

C^1 CONNECTING LEMMAS

LAN WEN AND ZHIHONG XIA

ABSTRACT. Like the closing lemma, the connecting lemma is of fundamental importance in dynamical systems. Hayashi recently proved the C^1 connecting lemma for stable and unstable manifolds of a hyperbolic invariant set. In this paper, we prove several very general C^1 connecting lemmas. We simplify Hayashi's proof and extend the results to more general cases.

1. INTRODUCTION

We give a simpler proof for the C^1 connecting lemma of Hayashi [1] in this paper. This is Theorem E below. The theorem is also more general than the original one of Hayashi. It provides certain answers to some old problems as consequences. Since these consequences are also various kinds of C^1 connecting lemmas, and their statements are easier to formulate than Theorem E itself, we first state them as Theorems A–D. Let M be a compact manifold without boundary, and let $f: M \rightarrow M$ be a diffeomorphism. Denote by $\text{Diff}^1(M)$ the set of diffeomorphisms of M , endowed with the C^1 topology. We state these results for diffeomorphisms, but the corresponding results are also true for flows. We further remark that these results generalize easily to symplectic and volume preserving diffeomorphisms.

Theorem A. *Let p and q be two points of M with $\omega(p) \cap \alpha(q) \neq \emptyset$. We also assume that $\omega(p) \cap \alpha(q)$ contains some non-periodic point z . Then for any C^1 neighborhood $\mathcal{U}$ of f , there is $g \in \mathcal{U}$ such that q is on the positive g -orbit of p . More precisely, for any C^1 neighborhood $\mathcal{U}$ of f , there is a positive integer L such that for any $\delta > 0$, there is a $g \in \mathcal{U}$ such that $g = f$ outside the tube $\bigcup_{i=1}^L B(f^{-i}(z), \delta)$ and such that q is on the positive g -orbit of p .*

This gives a partial answer to an old problem of Pugh [9]. Roughly, Theorem A says that two points p and q can get connected via C^1 perturbations if the original forward orbit of p and the original backward orbit of q are nearly connected at some non-periodic point z . Note that the C^0 connecting problem is trivial, just like the C^0 closing problem. Also note that the compactness of M is essential; it guarantees the *Lift Axiom* formulated in [11]. If M is not compact, then Theorem A and the other results of this paper are not true even if the Whitney strong topology is used. An instructive counterexample is given by Pugh [10]. We remark that the original problem of Pugh does not assume that $\omega(p) \cap \alpha(q)$ contains some non-periodic points. This assumption is a technical one. It is demanded by our method, but not

Received by the editors January 24, 1997 and, in revised form, April 13, 1998.

2000 *Mathematics Subject Classification.* Primary 37Cxx, 37Dxx.

Both authors are supported in part by National Science Foundation and the Chinese Natural Science Foundation.

by the nature of the problem. Thus a complete answer to Pugh's original problem needs a separate treatment for the case that every point of $\omega(p) \cap \alpha(q)$ is periodic, which we do not know how to deal with at this point. Similarly, the problem of C^1 closing the n^{th} order prologational recurrence, also raised by Pugh in [9], is solved under the same sort of non-periodicity assumption. More precisely, $p \in M$ is n^{th} order prolongationally recurrent if there are n points $p = p_1, p_2, \dots, p_n \in M$ such that $\omega(p_i) \cap \alpha(p_{i+1}) \neq \emptyset$ for each $i = 1, 2, \dots, n$, where $p_{n+1} = p_1$. We have:

Theorem B. *Let p be n^{th} order prolongationally recurrent. Assume that $\omega(p_i) \cap \alpha(p_{i+1})$ contains non-periodic points for each $i = 1, 2, \dots, n$, where $p_{n+1} = p_1$. Then for any C^1 neighborhood $\mathcal{U}$ of f , there is a $g \in \mathcal{U}$ such that p is a periodic point of g .*

The statement of the second half of Theorem A, which is more precise than the first half, has the advantage that, in case p is not negatively recurrent under f and q is not positively recurrent under f , then δ can be chosen small enough so that the support tube $\bigcup_{i=1}^L B(f^{-i}(z), \delta)$ is disjoint from $\text{Orb}^-(p, f)$ and from $\text{Orb}^+(q, f)$; hence $\text{Orb}^-(p, f) = \text{Orb}^-(p, g)$ and $\text{Orb}^+(q, f) = \text{Orb}^+(q, g)$. For instance, p could go backward to a hyperbolic fixed point p_0 under f , and q could go forward to a hyperbolic fixed point q_0 under f . Then this C^1 perturbation creates a heteroclinic connection from p_0 to q_0 , respecting g . This gives Theorem C, which answers another old problem of Liao [3] and Mañé [5] about creating homoclinic points.

Theorem C. *Let Λ be an isolated hyperbolic set of f . Assume $W^s(\Lambda) \cap \overline{W^u(\Lambda)} - \Lambda \neq \emptyset$. Then for any C^1 neighborhood $\mathcal{U}$ of f , there is a $g \in \mathcal{U}$ such that $g = f$ on a neighborhood U of Λ and such that $W^s(\Lambda, g) \cap W^u(\Lambda, g) - \Lambda \neq \emptyset$.*

Note that if the assumption in Theorem C is replaced by a weaker one, namely $\overline{W^s(\Lambda) \cap W^u(\Lambda)} - \Lambda \neq \emptyset$, the conclusion is still true, as long as $\overline{W^s(\Lambda) \cap W^u(\Lambda)} - \Lambda$ contains non-periodic points, as is easily seen from Theorem E below. Another similar problem of this type appears in the study of the C^1 stability conjecture. We state an answer as Theorem D.

Theorem D. *Let Λ be an isolated hyperbolic set of f . Assume that periodic orbits outside Λ accumulate on Λ . Then for any C^1 neighborhood $\mathcal{U}$ of f , there is a $g \in \mathcal{U}$ such that $g = f$ on a neighborhood of Λ and $W^s(\Lambda, g) \cap W^u(\Lambda, g) - \Lambda \neq \emptyset$.*

All these theorems are straightforward consequences of the following general version of the C^1 connecting lemma.

Theorem E. *Let $z \in M$ not be a periodic point of f . For any C^1 neighborhood $\mathcal{U}$ of f , there are $\rho > 1$, $L \in \mathbb{N}$ and $\delta_0 > 0$ such that for any $0 < \delta \leq \delta_0$, and any two points p and q outside the tube $\Delta = \bigcup_{n=1}^L f^{-n}B(z, \delta)$, if the positive f -orbit of p hits the ball $B(z, \delta/\rho)$ after p , and if the negative f -orbit of q hits the same ball, then there is $g \in \mathcal{U}$ such that $g = f$ off Δ and q is on the positive g -orbit of p .*

Remark. Here we require that the positive orbit of p hits the ball *after* p just to guarantee that the positive orbit of p goes through the tube at least once before that hit. We do not require this for the negative orbit of q , because the tube Δ has been taken along the negative orbit of z . Symmetricly, we can restate Theorem E for a tube along the positive orbit of z , and require that the negative orbit of q goes through the tube at least once before the hit. The same is true for Theorem F below.

Note that the formulation of Theorem E is more complicated than Theorems A–D. However, Theorem E assumes less. It does not assume any limit behavior for z except that it is not a periodic point, and neither ω -limit set nor hyperbolicity is involved in the statement of the theorem. Hence it is more general than Theorems A–D, and more flexible in applications. Note that δ can always be chosen so small that $B(z, \delta)$ is disjoint from Δ , because z is non-periodic and δ is independent of L , and because the tube Δ covers iterates from 1 to L , but not from 0 to L . Thus the special case of Theorem E for which the point q itself is in $B(z, \delta/\rho)$ will read as the following Theorem F, which is convenient for some applications (see §7).

Theorem F. *Let $z \in M$ be a non-periodic point of f . For any C^1 neighborhood $\mathcal{U}$ of f , there are $\rho > 1$, $L \in \mathbb{N}$ and $\delta_0 > 0$ such that for any $0 < \delta \leq \delta_0$, and for any point p outside the tube $\Delta = \bigcup_{n=1}^L f^{-n}B(z, \delta)$ and any point $q \in B(z, \delta/\rho)$, if the positive f -orbit of p hits $B(z, \delta/\rho)$ after p , then there is $g \in \mathcal{U}$ such that $g = f$ off Δ and q is on the positive g -orbit of p .*

The C^1 connecting lemma is a long-desired result. Many authors have made important contributions to this problem. For the case that M is the 2-sphere, Robinson [12] and Pixton [7] solved the problem for any C^r , $1 \leq r \leq \infty$. For volume-preserving diffeomorphisms, Takens [15] solved the problem for $r = 1$, and Oliveira [6] solved the problem for any C^r , $1 \leq r \leq \infty$, when M is the 2-torus. Mañé [5] solved the problem for $r = 1, 2$ with an additional measure theoretic assumption. Liao [3] solved the problem for $r = 1$ with an additional topological assumption. A recent surprising result came when Hayashi [1] proved a general C^1 connecting lemma which solved Theorems C and D. The present paper proves another general version of the C^1 connecting lemma (Theorem E), which enables us to also obtain Theorems A and B. What encourages us is that the proof of Theorem E presented in this paper turns out not to be very long. We first formulate a basic C^1 perturbation theorem, which can be extracted from the work of Liao, Pugh, and Robinson (see [2], [8] and [11]) on the C^1 closing lemma. This is Theorem 3.1 below, which serves as a fundamental preliminary for our proof of the C^1 connecting lemma. A key ingredient then added in is Hayashi’s brilliant idea of “cutting” [1]. Finally, with an intensive use of these beautiful ideas in a series of combinatorial selections (Xia [16]), we are able to cut short some disjoint orbital arcs with C^1 perturbations, and eventually get the two points p and q connected.

This work is an outgrowth of an earlier paper [16] by the second author, which contains a self-contained proof of the C^1 connecting lemma for an important case, and contains all the crucial ideas of the present paper. We wish to thank J. Palis, Charles Pugh, and C. Robinson for many good critical comments and suggestions.

In §2 we introduce Theorem 2.2, which is a linear version of Theorem E. In §3 we formulate a basic C^1 perturbation theorem. In §4 we give an arrangement of boxes. These two sections serve as preparations for proving Theorem 2.2. The proof of Theorem 2.2 itself will be given in §5. In §6 we describe a linearization process needed to realize Theorem 2.2 on manifolds to obtain Theorem E. §7 is devoted to applications of Theorem E.

2. A LINEAR VERSION OF THE C^1 CONNECTING LEMMA

We introduce a linear version of Theorem E in this section. This is Theorem 2.2 below. Its formulation uses a geometrical notion called ε -kernel avoiding transition, which is due to Mai [4] and is the basic pattern for the C^1 perturbation constructed

below. This way of constructing perturbations in proving the C^1 closing lemma actually appeared very early ([8]). It is just the unified notion of ε -kernel avoiding transition that appeared relatively late ([4], [14]). First we define ε -kernel lifts, which serve as the basic elements of our C^1 perturbations.

Let $B \subset \mathbb{R}^m$ be a closed ball with radius r and let $\varepsilon > 0$. We denote by εB the ball of the same center and of radius εr . We call εB the ε -kernel of B . Thus the number ε here gives a relative ratio but not an absolute size. For any x and y in the interior of B , there is a (in fact many) C^∞ diffeomorphism $h: \mathbb{R}^m \rightarrow \mathbb{R}^m$ that is the identity outside B , which takes x to y . If x and y are in εB , we call such an h an ε -kernel lift that lifts x to y , supported on B . The following simple but fundamental lemma tells how ε controls the first derivatives of $h - id$ for certain ε -kernel lifts h . The formal formulation of this fact with the proof on manifolds can be found in [11, Theorem 6.1].

Lemma 2.1. *For any $\beta > 0$, there is an $\varepsilon > 0$ such that for any closed ball B in $\mathbb{R}^m$, and any x and y in εB , there is an ε -kernel lift h that lifts x to y , supported on B , such that all partial derivatives of $h - id$ have absolute values less than β .*

Proof. The proof is easy. We only need to consider the case that x is the center of the ball, because the composition of two lifts of this type gives what we want. Fix a C^∞ bump function $\alpha: \mathbb{R}^m \rightarrow \mathbb{R}$ such that $0 \leq \alpha \leq 1$ on all $\mathbb{R}^m$, $\alpha = 1$ on $B(0, \frac{1}{3})$, $\alpha = 0$ off $B(0, \frac{2}{3})$, and such that all partial derivatives of α have absolute values less than or equal to 6. Let r be the radius of B , and let $y \in \varepsilon B$. Define $h: \mathbb{R}^m \rightarrow \mathbb{R}^m$ by

$$h(u) = u + \alpha\left(\frac{u-x}{r}\right)(y-x).$$

Then h is a diffeomorphism if ε is small enough, and for any i ,

$$\left| \frac{\partial}{\partial u_i} (h - id) \right| \leq \frac{1}{r} \cdot 6 \cdot \varepsilon r = 6\varepsilon.$$

This proves the lemma. □

Roughly, the number ε controls the size of the first derivatives of $h - id$. Note that the radius r of B is not mentioned in the statement of Lemma 2.1, which clearly controls the C^0 size of $h - id$. Therefore the ε -kernel lift h can be defined to be C^1 close to the identity if both ε and r are small, and the composition $h \circ f$ hence gives a C^1 perturbation of f . The C^1 perturbations used in this paper will be a composition of f with finitely many of this kind of ε -kernel lifts with disjoint supports. By virtue of Lemma 2.1, we will not mention the ε -kernel lift h explicitly, but only mention the ball B and the two points $x, y \in \varepsilon B$. Whenever such B, x , and y are specified, we can perform a suitable ε -kernel lift h at any time. In this way we define ε -kernel avoiding transitions, which are the basic patterns of C^1 perturbations used below. Let $V_0, V_1, \dots, V_n, \dots$, be a sequence of m -dimensional inner product spaces, and let $T_n: V_n \rightarrow V_{n-1}$, $n = 1, 2, \dots$, be a sequence of linear isomorphisms. Let $\varepsilon > 0$, $u, v \in V_0$, $L \in \mathbb{N}$, $Q \subset V_0$, and $G \subset V_0$ be given. By an ε -kernel transition of $\{T_n\}$ from u to v of length L contained in Q and avoiding G , we mean $L+1$ points c_n , $0 \leq n \leq L$, together with L balls $B_n \subset V_n$, $0 \leq n \leq L-1$, such that

1. $c_0 = v$ and $c_L = F_L^{-1}(u)$, where $F_n = T_1 \circ T_2 \circ \dots \circ T_n$.
2. $c_n \in \varepsilon B_n$ and $T_{n+1}(c_{n+1}) \in \varepsilon B_n$, $0 \leq n \leq L-1$.

3. $B_n \subset F_n^{-1}(Q)$, $0 \leq n \leq L-1$.

4. $B_n \cap F_n^{-1}(G) = \emptyset$, $0 \leq n \leq L-1$.

Two ε -kernel transitions $c_0, c_1, \dots, c_L$; $B_0, B_1, \dots, B_{L-1}$ and $c'_0, c'_1, \dots, c'_L$; $B'_0, B'_1, \dots, B'_{L-1}$, contained respectively in Q_1 and Q_2 , are *disjoint* if $B_n \cap B'_n = \emptyset$ for all $0 \leq n \leq L-1$.

Roughly, a transition of length L consists of $L+1$ points that form a pseudo-orbit with L jumps. The associated L balls indicate that these L jumps are ε -kernel lifts. The containing set Q and the avoidance set G are constraints put on the transition. Note that the terminologies defined here are abbreviated ones. Such an ε -kernel transition actually is from $F_L^{-1}(u)$ to v , and is contained in the tube $\bigcup_{n=1}^L F_n^{-1}(Q)$, and avoid a set of orbital arcs $\bigcup_{n=1}^L F_n^{-1}(G)$. We emphasize that, in the definition of disjointness of two transitions, we do not require that Q_1 and Q_2 are disjoint. We merely require that B_n and B'_n in V_n are disjoint (and hence the $2L$ balls are mutually disjoint, since the V_n 's are distinct vector spaces. Later on the V_n 's will correspond to disjoint neighborhoods on the manifold). This is sufficient for our purpose because it is these balls B_n that support our C^1 perturbation.

Remark. We insert an informal illustration here on what these V_n and T_n have to do with M and f . Applied to the manifold via some standard linearization along a finite orbit of length L , these V_n , $n = 0, 1, \dots, L-1$, simply correspond to disjoint neighborhoods of the iterates along a backward orbit of f , and these T_n simply correspond to f itself. To see the dynamics we mix them up just for illustration. Thus the transition transits a point from one orbit of f to another orbit of f via L lifts which form a pseudo-orbit. If Q is small (which bounds the C^0 size of the perturbation), and if ε is small too, then the transition gives a C^1 small perturbation. An important case is that u is on the positive orbit of v before perturbation, say $u = f^N(v)$. Then N must be much larger than L , since it needs more than L iterates for $v \in Q$ to get back, for the first time, to Q at some point $v_1 \in Q$. Then it again needs more than L iterates for $v_1 \in Q$ to get back to Q at some $v_2 \in Q$, etc. Since u is some return of v , say, $u = v_k$, $k \geq 1$, the L^{th} pull-back $f^{-L}(u)$ of u , which corresponds to $F_L^{-1}(u)$, is still on the positive orbit of v under f . Now in case the avoidance set G contains the set of intermediate returns $\{v_1, v_2, \dots, v_{k-1}\}$, then the old orbit from v to $f^{-L}(u)$ remains unchanged. Hence this perturbation creates a periodic orbit through v , which goes from v to $f^{-L}(u)$ via the old orbits, and from $f^{-L}(u)$ to v via the transition. This is the way the perturbations are constructed in proving the C^1 closing lemma. The novelty of the perturbation constructed in proving the C^1 connecting lemma is that, while the perturbation for the closing case consists of a single transition in a tube, the perturbation for the connecting case consists of finitely many disjoint transitions in the same tube. Besides, for all but one transition in the connecting case the situation is somewhat the opposite: v is on the positive orbit of u , and the avoidance set is the set of returns which are non-intermediate to all (not just one) of the transitions. See below for more comments about this. Thus what the transition from $F_L^{-1}(u)$ to v does is not a closing, but a shortcut. This is formulated in the following linear version of the C^1 connecting lemma, whose proof will be given in §5.

Theorem 2.2 (The linear version of the C^1 connecting lemma). *Given any sequence of isomorphisms $\{T_n\}$, and any $\varepsilon > 0$, there are $\sigma > 1$ and $L \in \mathbb{N}$ such that, for any two sequences $\{x_i\}_{i=1}^s$ and $\{y_i\}_{i=1}^t$ in V_0 , with an order $<$ defined on*

the union

$$X = \{x_i\}_{i=1}^s \cup \{y_i\}_{i=1}^t$$

as

$$x_1 < x_2 < \cdots < x_s < y_t < y_{t-1} < \cdots < y_1,$$

there exist two points

$$x \in \{x_i\}_{i=1}^s \cap B(x_s, \sigma|x_s - y_t|) \quad \text{and} \quad y \in \{y_i\}_{i=1}^t \cap B(x_s, \sigma|x_s - y_t|),$$

together with some ordered pairs $\{p_i, q_i\} \subset X \cap B(x_s, \sigma|x_s - y_t|)$, say, k of them, with the order

$$x_1 \leq p_1 \leq q_1 < p_2 \leq q_2 < \cdots < p_{k'} \leq q_{k'} < x \\ < y < p_{k'+1} \leq q_{k'+1} < \cdots < p_k \leq q_k \leq y_1$$

such that the following four conditions are satisfied.

1. There is an ε -kernel transition from x to y of length L that is contained in $B(x_s, \sigma|x_s - y_t|)$.
2. For each $i = 1, 2, \dots, k$, there is an ε -kernel transition from p_i to q_i of length L , contained in $B(x_s, \sigma|x_s - y_t|)$.
3. These $k+1$ transitions each avoid $X - [x, y] - [p_1, q_1] - \cdots - [p_k, q_k]$, where the closed interval notation $[x, y]$ denotes $\{z \in X \mid x \leq z \leq y\}$, etc.
4. These $k+1$ transitions are mutually disjoint.

The formulation of Theorem 2.2 is complicated, because it describes in detail how the connection is made. Note that the difference between the pair $\{x, y\}$ and the pairs $\{p_i, q_i\}$ is that x and y belong to different sequences, while p_i and q_i belong to the same sequence. While x and y are different points, p_i and q_i may be the same point. In this case the transition from p_i to q_i is understood as the trivial one, i.e. no lifts at all. For convenience we will call the pair $\{x, y\}$ the *connecting pair*, and $\{p_i, q_i\}$ the *cutting pairs*.

Remark. Let us give an informal illustration for Theorem 2.2. We visualize the sequence $\{x_i\}_{i=1}^s$ as returns of the positive orbit of p to a neighborhood of z and $\{y_i\}_{i=1}^t$ as returns of the negative orbit of q to the same neighborhood, where we think of p , q , and z as the three points in Theorem A. Note that while x_s and y_t can be arbitrarily close to z , the numbers σ and L are independent of $|x_s - y_t|$. Hence all the $k+1$ ε -kernel transitions can be put in an arbitrarily thin tube of length L . Roughly, the existence of such a σ guarantees control of the C^0 size of the perturbation. After perturbations, the positive orbit of p will go through the following points successively:

$$p, \cdots, x_1, \cdots, F_L^{-1}(p_1), \cdots, q_1, \cdots, F_L^{-1}(p_2), \cdots, q_2, \cdots, \\ F_L^{-1}(p_{k'}), \cdots, q_{k'}, \cdots, F_L^{-1}(x), \cdots, y, \cdots, \\ F_L^{-1}(p_{k'+1}), \cdots, q_{k'+1}, \cdots, F_L^{-1}(p_k), \cdots, q_k, \cdots, y_1, \cdots, q.$$

That is, it takes an old orbit from p to $F_L^{-1}(p_1)$, then takes a transition (a shortcut) from $F_L^{-1}(p_1)$ to q_1 , then takes an old orbit from q_1 to $F_L^{-1}(p_2)$, then takes a transition from $F_L^{-1}(p_2)$ to q_2 , etc. All the transitions here are shortcuts except for one: the transition associated with $\{x, y\}$ goes from the orbit of p to a *different*

orbit of q . This is how p and q get connected. We emphasize that, according to condition 3 in the theorem, the transition from p_i to q_i does not need to avoid the points of X that are intermediate to the other pairs (p_j, q_j) . Only those points of X that are non-intermediate to all of the pairs are to be avoided. In this case, as long as these transitions are disjoint, they make the connections.

3. A BASIC C^1 PERTURBATION THEOREM

In this section we formulate a basic C^1 perturbation theorem, which can be extracted from the work of Liao, Pugh, and Robinson on the C^1 closing lemma. This is Theorem 3.1 below, which serves as a fundamental preliminary for the proof of Theorem 2.2.

Let V be an m -dimensional inner product space and $e = (e_1, e_2, \dots, e_m)$ an orthonormal basis of V . An e -box Q of center $x \in V$ and of size $(\lambda_1, \lambda_2, \dots, \lambda_m)$ is defined as

$$Q = \{y \in V \mid |y_i - x_i| \leq \lambda_i, 1 \leq i \leq m\},$$

where x_i and y_i are coordinates of x and y , with respect to the basis e . For $\alpha > 0$, define

$$\alpha Q = \{y \in V \mid |y_i - x_i| \leq \alpha \lambda_i, 1 \leq i \leq m\}.$$

If $\alpha < 1$, we say that αQ is the α -kernel of Q . We say that a box Q' is of type Q , if

$$Q' = z + \alpha Q$$

for some $z \in V$ and some $\alpha > 0$.

Theorem 3.1. *For any sequence of isomorphisms $T_n: V_n \rightarrow V_{n-1}$, $n = 1, 2, \dots$, there is an orthonormal basis $e = (e_1, e_2, \dots, e_m)$ in V_0 such that for any $\varepsilon > 0$, and any $0 < \alpha < 1$, there exist an e -box A and an integer $L \in \mathbb{N}$ such that for any e -box Q of type A and any two points $x, y \in \alpha Q$, there is an ε -kernel transition $c_0, c_1, \dots, c_L; B_0, B_1, \dots, B_{L-1}$ of $\{T_n\}$ from x to y of length L , contained in Q . Moreover, the radius of B_n is less than or equal to half of the distance between $\partial(F_n^{-1}(Q))$ and $\partial(F_n^{-1}(\alpha Q))$.*

Note that Theorem 3.1 does not consider the avoidance of the transition. To actually prove the C^1 closing lemma using Theorem 3.1, one needs to carefully select a sequence of points in such a way that certain avoidance is realized. This requires some combinatorial considerations. For the connecting case some additional considerations are needed below on the disjointness of different transitions. In the statement of Theorem 3.1 the requirement that the support balls be uniformly small in ratio deserves a special attention. More precisely, in addition to the requirement that the ball B_n should be contained in $F_n^{-1}(Q)$, the last sentence of this theorem requires that the ball B_n should also be small enough relative to the parallelepiped $F_n^{-1}(Q)$ that, via a parallel translation, it can be inserted into the gap between the two parallelepipeds $F_n^{-1}(Q)$ and $F_n^{-1}(\alpha Q)$. This is crucial to the proof of the C^1 connecting lemma given below.

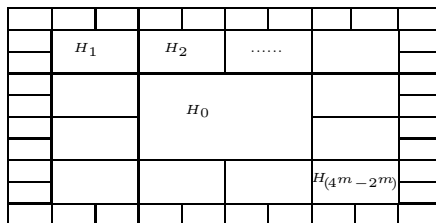
There is a beautiful proof for the C^1 closing lemma by Mai [4] with a different approach, which was later generalized to some non-invertible maps by Wen [14]. The proof is fairly simple. It does not yield Theorem 3.1, however, because the radii of the balls used there do not have to satisfy the last requirement of Theorem 3.1.

4. AN ARRANGEMENT OF BOXES

We describe in this section a simple arrangement of boxes in V_0 . The construction will be used in the next section to realize certain avoidance in the proof of Theorem 2.2.

Let e be an orthonormal basis of V_0 , and let A be an e -box. All boxes considered in this section will be e -boxes of the same type A .

Given a box H_0 of type A , there are $4^m - 2^m$ boxes H_i , $1 \leq i \leq 4^m - 2^m$, of the same type A , with size reduced by $1/2$, which fill out $2H_0 - H_0$ as the figure shows.

FIGURE 1. e -boxes

In other words, we can enclose H_0 with boxes of the same type A but of half-size to build up $2H_0$. Then we can enclose $2H_0$ with boxes ($10^m - 8^m$ of them) of the same type but of $1/4$ -size to build up $(2 + 1/2)H_0$. This process continues, and gives a sequence

$$H_0, H_1, H_2, \dots,$$

where H_1 through $H_{4^m - 2^m}$ (we may call them the boxes of the second generation) are those boxes of half-size, $H_{4^m - 2^m + 1}$ through $H_{10^m - 8^m}$ (we may call them the boxes of the third generation) are those boxes of $1/4$ -size, etc. The precise formulas for the subscripts like $4^m - 2^m$ are not essential to us, and will not be calculated explicitly below. It is clear that the union of all H_i is $\text{int}(3H_0)$. This open e -box will be somewhat important to us, because all ε -kernel transitions below will be contained in $\text{int}(3H_0)$ for a suitably chosen H_0 .

Now let

$$D_i = 2H_i, \quad i \geq 1.$$

The following three properties are clear.

D1) $\bigcup_{i \geq 0} D_i = \bigcup_{i \geq 0} H_i$.

D2) Each D_i is contained in $\text{int}(3H_0)$.

D3) There is a universal constant N^* , independent of e , A , and H_0 , such that each D_i intersects at most N^* of the other D_j .

Properties D1) and D2) are obvious. To see property D3), we first note that if H_i and H_j are three generations apart, then $D_i \cap D_j = \emptyset$. Then each D_i can intersect the other boxes D_j in at most five generations. For each of the five generations, it is easy to see that there is a constant N such that D_i intersects at most N of the other boxes D_j in this generation. Moreover, N can be chosen independent of D_i , and independent of e , A , and H_0 as well.

5. THE PROOF OF THEOREM 2.2

Now we prove Theorem 2.2.

Proof. Let $\{T_n\}$ and $\varepsilon > 0$ be given. Let e be the orthonormal basis given by Theorem 3.1. Let N^* be the universal constant in property D3) in §4. Let $0 < \alpha < 1$ be a number that satisfies the inequality

$$\left(\frac{1}{\alpha} - 1\right)\left(\frac{3}{\varepsilon}\right)^{2N^*+3} \leq 1.$$

For ε , α as chosen, take the e -box $A = (\lambda_1, \dots, \lambda_m)$ and the integer $L \in \mathbb{N}$ as Theorem 3.1 claims. Then let

$$\sigma = 6 \max\{\lambda_i/\lambda_j \mid 1 \leq i, j \leq m\},$$

that is, 6 times the *ballicity* of A . □

Now we verify that σ and L so chosen satisfy the conditions of Theorem 2.2. Given any two sequences $\{x_i\}_{i=1}^s$ and $\{y_i\}_{i=1}^t$ in V_0 , we need to select from X a connecting pair $\{x, y\}$ and some cutting pairs $\{p_i, q_i\}$ in $B(x_s, \sigma|x_s - y_t|)$ that satisfy the four conditions in Theorem 2.2. This will be done in the following three steps.

Step 1. A preliminary selection.

This step selects a pair $\{\xi, \zeta\}$, which is a candidate for $\{x, y\}$, and some pairs $\{u_i, v_i\}$, which are candidates for $\{p_i, q_i\}$. The final selection for $\{x, y\}$ and $\{p_i, q_i\}$ will be done in Step 3.

The selection of $\{\xi, \zeta\}$ and $\{u_i, v_i\}$ proceeds through a series of trial selections as follows.

Let H_0 be an e -box of type A such that $x_s, y_t \in H_0 \subset 3H_0 \subset B(x_s, \sigma|x_s - y_t|)$. Such an H_0 exists because of the choice of σ . This can be illustrated as follows. The ball $B(x_s, \sigma|x_s - y_t|)$ contains an e -cube C of size $\sigma|x_s - y_t|$ (in fact an e -cube of size $\sqrt{m}\sigma|x_s - y_t|$) centered at x_s . Shrinking with a suitable ratio on each side, one obtains in C an e -box of type A of center x_s with longest side $\sigma|x_s - y_t|$ and with shortest side $6|x_s - y_t|$. It is easy to check that this box could be our $3H_0$. Let $H_0, H_1, \dots$ be the sequence of boxes determined by H_0 , as arranged in §4. Everything we do below is in those D_n , and hence in $3H_0$ by D2), which in turn is contained in $B(x_s, \sigma|x_s - y_t|)$.

First we look at H_0 and choose ξ and ζ to be the smallest and the largest points of $X \cap H_0$. Here the terms smallest and largest are respecting the order $<$ of X . Thus $\xi \leq x_s < y_t \leq \zeta$ because $x_s, y_t \in H_0$. We emphasize that this selection is a trial one. It is subject to change.

Then we look at H_1 and let a and b be the smallest and the largest point of X within H_1 , subtracting the closed interval $[\xi, \zeta]$. In other words, a and b are the smallest and the largest points in $(X - [\xi, \zeta]) \cap H_1$. Note that a and b could be the same, and then $(X - [\xi, \zeta]) \cap H_1$ would reduce to a single point. This corresponds to the case that some cutting pairs $\{p_i, q_i\}$ are a single point and the corresponding transition is the trivial one. Also, $(X - [\xi, \zeta]) \cap H_1$ could be empty; in this case we simply go on to H_2 . There are two cases to consider.

Case 1. a and b belong to the same sequence. That is, $a \leq b < \xi < \zeta$, or $\xi < \zeta < a \leq b$.

In this case we make a trial selection of $\{a, b\}$ as one of the cutting pairs, say $\{u, v\}$. Note that $b \neq \xi$ and $\zeta \neq a$, because the closed interval $[\xi, \zeta]$ has been subtracted out from X .

Case 2. a and b belong to different sequences. That is, $a < \xi < \zeta < b$.

In this case we select $\{a, b\}$ as a better candidate for the connecting pair, and drop the open interval (a, b) , including the old candidates ξ and ζ , from our considerations, forever. Thus we rename a as ξ , and b as ζ . Note that this is still subject to change.

This finishes our observation in H_1 . We obtain a candidate connecting pair $\{\xi, \zeta\}$, and a (or no) candidate cutting pair $\{u, v\}$. We emphasize that the closed intervals determined by these pairs are mutually disjoint.

Then we look at H_2 . Let a and b be the smallest and the largest points in $(X - [\xi, \zeta] - [u, v]) \cap H_2$, where $\{u, v\}$ is the candidate cutting pair obtained in Case 1 above. For Case 2, we simply do not have this term. There are still two cases to consider.

Case 1. a and b belong to the same sequence.

In this case we select $\{a, b\}$ as a candidate cutting pair. Note that a and b do not belong to any of the closed intervals of the previously chosen pairs, since they have been subtracted out from X . Thus either $[a, b]$ is disjoint from all these intervals, or its interior (a, b) covers any of these intervals that intersect $[a, b]$. In the latter case, we drop (a, b) from our considerations. Thus the closed intervals of all pairs so obtained are mutually disjoint.

Case 2. a and b belong to different sequences.

In this case we select $\{a, b\}$ as a better candidate for $\{\xi, \zeta\}$, and drop (a, b) from our considerations.

Remark. This might be a good place to indicate Hayashi's brilliant idea of "cutting" [1]. As mentioned earlier, $\{x_i\}$ will be returns of the positive orbit of p to a neighborhood of z , and $\{y_i\}$ will be returns of the negative orbit of q to the same neighborhood of z , where p , q , and z are the three points in Theorem A. These returns are ordered in a way that fits our aim of connecting. Now a and b are two of the returns, and a is smaller than or equal to b . Whenever we can transit from a to b , or more precisely, from $F_L^{-1}(a)$ to b , directly via a transition, the old iterates after $F_L^{-1}(a)$ and before b (which could form a single orbital segment if $\{a, b\}$ is a cutting pair, or two orbital segments if $\{a, b\}$ is, a connecting pair), including in particular those returns between a and b in X , would be irrelevant to our aim of connecting. The farther a and b are apart in X , the better the transition would be, whatever $\{a, b\}$ is a connecting or cutting pair. This is why when a pair covers some other pairs in the selection process above, we simply drop those pairs from our considerations. This beautiful idea of Hayashi cutting will be used throughout the proof of Theorem 2.2.

This finishes our observation in H_2 , and we go on to H_3 , etc. After each stage, we obtain a unique candidate connecting pair, together with some candidate cutting pairs such that the closed intervals determined by these pairs are mutually disjoint. We may visualize the ordered set X as a line, and draw a bridge across each of these closed intervals, whether connecting type or cutting type, as Figure 2 shows.

Then the rule can be formulated as follows. Whenever a new bridge (the solid line in the figure) appears, its two end points must not be on any of the old closed bridges because they are subtracted before we choose a new pair, and we drop the

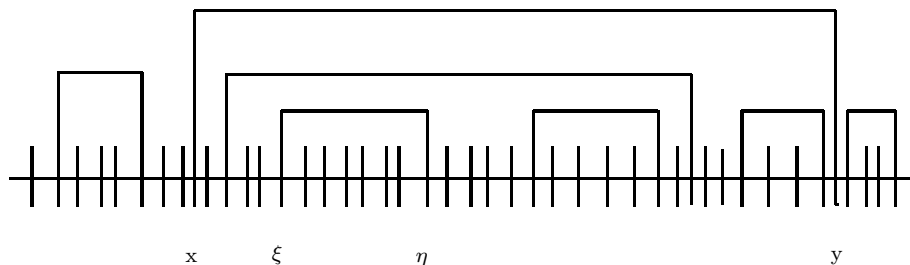


FIGURE 2. Connecting and cutting pairs

whole open interval down the new bridge, all the old bridges down the new bridge in particular, from our considerations. Note that the new bridge may have both end points chosen from a box H_i , but covers some points of X that are in some other boxes H_j , because the order in X does not reflect the location in V_0 . We must drop these points as well. This is important to keep bridges mutually disjoint.

This process terminates, because X is finite. This finishes Step 1 and gives us a connecting pair $\{\xi, \zeta\}$, which is a candidate for $\{x, y\}$, and some cutting pairs $\{u_i, v_i\}$, say, l of them, which are candidates for $\{p_i, q_i\}$, ordered as

$$x_1 \leq u_1 \leq v_1 < \cdots < u_{l'} \leq v_{l'} < \xi < \zeta < u_{l'+1} \leq v_{l'+1} < \cdots < u_l \leq v_l \leq y_1.$$

Note that the index i of $\{u_i, v_i\}$ is determined by the order $<$ of X , and not the order in which $\{u_i, v_i\}$ was chosen in the above selection process. Moreover, some boxes H_i may produce no new pairs. Thus the box from which $\{u_i, v_i\}$ was chosen may not be H_i at all. Let us denote it as H_i^* , and denote

$$D_i^* = 2H_i^*.$$

To keep the notations uniform, we use also $\{u_0, v_0\}$ to denote $\{\xi, \zeta\}$. The following four properties are clear.

- $D^*1)$ $u_i, v_i \in H_i^* \subset D_i^*$, $0 \leq i \leq l$.
- $D^*2)$ $D_i^* \subset B(x_s, \sigma|x_s - y_t|)$, $0 \leq i \leq l$.
- $D^*3)$ Each D_i^* intersects at most N^* of the other D_j^* .
- $D^*4)$ $(X - [\xi, \zeta] - [u_1, v_1] - \cdots - [u_l, v_l]) \cap \text{int}(3H_0) = \emptyset$.

Properties $D^*1)$, $D^*2)$, and $D^*3)$ are obvious by construction. Property $D^*4)$ holds because $\text{int}(3H_0)$ is the union of all H_i , $i = 0, 1, \dots$, and hence if the intersection were not empty, the selection process would still continue, and produce more new pairs.

Step 2. Basic balls and jumbo balls in V_0 .

We define the so-called jumbo balls in V_n for every n , which will be used to form the ε -kernel transitions required by Theorem 2.2. First we look at V_0 . The other V_n will be treated in the same way in Step 4 at the end of the proof of Theorem 2.2. Let

$$Q_i^* = \frac{1}{\alpha} H_i^*, \quad 0 \leq i \leq l,$$

where α is the number determined in the beginning of the proof of Theorem 2.2, which is less than but very close to one. Then

$$u_i, v_i \in \alpha Q_i^*, \quad 0 \leq i \leq l.$$

The box Q_i^* is between D_i^* and H_i^* . It is only a little bit larger than H_i^* . More precisely, $H_i^* = \alpha Q_i^*$. Since the type box A has been chosen according to ε and this α , Theorem 3.1 applies by treating Q_i^* as Q and H_i^* as αQ . That is, for each $i = 0, 1, \dots, l$, there is an ε -kernel transition

$$c_{i0}, c_{i1}, \dots, c_{iL}; \quad B_{i0}, B_{i1}, \dots, B_{i,L-1},$$

from u_i to v_i , contained in Q_i^* , where L is the number determined at the beginning of the proof of Theorem 2.2. Since the gap between Q_i^* and H_i^* is very narrow, the balls B_{in} are very small relative to $F_n^{-1}(D_i^*)$, and the precise ratio is given by the inequality that defines the number α , whose geometrical meaning will become clear shortly. This will be crucial to what follows. Recall that $u_0 = \xi$, $v_0 = \zeta$.

Note that the first two conditions in Theorem 2.2 are satisfied if we treat u_i , v_i as p_i , q_i . The third condition is also satisfied because of D^*4). Another condition in Theorem 2.2 about the locations of x , y , p_i and q_i is guaranteed by D^*2). The only problem now left is that the fourth condition in Theorem 2.2 is not satisfied. Clearly, these transitions are not necessarily disjoint. In fact they are not the transitions we want. The transitions we want will use the so-called jumbo balls defined below. The key to this will be the fact that the universal number N^* bounds all the multiplicities of overlaps of D_i^* .

We first define jumbo balls in V_0 . We will define jumbo balls in other V_n later. There are $l + 1$ balls

$$B_{00}, B_{10}, \dots, B_{l0}$$

in $3H_0 \subset V_0$. Let us call them *basic balls* in V_0 . Each basic ball B_{i0} contains two interesting points c_{i0} and $T_1(c_{i1})$ in its ε -kernel. Let us call these $l + 1$ pair of points $c_{i0}, T_1(c_{i1})$, $0 \leq i \leq l$, *basic points* in V_0 . Recall that $c_{i0} = v_i$.

Let b_i and r_i be the center and the radius of B_{i0} , respectively, $0 \leq i \leq l$. Consider the $2N^* + 3$ balls

$$B(b_i, (3/\varepsilon)^n r_i), \quad 1 \leq n \leq 2N^* + 3,$$

of the same center b_i . The largest of them is still in D_i^* , by Theorem 3.1 and the choice of α . Indeed, the inequality that defines α just means geometrically that the gap between H_i^* and $1/\alpha H_i^*$ should be so narrow that by enlarging $2N^* + 3$ times a ball B , which is contained in $1/\alpha H_i^*$ and is small enough that it can be inserted via a parallel translation into the gap, each time by a factor $3/\varepsilon$, we can never get out of $2H_i^*$. Then there is one of these $2N^* + 3$ balls, call it β_i , such that $\beta_i - (\frac{\varepsilon}{3})\beta_i$ does not contain any basic points. This is because each D_i^* contains at most $2N^* + 2$ basic points. In fact, each D_i^* intersects at most N^* of the other D_j^* by property D^*3), and $c_{j0}, T_1(c_{j1}) \in \varepsilon B_{j0} \subset Q_j^* \subset D_j^*$. Let

$$J_i = \beta_i/3.$$

We will call J_i the *prejumbo ball* in V_0 associated with B_{i0} . Thus each basic ball B_{i0} gives rise to a prejumbo ball J_i , $i = 0, 1, \dots, l$. Of course J_i is contained in D_i^* , and hence in $\text{int}(3H_0)$.

Claim 1. Two prejumbo balls either are disjoint, or the ε -kernel of one contains all the basic points contained in the other.

In fact, Let J_i and J_j be two prejumbo balls. Without loss of generality we assume that the radius of J_i is less than or equal to the radius of J_j . If $J_i \cap J_j \neq \emptyset$,

then $J_i \subset \beta_j$. So all basic points in J_i are in β_j , hence are in $\varepsilon J_j = \frac{\varepsilon}{3}\beta_j$, since $\beta_j - \frac{\varepsilon}{3}\beta_j$ contains no basic points. This proves the claim.

Let us call a collection $\mathcal{C}$ of prejumbo balls *regular*, if for every $i = 0, 1, \dots, l$ there is a prejumbo ball $J(i)$ in $\mathcal{C}$ such that the pair of basic points c_{i0} and $T_1(c_{i1})$ is contained in $\varepsilon J(i)$. For instance, the whole set of prejumbo balls $J_0, J_1, \dots, J_l$ is regular, because every such pair of basic points are contained in the ε -kernel of some basic ball, which in turn is contained in the ε -kernel of some prejumbo ball.

Claim 2. There is a regular collection of prejumbo balls

$$J_{i_1}, J_{i_2}, \dots, J_{i_d}$$

which are mutually disjoint.

In fact, if all prejumbo balls are mutually disjoint, we are done. Otherwise there is a prejumbo ball, say J_l , such that all basic points contained in J_l are contained in the ε -kernel of some other prejumbo ball, by Claim 1. Then we drop J_l from the collection. (Note that at this stage we do not drop any other prejumbo balls even if their basic points may be contained in J_l . For instance, J_l and J_{l-1} might contain each other's basic points, and do not intersect any other prejumbo balls. In this case we certainly do not want drop both of them.) Then the collection of the remaining prejumbo balls $J_0, J_1, \dots, J_{l-1}$ is still regular, and we go on to the next stage. In this way Claim 2 is proved by induction.

From now on we fix such a disjoint regular collection of prejumbo balls stated in Claim 2, and call them *jumbo balls* in V_0 . Clearly, each pair of basic points c_{i0} and $T_1(c_{i1})$ in V_0 is contained in the ε -kernel of a *unique* jumbo ball in V_0 . Note that all jumbo balls defined in V_0 are contained in $\text{int}(3H_0)$.

Step 3. A combining process.

We now combine the transitions whose basic points in V_0 are contained in the same jumbo ball in V_0 into a single ε -kernel transition, and further adjust the candidates for the connecting and cutting pairs. Since this process will appear later for other V_n as well, we first describe it in a more general way. Let $[a_1, b_1], [a_2, b_2], \dots, [a_k, b_k]$ be k disjoint closed intervals in X , ordered as

$$x_1 \leq a_1 \leq b_1 < a_2 \leq b_2 < \dots < a_k \leq b_k \leq y_1,$$

where X is the set in Theorem 2.2. Assume that for each $i = 1, \dots, k$ we have an ε -kernel transition

$$p_{i0}, p_{i1}, \dots, p_{iL}; \quad P_{i0}, P_{i1}, \dots, P_{i,L-1}$$

from a_i to b_i of length L . Here P_{in} could be any ball in V_n , not necessarily basic, nor jumbo (as mentioned before, we will define basic and jumbo balls in V_n for every n later). We allow such a generality here because later we will deal with transitions that use both basic balls and jumbo balls in a mixed way. After all, the balls used for transitions in Theorem 2.2 are not specified to be basic, nor jumbo.

Now let n be an integer with $0 \leq n \leq L-1$, and let J be a ball in V_n whose ε -kernel εJ contains the two points $T_{n+1}(p_{1,n+1})$ and $p_{k,n}$. We can combine these transitions into a single transition as follows. We first use the old ε -kernel lifts of index $i = 1$ from L up to $n+1$. When we get to the point $T_{n+1}(p_{1,n+1})$ in V_n , instead of jumping onto the point p_{1n} within P_{1n} , we jump onto the point p_{kn}

within J . Then we continue the rest of the old ε -kernel lifts of index $i = k$. That is, we make a new ε -kernel transition

$$p_{k0}, \dots, p_{k,n-1}, p_{kn}, p_{1,n+1}, \dots, p_{1,L}; \quad P_{k0}, \dots, P_{k,n-1}, J, P_{1,n+1}, \dots, P_{1,L-1}.$$

This new transition is from a_1 to b_k , or more precisely, from $F_L^{-1}(a_1)$ to b_k . Then Hayashi's cutting idea applies. That is, we make a new pair $\{a_1, b_k\}$, and drop everything in the open interval (a_1, b_k) . This includes in particular all the old pairs in (a_1, b_k) , together with all their transitions. We call this process *combining transitions and adjusting pairs via* $\{a_1, b_k\}$.

Now we do this process in V_0 . We start with the first jumbo ball J_{i_1} . Its ε -kernel εJ_{i_1} contains some basic points, but $J_{i_1} - \varepsilon J_{i_1}$ contains no basic points. Let s_1 and l_1 be the smallest and the largest index of basic points contained in εJ_{i_1} . By the regularity, basic points in εJ_{i_1} appear in pairs. In particular the two points $T_1(c_{s_1,1})$ and $c_{l_1,0}$ are in εJ_{i_1} . This gives a new ε -kernel lift that pushes $T_1(c_{s_1,1})$ onto $c_{l_1,0}$ within the jumbo ball J_{i_1} in V_0 . As described above, we make a new ε -kernel transition

$$c_{l_1,0}, c_{s_1,1}, c_{s_1,2}, \dots, c_{s_1,L}; \quad J_{i_1}, B_{s_1,1}, B_{s_1,2}, \dots, B_{s_1,L-1}.$$

This new transition is from u_{s_1} to v_{l_1} (more precisely, from $F_L^{-1}(u_{s_1})$ to v_{l_1}). It uses the original ε -kernel lifts in V_n , $n \geq 1$, but a new ε -kernel lift with a jumbo ball in V_0 . As described above, we simply take $\{u_{s_1}, v_{l_1}\}$ as a new candidate pair, and drop all pairs of index i with $s_1 \leq i \leq l_1$, *together with all their transitions*, from our further considerations. This is just the combining and adjusting process described above caused by a jumbo ball in V_0 . That is, we combine all the transitions associated with the pairs which are covered by the bridge of the new pair $\{u_{s_1}, v_{l_1}\}$, including the two transitions associated with the two pairs $\{u_{s_1}, v_{s_1}\}$ and $\{u_{l_1}, v_{l_1}\}$ themselves, into this single new transition. Note that in Step 1 we did not have transitions yet and what we dropped were pairs in V_0 , while now we drop the transitions. But this is not really a difference because, when we dropped a pair (a, b) in Step 1, we had ignored forever all possible transitions from a to b . Thus what are left now after this combination are still some transitions, one of which uses a jumbo ball at V_0 . Also note that, when we drop all pairs of index i with $s_1 \leq i \leq l_1$ together with their transitions, we may have dropped at the same time some transitions associated with (i.e. whose basic points in V_0 are contained in) some other jumbo balls rather than just with J_{i_1} , because the inequality $s_1 \leq i \leq l_1$ reflects the order in X but not the location in V_0 . This is similar to the situation we had in Step 1, and is important to keep bridges mutually disjoint. Finally, note that this new pair $\{u_{s_1}, v_{l_1}\}$ could be of the connecting type, and becomes a better candidate for $\{x, y\}$. Or, it could be of the cutting type, and becomes a better candidate for a $\{p_i, q_i\}$. In either case, we still have a unique candidate connecting pair and some candidate cutting pairs in V_0 with disjoint bridges, together with their transitions, one of which is the combined one. We remark that this is still subject to change.

Then we deal with the transitions that are left after the combination, and look at the second jumbo ball J_{i_2} in V_0 . In the same way, we combine the transitions whose basic points in V_0 are contained in J_{i_2} into a new transition. This gives a new pair, or what is the same, a new bridge in X , which either is disjoint from the old (here the word "old" means the ones that are just left after the last combination) bridges, or wholly covers some old bridges. Then we adjust the pairs in X as before.

We go on dealing with the transitions that are left and look at the third jumbo ball J_{i_3} , etc. In this way we will end up with a collection of disjoint bridges in X , together with their ε -kernel transitions of length L , each using one jumbo ball at V_0 . The number of transitions that are left could be equal to d , which is the number of jumbo balls in V_0 , or less than d , because when we deal with J_{i_1} , say, we may have dropped some transitions associated with some other jumbo balls as well, as noticed above. This finishes our work in V_0 .

Step 4. The other V_n with $n \geq 1$. The final selection of pairs.

Now we treat V_1 . By *basic ball* and *basic points* in V_1 we mean the balls B_{i1} and the points c_{i1} , $T_2(c_{i2})$, where the indices i are rearranged ones which run from one to however many transitions are left. These basic balls are contained in the parallelepiped $F_1^{-1}(3H_0)$. In the same way, we define prejumbo balls and jumbo balls in V_1 , then combine all the transitions (each has exactly one jumbo ball at V_0) that are associated with (i.e. whose basic points in V_1 are contained in) a jumbo ball in V_1 into a single new transition, and adjust the pairs in V_0 accordingly, as we did in V_0 . The only difference is that the boxes D_i^* , Q_i^* and H_i^* in V_0 become parallelepipeds $F_1^{-1}(D_i^*)$, $F_1^{-1}(Q_i^*)$, and $F_1^{-1}(H_i^*)$ in V_1 , respectively. But the parallelepipeds $F_1^{-1}(D_i^*)$ still overlap with multiplicity no more than N^* as the boxes D_i^* in V_0 do. Hence each $F_1^{-1}(D_i^*)$ contains at most $2N^* + 2$ basic points, since c_{j1} , $T_2(c_{j2}) \in F_1^{-1}(Q_j^*) \subset F_1^{-1}(D_j^*)$. Moreover, when we consider the $2N^* + 3$ concentric balls surrounding a basic ball B_{i1} , which is contained in $F_1^{-1}(Q_i^*)$ and is small enough so that it can be inserted into the gap between $F_1^{-1}(Q_i^*)$ and $F_1^{-1}(H_i^*)$ by Theorem 3.1, the largest one is still in $F_1^{-1}(D_i^*)$ by the choice of α . Thus all the arguments still go through, and we get a unique connecting pair and some cutting pairs in V_0 , together with their ε -kernel transitions of length L , each of which uses two jumbo balls at V_0 and V_1 .

Inductively, we treat V_2 , V_3 , $\dots$, V_L the same way. This eventually gives us a unique connecting pair x, y and some cutting pairs p_i, q_i in V_0 such that the balls used for ε -kernel transitions are all jumbo balls. All these transitions are contained in $3H_0$, and hence in $B(x_s, \sigma|x_s - y_t|)$. The four conditions of Theorem 2.2 are easily checked. In fact, the first two conditions are obvious. The third condition (the avoidance condition) is satisfied because, by the way those combinations are done in Step 3, the avoidance set $X - [x, y] - [p_1, q_1] - \dots - [p_k, q_k]$ is a subset of the avoidance set $X - [\xi, \zeta] - [u_1, v_1] - \dots - [u_l, v_l]$, which is outside $\text{int}(3H_0)$, and because all jumbo balls are contained in some pull-back of $\text{int}(3H_0)$. The fourth condition is satisfied because, being jumbo balls, those balls are mutually disjoint. Thus all conditions of Theorem 2.2 are satisfied, and Theorem 2.2 is proved.

6. AN ILLUSTRATION FOR THE LINEARIZATION PROCESS

Via a linearization process, Theorem E reduces to Theorem 2.2. This linearization process is rather standard, and the details can be found, for instance, in [14]. For convenience we give in this section a brief illustration of this process. Thus we are in a position of having the assumptions of Theorem E, and trying to prove Theorem E by using Theorem 2.2.

Let $\mathcal{U}$ be any C^1 neighborhood of f in $\text{Diff}^1(M)$. First we take a number $\eta > 0$ such that the η -ball of f in $\text{Diff}^1(M)$ is contained in $\mathcal{U}$. Then we take two numbers

$r > 0$ and $\varepsilon > 0$ such that hg is within the $\eta/2$ -ball of g for any g that is within the 1-ball of f in $\text{Diff}^1(M)$, where h is any ε -kernel lift supported on a ball of radius r as defined in Lemma 2.1.

Since z is not periodic of f , all terms in the negative orbit of z are hence distinct. Treating the tangent spaces $T_{f^{-n}(z)}M$ as V_n , and the tangent maps $T_{f^{-n}(z)}f$ as T_n , we get a sequence of isomorphisms $\{T_n\}$. Applying Theorem 2.2 to the sequence $\{T_n\}$ and the number ε determined above, we get two numbers $\sigma > 1$ and $L \in \mathbb{N}$. Since L is specified now, we can take a small number $\delta_0 > 0$ such that the ball $B(z, \delta_0)$ and its negative iterates $f^{-n}B(z, \delta_0)$, $0 \leq n \leq L$, are mutually disjoint, and are all of radius less than r . Actually δ_0 may have to be smaller to meet another requirement of some linearization process below. For $0 < a \leq \delta_0$, let us denote by $\Delta(a)$ the tube $\bigcup_{n=1}^L f^{-n}B(z, a)$.

Let $0 < \delta \leq \delta_0$ be given. As a preliminary C^1 perturbation we take a linearization of f , which is a diffeomorphism that agrees with f off the tube $\Delta(\delta)$, and agrees with the tangent maps $T_{f^{-n}(z)}f$ on the thinner tube $\Delta(\delta/2)$. Here we identify a neighborhood of the iterates $f^{-n}(z)$ in the manifold with a neighborhood of the origin in the tangent space $T_{f^{-n}(z)}M$, via the exponential map. Also, as noticed above, here we may assume that δ_0 has been chosen so small that the linearization along the δ -tube $\Delta(\delta)$ of length L is within the $\eta/2$ -ball of f . To keep the behavior of some orbits unchanged we need to cancel out the change brought in by this linearization. See [14] for details. For simplicity we still denote this linearization by f , and let $\rho = 2\sigma$.

Now let $0 \leq \delta \leq \delta_0$ be given, and let p and q be two points outside the tube $\Delta(\delta)$ such that the positive orbit of p hits the ball $B(z, \delta/\rho)$ after p , and the negative orbit of q hits the ball $B(z, \delta/\rho)$ too, at two points $f^a(p)$ and $f^{-b}(q)$, respectively. Collect the iterates of p from $f(p)$ up to $f^a(p)$ that are in the ball $B(z, \delta/2)$ as $x_1, x_2, \dots, x_s$, and the iterates of q from q up to $f^{-b}(q)$ that are in the same ball as $y_1, y_2, \dots, y_t$. Thus x_s and y_t are both in $B(z, (\delta/2)/\sigma)$. Clearly, the point $f^{-L}(x_1)$ is still on the positive orbit of p . (This is why we require that the positive orbit hits the ball $B(z, \delta/\rho)$ after p , and why we have collected the iterates of p from $f(p)$ to $f^a(p)$ but not from p to $f^a(p)$ that are in the ball $B(z, \delta/2)$.) Then a direct application of Theorem 2.2 gives Theorem E.

7. SOME APPLICATIONS OF THE C^1 CONNECTING LEMMA

In this section we are concerned with applications of the C^1 connecting lemma. First we note that it directly implies the celebrated C^1 closing lemma. In fact, the non-wandering point z in the assumption of the C^1 closing lemma can be assumed to be non-periodic, for otherwise it is already closed. Since z is non-wandering, there are a point $p \in B(z, \delta/\sigma)$ and an integer $n \geq 1$ such that $f^n(p) \in B(z, \delta/\sigma)$. Applying Theorem F by treating p itself as q yields a periodic orbit of g . Another perturbation takes this periodic orbit through z . This is because the proofs for the C^1 connecting lemma and the C^1 closing lemma both are in an essential way based on the basic C^1 perturbation theorem, Theorem 3.1, and the main difference is just that the former contains more complicated combinatorial considerations of avoidance discussed above.

Now we prove Theorems A–D. This is fairly straightforward. We only take Theorems B and D as examples, as they are more general than Theorems A and C, respectively.

Proof of Theorem B. Take $z_i \in \omega(p_i) \cap \alpha(p_{i+1})$ such that z_i is not periodic. We treat the case that the $2n$ orbits $Orb(p_i)$ and $Orb(z_i)$ are distinct and apply Theorem E. The other cases can be treated similarly. (For instance, if z_i and z_j , or p_i and p_j , share the same orbit, then p is actually a reduced order prolongationally recurrent. Also, if z_i and p_j share the same orbit, then the special connecting lemma, Theorem F, applies too.)

Given any C^1 neighborhood $\mathcal{U}$ of f , by Theorem E, for each z_i , there are three numbers ρ_i, L_i, δ_{0i} that satisfy the conditions of Theorem E. Denote $\rho = \max\{\rho_i\}$, $L = \max\{L_i\}$, and $\delta_0 = \min\{\delta_{0i}\}$. We may assume that δ_0 has been taken small enough so that the n tubes $\Delta_i(\delta_0) = \bigcup_{k=1}^L B(f^{-k}(z_i), \delta_0)$ are mutually disjoint, and that p_i and p_{i+1} are outside $\Delta_i(\delta_0)$ for all i . By the prolongational recurrence, for any $0 < \delta \leq \delta_0$, the positive orbit of p_i hits $B(z_i, \delta/\rho)$ after p_i , and the negative orbit of p_{i+1} hits $B(z_i, \delta/\rho)$ too.

Now for each $i = 1, \dots, n$ we choose suitable δ_i to get the desired tube $\Delta_i(\delta_i)$. For $i = 1$, we simply take δ_1 to be δ_0 itself. Then take two integers $a_1 \geq 1$ and $b_1 \geq 0$ so that $f^{a_1}(p_1)$ and $f^{-b_1}(p_2)$ are both in $B(z_1, \delta_1/\rho)$. Take $0 < \delta_2 \leq \delta_0$ small enough so that the tube $\Delta_2(\delta_2)$ is disjoint from the two finite orbits $\{p_1, f(p_1), \dots, f^{a_1}(p_1)\}$ and $\{p_2, f^{-1}(p_2), \dots, f^{-b_1}(p_2)\}$. Then take two integers $a_2 \geq 1$ and $b_2 \geq 0$ so that $f^{a_2}(p_2)$ and $f^{-b_2}(p_3)$ are both in $B(z_2, \delta_2/\rho)$. Then take $0 < \delta_3 \leq \delta_0$ small enough so that the tube $\Delta_3(\delta_3)$ is disjoint from the previously chosen four finite orbits, and so on. After these n tubes $\Delta_i(\delta_i)$ and the $2n$ finite orbits have been chosen, we start the connecting process as follows. First we apply Theorem E to the tube $\Delta_n(\delta_n)$ by treating z_n as z , p_n as p , and p_1 as q to make p_1 on the positive orbit of p_n under a new diffeomorphism g_1 . By construction this does not affect the other $2n-2$ finite orbits, one of which takes p_n backwards to hit the ball $B(z_{n-1}, \delta_{n-1}/\rho)$. Hence the negative g_1 -orbit of p_1 hits $B(z_{n-1}, \delta_{n-1}/\rho)$. Then we apply Theorem E to the tube $\Delta_{n-1}(\delta_{n-1})$ by treating z_{n-1} as z , p_{n-1} as p , and, still, p_1 as q , to make p_1 on the positive orbit of p_{n-1} for some g_2 . This does not affect the rest of the $2n-4$ finite orbits, and so on. This eventually makes p periodic. $\square$

Proof of Theorem D. We apply Theorem F, the special C^1 connecting lemma, twice as follows.

By assumption, there is a sequence of periodic orbits outside Λ that accumulate on Λ . Since Λ is isolated, there is a neighborhood W of Λ such that any periodic orbit which is not in Λ cannot be entirely in W . Then there must be $z^u \in W^u(\Lambda) - \Lambda$ and $z^s \in W^s(\Lambda) - \Lambda$ such that for any neighborhood U of z^u and any neighborhood V of z^s , there are a point $p \in U$ and an integer $n \in \mathbb{N}$ such that $f^n(p) \in V$. Note that z^u and z^s are not periodic points of f . Also, we may assume z^s is not on the positive orbit of z^u under f , as otherwise there is nothing to prove.

Let $\mathcal{U}$ be any C^1 neighborhood of f . By Theorem F, there are three numbers ρ^u, L^u , and δ_0^u for z^u that satisfy the conditions of Theorem F. Also, there are three numbers ρ^s, L^s , and δ_0^s for z^s that satisfy the conditions of Theorem F. Denote $\rho = \max(\rho^u, \rho^s)$, $L = \max(L^u, L^s)$, and $\delta_0 = \min(\delta_0^u, \delta_0^s)$. We may assume that δ_0 has been chosen small enough so that the tube $\Delta^s = \bigcup_{i=1}^L B(f^{-i}(z^s), \delta_0)$ is disjoint from the tube $\Delta^u = \bigcup_{i=1}^L B(f^i(z^u), \delta_0)$, and Δ^s is disjoint from $Orb^+(z^s, f)$ and Δ^u is disjoint from $Orb^-(z^u, f)$. Now there is a point $p \in B(z^u, \delta_0/\rho)$ and $n \geq 1$ such that $f^n(p) \in B(z^s, \delta_0/\rho)$. Applying Theorem F to Δ^s by treating z^s as $q = z$ makes z^s on the positive orbit of p for some g_1 . Let $f^k(p)$ be the first f -iterate of p that is in the tube Δ^s . We emphasize that while this perturbation does many cuttings

(and ε -kernel transitions too) to the finite orbit $\{f^k(p), f^{k+1}(p), \dots, f^n(p)\}$, it does not hurt the finite orbit $\{p, f(p), \dots, f^k(p)\}$. Thus p is on the negative orbit of z^s for g_1 . Then we apply Theorem F symmetrically, as mentioned in the remark after the statement of Theorem E in §1, to the tube Δ^u by treating z^u as $z = q$, z^s as p . This makes z^u on the negative orbit of z^s for some g , and proves Theorem D. $\square$

REFERENCES

- [1] S. Hayashi, *Connecting invariant manifolds and the solution of the C^1 -stability and Ω -stability conjectures for flows*, Ann. of Math. (2) **145** (1997), 81–137. MR **98b**:58096
- [2] S. T. Liao, *An extension of the C^1 closing lemma*, Acta Sci. Natur. Univ. Pekinensis, 1979, No. 2, 1-41. (Chinese) MR **81m**:58066
- [3] S. T. Liao, *On the perturbations of stable manifolds*, J. Sys. Sci. & Math. Scis., **8**(1988), 193-213 & 289-314. (Chinese) MR **90b**:58218
- [4] J. Mai, *A simpler proof of C^1 closing lemma*, Scientia Sinica, **10** (1986), 1021-1031. MR **88m**:58168
- [5] R. Mañé, *On the creation of homoclinic points*, Publ. Math. IHES, **66**(1988), 139-159. MR **89e**:58089
- [6] F. Oliveira, *On the generic existence of homoclinic points*, Ergod. Th. and Dynam. Sys., **7**(1987), 567-595. MR **89j**:58104
- [7] D. Pixton, *Planar homoclinic points*, J. Diff. Eq., **44**(1982), 365-382. MR **83h**:58077
- [8] C. Pugh, *The closing lemma*, Amer. J. Math., **89**(1967), 956-1009. MR **37**:2256
- [9] C. Pugh, *Against the C^2 closing lemma*, Journal of Differential Equations, **17**(1975), 435-443. MR **51**:4321
- [10] C. Pugh, *The C^1 connecting lemma*, Journal of Dynamics and Differential Equations, **4**(1992), 545-553. MR **93i**:58091
- [11] C. Pugh & C. Robinson, *The C^1 closing lemma, including Hamiltonians*, Ergod. Th. & Dynam. Sys., **3**(1983), 261-313. MR **85m**:58106
- [12] C. Robinson, *Closing stable and unstable manifolds on the two-sphere*, Proc. Amer. Math. Soc., **41**(1973), 299-303. MR **47**:9674
- [13] S. Smale, *Differentiable dynamical systems*, Bull. Amer. Math. Soc., **73**(1967), 747-817. MR **37**:3598
- [14] L. Wen, *The C^1 closing lemma for non-singular endomorphisms*, Ergod. Th. & Dynam. Sys., **11**(1991), 393-412. MR **92k**:58146
- [15] F. Takens, *Homoclinic points in conservative systems*, Invent. Math. **18**(1972), 267-292. MR **48**:9768
- [16] Z. Xia, *Homoclinic points in symplectic and volume-preserving diffeomorphisms*, Communications in Math. Phys. **177**(1996), 435-449. MR **97d**:58154

DEPARTMENT OF MATHEMATICS, PEKING UNIVERSITY, BEIJING, 100871, CHINA
E-mail address: lwen@math.pku.edu.cn

DEPARTMENT OF MATHEMATICS, NORTHWESTERN UNIVERSITY, EVANSTON, ILLINOIS 60208
E-mail address: xia@math.nwu.edu