

UNIVERSIDADE DE SÃO PAULO
Instituto de Ciências Matemáticas e de Computação

**A Necessary Condition for Formation of Internal
Transition Layer in a Reaction-Difusion System**

Arnaldo Simal do Nascimento
Janete Crema

Nº 88

NOTAS



São Carlos – SP
Abr./2000

SYSNO	1090428
DATA	/ /
ICMC - SBAB	

Resumo

Neste trabalho mostramos que uma condição de área generalizada é necessária para haver formação de camada limite interna em uma classe de sistemas elípticos singularmente perturbados, definidos em domínios variáveis ou não. Como aplicação conclui-se que tal condição é necessária para a existência de uma família de soluções estacionárias estáveis para um sistema de reação e difusão com condição de Neumann.

Palavras Chaves: Sistemas de reação e difusão, camada de transição interna, perturbações singulares.

A Necessary Condition for Formation of Internal Transition Layer in a Reaction-Difusion System

by

Arnaldo Simal do Nascimento⁽¹⁾ and *Janete Crema*⁽²⁾

⁽¹⁾ UFSCar - Dep. Matemática; C.P. 676; 13565-905 São Carlos, SP, Brasil

⁽²⁾ USP-ICMC; CP 668; 13560-970 São Carlos, SP, Brasil

Abstract

We prove that an equal-area like condition is necessary for the formation of internal transition layer in a singularly perturbed elliptic problem in a varying or fixed domain. As an application it is formally derived that such condition is also necessary for the existence of a family of stable stationary solutions to a reaction-diffusion system with Neumann boundary conditions on a fixed domain.

1

Key words. reaction-diffusion system, internal transition layer, singular perturbation

AMS subject classifications. 35B25, 35B35, 35K57, 35R35.

1 Introduction

In this paper we look for necessary conditions for development of internal transition layer for the following problem:

$$\begin{cases} \varepsilon \operatorname{div} (h_1(x) \nabla u_\varepsilon) + f(u_\varepsilon, v_\varepsilon) = 0, & x \in \Omega_\varepsilon, \\ \operatorname{div} (h_2(x) \nabla v_\varepsilon) + \delta(\varepsilon) g(u_\varepsilon, v_\varepsilon) = 0, & x \in \Omega_\varepsilon, \\ \frac{\partial v_\varepsilon}{\partial \hat{n}} = 0 \text{ (or } v_\varepsilon = 0), & \text{on } \partial\Omega_\varepsilon, \end{cases} \quad (1.1)$$

where ε is a small positive parameter, $\delta(\varepsilon)$ a real function continuous at $\varepsilon = 0$, f and g are functions in $C^1(\mathbb{R}^2)$ and $h_i \in C^{1,\alpha}(\Omega_\varepsilon)$, $0 < \alpha < 1$, satisfies $0 < m < h_i < M$ ($i = 1, 2$), for some constants m and M .

As for Ω_ε it is supposed to be a ε -family of uniformly bounded domains in $\mathbb{R}^N$, i.e., $|\Omega_\varepsilon| \leq M$, where M is a constant independent of ε and $|\cdot|$ denotes the N -dimensional Lebesgue measure. The boundary of Ω_ε , denoted by $\partial\Omega_\varepsilon$, is at least piecewise smooth.

The scenario depicted above aims at an application which will be presented in the last section and which motivated the present work. Roughly speaking, (1.1) is obtained from a reaction-diffusion system after performing a suitable scaling in the spatial variable and considering only its stationary solutions. If the original system has an ε -family of stationary patterns with a periodic structure then typically Ω_ε stands for a unit periodic cell.

Next we define what we mean by formation of internal transition layer.

Definition 1.1 *Let $\mathcal{O}$ be a open connected set in $\mathbb{R}^N$ and $\Gamma \subset \mathcal{O}$ a $(N - 1)$ -dimensional smooth compact manifold without boundary. We will say that a family of functions $\{(u_\varepsilon, v_\varepsilon)\}_{0 < \varepsilon \leq \varepsilon_0}$, with u_ε and v_ε in $C^1(\overline{\mathcal{O}}) \cap C^2(\mathcal{O})$ develops internal transition layer, as $\varepsilon \rightarrow 0$, in $\mathcal{O}$ with interface $\Gamma \subset \mathcal{O}$ and limit $(u_0(x), v_0(x))$ if:*

- $u_\varepsilon \xrightarrow{\varepsilon \rightarrow 0} u_0$, uniformly on any compact set $K \subset \{\mathcal{O} \setminus \Gamma\}$, where u_0 is given by

$$u_0(x) = \alpha(x)\chi_{\mathcal{O}_\alpha}(x) + \beta(x)\chi_{\mathcal{O}_\beta}(x) \quad (1.2)$$

for some functions α and β in $C(\mathcal{O})$ with $\alpha(x) < \beta(x)$ for $x \in \Gamma$, $\mathcal{O}_\alpha$ is the open region in $\mathcal{O}$ enclosed by Γ and $\mathcal{O}_\beta = \mathcal{O} \setminus \{\Gamma \cup \mathcal{O}_\alpha\}$

- $v_\varepsilon \xrightarrow{\varepsilon \rightarrow 0} v_0$ uniformly on any compact set $K \subset \mathcal{O}$.

Note that $\mathcal{O} = \mathcal{O}_\alpha \cup \Gamma \cup \mathcal{O}_\beta$. Also $\mathcal{O}$ can be either a narrow tubular neighborhood of Γ or a simply connected domain. In the latter case $\partial\mathcal{O}_\alpha = \Gamma$.

The above definition, which is local in space, suffices for our purposes and it allows for the existence of more layers.

2 Main Result

With the above definitions and notations we now state our main result in a slightly more general framework than we need for the application we have in mind.

Theorem 2.1 *Let $\{(u_\varepsilon, v_\varepsilon)\}_{\varepsilon>0}$ be a family of solutions to problem (1.1) which develops internal transition layer, as $\varepsilon \rightarrow 0$, in $\mathcal{U}$ with interface Γ and limit $(u_0(x), v_0(x))$ satisfying $f(u_0, v_0) = 0$, a.e. in $\mathcal{U}$, where $\mathcal{U}$ is an open tubular neighborhood of Γ such that $\mathcal{U} \subset \Omega_\varepsilon$, $0 < \varepsilon \leq \varepsilon_0$, for ε_0 small enough.*

Moreover suppose that the following hypotheses are satisfied.

$$(H.1) \quad \|u_\varepsilon\|_{L^\infty(\Omega_\varepsilon)} \leq C, \quad \|v_\varepsilon\|_{L^\infty(\Omega_\varepsilon)} \leq C, \quad \text{for } 0 < \varepsilon \leq \varepsilon_0 \text{ and some constant } C > 0,$$

$$(H.2) \quad \lim_{\varepsilon \rightarrow 0} \varepsilon \int_{\partial \mathcal{U}} |\nabla u_\varepsilon(x)|^2 dS = 0.$$

Then the N -vector equation

$$\int_{\Gamma} \int_{\alpha(x)}^{\beta(x)} f(s, v_0(x)) ds \quad \hat{\nu}(x) dS = 0 \quad (2.1)$$

holds. Here $\hat{\nu}$ stands for the outward unit normal vector on Γ .

Remark 2.1 *At the end of this section we will see that in some important cases which appear in the applications, (H.2) is automatically satisfied.*

Proof: By taking a narrower tubular neighborhood $\tilde{\mathcal{U}}$ of Γ , i.e., $\tilde{\mathcal{U}} \subset \mathcal{U}$, it follows from Definition 1.1 that u_ε and v_ε converge uniformly to u_0 and v_0 , respectively, on $\partial \tilde{\mathcal{U}}$. For simplicity we drop $\sim$ and write just $\mathcal{U}$. As in the Pohozaev procedure, multiplying the first equation in (1.1) by $x \cdot \nabla u_\varepsilon$ and integrating over $\mathcal{U}$ we obtain

$$\int_{\mathcal{U}} \{ \varepsilon \operatorname{div}(h_1(x) \nabla u_\varepsilon) (x \cdot \nabla u_\varepsilon) + f(u_\varepsilon, v_\varepsilon) x \cdot \nabla u_\varepsilon \} dx = 0 \quad (2.2)$$

Working with the first term of this equality and using that

$$\operatorname{div}[x \cdot \nabla u \, h_1 \nabla u] = x \cdot \nabla u \operatorname{div}(h_1 \nabla u) + h_1 [|\nabla u|^2 + x \cdot \nabla(|\nabla u|^2)/2]$$

along with the Divergence theorem, it follows that

$$\begin{aligned} & - \varepsilon \int_{\partial \mathcal{U}} h_1(x) \, x \cdot \nabla u_\varepsilon \, \frac{\partial u_\varepsilon}{\partial \hat{n}} dS - \frac{\varepsilon}{2} \int_{\partial \mathcal{U}} h_1(x) |\nabla u_\varepsilon|^2 x \cdot \hat{n} dS \\ & - \frac{\varepsilon}{2} \int_{\mathcal{U}} |\nabla u_\varepsilon|^2 x \cdot \nabla h_1 dx + \varepsilon(2 - N)/2 \int_{\mathcal{U}} h_1(x) |\nabla u_\varepsilon|^2 dx \\ & = \int_{\mathcal{U}} f(u_\varepsilon, v_\varepsilon) x \cdot \nabla u_\varepsilon dx. \end{aligned} \quad (2.3)$$

We claim that the left hand side of this equality goes to 0, as $\varepsilon \rightarrow 0$. In fact, this holds for the first and second terms by virtue of (H.2). We claim that the third and fourth terms approach zero, as $\varepsilon \rightarrow 0$, too. It suffices to show that

$$\lim_{\varepsilon \rightarrow 0} \varepsilon \int_{\mathcal{U}} |\nabla u_\varepsilon(x)|^2 dx = 0. \quad (2.4)$$

But multiplying the first equation in (1.1) by u_ε and integrating over $\mathcal{U}$ we obtain

$$\varepsilon \int_{\mathcal{U}} h_1 |\nabla u_\varepsilon|^2 dx = \varepsilon \int_{\partial \mathcal{U}} h_1 \frac{\partial u_\varepsilon}{\partial \nu} dS + \int_{\mathcal{U}} f(u_\varepsilon, v_\varepsilon) u_\varepsilon dx \xrightarrow{\varepsilon \rightarrow 0} 0, \quad (2.5)$$

by (H.1), (H.2) and the fact that $f(u_\varepsilon, v_\varepsilon) \rightarrow f(u_0, v_0) = 0$, a.e. in $\mathcal{U}$. Hence we have (2.4) and going back to (2.3) we conclude

$$\lim_{\varepsilon \rightarrow 0} \int_{\mathcal{U}} f(u_\varepsilon, v_\varepsilon) x \cdot \nabla u_\varepsilon dx = 0. \quad (2.6)$$

Note that

$$\begin{aligned} \operatorname{div} [x \int_{\theta}^u f(s, v) ds] &= (\nabla \int_{\theta}^u f(s, v) ds) \cdot x + \operatorname{div} x \cdot \int_{\theta}^u f(s, v) ds \\ &= f(u, v) \nabla u \cdot x + \int_{\theta}^u x \cdot \nabla_x f(s, v) ds + N \int_{\theta}^u f(s, v) ds \end{aligned}$$

Therefore

$$\begin{aligned} \int_{\mathcal{U}} f(u_\varepsilon, v_\varepsilon) x \cdot \nabla u_\varepsilon dx &= \int_{\mathcal{U}} \operatorname{div} [x \int_{\theta}^{u_\varepsilon} f(s, v_\varepsilon) ds] dx - \int_{\mathcal{U}} \left\{ \int_{\theta}^{u_\varepsilon} x \cdot \nabla_x f(s, v_\varepsilon) ds + N \int_{\theta}^{u_\varepsilon} f(s, v_\varepsilon) ds \right\} dx \\ &= \int_{\partial \mathcal{U}} x \cdot \hat{n} \int_{\theta}^{u_\varepsilon} f(s, v_\varepsilon) ds dS - \int_{\mathcal{U}} \left\{ \int_{\theta}^{u_\varepsilon} x \cdot \nabla_x f(s, v_\varepsilon) ds + N \int_{\theta}^{u_\varepsilon} f(s, v_\varepsilon) ds \right\} dx \end{aligned} \quad (2.7)$$

Recall that by hypothesis, $\{u_\varepsilon\}$ and $\{v_\varepsilon\}$ converge uniformly on $\partial \mathcal{U}$, so that the first integral of the last term in (2.13) satisfies

$$\int_{\partial \mathcal{U}} x \cdot \hat{n} \int_{\theta}^{u_\varepsilon} f(s, v_\varepsilon) ds dS \xrightarrow{\varepsilon \rightarrow 0} \int_{\partial \mathcal{U}} x \cdot \hat{n} \int_{\theta}^{u_0} f(s, v_0) ds dS \quad (2.8)$$

Also it is easy to see that

$$\begin{aligned} \int_{\mathcal{U}_\alpha} \int_{\theta}^{u_\varepsilon} f(s, v_\varepsilon) ds dx &\xrightarrow{\varepsilon \rightarrow 0} \int_{\mathcal{U}_\alpha} \int_{\theta}^{\alpha} f(s, v_0) ds dx \\ \int_{\mathcal{U}_\beta} \int_{\theta}^{u_\varepsilon} f(s, v_\varepsilon) ds dx &\xrightarrow{\varepsilon \rightarrow 0} \int_{\mathcal{U}_\beta} \int_{\theta}^{\beta} f(s, v_0) ds dx \end{aligned} \quad (2.9)$$

where $\mathcal{U}_\alpha$ and $\mathcal{U}_\beta$ are as in Definition 1.1.

Finally note that $\nabla_x f(s, v) = \partial_2 f(s, v) \nabla v(x)$ (here ∂_2 denotes the partial derivative with respect to the second argument) and $\partial_2 f$ is uniformly continuous in compact sets. Hence by (H.1),

$$x \cdot \int_{\theta}^{u_{\varepsilon}} \partial_2 f(s, v_{\varepsilon}(x)) ds \rightarrow x \cdot \int_{\theta}^{u_0} \partial_2 f(s, v_0(x)) ds \quad (2.10)$$

strongly in $L^2(\mathcal{U})$. Then if ∇v_{ε} converges weakly in $L^2(\mathcal{U})$ we have

$$\begin{aligned} \int_{\mathcal{U}} \int_{\theta}^{u_{\varepsilon}} x \cdot \nabla_x f(s, v_{\varepsilon}) ds dx &= \\ \int_{\mathcal{U}} x \cdot \nabla v_{\varepsilon}(x) \int_{\theta}^{u_{\varepsilon}} \partial_2 f(s, v_{\varepsilon}(x)) ds dx &\xrightarrow{\varepsilon \rightarrow 0} \int_{\mathcal{U}} x \cdot \nabla v_0(x) \int_{\theta}^{u_0} \partial_2 f(s, v_0(x)) ds dx = \\ \int_{\mathcal{U}} \int_{\theta}^{u_0} x \cdot \nabla_x f(s, v_0) ds dx. \end{aligned} \quad (2.11)$$

In order to establish the weak convergence of $\{v_{\varepsilon}\}$ in $H^1(\mathcal{U})$ note that multiplying the second equation in (1.1) by v_{ε} and integrating on Ω_{ε} we conclude that

$$\int_{\Omega_{\varepsilon}} h_2(x) |\nabla v_{\varepsilon}|^2 dx = \int_{\Omega_{\varepsilon}} \delta(\varepsilon) g(u_{\varepsilon}, v_{\varepsilon}) v_{\varepsilon} dx$$

By our hypotheses and since $\delta(\varepsilon)$ is continuous at $\varepsilon = 0$, the righthand term in the above equality is uniformly bounded in ε , and therefore so is $\|\nabla v_{\varepsilon}\|_{L^2(\Omega_{\varepsilon})}$. Moreover $v_{\varepsilon} \rightarrow v_0$ uniformly in compact sets of $\mathcal{U}$ hence v_{ε} is bounded in $H^1(\mathcal{U})$. It follows that $v_{\varepsilon} \rightharpoonup v_0$ in $H^1(\mathcal{U})$ and so $\nabla v_{\varepsilon} \rightharpoonup \nabla v_0$ in $L^2(\mathcal{U})$, as we wanted.

Passing to the limit in (2.7), as $\varepsilon \rightarrow 0$, and using (2.6), (2.8), (2.9) and (2.11) we obtain

$$\begin{aligned} 0 &= \int_{\partial \mathcal{U}} x \cdot \hat{n} \int_{\theta}^{u_0(x)} f(s, v_0(x)) ds dS \\ &- \int_{\mathcal{U}_{\alpha}} \left\{ \int_{\theta}^{\alpha(x)} x \cdot \nabla_x f(s, v_0) + N f(s, v_0) ds \right\} dx \\ &- \int_{\mathcal{U}_{\beta}} \left\{ \int_{\theta}^{\beta(x)} x \cdot \nabla_x f(s, v_0) + N f(s, v_0) ds \right\} dx. \end{aligned}$$

We will rewrite this expression remarking that if $F(\beta(x)) = \int_{\theta}^{\beta(x)} f(s, v_0(x)) ds$ then

$$\operatorname{div}(x F(\beta(x))) = N F(\beta(x)) + x \cdot \nabla_x F(\beta(x)) =$$

$$N F(\beta(x)) + x \cdot [f(\beta(x), v_0(x)) \cdot \nabla \beta(x) + \int_{\theta}^{\beta(x)} \nabla_x f(s, v_0(x)) ds],$$

with an analogous formulation for α . So we obtain

$$0 = \int_{\partial \mathcal{U}} x \cdot \hat{n} F(u_0(x)) dS - \int_{\mathcal{U}_\alpha} \{ \operatorname{div}(x F(\alpha(x)) - f(\alpha, v_0) \nabla \alpha \cdot x) dx - \\ \int_{\mathcal{U}_\beta} \{ \operatorname{div}(x F(\beta(x)) - f(\beta, v_0) \nabla \beta \cdot x) dx.$$

Since $f(\alpha, v_0) = f(\beta, v_0) = 0$ the Divergence theorem implies

$$\int_{\Gamma} \left\{ \int_{\alpha(x)}^{\beta(x)} f(s, v_0(x)) ds \right\} x \cdot \hat{\nu} dS = 0. \quad (2.12)$$

In order to obtain (2.1) it suffices to follow above procedure with $x \cdot \nabla v$ replaced with $(x + p) \cdot \nabla v$, with p being an arbitrary N -vector. We end up with the following equation

$$\int_{\Gamma} \left\{ \int_{\alpha(x)}^{\beta(x)} f(s, v_0(x)) ds \right\} (x + p) \cdot \hat{\nu} dS = 0.$$

Since p is an arbitrary vector, (2.1) is proved. ■

Remark 2.2 Condition (H.2) is satisfied in $\partial \mathcal{U}$ if $\{\|\nabla u_\varepsilon\|\}_{0 < \varepsilon \leq \varepsilon_0}$ can be controlled there. For example, if $\|\nabla u_\varepsilon(x)\| = o(\varepsilon^{-1/2})$, uniformly in $0 < \varepsilon \leq \varepsilon_0$, in some neighborhood of $\partial \mathcal{U}$, then (H.2) holds.

Remark 2.3 In [3], instead of (1.1) the following scalar problem was considered

$$\begin{cases} \varepsilon \operatorname{div}(h_1(x) \nabla v_\varepsilon) + f(x, v_\varepsilon) = 0 \\ \frac{\partial v_\varepsilon}{\partial \hat{n}} = 0, \text{ for } x \in \partial \Omega \end{cases} \quad (2.13)$$

where $\Omega \subset \mathbb{R}^N$. It was assumed that $f(x, \alpha) = f(x, \beta) = 0$, $\forall x \in \Omega$, for two constants α and β , $\alpha < \beta$.

Under these conditions it was proved that if the above problem possesses a ε -family of solutions which is bounded in Ω , uniformly in ε and develops an internal transition layer with interface $\Gamma \subset \Omega$ then

$$\lim_{\varepsilon \rightarrow 0} \varepsilon |\nabla u_\varepsilon(x)|^2 = 0, \text{ uniformly in } \partial \Omega.$$

It turns out that the proof about convergence of the gradient on the boundary provided in [3] can be of assistance here. Following the procedure set forth therein it can be shown

mutatis mutandis that if in the hypotheses of Theorem 2.1, $\alpha(x)$ and $\beta(x)$ are constant functions, then

$$\lim_{\varepsilon \rightarrow 0} \varepsilon |\nabla u_\varepsilon(x)|^2 = 0, \text{ uniformly in } \partial\mathcal{U}.$$

Therefore (H.2) is automatically satisfied.

The key point when adapting the proof provided in [3] to the case presented here is the observation that in [3] the above uniform convergence on $\partial\Omega$ was obtained using a blow-up technique and some Schauder-like estimates. This in turn was possible because the functions α and β were constant, u_ε was bounded in Ω uniformly in ε and u_ε converged to u_0 on $\partial\Omega$ uniformly. In the present case the last convergence on the boundary $\partial\mathcal{U}$ can be obtained, in view of Definition 1.1, through an argument like the one used in the beginning of the proof of Theorem 2.1.

Consequently, when α and β are constant, Theorem 2.1 can be restated without the awkward condition (H.2), as follows.

Corollary 2.1 *Let $\{(u_\varepsilon, v_\varepsilon)\}_{\varepsilon>0}$ be a family of solutions to problem (1.1) which develops internal transition layer, as $\varepsilon \rightarrow 0$, in $\mathcal{U}$ with interface Γ and limit $(u_0(x), v_0(x))$ satisfying $f(u_0, v_0) = 0$, a.e. in $\mathcal{U}$, where $\mathcal{U}$ is an open tubular neighborhood of Γ such that $\mathcal{U} \subset \Omega_\varepsilon$, $0 < \varepsilon \leq \varepsilon_0$, for ε_0 small enough.*

Moreover suppose that (H.1) is satisfied and that α and β are constant functions.

Then

$$\int_{\Gamma} \int_{\alpha}^{\beta} f(s, v_0(x)) ds \, \hat{\nu}(x) \, dS = 0.$$

Working towards the application that will be presented in the last section, we now state some consequences of the above result which can be obtained under some additional hypotheses.

Corollary 2.2 *Let $\{(u_\varepsilon, v_\varepsilon)\}_{\varepsilon>0}$ be a family of solutions to problem (1.1) which develops internal transition layer, as $\varepsilon \rightarrow 0$, in $\mathcal{U}$ with interface Γ and limit $(u_0(x), v_0(x))$ satisfying $f(u_0, v_0) = 0$, a.e. in $\mathcal{U}$, where $\mathcal{U}$ is an open tubular neighborhood of Γ such that $\mathcal{U} \subset \Omega_\varepsilon$, $0 < \varepsilon \leq \varepsilon_0$. Suppose that (H.1) is satisfied and*

- $\delta(\varepsilon) \rightarrow 0$, as $\varepsilon \rightarrow 0$.

Then v_0 is a constant function.

If in addition the nullcline of f ($\stackrel{\text{def}}{=} \{(u, v) : f(u, v) = 0\}$) intersects $\{(s, v_0) : \alpha(x) \leq s \leq \beta(x), x \in \mathcal{U}\}$ in a discrete set, then α and β are constant functions and

$$\int_{\alpha}^{\beta} f(s, v_0) ds = 0. \quad (2.14)$$

Proof: Multiplying the second equation in (1.1) by v_{ε} and integrating over Ω_{ε} , it yields

$$\int_{\Omega_{\varepsilon}} \{-h_2(x)|\nabla v_{\varepsilon}(x)|^2 + \delta(\varepsilon)g(u_{\varepsilon}, v_{\varepsilon})v_{\varepsilon}(x)\} dx = 0.$$

Our hypotheses imply right away that

$$\int_{\Omega_{\varepsilon}} |\nabla v_{\varepsilon}(x)|^2 dx \rightarrow 0$$

This by its turn implies that $v_{\varepsilon} \rightarrow v_0$ in $H^1(\mathcal{U})$ and so v_0 is a constant function a.e. in $\mathcal{U}$.

Recall that $u_{\varepsilon} \rightarrow u_0$ where u_0 is given by (1.2). By hypothesis $\{s; f(s, v_0) = 0\}$ is a discrete set. But α and β are continuous functions and therefore are constant on $\mathcal{U}$. Thus by (2.1) we will have

$$\int_{\Gamma} \int_{\alpha}^{\beta} \{f(s, v_0) ds\} x \cdot \hat{n} dS = 0$$

But $\int_{\Gamma} x \cdot \hat{n} dS = \int_{\mathcal{O}_{\Gamma}} \text{div } x dx = N |\mathcal{O}_{\Gamma}| \neq 0$, where $\mathcal{O}_{\Gamma}$ stands for the open region enclosed by Γ . Hence $\int_{\alpha}^{\beta} f(s, v_0) ds = 0$ ■

Remark 2.4 Equation (2.14) is sometimes referred to as the v_0 -level equal-area condition for $f(\cdot, v)$ with its obvious counterpart for scalar equations. It appears in several works which give sufficient conditions for the existence of internal transition layer either in systems or in scalar equations. See [1], [2], [4], [5], [8] and [11], for instance. In [3] it is shown that a generalization of the equal-area condition is actually a necessary condition for the formation of internal transition layer for solutions of (2.13). In particular, when f does not depend on the spatial variable we recover the equal-area condition $\int_{\alpha}^{\beta} f(s) ds = 0$. Thus when this condition does not hold (for scalar equations) we should have formation of boundary and/or spiky layers. This problem has been addressed in [6] and [7].

The present work seems to be the first attempt to investigate the necessity of such condition for systems to develop internal transition layers.

3 Application

In [9], Nishiura and Suzuki considered the following model system of reaction-diffusion

$$\begin{cases} u_t = \varepsilon^2 \Delta u + f(u, v), \\ \nu v_t = D \Delta v + g(u, v), \quad (x, t) \in \Omega \times (0, \infty) \\ \frac{\partial u}{\partial \hat{n}} = \frac{\partial v}{\partial \hat{n}} = 0, \quad (x, t) \in \partial \Omega \times (0, \infty) \end{cases} \quad (3.1)$$

where u is the activator, v is the inhibitor, Ω is a smooth domain in $\mathbb{R}^N$, $\nu > 0$, $D > 0$ and ε a small positive parameter. The nullcline of f is sigmoidal and consists of three smooth curves $u = h_-(v)$, $u = h_0(v)$ and $u = h_+(v)$ defined on the intervals I_- , I_0 and I_+ , respectively. Also if $\min I_- = \underline{v}$ and $\max I_+ = \bar{v}$ then the inequality $h_-(v) < h_0(v) < h_+(v)$ holds for $I^* = (\underline{v}, \bar{v})$ and $h_+(v)$ (resp., $h_-(v)$) coincides with $h_0(v)$ at only one point $v = \bar{v}$ (resp., $v = \underline{v}$), respectively.

Therein, (3.1) is assumed to satisfy a set of hypotheses, which we denote by $\mathcal{H}$, among which we mention the following one concerning f :

$$(Q.1) \quad J(v) = \int_{h_-(v)}^{h_+(v)} f(\xi, v) d\xi \text{ has one isolated zero at } v = v^* \in I^* \text{ such that } \frac{dJ}{dv} < 0 \text{ at } v = v^*.$$

Under the hypotheses $\mathcal{H}$, they prove that if (3.1) has an ε -family $(U_\varepsilon, V_\varepsilon)$ of stationary matched solutions of order 1 whose α^* -level surface S_ε of U_ε is smooth up to $\varepsilon = 0$, then it must become unstable for ε small.

Here α^* an intermediate value between the two stable branches of the nullcline of f . Also by “smooth up to $\varepsilon = 0$ ” it is meant that there exists an $(N - 1)$ -dimensional smooth compact connected manifold S_0 without boundary in $\mathbb{R}^N$ such that $S_\varepsilon \xrightarrow{\varepsilon \rightarrow 0} S_0$.

Therefore the question arises as to what happens with the configuration of a ε -family $(u_\varepsilon, v_\varepsilon)$ of matched stable stationary solutions to (3.1). As said in [9], their result suggests that stable patterns somehow must become very fine and/or complicated as $\varepsilon \rightarrow 0$.

Actually based on the balance of the bulk force and the mean-curvature effect they formally derive that the rate of shrinking of stable patterns is of order $\varepsilon^{1/3}$. See [9], p. 1103. Hence by a suitable scaling, the resulting rescaled equations capture the

morphology of the magnified patterns. This has been accomplished in [10], where after performing the change of variable $X = \frac{x-x^*}{\varepsilon^{1/3}}$ for a suitable $x^* \in R^N$, the rescaled system becomes

$$\begin{cases} \tilde{u}_t = \tilde{\varepsilon}^2 \Delta \tilde{u} + f(\tilde{u}, \tilde{v}), \\ \nu \tilde{\varepsilon} \tilde{v}_t = D \Delta \tilde{v} + \tilde{\varepsilon} g(\tilde{u}, \tilde{v}), \quad (X, t) \in \Omega_\varepsilon \times (0, \infty) \\ \frac{\partial \tilde{u}}{\partial \tilde{n}} = \frac{\partial \tilde{v}}{\partial \tilde{n}} = 0, \quad (X, t) \in \partial \Omega_\varepsilon \times (0, \infty) \end{cases} \quad (3.2)$$

where $\tilde{\varepsilon} = \varepsilon^{2/3}$ and the rescaled domain Ω_ε satisfies $\Omega_\varepsilon \xrightarrow{\varepsilon \rightarrow 0} \tilde{\Omega}$ with $\tilde{\Omega} < \infty$.

It follows from the above argument that if $\Gamma_\varepsilon \stackrel{\text{def}}{=} \{X \in \Omega_\varepsilon : \tilde{u}_\varepsilon = \alpha^*\}$ then there exists an $(N-1)$ -dimensional smooth compact connected manifold Γ_0 without boundary such that $\Gamma_\varepsilon \xrightarrow{\varepsilon \rightarrow 0} \Gamma_0$.

The convergence $\Omega_\varepsilon \xrightarrow{\varepsilon \rightarrow 0} \tilde{\Omega}$ may be any one as long as there is a tubular neighborhood $\mathcal{U}$ of Γ_0 such that $\mathcal{U} \subset \Omega_\varepsilon$, $0 < \varepsilon \leq \varepsilon_0$, for ε_0 small enough.

As a matter of fact, in [10] the rescaled system (3.2) is considered in the rescaled limiting domain $\tilde{\Omega}$, rather than in Ω_ε .

It turns out however that the system for the stationary solutions to (3.2) is just (1.1) when $h_1 = 1$, $h_2 = D$ and $\delta(\tilde{\varepsilon}) = \tilde{\varepsilon}$, namely,

$$\begin{cases} \tilde{\varepsilon}^2 \Delta \tilde{u} + f(\tilde{u}, \tilde{v}) = 0, & X \in \Omega_\varepsilon \\ D \Delta \tilde{v} + \tilde{\varepsilon} g(\tilde{u}, \tilde{v}) = 0, & X \in \Omega_\varepsilon \\ \frac{\partial \tilde{u}}{\partial \tilde{n}} = \frac{\partial \tilde{v}}{\partial \tilde{n}} = 0, & X \in \partial \Omega_\varepsilon \end{cases} \quad (3.3)$$

Thus Corollary (2.2) applies to (3.3) with $\tilde{v}_0 = v^*$, $\alpha = h_-(v^*)$ and $\beta = h_+(v^*)$, thus yielding $J(v_0) = 0$.

Summing up the heuristic argument shown in [9], p. 1103, along with Corollary 2.2 strongly suggest that *if $(u_\varepsilon, v_\varepsilon)$ is a family of matched stable stationary solutions to (3.1) which is bounded in Ω , uniformly in ε and such that $(\tilde{u}_\varepsilon, \tilde{v}_\varepsilon)$ develops an internal transition layer (in the sense of Definition 1.1), then necessarily $v_\varepsilon \rightarrow v_0$, as $\varepsilon \rightarrow 0$, where v_0 is a constant such that $J(v_0) = 0$.*

It is worthwhile to emphasize that the formation of internal transition layer, in the sense of Definition 1.1, implies that any μ -level hypersurface of u_ε , $\alpha < \mu < \beta$, which lies in $\mathcal{U}$, converges uniformly to the interface, i.e., $\Gamma_\varepsilon \xrightarrow{\varepsilon \rightarrow 0} \Gamma_0$, uniformly. The last assertion has been assumed in [9] and [10].

Note that the case in which $\Gamma_0 \cap \partial\Omega \neq \emptyset$ has been ruled out although when h_1 and h_2 are constant functions this is more likely to happen. The analysis poses some additional technical difficulties and will be done in another work.

References

- [1] Hale, J. K. and Sakamoto, K., "Existence and stability of transition layers". *Japan J. Appl. Math.* **5** (1988), 367-405.
- [2] do Nascimento, A. S., "Existence, stability and geometric structure of N-dimensional transition layers in a reaction-diffusion equation." (submitted)
- [3] do Nascimento, A. S., "Inner transitions layers in a elliptic boundary value problem: a necessary condition." *Nonlinear Analysis, T., M. & A.*, to appear.
- [4] do Nascimento, A.S., "Stable Stationary Solutions Induced by Spatial Inhomogeneity Via Γ -Convergence". *Bulletin of the Brazilian Mathematical Society*, **29**, No. 1 (1998), 75-97.
- [5] do Nascimento, A.S., "Stable multiple-layer stationary solutions of a semi-linear parabolic equation in two-dimensional domains". *Electronic J. Diff. Eqns.*, **1997**, No.21 (1997), 1-17.
- [6] do Nascimento, A.S., "Reaction-diffusion induced stability of a spatially inhomogeneous equilibrium with boundary layer formation". *J. Diff. Eqns.* **108**, No.2 (1994), 296-325.
- [7] do Nascimento, A.S., "Boundary and interior spike layer formation for an elliptic equation with symmetry". *SIAM J. Math. Anal.*, **24**, No. 2 (1993), 466-479.
- [8] Nishiura, Y., "Global structure of bifurcating solutions of some reaction-diffusion systems". *S.I.A.M. J. Math. Anal.*, **13**, no. 13, (1982).
- [9] Nishiura, Y., Suzuki, H., "Nonexistence of higher dimensional stable Turing patterns in the singular limit." *S.I.A.M. J. Math. Anal.*, **29**, no. 5 (1998), 1087-1105.

- [10] Suzuki, H., "Asymptotic characterization of stationary interfacial patterns for reaction diffusion systems". *Hokkaido Math.* , **26**, (1997), 631-667.
- [11] Taniguchi, M., Nishiura, Y., "Instability of planar interfaces in reaction-diffusion systems". *S.I.A.M. J. Math. Anal.*, **25**, no. 1 (1994), 99-134.

Arnaldo Simal do Nascimento[†]

Dep. de Matemática, Universidade Federal de São Carlos, S. Carlos, S.P., Brasil

email address: arnaldon@dm.ufscar.br

Janete Crema[‡]

Instituto de Ciências Matemáticas e de Computação - USP, S. Carlos, S.P., Brasil

email address: janete@icmc.sc.usp.br

NOTAS DO ICMC

SÉRIE MATEMÁTICA

- 087/99 PÉREZ, V.H.J. – Polar multiplicities and equisingularity of maps germs from C^3 to C^3 .
- 086/99 BRUCE, J.W.; TARI, F. – Families of surfaces in IR^4 .
- 085/99 MOCHIDA, D.R.H.; ROMERO-FUSTER, M.C.; RUAS, M.A.S. – Inflections points and nonsingular embeddings of surfaces in IR^5 .
- 084/99 CARVALHO, A.N.; CHOLEWA, J.W. – Strongly damped wave equations with critical nonlinearities I: case $\theta = \frac{1}{2}$.
- 083/99 ARRAUT, J.L. – Circle leaves of two dimensional foliations.
- 082/99 CARVALHO, A.N.; GENTILE, C.B. – Comparison results for nonlinear parabolic equations with monotone principal part.
- 081/99 CARVALHO, A.N.; GENTILE, C.B. – Asymptotic behaviour of nonlinear parabolic equations with monotone principal part.
- 080/99 DUFOUR, J.P.; KUROKAWA, Y. – Topological rigidity for certain families of differentiable plane curves.
- 079/99 CONDE, A.; MENDES, M.P. - O n – pacote de variedades flâmulas reais e seus poliedros.
- 078/99 MENEGATTO, V.A.; PERON, A.P. – Positive definite Kernels on complex spheres.