

Jorge Carballido-Landeira
Bruno Escribano *Editors*

Nonlinear Dynamics in Biological Systems

SEMA SIMAI Springer Series

Series Editors: Luca Formaggia • Pablo Pedregal (Editors-in-Chief) •
Amadeu Delshams • Jean-Frédéric Gerbeau • Carlos Parés • Lorenzo Pareschi •
Andrea Tosin • Elena Vazquez • Jorge P. Zubelli • Paolo Zunino

Volume 7

More information about this series at <http://www.springer.com/series/10532>

Jorge Carballido-Landeira • Bruno Escibano
Editors

Nonlinear Dynamics in Biological Systems

Editors

Jorge Carballido-Landeira
Université Libre de Bruxelles
Brussels, Belgium

Bruno Escribano
Basque Center for Applied Mathematics
Bilbao, Spain

ISSN 2199-3041

SEMA SIMAI Springer Series

ISBN 978-3-319-33053-2

DOI 10.1007/978-3-319-33054-9

ISSN 2199-305X (electronic)

ISBN 978-3-319-33054-9 (eBook)

Library of Congress Control Number: 2016945273

© Springer International Publishing Switzerland 2016

This work is subject to copyright. All rights are reserved by the Publisher, whether the whole or part of the material is concerned, specifically the rights of translation, reprinting, reuse of illustrations, recitation, broadcasting, reproduction on microfilms or in any other physical way, and transmission or information storage and retrieval, electronic adaptation, computer software, or by similar or dissimilar methodology now known or hereafter developed.

The use of general descriptive names, registered names, trademarks, service marks, etc. in this publication does not imply, even in the absence of a specific statement, that such names are exempt from the relevant protective laws and regulations and therefore free for general use.

The publisher, the authors and the editors are safe to assume that the advice and information in this book are believed to be true and accurate at the date of publication. Neither the publisher nor the authors or the editors give a warranty, express or implied, with respect to the material contained herein or for any errors or omissions that may have been made.

Printed on acid-free paper

This Springer imprint is published by Springer Nature
The registered company is Springer International Publishing AG Switzerland

Preface

Many biological systems, such as circadian rhythms, calcium signaling in cells, the beating of the heart and the dynamics of protein folding, are inherently nonlinear and their study requires interdisciplinary approaches combining theory, experiments and simulations. Such systems can be described using simple equations, but still exhibit complex and often chaotic behavior characterized by sensitive dependence on initial conditions.

The purpose of this book is to bring together the mathematical bases of those nonlinear equations that govern various important biological processes occurring at different spatial and temporal scales.

Nonlinearity is first considered in terms of the evolution equations describing RNA neural networks, with emphasis on analysing population asymptotic states and their significance in biology. In order to achieve a quantitative characterization of RNA fitness landscapes, RNA probability distributions are estimated through stochastic models and a first approach for correction of the experimental biases is proposed. Nonlinear terms are also present during transcription regulation, discussed here through the mathematical modeling of biological logic gates, which opens the possibility of biological computing, with implications for synthetic biology. The instabilities arising in the different enzymatic kinetics models as well as their spatial considerations help in understanding the complex dynamics in protein activation and in developing a generic model. Furthermore, understanding of the biological problem and knowledge of the mathematical model may lead to the design of selective drug treatments, as discussed in the case of chimeric ligand–receptor interaction.

Nonlinear dynamics are also manifested in human muscular organs, such as the heart. Most of the solutions available in generic excitable systems are also obtained with mathematical models of cardiac cells which exhibit spatio-temporal dynamics similar to those of real systems. The spatial propagation of cardiac action potentials through the heart tissue can be mathematically formulated at the macroscopic level by considering a monodomain or bidomain system. The final nonlinear phenomenon discussed in the book is the electromechanical cardiac alternans. Understanding of

the mechanisms underlying the complex dynamics will greatly benefit the study of this significant biological problem.

The works compiled within this book were discussed by the speakers at the “First BCAM Workshop on Nonlinear Dynamics in Biological Systems”, held from 19 to 20 June 2014 at the Basque Center of Applied Mathematics (BCAM) in Bilbao (Spain). At this international meeting, researchers from different but complementary backgrounds—including disciplines such as molecular dynamics, physical chemistry, bio-informatics and biophysics—presented their most recent results and discussed the future direction of their studies using theoretical, mathematical modeling and experimental approaches.

We are grateful to the Spanish Government for their financial support (MTM201018318 and SEV-2013-0323) and to the Basque Government (Eusko Jaurlaritza) for help offered through projects BERC.2014-2017 and RC 2014 1 107. We also express gratitude to the Basque Center of Applied Mathematics for providing valuable assistance in logistics and administrative duties and creating a good atmosphere for knowledge exchange. JC-L also acknowledges financial support from FRS-FNRS. Furthermore, we are thankful to Aston University, Georgia Institute of Technology and Potsdam University for providing financial support for their scientific representatives.

Brussels, Belgium
Bilbao, Spain
July 2015

Jorge Carballido-Landeira
Bruno Escribano

Contents

Modeling of Evolving RNA Replicators	1
Jacobó Aguirre and Michael Stich	
Quantitative Analysis of Synthesized Nucleic Acid Pools.....	19
Ramon Xulvi-Brunet, Gregory W. Campbell, Sudha Rajamani, José I. Jiménez, and Irene A. Chen	
Non-linear Dynamics in Transcriptional Regulation: Biological Logic Gates	43
Till D. Frank, Miguel A.S. Cavadas, Lan K. Nguyen, and Alex Cheong	
Pattern Formation at Cellular Membranes by Phosphorylation and Dephosphorylation of Proteins	63
Sergio Alonso	
An Introduction to Mathematical and Numerical Modeling of Heart Electrophysiology	83
Luca Gerardo-Giorda	
Mechanisms Underlying Electro-Mechanical Cardiac Alternans	113
Blas Echebarria, Enric Alvarez-Lacalle, Inma R. Cantalapiedra, and Angelina Peñaranda	

Modeling of Evolving RNA Replicators

Jacobo Aguirre and Michael Stich

Abstract Populations of RNA replicators are a conceptually simple model to study evolutionary processes. Their prime applications comprise molecular evolution such as observed in viral populations, SELEX experiments, or the study of the origin and early evolution of life. Nevertheless, due to their simplicity compared to living organisms, they represent a paradigmatic model for Darwinian evolution as such. In this chapter, we review some properties of RNA populations in evolution, and focus on the structure of the underlying neutral networks, intimately related to the sequence-structure map for RNA molecules.

1 Introduction: RNA as a Paradigmatic Model for Evolution

RNA molecules are a very well suited model for studying evolution because they incorporate, in a single molecular entity, both genotype and phenotype. RNA sequence represents the genotype and the biochemical function of the molecule represents the phenotype. Since in many cases the *spatial structure* of the molecule is crucial for its biochemical function, the structure of an RNA molecule can be considered as a minimal representation of the phenotype. A population of replicating RNA molecules serves as a model for evolution because there is a mechanism that introduces genetic variability (e.g., point mutations) and a selection process that differentiates the molecules according to their fitness. Based particularly on the seminal work by Eigen [2], these concepts have been developed over the decades (for some early work see Refs. [3–6]) and have been proven to be very successful to describe evolutionary dynamics (for a more recent review see [7]).

J. Aguirre

Centro Nacional de Biotecnología (CSIC), Madrid, Spain

Grupo Interdisciplinar de Sistemas Complejos (GISC), Madrid, Spain

e-mail: jaguirre@cnb.csic.es

M. Stich (✉)

Non-linearity and Complexity Research Group, Aston University, Aston Triangle, B4 7ET
Birmingham, UK

e-mail: m.stich@aston.ac.uk

For many RNA molecules, it is difficult to obtain with experimental techniques (like X-ray crystallography) the actual, three-dimensional structure in which a given single-strand RNA sequence folds. In practice we know to precision only the structures of a few short RNA sequences, like the tRNA, of typically 76 bases (nucleotides). However, the *secondary structure* (in first approximation the set of Watson-Crick base pairs) of a molecule can be determined experimentally more easily and can be predicted with computational algorithms to high fidelity [8]. Since furthermore a large part of the total binding energy of a folded structure is found within the secondary structure, the latter has been widely accepted as a minimal representation of the phenotype of a molecule. The map from sequences to structures constitutes a special case of a *genotype-phenotype map*, lying the basis to introduce suitable fitness functions [9].

A population of evolving RNA replicators is henceforth described by its set of molecules, and for each molecule we know the sequence and the secondary structure and therefore the phenotype. It is important to note that a large number of different sequences can share the same folded structure, as shown in Fig. 1. Each color stands for a different secondary structure and each small square for a different sequence. This is only a schematic view, since we do not show all possible structures or sequences. We draw a link between two sequences if they differ in only one nucleotide. In this way, *neutral networks* are defined. As will be shown in more detail below, neutral networks differ in size and other properties, and may be actually disconnected.

In this chapter, we will distinguish and discuss two particular cases: in the first one *all* molecules share the same folded secondary structure—and have the same fitness. Then, the population is constraint to the neutral network and the evolutionary

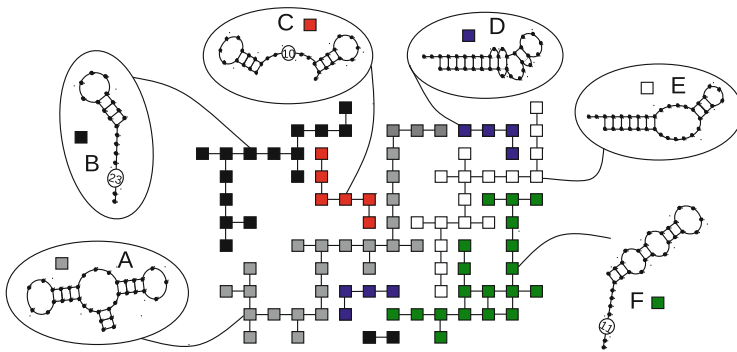


Fig. 1 Schematic view of RNA sequences and secondary structures. Each square represents a different RNA sequence of length 35. Squares with the same color fold into the same secondary structure. The RNA secondary structures are shown as *insets* (the number in the *circles* denotes the number of unpaired bases in that part of the molecule). The formed networks can have different sizes and may be disconnected. Missing squares do not correspond to missing RNA sequences, but to other, not displayed structures, including open structures. For a more detailed explanation, see main text. (Modified from [1])

dynamics of the population is largely determined by the topological properties of the network. In the second case, the molecules can be at any point in sequence space and hence can have different phenotypes—with therefore different fitness. We discuss the first case in Sect. 2 of this chapter, while Sect. 3 is devoted to the analysis of the second case. Section 4 concludes the chapter reflecting on the limitations and future developments in the modelization of evolving RNA replicators.

2 Evolution Constraint to Neutral Networks

2.1 RNA Neutral Networks: Definition and Main Properties

The idea of neutral evolution was first introduced by Kimura [10] in order to account for the known fact that a large number of mutations observed in proteins, DNA, or RNA, did not have any effect on fitness.

Neutrality becomes particularly important in the evolution of quasispecies [2], populations of fast mutating replicators which are formed by a large number of different phenotypes—and many more genotypes—and where high diversity and the concomitant steady exploration of the genome space happen to be an adaptive strategy. Relevant examples of quasispecies of RNA molecules are RNA viruses [11] (HIV, Ebola, flu, etc.) and error-prone replicators in the context of the prebiotic RNA world [12].

As it was already mentioned, RNA is a very fruitful model for the study of molecular evolution, and in fact RNA sequences folding into their minimum free energy secondary structures are likely the most used model of the genotype-phenotype relationship [8, 13, 14]. Here we assume that the RNA secondary structure can be used as a proxy for the phenotype—or biological function—of the molecule. Most models restrict their studies to the minimum free energy (MFE) structure. If this is the case, the mapping from sequence to secondary structure is many-to-one, i.e., there are many sequences that fold into the same structure. Assuming that all such sequences represent the same phenotype, they form a *neutral network* of genotypes. The number of different phenotypes gives the number of different neutral networks. The sequences that fold into the same secondary structure are the *nodes* of the neutral network. The *links* of the network connect sequences that are at a Hamming distance of one, i.e., that differ in only one nucleotide. Therefore, a neutral network may be connected—when all sequences are related to each other through single-point mutations—or disconnected. In the latter case, the neutral network is composed of a number of subnetworks. An example can be found in Fig. 2.

RNA is assembled as a chain of four nucleotides (adenine –A–, guanine –G–, cytosine –C– and uracil –U–). Therefore, the neutral networks associated to a certain sequence length l are subnetworks embedded in the macronetwork formed by all the possible networks of length l . The size of this macronetwork is 4^l , its dimension is

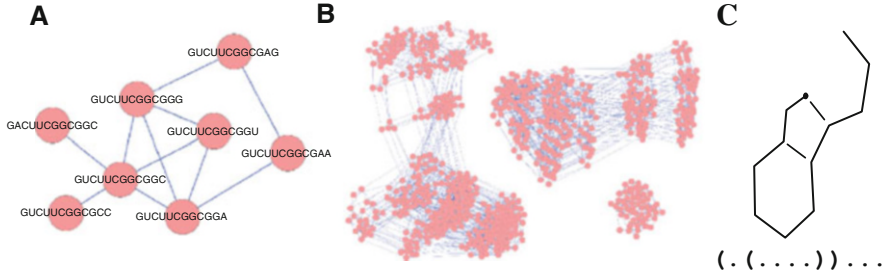


Fig. 2 Neutral network associated to an RNA secondary structure. (a) Sketch of the construction: the nodes of the network represent all of sequences that fold into the same secondary structure. They are connected if their Hamming distance is one, that is, if they differ in only one nucleotide. (b) Neutral network associated to the secondary structure of length 12 $(.)$, shown in (c). In this case, there are three disconnected subnetworks of size 404, 341 and 55. (Modified from [15])

l , and it shows a regular topology: each node has degree $3l$ because each nucleotide can suffer three different mutations.

Analytical studies of the number of sequences of length l compatible with a fixed secondary structure have revealed that the size of most neutral networks is astronomically large even for moderate values of the sequence length. For example, there should be about 10^{28} sequences compatible with the structure of a tRNA (which has length $l = 76$), while the smallest functional RNAs, of length $l \approx 14$ [16], could in principle be obtained from more than 10^6 . A rough upper bound to the number of different secondary structures S_l retrieved from sequences of length l , and valid for sufficiently large sequences and avoiding isolated base pairs, was derived in [5]: $S_l = 1.4848 \times l^{-3/2} (1.8488)^l$. This implies that the average size of a neutral network grows as $4^l/S_l = 0.673 \times l^{3/2} 2.1636^l$, which is a huge number even for moderate values of l . Let us note, however, that the actual distribution of neutral network sizes is a very broad function without a well-defined average and with a fat tail [5, 17]. In fact, the space of RNA sequences of length l is dominated by a relatively small number of common structures which are extremely abundant and happen to be found as structural motifs of natural, functional RNA molecules [18, 19]. Neutral networks corresponding to common structures percolate the space of sequences [20, 21] and thus facilitate the exploration of a large number of alternative structures. This is possible since different neutral networks are deeply interwoven: all common structures can be reached within a few mutational steps starting from any random sequence [21].

2.2 Dynamical Equations

We start assuming that, from all possible secondary structures into which RNA sequences of length l can fold, only one is functional (its specific function will depend on the biological context and is irrelevant now). Therefore, the fitness of the sequences of one network is maximized, while the fitness of all the sequences belonging to other neutral networks is zero. With this assumption, the evolution of a population of RNA through the space of genotypes due to mutations is limited to the functional neutral network (or to one of its subnetworks if it is not connected), as we consider that only the progeny that remains within the neutral network survives. Let's see how to model this process.

As it was already mentioned, the nodes of the neutral network associated to the functional secondary structure represent all the sequences that fold into it, while the links connect sequences that differ in only one nucleotide. Each node i in the network holds a number $n_i(t)$ of sequences at time t . There are $i = 1, \dots, m$ nodes in the network, each with a degree (number of nearest neighbors) k_i . The total population will be maintained constant through evolution, $N = \sum_i n_i(t)$, and we assume $N \rightarrow \infty$ to avoid stochastic effects due to finite population sizes. The initial distribution of sequences on the network at $t = 0$ is $n_i(0)$. Sequences of length l formed by 4 different nucleotides have at most $3l$ neighbors. We call $\{nn\}_i$ the set of actual neighbors of node i , whose cardinal is k_i . The vector $\mathbf{k}$ has as components the degree of the $i = 1, \dots, m$ nodes of the network. At each time step, the sequences at each node replicate. Daughter sequences mutate to one of the $3l$ nearest neighbors with probability μ , and remain equal to their mother sequence with probability $1 - \mu$. In our representation $0 < \mu \leq 1$. The singular case $\mu = 0$ is excluded to avoid trivial dynamics and guarantee evolution towards a unique equilibrium state. With probability $k_i/(3l)$, the mutated sequence exists in the neutral network and it adds to the population of the corresponding neighboring nodes. Otherwise, it falls off the network and disappears, this being the fate of a fraction $(1 - k_i/(3l))\mu$ of the total daughter sequences.

The mean-field equations describing the dynamics of the population on the network read

$$n_i(t+1) = (2 - \mu)n_i(t) + \frac{\mu}{3l} \sum_{j \in \{nn\}_i} n_j(t). \quad (1)$$

The dynamics can be written in matrix form as

$$\mathbf{n}(t+1) = (2 - \mu)\mathbf{In}(t) + \frac{\mu}{3l}\mathbf{Cn}(t), \quad (2)$$

where $\mathbf{I}$ is the identity matrix and $\mathbf{C}$ is the adjacency matrix of the network, whose elements are $C_{ij} = 1$ if nodes i and j are connected and $C_{ij} = 0$ otherwise. The transition matrix $\mathbf{M}$ is defined as

$$\mathbf{M} = (2 - \mu)\mathbf{I} + \frac{\mu}{3l}\mathbf{C}. \quad (3)$$

2.3 Analysis of the Asymptotic State

Let us call $\{\lambda_i\}$ the set of eigenvalues of $\mathbf{M}$, with $\lambda_i \geq \lambda_{i+1}$, and $\{\mathbf{u}_i\}$ the corresponding eigenvectors. Furthermore, its eigenvectors verify $\mathbf{u}_i \cdot \mathbf{u}_j = 0$, $\forall i \neq j$ and $|\mathbf{u}_i| = 1$, $\forall i$.

A matrix is irreducible when the corresponding graph is connected; in our case any pair of nodes i and j of the network are connected via mutations by definition. Irreducibility plus the condition $M_{ii} > 0$, $\forall i$ makes matrix $\mathbf{M}$ primitive. Since $\mathbf{M}$ is a primitive matrix, the Perron-Frobenius theorem assures that, in the interval of μ values used, the largest eigenvalue of $\mathbf{M}$ is positive, $\lambda_1 > |\lambda_i|$, $\forall i > 1$, and its associated eigenvector is positive (i.e., $(\mathbf{u}_1)_i > 0$, $\forall i$).

The dynamics of the system, Eq. (2) can be thus written as

$$\mathbf{n}(t) = \mathbf{M}^t \mathbf{n}(0) = \sum_{i=1}^m \lambda_i^t \alpha_i \mathbf{u}_i, \quad (4)$$

where we have defined α_i as the projection of the initial condition on the i th eigenvector of $\mathbf{M}$,

$$\alpha_i = \mathbf{n}(0) \cdot \mathbf{u}_i. \quad (5)$$

Furthermore, as $\lambda_1 > |\lambda_i|$, $\forall i > 1$, the asymptotic state of the population is proportional to the eigenvector that corresponds to the largest eigenvalue, $\mathbf{u}_1$:

$$\lim_{t \rightarrow \infty} \left(\frac{\mathbf{n}(t)}{\lambda_1^t \alpha_1} \right) = \mathbf{u}_1, \quad (6)$$

while the largest eigenvalue λ_1 yields the growth rate of the population at equilibrium. For convenience, in the following, and without any loss of generality, we normalize the population $\mathbf{n}(t)$ such that $|\mathbf{n}(t)| = 1$ after each generation. With this normalization, $\mathbf{n}(t) \rightarrow \mathbf{u}_1$ when $t \rightarrow \infty$.

Finally, let us remark the biological relevance of these results, as they yield that, independently of the initial condition, the sequence population evolves asymptotically towards a constant distribution over the network that can be obtained by the first eigenvector of the transition matrix.

2.4 Time to Asymptotic Equilibrium

A dynamical quantity with direct biological implications is the time that the population takes to reach the mutation-selection equilibrium. To obtain it, we start from Eq. (4), where we describe the transient dynamics towards equilibrium starting with an initial condition $\mathbf{n}(0)$.

The distance $\Delta(t)$ to the equilibrium state can be written as

$$\Delta(t) \equiv \left| \frac{\mathbf{M}^t \mathbf{n}(0)}{\lambda_1^t \alpha_1} - \mathbf{u}_1 \right| = \left| \sum_{i=2}^m \frac{\alpha_i}{\alpha_1} \left(\frac{\lambda_i}{\lambda_1} \right)^t \mathbf{u}_i \right|. \quad (7)$$

In order to estimate how many generations elapse before equilibrium is reached, we fix a threshold ϵ , and define the *time to equilibrium* t_ϵ as the number of generations required for $\Delta(t_\epsilon) < \epsilon$.

When $\alpha_2 \neq 0$, $\lambda_2 \neq 0$ and $\lambda_2 \neq \lambda_3$, t_ϵ can be approximated to first order by

$$t_\epsilon^1 \simeq \frac{\ln |\alpha_2/\alpha_1| - \ln \epsilon}{\ln |\lambda_1/\lambda_2|}. \quad (8)$$

This approximation turns out to be extremely good in most cases thanks to the exponentially fast suppression of the contributions due to higher-order terms (since $\lambda_i \geq \lambda_{i+1}$, $\forall i$). It may lose accuracy, however, when $\lambda_3 \approx \lambda_2$, when the initial condition $\mathbf{n}(0)$ is such that $\alpha_3 \gg \alpha_2$, or when ϵ is so large that the population is far from equilibrium and $\Delta(t)$ is still governed by λ_3 and higher order eigenvalues.

2.5 Influence of the Network Topology

Let us call $\{\gamma_i\}$ the set of eigenvalues of adjacency matrix $\mathbf{C}$, $\gamma_i \geq \gamma_{i+1}$, and $\{\mathbf{w}_i\}$ the set of corresponding eigenvectors. From Eq. (3),

$$\begin{aligned} \mathbf{M} \mathbf{w}_i &= (2 - \mu) \mathbf{I} \mathbf{w}_i + \frac{\mu}{3l} \mathbf{C} \mathbf{w}_i \\ &= \left[(2 - \mu) + \frac{\mu}{3l} \gamma_i \right] \mathbf{w}_i. \end{aligned} \quad (9)$$

The eigenvectors of the adjacency matrix are also eigenvectors of the transition matrix, $\mathbf{u}_i \equiv \mathbf{w}_i$, $\forall i$, demonstrating that the asymptotic state of the population only depends on the topology of the neutral network. The eigenvalues of both matrices are thus related through

$$\lambda_i = (2 - \mu) + \frac{\mu}{3l} \gamma_i, \quad (10)$$

where the set $\{\gamma_i\}$ does not depend on the mutation rate μ . The adjacency matrix contains all the information on the final states, while the transition matrix yields quantitative information on the dynamics towards equilibrium.

The minimal value of λ_1 is obtained in the limit of a population evolving at a very high mutation rate ($\mu \rightarrow 1$) on an extremely diluted matrix ($l \rightarrow \infty$). In this limit, all eigenvalues of $\mathbf{M}$ become asymptotically independent of the precise topology of the network and $\lambda_i \rightarrow 1, \forall i$. In this extreme case all daughter sequences fall off the network, but the population is maintained constant through the parental population. An extinction catastrophe (due to a net population growth below one [22]) never holds under this dynamics.

2.6 Evolution Towards Mutational Robustness

The *average degree* $K(t)$ of the population at time t is defined as

$$K(t) = \frac{\mathbf{k} \cdot \mathbf{n}(t)}{\sum_i n_i(t)}. \quad (11)$$

In the limit $t \rightarrow \infty$, we obtain the average degree at equilibrium

$$K(t \rightarrow \infty) = K = \frac{\mathbf{k} \cdot \mathbf{u}_1}{\sum_i (u_1)_i}. \quad (12)$$

We define as k_{min} , k_{max} , and $\langle k \rangle = \sum_i k_i / m$ the smallest, largest, and average degree of the network, respectively. The Perron-Frobenius theorem for non-negative, symmetric, and connected graphs, sets limits on the average degree $\langle k \rangle$: When $k_{min} < k_{max}$, that is, as far as the graph is not homogeneous,

$$k_{min} < \langle k \rangle < \gamma_1 < k_{max} \quad (13)$$

holds. A simple calculation yields that the average degree of the population at equilibrium, K , is equal to the largest eigenvalue γ_1 of the adjacency matrix, also known as the spectral radius of the network [23]. Therefore, from Eq. (13) we obtain that the average degree K of the population at equilibrium will be larger than the average degree $\langle k \rangle$ of the network, indicating that when all nodes are identical the population selects regions with connectivity above average. This demonstrates a natural evolution towards mutational robustness, because the most connected nodes are those with the lowest probability of suffering a (lethal) mutation that could push them out of the network. See Fig. 3 for a numerical example of this phenomenon.

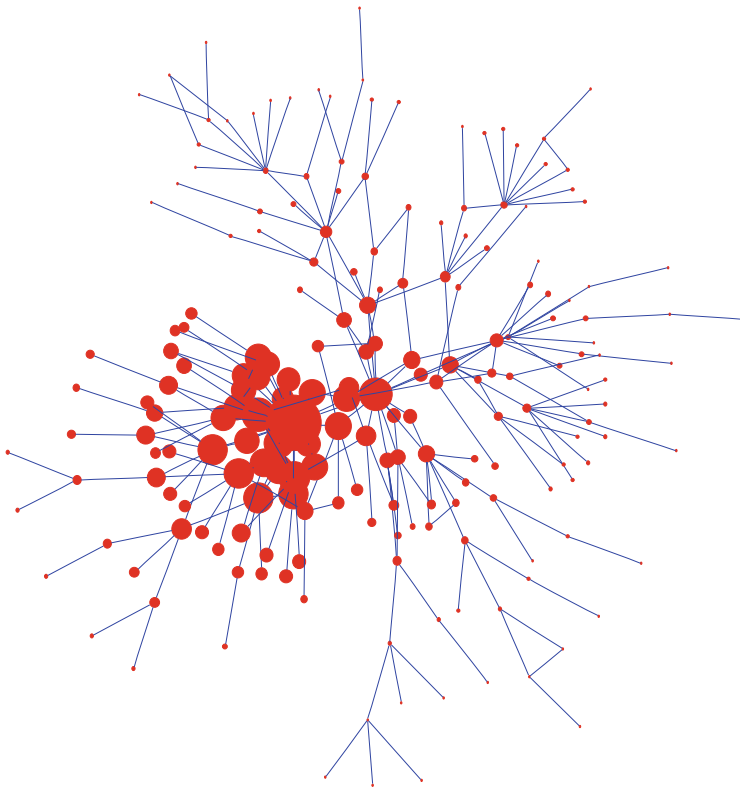


Fig. 3 Evolution towards mutational robustness. A population of replicators that evolves on a complex network tends to the most connected regions in the mutation-selection equilibrium. In the figure we show a population that has evolved following Eq. (2) over an artificial RNA neutral network. The size of each node is proportional to the population in the equilibrium. The natural tendency to populate the region of the network with the largest average degree is clear; in consequence, the robustness against mutations is enhanced. (Modified and reproduced with permission from [24])

3 Evolution Towards a Target Structure

Until now we have considered RNA populations confined to a given neutral network and hence secondary structure. Molecules that fold into any other secondary structure were simply discarded. This is of course only an approximation. In this section we present a theoretical framework that includes all possible RNA secondary structures and assigns non-zero fitness values to all of them. Such a landscape was introduced in Fig. 1, and we describe the evolutionary dynamics on it in the following.

3.1 Sequence-Structure Map

Above, we have already introduced RNA sequence and structure in the context of RNA neutral networks, but we have to explain in more depth some properties of the sequence-structure map. Two fundamental properties of the sequence-structure map are that (1) the number of different sequences is much higher than the number of structures and (2) not all possible structures are equally probable [5, 18]. In this context, *common* structures are those which have many different sequences folding into them (having large neutral networks) and *rare* structures are those which have only few sequences folding into them (having small neutral networks).

In this section, we are interested in the whole set of neutral networks and its size distribution. We describe some results of the folding of 10^8 RNA molecules of length 35 nt consisting of random sequences, following Ref. [17]. In total, the sequences folded into 5,163,324 structures, using the `fold()` routine from the Vienna RNA Package [25]. A way to visualize the distribution of sequences into structures is the frequency-rank diagram. In Fig. 4b the structures are ranked according to the number of sequences folding into them and we focus first on the thick black curve. One can see that there are around thousand common structures, each of them obtained from about 10^4 different sequences. On the other hand,

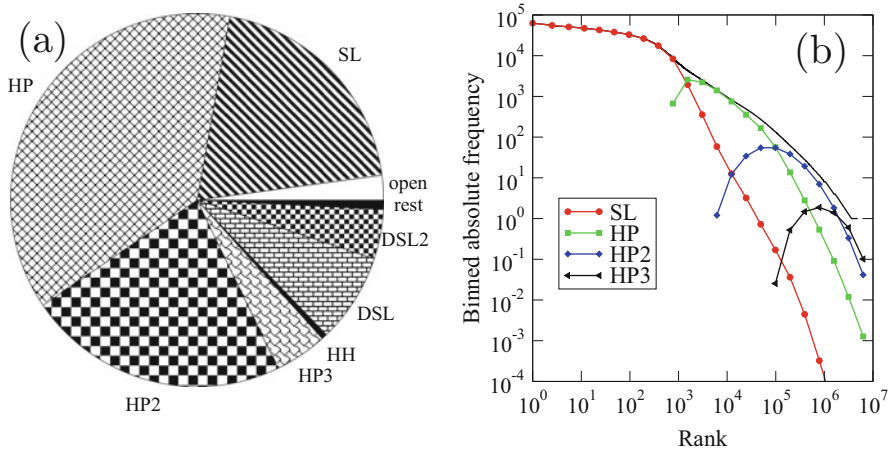


Fig. 4 Properties of the RNA sequence-structure map. **(a)** Distribution of the sequences in structure families according to their frequency. Higher-order hairpins, HP x , are defined as $(H,I,M)=(1,x,0)$, being $x \geq 2$, higher-order double stem-loops, DSL x , as $(H,I,M)=(2,x-1,0)$, and higher-order hammerheads, HH x , as $(2,x-1,1)$. **(b)** Frequency-rank diagram according to the structural family. We have binned in boxes of powers of 2 the total number of structures belonging to the interval and have determined the absolute frequency of the corresponding sequences for each family in the bin. The *thick black curve* denotes the whole set of structures (see text). (Modified from [17])

we also find a few million rare structures yielded by only one or two sequences. Although for smaller pools, this had already been reported before (e.g., [5]).

In order to study the distribution of common vs. rare structures in more detail, we have proposed a classification where an RNA secondary structure is characterized in terms of three numbers [17]: (a) the number of hairpin (terminal) loops, H , (b) the sum of bulges and interior loops, I , and (c) the number of multiloops, M . For example, a simple stem-loop structure, denoted as SL, is characterized by $(H,I,M)=(1,0,0)$, and all stem-loop structures found in the pool are grouped into that *structure family*. Other important families are the hairpin structure family, HP, with one interior loop or bulge $(1,1,0)$, the double stem-loop, DSL, represented by $(2,0,0)$, and the simple hammerhead structure, HH, by $(2,0,1)$. Of course, there exist more complicated structure families, as detailed in [17]. For the pool that we have folded, we find that only 21 structure families are enough to cover all the 5.2 million structures identified.

Our analysis, displayed in Fig. 4a, shows that the vast majority of sequences fold into simple structure families. For example, 79.0 % of all sequences belong to only three structure families (HP, HP2, SL, in decreasing abundance), and 92.1 % of all sequences fold into simple structures with at most three stems (HP, HP2, SL, DSL, DSL2, HH). Note that 2.1 % of all sequences remain open and do not fold. This data is in agreement with previous findings on the structural repertoire of RNA sequence pools where the influence of the sequence length [26, 27], the nucleotide composition [28, 29], and pool size [27] were studied.

With this classification, we can reconsider the frequency-rank diagram. We sum up all structures of a given structure family within a rank interval. Through this binning procedure, for each structure family a curve is obtained which describes its relative frequency compared to that of the other families. The curves for some families are shown in Fig. 4b. We immediately see that the most frequent structures belong to the stem-loop family, followed by the hairpin family. For low ranks, the SL curve is identical to the curve describing all structures. For ranks between 4×10^3 and 10^4 , it is the HP curve which practically coincides with the total curve.

The understanding of the diversity of structures, the (relative) sizes of the neutral networks, and their distribution in sequence space are relevant properties to take into account when populations are allowed to move across the whole sequence space, as considered in the following.

3.2 Evolutionary Algorithm

We assume that an RNA molecule of a given length can have any possible sequence, i.e., the population can access the whole sequence space. Furthermore, by the mapping introduced above, all secondary structures can be determined. We now identify one structure as *target structure*. Above it has been argued that structure is a proxy for biochemical function and therefore the target structure represents optimal biochemical function in a given environment. All molecules have now a fitness that

depends on their *structure distance* to the target structure—structure distances are defined below. Replication according to the fitness function subjected to mutations enables the population to search for, find, and finally fix the target structure. The evolution of this population can be expected to be a highly nonlinear process, not only because of the vast amount of different structures (and fitnesses), but also because two sequences that are just one mutation apart, may fold into structures very different from each other. At the same time, in a relatively small neighborhood of any sequence, almost all common structures can be found [18]. The scheme is as follows:

1. Set a target structure of length l . Any possible secondary structure can be chosen, but often biological relevant structures like hairpin or hammerhead ribozymes or tRNA structures are selected. Sequence space, imprecision of folding algorithms and computational resources increase strongly with the length of the molecule. Therefore, most computational studies use short molecules, up to a few hundred bases. In many examples, we use a 35 nt hairpin ribozyme structure.
2. Construct an initial population of size N . Typical choices are random sequences, sequences pre-evolved under different conditions, or sequences selected from an experiment or database.
3. Fold the structures with an appropriate routine, like `fold()` from the Vienna Package [25].
4. Determine the distance to the target structure for all molecules. To compare RNA secondary structures, there exists a range of different distance measures like base-pair, Hamming, or various kinds of tree-edit, all of which establish a metric [25]. Consequently, we can use any of these distance measures to define a fitness function.
5. Replicate the population according to the Wright-Fisher sampling, keeping the population size constant, and using the normalized probabilities

$$p(d_i) = \frac{\exp(-\beta d_i)}{\sum_{i=1}^N \exp(-\beta d_i)}. \quad (14)$$

We discuss the fitness function itself further below, but we already mention here two possibilities for β : (a) a constant, e.g., $\beta = 1/l$, or (b) a time-dependent function, e.g., $\beta = 1/d$, being $d = \sum_{i=1}^N d_i/N$ the average distance of the population.

6. In the replication event, a source of variability for the genotype must be included, being the simplest case point mutations with a rate μ .
7. Steps (3)–(6) are repeated as long as desired. Most of the evolutionary dynamics can be seen within $10^2 - 10^4$ generations.

The scheme as presented here can be modified in several ways: for example the selection function may not be an exponential function, or the comparison between structures can rely on other distance measures. Nevertheless, many properties of evolving RNA populations only depend quantitatively, not qualitatively on these variables. Obviously, the fitness/replication function can be modified to

include also the Hamming distance to a conserved sequence, or some energetic constraint [30]. Finally, even the condition of N constant can be relaxed to allow for bottlenecks [31].

3.3 Characteristic Quantities

In the following, it is assumed that the initial population consists of random sequences of length $l = 35$ and that the population size is in the range 10^2 – 10^3 . Then, even though a relatively common target structure is chosen, it is not likely that the initial population contains a molecule that folds into the target structure. Therefore, the first phase of evolution is a search phase where molecules that fold into a structure more similar to the target structure are slowly picked up and replicate with a higher probability. This phase terminates when the population finds the target structure for the first time (an event that defines the *search time*). Then, the second phase starts, in which the number of molecules folding into the target structure typically increases (we say that the structure is being *fixed*). However, due to the probabilistic nature of the system, and in particular for high mutation rates, the structure can be lost again. Then, no fixation takes place.

After fixation, the population enters the third phase, the time to approach the asymptotic state, characterized by a minimal average distance and a maximal number of molecules in the target structure. Due to the non-deterministic features of the model, the asymptotic state of the system can only be characterized after ensemble and time averaging. Furthermore, if N is too small, finite size effects can be observed.

In Fig. 5, we show an example of a typical simulation. We see for three different mutation rates the evolution of an RNA population, illustrated by the average

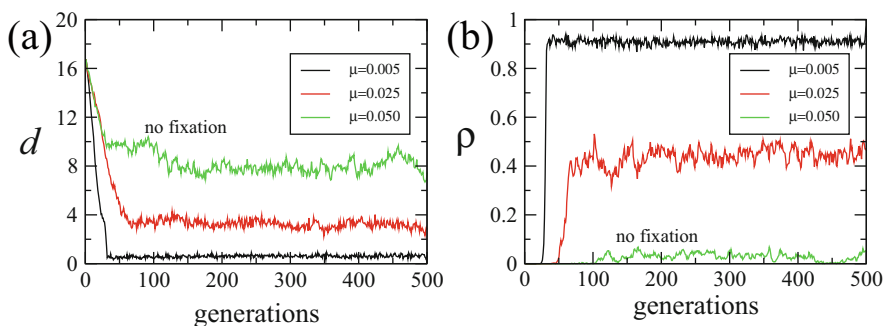


Fig. 5 Typical evolutionary processes for RNA populations with three different mutation rates. In (a), we show d , the average base-pair distance to the target structure and in (b), we show ρ , the density of correctly folded molecules in the population. Parameters: $N = 602$, $\beta = 1/d$, and the target structure is a hairpin structure with $l = 35$

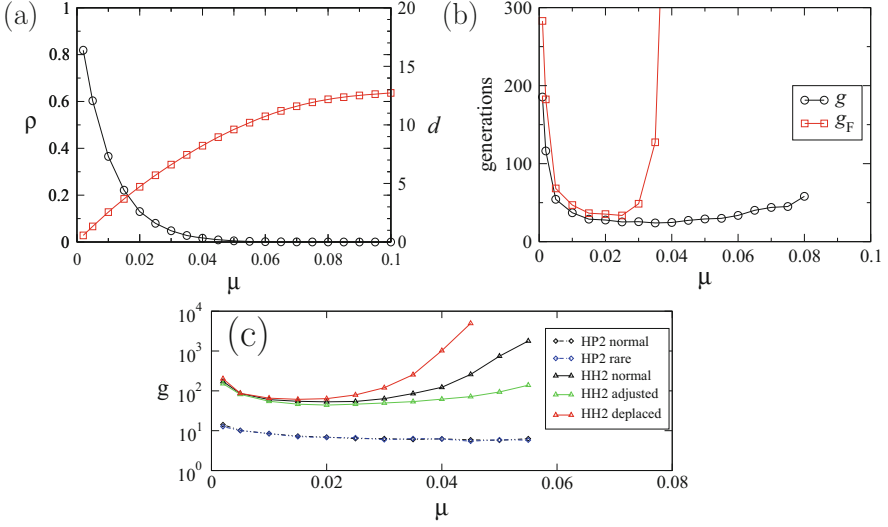


Fig. 6 Characteristic quantities for evolving RNA populations. In (a), we show the asymptotic values for ρ (the density of correctly folded molecules) in the population and d (the average base-pair distance to the target structure), in (b), we show the search and search-plus-fixation times g and g_F , and in (c), the search time for different target structures and nucleotide compositions

distance d and the density of correct structures ρ . If the mutation rate is small, the population ends up with a lot of correct structures and the average distance is small. If the mutation rate is higher, less correct structures are found and the distance increases. If the mutation rate is larger than a certain threshold, the target structure is not maintained steadily in the population and is not fixed. The fixation (or not) of target structures can be formulated in terms of the population being below (or above) the phenotypic error threshold.

Figure 6a plots the asymptotic values for ρ and d as a function of the mutation rate. In qualitative agreement with Fig. 5, the average distance increases and the average density decreases with μ . However, the transition to no-fixation cannot be seen in this graph and therefore, the search and search-plus-fixation times are plotted in Fig. 6b. These curves reveal the success of the evolutionary process: while the search time g is a function that decreases with μ (more mutations make it more likely to find the target structure), whereas the search-plus-fixation time g_F (note that $g_F \geq g$) is showing a strong increase for intermediate values of μ . The case of no fixation sets in where the graph of g_F diverges.

In [32], the dependence of the search time on several parameters of the model was studied in detail. For example, the search times do not only depend on the mutation rate, but also on the chosen structure. One particular interesting result is shown in Fig. 6c where we display the search time for two different target structures and various nucleotide compositions. The HP2 structure (the most abundant structure of the HP2 class) is relatively frequent, and we let a population evolve under

two different nucleotide compositions: one is the standard one (25 % of mutation probability for all types of bases A,C,G,U), and the other one is a nucleotide composition of $(A,C,G,U) = (0.23, 0.26, 0.30, 0.21)$, the average composition of rare structures [32]. Essentially, we do not see differences in the search time. This shows that for this structure (and for all frequent structures) the nucleotide composition is not crucial since molecules folding into that structures are found in all parts of sequence space.

The other structure being used is the most abundant structure of the HH2 class (short: HH2 structure), a rare structure compared to the HP2 structure. There, tuning the nucleotide pool into a direction that makes it more likely to find HH2 structures $(A,C,G,U) = (0.22, 0.26, 0.32, 0.20)$ is an efficient way to shorten the search (and search-plus-fixation) times. This shows that the neutral network of this HH2 structure is not uniformly distributed in sequence space. Note that we changed the nucleotide composition on the basis of the average over only seven found HH2 sequences (among 100 million random sequences), whereas the real size of the neutral network is much larger (thousands of different sequences are found by the evolutionary algorithm through the simulations leading to this figure). Then, the objective becomes to see whether the search process could be less efficient: we chose “opposite” nucleotide composition values, i.e., $(A,C,G,U) = (0.28, 0.24, 0.18, 0.30)$. Again, the search times are displaced—this time towards larger values, thus showing that indeed the neutral network of HH2 molecules is not distributed in a uniform way over the sequence space.

Discovering shorter search times for nucleotide compositions that are correlated with certain structures obtained by random folding looks like a circular argument, but it is not: the evolutionary algorithm finds the neutral network, and detects with higher probability the dominant parts of it (if there are), whereas the random folding samples the sequence space with equal probability.

4 Discussion: Limitations and Future Lines in the Modelization of Evolving RNA Replicators

We devote this last section to sketch some future lines of work related to the modelization of evolving RNA replicators. In particular, we will pay special attention to the potentialities and limitations in the applicability of the tools and results presented throughout the chapter.

Regarding the modelization of RNA neutral networks (i.e., Sect. 2), their extremely large size precludes systematic studies unless we are analysing very short RNAs, and poses a major computational challenge. Currently, extensive studies folding all RNA sequences have been done only for lengths below or around 20 nucleotides, where the number of sequences to be folded does not exceed $\sim 10^{12}$ [33–35]. For example, we systematically studied the topological properties of the RNA neutral networks corresponding to $l = 12$ in [15]. In this

work it was shown that the topology of RNA neutral networks is far from being random, and it was proved that the average degree, and therefore the robustness to mutations, grows with the size of the network. The latter reinforces the hypothesis that RNA structures with the largest networks associated seem to be the ones present in natural, functional RNA molecules. The reasons are two-fold: A sequence randomly evolving in the space of genomes will find an extensive network more easily, and once it has been found, the greater robustness to mutations will keep the sequences in the network with a larger probability [35]. On the other hand, the results obtained for neutral networks corresponding to short RNAs might not be directly applicable to longer RNAs such as tRNA ($l = 76$), and obviously we are still far from applying these techniques to RNA viruses, where genetic chains of several thousands of nucleotides or even more are usual. The computational, theoretical and experimental challenge is enormous.

We must not forget that in order to develop the calculations shown in Sect. 2, we have assumed that there is only one functional neutral network. In Nature, several different—perhaps similar—secondary structures might develop the same function, with possibly different efficiency depending on environmental factors. Each of the neutral networks associated with these structures can be connected in a network of networks, where sequences mutate and jump from one network to another while they vary their fitness. For instance, it has been recently shown that the probability that the population leaves a neutral network is smaller the longer the time spent on it, leading to a *phenotypic entrapment* [36]. On the other hand, if each of these networks is represented by a single node, then we get what is known as a *phenotype network* [33], a concept to be explored in more detail in the future.

If we lift the constraint of a population living on a single neutral network completely and allow all sequences and secondary structures (as considered in Sect. 3), we are a priori more realistic from a biological point of view, since the limitation on a single structure is a rough approximation. The price paid for this is the analytical (in the mathematical sense) intractability, since the sequence-structure map is highly nonlinear, and there is an intrinsic ambiguity in defining a fitness function. Furthermore, sequence space, imprecision of folding algorithms and computational resources increase strongly with the length of the molecule and the size of the population. Therefore, most computational studies use short molecules and small populations. However, most biologically relevant cases correspond not only to longer molecules than those considered in *in silico* studies, but also to larger populations. Besides, to map a sequence to a single minimum free energy structure is a strong approximation, since it does not take into account suboptimal structures (those with a slightly larger free energy) or kinetic folding processes. From this it follows that molecular evolution experiments can access a much larger part of the sequence space than simulations and hence that conclusions drawn by numerical studies have to be interpreted with caution.

On the other hand, the paradigm of RNA populations evolving *in silico* is still very useful since many relevant parameters—not only molecule length and population size—can be tuned individually. The main one is certainly the mutation rate: above, we have shown how the interplay between mutation rate, target

structure, or nucleotide content changes the evolutionary dynamics significantly. Note that we can also impose selection criteria on sequences rather than structures (e.g., if a certain structure is known to have a conserved sequence as well) or on the folding energy. The population size need not to be constant either, such that population bottlenecks can be studied. Other works show the importance of evolving RNA populations to study phylogenetic trees [37] or to establish a relationship between microscopic and phenotypic mutation rates [30]. Future work using this approach will probably see the development of models to study generic features of evolution and data-driven approaches, aiming at understanding natural RNA obtained from sequencing experiments or *in vitro* selected RNA from SELEX experiments.

Acknowledgements The authors are indebted to S. Manrubia for useful discussions and comments on the manuscript. JA acknowledges financial support from Spanish MICINN (projects FIS2011-27569 and FIS2014-57686). MS acknowledges financial support from Spanish MICINN (project FIS2011-27569).

References

1. Manrubia, S.C., Cuesta, J.A.: Neutral networks of genotypes: evolution behind the curtain. *ARBOR* **186**, 1051–1064 (2010)
2. Eigen, M.: Selforganization of matter and the evolution of biological macromolecules. *Naturwissenschaften* **58**, 465–523 (1971)
3. Fontana, W., Schuster, P.: A computer model of evolutionary optimization. *Biophys. Chem.* **26**, 123–47 (1987)
4. Huynen, M.A., Konings, D.A.M., Hogeweg, P.: Multiple coding and the evolutionary properties of RNA secondary structure. *J. Theor. Biol.* **165**, 251–267 (1993)
5. Schuster, P., Fontana, W., Stadler, P.F., Hofacker, I.L.: From sequences to shapes and back: a case study in RNA secondary structures. *Proc. R. Soc. Lond. B* **255**, 279–284 (1994)
6. Schuster, P.: Genotypes with phenotypes: adventures in an RNA toy world. *Biophys. Chem.* **66**, 75–110 (1997)
7. Takeuchi, N., Hogeweg, P.: Evolutionary dynamics of RNA-like replicator systems: a bioinformatic approach to the origin of life. *Phys. Life Rev.* **9**, 219–263 (2012)
8. Schuster, P.: Prediction of RNA secondary structures: from theory to models and real molecules. *Rep. Prog. Phys.* **69**, 1419 (2006)
9. Stadler, P.F.: Fitness landscapes arising from the sequence-structure maps of biopolymers. *J. Mol. Struct. THEOCHEM* **463**, 7–19 (1999)
10. Kimura, M.: Evolutionary rate at the molecular level. *Nature* **217**, 624–626 (1968)
11. Holland, J.J., De la Torre, J.C., Steinhauer, D.A.: RNA virus populations as quasispecies. *Curr. Top. Microbiol. Immunol.* **176**, 1–20 (1992)
12. Briones, C., Stich, M., Manrubia, S.C.: The dawn of the RNA world: toward functional complexity through ligation of random RNA oligomers. *RNA* **15**, 743–9 (2009)
13. Ancel, L.W., Fontana, W.: Plasticity, evolvability, and modularity in RNA. *J. Exp. Zool. Mol. Dev. Evol.* **288**, 242–283 (2000)
14. Fontana, W.: Modelling ‘evo-devo’ with RNA. *BioEssays* **24**, 1164–77 (2002)
15. Aguirre, J., Buldú, J.M., Stich, M., Manrubia, S.C.: Topological structure of the space of phenotypes: the case of RNA neutral networks. *PLoS ONE* **6**, e26324 (2011)

16. Anderson, P.C., Mecozzi, S.: Unusually short RNA sequences: design of a 13-mer RNA that selectively binds and recognizes theophylline. *J. Am. Chem. Soc.* **127**, 5290–1 (2005)
17. Stich, M., Briones, C., Manrubia, S.C.: On the structural repertoire of pools of short, random RNA sequences. *J. Theor. Biol.* **252**, 750–63 (2008)
18. Fontana, W., Konings, D.A.M., Stadler, P.F., Schuster, P.: Statistics of RNA secondary structures. *Biopolymers* **33**, 1389–404 (1993)
19. Gan, H.H., Pasquali, S., Schlick, T.: Exploring the repertoire of RNA secondary motifs using graph theory; implications for RNA design. *Nucleic Acids Res.* **31**, 2926–43 (2003)
20. Grüner, W., Giegerich, R., Strothmann, D., Reidys, C., Weber, J., Hofacker, I.L., Stadler, P.F., Schuster, P.: Analysis of RNA sequence structure maps by exhaustive enumeration II. Structures of neutral networks and shape space covering. *Monatsh. Chem.* **127**, 375–389 (1996)
21. Reidys, C., Stadler, P.F., Schuster, P.: Generic properties of combinatory maps: neutral networks of RNA secondary structures. *Bull. Math. Biol.* **59**, 339–97 (1997)
22. Bull, J.J., Meyers, L.A., Lachmann, L.: Quasispecies made simple. *PLoS Comput. Biol.* **1**, 450–460 (2005)
23. van Nimwegen, E., Crutchfield, J.P., Huynen, M.: Neutral evolution of mutational robustness. *Proc. Natl. Acad. Sci. U. S. A.* **96**, 9716–20 (1999)
24. Aguirre, J., Buldú, J., Manrubia, S.: Evolutionary dynamics on networks of selectively neutral genotypes: effects of topology and sequence stability. *Phys. Rev. E* **80**, 066112 (2009)
25. Hofacker, I.L., Fontana, W., Stadler, P.F., Bonhoeffer, L.S., Tacker, M., Schuster, P.: Fast folding and comparison of RNA secondary structures. *Monatsh. Chem.* **125**, 167–188 (1994)
26. Sabeti, P.C., Unrau, P.J., Bartel, D.P.: Accessing rare activities from random RNA sequences: the importance of the length of molecules in the starting pool. *Chem. Biol.* **4**, 767–74 (1997)
27. Gevertz, J., Gan, H.H., Schlick, T.: In vitro RNA random pools are not structurally diverse: a computational analysis. *RNA* **11**, 853–63 (2005)
28. Knight, R., De Sterck, H., Markel, R., Smit, S., Oshmyansky, A., Yarus, M.: Abundance of correctly folded RNA motifs in sequence space, calculated on computational grids. *Nucleic Acids Res.* **33**, 5924–35 (2005)
29. Kim, N., Gan, H.H., Schlick, T.: A computational proposal for designing structured RNA pools for in vitro selection of RNAs. *RNA* **13**, 478–92 (2007)
30. Stich, M., Lázaro, E., Manrubia, S.C.: Phenotypic effect of mutations in evolving populations of RNA molecules. *BMC Evol. Biol.* **10**, 46 (2010)
31. Stich, M., Lázaro, E., Manrubia, S.C.: Variable mutation rates as an adaptive strategy in replicator populations. *PLoS ONE* **5**, e11186 (2010)
32. Stich, M., Manrubia, S.C.: Motif frequency and evolutionary search times in RNA populations. *J. Theor. Biol.* **280**, 117–26 (2011)
33. Cowperthwaite, M., Economo, E., Harcombe, W., Miller, E., Meyers, L.: The ascent of the abundant: how mutational networks constrain evolution. *PLoS Comput. Biol.* **4**, e1000110 (2008)
34. Greenbury, S.F., Johnston, I.G., Louis, A.A., Ahnert, S.E.: A tractable genotype-phenotype map modelling the self-assembly of protein quaternary structure. *J. R. Soc. Interface* **11**, 20140249 (2014)
35. Schaper, S., Louis, A.A.: The arrival of the frequent: how bias in genotype-phenotype maps can steer populations to local optima. *PLoS ONE* **9**, e86635 (2014)
36. Manrubia, S., Cuesta, J.: Evolution on neutral networks accelerates the ticking rate of the molecular clock. *J. R. Soc. Interface* **12**, 20141010 (2015)
37. Stich, M., Manrubia, S.C.: Topological properties of phylogenetic trees in evolutionary models. *Eur. Phys. J. B* **70**, 583–92 (2009)

Quantitative Analysis of Synthesized Nucleic Acid Pools

Ramon Xulvi-Brunet, Gregory W. Campbell, Sudha Rajamani,
José I. Jiménez, and Irene A. Chen

Abstract Experimental evolution of RNA (or DNA) is a powerful method to isolate sequences with useful function (e.g., catalytic RNA), discover fundamental features of the sequence-activity relationship (i.e., the fitness landscape), and map evolutionary pathways or functional optimization strategies. However, the limitations of current sequencing technology create a significant undersampling problem which impedes our ability to measure the true distribution of unique sequences. In addition, synthetic sequence pools contain a non-random distribution of nucleotides. Here, we present and analyze simple models to approximate the true sequence distribution. We also provide tools that compensate for sequencing errors and other biases that occur during sample processing. We describe our implementation of these algorithms in the Galaxy bioinformatics platform.

1 Introduction

Selection experiments in biochemistry are becoming increasingly important because they allow us to discover and optimize molecules capable of specific desired biochemical functions. Nucleic acid based selections are also becoming increasingly easy for researchers to carry out, as high-throughput sequencing and oligo library synthesis technologies become ever more reliable and affordable, and as new, novel methods of selection become available. Additionally, increased computational capa-

R. Xulvi-Brunet (✉)

Departamento de Física, Facultad de Ciencias, Escuela Politécnica Nacional, Quito, Ecuador

Department of Chemistry and Biochemistry, University of California, Santa Barbara, CA, USA

e-mail: ramon.xulvi@epn.edu.ec

G.W. Campbell • I.A. Chen

Department of Chemistry and Biochemistry, University of California, Santa Barbara, CA, USA

S. Rajamani

Indian Institute of Science Education and Research, Pune, India

J.I. Jiménez

Faculty of Health and Medical Sciences, University of Surrey, Guildford, UK

bilities make analyzing the vast amounts of data produced by selection experiments much more tractable.

There are many examples of useful functional molecules that result from in vitro selections. New aptamers (short nucleic acid sequences that bind a particular target with high affinity and specificity) are being discovered at a phenomenal pace, many of which use their tight and specific binding properties in novel and valuable ways. For example, the *Spinach2* aptamer [1] regulates the fluorescence of a fluorophore via binding; in many ways, this aptamer/ligand pair is more advantageous than, and is often used in place of, traditional GFP tagging. It is also becoming clear that functions other than binding can be selected for, as is the case with catalytic nucleic acids (such as ribozymes [2], deoxyribozymes [3], and aptazymes [4]), or structure-switching regulators (such as riboswitches [5]). In any case, we are only just beginning to tap the massive potential of in vitro selection, especially in terms of generating powerful and precise molecular tools.

Recently, however, another benefit of such selections has become evident. Regardless of the type of molecule being selected, the selection data yield an unprecedented view of how a particular function is distributed across sequence space. That is, given a large pool of unique sequences, we can investigate which of those particular sequences are capable of performing the selected function. Not only that, but often we can also see how minor variations of functional sequences impact the function, which gives us important insight regarding the optimization of such molecules. Additionally, we can improve our understanding of potential evolutionary pathways, as we can look for networks of functional molecules that can be traversed in small mutational steps. These factors provide the impetus for creating and analyzing molecular fitness landscapes, which allow functional information to be mapped directly to sequence information. Fitness landscapes assign a quantitative measure of fitness (i.e., how well a sequence performs under particular selection conditions) to individual sequences across sequence space. These landscapes provide a means to relate genotypes (the nucleic acid sequence) to phenotypes (the functional activity).

For those interested in molecular fitness landscapes, the main goal of selection experiments is to compare two populations of molecules, namely, an initial pre-selection population, and a subsequent post-selection population that exhibits enrichment for certain sequences. This comparison allows us to infer how selective pressure affects different particular sequences by looking at the bias in the sequence distribution that is introduced by the selection process. Thus, the particular sequences best suited to performing the selected function can be found. A more in-depth approach to selections involves taking into account fitness values associated with each individual molecule's ability to perform the selected function, from which we can construct molecular fitness landscapes that quantitatively map sequence to activity [6]. Creating such landscapes provides a means of investigating how properties of particular sequences give rise to biochemical function, as well as how different functional sequences are related.

This study is concerned with the proper quantification of selection experiments, which is crucial for creating accurate molecular fitness landscapes, but can also be useful if the experimental goal is simply to identify a few functional molecules.

Our goal is two-fold: first, we aim to show that selection experiments may be significantly biased in some cases, due to experimental difficulties intrinsic to the processes required to exhaustively identify and accurately quantify sequences in a selection pool. These biases become especially important when the objective of the selection experiment is constructing a fitness landscape. And second, we aim to offer a primary set of tools, or quantitative procedures, that may help overcome, or at least reduce, biases inherent to current selection methodologies. Ultimately, the goal is either to estimate fitness landscapes with a high degree of confidence and accuracy, or to provide statistical support in determining which molecules in a selection pool perform best, depending on the goals of the experimentalist.

2 The Undersampling Problem

The first step in quantitatively characterizing in vitro selection experiments is determining the initial and final experimental population distributions. When the total number of molecules is small, this can be done by finding the copy number of each unique sequence in both populations. However, when the total number of molecules in the selection pool is much larger than the number of molecules that we can experimentally count (due to limitations in sequencing technology, for example), which is precisely the case with many complex biochemical selection experiments, we necessarily need to estimate the initial and final sequence distributions in a different way.

In DNA or RNA selection experiments, for instance, whether selecting the most functional sequences or constructing a fitness landscape, we ideally start with an initial pool of molecules that represents all possible unique sequences approximately equally. Here, we assume (consistently with most selection experiments) that the sequence length is constant (or at least similar) across the selection pool. For example, let us assume that the length of the sequences in a given selection pool is 40 nucleotides. This means that the number of unique sequences that would comprise an ideal, initial pool is around $4^{40} \simeq 10^{24}$. Let us further assume that, in order to extract accurate information from the experiment, we require approximately 10,000 copies of each of the 4^{40} unique sequences in the initial pool. In this case, the total number of molecules that we would need to count in order to obtain the true distribution of molecules would be around 10^{28} . Unfortunately, current sequencing technology is only able to identify roughly 10^8 , a number which is many orders of magnitude smaller than even the number of unique sequences, 10^{24} . Thus, we can easily see that statistical undersampling creates a significant obstacle in measuring the true distribution.

To fully appreciate the magnitude of the problem, let's consider the following, more intuitive, example. Suppose we want to find the distribution of favorite TV channels among US TV viewers, assuming that there are only 300 channels in the US. We want to ask N viewers what their favorite channel is. If we pose this question to the entire US TV-watching population (very close to $3 \cdot 10^8$ people), we will recover the exact distribution. If we ask the question to $N = 10^7$ people, we will

likely get a reasonable approximation to the exact distribution (provided that the survey respondents are randomly selected without repetition). However, if we ask only $N = 1000$ people, we can easily imagine that the distribution we will recover will be quite distorted compared to the true distribution. And finally, if we ask only $N = 10$ people, the resulting distribution will almost certainly be meaningless. (Of course, if there were only a single sharp peak—e.g., if everyone’s favorite TV station was PBS—then even these low numbers would still provide a reasonable approximation of the true distribution, but in reality, we do not expect either TV preference or sequence abundance to be represented by such an extreme case.)

Unfortunately, in real biochemical selection experiments, undersampling poses a significant problem. Essentially, the number of observations we can make is infinitesimally smaller than even the number of all possible different outcomes. Thus, the undersampling problem becomes a very serious issue, especially when the goal of the experiment is to quantitatively map the corresponding fitness landscape. All we can do when dealing with this kind of extreme undersampling is to try to figure out a way to estimate the distributions of molecular populations, other than just counting the molecules that are reported by a sequencing device.

The approach we propose here is to create a model that quantitatively describes how an initial pool of sequences is synthesized, so that we can use the model to estimate the abundance of each unique sequence in the initial pool. The model must be able to describe the synthesis from a set of just a few parameters, which themselves must be able to be estimated from general statistical properties of the sequences present in the pool. Armed with such a model, we will be able to compute the probability of finding any sequence in the initial pool, and therefore we can estimate the initial molecular distribution.

3 Modeling the Synthesis of Pools

Today, pools of sequences are typically generated via solid-phase synthesis carried out by fully automated oligonucleotide synthesizers. Protected nucleoside phosphoramidites are coupled to a growing nucleotide chain attached to a solid support, forming the desired sequence. Modern protocol [7] controls this process sufficiently well, such that any desired proportion of G, A, C, and U or T (for RNA and DNA pools, respectively), on average, can be obtained for each position in the sequence. A protected nucleoside is first attached to a solid support (this will be the 3’ end of the final sequence). That nucleoside is then deprotected and exposed to a solution containing particular concentrations (depending on the desired final distribution) of the four bases, G, A, C, U or T, in the form of protected nucleoside phosphoramidites. (The solution may contain only a single base if the goal is to generate a particular predetermined sequence). The chemical protecting groups, in conjunction with controlled deprotection steps and phosphite linkage chemistry, ensure that only one nucleotide at a time can be attached to the growing chain. The synthesis cycle of deprotection, exposure to activated monomer, and phosphite linkage is repeated

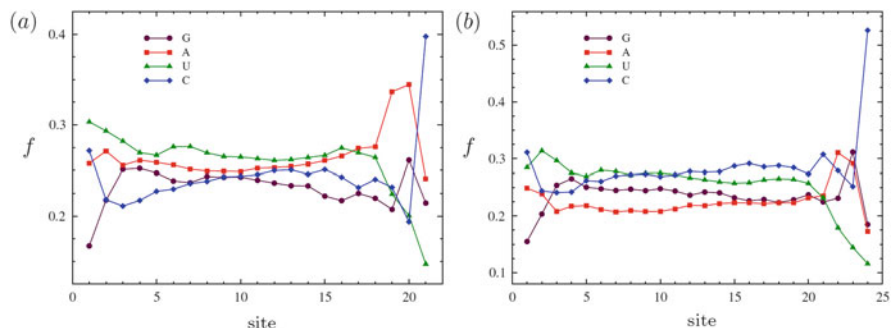


Fig. 1 Frequency, f , at which the four nucleotides G, A, C, U are found at each site of synthesized sequences. Panel (a) corresponds to a pool of synthesized sequences of length 21, while panel (b) corresponds to a pool of synthesized sequences of length 24. The pools were synthesized by different companies, and very probably, under slightly different conditions

as necessary, adding one base at a time, until sequences of a given length are formed. Finally, the oligonucleotide chains are released from the solid supports into solution, deprotected, and collected. Usually, commercial synthesis companies will then perform purification as specified by the end user before lyophilizing and shipping the oligonucleotide pool to the customer.

For some applications, the resulting oligonucleotides can be considered sufficiently random (e.g., if the goal is to identify an aptamer without regard to the fitness landscape). However, the pool generated by this procedure is not truly random. For a pool to be considered truly random, the probability that a given nucleotide is incorporated into the growing oligonucleotide chain must be $1/4$ at any step in the synthesis. Thus, if we compute the frequency at which a given nucleotide appears in a particular position of a synthesized sequence, the frequency must be $1/4$ for each nucleotide at each site. We measured this frequency for several synthesized pools and always found something qualitatively similar to what it is depicted in Fig. 1. We can see that the probability of finding a given nucleotide at a particular site is not $1/4$.

Similarly, the probability of finding any of the 16 possible nucleotide dimers within the pool of synthesized sequences should be $1/16$. The probability of finding any of the 64 trimers, 256 tetramers, 1024 pentamers, 4096 hexamers, and 16384 heptamers, etc., should be $1/64$, $1/256$, $1/1024$, $1/4096$, and $1/16384$, respectively. Once again, our investigations contradict the assumption of randomness (data not shown).

The fact that the synthesis process is not completely random should not be too surprising. Bearing the synthesis process in mind, the degree of randomness will depend on the relative concentrations of the nucleotides in solution. It is also likely to depend on the chemical reactivities among the different nucleotide species. In addition, the fact that nucleotides are not perfectly stable in water, and can break down spontaneously, could also play a role in the process by changing the

concentrations of nucleotides over time. Furthermore, we should not rule out the possibility that some of the nascent sequences may fold during synthesis, hindering the intended chemistry, despite the fact that the synthesis process is carried out at relatively high temperatures in order to avoid this as much as possible. Thus, we see that a number of factors may contribute to non-random nucleotide incorporation during synthesis.

Here, we will create quantitative models, based on some of the ideas listed above, and investigate which of these models is best able to replicate the statistical properties of synthesized sequences.

For the sake of completeness, we start with a perfectly random model, which we will call the 0-nucleotide model (so named because none of the nucleotide identities are differentiated in any way). As explained above, if the synthesis process follows this model, the probability of finding any monomer, dimer, trimer, etc., at any site in the synthesized sequences should be the same for all 4 monomers, 16 dimers, 64 trimers, etc.

The next model we investigate is the 1-nucleotide model, in which we only consider variation in the relative concentrations of the nucleotides in solution, C_G , C_A , C_C , and C_U (or C_T if working with DNA). The chemical reactivity between any two given nucleotides, according to this model, should be the same. Therefore, the probability that nucleotide i is incorporated into any forming sequence during the synthesis process is $\mathcal{P}(i) = C_i / \sum_{i'} C_{i'}$, where $i, i' \in \{G, A, C, U \text{ (or T)}\}$. Assuming, in addition, that the first nucleotide is attached to the platform with probability p_G , p_A , p_C , p_U (or p_T), for G, A, C, U (or T), respectively, we can easily calculate the probability of finding any sequence in the pool. We will not explain this calculation in detail here, since, as we will show, this 1-nucleotide model does not sufficiently explain the synthesis of nucleic acids.

We now introduce a set of models which consider, in addition to the concentrations C_G , C_A , C_C , and C_U (or C_T), different reactivities among nucleotides. The first model of this set, which we call the 2-nucleotide model, assumes that the probability for nucleotide $i \in \{G, A, C, U \text{ (or T)}\}$ to be incorporated in a forming chain depends on both nucleotide i and the nucleotide at the end of the forming chain, nucleotide $j \in \{G, A, C, U \text{ (or T)}\}$. The probability of incorporation of nucleotide i , conditional to nucleotide j , is $\mathcal{P}(i|j) = r_{ij} C_i / \sum_{i'} r_{i'j} C_{i'}$, where C_i and $C_{i'}$, $i, i' \in \{G, A, C, U \text{ (or T)}\}$, are again the concentrations of the nucleotides in solution, and r_{ij} are 16 chemical reactivity parameters that account for the likelihood of dimerization between each potential pairing of the 4 nucleotides G, A, C, and U (or T).

The second model of this set, the 3-nucleotide model, assumes that nucleotide i incorporates into a given nascent chain with probability $\mathcal{P}(i|j, k) = r_{ijk} C_i / \sum_{i'} r_{i'jk} C_{i'}$, where r_{ijk} describes the likelihood that a given nucleotide i attaches to a chain whose last and second-to-last nucleotides are, respectively, j and k . Again, $C_{i'}$, $i' \in \{G, A, C, U \text{ (or T)}\}$, are the nucleotide concentrations. The number of parameters associated with nucleotide reactivity in this model is $4^3 = 64$.

Similarly, the third model of the set, the 4-nucleotide model, assumes that the probability for nucleotide i attaching to a nascent sequence depends on nucleotide i

and the nucleotides in the last three sites of the sequence. The conditional probability is now given by $\mathcal{P}(i|j, k, l) = r_{ijkl}C_i / \sum_{i'} r_{ijkl}C_{i'}$, where j , k , and l are the last, second-to-last, and third-to-last nucleotides, respectively, of the forming sequence. The number of parameters associated with nucleotide reactivity in this model is $4^4 = 256$.

Following the same line of thought, the next models of the set, the 5-nucleotide model, 6-nucleotide model, etc., are those where we assume that nucleotide i attaches to a forming sequence depending on nucleotide i and the identity of nucleotides in the last 4, 5, etc., positions of the forming chain. This nested set of models can be arbitrarily extended, but selection-model criteria, such as the Bayesian Information Criterion and the Akaike Information Criterion, indicate that only the first models of this set are physically relevant. (This result is quite intuitive. It is clear that the more parameters a model has, the better we can fit the model to a given set of data. But, from a practical viewpoint, we are more interested in models with fewer parameters. Statistical Information Criteria allow us to identify the models that provide the best balance between fitting experimental data well and using the fewest number of parameters. Since the general k -nucleotide model requires at least the 4^k parameters associated with nucleotide reactivity, these statistical criteria penalized all k -nucleotide models with a relatively large k .)

In general, the above probabilities $\mathcal{P}(i)$, $\mathcal{P}(i|j)$, $\mathcal{P}(i|j, k)$, etc., can change at any step of the synthesis process, either by some external change in conditions (like temperature, for instance, which would directly affect the reactivities among nucleotides and, therefore, parameters r_{ij} , $r_{ij,k}$, etc), or by a change in the concentrations of nucleotides in solution, either intentionally or unintentionally (for instance, due to the natural degradation of nucleotides in solution). Here, we will assume that those probabilities, $\mathcal{P}(i)$, $\mathcal{P}(i|j)$, $\mathcal{P}(i|j, k)$, etc., do not significantly change in the actual synthesis process. In other words, our models will be based on the assumption that the synthesis conditions remain constant during all stages of the process.

The probabilities $\mathcal{P}(i)$, $\mathcal{P}(i|j)$, $\mathcal{P}(i|j, k)$, ..., themselves can actually be considered the parameters of our respective models. Indeed, for our goal, computing the abundance of a given unique sequence in a given pool, these probabilities are sufficient. Of course, if we are interested in studying the reactivities r_{ij} , r_{ijk} , etc., among nucleotides, they can be determined by using the above relationships between reactivities and probabilities.

Let us consider now, in some detail, the 2-nucleotide model. In order to compute the abundance of an arbitrary sequence, we need to know probabilities $\mathcal{P}(i|j)$, $\forall i, j \in \{G, A, C, U \text{ (or T)}\}$, and p_j , $j \in \{G, A, C, U \text{ (or T)}\}$ (the probabilities with which the first nucleotide is attached to the solid support). Knowing p_j and $\mathcal{P}(i|j)$, the probability that a sequence $i_1 i_2 i_3 \dots i_n$, $i_k \in \{G, A, C, U \text{ (or T)}\}$, is synthesized according to the 2-nucleotide model is exactly

$$P_{i_1 i_2 i_3 \dots i_n} = p_{i_1} \mathcal{P}(i_2|i_1) \mathcal{P}(i_3|i_2) \dots \mathcal{P}(i_n|i_{n-1}). \quad (1)$$

In practice, each probabilities p_j can be set to an arbitrary value with reasonable accuracy by manual mixing of resins. These probabilities are typically fixed at $p_j = 1/4$, $\forall j \in \{G, A, C, U \text{ (or T)}\}$ when the goal is to build random pools. However,

probabilities $\mathcal{P}(i|j)$ mainly derive from the chemical reactivities between different nucleotides pairs and, therefore, cannot be fixed by companies. If we want to compute $P_{i_1 i_2 i_3 \dots i_n}$, then we need to somehow estimate $\mathcal{P}(i|j)$.

This is done based on the relative fraction of dinucleotides found during sequencing of the synthesized pool. (Note that chemical reactivities may change as a function of temperature and other possible physical conditions, so we want to estimate $\mathcal{P}(i|j)$ for each individual pool, rather than estimating those probabilities once and assuming they will remain constant in future syntheses). In order to estimate these probabilities, we must first realize that the underlying probability distribution of unique sequences in a given pool always follows a multinomial distribution

$$\frac{n!}{n_{s_1}! n_{s_2}! \dots n_{s_{4L}}!} P_{s_1}^{n_{s_1}} P_{s_2}^{n_{s_2}} \dots P_{s_{4L}}^{n_{s_{4L}}}, \quad (2)$$

which gives the probability that a pool (of length L) with 4^L unique sequences, and n_{s_k} copies of each sequence s_k , can be generated when the total number of molecules created in the synthesis is n . Here, k is an index that runs from 1 to 4^L . Probability P_{s_k} , the probability that sequence s_k is created by the synthesis process is, according to our 2-nucleotide model, given by Eq. (1). (In Eq. (2) above, of course, $\sum_k n_{s_k} = n$ is satisfied).

Taking into account that in the 2-nucleotide model, all probabilities P_{s_k} , for any sequence s_k , depend at most on the 4 probabilities p_j and the 16 probabilities $\mathcal{P}(i|j)$, we can make use of the maximum likelihood estimation technique [8] to estimate $\mathcal{P}(i|j)$ and p_j . Since p_j can be obtained from the distribution of nucleotides at the 3' end of the sequences, we focus here on how $\mathcal{P}(i|j)$ can be estimated. The maximum likelihood estimation technique consists of maximizing the so-called log-likelihood function [8]. In our case, the log-likelihood function that needs to be maximized is

$$\Lambda = n_{s_1} \ln P_{s_1} + n_{s_2} \ln P_{s_2} + \dots + n_{s_{4L}} \ln P_{s_{4L}} + \sum_{j=1}^4 \lambda_j \left(\sum_{i=1}^4 \mathcal{P}(i|j) - 1 \right). \quad (3)$$

This expression, after some manipulation, can be represented as

$$\Lambda = \sum_{k=1}^4 m_k + \sum_{i=1}^4 \sum_{j=1}^4 d_{ji} \ln \mathcal{P}(i|j) + \sum_{j=1}^4 \lambda_j \left(\sum_{i=1}^4 \mathcal{P}(i|j) - 1 \right), \quad (4)$$

where m_k is the number of monomers of type $k \in \{G, A, C, U \text{ (or T)}\}$ present at the first site of the sequences of the actual pool, and d_{ji} is the number of dimers present in the sequences whose first monomer is j and second monomer is i , $i, j \in \{G, A, C, U \text{ (or T)}\}$.

In both Eqs. (3) and (4), the last term $\sum_{j=1}^4 \lambda_j \left(\sum_{i=1}^4 \mathcal{P}(i|j) - 1 \right)$ accounts for the constraints imposed by the fact that $\sum_{i=1}^4 \mathcal{P}(i|j) = 1$, $\forall j \in \{G, A, C, U \text{ (or T)}\}$, are the Lagrange multipliers that are used to deal with optimization problems

with constraints. (In addition, we should have added to both Eqs. (3) and (4), the term $\ln(n!) - \sum_k \ln(n_{s_k}!)$, if the goal was to take the logarithm of the probability distribution defined by Eq. (2). We dropped this additional term because it does not depend on $\mathcal{P}(i|j)$ and, therefore, it does not play any role in the log-likelihood maximization process [8].)

Once the log-likelihood function has been found, we can obtain $\mathcal{P}(i|j)$ by computing the following partial derivatives [8]:

$$\frac{\partial \Lambda}{\partial \mathcal{P}(i|j)} = \frac{d_{ji}}{\mathcal{P}(i|j)} + \lambda_j = 0 \quad (5)$$

$$\frac{\partial \Lambda}{\partial \lambda_j} = \sum_{i=1}^4 \mathcal{P}(i|j) - 1 = 0. \quad (6)$$

From these last equations we can easily obtain what we need, an estimation of the probabilities $\mathcal{P}(i|j)$ based on statistical properties of the actual sequences, in this case, based on d_{ji} :

$$\mathcal{P}(i|j) = \frac{d_{ji}}{\sum_{k=1}^4 d_{jk}}. \quad (7)$$

Let us now consider the 3-nucleotide model. We use it to estimate the abundance of each unique sequence in a synthesized pool by following a line of reasoning analogous to the 2-nucleotide model. We start with the multinomial distribution Eq. (2), and use the maximum likelihood estimation technique to estimate the 64 probabilities $\mathcal{P}(i|j, k)$. For the sake of brevity, we will not derive the log-likelihood formula used to estimate probabilities $\mathcal{P}(i|j, k)$; it is done exactly as described above for the 2-nucleotide model. It suffices to say that probabilities $\mathcal{P}(i|j, k)$ can be obtained by observing the frequency of trimers in the synthesized pool of sequences. (In this case, however, the 3-nucleotide model can be used to describe the synthesis only after the nascent chains contain two nucleotides. Thus, we should avoid counting those trimers containing the first or second monomers of each sequence when computing $\mathcal{P}(i|j, k)$. To compute the probability of incorporating the second nucleotide into a forming chain, we should necessarily use the 2-nucleotide model).

Following this logic further, we can estimate $P_{i_1 i_2 i_3 \dots i_n}$ for each sequence $i_1 i_2 i_3 \dots i_n$, provided that the best model to describe the synthesis is a k -nucleotide model, whatever the value of k may be.

Next, we will investigate which model best fits the statistical properties of a given pool. To do so, we use each of the different models (after estimating the corresponding probabilities $\mathcal{P}(i|j)$, $\mathcal{P}(i|j, k)$, etc., from an experimental sequence pool) to generate a large pool of sequences, *in silico*. Then, we count how often each monomer, dimer, trimer, tetramer, pentamer, etc., is found within the generated sequences; this allows us to compute the relative abundance of each n -mer in the theoretical pool. Next, we repeat the process for the actual, synthesized pool. We count the number of each n -mer present in the synthesized experimental pool

and calculate their relative abundances. Finally, we compare the theoretical and experimental results. The model that best matches the data is the most appropriate choice for subsequent data analysis.

We adopted the following specific procedure to validate the models. We divided the experimental sequences into two groups. One group was used to estimate, by counting the relevant n -mers, probabilities $\mathcal{P}(i|j)$, $\mathcal{P}(i|j, k)$, etc. In other words, the first group was used to infer the parameters from which the models are built. From the second group of sequences, we obtained the experimental n -mer ($n \in \{1, 2, 3, 4, 5, 6\}$) abundances, in order to compare them to the abundances computed from the respective models. In other words, the second group was used to assess how well the models fit the actual data.

As a final remark, when counting n -mers, whether for building or assessing the models, we always avoided including the first and last six nucleotides of the synthesized sequences. We will discuss this in more detail later.

Figure 2 depicts this comparison. In all panels of the figure, the x -axis represents the relative abundance of each of the 4 monomers, 16 dimers, 64 trimers, 256 tetramers, 1024 pentamer, and 4096 hexamers, all obtained from sequences in the actual, synthesized pool. The y -axis represents the same data, but obtained from computationally generated sequences according to the respective theoretical model. That is, each panel illustrates the correlation between the experimental and theoretical relative abundances of each of the n -mers, $n \in \{1, 2, 3, 4, 5, 6\}$. In the figure, we show the results for four of our models: the random model (panel *a*), the 1-nucleotide model (panel *b*), the 2-nucleotide model (panel *c*), and the 3-nucleotide model (panel *d*).

In such a graphical representation, a model is considered good if most of its points lie on, or near, the $x = y$ diagonal. Conversely, models displaying points far away from the diagonal line should be excluded from consideration. Figure 2 shows that both the random and the 1-nucleotide models do not explain with sufficient accuracy what really happens during synthesis. Thus, any final probabilities $P_{i_1 i_2 i_3 \dots i_n}$ computed based on these two models will be quite inaccurate. Better results are obtained with the 2-nucleotide model, and the reason is clear: This is the first of the n -nucleotide models that considers the different chemical reactivities of the nucleotides. Finally, panel *d* of the figure shows that the 3-nucleotide model best describes the data.

We analyzed the modeling power of the four models described above using several synthesized pools, and always found that the 3-nucleotide model was the best of the first four k -nucleotide models. We also fit the 4-, 5-, 6-, and 7-nucleotide models to each of the synthesized pools, and observed that none of them significantly improved the fit to the data. Moreover, when we used statistical information criteria, such as the Bayesian Information Criterion and the Akaike Information Criterion, to select the best n -nucleotide model, $n \in \{1, 2, 3, 4, 5, 6, 7\}$, we found that the 3- and 4-nucleotide models were best (depending on the particular synthesized pool), and that there was very little difference between them.

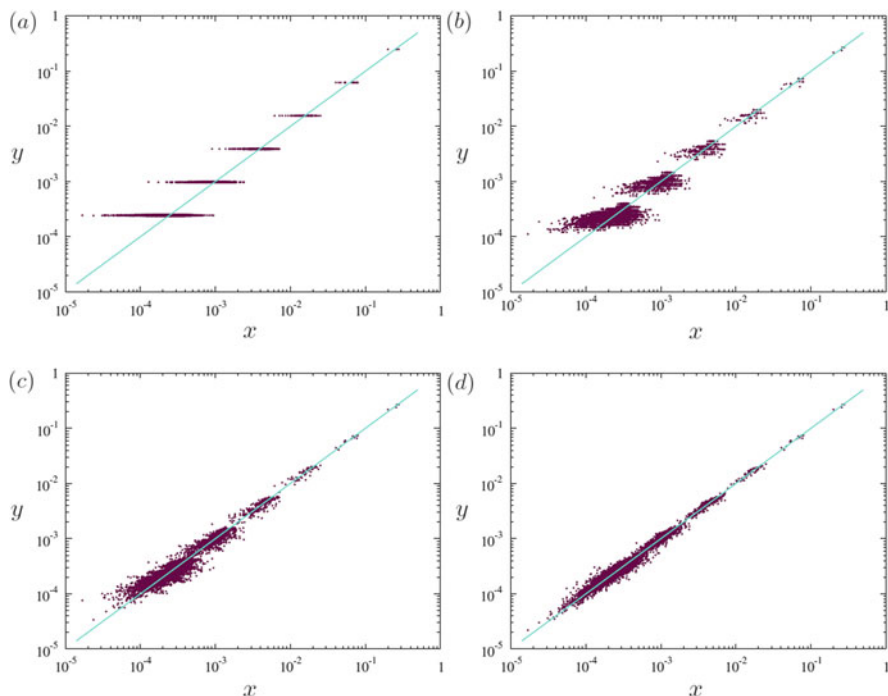


Fig. 2 Validation of models. The x -axis of all four panels represents the probability of finding each of the 4 monomers, 16 dimers, 64 trimers, 256 tetramers, 1024 pentamer, and 4096 hexamers, within the sequences of the actual, synthesized pool. The y -axis of the panels represents the probability of finding each of the 4 monomers, 16 dimers, 64 trimers, 256 tetramers, 1024 pentamer, and 4096 hexamers, within sequences computationally generated according to the respective theoretical model. Panel (a): random model. Panel (b): 1-nucleotide model. Panel (c): 2-nucleotide model. Panel (d): 3-nucleotide model

Our results therefore suggest that during sequence synthesis, the incorporation of a nucleotide depends on its species, as well as the species of the last two or three nucleotides that were previously incorporated.

Despite the fact that the 3- and 4- nucleotide models have been proven to model data better than the random and 1-nucleotide models (which, incidentally, are still commonly used in literature), it is important to remember that these models are still only approximations, and there is room for improvement with respect to their ability to perfectly describe the synthesis of sequences. In order to achieve an even more accurate description of how pools are synthesized, we contend that further research is necessary, specifically research that investigates aspects of synthesis that we have not included in our models here.

Before closing this section, we feel it is prudent to ask how the probabilities $P_{i_1 i_2 i_3 \dots i_n}$ predicted by the different models compare to one another. Figure 3 illuminates this issue. The figure reproduces the sequence probability distributions estimated for an experimental pool according to our first seven k -nucleotide models.

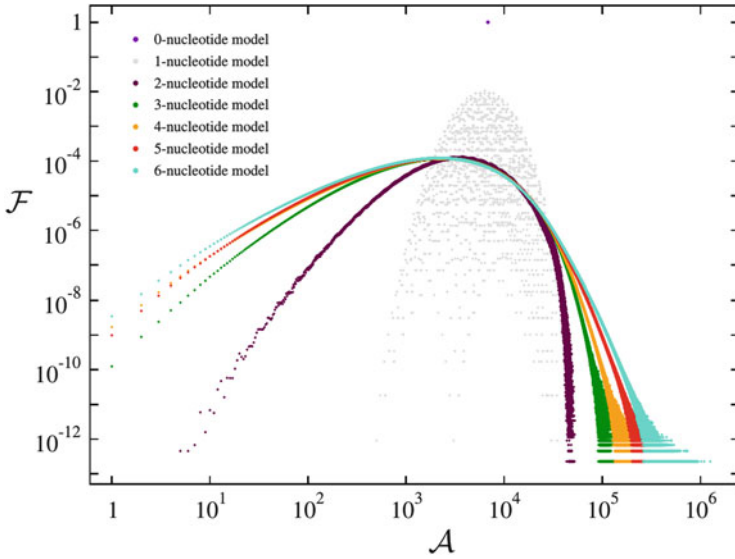


Fig. 3 Fraction of sequences $\mathcal{F}$ (computed according to different models) that are present in a synthesized pool with a certain abundance $\mathcal{A}$. The models considered are the seven first k -nucleotide models

The x -axis represents the abundance of each sequence in the pool; the y -axis describes the fraction of sequences in the pool that are present at a given abundance.

In the figure, the sequence probability distribution corresponding to the random model is only a single point. Obviously, if the correct synthesis model is the random one, then all sequences must have the same probability of appearing in the pool; that is, the number of copies of each unique sequence will be (approximately) the same. In this random case, the number of copies of each sequence is (approximately) $n/4^L$, where n is the total number of molecules in the pool and L the length of the sequences. Also, the probability that a unique sequence has (approximately) $n/4^L$ copies will be 1, since all sequences have the same abundance.

The probability distributions corresponding to the other models show that not all unique sequences are predicted to have the same abundance in the pool. In particular, for those n -nucleotide models with $n > 2$ (the more realistic ones), we see that some sequences appear hundreds of thousands of times in the pool, while others appear just a few times. It is true that the abundance of the vast majority of sequences is not that dissimilar, but the difference in the number of copies between the most and least abundant sequences can span several orders of magnitude. Note that this behavior strongly differs from the random model's predictions, namely that all sequences have (approximately) the same abundance $n/4^L$.

This significant difference between the frequencies of the most and least abundant sequences in a pool also plays an important role in determining when we can reasonably assume that all possible 4^L unique sequences are present. Assuming there

is no mutagenesis (for example, via PCR amplification), then if some sequences are initially absent, they cannot be selected upon by any selection experiment, even if they happen to be the best (most functional or selectable) sequences. Our results indicate that, when n is larger than 4^L by several orders of magnitude, all sequences tend to appear in the synthesized pool (despite the fact that some of them may appear just a few times and other hundreds of thousand of times). However, when n is larger than 4^L by just one or two orders of magnitude (or when $n \leq 4^L$), then the starting pool tends to contain only a fraction of all unique sequences; that is, a significant number of unique sequences may be missing.

Most experimentalists accept that synthetic pools are sufficiently random for their purpose, and therefore tend to assume that the number of copies of any sequence in the pool is approximately $n/4^L$. We have seen, however, that synthetic pools are often not completely random; the abundance of the different sequences in the pool may differ by as much as several orders of magnitude. To assume that a pool is truly random when, actually, it is not, may potentially lead to incorrect conclusions, depending on the goals of the experiment.

We created a suite of computational tools that, among other things, estimates the conditional probability parameters needed by our models, based on sequencing data for an initial, pre-selection experimental pool. These tools are based on the 3-nucleotide model, given that it is the simplest model to describe experimental pools with reasonable accuracy, and based on this model, they can also compute the predicted abundances for any arbitrary sequence that may appear in the pool. These tools were incorporated into the Galaxy bioinformatics platform (see <http://galaxyproject.org> for general information about Galaxy), in order to facilitate ease of access, program automation, and an intuitive user-friendly interface. Researchers interested in using these tools are free to visit <http://galaxy-chen.cnsi.ucsb.edu:8080/> or contact us for more information.

4 Sequencing Biases

In the preceding section, we discussed the problem of estimating the initial frequency of specific sequences and offered a practical solution. Now we address a second issue, that of biases introduced by the process of sequencing itself.

In any selection experiment, we need to *observe* the sequences that are present in a pool both before and after the selection process. But, how exactly can we *observe* the sequences accurately, given their nanoscopic nature? In order to *see* them, selection experiments make use of various high-throughput sequencing (HTS) protocols, but all of them, in essence, follow the general theme outlined below.

- 1) 5' and 3' adaptor ligation: If needed, small segments of RNA or DNA (depending on the type of nucleic acid comprising the pool) with a particular known sequence are attached to the 5' and 3' ends of the sequences in the pool. These segments are needed as constant regions, designed to be complementary to PCR

primers, or to couple the sequences to a slide for sequencing. This step may not be necessary in all protocols, such as those in which the constant regions were already incorporated into the selection library.

- 2) Reverse transcription: When the nucleic acid is RNA, reverse transcription is carried out to convert the RNA sequence into DNA. This is done because DNA sequences are inherently more stable, and therefore degrade less easily during analysis. Also, DNA is required for currently used high-throughput sequencing methods.
- 3) *PCR* amplification: The total number of molecules that can be experimentally manipulated is often somewhat limited, so it is sometimes necessary to amplify each of the molecules we want to observe. In addition, PCR is often used to introduce additional constant regions required for sequencing, if they are not already present in the selection library. Ideally, all molecules would be increased by the same factor, such that the initial sequence distribution is not artificially distorted. However, several investigations that we carried out to investigate this issue show that, despite the fact that most sequences are indeed multiplied by approximately the same factor, a small percentage tend to be multiplied much more or much less than expected (even by an order of magnitude). We do not show the results of these investigations here since other groups have already shown the existence of PCR artifacts and bias [9]. Therefore, our recommendation is to avoid amplification as much as possible.
- 4) Sequencing: Finally, DNA sequences consisting of known constant regions flanking the initially randomized region are sequenced. There are several mechanisms by which this takes place, depending on the type of sequencing, but ultimately, each nucleotide of every sequence submitted to the sequencer is identified, and a list of all sequences in the sample is reported.

We will deal with steps 1 and 2 in the next section. Step 3 will not be considered here, because it is sometimes avoidable [6, 10], and the biases of PCR can be minimized by reducing the number of cycles. In practice, the bias itself is difficult to disentangle from bias and selection introduced at other steps; thus, we effectively consider the PCR amplification biases to be a component of the fitness of a sequence during selection. Here, we will focus on the quantitative issues that arise from the inaccuracies in the sequencing process. The specific sequencing protocol we investigated is the widely used Illumina protocol [11].

The Illumina reading mechanism infrequently, but occasionally, misreads nucleotides in different positions of a given sequence. Several studies indicate that the probability of a nucleotide being misread mostly depends on both its position in the sequence and the identity of its neighbors (please see references [12–15] for advanced studies on this topic). But, as a first approximation, it seems to be reasonably accurate to assume that the probability of misreading is roughly constant. This constant probability is gradually being reduced as sequencing technologies improve, but as of yet, it is not zero. For instance, current Illumina technology has an accuracy of about 99.5 % (up from 98–99 % a few years ago).

The fact that the reading process is not completely accurate tends to distort measurements. For some types of biostatistical studies, those distortions are irrelevant or easy to overcome. For example, the accuracy can be improved (squared) by reading the same segment in opposite directions (paired-end reading). For others, such as mutation studies or measuring fitness landscapes, these distortions may significantly affect the results; in mutations studies, by masking the probabilities of mutation, and in selection experiments, by reducing the resolution of fitness peaks in a landscape.

In order to quantitatively understand how this distortion affects our observations, let us first consider a simple example. Assume that the distribution of sequences that we want to measure consists of one unique sequence of length L that appears in a pool N times. Let us assume that the probability of reading a nucleotide correctly, p , is constant. If $p < 1$, then there is a non-zero probability that the sequence we think we are observing differs from the sequence that is actually in the pool. Smaller values of p , lead to larger numbers of erroneous unique sequences generated by misreadings (or mutations). To quantify the distortion introduced by the sequencing device, we need to answer the two following questions: (1) What is the probability of misreading n nucleotides in a sequence of length L ? And (2) What is the probability of finding a particular sequence containing n misreads?

To answer the first question, we must remember that, under our assumptions, the probability of misreading a nucleotide at a certain position is independent of the success or failure in reading nucleotides at other positions. This condition naturally leads us to the binomial probability distribution. Therefore, the probability of misreading n nucleotides is given by

$$\mathcal{P}(n) = \frac{L!}{n!(L-n)!} p^{L-n} (1-p)^n . \quad (8)$$

To answer the second question, we first need to know the probability that each of the other three possible nucleotides is reported when an erroneous read occurs. Again, as a first approximation, we can make the additional assumption that the probability is $1/3$; that is, any of the other three nucleotides are equally likely to be falsely reported. Under this extra assumption, any possible sequence containing n mutations has the same probability of occurring. Therefore, to answer the second question above, we just need to know how many different, unique sequences containing n mutations are possible. The number of different, unique sequences that are one mutation off from the original sequence is $3L$. The number of unique sequences containing two mutations is $9L(L-1)/2$; the number with three mutations is $27L(L-1)(L-2)/6$. The total number of unique sequences containing n mutations is given by

$$\mathcal{N}(n) = \frac{3^n L!}{n!(L-n)!} . \quad (9)$$

Therefore, the probability of finding a particular sequence at distance n from the original is

$$P(n) = \frac{\mathcal{P}(n)}{\mathcal{N}(n)} = p^{L-n} \left(\frac{1-p}{3} \right)^n. \quad (10)$$

Once we understand the quantitative effects of sequencing failures, we can ask how to overcome them. Several algorithms may be devised by which these unwanted sequencing effects might be reverted. Here we present one such algorithm, which, despite not being the most rigorous, is very easy to implement, relatively fast, and capable of yielding reasonable results.

The algorithm is based on the simple fact that, if $p < 1$, the output of any sequencing process—the distribution of sequences that we observe—will always differ from the input—the real distribution of sequences. That is, the observed abundance of a given sequence i will, in general, differ from its true abundance. The relationship between these contrasting abundance values is given by the following formula:

$$\langle n_i^{ob} \rangle = n_i^{re} P(0) + \sum_j n_j^{re} P(d_{ij}), \quad (11)$$

where n_i^{ob} is the observed abundance of sequence i , n_i^{re} is the true abundance of i , and $P(d_{ij})$ is the probability given by Eq. (10), where d_{ij} is the distance between sequences i and j . Equation (11) says that, for any sequence i , the expected number of copies we are likely to observe is equal to the true number of copies multiplied by the probability that it is correctly read, plus the number of erroneous variants of other sequences that are identified as sequence i . Note that it is possible to (1) observe a sequence that was not present in the input distribution—because it is a mutational variant of an initial sequence—and (2) miss a sequence that was actually present—which could happen if the number of copies of an initially present sequence was very low and every copy was misread. In the latter case, we cannot tell, just by looking at the output distribution, whether or not the sequence was present in the true distribution. Fortunately, that situation is highly improbable unless the sequence appears just one or two of times, in which case it is probably not very relevant for the selection process.

The algorithm has two main stages:

- 1) First, we sort the sequences according to their observed abundance. This results in a list of sequences such that the first sequence is the most abundant, the second is the next most abundant, and so on.
- 2) The second stage of the algorithm is an iterative procedure that performs the following steps on each item in the sorted list described above. At each iteration i , the algorithm estimates the true abundance of the i th sequence in the list as described below:

2.1) For each sequence $j, j \neq i$, the quantity

$$c_{ji} = \text{int} \left(\frac{n_j^{ob}}{p^L} \left(\frac{1-p}{3} \right)^{d_{ij}} p^{L-d_{ij}} \right) \quad (12)$$

is computed. $n_j^{ob}/P(0) = n_j^{ob}/p^L$ can be seen as the expected true abundance, n_j^{re} , of sequence j , provided that it is sufficiently different from other sequences. This is certainly an approximation, since sequence j might actually have some real neighboring sequences. However, this approximation, $n_j^{re} \simeq n_j^{ob}/p^L$, is actually the key assumption in this heuristic algorithm. $((1-p)/3)^{d_{ji}} p^{(L-d_{ji})}$ gives the probability that sequence j mutates into sequence i . Finally, $\text{int}(x)$ is just the nearest integer function, which converts any real number x into its nearest integer. Thus, the quantity c_{ji} is an approximation to the expected fraction of the observed abundance of sequence i that arises due to misreads of sequence j .

- 2.2) The quantity $c_i = \sum_j c_{ji}$ is computed. To a first approximation, c_i gives the expected increase in abundance of sequence i due to erroneous readings of the rest of the sequences.
- 2.3) The observed abundance n_i^{ob} is corrected by subtracting c_i from it; thus, we can write $n_i^c = n_i^{ob} - c_i$. Effectively, this step simply reduces the abundance of n_i^{ob} by the expected contribution from all other sequences of the distribution.
- 2.4) If the corrected abundance n_i^c satisfies $n_i^c \geq 0$, then n_i^c is updated by the next formula: $n_i^u = \text{int}(n_i^c/p^L)$. Note that the effect of this final correction simply converts the copies of sequence i that were lost due to misreading back into sequence i . If $n_i^c \leq 0$, then final updated n_i^u is simply set to zero, $n_i^u = 0$. The role of function $\text{int}(x)$ is to yield a final integer number of copies for each sequence.

Now, we proceed with some tests to assess the quality of this algorithm.

In our first example, let us assume that the output of a given experiment consists of exactly 200,000 copies of a unique sequence, for example, TAAGGCTATGAAGAGATACTG. We can simulate the effect of misreading nucleotides by imposing that nucleotides in each position of any sequence can be replaced with any of the other three nucleotides with probability $q = 0.01$. Such a simulation will show results like the ones we describe next (corresponding to a particular simulation): 1833 different, unique sequences were generated. Sequence TAAGGCTATGAAGAGATACTG has 162,056 copies (instead of the original 200,000); $63 = 3 \cdot 21 = 3 \cdot (\text{length of TAAGGCTATGAAGAGATACTG})$ unique first neighbors to the original sequence have about 540 copies each; and 1769 unique sequences will have only a few copies each, most of which differ from the original by two nucleotides; and the rest, sequences that differ by three or more nucleotides. From those 1769 new, unique sequences, 752 appear only once, 520 appear twice, 283 appear three times, 132, four times, 55, five times, 20, six times, and 5, seven times.

The goal of our algorithm is to use this new distribution of sequences to recover the original distribution, containing 200,000 copies of the original sequence TAAGGCTATGAAGAGATACTG. Of course, we do not expect the algorithm to recover exactly one sequence with 200,000 copies (because of the stochasticity of the process), but we do expect something close to that. For the particular simulation described in the paragraph above, this is the output of our algorithm: there are now 701 different, unique sequences. Our original sequence, sequence TAAGGCTATGAAGAGATACTG, has 199,981 copies (very close to 200,000), and the rest of the unique sequences appear just a few times. Concretely, 487 sequences appear just once, 133 appear twice, 55 appear three times, 20 appear five times, and 5 appear six times. All 63 sequences that appeared around 500 times before the correction now have significantly reduced copy numbers (6 copies or less). The total number of molecules is 201,084, only slightly larger than 200,000. This degree of recovery is typical for the algorithm. It is capable of closely approximating the correct number of copies of the original sequences, and it removes practically all the significantly abundant, erroneous sequences. It also reduces the total number of unique erroneous sequences (in this case, from 1833 to 701), and slightly increases the total number of copies (from 200,000 to 201,084, in this case).

For our second example, suppose that our experimental output consists of the following seven sequences: TAAGGCTATGAAGAGATACTG, AAAGTGCAGACAGGCCTGGTC, CACCCATGCCTCGACCATCCT, CACGGCACTGTACATTGGTTT, TCGCACCTTTCCGGCCATTG, AAATCGTTTCGAAAGCGCCGAT, and AATTTCGAGCCGTGAGCGTATG. And let us suppose that there are 200,000, 185,794, 180,592, 175,872, 163,603, 95,617, and 1042 copies, respectively. The total number of molecules, in this case, is 1,002,520. Now, we simulate the “sequencing bias effect”, again with $q = 0.01$, and get the following: the number of copies of each real sequence is 162,102, 150,470, 146,304, 142,570, 132,241, 77,543, and 819, respectively. In addition, 378 unique artifactual sequences appear, whose abundances range from 227 to 592 copies, together with 9997 other unique sequences that appear a only few times (most of them, just once or twice). After applying our algorithm to this new, false distribution, we obtain: for the seven original sequences, copy numbers of 200,038, 185,672, 180,527, 175,922, 163,192, 95,686, and 1011, respectively. All 378 fictitious sequences whose abundances were between 227 and 592 copies have been removed. We still have a number of unique sequences that appear a few times, most of them just once or twice, but the number has been reduced to 3878. In total, we now have 1,007,650 molecules, instead of the original 1,002,520 (again, slightly higher than the actual starting number of molecules, but only around a 0.5 % increase). Taken together, (1) we have recovered decent copy numbers for the seven real, original sequences, (2) all the erroneous unique sequences with relatively large abundances are gone, and (3) the number of the erroneous unique sequences that still remain appear only a few times.

After this discussion on sequencing bias, it should be clear that the resolution with which we can *observe* any experimental distribution of sequences is finite, and that caution must be employed when drawing conclusions about sequences with very low copy numbers. If those sequences that (apparently) appear only a few times

in the sequencing output happen to be neighbors of highly abundant sequences, there is a strong probability that they may be artifacts of the sequencing process, sequences that were not actually present in the experimental pool. We should only assume that they may be real when they are many mutations away from the most abundant sequences. Despite its minor inaccuracy, our algorithm makes great strides toward improving the resolution with which we observe experimental pools.

Since we will never be able to make *observations* with infinite resolution due to the inherent stochasticity of sequencing biases, it seems reasonable to ignore those sequences that are output by the sequencer only a few times. (Indeed, even if some of those sequences happen to be real sequences from the experimental pool, they are likely sequences of lower selective significance than those that appear many times, especially during later rounds of selection). Applying the correcting algorithm and ignoring sequences of very low abundance, aside from being relatively harmless to the analysis of a selection experiment, has the important consequence of eliminating many of the minor sequences output from the sequencing device, which notably reduces experimental noise and the computational time required for subsequent data analysis.

5 Sequence Ligation and Reverse Transcription

We now address the other noteworthy feature seen in Fig. 1. The figure shows that the distribution of unique sequences is not random; it also displays peculiar behavior at both terminal regions of the displayed curves, indicating an exceptional distortion of monomer frequency at the first and last approximately five sequence positions. Additionally, the qualitative features of this remarkable distortion are the same in any synthesized pool, regardless of the length of synthesized sequences or which company has performed the synthesis.

This errant behavior should not be too surprising. Biochemists have known for several decades that the efficiency of ligating adaptors to RNA or DNA sequences strongly depends on the nucleotides at the terminal regions of the sequence to which the adaptor is ligated. For instance, England and Uhlenbeck [16] and Middleton et al. [17], in the 1970s and 1980s, respectively, showed that sequences containing GC, AC, and GA at the 3' end, and TT, TA, and AT at the 5' end were observed more frequently after ligation, while sequences with UU at the 3' end and GG at the 5' end were observed less frequently.

Therefore, the distortion found at ends of the curves in Fig. 1 seems to be consistent with the idea that adaptors ligate onto different unique sequences with different probabilities. Also, we must keep in mind that the reverse transcription process, which is initiated by a *primer* annealing to the 3' end of a sequence, may also play a role in these distortions. In fact, we do not know the exact mechanism that makes sequences which begin and end with a particular nucleotide combination more or less likely to be selected after adaptor ligation and reverse transcription. All we can claim, according to Fig. 1 (and other pools that we have analyzed), is that

(1) the effect is real and quantitatively important, and (2) it seems to depend on the particular terminal pentamers found in the sequences.

This ligation/reverse transcription effect may significantly hinder the estimation of the true sequence distribution in a given synthesized pool. This effect is not important when estimating the distribution of an initial (pre-selection) pool, because, in this case, the distribution can be estimated using the models developed in Sect. 2. However, the effect becomes important when we want to estimate the distribution of a pool that has already undergone selection, since, in this case, we have no way of estimating the probability distribution without knowing how the selection works, which is precisely what we aim to uncover.

So, how can we correct for the inaccuracies that result from this ligation/reverse transcription effect in post-selection pools, such that we might recover the real distribution?

The approach we propose here is as follows:

- 1) First, examine the relative abundance of all pre-selection sequences reported by the sequencing device and correct the observed abundance for sequencing errors (e.g., using the method described in the previous section).
- 2) Classify the sequences of the initial pool according to the first and last 5 nucleotides located at the terminal regions of each sequence. That is, arrange the sequences into the $1024 \cdot 1024 = 1,048,576$ potential sequence classes, such that all sequences in each class have the same first and last 5 nucleotides. (In view of the adaptor ligation and reverse transcription effects, it should now be clear why, in previous sections, we excluded the first and last six nucleotides of each sequence when estimating the various probabilities/parameters used by our models). Then, for each class l , compute the total abundance of that class; that is, sum the abundances (already corrected for sequencing in step 1) of all sequences that belong to the class. Finally, divide the total abundance of each class l by the total number of observed molecules in order to get the *observed* probability p_l^{ob} that a pre-selection sequence belongs to class l .
- 3) Making use of our previously developed models (section 3) and the sequencing corrected abundances (from step 1), compute the true abundance of the sequences in the initial pool. Next, sort them into the same 1,048,576 sequence classes described above and sum the true abundances of all pre-selection sequences that belong to each class l . Divide the total abundance of each class l by the total number of molecules to obtain the probability p_l^{true} that a pre-selection sequence belongs to class l .
- 4) For each sequence class l , determine the correcting factor $f_l = p_l^{true} / p_l^{ob}$.
- 5) Examine the relative abundance of all post-selection sequences reported by the sequencing device, and correct the observed abundance of each sequence for sequencing errors.
- 6) Adapter ligation and reverse transcription biases of post-selection data can be corrected as follows: Identify the class to which a post-selection sequence belongs; this determines which correcting factor f_l should be used. Then,

multiply the *observed* abundance of the sequence by the appropriate f_i . This product yields the corrected abundance of the sequence.

Another issue we should discuss here is the size of the correcting factors, f_i . Most studies ignore these types of corrections, which is tantamount to assuming that all f_i values are 1, or very close to 1. Our results indicate that most of these factors do indeed have values that are close to 1; however, they also indicate that some factors may differ from 1 by more than an order of magnitude. We saw some factors with values close to 10, and others close to 0.1. This means that, for some fraction of sequences, we might expect inaccuracy in abundance values of up to two orders of magnitude, if these correction techniques are not applied.

Finally, we want to point out that the precise value of these correction factors, f_i , depend on the particular conditions (temperature, etc.) in which adaptor ligation and reverse transcription are carried out. Therefore, since f_i are, to some extent, dependent on a particular experiment, we recommend that the f_i factors are recalculated every time adaptor ligation and reverse transcription are done. In our suite of tools on the Galaxy platform, the reader can find programs that compute the true abundances after calculating f_i . These tools are fully automated, and require minimal user input, aside from setting relevant run-time parameters (such as the minimum and maximum allowed sequence lengths). They have been extensively tested and shown to apply reasonable and realistic corrections to sequencing data. They also create and analyze fitness landscapes based on the corrected sequencing data (which, although we have not discussed it in detail here, is a major feature of the software).

6 Conclusions

We can conclude that *observing* the true distributions of sequences from selection experiments, both before and after selection occurs, is an important, non-trivial task that is necessary for deriving quantitative insights about the evolution of sequences during selection experiments.

The first thing to remember is that purportedly random sequence pools, on which selection experiments are based, are not completely random. It is true that they may be nearly random, but this is not sufficient if the goal is to accurately quantify the evolution of a sequence pool over the course of a selection experiment. To this end, we provide rigorous evidence that most synthesized pools tend to contain sequences that appear several orders of magnitude more or less often than others. This may be a serious issue if the goal of the study is to draw quantitative conclusions about the relative fitness of sequences during selection.

The second thing to consider is that selection experiments are not free from bias. In order to make observations, we may need to attach adaptors and reverse transcribe sequences before submitting them for sequencing. We have seen that adaptor ligation and reverse transcription do not affect all sequences equally, but rather, in

some cases, affect different sequences very differently. Also, sequencing errors tend to blur the real distribution of sequences. Both bias types may significantly distort the quantitative results of a selection experiment.

These issues can be overcome by better characterizing our observations, quantitatively speaking. To do that, it is crucial to realize that biochemical selection experiments aimed at constructing global fitness landscapes suffer from significant undersampling problems. Here, we propose avoiding such problems by constructing reliable models that describe how sequence pools are synthesized, which can then be used to estimate the abundance of any arbitrary sequence in any initial, pre-selection pool. These initial abundance calculations are crucial to rigorously assigning a fitness value to a given sequence. These models, as shown in Sect. 5, can also be used to correct biases which result from adaptor ligation and reverse transcription, improving the precision of our observations. Similarly, the construction of a reliable model that quantitatively describes the biases associated with sequencing errors, helps to design algorithms aimed at correcting those biases, again furthering our observational precision.

Here, we have proposed some simple models to describe relevant biochemical processes, models that we consider first attempts in quantitatively describing those processes. As researchers perform more selection experiments, we hope that more accurate models will be developed, such that we may increasingly trust the analysis of those experiments to reflect reality.

References

1. Strack, R.L., Disney, M.D., Jaffrey, S.R.: A superfolder Spinach2 reveals the dynamic nature of trinucleotide repeat-containing RNA. *Nat. Methods* **10**, 1219–1224 (2013)
2. Dolan, G.F., Akoopie, A., Muller, U.F.: A faster triphosphorylation ribozyme. *PLoS ONE* **10**, e0142559 (2015)
3. Chan, L., Tram, K., Gysbers, R., Gu, J., Li, Y.: Sequence mutation and structural alteration transform a noncatalytic DNA sequence into an efficient RNA-cleaving DNzyme. *J. Mol. Evol.* **85**, 245–253 (2015)
4. Sefah, K., Phillips, J.A., Xiong, X., Meng, L., Van Simaey, D., Chen, H., Martin, J., Tan, W.: Nucleic acid aptamers for biosensors and bio-analytical applications. *Analyst* **134**, 1765–1775 (2009)
5. Martini, L., Meyer, A.J., Ellefson, J.W., Milligan, J.N., Forlin, M., Ellington, A.D., Mansy, S.S.: In vitro selection for small-molecule-triggered strand displacement and riboswitch activity. *ACS Synth. Biol.* **4**, 1144–1150 (2015)
6. Jimenez, J.I., Xulvi-Brunet, R., Campbell, G.W., Turk-MacLeod, R., Chen, I.A.: Comprehensive experimental fitness landscape and evolutionary network for small RNA. *Proc. Natl. Acad. Sci. U. S. A.* **110**, 14984–14989 (2013)
7. Ellington, A., Pollard Jr., J.D.: Synthesis and purification of oligonucleotides. In: *Current Protocols in Molecular Biology*, Chap. 2, Unit 2.11. Wiley, New York (2001)
8. Edwards, A.W.F.: *Likelihood*. Johns Hopkins University Press, Baltimore (1992). Expanded edition
9. Acinas, S.G., Sarma-Rupavtarm, R., Klepac-Ceraj, V., Polz, M.F.: PCR-induced sequence artifacts and bias: insights from comparison of two 16S rRNA clone libraries constructed from the same sample. *Appl. Environ. Microbiol.* **71**, 8966–8969 (2005)

10. Aird, D., Ross, M.G., Chen, W., Danielsson, M., Fennell, T., Russ, C., Jaffe, D.B., Nusbaum, C., Gnirke, A.: Analyzing and minimizing PCR amplification bias in Illumina sequencing libraries. *Genome Biol.* **12**, R18 (2011)
11. Bennett, S.: Solexa Ltd. *Pharmacogenomics* **5**, 433–438 (2004)
12. Meacham, F., Boffelli, D., Dhahbi, J., Martin, D.I.K., Singer, M., Pachter, L.: Identification and correction of systematic error in high-throughput sequence data. *MBC Bioinform.* **12**, 451 (2011)
13. Nakamura, K., Oshima, T., Morimoto, T., Ikeda, S., Yoshikawa, H., Shiwa, Y., Ishikawa, S., Linak, M. C., Hirai, A., Takahashi, H., Altaf-Ul-Amin, Md., Ogasawara, N., Kanaya, S.: Sequence-specific error profile of Illumina sequencers. *Nucl. Acids Res.* **39**(13), e90 (2011)
14. Dohm, J.C., Lottaz, C., Borodina, T., Himmelbauer, H.: Substantial biases in ultra-short read data sets from high-throughput DNA sequencing. *Nucl. Acids Res.* **36**(16), e105 (2008)
15. Schirmer, M., Ijaz, U.Z., D'Amore, R., Hall, N., Sloan, W.T., Quince, C.: Insight into biases and sequencing errors for amplicon sequencing with the Illumina MiSeq platform. *Nucl. Acids Res.* **43**(6), e37 (2015)
16. England, T.E., Uhlenbeck, O.C.: Enzymatic oligoribonucleotide synthesis with T4 RNA ligase. *Biochemistry* **17**, 2069–2076 (1978)
17. Middleton, T., Herlihy, W.C., Schimmel, P.R., Munro, H.N.: Synthesis and purification of oligoribonucleotides using T4 RNA ligase and reverse-phase chromatography. *Anal. Biochem.* **144**, 110–117 (1985)

Non-linear Dynamics in Transcriptional Regulation: Biological Logic Gates

Till D. Frank, Miguel A.S. Cavadas, Lan K. Nguyen, and Alex Cheong

Abstract Gene expression relies on the interaction of numerous transcriptional signals at the promoter to elicit a response—to read or not to read the genomic code, and if read, the strength of the read. The interplay of transcription factors can be viewed as nonlinear dynamics underlying the biological complexity. Here we analyse the regulation of the cyclooxygenase 2 promoter by NF- κ B using thermostatistical and quantitative kinetic modelling and propose the presence of a genetic Boolean AND logic gate controlling the differential expression of cyclooxygenase 2 among species.

1 Introduction

The unravelling of the DNA sequence of the human genome [1, 2] has been a remarkable milestone in the scientific world, and has provided us with a book of knowledge from which we are still only deciphering the contents of. However, knowing the sequence of genes is only the beginning. The cell appears to know how to read the genome and crucially what to read and when to read portions of it, through integration of signals coming from both inside and outside the cell. This complexity of signal inputs is mirrored by the multitude of potential outputs from reading the genetic code. In the light of these input/output signalling processes, we can view DNA sequences and its regulation as electronic operations based on Boolean logic gates [3]. Through mathematical modelling of Boolean circuits and logic gates at molecular level using either thermostatistical or ODE-based

T.D. Frank

Center for the Ecological Study of Perception and Action, University of Connecticut, Storrs, CT 06269, USA

M.A.S. Cavadas

Instituto Gulbenkian de Ciencia, Rua da Quinta Grande, 2780-156 Oeiras, Portugal

L.K. Nguyen

Systems Biology Ireland, University College Dublin, Dublin 4, Ireland

A. Cheong (✉)

Life and Health Science, Aston University, Birmingham B4 7ET, UK

e-mail: a.cheong@aston.ac.uk

approaches, it is possible to gain insights into the rules and mechanisms by which a subset of the genes is selectively expressed in each cell—for example, how does a signalling event trigger a specific gene to be active? Furthermore, our understanding of Boolean logic gates can be applied for other uses, namely the encoding and archiving of digital information into synthetic DNA [4] and the construction of fuzzy logic biological computer [5]. Soon, there will be equivalences between electronic and biochemical operations.

1.1 *The Genome*

An organism's DNA encodes the necessary information required to construct its cell. Knowing the DNA sequencing of a bacterium (a mere few million nucleotide [6]) or human (a few billion nucleotides [1, 2]) does not mean that we can rebuild a bacterial or human cell, although attempts are in progress to construct synthetic cell [7]. However understanding the intricacies of the base pairs forming the genome can be akin to the discovery of the Rosetta stone: it is only the first steps towards learning how to read and write the genetic code. Here we will review the basics of gene regulation.

1.2 *Transcription Factors Promoter, Activators and Repressors*

Transcription factors (TF) are important proteins involved in the control of gene networks. They are generally described as either activators which increase transcriptional activity or as repressors which decrease the activity. The interaction between transcription factors and DNA is integral to the regulation of transcription. Thus the ability to predict and identify their binding sites in the promoter region of the gene of interest is crucial to understanding the details of gene regulation and for inferring regulatory networks. From experimental data, a set of validated transcription factor binding sites (TFBSs) for a given TF can be constructed to explore the binding preference of the TF. These DNA sequences are aligned and the occurrence of each nucleotide at each position is noted and scored to produce a consensus TF binding sequence. For example, the consensus binding motif of the ubiquitous transcription factor Nuclear factor kappa B ($\text{NF}\kappa\text{B}$) is described as GGG RNN YYC C, where G is guanine, R is a purine (i.e. either adenine or guanine) and Y is a pyrimidine (i.e. either uracil, cytosine or thymine) and N is any of the bases [8]. The different DNA sequences which can be bound by $\text{NF}\kappa\text{B}$ probably arose from evolutionary changes and provides differential binding strength leading to differential transcriptional activities [9]. Duplicate consensus sequences can also be present on the promoter, synergising to produce a stronger activity as more TFBS are occupied or they could be spare TFBS leftover from evolution. Furthermore, the interaction among

transcription factors can also include both activators and repressors binding to the promoter, leading to an increasing complexity in transcriptional regulation.

1.3 The Transcriptome: Alternative Splicing and Transcriptomics

There may be more than one type of protein arising from a specific mRNA. The reason for this lies in a process called splicing, which occurs between transcription and translation. There exists, within the sequence of a gene, parts that will be present in the final translated mRNA (the exons) and parts that will not (the introns). The excision of the introns and the ligation of the exons to form a continuous strand of mRNA is the basic definition of splicing [10]. Splicing is not thought to always occur in the same way for a specific gene. Sometimes, introns appear to be included in a final mRNA, while in other mRNAs, one or more exons may have been spliced out [11, 12]. This variation leads to variations in the mRNA pool, and subsequently in the protein population of the cell. Moreover, alternative splicing is proposed to be a significant process in the cell, with 40–60 % of human genes displaying alternative splice forms [13]. Thus there is not always a correlation between the expression level of the gene (the abundance of its mRNA) and the number of proteins present in the sample of cells one is measuring. Investigating the proteins gives, for all these reasons, a wealth of important information that cannot be inferred from a study of mRNA alone. And lastly, working with proteins has the additional advantage that proteins are generally more stable than mRNA, which greatly facilitates the analysis.

1.4 Transcriptional Assays

Knowing the sequence of the genes is only the beginning. What we strive to decipher is to understand what these genes do and how they do it. In other words, when are they active and how active are they under any specific set of conditions? The simplest answer to that kind of question is given by the study of the abundance and variety of mRNA molecules in a cell. Typically, the more active a gene is, the more copies of the corresponding mRNA are produced.

Using microarray, one of the most reliable and widely used workhorses of transcriptomics, it is possible to follow the expression patterns of all identified genes in an organism simultaneously. However, despite its power, years of experimentation have revealed major drawbacks, including the availability of gene microarrays for only a few of the better studied organisms (e.g. human, mice, rat) and the inability to detect splicing events, i.e. to distinguish differentially spliced mRNA binding to a probe. The latter can be overcome by the use of tiling arrays, which unfortunately substantially increases the operational cost [14]. Furthermore, recent developments

in deep sequencing technology, stimulated by the world-wide effort to sequence the human genome, has allowed the development of RNA-sequencing technology, a revolutionary platform that allows the quantification of mRNA expression, non-coding RNA and splicing events in non-annotated genomes [15].

While microarrays provide a global picture of the epigenetic differences, they are limited in their ability to explore differences in specific genes. At the individual levels, other methods exist, such as chromatin immunoprecipitation—which can reveal the interaction of transcription factors to their binding sites—and luciferase assays—which can show function of specific DNA sequences in activating gene transcription [16]. The ability to mutate DNA base pairs makes these methods particularly suitable to investigate nonlinear dynamics at the individual sequence.

2 Mechanisms for Nonlinear Gene Expression

Gene expression is tightly regulated in a stimuli-dependent and cell dependent way. Genes can be expressed in response to specific stimuli in a linear, graded way, or in a switch-like, binary manner. Notable examples of mechanisms that can impart a switch-like or binary behaviour to the expression of a particular gene, include, but are not limited to: transcription factor synergism [17, 18], competition between an activator and a repressor [19], Boolean-logic within promoter regulatory elements [20] and chromatin remodelling surrounding genetic-regulatory elements [21, 22]. Further details are provided in Fig. 1. Emerging data from transcriptomics has revealed that gene expression is also largely bimodal across different tissues within the same organism [23, 24]. In addition, during evolution, different species seem to have chosen opposing regulatory mechanism (linear vs nonlinear) for the regulation of the same gene, and species-specific regulatory mechanism can contribute to physiological differences [25].

2.1 *Examples of Regulatory Mechanism of Nonlinear Mammalian Gene Expression*

A single gene can be expressed with very different dynamics in response to the same graded stimuli in different cell types. The mechanisms allowing for these differences are only now emerging. Work from our labs [20] has revealed that cyclooxygenase 2 (COX2), a key enzyme of the prostaglandin/eicosanoid pathway, is expressed in a linear dose dependent manner in response to a gradient of tumour necrosis factor alpha (TNF α) stimulation in mouse embryonic fibroblasts (MEF), while in human colorectal cancer cells (HT29) it follows a switch-like expression pattern with low levels being expressed under low grade concentrations of the inflammatory stimuli (TNF α), dramatically increasing at a defined threshold concentration to a

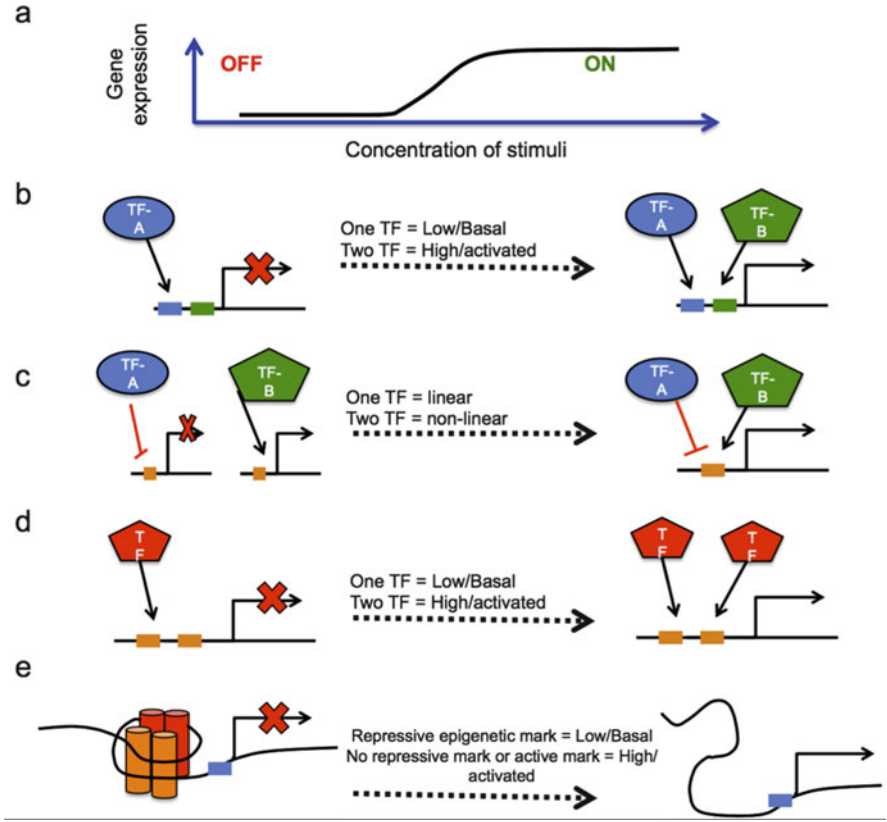


Fig. 1 Notable examples of regulatory mechanism of nonlinear gene expression. Multiple mechanisms have been shown to biologically regulate gene expression transition from an ON to OFF state, or from a linear to a nonlinear response (a). These include TF synergism, whereby gene transcription is maximally induced under the presence of two activators (b). TF competition, whereby in the presence of the activator or repressor alone, linear gene responses are observed, while in the presence of both the activator and repressor a nonlinear dose-response of gene activation arises in response to a particular stimuli. (c) Epigenetic mechanisms can silence (OFF) transcription, and when removed, or in the presence of epigenetic marks of active gene transcription (ON), transcription is induced. (d) Boolean logic within regulatory promoter elements, whereby occupancy of both DNA-regulatory elements by their cognate transcription factor is necessary for maximal gene induction

final plateau of expression (Fig. 2). In our study, we have shown that the switch like behaviour could be explained by the presence of regulatory NF κ B DNA sequences within the COX2 promoter that function according to the rules of Boolean logic, allowing on and off gene expression (see Sect. 2).

NF κ B is a ubiquitous transcription factor and a prime example where the complexity of the DNA regulatory sequences in the promoter can lead to a switch like behaviour [26]. Genes shown to have a NF κ B switch-like behaviour

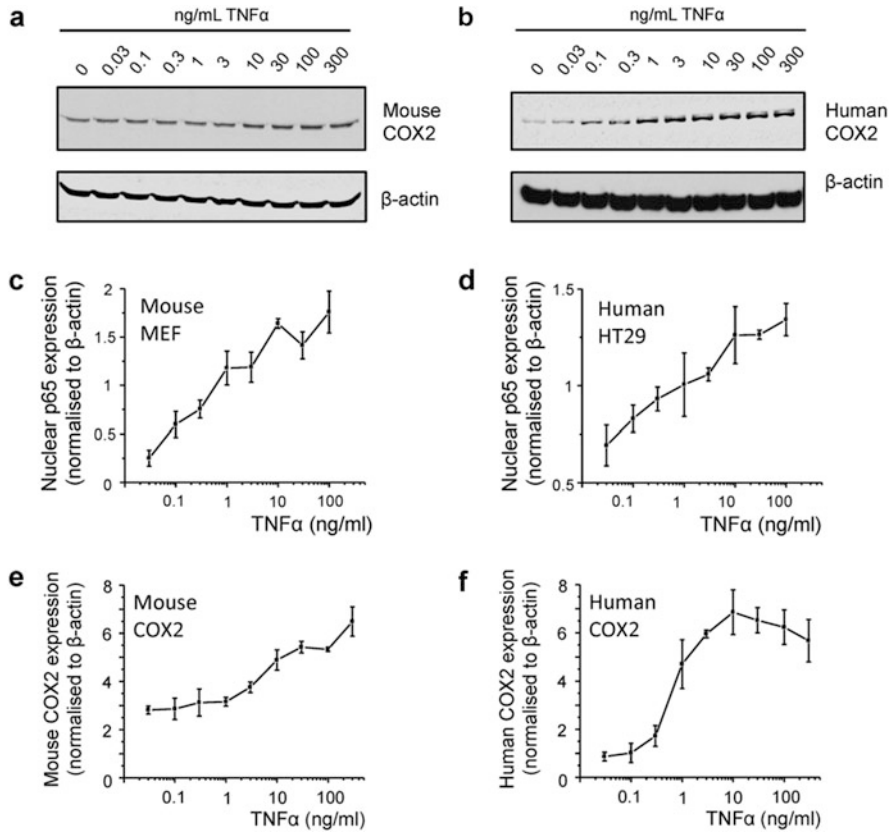


Fig. 2 Linear and switch like expression of COX2 in mouse and human cells. Expression of COX2 protein in mouse MEF cells and human HT29 cells stimulated with TNF α , representative western blots are shown in (a), (b) and their densitometry in (c), (d). TNF α -induced luciferase activity under the control of mouse (d) and human (e) COX2 promoters expressed in HEK293 cells (e), (f). (Reproduced with permission from [20])

include the COX2 gene in response to TNF α (see above), but also the RAGE (for receptor for advanced glycation end products) which requires 2 κ B binding sites for lipopolysaccharide (LPS) induction [27] and the C-X-C motif chemokine 10 (CXCL10), which requires 2 κ B binding sites for LPS or interferon- γ induction [28].

NF κ B is not the only transcription factor reported to generate a switch-like or a linear gene activation dependent on the gene promoter structure, for example, Joers and colleagues have shown that depending on the promoter structure of the target gene, using transfected plasmids containing different parts of the endogenous promoter, p53—dependent transcription can be binary or graded [22]. The transcription factor nuclear factor of activated T-cells (NF-AT) has also been

shown to provide a tightly regulated threshold controlled activation of artificial gene promoters exhibiting a switch like behaviour in response to T cell activation [29].

Binary changes in mammalian gene expression can also result from transcriptional activator synergism, with minimal or no expression when one factor binds to the promoter, but with an unexpectedly high expression level when multiple transcription factors bind to the promoter [17, 21, 30]. Another mechanism for binary gene expression involving transcription factors is the competition for a binding site between an activator and a repressor. Rossi and colleagues infected mammalian cells with the repressor TetR and the activator protein TetR-VP16 fusion protein [19]. They found that cells expressing the activator only (TetR-VP16) produce a graded change in gene expression, while cells infected with both the activator and the repressor (TetR) generate binary patterns of expression [19].

Promoters in a repressive, condensed chromatin environment are unresponsive to the transcriptional machinery [31], and the change between the repressed and open chromatin states may confer the properties of a switch like gene expression profile [21]. For example, Ertel and colleagues have shown that genes with a bimodal gene expression profile in different tissues have an association with histone methylation (H3K4me3) [23], a marker for transcriptionally active genes [32]. Histone methylation, along with DNA methylation, is a key player in cell differentiation during development [33] and aberrant histone methylation patterns are among the epigenetic modifications that give rise to cancer [34]. Integrated experimental and model-based analysis also revealed that gene expression mediated by H2A may follow bistable, switch-like dynamics under specific conditions due to bistable H2A ubiquitination controlled by the master gene silencer Polycomb complex 1 [35]. Thus epigenetic mechanisms such as histone methylation are important determinants of tissue specific switch-like expression of different genes, and deregulation of this mechanism of gene regulation may have pathological consequences.

2.2 Species-Specific Differential Regulatory Mechanism and Their Implications in Comparative Physiology

Interestingly, during evolution, several homologue genes seem to have selected to diverge in the linear or nonlinear gene expression patterns. Species-specific differential regulatory mechanisms of gene expression, as demonstrated in our work for the COX2 promoter [20], may contribute to the recently reported discordance observed between human and mouse models of inflammatory diseases [36]. Species-difference has also been observed for other genes, such as for ETS1 [37] and Waf1/p21 [22]. The transcription factor ETS1 is responsible for the mouse specific expression of the T cell factor Thy-1 in the thymus. ETS1 is found to be preferentially to the proximal promoter of human genes but not mice genes and has been suggested to contribute to the immune system differences between mice and

human [37]. A comparison of the expression in response to DNA damaging agents of the Waf1/p21 protein, a p53 target gene revealed a graded response in mouse NIH 3T3 cells, but a binary, switch like response in human MCF7 cells [22]. In this example, the underlying molecular mechanism was not studied in detail, but the authors suggest, having in consideration their analysis of recombinant versions of the Waf1 promoter, that different regulatory elements within the promoter might be masked by epigenetic regulation in different cell types [22].

Epigenetic regulation, as evidenced by differential patterns of histone methylation and acetylation, is also a key mechanism leading to differential gene expression in mice and human tissues [38]. In addition to epigenetic marks and differential binding of transcription factors, differences in the cis-regulatory elements (CRE, i.e. TF binding sites and associated sequences required for transcription) can also contribute to species-specific expression pattern of a particular gene [39]. Novel CRE, but also mutations in pre-existing CRE including insertions, deletions, duplications and changes in the DNA strand of the TF binding site can lead to species-specific differences in gene expression [39].

Thus species-specific gene expression can be achieved by several mechanisms, as outlined above, and the linear or switch-like behaviour of a particular gene can be an important contribution to differences in animal physiology, that can be easily overlooked and impair the translational application of biomedical research. A better understanding of these mechanisms will allow a better prediction of when the findings from biomedical research in mouse models are relevant to human disease [37].

2.3 Tissue Specific Binary Gene Regulation

One of the most exciting questions in biology is to understand how, within a given individual, cells with the same genetic code can give rise to different tissues, with diverse physiological properties. It is clear that gene expression varies between graded and switch like, binary expression pattern across different tissues, being highly expressed in one tissue type, and with low expression in other. Genome-wide identification and annotation of genes with switch-like expression at the transcript level in mouse and human using large transcriptomics datasets for several different healthy tissue (e.g. brain, heart, and skeletal muscle) have been used to study the cellular pathways and regulatory mechanisms involving this class of genes [23, 24]. Genes with binary expression across different tissues seem to have higher than average number of transcription start sites per gene, differentially methylated histones, enriched TATA box regulatory motifs and alternative promoters [23, 24]. What is lacking here is a systematic modelling of the complex nonlinear dynamics to unravel the transcriptional signalling cascades.

2.4 Mathematical Modelling of Transcription Factor-Promoter Interaction

2.4.1 Thermostatistical Modelling

Transcriptional activity is in general a nonlinear function of transcription factor abundance. There are several approaches to understand the nonlinear impact of transcription factors on transcriptional activity, one of which uses thermostatistical principles [40–45] and was introduced by Shea and Ackers [46]. In what follows, a recent study on the thermostatistical mathematical modelling of transcription-factor induced promoter activity will be reviewed [20] and emphasis will be put on the nonlinear aspects.

The thermostatistical modelling makes two a priori assumptions. The first assumption is that the rate of transcription is proportional to the probability of transcription initiation. The probability of transcription initiation in turn is assumed to be equal with the probability that RNAP binds at the promoter. Consequently, this first assumption may be regarded as an assumption involving two sub-assumptions. Mathematically speaking, let r denote the rate of transcription of a gene and P denote the probability of RNAP binding. Let κ denote a positive constant. Then, the first assumption state that

$$r = \kappa P. \quad (1)$$

In order to determine the probability P , a second assumption is made. Accordingly, it is assumed that the principles of statistical physics hold and determine the value of P . In more detail, P can be derived by looking at three levels: macro, meso, and micro [47]. On a macro level, there are only two possibilities: there is a transcription initiation or not. In line with the aforementioned first assumption, this is phrased like RNAP can be bound to the promoter or not. Let $P(\text{on})$ and $P(\text{off})$ denote the probabilities that RNAP is bound or not bound at the promoter. Then we have $P = P(\text{on})$ and

$$P(\text{on}) + P(\text{off}) = 1. \quad (2)$$

In order to determine $P(\text{on})$ all so-called microstates or possibilities that RNAP can be bound at the promoter must be considered. $P(\text{on})$ is the sum of the probabilities of these microstates. Let $p(j, \text{on}, \text{micro})$ denote the probability that RNAP is bound at the promoter in a particular way described by the index j . Let us assume there are N different ways that RNAP can be bound. Then, $P(\text{on})$ is defined by

$$P(\text{on}) = \sum_j p(j, \text{on}, \text{micro}). \quad (3)$$

Statistical physics claims that the probability of a state is determined by its energy. On the microlevel, the so-called Boltzmann distribution holds. Accordingly, the probability is an exponential function of the energy such that high energetic states have low probability, while states with low energy have high probability. Let $E(j, \text{on})$ denote the energy of a microstate j . Then, we have

$$p(j, \text{on}, \text{micro}) = \exp(-E(j, \text{on})/RT)/Z, \quad (4)$$

where R is the gas constant, T is the temperature and Z is a normalization factor. The energy values must be determined from a priori knowledge or can be fitted from data. In contrast, the normalization factor Z can be determined by theoretical reasoning. An equation similar to Eq. (2) also holds on the microlevel and can be used to determine Z . To this end, the microstates k for which RNAP is not bound at the promoter are considered. Let $p(k, \text{off}, \text{micro})$ denote the probabilities of these states. Then, we have

$$p(k, \text{off}, \text{micro}) = \exp(-E(k, \text{off})/RT)/Z, \quad (5)$$

where $E(k, \text{off})$ is the energy of the microstate k . All probabilities taken together must sum up to unity like

$$\sum_j p(j, \text{on}, \text{micro}) + \sum_k p(k, \text{off}, \text{micro}) = 1, \quad (6)$$

which can be used to determine Z .

Having discussed the microstate picture, we note that there is typically a large number of microstates, which makes the counting of the states a tedious task. The thermo-statistical approach can be applied more effectively in the mesoscale picture. To this end, all microstates that exhibit exactly the same energy are group together into classes. The number of microstates that have a particular energy level $E(j, \text{on})$ or $E(k, \text{off})$ in common is the degeneracy of that energy level. The degeneracy is expressed as a positive number and typically denoted by g . On the mesolevel, we are concerned with the possibilities that RNAP is bound at the promoter with certain energies $E(j, \text{on})$. The number of different possibilities for a given energy level $E(j, \text{on})$ is the degeneracy of that energy level j . Likewise, we are concerned with energetically different possibilities that the promoter is not occupied by RNAP. Again, for each energy level $E(k, \text{off})$ the number of different ways to achieve this corresponds to the degeneracy of that energy level k . In statistical physics, the probabilities are calculated as Boltzmann probabilities involving degeneracy factors like

$$p(j, \text{on}, \text{meso}) = g(j, \text{on}) \exp(-E(j, \text{on})/RT)/Z \quad (7)$$

and

$$p(k, \text{off}, \text{meso}) = g(k, \text{off}) \exp(-E(k, \text{off})/RT)/Z \quad (8)$$

The normalization constant Z is the same as before. Therefore, it could be determined from Eq. (6). However, usually, Z is determined from the mesoscale picture using the fact that all probabilities on the mesolevel must sum up to unity. That is, Z can be determined from

$$\sum_j p(j, on, meso) + \sum_k p(k, off, meso) = 1. \quad (9)$$

By analogy to Eq. (3), $P(on)$ can be determined from the mesoscale RNAP binding probabilities like

$$P(on) = \sum_j p(j, on, meso). \quad (10)$$

The degeneracy factors are functions of the RNAP and transcription factor numbers that are available for binding. They also depend on the number of binding sites for the transcription factors. In order to discuss the explicit rules for the degeneracy factors, we present the simplest non-trivial example and then discuss generalization. This example is a promoter that is regulated by a single transcription factor and involves a single binding site for the transcription factor.

Case 1: Nonlinear Regulation of Transcriptional Activity of a Promoter with a Single Transcription Factor and Single Binding Site

Let TF1 denote the transcription factor. There are four mesoscale states. RNAP is or is not bound at the promoter. For each of these two cases, TF1 is or is not bound at its binding site. Let n_R and n_1 denote the molecule numbers of RNAP and TF1 available for binding at the promoter. Then, the degeneracy factors of the four states listed in Table 1.

According to Table 1, the degeneracy is proportional to the number of available molecules. For example, let us consider the case in which RNAP is not bound at the promoter but a transcription factor is bound at the promoter. Let us assume there are ten transcription factor molecules labelled with A, B, C, D, E, F, G, H, J, K, L available for binding at the transcription factor binding site. Then the binding of molecule A involves the same energy as the binding of any other molecule B, ...,

Table 1 Mesoscale characterization and degeneracy factors of a promoter regulated by a single transcription factor

Transcriptional activity	RNAP binding (On = Yes, Off = No)	TF1 binding (On = Yes, Off = No)	Degeneracy g
Yes	On	On	$g(1, on) = n_R * n_1$
	On	Off	$g(2, on) = n_R$
No	Off	On	$g(1, off) = n_1$
	Off	Off	$g(2, off) = 1$

Table 2 Mesoscale probabilities of a promoter regulated by a single transcription factor

RNAP	TF1	Energy E	Probabilities p
On	On	E_{R1}	$n_R * n_1 * w(R1)/Z$
On	Off	E_R	$n_R * w(R)/Z$
Off	On	E_1	$n_1 * w(1)/Z$
Off	Off	E_0	$w(0)/Z$

Table 3 Mesoscale probabilities defined on an appropriately shifted energy scale

RNAP	TF1	Relative energy $E(\text{rel})$	Probabilities p
On	On	$E_{R1}-E_0$	$p(1, \text{on}) = n_R * n_1 * \Omega(R1)/Z^*$
On	Off	E_R-E_0	$p(2, \text{on}) = n_R * \Omega(R)/Z^*$
Off	On	E_1-E_0	$p(1, \text{off}) = n_1 * \Omega(1)/Z^*$
Off	Off	0	$p(2, \text{off}) = 1/Z^*$

L. Therefore, all ten possibilities (all ten microstates) have the same energy. The degeneracy equals 10. Similarly, the other cases in Table 1 can be derived.

Taking the energy levels of four different mesostates into account, formally, the probabilities of the mesostates can be written down. This is shown in Table 2, with

$$\begin{aligned} w(R1) &= \exp(-E_{R1}/RT), & w(R) &= \exp(-E_R/RT), \\ w(1) &= \exp(-E_1/RT), & w(0) &= \exp(-E_0/RT). \end{aligned} \quad (11)$$

A detailed calculation shows that Boltzmann probabilities in statistical physics such as those listed in Table 2 only depend on relative energy levels. Therefore, we may shift all energy levels by E_0 . Table 2 then becomes Table 3, with

$$\begin{aligned} \Omega(R1) &= \exp(-(E_{R1} - E_0)/RT), & \Omega(R) &= \exp(-(E_R - E_0)/RT), \\ \Omega(1) &= \exp(-(E_1 - E_0)/RT) \end{aligned} \quad (12)$$

and $Z^* = Z/\exp(-E_0/RT)$.

The sum of the four probabilities equals unity. This implies that Z^* reads like

$$Z^* = 1 + n_1 \Omega(1) + n_R \Omega(R) + n_1 n_R \Omega(R1). \quad (13)$$

$P(\text{on})$ is given by the sum of $p(1, \text{on})$ and $p(2, \text{on})$. Consequently, we have

$$P(\text{on}) = (n_R \Omega(R) + n_1 n_R \Omega(R1))/Z^*. \quad (14)$$

We are interested in the nonlinear regulation of transcriptional activity via the abundance of TF1. Therefore, $P(\text{on})$ may be expressed like

$$P(\text{on}) = \frac{A + B n_1}{C + D n_1}, \quad (15)$$

where A, B, C, D can be determined from Eq. (14). For sake of completeness let us write down the equation for the rate of transcription. From Eq. (1) it follows that

$$r = \kappa \frac{A + B n_1}{C + D n_1}. \quad (16)$$

We see that the probability $P(\text{on})$ of RNAP binding and the rate of transcription r are nonlinear function of n_1 . For $n_1 = 0$ we have the baseline binding probability $A/C < 1$ related to the baseline rate of transcription $r = \kappa A/C$. In contrast, for n_1 to infinity we have the saturation binding probability $P(\text{on}) = B/D < 1$ related to the saturation rate of transcription $r = \kappa B/D$. Note that since the transcription factor is assumed to be an activator it lowers the binding energy of RNAP at the promoter. This implies that the saturation binding probability B/D is larger than the baseline binding probability A/C . Likewise, we have $r(0) < r(n_1 \rightarrow \infty)$. Moreover, one can show that $P(\text{on})$ and consequently r are monotonically increasing functions of the transcription factor molecule number n_1 . Having discussed the nonlinear regulation of the transcriptional activity by means of a single transcription factor, we can easily generalize the approach to account for several transcription factors. Again, let us assume that each transcription factor only exhibits a single binding site in the promoter region.

Case 2: Nonlinear Regulation of Transcriptional Activity of a Promoter with Multiple Transcription Factors Each Having a Single Binding Site

Let us rewrite the four mesostate probabilities $p(j, \text{on/off})$ listed in Table 3 like

$$p(j, \text{on/off}) = n_R^{m_R} n_1^{m_1} \Omega(j, \text{on/off}) / Z^* . \quad (17)$$

The parameters m_R and m_1 are integer that can assume only the values 0 and 1. Then in the case of a promoter regulated by L activators, we have

$$p(j, \text{on/off}) = n_R^{m_R} n_1^{m_1} \cdots n_L^{m_L} \Omega(j, \text{on/off}) / Z_L, \quad (18)$$

with $\Omega(j, \text{on/off}) = \exp(-(E(j, \text{on/off}) - E_0)/RT)$.

The variables $n_1, \dots, n_L$ denote the number of transcription factors TF1, $\dots$, TFL available for binding. The parameters $m_1, \dots, m_L$ are 0 or 1 depending whether the respective transcription factor is not bound (0) or bound (1) at the transcription-factor specific binding site. The parameter Z_L is a normalization constant that can be determined from the requirement that all mesostate probabilities must add up to unity. From Eqs. (1), (17), (20), we then obtain the rate of transcription as a nonlinear function $f(\dots)$ of the transcription factor abundances $n_1, \dots, n_L$:

$$r = f(n_1, \dots, n_L). \quad (19)$$

It is important to note that Eq. (18) may exhibit nonlinear terms of the form $n_1 * n_2, n_1 * n_2 * n_3$, etc. This can lead to synergy effects of the transcription factor concentrations. Note that alternatively to the number of transcription factor molecules,

molar units or other measures of concentration can be used. Let $[TF_j]$ denote the concentration of the transcription factor j measured in a suitable unit. Then, Eq. (18) is replaced by

$$r = h([TF1, TF2, \dots, TFL]), \quad (20)$$

where $h(\dots)$ is a nonlinear function of the concentrations. Again, the function may involve products of the form $[TF1]*[TF2]$, $[TF1]*[TF2]*[TF3]$, etc.

Case 3: Transcription Factors with More Than 1 Binding Site

In the considerations made above, it was assumed that each transcription factor exhibits only one binding site. More complicated promoters can be addressed with the thermostistical approach. In order to illustrate this issue, let us return to a promoter regulated by a single transcription factor called TF. The abundance is measured by n . If there are two binding sites for the transcription factor under consideration, then the degeneracy increases quadratic with the number of transcription factor molecules competing for binding. In general, if there are s binding sites then the degeneracy is given by

$$g = n^s. \quad (21)$$

Note that this rule holds for not too small numbers of transcription factors. If there are only a few transcription factors available for binding, then a more precise counting scheme should be used [48, 49]. In analogy to Eq. (16), the rate of transcription for a promoter with s binding sites for its activator reads

$$r = \kappa \frac{A + B n_1^s}{C + D n_1^s}. \quad (22)$$

In general, increasing the number of binding sites makes the sigmoid character of Eq. (22) more pronounced. That is, the function increases more rapidly. This property of s becomes obvious in the special case when the baseline transcription rate is negligible. That is, let us assume that the RNAP binding energy is relatively high such that the binding probability of RNAP without the interaction of the transcription factor is almost zero. That is, $\Omega(R)$ is approximately equal to zero which implies that we can put $A = 0$. In this case, Eq. (22) reduces to

$$r = \kappa B \frac{n_1^s}{C + D n_1^s}. \quad (23)$$

This is the standard Hill function used in systems biology. The Hill function has its 50 percent point at $n^s = C/D$. For the Hill function it is well known that around this 50 percent point, the function increases more sharply when increasing the exponent s .

Let us summarize this section by pointing out some key aspects of the thermostistical approach for understanding the nonlinear regulation of transcriptional

activity by means of transcription factors. First, the nonlinearity is related to the degeneracy of energetically equivalent states. That is, for understanding the nonlinear properties the pivotal element is the degeneracy of states, whereas the binding energies “only” show up as parameters (in our notation: the omega parameters). Second, the thermostatistical approach naturally explains that transcriptional activity is regulated by products of transcription factor concentrations. Since such products under suitable circumstances are the mathematical expressions of synergy expressions, some synergy effects arising on the transcriptional level may be identified as a consequence of the thermostatistical nature of the transcriptional machinery. In other words, when in asking the question, why the transcription of a gene in a particular cell line is regulated synergistically by two transcription factors, then we may answer that this is due to the fact that the transcriptional machinery obeys the laws of thermostatics. Third, the thermostatistical approach reveals that the number of binding sites for a given transcription factor has a crucial impact on the nature of the nonlinear regulation by means of this transcription factor. Roughly speaking, the degree of nonlinearity becomes stronger when the number of binding sites is increased.

2.4.2 Quantitative Kinetic Modelling of Promoter Activity Mediated by Transcriptional Factors

There have been considerable efforts to construct models for transcriptional regulation, which result in different modelling formulations [50–52]. Beside the thermostatistical approach discussed in the above Sect. 2.4.1, quantitative models of transcriptional regulation based on ordinary differential equations (ODEs) have been a major focus [53]. In a recent publication, we have employed ODE-based modelling to investigate the dynamical properties of the COX2 gene’s promoter structure consisting of multiple binding sites and its interaction with the cognate transcriptional factor NF κ B across different species [20].

Instead of assuming or estimating transcriptional activity as a lumped, nonlinear function of the transcriptional factor, our model formulation were based on the law of elementary mass-action kinetics that explicitly accounts for the association and dissociation of the transcriptional factor and its binding sites. The total transcriptional activity is then calculated as a weighted aggregated sum of the active promoter-TF complexes formed, which may possess different levels of transcriptional activation potential. For example, a fully bound complex where all the promoter’s binding sites are occupied by a TF molecule may elicit a substantially stronger gene activation capacity than partially occupied promoters where only one or a subset of the binding sites are TF-bound. Depending on the kinetics between the TF and different binding sites, the rate of which could be experimentally measured, we may have different distribution of the promoter-TF complexes. Furthermore, such formulation allowed us to make assumptions as to how the binding sites could interplay (e.g. independent or cooperative) in mediating the transcriptional activity, thereby embedding different assumptions of modes of regulation such

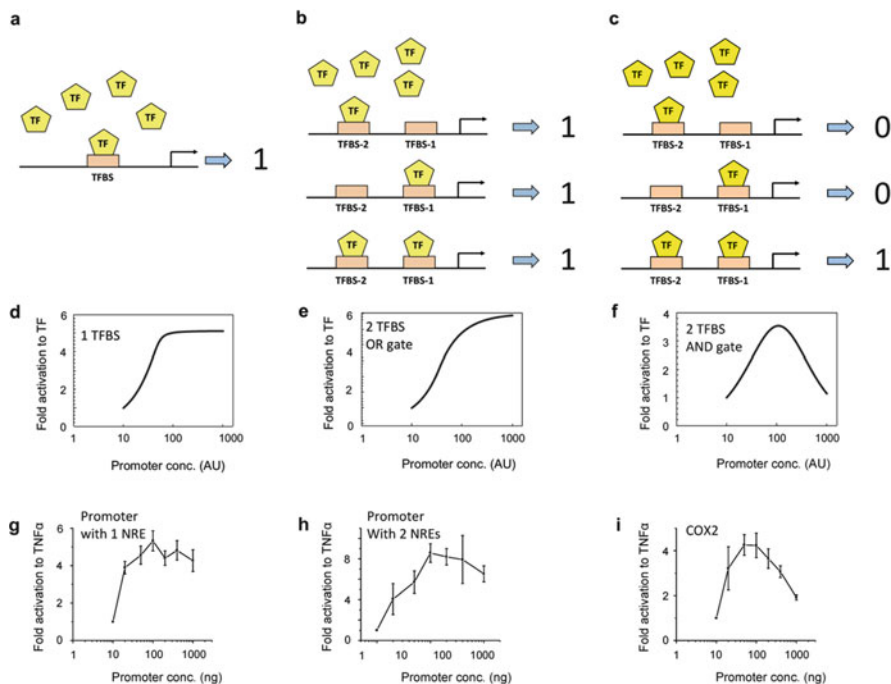


Fig. 3 Mathematical models of promoter activity. (a)–(c) Simplified schemes showing binding of TFs on promoters with one transcription binding site (TFBS) (a), or 2 TFBS arranged as OR (b) or AND gates (c). Promoter is either activated (1) or not (0). (d)–(f) Mathematical predictions of the transcriptional activity for each scheme, at a fixed concentration of TF but varying concentration of promoters. Model association/dissociation rates used for plotting panel (f) are 0.001 and 0.2 respectively. (g)–(i) Experimental validation of predictions using artificial promoters (g), (h) or the human COX2 promoter (i) expressed in HEK293 cells in response to TNF α (1 ng/ml). Data is shown as fold activation over unstimulated ($n = 4 - 5$). (Reproduced with permission from [20])

as OR and AND logics between the sites [20]. The differential transcriptional activation capacity of the formed promoter-TF complexes, possible to obtain from experiments, will serve to guide the choice of the weigh coefficients that make up the overall total transcriptional activity.

Following such approach, we developed three kinetic models to quantitatively analyze and predict steady-state dynamics of the COX2 gene expression under different modes of regulation by NF κ B. These models, schematically shown in Fig. 3 are referred to as the “1-site”, “2-site OR-gate” and “2-site AND-gate” models, describe: (1) a promoter that is regulated by a transcriptional factor through a single TF-Promoter binding site; (2) a promoter regulated by a transcriptional factor through two TF-Promoter binding sites following an OR gate and (3) a promoter regulated by a transcriptional factor through two TF-Promoter binding sites following an AND gate regulation. Although the models were kept to the minimal level of complexity, they were able to generate important predictions and

insights that guided useful experiments and distinguish the model that best match experimental data. Interestingly, the models suggested an unintuitive experiment where a graded increase of the promoter concentration led to distinct responses of transcriptional activity between the three models [20]. While the 1-site and 2-site OR-gate models showed a nonlinear monotonic response, only the 2-site AND-gate model showed a biphasic response where either too low or too high abundance of the promoter would suppress transcriptional activity. Our followed up experiment indeed showed a biphasic response, which suggest the AND-gate effect between the NF κ B binding sites. It is worth noting here that although varying the promoter concentration may be considered as an unintuitive experiment from an experimental perspective, such data allowed us to gain crucial insights into the dynamic behaviour of the system, suggesting model-led experiments may not need to answer a specific biological question but still are useful.

3 Conclusions

In this review, we discussed two main modelling approaches for constructing models of transcriptional regulation and the consequent nonlinearity arising from interactions between the transcriptional factors and binding sites. We have demonstrated the application of thermostistical approach using a number of case-study scenarios, and showed how ordinary differential equations based approach may complement the former method in generating experimentally testable hypotheses. Choosing which modelling framework for a specific study would largely depend on the questions asked and the data available at hand. It should be noted that these approaches are complementary rather than competitive, and thus one could see their combined usage in examining the same biological questions.

While transcriptional networks can be observed to behave as logic circuits, and can be more or less described mathematically, they are likely to be very different from a neatly constructed electronic logic circuit designed for maximal efficiency. Herein lies the problem when attempting to reverse-engineer biological networks. Real-life transcriptional networks arose from evolution, with sequences and design which may seem inefficient or redundant, but which are/were important for certain situation [54]. It could well be that the nonlinearity of the transcriptional network have other undetermined functions to cope with the unknown.

References

1. Venter, J.C., Adams, M.D., Myers, E.W., Li, P.W., Mural, R.J., et al.: The sequence of the human genome. *Science* **291**, 1304–1351 (2001)
2. Lander, E.S., Linton, L.M., Birren, B., Nusbaum, C., Zody, M.C., et al.: Initial sequencing and analysis of the human genome. *Nature* **409**, 860–921 (2001)

3. Okamoto, A., Tanaka, K., Saito, I.: DNA logic gates. *J. Am. Chem. Soc.* **126**, 9458–9463 (2004)
4. Goldman, N., Bertone, P., Chen, S., Dessimoz, C., LeProust, E.M., et al.: Towards practical, high-capacity, low-maintenance information storage in synthesized DNA. *Nature* **494**, 77–80 (2013)
5. Fisher, J., Henzinger, T.A.: Executable cell biology. *Nat. Biotechnol.* **25**, 1239–1249 (2007)
6. Kunst, F., Ogasawara, N., Moszer, I., Albertini, A.M., Alloni, G., et al.: The complete genome sequence of the gram-positive bacterium *Bacillus subtilis*. *Nature* **390**, 249–256 (1997)
7. Gibson, D.G., Glass, J.I., Lartigue, C., Noskov, V.N., Chuang, R.Y., et al.: Creation of a bacterial cell controlled by a chemically synthesized genome. *Science* **329**, 52–56 (2010)
8. Wan, F., Lenardo, M.J.: Specification of DNA binding activity of NF-kappaB proteins. *Cold Spring Harb. Perspect. Biol.* **1**, a000067 (2009)
9. Wong, D., Teixeira, A., Oikonomopoulos, S., Humburg, P., Lone, I.N., et al.: Extensive characterization of NF-kappaB binding uncovers non-canonical motifs and advances the interpretation of genetic functional traits. *Genome Biol.* **12**, R70 (2011)
10. Fedorova, L., Fedorov, A.: Introns in gene evolution. *Genetica* **118**, 123–131 (2003)
11. Matlin, A.J., Clark, F., Smith, C.W.: Understanding alternative splicing: towards a cellular code. *Nat. Rev. Mol. Cell Biol.* **6**, 386–398 (2005)
12. Black, D.L.: Mechanisms of alternative pre-messenger RNA splicing. *Annu. Rev. Biochem.* **72**, 291–336 (2003)
13. Modrek, B., Lee, C.: A genomic view of alternative splicing. *Nat. Genet.* **30**, 13–19 (2002)
14. Mockler, T.C., Chan, S., Sundaresan, A., Chen, H., Jacobsen, S.E., et al.: Applications of DNA tiling arrays for whole-genome analysis. *Genomics* **85**, 1–15 (2005)
15. Wang, Z., Gerstein, M., Snyder, M.: RNA-Seq: a revolutionary tool for transcriptomics. *Nat. Rev. Genet.* **10**, 57–63 (2009)
16. Cavadas, M.A.S., Cheong, A.: Monitoring of transcriptional dynamics of HIF and NF kappa B activities. *Bioluminescent Imaging Methods Protoc.* **1098**, 97–105 (2014)
17. Carey, M., Lin, Y.S., Green, M.R., Ptashne, M.: A mechanism for synergistic activation of a mammalian gene by GAL4 derivatives. *Nature* **345**, 361–364 (1990)
18. Bruning, U., Fitzpatrick, S.F., Frank, T., Birtwistle, M., Taylor, C.T., et al.: NFkappaB and HIF display synergistic behaviour during hypoxic inflammation. *Cell. Mol. Life Sci.* **69**, 1319–1329 (2012)
19. Rossi, F.M., Kringstein, A.M., Spicher, A., Guicherit, O.M., Blau, H.M.: Transcriptional control: rheostat converted to on/off switch. *Mol. Cell* **6**, 723–728 (2000)
20. Nguyen, L.K., Cavadas, M.A., Kholodenko, B.N., Frank, T.D., Cheong, A.: Species differential regulation of COX2 can be described by an NFkappaB-dependent logic AND gate. *Cell. Mol. Life Sci.* **72**, 2431–2443 (2015)
21. Biggar, S.R., Crabtree, G.R.: Cell signaling can direct either binary or graded transcriptional responses. *EMBO J.* **20**, 3167–3176 (2001)
22. Joers, A., Jaks, V., Kase, J., Maimets, T.: p53-dependent transcription can exhibit both on/off and graded response after genotoxic stress. *Oncogene* **23**, 6175–6185 (2004)
23. Ertel, A., Tozeren, A.: Human and mouse switch-like genes share common transcriptional regulatory mechanisms for bimodality. *BMC Genom.* **9**, 628 (2008)
24. Ertel, A., Tozeren, A.: Switch-like genes populate cell communication pathways and are enriched for extracellular proteins. *BMC Genom.* **9**, 3 (2008)
25. Nguyen, L.K., Cavadas, M.A., Scholz, C.C., Fitzpatrick, S.F., Bruning, U., et al.: A dynamic model of the hypoxia-inducible factor 1-alpha (HIF-1alpha) network. *J. Cell Sci.* **126**, 1454–63 (2013)
26. Leung, T.H., Hoffmann, A., Baltimore, D.: One nucleotide in a kappaB site can determine cofactor specificity for NF-kappaB dimers. *Cell* **118**, 453–464 (2004)
27. Li, J., Schmidt, A.M.: Characterization and functional analysis of the promoter of RAGE, the receptor for advanced glycation end products. *J. Biol. Chem.* **272**, 16498–16506 (1997)
28. Ohmori, Y., Hamilton, T.A.: Cooperative interaction between interferon (IFN) stimulus response element and kappa B sequence motifs controls IFN gamma- and lipopolysaccharide-

- stimulated transcription from the murine IP-10 promoter. *J. Biol. Chem.* **268**, 6677–6688 (1993)
29. Fiering, S., Northrop, J.P., Nolan, G.P., Mattila, P.S., Crabtree, G.R., et al.: Single cell assay of a transcription factor reveals a threshold in transcription activated by signals emanating from the T-cell antigen receptor. *Genes Dev.* **4**, 1823–1834 (1990)
 30. Lin, Y.S., Carey, M., Ptashne, M., Green, M.R.: How different eukaryotic transcriptional activators can cooperate promiscuously. *Nature* **345**, 359–361 (1990)
 31. Weintraub, H.: Assembly and propagation of repressed and depressed chromosomal states. *Cell* **42**, 705–711 (1985)
 32. Santos-Rosa, H., Schneider, R., Bannister, A.J., Sherriff, J., Bernstein, B.E., et al.: Active genes are tri-methylated at K4 of histone H3. *Nature* **419**, 407–411 (2002)
 33. Margueron, R., Trojer, P., Reinberg, D.: The key to development: interpreting the histone code? *Curr. Opin. Genet. Dev.* **15**, 163–176 (2005)
 34. Kurdistani, S.K.: Histone modifications as markers of cancer prognosis: a cellular view. *Br. J. Cancer* **97**, 1–5 (2007)
 35. Nguyen, L.K., Munoz-Garcia, J., Maccario, H., Ciechanover, A., Kolch, W., et al.: Switches, excitable responses and oscillations in the Ring1B/Bmi1 ubiquitination system. *PLoS Comput. Biol.* **7**, e1002317 (2011)
 36. Seok, J., Warren, H.S., Cuenca, A.G., Mindrinos, M.N., Baker, H.V., et al.: Genomic responses in mouse models poorly mimic human inflammatory diseases. *Proc. Natl. Acad. Sci. U. S. A.* **110**, 3507–3512 (2013)
 37. Cheng, Y., Ma, Z., Kim, B.H., Wu, W., Cayting, P., et al.: Principles of regulatory information conservation between mouse and human. *Nature* **515**, 371–375 (2014)
 38. Lin, S., Lin, Y., Nery, J.R., Urlich, M.A., Breschi, A., et al.: Comparison of the transcriptional landscapes between human and mouse tissues. *Proc. Natl. Acad. Sci. U. S. A.* **111**, 17224–17229 (2014)
 39. Wittkopp, P.J., Kalay, G.: Cis-regulatory elements: molecular mechanisms and evolutionary processes underlying divergence. *Nat. Rev. Genet.* **13**, 59–69 (2012)
 40. Ackers, G.K., Johnson, A.D., Shea, M.A.: Quantitative model for gene regulation by lambda phage repressor. *Proc. Natl. Acad. Sci. U. S. A.* **79**, 1129–1133 (1982)
 41. Wang, J., Ellwood, K., Lehman, A., Carey, M.F., She, Z.S.: A mathematical model for synergistic eukaryotic gene activation. *J. Mol. Biol.* **286**, 315–325 (1999)
 42. Buchler, N.E., Gerland, U., Hwa, T.: On schemes of combinatorial transcription logic. *Proc. Natl. Acad. Sci. U. S. A.* **100**, 5136–5141 (2003)
 43. Segal, E., Raveh-Sadka, T., Schroeder, M., Unnerstall, U., Gaul, U.: Predicting expression patterns from regulatory sequence in *Drosophila* segmentation. *Nature* **451**, 535–540 (2008)
 44. Dresch, J.M., Liu, X., Arnosti, D.N., Ay, A.: Thermodynamic modeling of transcription: sensitivity analysis differentiates biological mechanism from mathematical model-induced effects. *BMC Syst. Biol.* **4**, 142 (2010)
 45. Frank, T.D., Carmody, A.M., Kholodenko, B.N.: Versatility of cooperative transcriptional activation: a thermodynamical modeling analysis for greater-than-additive and less-than-additive effects. *PLoS One* **7**, e34439 (2012)
 46. Shea, M.A., Ackers, G.K.: The OR control system of bacteriophage lambda. A physical-chemical model for gene regulation. *J. Mol. Biol.* **181**, 211–230 (1985)
 47. Frank, T.D., Cheong, A., Okada-Hatakeyama, M., Kholodenko, B.N.: Catching transcriptional regulation by thermostistical modeling. *Phys. Biol.* **9**, 045007 (2012)
 48. Bintu, L., Buchler, N.E., Garcia, H.G., Gerland, U., Hwa, T., et al.: Transcriptional regulation by the numbers: applications. *Curr. Opin. Genet. Dev.* **15**, 125–135 (2005)
 49. Bintu, L., Buchler, N.E., Garcia, H.G., Gerland, U., Hwa, T., et al.: Transcriptional regulation by the numbers: models. *Curr. Opin. Genet. Dev.* **15**, 116–124 (2005)
 50. Kulasiri, D., Nguyen, L.K., Samarasinghe, S., Xie, Z.: A review of systems biology perspective on genetic regulatory networks with examples. *Curr. Bioinform.* **3**, 197–225 (2008)
 51. Smolen, P., Baxter, D.A., Byrne, J.H.: Modeling transcriptional control in gene networks—methods, recent results, and future directions. *Bull. Math. Biol.* **62**, 247–292 (2000)

52. De Jong, H.: Modeling and simulation of genetic regulatory systems: a literature review. *J. Comput. Biol.* **9**, 67–103 (2002)
53. Chen, T., He, H.L., Church, G.M.: Modeling gene expression with differential equations. *Pac. Symp. Biocomput.* **4**, 29–40 (1999)
54. Babu, M.M., Luscombe, N.M., Aravind, L., Gerstein, M., Teichmann, S.A.: Structure and evolution of transcriptional regulatory networks. *Curr. Opin. Struct. Biol.* **14**, 283–291 (2004)

Pattern Formation at Cellular Membranes by Phosphorylation and Dephosphorylation of Proteins

Sergio Alonso

Abstract We consider a classical model on activation of proteins, based in two reciprocal enzymatic biochemical reactions. The combination of phosphorylation and dephosphorylation reactions of proteins is a well established mechanism for protein activation in cell signalling. We introduce different affinity of the two versions of the proteins to the membrane and to the cytoplasm. The difference in the diffusion coefficient at the membrane and in the cytoplasm together with the high density of proteins at the membrane which reduces the accessible area produces domain formation of protein concentration at the membrane. We differentiate two mechanisms responsible for the pattern formation inside of living cells and discuss the consequences of these models for cell biology.

1 Introduction

Cascades of biochemical reactions and interactions regulate multiple processes inside living cells [1]. Proteins, enzymes and small molecules strongly interact and participate in genetic and metabolic networks.

Biochemical processes inside cells are highly nonlinear and their dynamics complex. Two characteristics examples of such complexity in the cell are genetic regulatory networks and cell signalling [2]. Regulatory pathways involve different types of proteins, which control the transcription of the genes. They form the genetic networks in cell biology and governs processes at large times scales (hours). A protein that represses the transcription of its own gene is a simple example of a regulatory network. The repression produces a negative feedback and under certain conditions it induces a periodic synthesis of the protein [3].

S. Alonso (✉)

Departament de Física, Universitat Politècnica de Catalunya, Av. Dr. Marañón 44-50, 08028 Barcelona, Spain

e-mail: s.alonso@upc.edu

On the other hand, interactions among different proteins can also produce complex biochemical oscillations of protein concentration inside signalling pathways [4]. Cell signalling controls the rapid response of the cell to external variations of the environment, and corresponds to fast processes at small time scales (minutes). One of the simplest components of such pathways are enzymatic reactions. An enzyme, is a protein which catalyzes and accelerates a particular biochemical reaction, which under other circumstances would be much slower [5]. An important characteristics of the enzyme is that it takes part in the reaction, however, after the reaction occurs, the enzyme is completely recovered and can catalyze again another reaction.

A phosphorylation reaction consists in the incorporation of a phosphate group to a certain protein, producing a phosphorylated version of the protein. Such type of biochemical reactions are catalyzed by kinases. Protein kinases are enzymes and the phosphorylation reaction is an enzymatic reaction. The opposite reaction is also possible. The dephosphorylation reaction removes the phosphate group from the phosphorylated protein, and it is catalyzed by the enzyme phosphatase. Protein kinases and phosphatases are particularly active in signalling processes. Important parts of biochemical pathways consist on multiple phosphorylations and dephosphorylations of diverse proteins [6].

When both processes occurs simultaneously, phosphorylation and dephosphorylation may control the activity of a particular protein in a pathway. Activation and deactivation are important mechanisms in the regulation of many cellular processes. Both reactions are usually described in well-mixed environments by the Goldbeter-Koshland model of reciprocal covalent modifications [7]. The stiff response of the system to small changes in the kinase or the phosphatase concentrations makes the Goldbeter-Koshland mechanism a good model for activation of proteins in signalling pathways. The response of the model to the change of the enzyme concentration is stiff but monotonous and, therefore, no bistability or another kind of pattern formation mechanisms are accessible from this model. The complexity necessary for a non-monotonous behaviour can be incorporated by positive and negative interactions among the two versions of the proteins and the enzymes [8].

An alternative strategy for the formation of complex dynamics is the incorporation of spatial restrictions. Living cells are not always well-mixed environments and active or passive transport is crucial for the organization of some metabolic and genetic processes [9]. Furthermore, there is a high degree of compartmentalization in the cell and different types of proteins are located in different parts of the cell. Thus, the spatial aspects of the interior of the cells become relevant in intracellular communication [10] and self-organization may rule many processes in cell biology [11]. In particular, kinases and phosphatases may locate in opposite positions inside the cell, e.g. membrane/cytoplasm [12, 13], nucleus/cytoplasm [14] or anterior/posterior [15, 16]. It may produce the formation of spatial gradients in metabolic reactions. The formation of biochemical gradients may induce the polarization of the cell [17], and the definition of a preference direction for a posterior motion [18] or division [19] of a living cell.

The equations employed for the mathematical modeling of cell polarity typically consist in two components, a membrane protein with very slow diffusion in comparison with a second protein which diffuses faster through the cytoplasm. The interactions between these two components, which are highly nonlinear, produce a Turing-like instability of the homogeneous solution [20–22] or a wave-pinning dynamics due to the frustration under bistable conditions of the wave between both stable solution due to a mass-conservation condition [19, 23–25].

Here, we derive from the basis dynamics of enzymatic reactions a reaction-diffusion system of three equations representing the concentrations of the same protein at the membrane, phosphorylated and dephosphorylated in the cytoplasm. A similar set of equations have been previously successfully employed for the modeling of experimental observations on protein translocation results in an insulin-secreting cell [26, 27].

This chapter is organized as follows, first we review in Sect. 2 the derivation of the simplest model on reciprocal covalent modification composed by the phosphorylation and dephosphorylation processes. Second, we introduce in Sect. 3 the effects of compartemization and the effects of saturation at the membrane where the large amount of membrane proteins restricts the accessible area. Finally, in Sect. 4 the transport by diffusion is incorporated to the model of the biochemical reactions to generate the final reaction-diffusion model. The different mechanisms of pattern formation are described and analyzed.

2 Modelling Enzymatic Kinetics

The basis of any enzymatic reaction is the fast conversion of a substrate S into a product P . One is tempt to consider the next linear conversion



for the modeling of an enzymatic reaction, see the simple sketch in Fig. 1a of the reaction in Eq. (1). The velocity of reaction, which corresponds to the rate of production of $[P]$, has a simple linear relation with the concentration of the substrate and it is linearly proportional to the concentration of the enzyme:

$$\frac{\partial [P]}{\partial t} = k[E][S]; \quad (2)$$

where k is the rate of the reaction in Eq. (1). This simple dynamics holds when the number of substrate molecules is small in comparison with the capacity of enzymes to induce the reactions. If the number of substrate molecules is large, there is a delay due to the lack of available enzymes to perform the reaction. In this case the linear approximation shown in Eq. (1) is not correct.

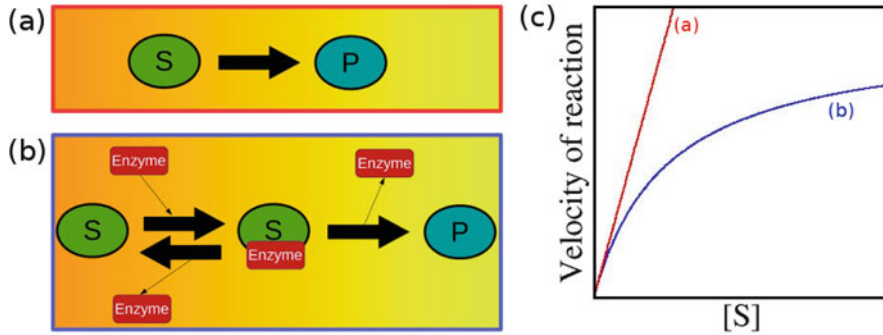


Fig. 1 (a) Sketch of the conversion of a substrate into a product through a linear reaction. (b) Sketch of the conversion of a substrate into a product through an enzymatic reaction. (c) Dependence of the velocity of reaction in the concentration of the substrate for the two types of reactions

2.1 Michaelis-Menten Model

Although the number of enzymes is the same before and after a reaction event, it participates in the reaction. Enzymes change the structure of the substrate to enhance the affinity to generate the product. We consider an intermediate step: the enzyme reacts with the substrate giving rise to a complex molecule C . This complex state may react and give rise to the product together with the original enzyme, however, there is a small probability than the complex C reacts in the opposite direction giving rise to the substrate and the enzyme. The three reactions together read:



and a schematic description of the reaction is shown in Fig. 1b. If we apply the law of mass action to the three reactions we obtain that the four concentrations follow [5]:

$$\begin{aligned} \frac{\partial[S]}{\partial t} &= k_{-1}[C] - k_1[S][E], \\ \frac{\partial[E]}{\partial t} &= (k_{-1} + k_2)[C] - k_1[S][E], \\ \frac{\partial[C]}{\partial t} &= -(k_{-1} + k_2)[C] + k_1[S][E], \\ \frac{\partial[P]}{\partial t} &= k_2[C]; \end{aligned} \quad (4)$$

where k_1 , k_{-1} and k_2 are the reaction rates for the three reactions shown in Eq. (3). Note that the total number of enzymes $[C] + [E] = [E_0]$ is conserved and that the

product is removed immediately from the system and, therefore, the reverse reaction $C \leftarrow P + E$ is not considered.

Assuming a quasi-steady approximation $\partial[C]/\partial t = 0$, the steady concentration of the complex C is derived from the equation of its evolution:

$$[C] = \frac{[S][E_0]}{K + [S]}; \quad (5)$$

where we have defined

$$K = \frac{k_{-1} + k_2}{k_1}; \quad (6)$$

which can be used in Eq. (4) to obtain the velocity of reaction as function of the concentration of the substrate:

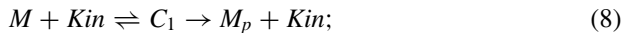
$$\frac{\partial[P]}{\partial t} = k_2[E_0] \frac{[S]}{K + [S]}; \quad (7)$$

within the condition $[S] \ll K$ we recover the prediction of the linear model, compare with Eq. (2) using $k = k_2[E_0]/K$. For large values of the substrate concentration, $[S] \gg K$, the velocity of the reaction saturates to a maximum velocity $V_{max} = k_2[E_0]$. For a comparison between the linear and the Michaelis-Menten models see Fig. 1c. While both types of dynamics coincide for small concentrations of the substrate, they differ for intermediate and large values.

2.2 Goldbeter-Koshland Model

There are multiple examples of enzymatic reactions in cell biology, but two of the most characteristics are the phosphorylation and dephosphorylation of proteins. They are close related because the product of the first is the substrate for the second reaction and the product of the second is the substrate for the first reaction. The protein develops a reciprocal covalent modification [7].

The Goldbeter-Koshland model incorporates the enzymatic dynamics to the mechanism of phosphorylation and dephosphorylation. Therefore, assuming Michaelis-Menten dynamics for both enzymes we arrive to the next set of biochemical reactions:



for the phosphorylation by an enzyme Kinase Kin of a protein M into a phosphorylated protein M_p , and



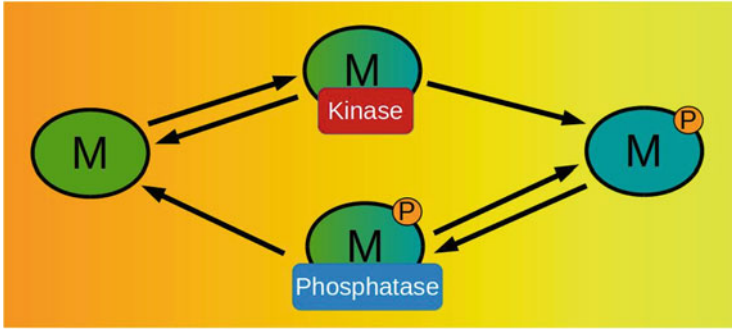


Fig. 2 Sketch of the reciprocal phosphorylation-dephosphorylation process of proteins. An enzyme kinase binds to the unphosphorylated protein to add a phosphate group. An enzyme phosphatase binds to the phosphorylated protein to remove the phosphate group

for the dephosphorylation by an enzyme phosphatase *Phos* of a protein M_p . A sketch of both reactions is shown in Fig. 2.

Finally, the evolution of the concentrations of the two types of proteins M and M_p can be expressed with two nonlinear equations after the assumption of quasi-static conditions for both complex $\partial[C_1]/\partial t = 0$ and $\partial[C_2]/\partial t = 0$:

$$\begin{aligned}\frac{\partial[M]}{\partial t} &= -G_1 \frac{[M]}{K_1 + [M]} + G_2 \frac{[M_p]}{K_2 + [M_p]} \\ \frac{\partial[M_p]}{\partial t} &= +G_1 \frac{[M]}{K_1 + [M]} - G_2 \frac{[M_p]}{K_2 + [M_p]}\end{aligned}\quad (10)$$

with $G_1 = k_2[Kin]$ and $G_2 = k_4[Phos]$ for the kinase and phosphatase controlled reaction rates, and $K_1 = (k_{-1} + k_2)/k_1$ and $K_2 = (k_{-3} + k_4)/k_3$ for the equilibrium reactions. Furthermore, the total number of proteins is conserved: $[T] = [M] + [M_p]$.

The steady state condition is obtained when $\partial[M]/\partial t = 0$ or equivalently $\partial[M_p]/\partial t = 0$. For a given set of parameter values only a single combination of values $[M]$ and $[M_p]$ is possible. It means that there is only a single solution.

With the tuning of a control parameter we may obtain large changes in the response. It permits the definition of two activation states. In this case, we consider G_1 as control parameter (equivalent analysis is possible with G_2), see for example Fig. 3. Two different states are obtained corresponding to high concentration of $[M]$ or to high concentration of $[M_p]$. Depending on the relative activity between the kinase and the phosphatase, see Fig. 3a, b, the solution of the steady state can be very different, see Fig. 3c–d. More important, the transition between these two states is not gradual but abrupt, a small change on the control parameter G_1 implies a big change in the response. This particular dynamics is employed to explain the activation of certain proteins in cell biology. For example, in the case of Fig. 3, if we assume that the active form of the protein is M_p , for a value of $G_1 = 9$ (low values

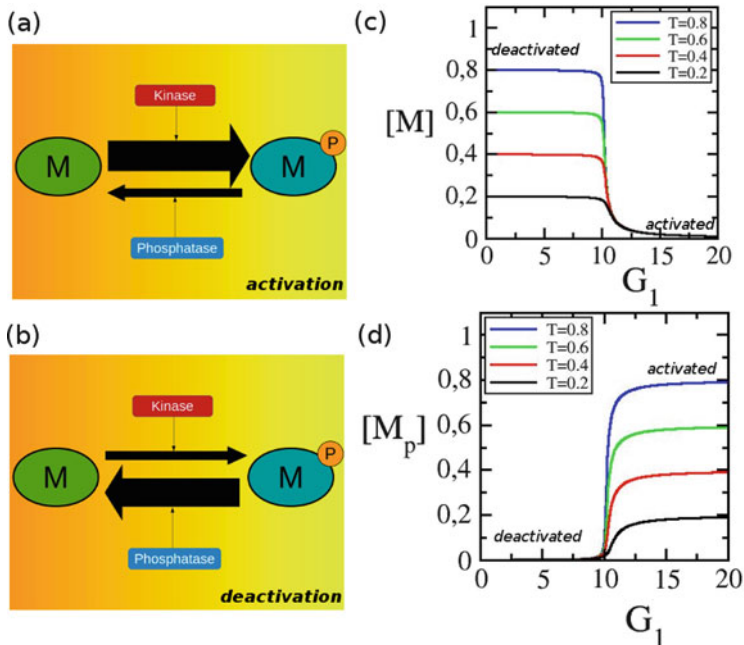


Fig. 3 (a) Sketch of a high rate of phosphorylation of a protein (large value of G_1). (b) Sketch of a low rate of phosphorylation of a protein (low value of G_1). (c), (d) Dependence of the equilibrium concentration of the inactive (c) and the active (d) version of the protein in the value of the rate G_1 for different values of the total concentration of protein ($T = [M] + [M_p]$)

of $[kin]$) the concentration of M_p is small, however, if there is a slightly increase of $[kin]$ the parameter G_1 changes ($G_1 = 11$) and the concentrations of M_p strongly increases, which is the activate state in this example.

2.3 Bistability and Complex Dynamics in Protein Activation

The Goldbeter-Koshland mechanism produces a monotonic dependence on the response. Although the abrupt change in this dependence can be use as activation mechanism in cell signalling, more complex dynamics are needed for the explanation of some other observations in cell biology [2].

Interactions among enzymes, substrates and products may induce a bistable switch [13]. For example, a positive feedback of the phosphorylated protein M_p in the enzyme kinase, modifies the Michaelis-Menten kinetics into a system with

higher order terms:

$$\begin{aligned}\frac{\partial[M]}{\partial t} &= -G_1 \frac{[M]}{K_1 + [M]} \left(1 + A \frac{[M_p]}{K_B + [M_p]} \right) + G_2 \frac{[M_p]}{K_2 + [M_p]}, \\ \frac{\partial[M_p]}{\partial t} &= +G_1 \frac{[M]}{K_1 + [M]} \left(1 + A \frac{[M_p]}{K_B + [M_p]} \right) - G_2 \frac{[M_p]}{K_2 + [M_p]};\end{aligned}\quad (11)$$

which induces bistability for a wide parameter range.

Although this interaction are observed in some experimental cases, equivalent terms are also obtained due to a saturation at the different compartments of the cell, as we discuss in the following sections.

3 Spatial Aspects of Enzymatic Kinetics

The interior of living cells is outside of the well-mixed approach because the cytoplasm is heterogeneous. Therefore, the concentrations of enzymes and proteins are not be homogeneous. Proteins may have tendency to accumulate in some parts of the cell. A typical example of this heterogeneous distributions is the effect of membranes. Enzymes can interact with the membranes and accumulate, for example in opposite regions of a bacteria [16], see Fig. 4a. A different case corresponds to the accumulation at the membrane of only one type of enzyme, for example the kinases in Fig. 4b, such inhomogeneous distribution induces a gradient between the interior and the exterior of the cytoplasm [13].

On the other hand, the spatial location of the protein implies important limitations in the protein distribution. For example, the space at the membrane is limited because of the high density of proteins and structures.

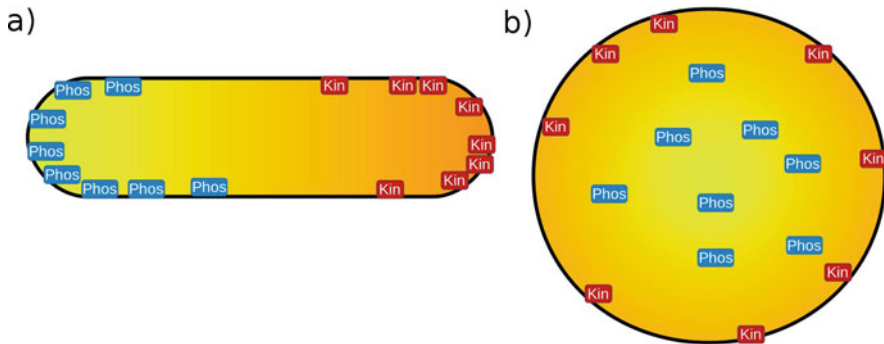


Fig. 4 Inhomogeneous affinity of enzymes. (a) Location of kinases and phosphatases at opposite poles in bacteria. (b) Location of kinases at cellular membrane and of phosphatases in the cytosol in Eukaryote

To model the effects of the compartments in the Goldbeter-Koshland mechanism, first, we introduce a constrain related with the accessible space at the membrane in the reaction equations for the protein concentrations. Second, we consider the intermediate concentration to account for the time of the unphosphorylated protein to diffuse to the membrane.

3.1 Saturation at the Membrane

The cellular membrane is a busy environment where the addition of large quantity of a new protein may occupy an extended region of the available area at the inner part of the membrane. It incorporates an extra constrain in the modeling of the dynamics of the protein concentration because the saturation of the membrane prevents the binding of new proteins from the cytoplasm.

3.1.1 Membrane-Controlled Binding

If the concentration of protein at the membrane approaches the saturation concentration, the binding rate decreases to zero. Assuming that the concentration of kinases is low, we renormalized the binding rate G_2 with a factor which accounts for the available space.

$$\begin{aligned}\frac{\partial[M]}{\partial t} &= -G_1 \frac{[M]}{K_1 + [M]} + G_2 \left(1 - \frac{[M]}{[M_S]}\right) \frac{[M_p]}{K_2 + [M_p]}, \\ \frac{\partial[M_p]}{\partial t} &= +G_1 \frac{[M]}{K_1 + [M]} - G_2 \left(1 - \frac{[M]}{[M_S]}\right) \frac{[M_p]}{K_2 + [M_p]};\end{aligned}\quad (12)$$

where M_S is the saturation concentration at the membrane. If $[M] = [M_S]$ the membrane is full, the prefactor $(1 - [M]/[M_S])$ is zero, and new proteins cannot bind to the membrane.

The linear stability analysis of Eqs. (12) shows the existence of three different solutions. However, for a give value of the control parameter G_1 only one of the three solutions is stable. Therefore, there is a monotonous behaviour on the control parameter, see Fig. 5a. In the limit $G_1 \rightarrow 0$ the concentration of proteins at the membrane approaches to $[M] = T$ for $T < [M_S]$, e.g. all the proteins are at the membrane, or to $[M] = [M_S]$ for $T > [M_S]$, because not all proteins can bind to the membrane. In the opposite limit, large values of G_1 , the membrane is empty $[M] = 0$.

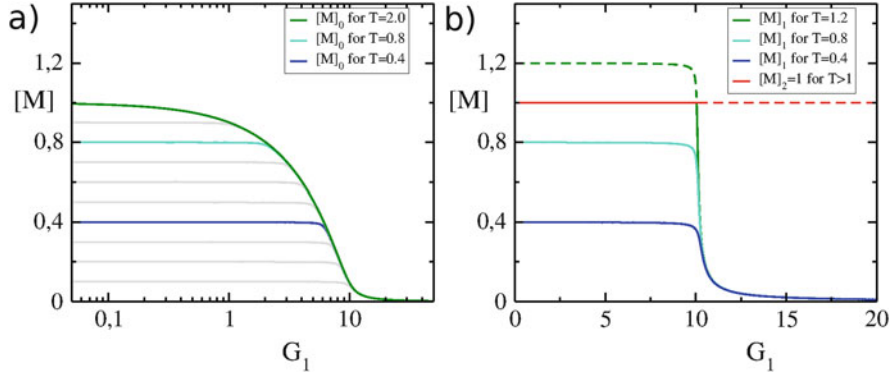


Fig. 5 Dependence of the physically relevant solution ($[M]_0$) of Eqs.(12) (a) and of the two physically relevant solutions ($[M]_1$ and $[M]_2$) of Eqs. (13) (b) of the inactive version of the protein in the value of the rate G_1 for different values of the total concentration of protein ($T = [M] + [M_p]$). Dashed lines correspond to unphysical or unstable solutions

3.1.2 Membrane-Controlled Reaction

A saturated membrane precludes the binding of all types of proteins including the enzymes. The high concentration of membrane-bound proteins incorporates an additional constrain to the binding of the kinase at the membrane. In such case, the unbinding rate of M , related with the enzymatic reaction leaded by the kinase, incorporates an equivalent constrain term than the binding rate, see the following set of nonlinear equations for the concentrations:

$$\begin{aligned} \frac{\partial[M]}{\partial t} &= -G_1 \left(1 - \frac{[M]}{[M_S]}\right) \frac{[M]}{K_1 + [M]} + G_2 \left(1 - \frac{[M]}{[M_S]}\right) \frac{[M_p]}{K_2 + [M_p]}, \\ \frac{\partial[M_p]}{\partial t} &= +G_1 \left(1 - \frac{[M]}{[M_S]}\right) \frac{[M]}{K_1 + [M]} - G_2 \left(1 - \frac{[M]}{[M_S]}\right) \frac{[M_p]}{K_2 + [M_p]}; \end{aligned} \quad (13)$$

where M_S is the saturation concentration at the membrane. The high concentration of proteins at the membrane inhibits both binding and unbinding by phosphorylation of proteins.

The linear stability analysis of Eqs. (13) reveals the existence of two physical solutions. One of the solutions corresponds to the complete saturation of the membrane with $[M]_2 = M_S$. Increasing the control parameter G_1 , the complete saturation condition becomes unstable and the new stable solution with $[M]_1 < M_S$ decreases to $[M]_1 = 0$ at large values of G_1 . The combination of the two stable solutions produces a monotonous response and no bistability is obtained.

In Fig. 5b three cases are studied for different values of T , keeping $M_S = 1$. If $T < 1$ there is a unique solutions $[M]_1$ for all the values of G_1 . In this case, for $G_1 < 10$ the value of $[M]_1$ saturates to $M = T$. Such saturation is unphysical for $T > 1$ and the solution $M = T$ is not possible. In this case, the second solution $[M]_2 = M_S = 1$ is stable for $G_1 < 10$ and exchanges stability with the solution $[M]_1$ at $G_1 = 10$. For large values of G_1 the solutions $[M]_2 = M_S$ is not stable, see Fig. 5b.

3.2 Cytosolic Diffusion

The models previously discussed relate unphosphorylated proteins to the membrane and phosphorylated proteins to the cytoplasm. With these assumptions we neglect the concentration of phosphorylated proteins at the membrane and unphosphorylated proteins in the cytosol. The first assumption seems adequate because there is an immediate lose of affinity of the proteins to the membrane after phosphorylation. However, after the dephosphorylation reaction the resulting protein needs to diffuse to the membrane and the binding to the membrane is not immediate. A third concentration of unphosphorylated cytosolic protein $[M_c]$ can be considered.

In summary, the protein is translocated from the membrane when it is phosphorylated by a kinase. Back in the cytoplasm, the translocated proteins are dephosphorylated by a phosphatase. The resulting unphosphorylated proteins diffuse and bind again at the membrane. These three processes give rise to a cyclic dynamics, see Fig. 6. We derive the next set of equations for the three concentrations:

$$\begin{aligned}\frac{\partial[M]}{\partial t} &= -G_1 \left(1 - \frac{[M]}{[M_S]}\right) \frac{[M]}{K_1 + [M]} + G_3 \left(1 - \frac{[M]}{[M_S]}\right) [M_c], \\ \frac{\partial[M_p]}{\partial t} &= +G_1 \left(1 - \frac{[M]}{[M_S]}\right) \frac{[M]}{K_1 + [M]} - G_2 \frac{[M_p]}{K_2 + [M_p]}, \\ \frac{\partial[M_c]}{\partial t} &= -G_3 \left(1 - \frac{[M]}{[M_S]}\right) [M_c] + G_2 \frac{[M_p]}{K_2 + [M_p]},\end{aligned}\tag{14}$$

where G_3 is the binding rate to the membrane of the cytosolic proteins. We assume that the affinity to the membrane is linear on $[M_c]$ and it is penalized by a possible saturation of the membrane. Note also that the total number of proteins is conserved $[M] + [M_c] + [M_p] = T$, giving rise to a mass-conserved model [25].

The linear stability analysis of Eqs. (14) shows the simultaneous existence of three different physically relevant steady solutions for a window of values of the control parameter. There is non-monotonous dependence on the parameter G_1 and bistability appears, see Fig. 7.

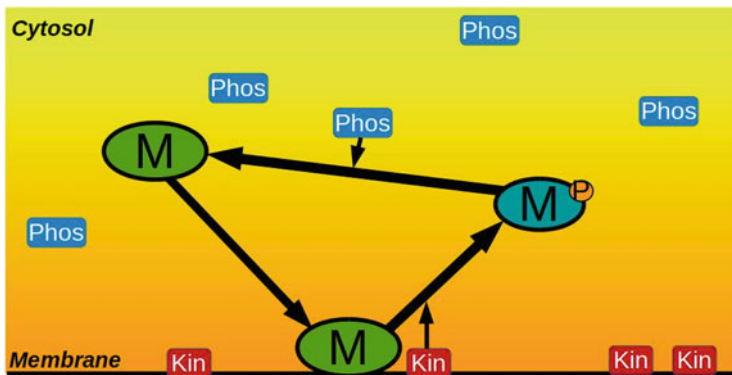


Fig. 6 Cyclic dynamics of the protein, from the membrane to the cytoplasm, phosphorylation, and from the cytoplasm back to the membrane

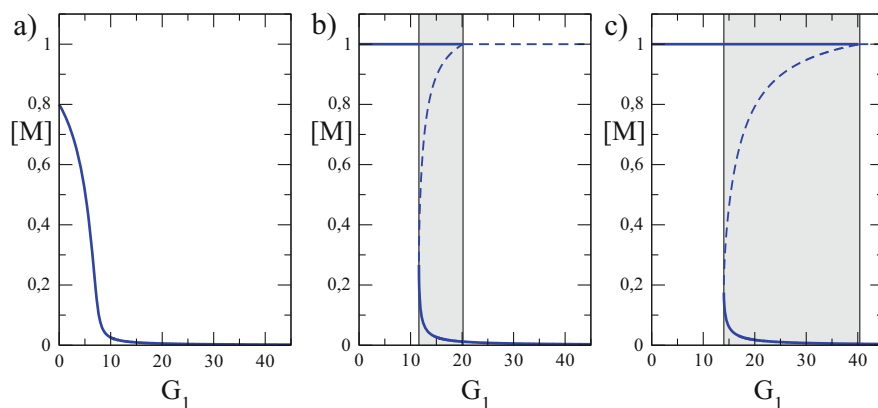


Fig. 7 Dependence of the physically relevant solutions of the inactive version of the protein in the value of the rate G_1 for different values of the total concentration of protein ($T = [M] + [M_p] + [M_c]$): $T = 0.8$ (a), $T = 1.2$ (b), and $T = 1.4$ (c). Solid and dashed lines correspond, respectively, to stable and unstable solutions. Gray areas mark region of bistability where two possible values of $[M]$ are stable

For a total number of proteins below the saturation value $T < [M_s]$ there is only one solution and its dependence on G_1 is monotonous, see Fig. 7a. The concentration at the membrane decreases to zero for large values of G_1 .

For $T > [M_s]$ a new solution is possible. It consists in a completely saturated membrane, full of proteins $[M] = [M_s]$, and the excess of proteins are located in the cytoplasm, see Fig. 7b, c. It produces a bistable condition for a window of values of the parameter G_1 . The saturation condition for the membrane is not stable for large values of G_1 , see Fig. 7b, c.

4 Reaction-Diffusion Model of Phosphorylation and Dephosphorylation of Proteins

Now, we explicitly consider the spatial distribution of the proteins with the use of spatio-temporal equations for the three concentrations. We take into account the diffusion (D_c) of the unphosphorylated (M_c) and phosphorylated (M_p) proteins at the cytoplasm and the diffusion of the unphosphorylated bound proteins at the membrane (D_m). The set of reaction-diffusion equations read:

$$\begin{aligned}\frac{\partial[M]}{\partial t} &= -G_1 \left(1 - \frac{[M]}{[M_S]}\right) \frac{[M]}{K_1 + [M]} + G_3 \left(1 - \frac{[M]}{[M_S]}\right) [M_c] + \nabla \cdot D_m \nabla [M], \\ \frac{\partial[M_p]}{\partial t} &= +G_1 \left(1 - \frac{[M]}{[M_S]}\right) \frac{[M]}{K_1 + [M]} - G_2 \frac{[M_p]}{K_2 + [M_p]} + \nabla \cdot D_c \nabla [M], \\ \frac{\partial[M_c]}{\partial t} &= -G_3 \left(1 - \frac{[M]}{[M_S]}\right) [M_c] + G_2 \frac{[M_p]}{K_2 + [M_p]} + \nabla \cdot D_c \nabla [M];\end{aligned}\quad (15)$$

where the diffusion of the proteins in the cytosol is higher than at the membrane ($D_m \ll D_c$). The characteristic values of the diffusion coefficient in the cytosol are around two orders of magnitude larger than the value of a equivalent molecule at the membrane [28].

For the integration of Eqs. (15) we have to define adequate boundary conditions. One possibility is the use of non-flux boundaries in a one-dimensional approach of the cell, see Fig. 8a. This type of models are commonly used for the description of cell polarity [19, 20, 23, 25]. We employ this model for the calculation of the linear stability analysis in Sect. 4.1. On the other hand, we can employ the cell membrane

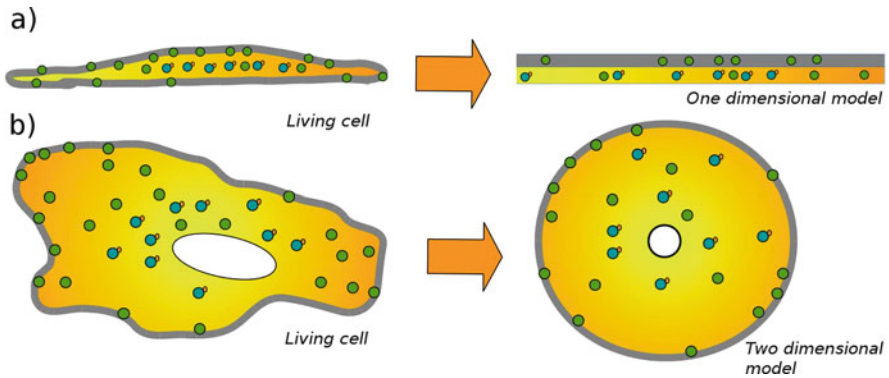


Fig. 8 Geometry reduction for the implementation of the model on Pattern formation inside living cells. (a) Reduction to a simple one dimensional geometry using to different values for the concentration at the membrane and in the cytoplasm. (b) Simplification of the cell to symmetric two-dimensional geometry with no-flux boundary conditions and binding to the membrane

as boundary conditions in a two-dimensional domain, see Fig. 8b. Such type of models has been employed to simulate insulin-secreting cells [27] and in polarity of yeast cell [21]. We use such approach to perform the numerical simulations appearing in Sects. 4.2 and 4.3.

4.1 Linear Stability Analysis

The number of homogeneous solutions depends on the parameter values. Changing the parameters T and G_1 two different zones appear: a region where a single solution is possible and a region where three solutions appear. To study its stability we calculate the linear stability analysis of Eqs. (15). For a particular homogeneous steady state, composed by the concentrations $[M]_0$, $[M_c]_0$, and $[M_p]_0$, we introduce a perturbation:

$$\begin{aligned} [M] &= [M]_0 + (\delta M)e^{\omega t + ikx}, \\ [M_c] &= [M_c]_0 + (\delta M_c)e^{\omega t + ikx}, \\ [M_p] &= [M_p]_0 + (\delta M_p)e^{\omega t + ikx}, \end{aligned} \quad (16)$$

and evaluate if the perturbations grow or decrease with time using Eqs. (15). The variable ω is the growing rate and indicates the stability of the solution to small perturbations.

The results of the linear stability analysis of Eqs. (15) is plotted in Fig. 9. It shows a region where only two of the three solutions are stable (bistability),

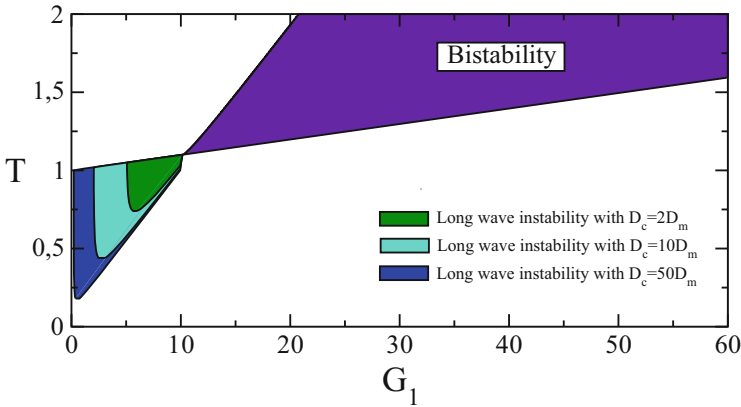


Fig. 9 Analytic phase diagram of the reaction-diffusion system, see Eqs. (15). The linear stability problem of the homogeneous solutions is solved in the parameter space defined by G_1 and T for three different ratios D_m/D_c

corresponding to the situation described in the previous Sect. 3.2, and a region where the unique physically relevant solution is unstable due to a long-wave instability. This instability produces spontaneous domain formation and, because of the conserved protein concentration, the posterior coarsening into a reduced number of domains.

While bistability is independent of the diffusion coefficients, the long wave instability appears initially at a given value of wavenumber which will depend on the quotient D_m/D_c . The area of the phase diagram where the solution is unstable changes with the diffusion as it is shown in Fig. 9 for three values of the quotient D_m/D_c .

Two different types of pattern formation, with different dynamics, are expected depending on the parameter values. Next we analyze separately both mechanisms.

4.2 Long-Wave Instability

For the parameter values inside the region of long-wave instability, a one-dimensional system, see Sketch in Fig. 8a, spontaneously develops the formation of domains as predicted by the linear stability analysis. If we change the symmetry of the integration domain, see Sketch in Fig. 8b, the linear stability analysis shown in the previous section cannot be applied directly. However, it is known [26] that the parameter values can be renormalized considering the size of the two-dimensional cytoplasm in comparison with the one-dimensional membrane.

Numerical simulations for a convenient choice of the parameter values are shown in Fig. 10. First, the spatio-temporal plot in Fig. 10a shows the evolution of the concentration of membrane-bound proteins. From an initially homogeneous condition with a small spatially distributed random perturbation, two evolving maxima appear. While one of the domains grows the other one decreases, and finally, only one single large domain survives. Such competition among the domains is a typical signature of coarsening.

In panel (b) of Fig. 10 the initial condition is plotted, an homogeneous concentration with a small random perturbation around the unstable value. Two-dimensional panels with the distribution of free protein concentration and phosphorylated concentration at the membrane and in the cytosol are also shown.

In the other two panels (c) and (d) of Fig. 10 the spatial distributions of the concentrations are shown at two different times. Note that the concentrations of $[M]_c$ and $[M]_p$ are complementary: large (small) values of $[M]_c$ coincide with small (large) $[M]_p$. Furthermore, $[M]_p$ accumulates in the region of the cytoplasm close to regions of the membrane where no proteins are bound.

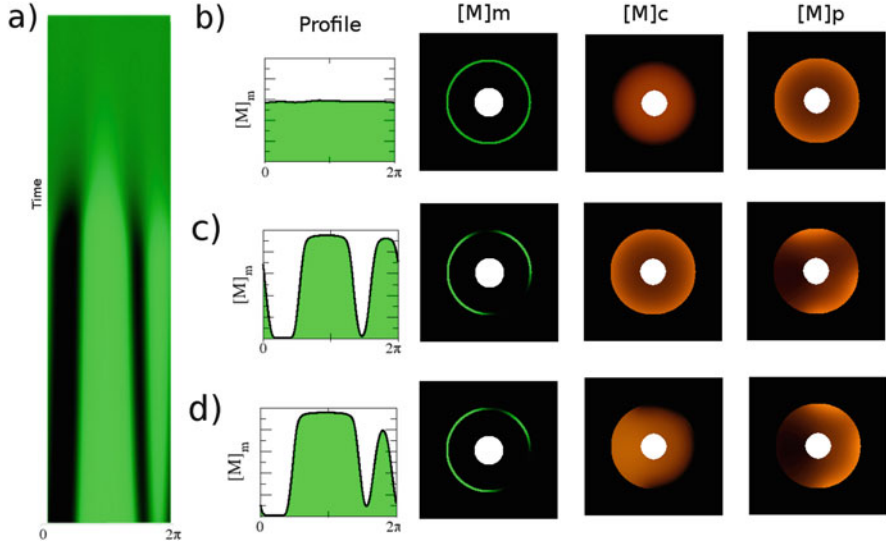


Fig. 10 Pattern formation by the a long wave instability mechanism. A circular domain is employed, representing the interior and the membrane of a living cell. **(a)** Spatio-temporal plot of the concentration of membrane-bound protein. **(b)–(d)** Profiles of the membrane protein and spatial distribution of the membrane-bound, cytosolic unphosphorylated and phosphorylated protein at times: $t = 0$ **(b)**, $t = t_{tot}/2$ **(c)**, and $t = t_{tot}$ **(d)**, where t_{tot} is the total time of the numerical simulation

4.3 Bistability-Induced Instability

Numerical simulations in the bistable region of the parameter space produce similar final domains than numerical results shown in the previous section, see a representative example in Fig. 11. However, the mechanism and the conditions are different. Under bistability the two homogeneous solution are stable, and, therefore, an small perturbation of the homogeneous solution decreases, and eventually, the homogeneous condition is recovered.

However, if the two solutions are connected by a front, see Fig. 11b, it moves as it is shown in the spatio-temporal plot in Fig. 11a. In contrast to a classical bistable system, here there is a mass-conserved condition which precludes the complete translocation to the membrane of the proteins in the cytoplasm. Finally, the two solutions are separated by an pinned wave [23], see Fig. 11c, d.

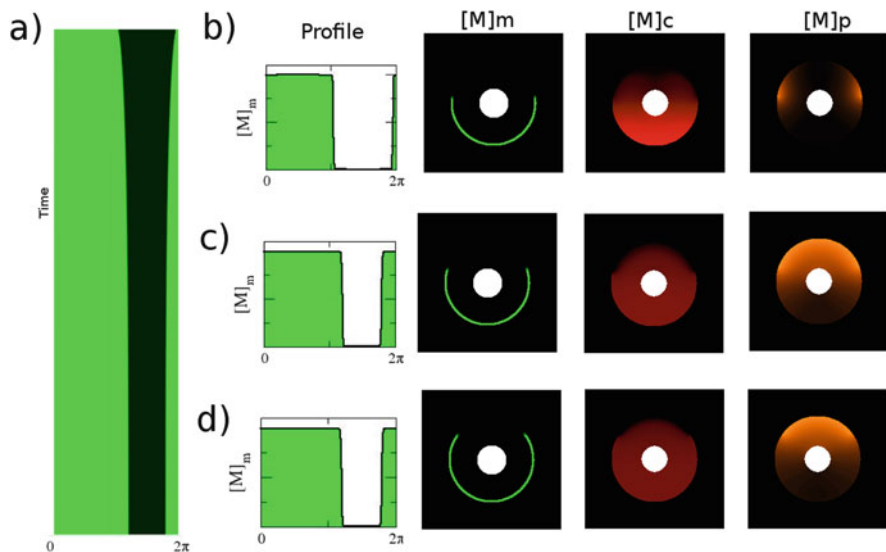


Fig. 11 Pattern formation by the bistable-induced mechanism. A circular domain is employed, representing the interior and the membrane of a living cell. (a) Spatio-temporal plot of the concentration of membrane-bound protein. (b)–(d) Profiles of the membrane protein and spatial distribution of the membrane-bound, cytosolic unphosphorylated and phosphorylated protein at times: $t = 0$ (b), $t = t_{tot}/2$ (c), and $t = t_{tot}$ (d), where t_{tot} is the total time of the numerical simulation

5 Conclusions

We employ an extended version of the classical Goldbeter-Koshland model on covalent modification in biological systems for phosphorylation and dephosphorylation of proteins. From a simple activation model based in two opposite enzymatic reactions following Michaelis-Menten kinetics, we have generated a scenario for pattern formation in the interior of a living cell. The mechanism is based on the different localization inside the cell of the two enzymatic reactions: while phosphorylation only occurs at the membrane of the cell, the opposite reaction occurs in the cytoplasm. After dephosphorylation the resulting proteins diffuse to the membrane where they bind and the cycle can start again.

We analyze two different mechanisms of pattern formation for protein at membranes: A long wave instability and a bistability-related mechanism, previously described in models of cell polarization [20] and [23] respectively. The conditions of the model to fulfill the requirements for pattern formation are simple and generic:

- *Difference on the diffusion coefficients.* The first important condition is the large diffusion in the cytoplasm in comparison with the diffusion in the membrane $D_c \gg D_m$ [28]. This constrain is naturally achieved by living cells where

diffusion of proteins in the cytoplasm is estimated to be around $D_c \sim 1 - 10 \mu\text{m}^2/\text{s}$ [29], while membrane diffusion coefficient of proteins in mammalian cells has been estimated to be $D_m \sim 0.1 \mu\text{m}^2/\text{s}$ [30].

- *Compartmentalization of the reaction.* The separation of the two enzymatic processes in two different compartments of the cell (phosphorylation at the membrane and dephosphorylation in the cytoplasm) introduces a delay between dephosphorylation and the re-binding of the protein to the membrane. This temporal delay accounts for the diffusion time of the proteins in the cytoplasm.
- *Mass-conservation.* The conservation of the total number of proteins is an important condition for the bistability-related mechanism, this constrain stops the bistable front going from the metastable to a stable solution [23].

Similar extensions of the Goldbeter-Koshland mechanism may be applicable to a large variety of biological systems. Phosphorylation by kinases regulates multiple processes in living cells, e.g. the formation of polarity of cells induced by PAR proteins [19], the cyclic dynamics of Rho GTPases [20] or the regulation of the cell division of *E. Coli* controlled by the Min proteins [31].

The approach described here has been employed in the modeling of the myristoyl-electrostatic switch [32] composed by a cyclic binding and unbinding dynamics of the myristoylated alanine-rich C kinase substrate proteins (MARCKS). After phosphorylation of MARCKS proteins by protein kinase C (PKC), MARCKS proteins lose their affinity to membranes. The phosphates reduce the positive charge of the protein and cause the unbinding from the membrane. In the cytoplasm, phosphatases remove the phosphates from the protein and, consequently, MARCKS can bind again at the membrane. This system has been recently described in terms of mass-conserved reaction-diffusion equations [26, 27] and the resulting equations have been employed for the calibration of the binding of the MARCKS at lipid monolayers [33, 34].

In our approach stochastic effects due to a low number of proteins have been neglected. The typical large concentration of proteins in cells permits the use of deterministic dynamics, however, the concentration of enzymes is smaller and stochastic effects may become relevant. The use of a stochastic model may enhance domain formation [35, 36].

Living cells are three-dimensional and future models will take this condition into account. The mechanisms are, however, equivalent at different levels of spatial complexity. One-dimensional approximation is employed to calculate linear stability analysis and identify the instabilities. The two-dimensional view is, however, sufficient to perform numerical simulations and describe proteins concentrations at the membrane and in the different regions of the cytoplasm, see Figs. 10 and 11. Such two-dimensional modeling considers the cytoplasm volume and diffusion orthogonal to the membrane. The proteins can diffuse from the membrane and are, hence, diluted near the membrane. As the cytosolic volume increases, the concentration of proteins close to the membrane decreases and the binding process is affected. However, the change of the cytosolic volume is equivalent to a renormalization of some of the reaction rates in the model [26].

In summary, we have developed a simple model of binding, phosphorylation and desphosphorylation for membrane proteins, which predicts the spontaneous appearance of domains of high protein concentration at the membrane of living cells.

Acknowledgements I acknowledge fruitful discussions with professor Markus Bär. Support by MINECO of Spain under the Ramon y Cajal program with the grant number RYC-2012-11265 is acknowledged.

References

1. Alberts, B., Johnson, A., Lewis, J., Raff, M., Roberts, K., Walter, P.: *Molecular Biology of the Cell*. Garland Science, New York (2002)
2. Tyson, J.J., Chen, K.C., Novak, B.: Sniffers, buzzers, toggles and blinkers: dynamics of regulatory and signaling pathways in the cell. *Curr. Opin. Cell Biol.* **15**, 221–231 (2003)
3. Novák, B., Tyson, J.J.: Design principles of biochemical oscillators. *Nat. Rev. Mol. Cell Biol.* **9**, 981–991 (2008)
4. Tyson, J.J., Novák, B.: Functional motifs in biochemical reaction networks. *Annu. Rev. Phys. Chem.* **61**, 219 (2010)
5. Keener, J., Sneyd, J.: *Mathematical Physiology*. Springer, New York (2009)
6. Kholodenko, B.N., Hancock, J.F. and Kolch, W.: Signalling ballet in space and time. *Nat. Rev. Mol. Cell Biol.* **11**, 414–426 (2010)
7. Goldbeter, A., Koshland, D.E.: An amplified sensitivity arising from covalent modification in biological systems. *Proc. Natl. Acad. Sci.* **78**, 6840–6844 (1981)
8. Straube, R., Conradi, C.: Reciprocal enzyme regulation as a source of bistability in covalent modification cycles. *J. Theor. Biol.* **330**, 56–74 (2013)
9. Bressloff, P.C., Newby, J.M.: Stochastic models of intracellular transport. *Rev. Mod. Phys.* **85**, 135 (2013)
10. Dehmelt, L., Bastiaens, P.I.: Spatial organization of intracellular communication: insights from imaging. *Nat. Rev. Mol. Cell Biol.* **11**, 440–452 (2010)
11. Karsenti, E.: Self-organization in cell biology: a brief history. *Nat. Rev. Mol. Cell Biol.* **9**, 255–262 (2008)
12. Kholodenko, B.N., Hoek, J.B., Westerhoff, H.V.: Why cytoplasmic signalling proteins should be recruited to cell membranes. *Trends Cell Biol.* **10**, 173–178 (2000)
13. Kholodenko, B.N.: Cell-signalling dynamics in time and space. *Nat. Rev. Mol. Cell Biol.* **7**, 165–176 (2006)
14. Maeder, C.I., Hink, M.A., Kinkhabwala, A., Mayr, R., Bastiaens, P.I., Knop, M.: Spatial regulation of Fus3 MAP kinase activity through a reaction-diffusion mechanism in yeast pheromone signalling. *Nat. Cell Biol.* **9**, 1319–1326 (2007)
15. Suzuki, A., Ohno, S.: The PAR-aPKC system: lessons in polarity. *J. Cell Sci.* **119**, 979–987 (2006)
16. Tropini, C., Rabbani, N., Huang, K.C.: Physical constraints on the establishment of intracellular spatial gradients in bacteria. *BMC Biophys.* **5**, 17 (2012)
17. Mogilner, A., Allard, J., Wollman, R.: Cell polarity: quantitative modeling as a tool in cell biology. *Science* **336**, 175–179 (2012)
18. Van Haastert, P.J., Devreotes, P.N.: Chemotaxis: signalling the way forward. *Nat. Rev. Mol. Cell Biol.* **5**, 626–634 (2004)
19. Goehring, N.W., Trong, P.K., Bois, J.S., Chowdhury, D., Nicola, E.M., Hyman, A.A., Grill, S.W.: Polarization of PAR proteins by advective triggering of a pattern-forming system. *Science* **334**, 1137–1141 (2011)

20. Otsuji, M., Ishihara, S., Kaibuchi, K., Mochizuki, A., Kuroda, S.: A mass conserved reaction-diffusion system captures properties of cell polarity. *PLoS Comput. Biol.* **3**, e108 (2007)
21. Goryachev, A.B., Pokhilko, A.V.: Dynamics of Cdc42 network embodies a Turing-type mechanism of yeast cell polarity. *FEBS Lett.* **582**, 1437–1443 (2008)
22. Orlandini, E., Marenduzzo, D., Goryachev, A.B.: Domain formation on curved membranes: phase separation or Turing patterns?. *Soft Matter* **9**, 9311–9318 (2013)
23. Mori, Y., Jilkine, A., Edelstein-Keshet, L.: Wave-pinning and cell polarity from a bistable reaction-diffusion system. *Biophys. J.* **94**, 3684–3697 (2008)
24. Beta, C., Amselem, G., Bodenschatz, E.: A bistable mechanism for directional sensing. *N. J. Phys.* **10**, 083015 (2008)
25. Trong, P.K., Nicola, E.M., Goehring, N.W., Kumar, K.V., Grill, S.W.: Parameter-space topology of models for cell polarity. *N. J. Phys.* **16**, 065009 (2014)
26. Alonso, S., Bär, M.: Phase separation and bistability in a three-dimensional model for protein domain formation at biomembranes. *Phys. Biol.* **7**, 046012 (2010)
27. Alonso, S., Bär, M.: Modeling domain formation of MARCKS and protein kinase C at cellular membranes. *EPJ Nonlinear Biomed. Phys.* **2**, 1–18 (2014)
28. Edidin, M., Zagyansky, Y., Lardner, T.J.: Measurement of membrane protein lateral diffusion in single cells. *Science* **191** 466–468 (1976)
29. Luby-Phelps, K.: Cytoarchitecture and physical properties of cytoplasm: volume, viscosity, diffusion, intracellular surface area. *Int. Rev. Cytol.* **192**, 189–221 (1997)
30. Valdez-taubas, J., Pelham, R.B.: Slow diffusion of proteins in the yeast plasma membrane allows polarity to be maintained by endocytic cycling. *Curr. Biol.* **13** 1636–1640 (2003)
31. Howard, M., Kruse, K.: Cellular organization by self-organization: mechanisms and models for Min protein dynamics. *J. Cell Biol.* **168**, 533–36 (2005)
32. Mogami, H., Zhang, H., Suzuki, Y., Urano, T., Saito, N., Kojima, I., Petersen, O.H.: Decoding of short-lived Ca^{2+} influx signals into long term substrate phosphorylation through activation of two distinct classes of protein kinase C. *J. Biol. Chem.* **278**, 9896–9904 (2003)
33. Alonso, S., Dietrich, U., Händel, C., Käs, J.A., Bär, M.: Oscillations in the lateral pressure of lipid monolayers induced by nonlinear chemical dynamics of the second messengers MARCKS and protein kinase C. *Biophys. J.* **100**, 939–947 (2011)
34. Lippoldt, J., Händel, C., Dietrich, U., Käs, J.A.: Dynamic membrane structure induces temporal pattern formation. *Biochim. Biophys. Acta Biomembr.* **1838**, 2380–2390 (2014)
35. Lawson, M.J., Drawert, B., Khammash, M., Petzold, L., Yi, T.M.: Spatial stochastic dynamics enable robust cell polarization. *PLoS Comput. Biol.* **9**, e1003139 (2013)
36. McKane, A.J., Biancalani, T., Rogers, T.: Stochastic pattern formation and spontaneous polarisation: the linear noise approximation and beyond. *Bull. Math. Biol.* **76**, 895–921 (2014)

An Introduction to Mathematical and Numerical Modeling of Heart Electrophysiology

Luca Gerardo-Giorda

Abstract The electrical activation of the heart is the biological process that regulates the contraction of the cardiac muscle, allowing it to pump blood to the whole body. In physiological conditions, the pacemaker cells of the sinoatrial node generate an action potential (a sudden variation of the cell transmembrane potential) which, following preferential conduction pathways, propagates throughout the heart walls and triggers the contraction of the heart chambers.

The action potential propagation can be mathematically described by coupling a model for the ionic currents, flowing through the membrane of a single cell, with a macroscopical model that describes the propagation of the electrical signal in the cardiac tissue. The most accurate model available in the literature for the description of the macroscopic propagation in the muscle is the Bidomain model, a degenerate parabolic system composed of two nonlinear partial differential equations for the intracellular and extracellular potential. In this paper, we present an introduction to the fundamental aspects of mathematical modeling and numerical simulation in cardiac electrophysiology.

1 Introduction

Cardiac-specific diseases account for 700,000 deaths each year in Europe, half of this mortality being due to heart failure (ineffective contraction, principally due to ventricular dyssynchrony). The other half of cardiac mortality occurs suddenly, essentially due to ventricular tachyarrhythmias. Although the vast majority of these cases is associated with chronic cardiac disease, sudden cardiac death can also occur in seemingly healthy and sometimes very young people. Knowledge about the underlying causes and options for diagnosis and prevention is still very limited. Still, with six million individuals suffering from atrial fibrillation, nine million people affected by heart failure, and 350,000 sudden deaths every year, the human and economic burden of cardiac electrical diseases skyrocketed in Europe. The

L. Gerardo-Giorda (✉)

Basque Center for Applied Mathematics (BCAM), Alameda Mazarredo 14, 48009 Bilbao, Spain

e-mail: lgerardo@bcamath.org

estimated annual direct costs for the health care system in Europe is topping 1100 billion Euros.

The contraction of the heart is orchestrated by a complex mechanism of electrical activation. As a consequence, an important part of heart failure occurrences is caused or aggravated by electrical dysfunctions, and heart electrophysiology has become in the recent years the subject of a vast interdisciplinary literature, from medical sciences through bio-engineering, physiology, chemistry and physics.

The electrical activity of the heart as a whole is characterized by a complex multiscale structure, ranging from the microscopic activity of ion channels in the cellular membrane to the macroscopic properties of the anisotropic propagation of the excitation and recovery fronts in the whole heart. Cardiac arrhythmias are complex disruptions of this organisation.

Cardiac cells, called myocytes, are a particular type of excitable cells. At resting condition, they feature a negative transmembrane potential, resulting from the difference between internal and external concentrations of charged ions ($[Na]^+$, $[K]^+$, $[Ca]^{++}$). A small change of the cell's environment from its rest state produces a very fast depolarization, followed by a slower repolarization process towards the resting state. This cellular activity is called Action Potential (AP), and is mathematically described by an ionic ODE model which is the basis of dynamic behaviour in the model. Modern detailed ionic models take into account transmembrane current flows, intracellular calcium handling, and can include energetics and force production [8, 41].

Higher levels of complexity govern the propagation of the electrical impulse for optimal contraction at the tissue and organ levels. The global activation sequence of the heart follows from the physical organization of a special conduction network that is essential for the synchronization of the whole heart. In healthy conditions, atria and ventricles are electrically insulated from each other and are connected only through the atrioventricular node (AV). The AP originates in the sinoatrial node (SA), propagates in the atria through Bachmann's bundle, it is modulated in frequency by the AV, and proceeds to the ventricles, where the His-Purkinje system (PS) provides a preferential pathway for the AP to propagate through the lower chambers.

The difficulty in having access to direct measures on real patients fueled the interest in mathematical modeling and numerical simulations, which have been supporting cardiovascular science for more than 20 years. In this respect, cardiac modeling in medicine has significantly evolved in the recent decades, providing the best and highly detailed mathematical description of any organ system in the body. Many fundamental insights have been gained from in-silico experiments [65], even ahead of experimental evidence [55]. Numerical models and fast dedicated solvers already exist and allow in-silico exploration of the mechanisms underlying these pathologies at the cost of large-scale simulations.

If on the one hand, modern imaging techniques, such as high-resolution magnetic resonance imaging (MRI), allows high level of accuracy in the description of both the microstructure of the tissue and the global anatomy of the organ, current mathematical models are based on a formalism (the Bidomain equations, [16, 38]) whose

derivation is based on a very simplified geometry with respect to the nowadays available structural data. The heart is, in fact, a collection of interconnected excitable tissues, each of which has specific modelling requirements.

Each of these tissues is a complex network of cardiomyocytes connected to each other by gap-junctions, together with other types of cells (fibroblasts) and collagen. APs propagate from cell to cell, resulting in AP waves at the macroscopic tissue scale. From a mathematical standpoint, a multiscale technique allows to model the tissue, at the macroscopic level, as a continuum where the intra- and extra-cellular media are superimposed and the corresponding potentials are the solutions to a system of degenerate partial differential equation of reaction-diffusion type coupled over space with the system of ODEs (the ionic model). Such model is known as the Bidomain system of equations. In this model, the anisotropies of the intra- and extra-cellular conductivities differ. In case of an insulated tissue and under the so-called equal anisotropy ratio assumption, the system reduces to a single reaction-diffusion equation: the Monodomain equation. The Monodomain equation is no longer degenerate, thus far cheaper to solve numerically [34].

The numerical approximation of the Bidomain model is often based on a finite element discretization in space and on implicit-explicit time advancing schemes (IMEX): the ionic variables are advanced to the current time step, and inserted in the nonlinear term, while the latter one is then linearized around the value of the membrane potential at the previous time step. The degenerate parabolic nature of the Bidomain system, however, entails a very ill conditioning for the linear system associated to its discretization. From the mathematical and numerical standpoint, many efforts have been devoted in the recent years to set up efficient solvers and preconditioners to reduce the high computational costs associated to its numerical solution [4, 15, 44, 47, 62, 63]. Many proposed preconditioning strategies have been based on multigrid approaches [45, 54, 64] or suitable approximations of the equations [22, 24]. Among these works, most are based on a proper decomposition of the computational domain in order to set up parallel preconditioners, or on suitable multigrid schemes still coupled with parallel architectures [42, 60]. In particular, a Classical Schwarz Method coupled with a multigrid approach has been proposed in [43], while an Optimized Schwarz method has been introduced in [23]. The stiffness of the problem, due to the presence of a steep propagation front, led to the introduction of adaptive schemes, in both time [46], and time and space [7]. Another approach has been aiming at a simplification of the original problem, by using a somehow optimized Monodomain model [39], and by developing model adaptive techniques, where the costly Bidomain model is replaced by the Monodomain one (or an extended version of it) far from the depolarization front and the recovery tail of the action potential [23, 25, 26]. In the rest of the paper we provide a general survey of the mathematical and numerical aspects of the cardiac electrophysiology modeling. Section 2 is devoted to the modeling of an excitable myocardial cell, and some ionic models are presented: a phenomenological one, a model for atrial myocytes and a model for ventricular ones. In Sect. 3 we present a derivation of the macroscopic Bidomain model for propagation and its simplified

version, the Monodomain. Section 4 provides an introduction to the numerical approximation of both Bidomain and Monodomain models.

2 Mathematical Modeling of an Excitable Myocardial Cell

The basic property of neural cells to produce signals is called Action Potential (AP). It consists of a sudden variation in the transmembrane potential, called upstroke, followed by a recovering of the resting condition. It shows different shapes and amplitudes according to the different kind of excitable media to which the cells belong to, and in the large muscle cells makes it possible the simultaneous contraction of the whole cell. An action potential propagates keeping the same shape and amplitude all along an entire neural or muscular fiber. The Action Potential propagates across the heart in an heterogeneous way. The pulse moves from the Sinoatrial Nodus (SA), and propagates through the ordinary myocardic fibers of the right atrium, while the Buchmann's bundle drives the pulse towards the left atrium. Some action potentials propagate downwards and reach the Atrioventricular Nodus (AV), which is, under normal conditions, the only gate for the pulse to propagate from atria to ventricles, where the conduction is quicker (4 ms^{-1} versus 1 ms^{-1}).

Cardiac cells are characterized by a transmembrane potential that is negative at rest, owing to the fact that the concentration of potassium ions $[\text{K}^+]_i$ inside the cardiac cell is remarkably higher than the outside concentration $[\text{K}^+]_e$, and show two kinds of action potentials: the quick and the slow response.

The quick response is typical in the myocardium fibers (both atrial and ventricular) and in the Purkinje fibers, which are fibers specialized in the conduction. The quick response cells are characterized by a negative transmembrane potential at rest (around -90 mV), and by a rapid depolarization (positive overshoot), where the potential difference changes sign and the internal potential overtakes the external one of around 20 mV : such phase is called Phase 0. Immediately after that (Phase 1) a short period of partial repolarization takes place, followed by a plateau (Phase 2) which lasts for around 0.2 s . The potential gets progressively more negative (Phase 3) until it reaches again the resting value. The repolarization procedure is far slower than the depolarization one, and the interval between the end of the repolarization and the next action potential is called Phase 4.

The slow response is the one taking place in the Sinoatrial Nodus (SA), the natural pacemaker of the heart, and in the Atrioventricular Nodus (AV), the tissue meant to transfer the pulse from atria to ventricles. The slow response cells are characterized by a resting potential less negative (around -50 mV), by a smaller slope and amplitude in the overshoot of the action potential, by the absence of the Phase 1, and by a relative refractory period that continues during Phase 4.

From the modeling standpoint, the electrical activity in myocytes is characterised by transmembrane ionic currents and voltage changes, whose temporal dynamics are governed by the presence of various players at the molecular level (ion channels, pumps, concentrations), as well as many different proteins (such as transporters)

that are spatially organized at the cellular scale to generate action potentials (AP). The cell membrane is modeled as a capacitor separating the intra- and extra-cellular media, two ionic solutions. In the framework of Hodgkin-Huxley (HH) formalism [28], state variables are associated with the membrane potential, ionic concentrations, and molecular actors such as gating variables, which handle opening and closing of ionic channels. The system dynamics is thus described by a set of differential equations which depend on time, voltage, ion concentrations and the gating variables. Some recent models include additional differential equations to describe calcium regulation within the cell and possibly mitochondrial activity or force generation. Ionic models consist generally of 10–50 ODEs [57], but if molecular actors are modelled by Markov processes, such number can grow up to 100 ODEs, [29]. These systems are highly nonlinear and extremely stiff because of the large range of time-scales necessary to represent the various phenomena involved (from 100 ms to 1 μ s).

The earliest model for AP appeared in the work on nerve action potential by Hodgkin and Huxley [28], which earned them the Nobel prize in Medicine in 1963. Models of this type have successively been developed for the cardiac action potential, where the variation in time of the membrane potential u (under the assumption of an equipotential cell) is given by

$$C_m \frac{du}{dt} = -I_{ion}(u, w) + I_{st}, \quad (1)$$

where I_{ion} and I_{st} are the total ionic current and stimulus current across the membrane, respectively, and C_m is the total membrane capacitance. In (1) the ionic current through the channels in the membrane depends on the transmembrane potential u and on M gating and concentration variables $w \in \mathbf{R}^M$:

$$I_{ion}(u, w) = \sum_{k=1}^N G_k(u) \prod_{j=1}^M w_j^{p_{jk}} (u - E_k(w)),$$

$G_k(u)$ being the membrane conductance, E_k being the reversal potential for the k th current and p_{jk} being integers, and where the dynamics of the gating and concentration variables is described by a system of ODE's

$$\frac{dw}{dt} = R(u, w), \quad w(\mathbf{x}, 0) = w_0(\mathbf{x}). \quad (2)$$

In such models, if w_j is a gating variable, the right hand side $R_j(u, w)$ has a special structure and the corresponding ODE is given by

$$\frac{dw_j}{dt} = R_j(u, w) = R_j(u, w_j) = \alpha_j(u)(1 - w_j) - \beta_j(u)w_j, \quad (3)$$

with $\alpha_j(u), \beta_j(u) > 0$, $0 < w_j < 1$. Within this formalism, the temporal variation of a gating variable y equals the difference between the opening rate of closed gates and the closing rate of open gates, with rates that are voltage-dependent. If we introduce the steady state of the gating variable with the cell at rest, y^∞ , and the time constant associated with the gating variable τ_y , defined as

$$y^\infty = \alpha_y(u)\tau_y(u) \quad \tau_y(u) = \frac{1}{\alpha_y(u) + \beta_y(u)},$$

we observe that the generic gating variable y satisfies the equivalent ordinary differential equation

$$\frac{dy}{dt} = \frac{y^\infty - y}{\tau_y}. \quad (4)$$

Concerning the modeling of ventricular cells, the fitting of improved experimental data with more complex models led to the development of many refinements of the original Hodgkin-Huxley model: among them, we recall the model by Beeler and Reuter (1977, with four ionic currents and seven gating and concentrations variables), and the phase-I Luo-Rudy (1991, with $N = 6$ and $M = 7$). In this direction, the most used model of mammalian ventricular cells is the phase-II Luo-Rudy (1994, [35]), which is based on measurements from guinea pig. Simpler models of reduced complexity have also been proposed, where only 1 or 2 gating variables are considered. In the remainder of this section we present, as an example three ionic models: a 2 variables, phenomenological, model, a detailed model for atrial cells, and a detailed model for ventricular cells.

2.1 *The FitzHugh-Nagumo Cell Model and Its Rogers-McCulloch Variant*

The simplest ionic model for an excitable cell is the phenomenological FitzHugh-Nagumo (FHN, [13]) model, consisting of one ionic current and one gating variable. The latter is a simplified version of the Hodgkin-Huxley model. Assuming the potential v to be zero at rest, the ionic current uses only one recovery variable:

$$I_{ion}(u, w) = u - \frac{u^3}{3} - w + I,$$

where I is a stimulus current, and w satisfies

$$\frac{\partial w}{\partial t} = u + a - b w,$$

with $a, b > 0$.

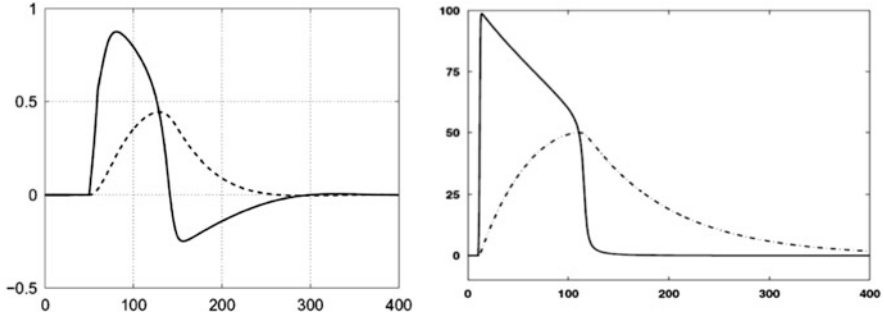


Fig. 1 Time evolution of the membrane potential u (solid line) and the recovery variable w (dashed line) in the FHN model (left) and its Rogers-McCulloch variant (right)

If the stimulus current I exceeds a given threshold, the system exhibits a spike and recovery dynamics. In Fig. 1 (left) we plot the temporal evolution of the membrane potential u (solid) and the recovery variable w (dashed).

An improvement of this model is given by the variant introduced in [49] by Rogers and McCulloch, where the ionic current and the recovery variable w dynamics are given by

$$I_{ion}(u, w) = Gu \left(1 - \frac{u}{u_{th}}\right) \left(1 - \frac{u}{u_p}\right) + \eta_1 uw,$$

and

$$\frac{\partial w}{\partial t} = \eta_2 \left(\frac{u}{u_p} - \eta_3 w\right),$$

where $G, \eta_1, \eta_2, \eta_3$ are positive coefficients, u_{th} is a threshold potential, and u_p is the peak potential. If the membrane potential does not exceed the threshold u_{th} , the AP is not triggered and the system gets back to the resting state. In Fig. 1 (right) we report the time evolution of the potential u (solid line) and of the gating variable w (dashed line) for the Rogers-McCulloch variant of the FHN model.

The recovery variable w ensures that, once an AP is triggered, the system cannot be excited again, unless a refractory period has passed. When a stimulus is introduced, the response of the system depends on the elapsed time since its spiking: if the recovery variable is small enough (or, equivalently, enough time has elapsed) another AP is created, with the same shape and amplitude of the first one. If the elapsed time does not outlast the refractory period, the generated AP can be shorter in duration, and smaller in amplitude, or just not being triggered, in the case too little time since spiking has elapsed.

Other all-or-nothing response models have been introduced in the literature, among which we recall the one proposed by Panfilov, Ten-Tusscher, and col-

laborators in [57]. The great simplicity of such model, and yet its ability in capturing significant aspects of the electrocardiac dynamics, is behind its wide use in literature. However, if on the one hand, such model is well suited to describe the positive overshoot in the quick depolarization phase, on the other hand it provides only a coarse approximation in the plateau and repolarization phases of the action potential, and behaves too poorly when accuracy in the description of the action potential is needed.

2.2 The Courtemanche, Ramirez and Nattel Atrial Cell Model

One of the most accurate models for atrial cells is the one proposed by Courtemanche, Ramirez and Nattel, (CRN, [11]). The total ionic current for the CRN model is given by the sum

$$I_{\text{ion}} = I_{\text{Na}} + I_{\text{K}} + I_{\text{Ca}} + I_{\text{b}} + I_{\text{p}}. \quad (5)$$

The above expression takes into account several aspects of the action potential generation. In (5), I_{Na} is the fast depolarizing Na^+ current, while the quantity I_{K} is the total rectifier K^+ current, given by

$$I_{\text{K}} = I_{\text{K1}} + I_{\text{to}} + I_{\text{Kur}} + I_{\text{Kr}} + I_{\text{Ks}},$$

where I_{K1} is the inward rectifier K^+ current, playing a major role in the late repolarization phase of the AP and in determining resting membrane potential and resistance, I_{to} is the transient outward K^+ current, I_{Kur} , I_{Kr} , and I_{Ks} are the ultrarapid, rapid, and slow rectifier currents. The quantity $I_{\text{Ca}} = I_{\text{Ca,L}}$ is the L-type Ca^{2+} current, while I_{b} is the background current for sodium Na^+ and calcium Ca^{2+}

$$I_{\text{b}} = I_{\text{b,Na}} + I_{\text{b,Ca}}.$$

Finally, I_{p} collects the actions of pumps and ion exchangers, designed to put back into balance the ion concentrations at rest:

$$I_{\text{p}} = I_{\text{NaCa}} + I_{\text{NaK}} + I_{\text{p,Ca}},$$

where I_{NaCa} is the sodium-calcium pump, I_{NaK} is the sodium-potassium pump, and $I_{\text{p,Ca}}$ is the calcium exchanger.

The model handles the intracellular concentrations $[\text{Na}^+]_i$, $[\text{K}^+]_i$, and $[\text{Ca}^{2+}]_i$. The intracellular calcium buffering by the sarcoplasmic reticulum system (SR) is described by the calcium concentrations in the uptake (NSR), and release (JSR) compartments of the sarcoplasmic reticulum, denoted by $[\text{Ca}^{2+}]_{\text{up}}$ and $[\text{Ca}^{2+}]_{\text{rel}}$ respectively.

In the model, no extracellular cleft space is included, the membrane capacitance is $c_m = 100\text{pF}$, the length and diameter of the cells are set to 100 and $16\text{ }\mu\text{m}$, respectively, and the cell compartment volumes are the same ones used in the phase-II Luo-Rudy model (LR2, [35]). Denoting by E_X the equilibrium potential for ion X , and with g_X its maximal conductance, from Nernst equation, E_X is given by

$$E_X = \frac{RT}{zF} \log \frac{[X]_e}{[X]_i},$$

where R is the gas constant, T is the absolute temperature, F is the Faraday constant, $z = 1$ for Na^+ and K^+ , $z = 2$ for Ca^{2+} , and $[X]_e$ and $[X]_i$ denote the external and internal concentration of ion X .

The dynamics of the concentration variables is governed by the following equations

$$\frac{d[\text{Na}^+]_i}{dt} = \frac{-3I_{\text{NaK}} - 3I_{\text{NaCa}} - I_{\text{b,Na}} - I_{\text{Na}}}{FV_i} \quad (6)$$

$$\frac{d[\text{K}^+]_i}{dt} = \frac{2I_{\text{NaK}} - I_{\text{K1}} - I_{\text{to}} - I_{\text{Kur}} - I_{\text{Kr}} - I_{\text{Ks}}}{FV_i} \quad (7)$$

$$\begin{aligned} \frac{d[\text{Ca}^{2+}]_i}{dt} = & \left[\frac{2I_{\text{NaCa}} - I_{\text{p,Ca}} - I_{\text{Ca,L}} - I_{\text{b,Ca}}}{2FV_i} + \frac{V_{\text{up}}(I_{\text{up,leak}} - I_{\text{up}}) + I_{\text{rel}}V_{\text{rel}}}{V_i} \right] \times \\ & \times \left[1 + \frac{\alpha_i\beta_i}{([\text{Ca}^{2+}]_i + \beta_i)^2} + \frac{\gamma_i\delta_i}{([\text{Ca}^{2+}]_i + \delta_i)^2} \right]^{-1} \end{aligned} \quad (8)$$

$$\frac{d[\text{Ca}^{2+}]_{\text{up}}}{dt} = I_{\text{up}} - I_{\text{up,leak}} - I_{\text{tr}} \frac{V_{\text{rel}}}{V_{\text{up}}} \quad (9)$$

$$\frac{d[\text{Ca}^{2+}]_{\text{rel}}}{dt} = (I_{\text{tr}} - I_{\text{rel}}) \left[1 + \frac{\alpha_{\text{rel}}\beta_{\text{rel}}}{([\text{Ca}^{2+}]_{\text{rel}} + \beta_{\text{rel}})^2} \right]^{-1}, \quad (10)$$

where V_i is the intracellular volume, V_{up} and V_{rel} are the volumes of the uptake (NSR) and release (JSR) compartments of the sarcoplasmic reticulum, α_i , γ_i , and α_{rel} depend on the total concentrations of troponin and calmodulin in myoplasm, and of calsequestrin in JSR, while β_i , δ_i , and β_{rel} depend on their half saturation constants, respectively. All these three proteins are responsible of the contraction of the cell.

In (8) and (9), $I_{\text{up,leak}}$ is the Ca^{2+} leak current by the JSR, I_{up} is the Ca^{2+} uptake current by the JSR, while I_{rel} is the Ca^{2+} release current from the JSR. Finally, in (9) and (10), I_{tr} is the transfer current from NSR to JSR.

The model consists globally of five concentration variables and 15 gating variables. In Figs. 2 and 3 we plot the time evolution of the potential and of the gating and concentration variables. For a more detailed description of the model we

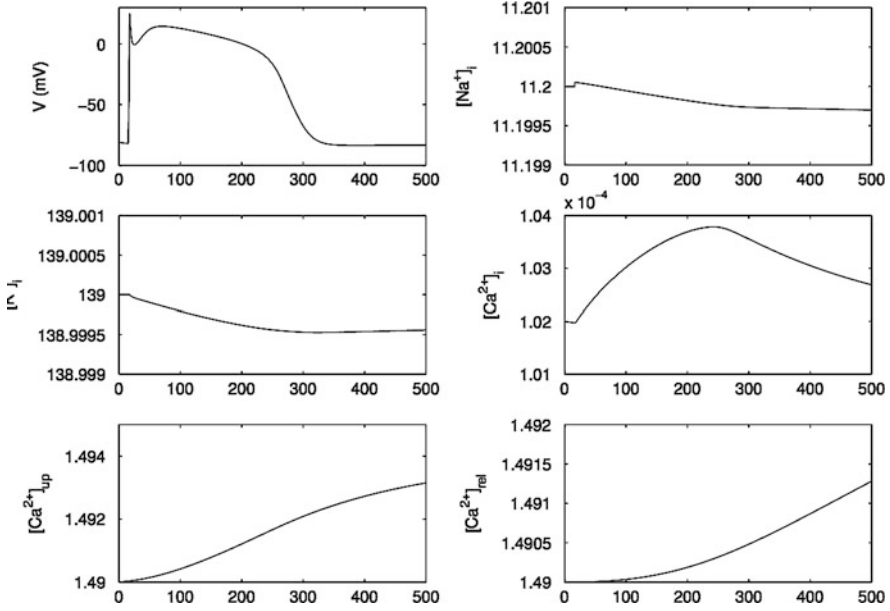


Fig. 2 CRN model: time evolution of the membrane potential and the concentration variables

refer the interested reader to the original paper by Courtemanche et al. [11]. Among other popular human atrial models, we recall the ones proposed by Earm and Noble in [12], and the one proposed by Nygren and his collaborators in [40].

2.3 The Luo-Rudy 1 Ventricular Cell Model

The Luo-Rudy Phase 1 model is among the most popular ionic model used in literature to model ventricular myocytes. It consists of six ionic currents, seven gating variables, and one concentration variable for the intracellular calcium, whose dynamics plays a pivotal role in the heart contraction. The total current is given by

$$I_{ion} = I_{Na} + I_{si} + I_K + I_{K1} + I_{Kp} + I_b, \quad (11)$$

where the ionic currents are given by

$$\begin{aligned} I_{Na} &= g_{Na} m^3 h j(u - E_{Na}) & I_{si} &= g_{si} d f(u - E_{si}) \\ I_K &= g_K ([K]_e) X(u - E_K) & I_{K1} &= f_{K1}([K]_e, u) (u - E_{K1}) \\ I_{Kp} &= f_{Kp}(u) K_p(u - E_{Kp}) & I_b &= b_1(u + b_2). \end{aligned} \quad (12)$$

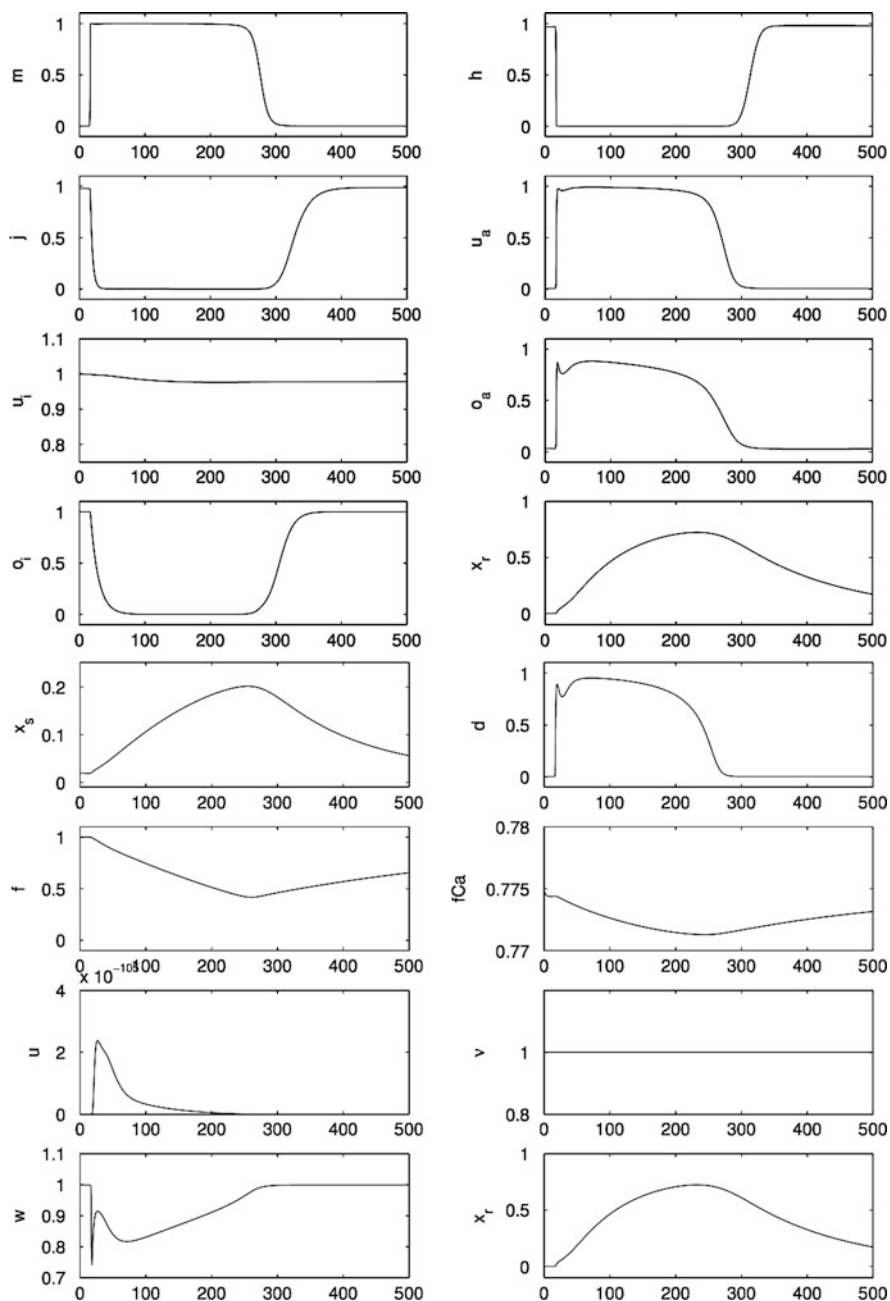


Fig. 3 CRN model: time evolution of the gating variables

The calcium concentration satisfies the differential equation

$$\frac{d[Ca]_i}{dt} = -c_3 I_{si} + c_2 (c_3 - [Ca]_i), \quad (13)$$

while the gating variables are described, within the HH formalism, as

$$\frac{dw}{dt} = \alpha_w(1 - w) - \beta_w w, \quad \text{with } w \in \{m, h, j, d, f, X\}. \quad (14)$$

In (12)–(14), I_{K_p} is the plateau current, I_b is the background current, f_{K_1} and f_{K_p} are rational exponentials of the membrane potential, $g_K([K]_e)$ is a function of the extracellular potassium concentration $[K]_e$, E_{si} is linearly dependent on the natural logarithm of the intracellular calcium concentration $[Ca]_i$, while g_{Na} , g_{si} , b_1 , b_2 , c_2 , and c_3 are constants determined by fitting with experimental data. In the LR1 model, α_h , β_h , α_j , β_j and X depend on the membrane potential u through functions that show different behavior with respect to a threshold. In Fig. 4 we plot the temporal

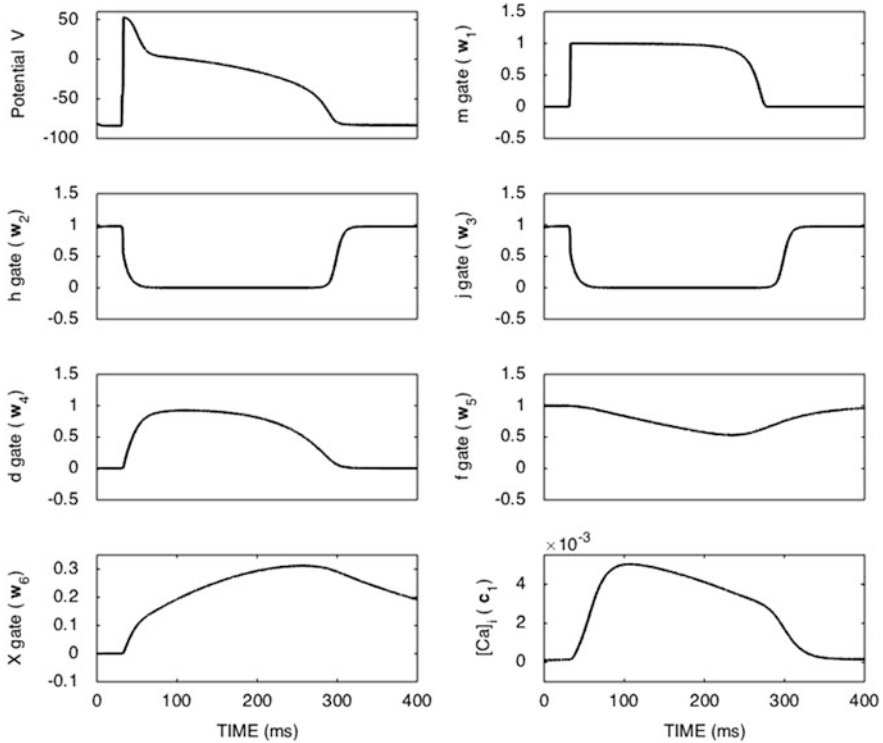


Fig. 4 Time evolution of the membrane potential, the gating variables and the $[Ca]_i$ concentration in the LR1 model

dynamics of the membrane potential, the calcium concentration, and the gating variables for the LR1 model.

The LR1 model was extended, in 1994, to the Luo-Rudy phase 2 model, that includes six ionic currents, and five ionic concentrations, which allow (as in the CRN model for atrial cells) the calcium handling by the sarcoplasmic reticulum in the interior of the cell. The LR2 model was further extended by Winslow in 1999, by including experimental data from a canine heart and a more detailed modeling of the calcium dynamics, to a model featuring 25 ionic currents and six ionic concentrations.

3 The Macroscopic Bidomain Model for Electrophysiological Propagation

The Bidomain model is commonly considered one of the most complete and accurate models to describe the propagation of the electrical potential in the myocardium tissue (see e.g. [27, 50, 52]). Such model has been derived, by an homogenization technique, starting from a periodic assembling of elongated cells surrounded by extracellular space and connected by end-to-end or side-to-side junctions (for the mathematical details we refer to [16, 32]). The mathematical problem is naturally set in a bounded region $\Omega \subset \mathbf{R}^3$, which represents a portion of the heart tissue.

3.1 Tissue and Conductivities Modeling

The Bidomain model relies on representing the cardiac tissue as the superposition of two media which are both continuous and anisotropic. The intra-cellular and the extra-cellular media coexist at each point $\mathbf{x} \in \Omega$ and are separated by a cell membrane. In a natural manner, the intracellular and extracellular potential are denoted by u_i and u_e , respectively, while their difference $u = u_i - u_e$ expresses the membrane potential.

The conductivity of the cardiac cells depends upon their orientation, featuring preferential pathway along *gap junctions* (see Fig. 5, left), and in the most general case the conductivity tensor is anisotropic. The structure of the cardiac cells can be modeled, following Le Grice et al. [33] as a sequence of muscular layers going from endocardium to epicardium (see also [56]). Anatomical studies show that the fibers direction rotates counterclockwise from epicardium to endocardium and that they are arranged in sheets, running across the myocardial wall [5, 52]. In any point $\mathbf{x}$ it is then possible to identify an orthonormal triplet of directions, $\mathbf{a}_f(\mathbf{x})$ along the fiber, $\mathbf{a}_t(\mathbf{x})$ orthogonal to the fiber direction and in the fiber sheet and $\mathbf{a}_n(\mathbf{x})$ orthogonal to the sheet (see Fig. 5, right, for a schematic representation). The intra

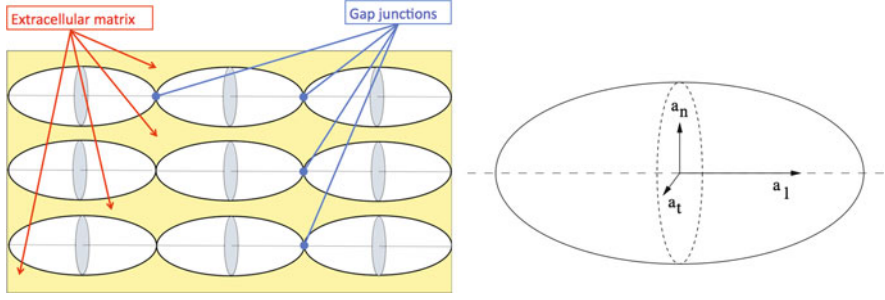


Fig. 5 *Left*: schematic representation of the fibers structure. *Right*: the cell reference frame

and extracellular media present different conductivity values in each direction. At point $\mathbf{x} \in \Omega$, we denote by $\sigma_\tau^l(\mathbf{x})$, $\sigma_\tau^n(\mathbf{x})$, and $\sigma_\tau^t(\mathbf{x})$ with $\tau = i, e$ the intra and extracellular conductivities in the $\mathbf{a}_l(\mathbf{x})$, $\mathbf{a}_t(\mathbf{x})$ and $\mathbf{a}_n(\mathbf{x})$ direction, respectively.

The intra and extracellular local anisotropic conductivity tensors read, for $\tau = i, e$,

$$\mathbf{D}_\tau(\mathbf{x}) = \sigma_\tau^l(\mathbf{x})\mathbf{a}_l(\mathbf{x})\mathbf{a}_l^T(\mathbf{x}) + \sigma_\tau^t(\mathbf{x})\mathbf{a}_t(\mathbf{x})\mathbf{a}_t^T(\mathbf{x}) + \sigma_\tau^n(\mathbf{x})\mathbf{a}_n(\mathbf{x})\mathbf{a}_n^T(\mathbf{x}). \quad (15)$$

We assume that $\mathbf{D}_\tau$ fulfills in Ω a uniform elliptic condition.

A common practical hypothesis, followed by many authors (see, e.g. [18]) is axial isotropy: the myocardium tissue is assumed to feature the same conductivity in both the tangential and normal direction ($\sigma_{i,e}^t = \sigma_{i,e}^n$). Under this hypothesis, the conductivity tensors, for $\tau = i, e$, simplify in

$$\mathbf{D}_\tau(\mathbf{x}) = \sigma_\tau^l \mathbf{I} + (\sigma_\tau^t - \sigma_\tau^l)\mathbf{a}_l(\mathbf{x})\mathbf{a}_l^T(\mathbf{x}). \quad (16)$$

3.2 Quasi-Static Electromagnetic Field

The propagation of an action potential across the myocardium generates, at the macroscopic level, an electrical signal that can be measured, and whose temporal dynamics is described by Maxwell's equations.

In a conducting body, the electrical current density is governed by Ohm's Law,

$$\mathbf{J} = \mathbf{D}\mathbf{E}, \quad (17)$$

where $\mathbf{D}$ is the conductivity of the medium, and $\mathbf{E}$ is the electrical field.

Faraday's law relates the time derivative of the magnetic field with the rotational of the electrical field:

$$\frac{\partial \mathbf{B}}{\partial t} + \nabla \times \mathbf{E} = 0.$$

Given the temporal scale of the AP and the spatial scale of the heart, variations of the magnetic field can be neglected, leading to the *quasi-static* assumption

$$\frac{\partial \mathbf{B}}{\partial t} \sim 0 \quad \implies \quad \nabla \times \mathbf{E} = 0. \quad (18)$$

Since the electric field $\mathbf{E}$ can be assumed irrotational, there exists some potential u such that $\mathbf{E} = -\nabla u$, and the current in the conducting medium can then be expressed in terms of such potential, as

$$\mathbf{J} = \mathbf{D} \nabla u.$$

3.3 The Bidomain Model

The assumption of coexistence of intra- and extracellular media entails that in each point of the domain $\mathbf{x} \in \Omega$, two currents exist and are given by

$$\mathbf{J}_i = \mathbf{D}_i \nabla u_i \quad \text{and} \quad \mathbf{J}_e = \mathbf{D}_e \nabla u_e, \quad (19)$$

where $\mathbf{D}_i$ and $\mathbf{D}_e$ represent the conductivities of the intra- and extra-cellular medium, respectively. For any given small volume V , the charge conservation principle entails that the total current entering the volume must equal the total current leaving it. Within our framework, this principle amounts to balance the current flowing between the intracellular and extracellular space, as

$$\int_{\partial V} \mathbf{n} \cdot (\mathbf{J}_i + \mathbf{J}_e) \, ds = 0. \quad (20)$$

By a straightforward application of the divergence theorem, we have, for any given small volume V , that

$$\nabla \cdot (\mathbf{J}_i + \mathbf{J}_e) = 0, \quad (21)$$

and, from (19), the current balance is given by

$$\nabla \cdot (\mathbf{D}_i \nabla u_i) + \nabla \cdot (\mathbf{D}_e \nabla u_e) = 0. \quad (22)$$

The current flowing from one domain to the other must equal the cell membrane current, that is given in (1). We thus obtain

$$\nabla \cdot (\mathbf{D}_i \nabla u_i) = -\nabla \cdot (\mathbf{D}_e \nabla u_e) = \chi \left(C_m \frac{\partial u}{\partial t} + I_{ion} \right) \quad (23)$$

where χ is the surface to volume ratio of the cell.

Depending on the way the three terms in system (23) are grouped, two different formulation of the Bidomain model emerge.

3.3.1 Parabolic-Parabolic Formulation of the Bidomain Model

By equaling the first and second terms to the third one in (23), the Bidomain model results in a system of two nonlinear parabolic reaction-diffusion equations:

$$\begin{aligned} \chi C_m \frac{\partial u}{\partial t} - \nabla \cdot (\mathbf{D}_i \nabla u_i) + \chi I_{ion} &= 0 \\ -\chi C_m \frac{\partial u}{\partial t} - \nabla \cdot (\mathbf{D}_e \nabla u_e) - \chi I_{ion} &= 0. \end{aligned} \quad (24)$$

This formulation is known in literature as *Parabolic-Parabolic (PP)*. The problem is completed by suitable initial conditions, and by homogeneous Neumann boundary conditions on $\partial\Omega$, modeling an insulated myocardium.

The complete PP formulation of the Bidomain model reads:

$$\left\{ \begin{array}{ll} \chi C_m \frac{\partial u}{\partial t} - \nabla \cdot (\mathbf{D}_i \nabla u_i) + \chi I_{ion} = I_i^{app} & \text{in } \Omega \times (0, T) \\ -\chi C_m \frac{\partial u}{\partial t} - \nabla \cdot (\mathbf{D}_e \nabla u_e) - \chi I_{ion} = I_e^{app} & \text{in } \Omega \times (0, T) \\ u = u_i - u_e & \text{in } \Omega \times (0, T) \\ \frac{\partial w}{\partial t} = R(u, w) & \text{in } \Omega \times (0, T) \\ \mathbf{n}^T \mathbf{D}_i \nabla u_i = 0, \quad \mathbf{n}^T \mathbf{D}_e \nabla u_e = 0 & \text{on } \partial\Omega \times (0, T) \\ u_i(\mathbf{x}, 0) = u_{i,0}, \quad u_e(\mathbf{x}, 0) = u_{e,0}, \quad w(\mathbf{x}, 0) = w_0 & \text{in } \Omega. \end{array} \right. \quad (25)$$

In (25), $\mathbf{n}$ is the unit normal outward-pointing vector on the surface. As a consequence of the Gauss theorem, the applied external stimuli must fulfill the compatibility condition

$$\int_{\Omega} I_i^{app} d\mathbf{x} = \int_{\Omega} I_e^{app} d\mathbf{x}. \quad (26)$$

System (24) consists of two parabolic reaction diffusion equations for u_i and u_e where the vector of time derivatives is multiplied by a singular matrix. The system is thus said to be *degenerate*. The transmembrane potential u is uniquely determined, while the intra and extracellular potentials u_i and u_e are determined up to the same function of time, whose value is usually obtained by imposing that the extracellular potential u_e has zero mean on Ω ($\int_{\Omega} u_e dx = 0$).

The PP formulation has been commonly used by several scientists. In particular, this formulation has been particularly popular in theoretical studies for well-posedness analysis of the problem. Little is known on degenerate reaction-diffusion systems such as the Bidomain model. We refer the reader to [16] for existence, uniqueness and regularity results, both at the continuous and the semi-discrete level, and to [53] for a convergence analysis of finite elements approximations. Both papers deal with the FitzHugh-Nagumo (FHN) model of the gating system. For well-posedness analysis of the Bidomain problem associated with different ionic models see [2, 3], and [61].

More results are known on the related eikonal approximation describing the propagation of excitation front (see for instance [14, 17, 30]), and a mathematical analysis of the Bidomain model using Γ -convergence theory can be found in [1].

3.3.2 Parabolic-Elliptic Formulation of the Bidomain Model

By equaling the first term to both the second and the third one in (23), and using the fact that $u_i = u + u_e$, the Bidomain model results in a system of one nonlinear parabolic reaction-diffusion equation, and a linear elliptic equation:

$$\chi C_m \frac{\partial u}{\partial t} - \nabla \cdot (\mathbf{D}_i \nabla u_i) - \nabla \cdot (\mathbf{D}_e \nabla u_e) + \chi I_{ion} = 0 \quad (27)$$

$$\nabla \cdot (\mathbf{D}_i \nabla u) + \nabla \cdot ((\mathbf{D}_i + \mathbf{D}_e) \nabla u_e) = 0.$$

This formulation is known in literature as *Parabolic-Elliptic (PE)*. As for the PP formulation, the problem is completed by suitable initial and boundary conditions, and the complete PE formulation of the Bidomain model reads:

$$\left\{ \begin{array}{ll} \chi C_m \frac{\partial u}{\partial t} - \nabla \cdot (\mathbf{D}_i \nabla u_i) + \chi I_{ion} = I_i^{app} & \text{in } \Omega \times (0, T) \\ \nabla \cdot (\mathbf{D}_i \nabla u) + \nabla \cdot ((\mathbf{D}_i + \mathbf{D}_e) \nabla u_e) = 0 & \text{in } \Omega \\ \frac{\partial w}{\partial t} = R(u, w) & \text{in } \Omega \times (0, T) \\ \mathbf{n}^T \mathbf{D}_i \nabla u_i = 0, \quad \mathbf{n}^T \mathbf{D}_e \nabla u_e = 0 & \text{on } \partial\Omega \times (0, T) \\ u_i(\mathbf{x}, 0) = u_{i,0}, \quad u_e(\mathbf{x}, 0) = u_{e,0}, \quad w(\mathbf{x}, 0) = w_0 & \text{in } \Omega. \end{array} \right. \quad (28)$$

Differently from the PP formulation, system (27) does not consist of two parabolic equations for u_i and u_e where the vector of time derivatives is multiplied by a singular matrix. Nevertheless, also system (27) is *degenerate*, since the elliptic equation in (27) is in practice a Laplacian with homogeneous Neumann boundary conditions, whose solution is known only up to a constant. Also in this formulation, thus, the transmembrane potential u is uniquely determined, while the intra and extracellular potentials u_i and u_e are determined up to the same function of time, whose value is again obtained by imposing zero mean to the extracellular potential on Ω ($\int_{\Omega} u_e dx = 0$).

By letting $\lambda_m = \min \{ \sigma_e^l / \sigma_i^l, \sigma_e^t / \sigma_i^t \}$ and $\lambda_M = \max \{ \sigma_e^l / \sigma_i^l, \sigma_e^t / \sigma_i^t \}$, an alternative PE formulation can be obtained by linear combinations of the equations in (24), with coefficients $(\frac{\lambda}{1+\lambda}, -\frac{1}{1+\lambda})$, $\lambda_m \leq \lambda \leq \lambda_M$, and $(1, 1)$:

$$\begin{aligned} \chi C_m \frac{\partial u}{\partial t} - \nabla \cdot \left[\frac{\lambda \mathbf{D}_i}{1+\lambda} \nabla u \right] - \nabla \cdot \left[\frac{\lambda \mathbf{D}_i - \mathbf{D}_e}{1+\lambda} \nabla u_e \right] + \chi I_{ion}(u) &= I^{app} \\ -\nabla \cdot [\mathbf{D}_i \nabla u + (\mathbf{D}_i + \mathbf{D}_e) \nabla u_e] &= \tilde{I}^{app}, \end{aligned} \quad (29)$$

where we have set $I^{app} = \frac{\lambda I_i^{app} + I_e^{app}}{1+\lambda}$ and $\tilde{I}^{app} = I_i^{app} - I_e^{app}$.

The (PE) formulation of the Bidomain problem has been widely used for numerical simulations, in particular among the Bioengineering community, serving as the basis for the development of efficient preconditioners.

3.4 A Simplified Model: The Monodomain

If we assume the anisotropy ratio to be the same in the two media, the Bidomain model reduces to a simpler one, called Monodomain. Its derivation can be obtained in different ways, the common underlying hypothesis being a proportionality assumption between the intracellular and the extracellular conductivity tensors, namely $\mathbf{D}_e = \lambda \mathbf{D}_i$, where λ is a constant to be properly chosen. For instance, under assumption (16), λ can be devised through a minimization procedure, as

$$\lambda = \operatorname{argmin} J(\lambda), \quad J(\lambda) = \left(\frac{\sigma_e^l - \lambda \sigma_i^l}{1 + \lambda} \right)^2 + 2 \left(\frac{\sigma_e^t - \lambda \sigma_i^t}{1 + \lambda} \right)^2$$

for given values of the conductivities. A time dependent choice of the parameter λ has been proposed in [39].

After defining $\mathbf{D} := \mathbf{D}_i + \mathbf{D}_e$ and $\mathbf{D}_M := \mathbf{D}_e \mathbf{D}^{-1} \mathbf{D}_i$, the first equation in (29) can be rearranged as

$$\chi C_m \frac{\partial u}{\partial t} - \nabla \cdot \mathbf{D}_M \nabla u + \nabla \cdot \left[\left(\mathbf{D}_e \mathbf{D}^{-1} - \frac{\lambda}{1 + \lambda} \mathbf{I} \right) (\mathbf{D}_i \nabla u + \mathbf{D} \nabla u_e) \right] + \chi I_{ion}(u, \mathbf{w}) = I^{app} \quad (30)$$

and, since the proportionality assumption $\mathbf{D}_e = \lambda \mathbf{D}_i$ entails $\mathbf{D}_e \mathbf{D}^{-1} - \frac{\lambda}{1 + \lambda} \mathbf{I} = \mathbf{0}$, a formulation of the Monodomain model (see [6, 31]) is then obtained from (30) as

$$\chi C_m \frac{\partial u}{\partial t} - \nabla \cdot \mathbf{D}_M \nabla u + \chi I_{ion}(u, \mathbf{w}) = I^{app}. \quad (31)$$

The problem is completed by suitable initial conditions, and by homogeneous Neumann boundary conditions on $\partial\Omega$, and reads:

$$\begin{cases} c_m \frac{\partial u}{\partial t} - \operatorname{div}(\mathbf{D}_M \nabla u) + I_{ion}(u, w) = I^{app} & \text{in } \Omega \times (0, T) \\ \frac{\partial w}{\partial t} - R(u, w) = 0 & \text{in } \Omega \times (0, T) \\ \mathbf{n}^T D \nabla u = 0 & \text{in } \partial\Omega \times (0, T) \\ u(\mathbf{x}, 0) = u_0(\mathbf{x}), \quad w(\mathbf{x}, 0) = w_0(\mathbf{x}), & \text{in } \Omega. \end{cases} \quad (32)$$

Such model consists of a single parabolic reaction-diffusion equation for the transmembrane potential u coupled with an ODE system for the gating and concentration variables. Differently from the Bidomain, several theoretical results on reaction-diffusion equations can be applied to the Monodomain model, which features a unique solution (u, w) , resulting in a much easier to solve problem after numerical discretization. In many applications the Monodomain model is

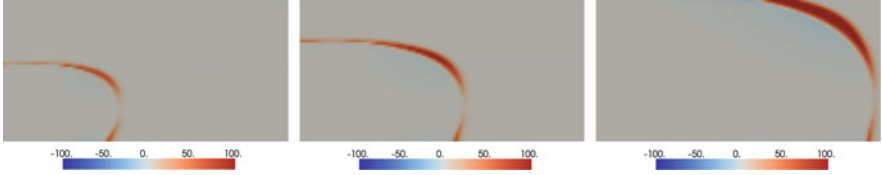


Fig. 6 Difference in the propagation of the membrane potential between Bidomain and Monodomain simulation with fibers oriented along the x axis. (Reproduced with permission from [25])

accurate enough to capture the desired dynamics and effects of the action potential propagation. In [47], Potse and his collaborators compare the action potential propagation velocities using Monodomain and Bidomain, observing that the Monodomain solution propagates a bit slower (2 %) than the Bidomain one, and conclude that “in absence of applied currents, propagating of action potentials on the scale of a human heart can be studied with a Monodomain model”. However, the Bidomain model becomes necessary when current stimuli are applied in the extracellular space. As a consequence, the Monodomain has been long considered inadequate to simulate defibrillation [59]. In recent works by Y. Coudiere and his collaborators, [9, 10], a proper generalisation of the boundary conditions has been introduced, that allows an external stimuli to be applied directly at the Monodomain model.

In Fig. 6 we plot the difference in the propagation between Bidomain and Monodomain on a slab with the principal fibers axis oriented horizontally from left to right. The difference in the propagation speed between the two models can be clearly appreciated. In addition, it is evident that the Monodomain error, although globally significant, is minimal in both directions along and across the fibers. An accurate knowledge of the fibers arrangements would strongly reduce the error when the Monodomain model is used for patient-specific simulation. Still, such arrangement, although having common features, is highly individual, and, unfortunately, a definitive knowledge of the fibers distribution is not available yet. Modern advanced medical imaging techniques, such as DTI (*Diffusion Tensor Imaging*) allow an accurate mapping of the fibers direction, making the Monodomain a viable and cheaper alternative to heavy Bidomain simulations.

4 Numerical Approximation

We give here a brief introduction to the numerical approximation of the models presented in the previous Section. In what follows we do not rely on a specific choice for the ionic model describing the cell membrane currents. Thus, from now on, we will simply denote by $I_{ion}(u, w)$ the ionic current. For more detailed description, we refer the interested reader to [15, 19, 63].

4.1 Time Marching Scheme

For the sake of simplicity in presentation, we consider a fixed time step Δt , even if effective time adaptive scheme have been developed in the literature (see e.g. [7, 46]), and we denote with superscript n the variables computed at time $t^n = n\Delta t$. The Bidomain and Monodomain systems can be advanced in time by an implicit-explicit (IMEX) scheme: moving from time step t^n to t^{n+1} , the gating and concentration variables w are updated first, and used to compute the new values for the electric potentials.

Let the ionic variables vector $w \in \mathbf{R}^m$ be arranged as $w = (g, c)^T$, where $g \in \mathbf{R}^p$ represent the gating variables, while $c \in \mathbf{R}^q$ represent the concentration variables ($p + q = m$). Owing to the Hodgkin-Huxley formalism (3)–(4), the components of g are first integrated exactly in time on $(0, \Delta t)$, upon an appropriate linearisation around the membrane potential at the previous time step

$$g_j^{n+1} = g_{j\infty}(u^n) + (g_j^n - g_{j\infty}(u^n)) \exp\left(-\frac{\Delta t}{\tau_{g_j}(u^n)}\right),$$

then the concentration variables c are integrated by a backward Euler scheme, taking into account the updated values g^{n+1} ,

$$\frac{c^{n+1} - c^n}{\Delta t} = R_c(u^n, g^{n+1}, c^n),$$

where R_c are the rows in (2) associated with c . The time step is selected to guarantee stability to the time advancing scheme.

The electric potentials are then updated by solving on Ω a semi-implicit problem, where the linear diffusion term is discretized implicitly, while the nonlinear reaction term (the ionic current $I_{ion}(u, w)$) is treated explicitly with respect to the membrane potential u . While taking into account the updated values of the gating and concentration variables w^{n+1} , this allows to skip the computationally expensive solution of nonlinear systems.

Within this framework, the semi-discrete version of the Bidomain PP formulation (24), solves for $0 < n \leq N = T/\Delta t$,

$$\left\{ \begin{array}{l} \chi C_m \frac{u^{n+1} - u^n}{\Delta t} - \nabla \cdot (\mathbf{D}_i \nabla u_i^{n+1}) = I_i^{app} - \chi I_{ion}(u^n, w^{n+1}) \\ -\chi C_m \frac{u^{n+1} - u^n}{\Delta t} - \nabla \cdot (\mathbf{D}_e \nabla u_e^{n+1}) = I_e^{app} + \chi I_{ion}(u^n, w^{n+1}) \\ u_i^0(\mathbf{x}) = u_{i,0}(\mathbf{x}) \quad u_e^0(\mathbf{x}) = u_{e,0}(\mathbf{x}) \\ \mathbf{n}^T \mathbf{D}_i \nabla u_i^{n+1}|_{\partial\Omega} = 0 \quad \mathbf{n}^T \mathbf{D}_e \nabla u_e^{n+1}|_{\partial\Omega} = 0. \end{array} \right. \quad (33)$$

Similarly, the semi-discrete version of PE formulation, solves for (27)

$$\left\{ \begin{array}{l} \chi C_m \frac{u^{n+1} - u^n}{\Delta t} - \nabla \cdot (\mathbf{D}_i \nabla u^{n+1} + \mathbf{D}_e \nabla u_e^{n+1}) = I^{app} - \chi I_{ion}(u^n, w^{n+1}) \\ -\nabla \cdot [\mathbf{D}_i \nabla u^{n+1} + (\mathbf{D}_i + \mathbf{D}_e) \nabla u_e^{n+1}] = \tilde{I}^{app} \\ u^0(\mathbf{x}) = u_0(\mathbf{x}) \quad u_e^0(\mathbf{x}) = u_{e,0}(\mathbf{x}) \\ \mathbf{n}^T \mathbf{D}_i (\nabla u^{n+1} + \nabla u_e^{n+1})|_{\partial\Omega} = 0 \quad \mathbf{n}^T \mathbf{D}_e \nabla u_e^{n+1}|_{\partial\Omega} = 0 \end{array} \right. \quad (34)$$

and, for (29)

$$\left\{ \begin{array}{l} \chi C_m \frac{u^{n+1} - u^n}{\Delta t} - \nabla \cdot \left(\frac{\lambda \mathbf{D}_i}{1 + \lambda} \nabla u^{n+1} + \frac{\lambda \mathbf{D}_i - \mathbf{D}_e}{1 + \lambda} \nabla u_e^{n+1} \right) = I^{app} - \chi I_{ion}(u^n, w^{n+1}) \\ -\nabla \cdot [\mathbf{D}_i \nabla u^{n+1} + (\mathbf{D}_i + \mathbf{D}_e) \nabla u_e^{n+1}] = \tilde{I}^{app} \\ u^0(\mathbf{x}) = u_0(\mathbf{x}) \quad u_e^0(\mathbf{x}) = u_{e,0}(\mathbf{x}) \\ \mathbf{n}^T \mathbf{D}_i (\nabla u^{n+1} + \nabla u_e^{n+1})|_{\partial\Omega} = 0 \quad \mathbf{n}^T \mathbf{D}_e \nabla u_e^{n+1}|_{\partial\Omega} = 0. \end{array} \right. \quad (35)$$

In a similar manner, the semi-discrete formulation of the Monodomain model updates the membrane potential by solving, at each time step:

$$\left\{ \begin{array}{l} \chi C_m \frac{u^{n+1} - u^n}{\Delta t} - \nabla \cdot \mathbf{D}_M \nabla u^{n+1} = I^{app} - \chi I_{ion}(u^n, w^{n+1}) \\ u^0(\mathbf{x}) = u_0(\mathbf{x}) \\ \mathbf{n}^T \mathbf{D}_M \nabla u^{n+1}|_{\partial\Omega} = 0. \end{array} \right. \quad (36)$$

4.2 Spatial Discretization

The most common approach in the literature is to look for approximate solutions to the Bidomain and Monodomain models in a finite element space. The computational domain $\Omega \subset \mathbf{R}^3$ is discretized in space with a regular triangulation $\mathcal{T}_h$, namely $\Omega = \bigcup_{j=1}^N T_j$, where each $T_j \in \mathcal{T}_h$ is obtained through an invertible affine map from a reference element E , a simplex (namely the tetrahedron with vertices $(0, 0, 0)$,

$(1, 0, 0)$, $(0, 1, 0)$, and $(0, 0, 1)$) or the unit cube $[0, 1]^3$. The associated finite element spaces X_h and Y_h (see e.g. [48] for an introduction to finite element methods) are defined as

$$X_h = \left\{ \varphi_h \in C^0(\Omega) \mid \varphi_h|_{T_j} \in \mathbf{P}_k(T_j) \right\}, \quad Y_h = \left\{ \varphi_h \in C^0(\Omega) \mid \varphi_h|_{T_j} \in \mathbf{Q}_k(T_j) \right\},$$

where $\mathbf{P}_k(T_j)$ is the space of polynomials of degree at most k on T_j , whereas $\mathbf{Q}_k(T_j)$ is the space of polynomials of degree at most k with respect to each variable on T_j .

A fully discrete problem for (33)–(36) is then obtained by applying a Galerkin procedure on their variational formulations, using as finite dimensional space $V_h = X_h$ or $V_h = Y_h$, and choosing a basis for V_h . Let then $\Phi = \{\varphi_j\}_{j=1}^{N_h}$ be a basis for V_h . We denote by M the mass matrix, and by K_τ ($\tau = i, e, M$) the stiffness matrices with entries

$$M^{ij} = \sum_{p=1}^N \int_{T_p} \varphi_i \varphi_j dx, \quad K_\tau^{ij} = \sum_{p=1}^N \int_{T_p} (\nabla \varphi_j)^T D_\tau(\mathbf{x}) \nabla \varphi_i dx.$$

Numerical evaluation of such integrals is obtained by means of suitable quadrature rules.

4.3 Algebraic Formulation

The unknowns of the fully discrete problem are represented by vectors $\mathbf{u}$, $\mathbf{u}_i$, $\mathbf{u}_e$, and $\mathbf{w}$, storing the nodal values of u , u_i , u_e , and w , respectively.

Advancing the potentials from time t^n to t^{n+1} amount eventually to solve, at each step, a linear system. Since the Bidomain system is degenerate, the matrix associated to its discrete formulation is singular, with a one dimensional kernel spanned by $(\mathbf{1}, \mathbf{1})^T$, independently from its formulation. As a consequence, the transmembrane potential $\mathbf{u}^{n+1}$ is uniquely determined, as in the continuous model, while $\mathbf{u}_i^{n+1}$ and $\mathbf{u}_e^{n+1}$ are determined up to the same additive time-dependent constant with respect to a reference potential. Such constant can be determined by imposing the condition $\mathbf{1}^T M \mathbf{u}_e^{n+1} = 0$. On the other hand, the matrix associated with the full discrete version of the Monodomain system is positive definite, due to the uniform ellipticity assumption on the conductivity tensors. As a consequence, the transmembrane potential $\mathbf{u}^{n+1}$ is uniquely determined for the discrete Monodomain system.

4.3.1 Parabolic-Parabolic Bidomain Formulation

The full discretization of the PP Bidomain system (24) reads:

$$\begin{bmatrix} \frac{\chi C_m}{\Delta t} M + K_i & -\frac{\chi C_m}{\Delta t} M \\ -\frac{\chi C_m}{\Delta t} M & \frac{\chi C_m}{\Delta t} M + K_i \end{bmatrix} \begin{bmatrix} \mathbf{u}_i \\ \mathbf{u}_e \end{bmatrix}^{n+1} = \frac{\chi C_m}{\Delta t} \begin{bmatrix} M & -M \\ -M & M \end{bmatrix} \begin{bmatrix} \mathbf{u}_i \\ \mathbf{u}_e \end{bmatrix}^n + \begin{bmatrix} \chi M I_{ion}^h(\mathbf{u}^n, \mathbf{w}^{n+1}) + M I_i^{app} \\ -\chi M I_{ion}^h(\mathbf{u}^n, \mathbf{w}^{n+1}) + M I_e^{app} \end{bmatrix}.$$

The above linear system features a symmetric positive semidefinite matrix, Due to its symmetry, the system can be solved by a Preconditioned Conjugate Gradient algorithm (PCG, see e.g. [51]), using as initial guess the solution at the previous time step.

4.3.2 Parabolic-Elliptic Bidomain Formulation

The full discretization of the PE Bidomain system (28) reads:

$$\begin{bmatrix} \frac{\chi C_m}{\Delta t} M + K_i & K_i \\ K_i & K_i + K_e \end{bmatrix} \begin{bmatrix} \mathbf{u} \\ \mathbf{u}_e \end{bmatrix}^{n+1} = \begin{bmatrix} \frac{\chi C_m}{\Delta t} M \mathbf{u}^n + \chi M I_{ion}^h(\mathbf{u}^n, \mathbf{w}^{n+1}) + M I_i^{app} \\ M I_i^{app} - M I_e^{app} \end{bmatrix},$$

while the full discretization of the (29) reads:

$$\begin{bmatrix} \frac{\chi C_m}{\Delta t} M + \frac{\lambda}{1+\lambda} K_i & \frac{\lambda}{1+\lambda} K_i - \frac{1}{1+\lambda} K_e \\ K_i & K_i + K_e \end{bmatrix} \begin{bmatrix} \mathbf{u} \\ \mathbf{u}_e \end{bmatrix}^{n+1} = \begin{bmatrix} \frac{\chi C_m}{\Delta t} M \mathbf{u}^n + \chi M I_{ion}^h(\mathbf{u}^n, \mathbf{w}^{n+1}) + \frac{\lambda}{1+\lambda} M I_i^{app} + \frac{1}{1+\lambda} M I_e^{app} \\ M I_i^{app} - M I_e^{app} \end{bmatrix}. \quad (37)$$

While the fully discrete formulation of (28) is symmetric, and can be solved by PCG, the fully discrete formulation of (29) system results naturally in a non-symmetric matrix at the discrete level. The resulting linear system can then be solved with an iterative method such GMRES or Bi-CGSTAB (see [51]), using as initial guess the solution at the previous time step.

4.3.3 Monodomain

The full discretization of the Monodomain system (32) reads:

$$\left[\frac{\chi C_m}{\Delta t} M + K_m \right] [\mathbf{u}]^{n+1} = [\mathbf{f}]^{n+1}.$$

The associated matrix is naturally symmetric, and the linear system can be solved by a PCG, using as initial guess the solution at the previous time step.

4.3.4 Computational Aspects

The numerical solution of the Bidomain system is an expensive computational task. If, on the one hand, the IMEX approach described above allows to avoid the costly solution of nonlinear systems, the degenerate nature of the Bidomain itself entails a very ill conditioning for the linear system associated to its full discretization. To cope with such computational complexity, several scientists have developed in the recent years effective preconditioning strategies to reduce the high computational costs associated to its numerical solution [15, 37, 43–45, 54, 62, 63]. Among these works, most are based on a proper decomposition of the computational domain in order to set up parallel preconditioners, or on suitable multigrid schemes still coupled with parallel architectures. The PE formulation is a popular choice in the Bioengineering community, as it is computationally more stable and allows for decoupled approaches in the solution. In fact, performing the space discretization first, results in a Differential-Algebraic system in time. It can then be natural to decouple the differential part from the algebraic one, and use the elliptic equation of the PE formulation as a corrector step in a two-level scheme (see, e.g. [20]). An efficient serial preconditioner for the PE formulation has been proposed in [22, 24] stemming from a suitable extension of the Monodomain model, and resulting in a lower block-triangular preconditioner for (29). Adaptive techniques have been proposed as well, to better capture the front propagation of the electrical excitation. An efficient adaptive strategy, reshaping the computational grid according to some suitable a posteriori error estimate has been introduced for instance in [7]. However, since the problem is time dependent, such approach requires to frequently recompute the mesh and interpolate between the old and the new grid, a feature that can become a serious bottleneck when simulating reentrant waves or fibrillation. Stemming from the observation that the Monodomain provides

an accurate approximation of the potential in most of the region of interest, a *model adaptive strategy* has been proposed as well, which aims at reducing computational costs and to maintain the accuracy by solving the Bidomain problem only in (hopefully) small, critical (in physiopathological terms) regions of the domain, while the Monodomain equation is solved in the remaining regions, where the potential propagation dynamics does not require the most sophisticated model. A first version of this approach was presented in [36], where a suitable a posteriori model estimator was introduced, and an hybrid model called *Hybridomain* was advocated. The latter assembles the block (1,2) in (37) only in correspondence with the nodes identified as Bidomain ones by the model estimator, while the second equation stays untouched, and a problem the same size of the original Bidomain model has to be solved. If the constant λ in (29) is properly chosen, the block (1,1) of (37) is actually the discretization of the Monodomain model. Following this consideration, an improved version of the model adaptive strategy has been introduced in [26], where only the block (1,1) of (37) is solved in the Monodomain regions. The coupling between regions was based on Optimized Schwarz Methods [23, 25], a popular technique in the field of Domain Decomposition algorithms (see e.g. [21, 58]), which relies on Robin transmission conditions on the interface between subdomains.

Acknowledgements This work was supported by Basque Government through the BERC 2014-2017 program, and by the Spanish Ministry of Economics and Competitiveness MINECO: BCAM Severo Ochoa excellence accreditation SEV-2013-0323.

References

1. Ambrosio, L., Colli Franzone, P., Savaré, G.: On the asymptotic behaviour of anisotropic energies arising in the cardiac bidomain model. *Interfaces Free Bound* **2**, 213–266 (2000)
2. Bendahmane, M., Karlsen, K.H.: Analysis of a class of degenerate reaction-diffusion systems and the bidomain model of cardiac tissue. *Netw. Heterog. Media* **1**, 185–218 (2006)
3. Bourgault, Y., Coudière, Y., Pierre, C.: Existence and uniqueness of the solution for the bidomain model used in cardiac electrophysiology. *Nonlinear Anal. Real World Appl.* **10**, 458–482 (2009)
4. Clayton, R.H., Panfilov, A.V.: A guide to modelling cardiac electrical activity in anatomically detailed ventricles. *Prog. Biophys. Mol. Biol.* **96**, 19–43 (2008)
5. Clayton, R.H., Bernus, O.M., Cherry, E.M., Dierckx, H., Fenton, F.H., Mirabella, L., Panfilov, A.V., Sachse, F.B., Seemann, G., Zhang, H.: Models of cardiac tissue electrophysiology: progress, challenges and open questions. *Prog. Biophys. Mol. Biol.* **104**, 22–48 (2011)
6. Clements, J.C., Nenonen, J., Li, P.K.J., Horacek, M.: Activation dynamics in anisotropic cardiac tissue via decoupling. *Ann. Biomed. Eng.* **32**, 984–990 (2004)
7. Colli Franzone, P., Deuffhard, P., Erdmann, B., Lang, J., Pavarino, L.F.: Adaptivity in space and time for reaction-diffusion systems in electrocardiology. *SIAM J. Sci. Comput.* **28**(3), 942–962 (2006)
8. Cortassa, S., Aon, M.A., Marbán, E., Winslow, R.L., O'Rourke, B.: An integrated model of cardiac mitochondrial energy metabolism and calcium dynamics. *Biophys. J.* **84**, 2734–2755 (2003)

9. Coudière, Y., Rioux, M.: Virtual electrodes mechanisms predictions with a current-lifted monodomain model. *Comput. Cardiol.* **39**, 837–840 (2012)
10. Coudière, Y., Bourgault, Y., Rioux, M.: Optimal monodomain approximations of the bidomain equations used in cardiac electrophysiology. *Math. Models Methods Appl. Sci.* **24**, 1115–1140 (2014)
11. Courtemanche, M., Ramirez, R.J., Nattel, S.: Ionic mechanisms underlying human atrial action potential properties: insights from a mathematical model. *Am. J. Physiol. Heart Circ. Physiol.* **275**(44), H301–H321 (1998)
12. Earm, Y.E., Noble, D.: A model of the single atrial cell: relation between calcium current and calcium release. *Pro. R. Soc. Lond. B Biol. Sci.* **240**, 83–96 (1990)
13. FitzHugh, R.: Impulses and physiological states in theoretical models of nerve membrane. *Biophys. J.* **1**, 445–466 (1961)
14. Franzone, P.C., Guerri, L.: Spread of excitation in 3-D models of the anisotropic cardiac tissue, I: validation of the eikonal approach. *Math. Biosci.* **113**, 145–209 (1993)
15. Franzone, P.C., Pavarino, L.F.: A parallel solver for reaction-diffusion systems in computational electrocardiology. *Math. Models Methods Appl. Sci.* **14**(6), 883–911 (2004)
16. Franzone, P.C., Savaré, G.: Degenerate evolution systems modeling the cardiac electric field at micro and macroscopic level. In: Lorenzi, A., Ruf, B. (eds.) *Evolution Equations, Semigroups and Functional Analysis*, pp. 49–78. Birkhauser, Basel (2002)
17. Franzone, P.C., Guerri, L., Pennacchio, M., Taccardi, B.: Spread of excitation in 3-D models of the anisotropic cardiac tissue, II: effect of the fiber architecture and ventricular geometry. *Math. Biosci.* **147**, 131–171 (1998)
18. Franzone, P.C., Pavarino, L.F., Taccardi, B.: Simulating patterns of excitation, repolarization and action potential duration with cardiac Bidomain and Monodomain models. *Math. Biosci.* **197**, 35–66 (2005)
19. Franzone, P.C., Pavarino, L.F., Savaré, G.: Computational electrocardiology: mathematical and numerical modeling. In: Quarteroni, A., Formaggia, L., Veneziani, A. (eds.) *Complex Systems in Biomedicine*. Springer, Milan (2006)
20. Franzone, P.C., Pavarino, L.F., Scacchi, S.: Parallel coupled and uncoupled multilevel solvers for the bidomain model of electrocardiology. In: *Domain Decomposition Methods in Science and Engineering XXI*, pp. 257–264. Springer, Cham (2014)
21. Gander, M.J.: Optimized Schwarz methods. *SIAM J. Numer. Anal.* **44**, 699–731 (2006)
22. Gerardo-Giorda, L., Mirabella, L.: Spectral analysis of a block-triangular preconditioner for the Bidomain system in electrocardiology. *Electron. Trans. Numer. Anal.* **39**, 186–201 (2012)
23. Gerardo-Giorda, L., Perego, M.: Optimized Schwarz methods for the bidomain system in electrocardiology. *Math. Model. Numer. Anal.* **47**(2), 583–608 (2013)
24. Gerardo-Giorda, L., Mirabella, L., Nobile, F., Perego, M., Veneziani, A.: A model-based block-triangular preconditioner for the Bidomain system in electrocardiology. *J. Comput. Phys.* **228**, 3625–3639 (2009)
25. Gerardo-Giorda, L., Perego, M., Veneziani, A.: Optimized Schwarz coupling of bidomain and monodomain models in electrocardiology. *Math. Model. Numer. Anal.* **45**(2), 309–334 (2011)
26. Gerardo-Giorda, L., Mirabella, L., Perego, M., Veneziani, A.: Optimized Schwarz Methods and model adaptivity in electrocardiology simulations. In: *Domain Decomposition Methods in Science and Engineering XXI*, pp. 265–272. Springer, Cham (2014)
27. Henriquez, C.S.: Simulating the electrical behavior of cardiac tissue using the bidomain model. *Crit. Rev. Biomed. Eng.* **21**, 1–77 (1993)
28. Hodgkin, A.L., Huxley, A.F.: A quantitative description of membrane current and its application to conduction and excitation in nerve. *J. Physiol.* **117**, 500–544 (1952)
29. Iyer, V., Mazhari, R., Winslow, R.: A computational model of the human left-ventricular epicardial myocyte. *Biophys. J.* **87**, 1507–1525 (2004)
30. Keener, J.P.: An eikonal-curvature equation for the action potential propagation in myocardium. *J. Math. Biol.* **29**, 629–651 (1991)
31. Keener, J.P.: Direct activation and defibrillation of cardiac tissue. *J. Theor. Biol.* **178**, 313–324 (1996)

32. Keener, J.P., Sneyd, J.: *Mathematical Physiology*. Springer, New York (1998)
33. Le Grice, J., Smaill, B.H., Hunter, P.J.: Laminar structure of the heart: a mathematical model. *Am. J. Physiol. Heart Circ. Physiol.* **272**(41), H2466–H2476 (1995)
34. Lines, G.T., Buist, M.L., Grottum, P., Pullan, A.J., Sundnes, J., Tveito, A.: Mathematical models and numerical methods for the forward problem in cardiac electrophysiology. *Comput. Vis. Sci.* **5**, 215–239 (2003)
35. Luo, C., Rudy, Y.: A dynamic model of the cardiac ventricular action potential. *Circ. Res.* **74**, 1071–1096 (1994)
36. Mirabella, L., Nobile, F., Veneziani, A.: An a posteriori error estimator for model adaptivity in electrocardiology. *Comput. Methods Appl. Mech. Eng.* **200**(37–40), 2727–2737 (2011)
37. Munteanu, M., Pavarino, L.F., Scacchi, S.: A scalable Newton–Krylov–Schwarz method for the Bidomain reaction-diffusion system. *SIAM J. Sci. Comput.* **31**, 3861–3883 (2009)
38. Neu, J.C., Krassowska, W.: Homogenization of syncytial tissues. *Crit. Rev. Biomed. Eng.* **21**, 137–199 (1993)
39. Nielsen, B.F., Ruud, T.S., Lines, G.T., Tveito, A.: Optimal monodomain approximation of the bidomain equations. *Appl. Math. Comput.* **184**, 276–290 (2007)
40. Nygren, A., Fiset, C., Firek, L., Clark, J.W., Lindblad, D.S., Clark, R.B., Giles, W.R.: Mathematical model of an adult human atrial cell: the role of K⁺ currents in repolarization. *Circ. Res.* **82**, 63–81 (1998)
41. O'Hara, T., Virág, L., Varró, A., Rudy, Y.: Simulation of the undiseased human cardiac ventricular action potential: model formulation and experimental validation. *PLoS Comput. Biol.* **7**, e1002061 (2011)
42. Pashaei, A., Romero, D., Sebastian, R., Camara, O., Frangi, A.F.: Fast Multiscale modeling of cardiac electrophysiology including purkinje system. *IEEE Trans. Biomed. Eng.* **58**, 2956–2960 (2011)
43. Pavarino, L.F., Scacchi, S.: Multilevel additive Schwarz preconditioners for the Bidomain reaction-diffusion system. *SIAM J. Sci. Comput.* **31**, 420–443 (2008)
44. Pennacchio, M., Simoncini, V.: Efficient algebraic solution of reaction-diffusion systems for the cardiac excitation process. *J. Comput. Appl. Math.* **145**, 49–70 (2002)
45. Pennacchio, M., Simoncini, V.: Algebraic multigrid preconditioners for the bidomain reaction-diffusion system. *Appl. Numer. Math.* **59**, 3033–3050 (2009)
46. Perego, M., Veneziani, A.: An efficient generalization of the Rush-Larsen method for solving electro-physiology membrane equations. *Electron. Trans. Numer. Anal.* **35**, 234–256 (2009)
47. Potse, M., Dubé, B., Richer, J., Vinet, A.: A comparison of monodomain and bidomain reaction-diffusion models for action potential propagation in the human heart. *IEEE Trans. Biomed. Eng.* **53**(12), 2425–2435 (2006)
48. Quarteroni, A., Valli, A.: *Numerical Approximation of Partial Differential Equations*. Springer, Berlin (1994)
49. Rogers, J., McCulloch, A.: A collocation-Galerkin finite element model of cardiac action potential propagation. *IEEE Trans. Biomed. Eng.* **41**, 743–757 (1994)
50. Roth, B.J.: Action potential propagation in a thick strand of cardiac muscle. *Circ. Res.* **68**, 162–173 (1991)
51. Saad, Y.: *Iterative Methods for Sparse Linear Systems*, 2nd edn. Society for Industrial and Applied Mathematics, Philadelphia (2003)
52. Sachse, F.B.: *Computational Cardiology*. Springer, Berlin (2004)
53. Sanfelici, S.: Convergence of the Galerkin approximation of a degenerate evolution problem in electrocardiology. *Numer. Methods Partial Differ. Equ.* **18**, 218–240 (2002)
54. Scacchi, S.: A hybrid multilevel Schwarz method for the bidomain model. *Comput. Methods Appl. Mech. Eng.* **197**, 4051–4061 (2008)
55. Sepulveda, N.G., Roth, B.J., Wikswo, J.P.: Current injection into a two-dimensional anisotropic bidomain. *Biophys. J.* **55**, 987–999 (1989)
56. Streeter, D.: Gross morphology and fiber geometry in the heart. In: Berne, R.M. (ed.) *Handbook of Physiology*, vol. 1, pp. 61–112. Williams and Wilkins, Baltimore (1979)

57. Ten Tusscher, K.H., Noble, D., Noble, P.J., Panfilov, A.V.: A model for human ventricular tissue. *Am. J. Physiol. Heart. Circ. Physiol.* **286**, H1573–H1589 (2004)
58. Toselli, A., Widlund, O.: *Domain Decomposition Methods*, 1st edn. Springer, Heidelberg (2004)
59. Trayanova, N.: Defibrillation of the heart: insights into mechanisms from modelling studies. *Exp. Physiol.* **91**, 323–337 (2006)
60. Vazquez, M., Aris, R., Houzeaux, G., Aubry, R., Villar, P.: A massively parallel computational electrophysiology model of the heart. *Int. J. Numer. Methods Biomed. Eng.* **27**(12), 1911–1929 (2011)
61. Veneroni, M.: Reaction-diffusion systems for the macroscopic bidomain model of the cardiac electric field. *Nonlinear Anal. Real World Appl.* **10**, 849–868 (2009)
62. Vigmond, E.J., Aguel, F., Trayanova, N.A.: Computational techniques for solving the bidomain equations in three dimensions. *IEEE Trans. Biomed. Eng.* **49**(11), 1260–1269 (2002)
63. Vigmond, E.J., Weber dos Santos, R., Prassl, A.J., Deo, M.D., Plank, G.: Solvers for the cardiac bidomain equations. *Prog. Biophys. Mol. Biol.* **96**(1–3), 3–18 (2008)
64. Weber dos Santos, R., Planck, G., Bauer, S., Vigmond, E.J.: Parallel multigrid preconditioner for the cardiac bidomain model. *IEEE Trans. Biomed. Eng.* **51**(11), 1960–1968 (2004)
65. Winfree, A.T.: Varieties of spiral wave behavior: an experimentalist's approach to the theory of excitable media. *Chaos* **1**, 303–334 (1991)

Mechanisms Underlying Electro-Mechanical Cardiac Alternans

**Blas Echebarria, Enric Alvarez-Lacalle, Inma R. Cantalapiedra,
and Angelina Peñaranda**

Abstract Electro-mechanical cardiac alternans consists in beat-to-beat changes in the strength of cardiac contraction. Despite its important role in cardiac arrhythmogenesis, its molecular origin is not well understood. The appearance of calcium alternans has often been associated to fluctuations in the sarcoplasmic reticulum calcium level (SR Ca load). However, cytosolic calcium alternans observed without concurrent oscillations in the SR Ca content suggests an alternative mechanism related to a dysfunction in the dynamics of the ryanodine receptor (RyR2). In this chapter we review recent results regarding the relative role of SR Ca content fluctuations and SR refractoriness for the appearance of alternans in both ventricular and atrial cells.

1 Introduction

Calcium is the cellular messenger that drives the contraction of cardiac cells [1]. Problems in the regulation of intracellular calcium underlie a large number of cardiac dysfunctions [2].

For example, a variety of cardiac arrhythmias are initiated by a focal excitation whose origin is a large release of calcium from the sarcoplasmic reticulum (SR). If this excitation is not timed with the external signaling provided by the depolarization of the tissue, then it generates either delayed (DAD) or early afterdepolarizations (EAD) that may provide a proarrhythmogenic substrate and lead to reentry and conduction blocks [3–6]. There is also increasing evidence that problems in calcium handling dynamics are the main origin of cardiac alternans [7], a well known proarrhythmic condition in which there is a beat-to-beat alternation in the strength of cardiac contraction [8]. Calcium alternans has been linked to the appearance and break-up of spiral waves, giving rise to tachycardia and ventricular fibrillation [9, 10]. This has prompted both modeling and experimental work on the dynamics of intracellular calcium, with the aim to unveil the molecular and dynamical mechanisms behind the arrhythmic pathways.

B. Echebarria (✉) • E. Alvarez-Lacalle • I.R. Cantalapiedra • A. Peñaranda
Departament de Física, Universitat Politècnica de Catalunya, 08028 Barcelona, Spain
e-mail: blas.echebarria@upc.edu

2 Electromechanical Alternans

Cardiac alternans has long been recognized as an important proarrhythmic factor [9]. It typically occurs at fast pacing rates, and is characterized by a change in the duration of the action potential (APD) from beat to beat, that becomes more pronounced the faster the pacing rate is. For very short stimulation periods the amplitude of the oscillations of APD may become so large that the cells are not able to recover their electrical properties before the next stimulation, making them unable to elicit a new action potential.

If the distribution of APDs in tissue is heterogeneous (that may occur because of gradients of electrophysiological properties, or it can be dynamically generated [11]), then when an action potential wave propagates, it may encounter a region that is still refractory, due to a previous large action potential duration. This results in localized conduction block that can, in turn, induce the formation of reentry, with the initiation of rotors that impose a rhythm of contraction much faster than the typical sinus frequency. The resulting state is that of ventricular tachycardia (VT) or atrial flutter (AFL), depending on the location of the rotors. If the rotors are unstable, then multiple wavelets are created, giving rise to a state of ventricular (VF) or atrial (AF) fibrillation, in which different parts of tissue are not able to contract synchronously. When this happens in the ventricles, the pumping of blood to the body is impeded, resulting in death in a few minutes unless a successful defibrillatory shock is provided. Atrial fibrillation, on the other hand, despite not being mortal, results in an important decrease in life quality. Due to its big prevalence, specially in people over 65 year old, it has become a big health concern.

The origin of alternans has been linked to problems in the transmembrane currents, that give rise to a steep restitution curve (relation between the duration of the action potential and the time elapsed since the end of the previous excitation, or diastolic interval (DI)). When the restitution curve is shallow, a cell (or tissue) is able to adapt the duration of its action potential to a decrease in the pacing period. If the curve is steep, though, any small change in the DI produces a big change in the APD, resulting in an unstable situation that gives rise to a 2:2 response, with an alternating sequence of long and short action potential durations, i.e., alternans [8]. In the simplest models this occurs when the slope of the restitution curve is larger than one.

Due to the coupling of transmembrane potential and intracellular calcium concentration through the action of the L-type calcium current and the sodium-calcium exchanger, a change in the duration of the action potential results in a change in the calcium transient, originating thus calcium alternans. In this view, calcium alternans would just be a consequence of action potential alternans. Experiments with action potential clamps, though, have shown that often calcium alternans is the origin and not the consequence of action potential alternans [12]. This agrees with experiments on intact heart [13], where transmembrane potential and intracellular calcium were simultaneously measured, and it was shown that calcium alternans occurs at slower pacing rates than APD alternans and always precedes it. Also, in situations where alternans has been observed to precede the transition to atrial fibrillation (AF),

the slope of the restitution curve was measured to be smaller than one [14]. This does not completely preclude AP alternans as the primary mechanism (there could be effects of memory, etc.), but all in all, there is a strong evidence that calcium alternans are the main alternans mechanism in many situations.

3 Mechanisms Explaining the Appearance of Calcium Alternans

Under sinus stimulation, a change in the transmembrane potential of cardiac myocytes opens the L-type calcium channels (LCC) of the cell membrane producing an influx of calcium ions. Within the cell, most of the calcium is sequestered in a bag known as the sarcoplasmic reticulum (SR). At rest, the concentration of calcium in the cytosol is around a hundred nanomolar, while in the SR, it is of the order of the millimolar (in the extracellular medium it is also typically of the order of 2 mM). The membrane of the sarcoplasmic reticulum is spotted by calcium sensitive channels, the ryanodine receptors (RyR2). Both LCCs and RyR2 appear grouped in clusters, of 1–5 LCCs and 50–100 RyRs in each. In ventricular cells, the cellular membrane has intubulations (t-tubules), with the consequence that LCC and RyR2 are always confronted (Fig. 1). In atrial myocytes, the presence of t-tubules is controversial [15]. In their absence, only RyR2s that are close to the cell membrane will be next

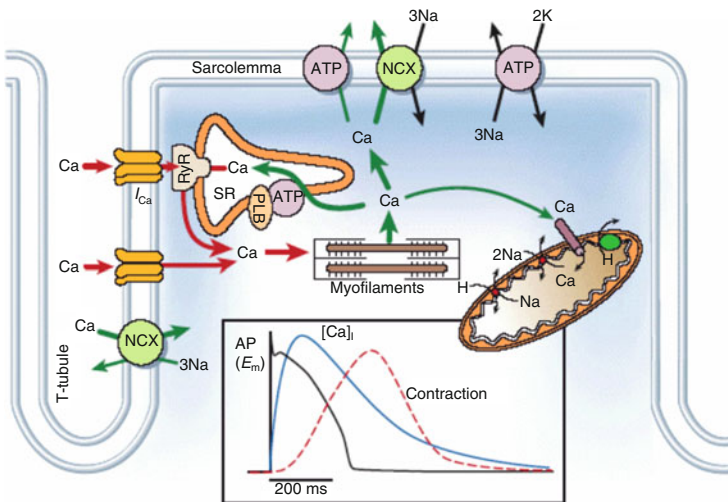


Fig. 1 Electrical excitation opens the voltage-gated Ca^{2+} channels (LCC), resulting in Ca^{2+} entry that induces Ca^{2+} release from the sarcoplasmic reticulum (SR) through the opening of the ryanodine receptors (RyRs), giving rise to cell contraction. *Inset* shows the time evolution of the transmembrane potential, the Ca^{2+} transient and contraction. (Reproduced with permission from Bers, Nature, 2002 [1])

to LCCs. This distinction between the geometry of ventricular and atrial cells has in fact a significant relevance for the dynamics of calcium inside the cell.

Once the RyR2s open due to the binding of the calcium ions that enter the cell through the LCC, the calcium of the SR is released, in a process known as calcium induced calcium release (CICR), resulting in an elevation of the concentration of calcium in the cytosol, up to $\sim 50 - 200 \mu\text{M}$ in the dyadic space corresponding to the LCC-RyR2 junction, and to $\sim 1 \mu\text{M}$ in the bulk cytosol. Then, calcium binds to several buffer proteins, including Troponin C, that drives the tropomyosin complex off the actin binding site allowing the binding of myosin and producing a shortening in the actin filaments and the contraction of the cardiomyocyte. Calcium in the cytosol is finally pumped out of the cell and into the SR by the sodium calcium exchanger (NCX) and the SERCA pump, respectively. Besides driving contraction, the calcium transient also has an influence in the form and duration of the action potential, both due to the NCX pump and to the L-type calcium current, since the LCC are also inactivated by high calcium concentrations.

For calcium alternans to develop there must be some effect that reduces the amount of calcium released from the SR in alternate beats. Two possible explanations for this are that either the SR is not completely full (because of a slow SERCA, for instance) or, if it is full, the RyR2s do not open completely (because of a long refractory time). Alternans then appears when two conditions are met: there exists a nonlinear dependence of release on either RyR2 recovery or SR load [16] together with a slow release recovery time scale (Fig. 2). Alternations in the strength of the calcium current I_{CaL} has also been cited as a possible cause, but calcium alternans has been detected without L-type calcium alternations [17, 18], suggesting that I_{CaL} modulations are the result of calcium alternans and not its cause. Still, a stronger I_{CaL} current, as well as a higher rate of spontaneous calcium release (related probably to fast activation kinetics of the RyR2) have also been related to the appearance of alternans in atrial cells [19].

3.1 *SR Calcium Fluctuations*

The relevance of a partial refill of the SR was supported by experiments by Diaz et al. [20], where they observed fluctuations in the content of the SR during calcium alternans. This mechanism needs a strong nonlinear dependence of calcium release with SR calcium load and a stimulation period that is fast compared with the typical refilling times given by SERCA (Fig. 2). Then, when the SR is filled, a large release is produced so SERCA is not able to completely refill the SR by the time the next stimulation is given. This results in a small depletion in SR calcium concentration, that is now able to recover the original concentration before the following stimulation (see Fig. 2).

Theoretical work [21, 22] has explained the appearance of alternans under this mechanism with the help of a map relating calcium release and calcium load, showing that alternans appears when the slope of this map is larger than

Alternans due to

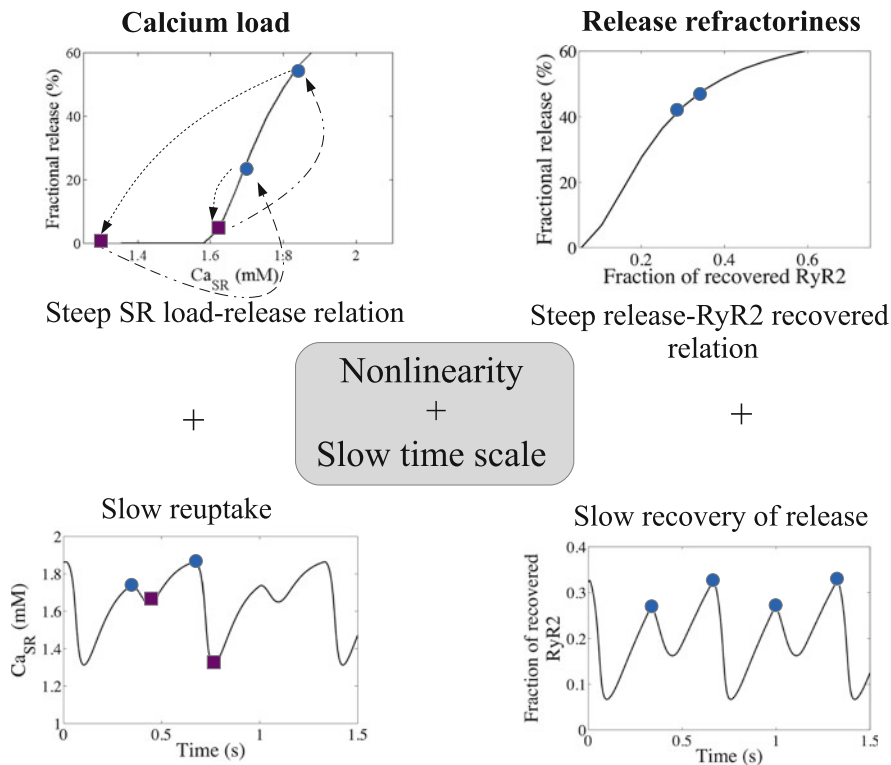


Fig. 2 Mechanisms for the appearance of calcium alternans. It is necessary a nonlinear dependence of calcium release (either on calcium load or in the number of recovered RyR2s) and a slow time scale (related to SR refilling in one case, to RyR2 recovery, in the other)

one, denoting the onset of a period doubling bifurcation. Alternatively, the strong dependence of calcium released from the SR with SR calcium load has also been linked to the binding of calcium to the buffer calsequestrin [22].

Although this mechanism can underlie calcium alternans in some situations, observation by Pitch et al. [23] of cytosolic calcium alternans without concurrent SR fluctuations suggests that in some cases the mechanism must be linked to refractoriness in the release.

3.2 RyR Refractoriness

A possible explanation for the occurrence of alternans at constant diastolic SR calcium load is the presence of a long refractory time of the RyR2, such that, after a large release, it takes a long time to recover and be able to open. Thus, even if

the SR is completely refilled, release is weak since most of the RyR2s have not recovered [17, 18, 22, 24]. For this mechanism to work there must exist a nonlinear dependence between calcium released from the SR and the number of recovered RyR2s. This is provided by the cooperativity in RyR2 opening by calcium [25].

Recently, techniques have been developed that allow the simultaneous measurement of transmembrane voltage and intracellular or SR calcium concentrations in an ex vivo heart [13]. From these experiments one can conclude that calcium release alternans appears with or without diastolic sarcoplasmic reticulum calcium alternans [13]. In the first case, SR release alternans occurs at slower rates than diastolic SR calcium alternans. Sensitization of RyR2 with low doses of caffeine decreases the magnitude and the threshold for induction of alternans, suggesting RyR2 refractoriness as the underlying mechanism. Interestingly, minimum release is produced in VF apparently due to continuous RyR2 refractoriness. This agrees with in silico analysis of alternans preceding AF [26], that showed inactivation kinetics of RyR2 as the only parameter able to reproduce experimental results.

3.3 Relevance of Each Mechanism

Depending on the working state of the RyR2, SR calcium load, strength of SERCA, etc., SR calcium load or RyR refractoriness can be the molecular mechanism giving rise to alternans. This is important to know, because a drug that suppresses alternans in one case may not work necessarily for the other mechanism, and could even have negative consequences. Using mathematical models of calcium handling, a possible way to elucidate which is the responsible mechanism to perform clamps of the appropriate factors. To this end, in [27] clamps were performed on either the prediastolic SR content or the number of recovered RyR2s, using a rabbit ventricular model [28]. The former was achieved increasing the strength of SERCA during a few milliseconds previous to an excitation, so it reached the same level in all beats (Fig. 3a). In the latter, the kinetics of the RyR2 was modified, again during the last

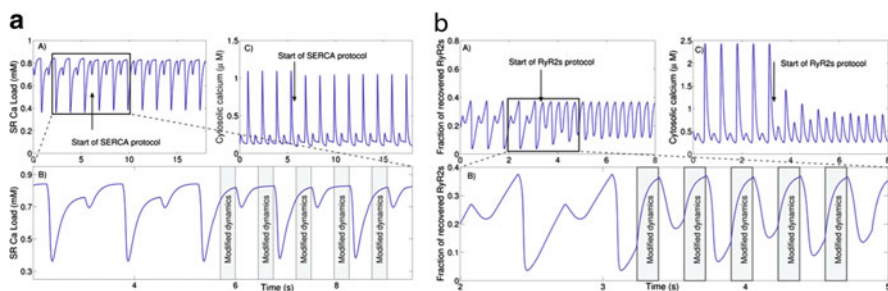


Fig. 3 Example of (a) SR and (b) RyR clamps. The clamp is always performed during the recovery period, not to affect release from the SR (from [27])

milliseconds previous to the stimulation, until the same number of recovered RyR2 was obtained as in the previous beat (Fig. 3b).

If alternans persists when the SR is clamped, then one can conclude that the instability is due to RyR2 refractoriness. On the contrary, if it persists when the RyR2 is clamped, then we conclude that it is due to SR calcium fluctuations. There are some cases where alternans persists with either clamp. In this case either mechanism can give rise to alternans. In some other cases, both clamping protocols make alternans disappear. Then, both have to collaborate to be able to sustain alternans (Fig. 4).

Depending on the conditions of the RyR2, it was obtained that at low activation of the RyR2, typically the responsible mechanism was RyR2 refractoriness (specifically, a slow recovery from inactivation), while at low inactivation, the predominant mechanism was related to fluctuations in SR load (Fig. 4, for more details see [27]).

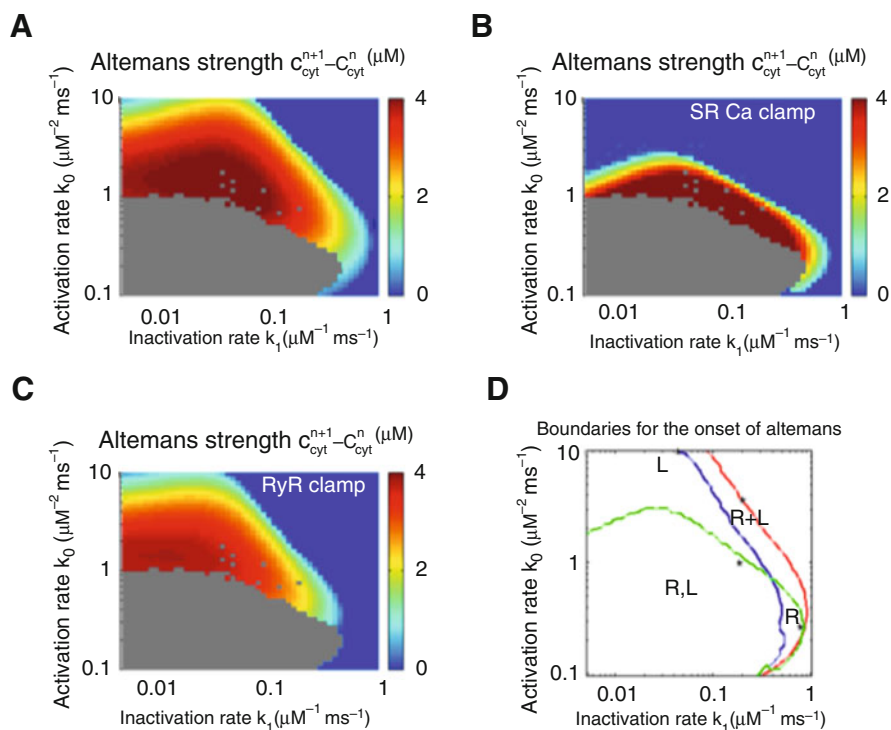


Fig. 4 Colormaps showing the difference in peak intracellular calcium in two consecutive beats as a function of RyR activation and inactivation rates, using a rabbit ventricular cell model [28]. (a) Under normal conditions and with a (b) SR and (c) RyR clamp. The displacement of the lines denoting the transition to alternans in each case delimitates the regions where each mechanism is relevant (d) [27]

4 Nature of the Instability

Besides the molecular mechanism responsible for the transition, it is important to understand the nature of the transition to alternans. Using whole cell models, that consider average concentrations, the transition to alternans typically appears as a period doubling bifurcation of the original periodic solution. This can be understood constructing a map. In the case where alternans is due to a slow recovery of the RyRs from inactivation, one can, for instance, draw a map with the relation between the number of recovered RyR2s at two consecutive stimulations (Fig. 5). As the stimulation period is decreased, this map develops a region with a negative slope larger than -1 , denoting the onset of the period doubling instability [27].

In deterministic calcium models, the transition to alternans is sharp (Fig. 5) and corresponds to a period doubling (PD) bifurcation [27]. An example of alternans appearing in an atrial model is shown in Fig. 6, together with the bifurcation diagram, showing that alternans appears through a supercritical period doubling

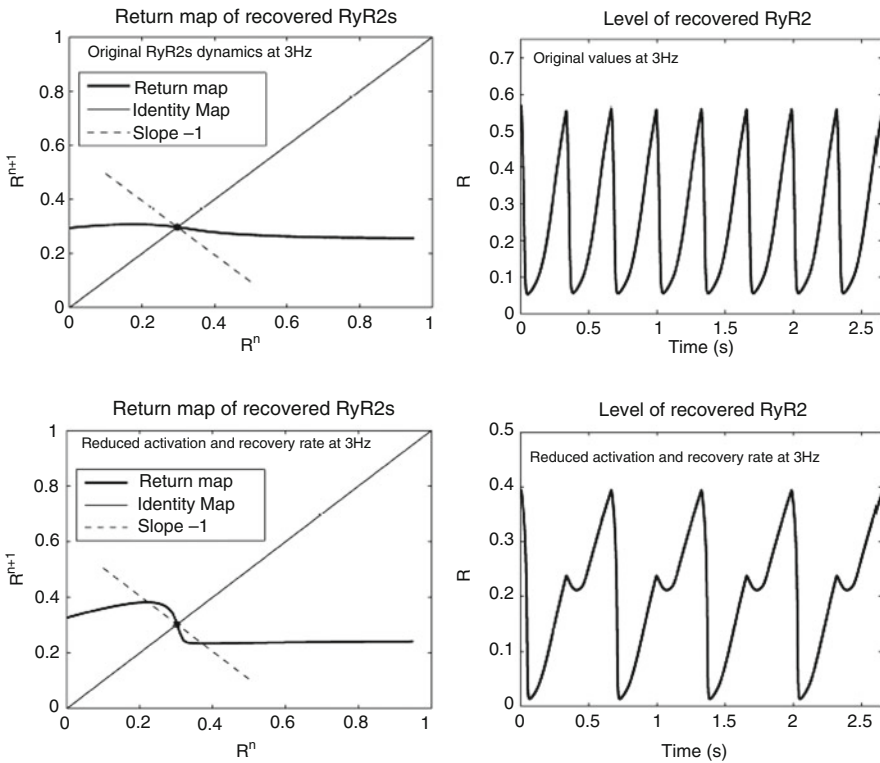


Fig. 5 Maps giving the relation between the number of recovered RyR2s at consecutive beats, as well as time traces of the number of recovered RyR2s

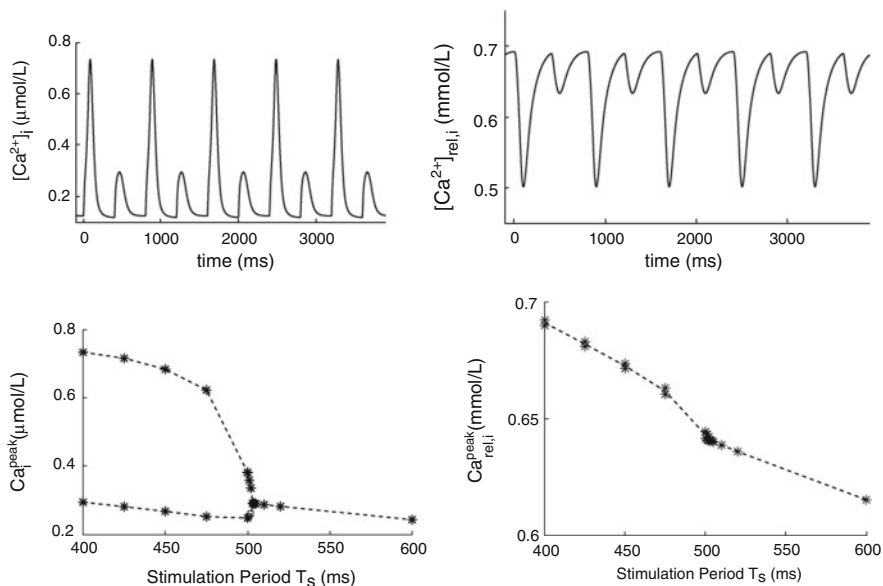


Fig. 6 Alternans appearing in a human atrial model [29, 32]. *Top*: Time traces of intracellular calcium, SR Ca concentration and transmembrane potential. *Bottom*: Bifurcation diagrams for the peak intracellular and SR calcium concentrations. During alternans peak SR concentration does not vary from beat to beat (reproduced with permission from C.A. Lugo et al., Am j Physiol 306, H1540, 2014 [29])

(PD) bifurcation. In this case, SR fluctuations are absent, so the mechanism is linked to RyR refractoriness [29].

In reality, the situation is a bit more complex. In the cell, the RyR2s are distributed in clusters of 50–100 elements and their dynamics is stochastic. So it is for the opening of the LCC channels. The distance among RyR2 clusters (or calcium release units, CaRUs) is of the order of the micrometer, so a typical cardiac cell is composed of $\sim 20,000$ CaRUs. In ventricular cells each CaRU is composed of a cluster of 50–100 RyR confronted to 1–5 LCC, with calcium diffusing between different CaRUs. Thus, there are thousands of stochastic elements diffusively coupled. Locally the dynamics is dominated by random calcium releases (sparks). Globally, this gives rise to a well defined alternating state and bifurcation diagrams (Fig. 7). In the limit of very strong coupling or large clusters (many RyR2s in any CaRU) one recovers the deterministic limit and the PD bifurcation. The relevance of stochastic effects has been stressed by Rovetti et al. [30], that propose that the onset of whole cell alternans from random calcium release events (sparks) is due to the interplay of three effects: randomness, recruitment and refractoriness. Ion channel stochasticity at the level of single calcium release units has also been shown to influence the whole-cell alternans dynamics by causing phase reversals over many beats during fixed frequency pacing close to the alternans bifurcation [31].

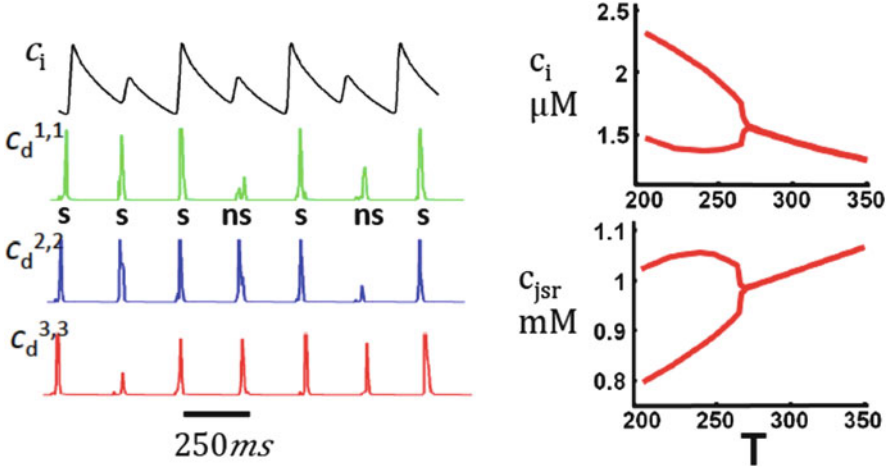


Fig. 7 *Left:* Times traces of global intracellular calcium concentration (*top*) and dyadic calcium concentration at given CaRUs. *Right:* Bifurcation diagram for the global intracellular (*top*) and SR (*bottom*) calcium concentrations. (Reproduced with permission from Alvarez-Lacalle et al., PRL 2015 [33])

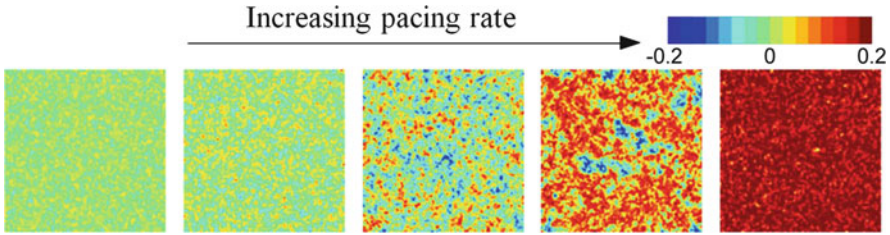


Fig. 8 Beat-to-beat difference in peak calcium concentration. From left to right, snapshots with characteristic spatial distribution of alternans as the pacing period approaches the critical point

Considering in more detail the transition, one observes that alternans first appears at a local level, but with a very short persistence both in space and time (Fig. 8). When one measures the average calcium concentration over the entire cell, these local alternations thus disappear. The appearance of global alternans needs the synchronization among different clusters with calcium oscillating in opposite phase [33]. To study this synchronization behavior, one can define an order parameter, that takes values 1 or -1 , depending on the phase of oscillation of calcium at a given point. Then:

$$m_{kl}(n) = \text{sgn}[(-1)^n(c_{jsr}^{kl}(n) - c_{jsr}^{kl}(n-1))] \quad (1)$$

and thus define the degree of synchronization over the entire cell as $\langle m \rangle = \sum_{kl} m_{kl}$. The first thing to notice is that the global order parameter fluctuates in time close

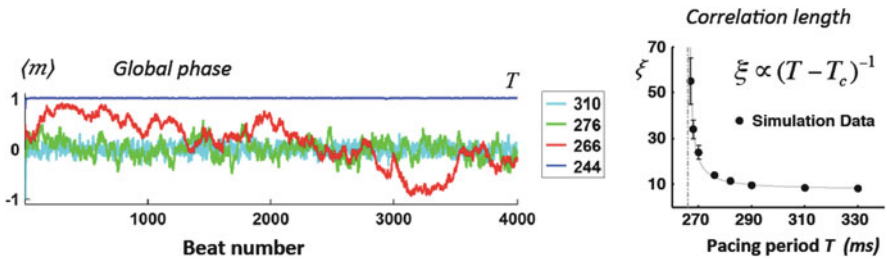


Fig. 9 *Left:* Time evolution of the average synchronization $\langle m \rangle$ at different pacing periods. *Right:* correlation length, showing a divergence at a critical pacing period (Adapted with permission from Alvarez-Lacalle et al., PRL 2015 [33])

Table 1 Exponents of the transition to alternans [33]

	γ/ν	β/ν	ν
Ventricular model	1.75 ± 0.08	0.18 ± 0.11	0.91 ± 0.14
Ising model	1.75	0.125	1

to the transition point and presents a diverging correlation length (Fig. 9). This resembles a second order phase transition in equilibrium systems. To ascertain the nature of the transition, one can study the critical exponents as the system approaches the critical point (the critical stimulation period, in this case). From the theory of critical phenomena, it is known that $\langle m \rangle \sim (T - T_c)^{-\beta/\nu}$, and the susceptibility $\chi \sim (T - T_c)^{\gamma/\nu}$. These exponents have been obtained performing a finite size scaling [33]. Remarkably, it was found that they are consistent with an order-disorder second order transition within the Ising universality class (Table 1). This is all the more remarkable since it occurs in a nonequilibrium system. Similar behavior has been also observed in nonequilibrium coupled maps [34].

5 Differences Between Atrial and Ventricular Cells

The main structural difference between atrial and ventricular cells is that the latter present t-tubules while the former do not. This has an important effect for the dynamics of calcium inside the cell. In ventricular myocytes, LCC and RyR2 are always confronted and the rise in calcium occurs in the bulk of the cell. In atrial cells, on the contrary, calcium concentration increases first close to the cell membrane and then propagates into the cell as a wave. As we have discussed in the previous section, in ventricular myocytes global alternans appears as a synchronization effect among local alternans in different microdomains. In atrial myocytes, calcium alternans may have fundamentally the same origin as in ventricle, with a coordination in the CaRU units close to the membrane then propagated through calcium waves of higher or lower amplitude. However, a new scenario is also possible: CaRUs close to the mem-

brane might not present alternans but the wave propagation towards the center might be completely or partially blocked in alternative beats. This new mechanism for alternans, not possible in ventricles, requires the analysis of wave-like calcium propagation in the cell. Particularly, the presence or absence of t-tubules has effects on the total exchanger strength and the homeostatic balance in the cell. This, together with buffering, might produce different loading of the SR in alternative beats leading to different wave propagation properties in the atrial cell at consecutive beats.

Even if this new mechanism can not be ruled out, models seem to indicate that some type of coordination at the surface level will always be, at least partially, responsible for alternans in atria. In whole-cell models, where wave-like events can not be studied, one can test whether important differences in loading of SR appear whenever there are large differences in the calcium transient. If after any given beat the total content of calcium in the cell is roughly the same, the SERCA pump will generate basically a constant calcium load at each beat. Under the same calcium load it seems unlikely to have very different wave-propagation dynamics. On the other hand, if the balance of the NaCa exchanger and the I_{CaL} current is extremely sensitive to the calcium transient, consecutive beats may have very different calcium loads. This could lead to different wave-propagation, producing a different calcium transient which results in a different SR calcium load.

In a model of human atrial cell without t-tubules it was observed that the total number of calcium ions transported across the membrane was roughly the same every beat [29] for very different calcium transients (Fig. 10). In this case the transients were different because RyR2 refractoriness was behind the presence of coordinated global alternans. It was also found that SR Ca content oscillations during alternans, even if SR Ca fluctuations are not an essential ingredient for it, seemed more common in a ventricle model [28] than in the atrial model. For instance, in Fig. 10c the alternans shown in ventricle is caused by RyR2 refractoriness but, nevertheless, results in a strong SR Ca content fluctuation. On the contrary, in the atrial model alternans appears first in the intracellular calcium transient, whereas the junctional SR Ca release presents very small beat-to-beat oscillations. Hence, at the subsarcolemmal space the oscillations are typically smaller than those found in ventricular cells and mostly due to the beat-to-beat change in concentration gradient with the interior of the cell. In this situation, the calcium transport across the cell membrane (Fig. 10a) seem to suggest that wave-like scenario in atria should be difficult to observe when alternans is due to changes in presystolic calcium load, although we cannot discard the presence of wave alternans due to alternations in presystolic RyR availability.

It is important to notice that the beat-to-beat difference in the amount of available calcium within the cell is not only dependent on the general structure of the compartments but also on the properties of the exchanger and LCC currents, not to mention that SERCA must be fast enough to refill the SR. Therefore, one can not rule out that different species with different characteristics in the transmembrane proteins might present very different behavior. In any case, our results clearly hint at a very important role of transmembrane currents, and homeostasis more generally, in determining the possibility of alternans in atrial cells due to wave-like alternation.

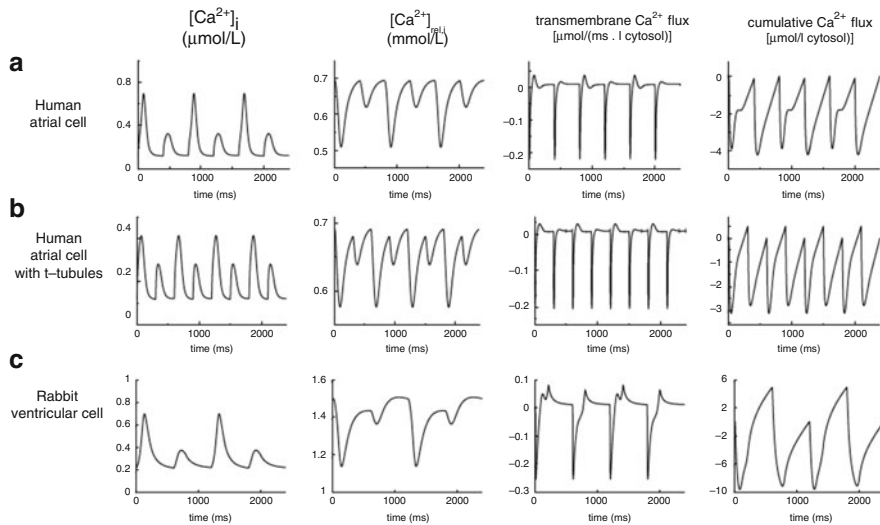


Fig. 10 Comparison between alternans in atrial, ventricular, and atrial cells with t-tubules, using (a) and (b) a human atrial model [29] and (c) a rabbit ventricular cell model [28]. (Reproduced with permission from C.A. Lugo et al, *Am j Physiol* 306, H1540, 2014 [29])

6 Conclusions

Dysfunctions in calcium handling are behind a whole range of cardiac arrhythmogenic behaviors. A good understanding of calcium dynamics is thus of crucial importance to develop ways to prevent these arrhythmias. In this chapter we have presented recent modeling efforts on this line, with the study of the mechanisms behind electromechanical alternans, both in atrial and ventricular myocytes. Even if much has been learned about the origin of alternans, many aspects of calcium dynamics are still not completely understood.

The main open problem, in our understanding, is how different species present different characteristic calcium dynamics under fast beating. Some species, like rabbits, have a strong tendency to generate alternans as the first instability as the heart-beat rate is increased, while others present wave-like phenomena. We have seen that currents which regulate calcium homeostasis are very important in fixing the response of SR load to different calcium transients. It seems that understanding homeostasis must be a key future line of work, not only on the effect of the Na-Ca exchanger and the ICaL, but also how different RyR2 models interact with the homeostatic function of membrane currents. In this regard, understanding whether it is possible to generate an in-silico myocyte model where both calcium alternans and wave-like arrhythmias appear at different heart beat rates will clarify strongly the main driving forces behind one behavior or the other. If this is possible, it will indicate that regulatory mechanisms and species specific properties produce the

different outcomes. If not, it can point to fundamental structural differences in the RyR2 functioning.

In any case, one can not disregard the possibility that the different prevalence of calcium dysfunction, alternans or wave-like off-beat responses, might be due to the different properties of RyR2 clustering together with the influence of beta-adrenergic stimulation. Models with large differences in intracellular clustering size and rate probabilities should be investigated. In this line of work, analyzing the interaction of RyR2 clusters in different network configuration can shed light on its relevance in the origin and maintenance of calcium waves; this is particularly relevant for the case of wave initiation in species where it appears as the main calcium dysfunction. Regarding this point, we have already addressed in this review whether atrial and ventricular myocytes might present different mechanisms for calcium alternans in case atria cells presents wave-blockage as an alternate mechanism.

Finally, although two possible mechanisms, SR load and RyR2 refractoriness, have been uncovered for the onset of calcium alternans, recent experiments [13] have shown a hierarchic tendency in their appearance in healthy heart, beginning by RyR2 refractoriness, following or not by SR load mechanism and finally showing as APD alternans. Pathological conditions can provide proarrhythmic substrates that could sustain or alter this hierarchy. Besides that, the two different local nonlinear mechanism at the CaRU presented here (calcium SR load or RyR2 refractoriness) can be present in different degrees in different species. So, future experimental work can clarify the relevance of the different mechanisms under different scenarios and circumstances. Our believe is that mathematical models of calcium dynamics will play a very important role in understanding the relevances and differences of these mechanisms and their physiological and clinical implications. Hopefully, it will allow new targets for pharmacological treatment of arrhythmias.

Acknowledgements B.E., I.R.C., A.P. and E.A-L. would like to acknowledge financial support from MICINN (Spain) through coordinated projects FIS2011-28820-C02-01 and SAF2014-58286-C2-2-R. We would like to thank L. Hove-Madsen, Y. Shiferaw and C. Lugo for fruitful discussions.

References

1. Bers, D.M.: Cardiac excitation–contraction coupling. *Nature* **415**(6868), 198–205 (2002)
2. Wagner, S., Maier, L.S., Bers, D.M.: Role of sodium and calcium dysregulation in tachyarrhythmias in sudden cardiac death. *Circ. Res.* **116**, 1956–1970 (2015)
3. Antzelevitch, C., Sicouri, S.: Clinical relevance of cardiac arrhythmias generated by afterdepolarizations. Role of M cells in the generation of u waves, triggered activity and torsade de pointes. *J. Am. Coll. Cardiol.* **23**(1), 259–277 (1994)
4. Shiferaw, Y., Aistrup, G.L., Wasserstrom, J.A.: Intracellular Ca²⁺ waves, afterdepolarizations, and triggered arrhythmias. *Cardiovasc. Res.* **95**(3), 265–268 (2012)
5. Xie, Y., Sato, D., GarfiAnkel, A., Qu, Z., Weiss, J.N.: So little source, so much sink: requirements for afterdepolarizations to propagate in tissue. *Biophys J.* **99**(5), 1408–1415 (2010)

6. Alvarez-Lacalle, E., Peñaranda, A., Cantalapiedra, I.R., Hove-Madsen, L., Echebarria, B.: Effect of RyR2 refractoriness and hypercalcemia on calcium overload, spontaneous release, and calcium alternans. *Comput. Cardiol.* **40**, 683–686 (2013)
7. Laurita, K.R., Rosenbaum, D.S.: Cellular mechanisms of arrhythmogenic cardiac alternans. *Prog. Biophys. Mol. Biol.* **97**(2–3), 332–347 (2008)
8. Weiss, J.N., Karma, A., Shiferaw, Y., Chen, P.S., Garfinkel, A., Qu, Z.: From pulsus to pulseless: the saga of cardiac alternans. *Circ. Res.* **98**(10), 1244–1253 (2006)
9. Pastore, J.M., Girouard, S.D., Laurita, K.R., Akar, F.G., Rosenbaum, D.S.: Mechanism linking T-wave alternans to the genesis of cardiac fibrillation. *Circulation* **99**(10), 1385–1394 (1999)
10. Karma, A.: Electrical alternans and spiral wave breakup in cardiac tissue. *Chaos Interdiscip. J. Nonlinear Sci.* **4**(3), 461–472 (1994)
11. Echebarria, B., Karma, A.: Mechanisms for initiation of cardiac discordant alternans. *Eur. Phys. J. Spec. Top.* **146**(1), 217–231 (2007)
12. Chudin, E., Goldhaber, J., Garfinkel, A., Weiss, J., Kogan, B.: Intracellular Ca²⁺ dynamics and the stability of ventricular tachycardia. *Biophys. J.* **77**(6), 2930–2941 (1999)
13. Wang, L., Myles, R.C., De Jesus, N.M., Ohlendorf, A.K., Bers, D.M., Ripplinger, C.M.: Optical mapping of sarcoplasmic reticulum Ca²⁺ in the intact heart ryanodine receptor refractoriness during alternans and fibrillation. *Circ. Res.* **114**(9), 1410–1421 (2014)
14. Narayan, S.M., Franz, M.R., Clopton, P., Pruvot, E.J., Krummen, D.E.: Repolarization alternans reveals vulnerability to human atrial fibrillation. *Circulation* **123**(25), 2922–2930 (2011)
15. Richards, M., Clarke, J., Saravanan, P., Voigt, N., Dobrev, D., Eisner, D., Trafford, A., Dibb, K.: Transverse tubules are a common feature in large mammalian atrial myocytes including human. *Am. J. Physiol. Heart Circ. Physiol.* **301**(5), H1996–2005 (2011)
16. Peñaranda, A., Alvarez-Lacalle, E., Cantalapiedra, I.R., Echebarria, B.: Nonlinearities due to Refractoriness in SR Ca Release. *Comput. Cardiol.* **39**, 297–300 (2012)
17. Shkryl, V.M., Maxwell, J.T., Domeier, T.L., Blatter, L.A.: Refractoriness of sarcoplasmic reticulum Ca²⁺ release determines Ca²⁺ alternans in atrial myocytes. *Am. J. Physiol. Heart Circ. Physiol.* **302**(11), H2310–H2320 (2012)
18. Hüser, J., Wang, Y.G., Sheehan, K.A., Cifuentes, F., Lipsius, S.L., Blatter, L.A.: Functional coupling between glycolysis and excitation-contraction coupling underlies alternans in cat heart cells. *J. Physiol.* **524**(3), 795–806 (2000)
19. Llach, A., Molina, C.E., Fernandes, J., Padro, J., Cinca, J., Hove-Madsen, L.: Sarcoplasmic reticulum and L-type Ca channel activity regulate the beat-to-beat stability of calcium handling in human atrial myocytes. *J. Physiol.* **589**, 3247–3262 (2011)
20. Díaz, M.E., O'Neill, S.C., Eisner, D.A.: Sarcoplasmic reticulum calcium content fluctuation is the key to cardiac alternans. *Circ. Res.* **94**(5), 650–656 (2004)
21. Shiferaw, Y., Watanabe, M.A., Garfinkel, A., Weiss, J.N., Karma, A.: Model of intracellular calcium cycling in ventricular myocytes. *Biophys. J.* **85**(6), 3666–3686 (2003)
22. Restrepo, J.G., Weiss, J.N., Karma, A.: Calsequestrin-mediated mechanism for cellular calcium transient alternans. *Biophys. J.* **95**(8), 3767–3789 (2008)
23. Picht, E., DeSantiago, J., Blatter, L.A., Bers, D.M.: Cardiac alternans do not rely on diastolic sarcoplasmic reticulum calcium content fluctuations. *Circ. Res.* **99**(7), 740–748 (2006)
24. Sobie, E.A., Song, L.S., Lederer, W.J.: Restitution of Ca²⁺ release and vulnerability to arrhythmias. *J. Cardiovasc. Electrophysiol.* **17**(Suppl 1), S64–S70 (2006)
25. Fill, M., Copello, J.A.: Ryanodine receptor calcium release channels. *Physiol. Rev.* **82**(4), 893–922 (2002)
26. Chang, K.C., Bayer, J.D., Trayanova, N.A.: Disrupted calcium release as a Mechanism for atrial alternans associated with human atrial fibrillation. *PLoS Comput. Biol.* **10**(12), e1004011 (2014)
27. Alvarez-Lacalle, E., Cantalapiedra, I.R., Peñaranda, A., Cinca, J., Hove-Madsen, L., Echebarria, B.: Dependency of calcium alternans on ryanodine receptor refractoriness. *PLoS One* **8**(2), e55042 (2013)

28. Shannon, T.R., Wang, F., Puglisi, J., Weber, C., Bers, D.M.: A mathematical treatment of integrated Ca dynamics within the ventricular myocyte. *Biophys. J.* **87**(5), 3351–3371 (2004)
29. Lugo, C.A., Cantalapiedra, I.R., Peñaranda, A., Hove-Madsen, L., Echebarria, B.: Are SR Ca content fluctuations or SR refractoriness the key to atrial cardiac alternans?: insights from a human atrial model. *Am. J. Physiol. Heart Circ. Physiol.* **306**(11), H1540–H1552 (2014)
30. Rovetti, R., Cui, X., Garfinkel, A., Weiss, J.N., Qu, Z.: Spark-induced sparks as a mechanism of intracellular calcium alternans in cardiac myocytes. *Circ. Res.* **106**(10), 1582–1591 (2010)
31. Restrepo, J.G., Karma, A.: Spatiotemporal intracellular calcium dynamics during cardiac alternans. *Chaos Interdiscip. J. Nonlinear Sci.* **19**, 037115 (2009)
32. Cantalapiedra, I.R., Lugo, C.A., Peñaranda, A., Echebarria, B.: Calcium alternans produced by increased sarcoplasmic reticulum refractoriness. *Comput. Cardiol.* **38**, 645–648 (2011)
33. Alvarez-Lacalle, E., Echebarria, B., Spalding, J., Shiferaw, Y.: Calcium alternans is due to an order-disorder phase transition in cardiac cells. *Phys. Rev. Lett.* **114**(10), 108101 (2015)
34. Marcq, P., Chaté, H., Manneville, P.: Universality in Ising-like phase transitions of lattices of coupled chaotic maps. *Phys. Rev. E* **55**(3), 2606 (1997)