *placas de U* , ou ainda *placas de $\mathcal{F}$* . Fixado $c \in U_2$, a aplicação $f = \phi^{-1}|_{U_1 \times \{c\}} : U_1 \times \{c\} \rightarrow U$ é um mergulho de classe C^r , portanto as placas são subvariedades conexas de dimensão n e classe C^r de M . Além disso, se α e β são duas placas de U , então $\alpha \cap \beta \neq \emptyset$ ou $\alpha = \beta$.

Um *caminho de placas de $\mathcal{F}$* é uma seqüência $\alpha_1, \dots, \alpha_k$ de placas de $\mathcal{F}$ tal que $\alpha_j \cap \alpha_{j+1} \neq \emptyset$ para todo $j \in \{1, \dots, k-1\}$. Como M é recoberta pelas placas de $\mathcal{F}$, podemos definir em M a seguinte relação de equivalência: $p \sim q$ se existe um caminho de placas $\alpha_1, \dots, \alpha_k$ tal que $p \in \alpha_1$ e $q \in \alpha_k$. As classes de

equivalência dessa relação “ $\sim$ ” são chamadas *folhas*. Em outras palavras, uma folha é um “pedaço” conexo uma das “linhas” da Figura B.1.

Da definição segue que uma folha de $\mathcal{F}$ é um subconjunto conexo por caminhos. Com efeito, se F é uma folha de $\mathcal{F}$ e $p, q \in F$, então existe um caminho de placas $\alpha_1, \dots, \alpha_k$ tal que $p \in \alpha_1$ e $q \in \alpha_k$. Como as placas α_i são conexas por caminhos e $\alpha_j \cap \alpha_{j+1} \neq \emptyset$, é imediato que $\alpha_1 \cup \dots \cup \alpha_k \subset F$ é conexo por caminhos, logo existe um caminho contínuo conexo ligando p e q .

Índice Remissivo

órbita

caótica, 74

homoclínica, 22

A-bloco, 72

diâmetro de um, 72

aplicação

topologicamente conjugada, 43

topologicamente transitiva, 47

bifurcação

de Andronov-Hopf, 87

de duplicação de período, 87, 92

desdobra genericamente, 105, 106

flip, 87

homoclínica, 51

sela-nó, 87, 88

cadeia

ε -cadeia, 135

caminho de Möbius, 111

campo

de cones

contínuo, 63

instável, 64

direcional contínuo, 64

conjunto

minimal, 134

condições de não-ressonância, 119

conjunto

α -limite, 134

cadeia limite, 135

de Cantor da reta, 38

ε -cadeia limite, 135

hiperbólico, 61, 62

invariante maximal, 59

não errante, 136

ω -limite, 134

recorrente por cadeias, 32, 135

coordenadas linearizantes, 54

difeomorfismo

dissipativo, 81

localmente dissipativo, 81

espaço

das sequencias de dois símbolos, 40

estrutura hiperbólica, 61

faixa

horizontal, 59

vertical, 59

família quadrática, 82

ferradura

condições iniciais para a criação de
uma, 107

criação de uma, 107

de Smale, 28

geométrica, 28

folha, 141

folheação, 69, 139

caminho de placas de uma, 140

- cartas trivializadoras de uma, 140
- estável, 71
- instável, 69
- placas de uma, 140
- ponto
 - α -recorrente, 134
 - de sela, 54
 - homoclínico, 21
 - homoclínico primário, 55
 - homoclínico transversal, 21
 - não errante, 136
 - ω -recorrente, 134
- ponto fixo hiperbólico, 20
- propriedade genérica, 83
- representação simbólica, 40
- shift, 40, 41
 - bilateral, 40
- subespaço
 - contrativo, 62
 - expansor, 62
- tangência homoclínica, 82
 - desdobra genericamente, 83
 - quadrática, 83
- teorema
 - bifurcação de duplicação de período,
94
 - bifurcação sela-nó, 89
 - da variedade central, 100
- variedade
 - central, 101
 - central local, 101
 - centro-estável local, 101
 - centroinstável local, 101