

Is the coexistence of Catalan and Spanish possible in Catalonia?

Luís F Seoane,^{1,2,3} Xaquín Loredo,⁴ Henrique Monteagudo,^{5,6} and Jorge Mira^{*7}

¹*Instituto de Física Interdisciplinar y Sistemas Complejos IFISC (CSIC-UIB), Campus UIB, 07122, Palma de Mallorca, Spain*

²*ICREA-Complex Systems Lab, Universitat Pompeu Fabra (GRIB), Dr Aiguader 80, 08003 Barcelona, Spain.*

³*Institut de Biologia Evolutiva, CSIC-UPF, Pg Marítim de la Barceloneta 37, 08003 Barcelona, Spain.*

⁴*Seminario de Sociolingüística da Real Academia Galega, Santiago de Compostela, Spain.*

⁵*Real Academia Galega and Instituto da Lingua Galega.*

⁶*Universidade de Santiago de Compostela, 15782 Santiago de Compostela, Spain.*

⁷*Departamento de Física Aplicada, Universidade de Santiago de Compostela, 15782 Santiago de Compostela, Spain.*

Supporting information

S1 Appendix. Fitting coupled non-linear differential equations with dummy polynomials.

We deal with a set of coupled, non-linear differential equations with 4 parameters $\{a, c, k, s\}$ and two initial conditions $(x(t = t^0), y(t = t^0))$. We estimate all these six free parameters from the data. The most straightforward way to do this is by minimizing the error function:

$$\epsilon(X, \alpha) = \frac{1}{N} \left(\sum_{i=1}^N \left[(x_i - x_P(t_i))^2 + (b_i - b_P(t_i))^2 + (y_i - y_P(t_i))^2 \right] \right)^{1/2}, \quad (1)$$

where N is the number of data points; $X \equiv \{X_i = (t_i, x_i, b_i, y_i), i = 1, \dots, N\}$ is the empirical data series that includes the times at which each observation is made (t_i) and the fractions of each linguistic group at that time (x_i , b_i , and y_i); and $(x_P(t), b_P(t), y_P(t))$ are the result of evaluating Eqs (1) with parameters $\alpha \equiv \{a, c, k, s; x(t^0), y(t^0)\}$.

To evaluate Eq (1) it is necessary to integrate the model numerically with an appropriate time step, which makes the fitting process very demanding. Advances in model validation often rely on bootstrapping techniques for which it is necessary to perform several stochastic fits upon the same data (Shalizi, 2013). It is unrealistic to apply these methods if we need to integrate Eqs (1) numerically. We decided to sacrifice some accuracy for more powerful statistical tools. To avoid these costly integrations we used an alternative fitting process that we describe below. The results obtained at different stages with this technique are illustrated in 1. To elaborate this figure we used the bilingualism thresholds for which best ($r = 9$, **1a**) and worst ($r = 35$, **1b**) results were obtained.

We started by approximating (t_i, x_i) and (t_i, y_i) data to a pair of dummy polynomials $P_{x,y}(t)$ controlling for overfitting with 5-fold cross-validation (**1a1** and **1b1**). Then we extracted the time derivatives $dP_{x,y}/dt \equiv \dot{P}_{x,y}(t)$ which are a function of time alone. These derivatives were compared to a direct evaluation of the right hand side of Eqs (1) given a set of parameters (α). We introduce the error function:

$$\epsilon_{\partial}(X, P_{x,y}; \alpha) = \frac{1}{N} \left(\sum_i \left[(\dot{P}_x(t_i) - \dot{x}(X_i; \alpha))^2 + (\dot{P}_y(t_i) - \dot{y}(X_i; \alpha))^2 \right] \right)^{1/2}. \quad (2)$$

This measures the deviation of the estimated derivatives of the variables, instead of the deviation of the model with respect to the actual data.

We do not need to integrate Eqs (1) to perform this analysis, so we could implement several ($N_f = 100$) fits for each data set. Different combinations of parameters might render equally good fits, so we aimed at capturing and characterizing several local optima in parameter space. For each of the N_f fits, the parameters were random and uniformly initialized in the intervals: $a \in [1, 2]$, $c \in [0, 2]$, $k \in [0, 1]$, and $s \in [0, 1]$; but they were left to evolve

* jorge.mira@usc.es

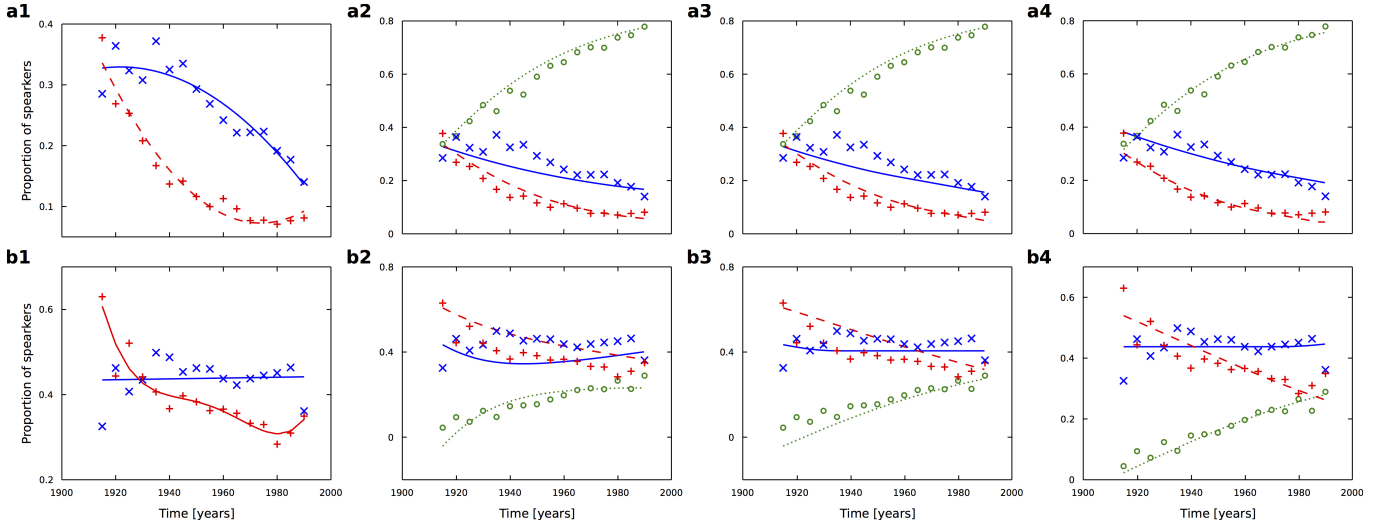


FIG. 1 **Different stages of the fitting method for the best and worst cases.** The best **a** ($r = 9$) and worst **b** ($r = 35$) fits obtained are used to exemplify the different stages of our regression technique and how they improve in explaining the original data. **a1** and **b1** Dummy polynomials are adjusted to the data using five-fold cross-validation. **a2** and **b2** Eq (2) is used as an approximation to the actual cost function (Eq (1)). It measures the difference between the right hand side of Eqs (1) and the derivatives of the dummy polynomials. The fit that minimizes this proxy cost is shown. **a3** and **b3** Good fits to the dummy polynomials are integrated with arbitrary initial conditions so that the actual residuals (Eq (1)) can be measured. We show the best fit with respect to this (actual) cost function. For **a**, the parameters minimizing both costs happen to be the same, so panels **a2** and **a3** are identical. For **b** we can observe some improvement. Most other cases (other values of r and different data sets) also improve after this step. **a4** and **b4** Fitting the derivatives of the model leaves the initial conditions undetermined. Finding better initial conditions is the last operation of our fitting procedure. In some cases it can greatly improve the results.

outside these boundaries as a greedy gradient descent minimized the cost function in Eq (2). Runs that ended up with $k, s \notin [0, 1]$ were discarded since they imply negative numbers of speakers. For each dataset we recovered a collection $\{\alpha_j, j = 1, \dots, N_v < N_f\}$, with each α_j a valid set of parameters.

The fits produced by the best parameters at this stage are shown in **1a2** and **b2**. These curves do not match the dummy polynomials – often they disagree notably. This is so because the curves in **1a2** and **b2** are constrained to belong to the family of curves that the model is capable of producing. These are defined by Eqs (1) and are usually less versatile than even low degree polynomials (note also that the low degree polynomials are not coupled with each other). This justifies that the cross-validation step is introduced to obtain $P_{x/y}(t)$. We could spare the cross-validation for a later stage, but this would gain us little in terms of avoiding overfitting and every computation would be repeated as many times as folds in our cross-validation.

The operations described to far gained us some speed but we have lost accuracy, so we re-evaluated all the valid sets of parameters (all α_j) derived for each data series by numerically integrating Eqs (1). Note that it is not as costly to integrate those equations a limited number (N_v) of times as it would be to perform a greedy algorithm – e.g. a genetic algorithm or simulated annealing – with original cost function, which would demand a much larger number of integrations. The best parameters after this step are shown in **1a3** and **b3**. Note that for $r = 9$ both Eqs (1) and (2) render the same best fit – hence **1a2** and **1a3** are identical. For $r = 35$ and virtually every other choice of bilingualism threshold we do appreciate some improvement though.

Because this fitting procedure relies on the derivatives of the time series, it does not determine the best initial conditions. For **1a1-3** and **1b1-3**, the initial conditions are extracted from the polynomial fit. This is just one valid option, but there are reasons why we would like to explore alternatives. First, these initial conditions are highly influenced by the first points in the data series. These are extracted from the groups with more age from the original surveys (Generalitat de Catalunya, 2008a,b, 2013), which may include larger biases than other groups. Secondly, and bearing on this, non-linear differential equations may be very sensitive to their boundary conditions. Altogether it seems worth it to loose yet another degree of freedom to obtain a little bit more accuracy. Hence, to further refine our results we tried each one of the empirical data points in each time series and took as initial conditions those that minimized the cost function in Eq (1). This is the last step in our fitting technique, so **1a4** and **b4** show the best results for $r = 9$ and $r = 35$ respectively.

Finally, we used this cost function to assign a weight to each one of the different sets of parameters (α_j) that

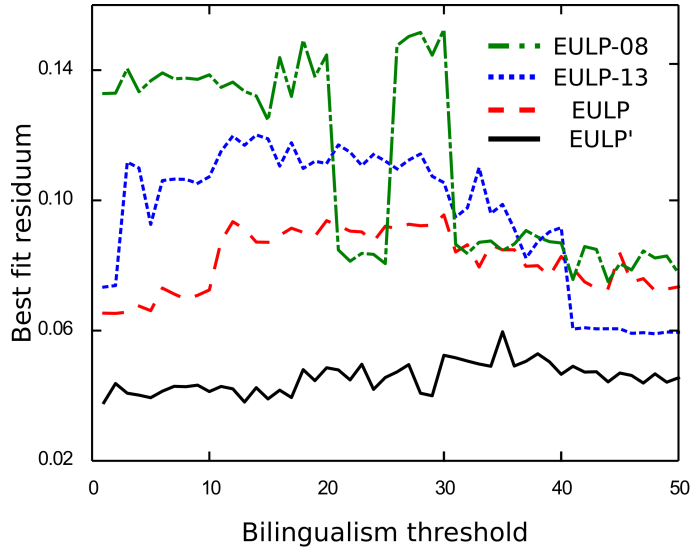


FIG. 2 **Comparing regressions for the different datasets.** The residuum normalized by the number of data points is shown for all different datasets available and for integer $r \in [1, 50]$. The EULP' time series allows the best parsimonious description of the data using our model Eqs (1). This is so because this dataset aggregates a larger population, hence yielding a better signal to noise ratio.

resulted from the analysis:

$$p(\alpha_j) = \frac{1}{Z} e^{-\epsilon(X, \alpha_j)}, \quad (3)$$

with $Z = \sum_j \exp(-\epsilon(X, \alpha_j))$. These weights were used to produce averages across all α_j . They confer more importance to parameters with small residua and minimize the importance of bad fits.

S2 Appendix. Exploring the different datasets.

Two relevant datasets (dubbed *EULP-2008* and *EULP-2013*) are derived from the original surveys (Generalitat de Catalunya, 2008a,b, 2013). These two datasets were combined in two different ways: i) aggregating all data, including those time periods for which there is data in one survey but not in the other; and ii) aggregating only the data that overlapped in time. We term these new series respectively *EULP* and *EULP'*. For all datasets we implemented the fitting procedure described in the previous appendix except for the last step – i.e. we did not correct for the initial conditions yet. This was relatively costly and we did not estimate it necessary for the more noisy time series.

At this point, we compared the average residuum per data point obtained for the best fit of each dataset. This is shown in 2 as a function of $r \in [1, 50]$. *EULP'* is the dataset for which the model Eqs (1) work better. This dataset aggregates more population so we naturally expect less fluctuations. As argued in (Seoane and Mira, 2017), we hope that mathematical theories of social systems perform better as our observations are averaged over larger populations – just like thermodynamics works so well because of the astronomical number of particles involved. Hence our preference for datasets with larger samples. Actually, *EULP* does consider a larger population than *EULP'*, but this is dispersed into relatively ill-represented groups that correspond to the extreme time periods.

For the best data set (*EULP'*) we also found the best initial conditions as described in the previous appendix. These are the results that the body of the paper is based on. Also, as reported above, this dataset was also processed holding $a = 1.31$ fixed. Finally, *EULP'* data was spatially separated into speakers that live in Barcelona or its metropolitan area versus speakers living in the rest of Catalonia. For these extra datasets (dubbed *EULP-BCN* and *EULP-notBCN*) the whole fitting procedure was carried out but always with $a = 1.31$. These results are also discussed in the main text.

Together with this paper we attach a collection of figures showing the best fit attained for each one of the datasets and for each integer value of $r \in [1, 50]$. Note that for *EULP-2008*, *EULP-2013*, and *EULP* these fits were not corrected to the best initial conditions; while for *EULP'*, *EULP-BCN*, and *EULP-notBCN* we have the best fits possible (in the case of *EULP'*, this is so both when all parameters are considered and when $a = 1.31$ is held fixed). This collection of figures can be explored to have an idea of the difference that correcting for the initial conditions often supposes.

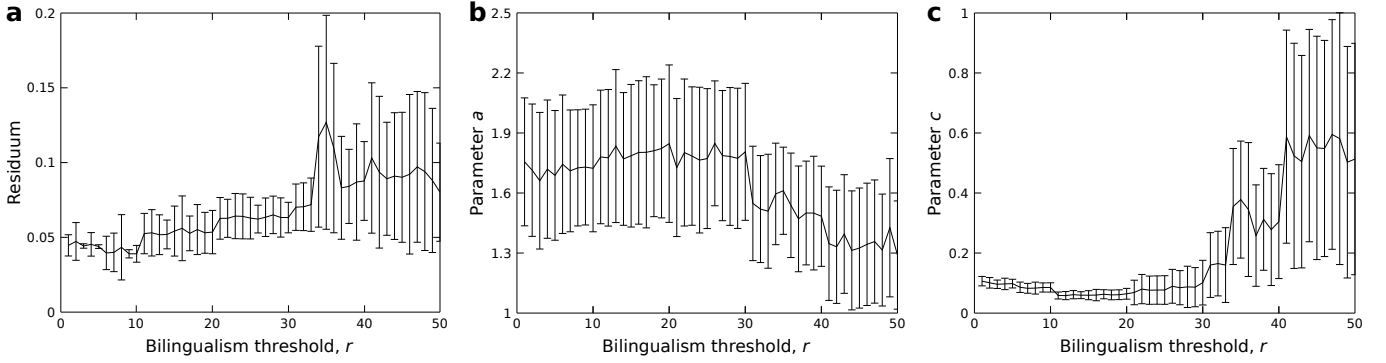


FIG. 3 Overview of the fit results for the less relevant parameters. **a** The residuum (as given by Eq (1)) measures the quality achieved by our fitting techniques. It reveals a low deviation per data point, especially for low values of r , for which most speakers score as bilinguals. For higher (more restrictive) values of r , we find a large difference in the residua of different realizations of our fitting process. This indicates a less consistent convergence and raises a warning about the quality of the parameters estimated in that regime, since some of the fits have a low quality. We attempt to correct for this by weighting the different results according to the residuum as indicated in the text ((3)). We refer to parameters a (panel **b**) and c (panel **c**) as less relevant because they have a lesser impact on the stability of the model. **b** The parameter a is larger than 1 in average for any r . It declines slightly for $r > 30$. It also presents a large variability (meaning that different realizations of our fitting procedure with the same r yield different values of a compatible with the data. This variability subtracts importance from the decline observed for larger r . **c** The parameter c is low and consistent for fits with $r \lesssim 30$. This parameter is related to how fast the dynamics unfold, so they are relatively slow in this regime. For $r > 30$ the dynamics get slightly accelerated, but also the fits are less consistent. This indicates, again, that dynamics with different speeds are similarly compatible with the data.

For EULP-2008, EULP-2013, and EULP we also include figures showing the average and standard deviation of a , c , k , and s as a function of $r \in [1, 50]$; as well as a figure showing average and standard deviations in the $k - s$ plane. For EULP-BCN and EULP-notBCN we include the corresponding plots for c , k , and s as a function of r (a is constant, so not worth showing, and the $k - s$ plots have already been shown in the main text). All these plots are consistent with the discussed results.

S3 Appendix. Overview of the lesser parameters.

3 reflects how the residuum of the fits and the less relevant parameters (a and c) change as we vary our definition of bilingualism by varying $r \in [1, 50]$. The mean and standard deviation are represented. They indicate that the model is a better fit for lower values of r . The average and standard deviation of the model parameters (a and c in 3b and c, and k and s to be discussed later) have been calculated weighting the results of each fit (α_j) with Eq (3), which assigns more importance to good fits and penalizes fits that deviate from the real data.

The parameter a appears broadly constant in average, but always with a large variation (3b). A consistent decline is appreciated for $r > 30$. This parameter has been discussed as a proxy for *volatility*, or a tendency of linguistic groups to dissolve easily (the lower a , the larger volatility) (Castelló et al., 2013). 3b shows how more stringent definitions of bilingualism require slightly more volatile dynamics to account for the real data. For $a < 1$ all parameter combinations imply the stable coexistence of both languages (Otero-Espinar et al., 2013), but for $a \geq 1$ all the dynamics accounted for by the model (i.e. both extinction and coexistence) are possible depending on k , s , and the initial conditions. Since our data show $a > 1$ in average always, more relevant information will be given by k and s .

The other lesser parameter (c) is a normalization constant related to the speed at which the dynamics are resolved. It is consistently low for $r < 30$, indicating relatively slower dynamics; and it increases abruptly after that. Faster dynamics are intuitively consistent with more volatile regimes. This increase in c is also accompanied by less consistent fits – i.e. very different values of c are obtained for the same data, as indicated by the error bars in 3c. (Note that, meanwhile, the deviation around $\langle a \rangle$ is fairly constant throughout $r \in [1, 50]$). The parameter c does not affect the stability of the system.

S4 Appendix. Counter-intuitive dynamics in the spatially segregated systems.

The dynamics reported for segregated data in Barcelona and its metropolitan area versus the rest of Catalonia can result counter-intuitive. We append two figures for clarification.

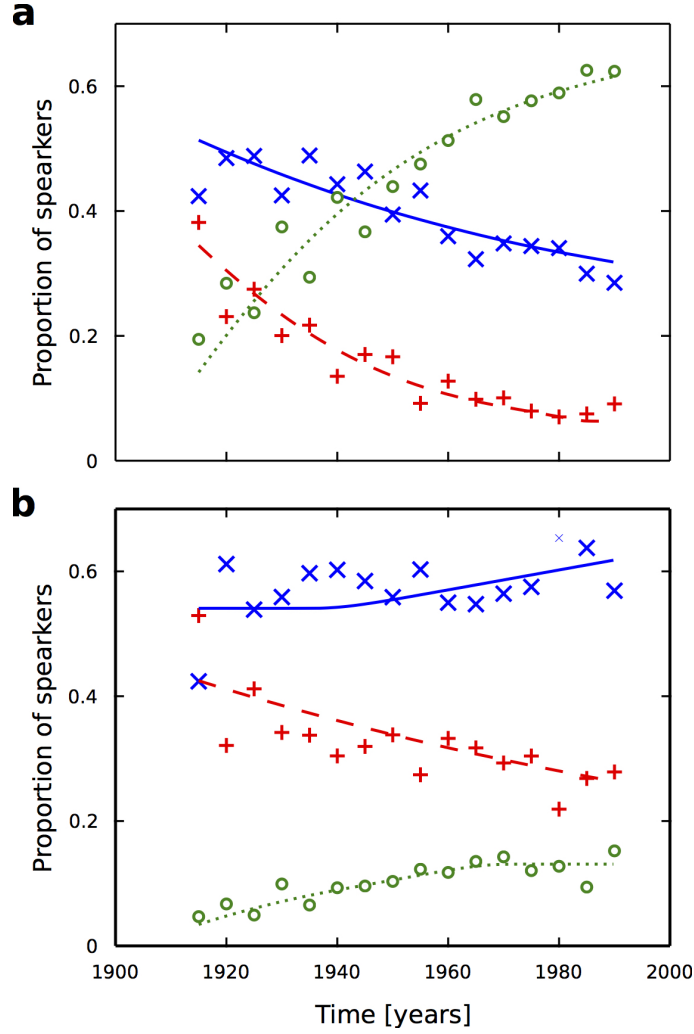


FIG. 4 **Informative fits in Barcelona and its metropolitan area.** **a** $r = 15$ for which coexistence of both languages is predicted in the long term, and **b** $r = 50$ whose parameters k and s imply the extinction of one of the languages – with projections up to the end of the century often compatible with the extinction of Catalan.

4a shows the best fit for Barcelona and its metropolitan area for $r = 15$. We appreciate a moderate and roughly parallel (while slightly steeper for Catalan) decay of both monolingual groups with a sustained grow of the bilingual group. The symmetry of the dynamics indicates a balanced prestige, as consistently found for this region ($s_S \sim 0.51$, $s_C \sim 0.49$), regardless of the historical majority of Spanish speakers. This majority, though, would be enough to strongly favor Spanish in the long term if the competition were harshened by a narrow definition of bilingualism (4b). For $r = 50$, the results of the fits are very often compatible with the extinction of Catalan in Barcelona within this century. Thinking towards the preservation of Catalan, strategies should seek to widen the perception of bilingualism because those are the scenarios in which both languages are more likely to subsist.

The most challenging situation to interpret is presented in the rest of Catalonia, where Catalan was historically the most spoken language but whose situation is currently turning over. 5a shows a very steep decline of Catalan language when the definition of bilingualism is relatively broad (again $r = 15$). Meanwhile, the number of Spanish speakers is more or less sustained over time. This is compatible with a low prestige parameter for Catalan ($s_C \sim 0.4$, the lowest across all datasets, while $s_S \sim 0.6$) even if the number of Catalan speakers remains larger than the number of Spanish speakers nowadays. Thinking carefully about it, a small population being able to cause such a large drop of speakers is a group with a language that is perceived as very *attractive*, and this is indeed what the prestige parameter of this model captures.

Notwithstanding this advantage for Spanish (as perceived by the Catalan population itself); would the dynamics become harshly competitive, the Spanish would suffer the most because the current number of monolingual Catalan speakers is enough to favor this later tongue asymptotically. 5b still shows a growing Spanish monolingual group and

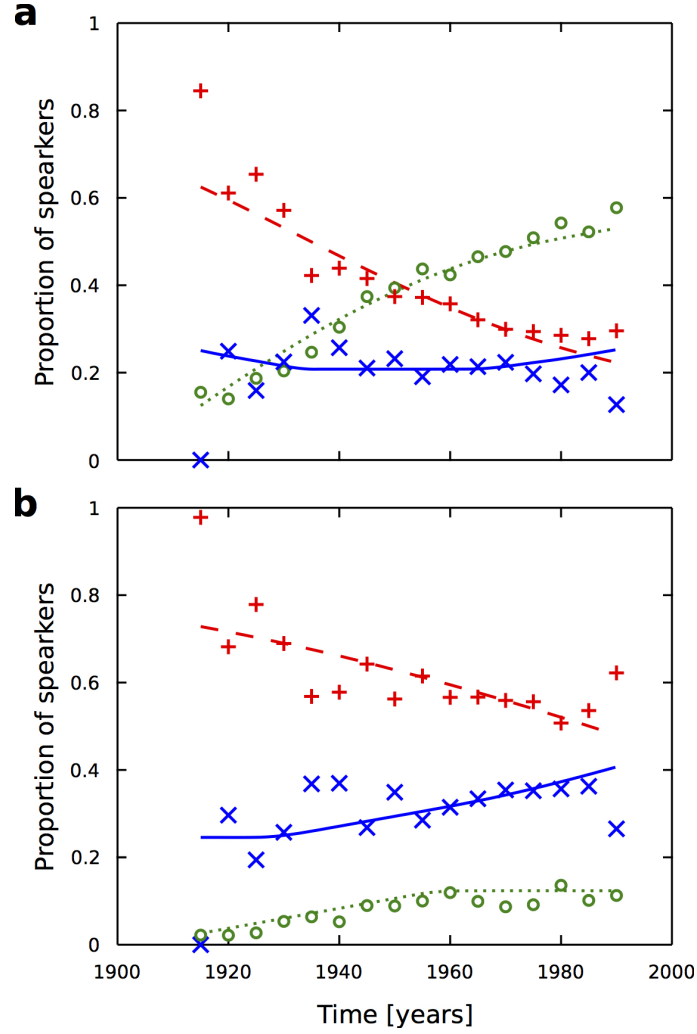


FIG. 5 **Informative fits in the rest of Catalonia.** **a** $r = 15$ for which coexistence of both languages is again predicted in the long term, and **b** $r = 50$ for which, again, the k and s parameters imply the extinction of one of the languages – with Spanish the more likely tongue to go extinct towards the end of the century, should one of them disappear.

a seemingly declining Catalan monolingual population. Averaging over many simulations shows that most of them are compatible with the coexistence of both languages with similar numbers of speakers towards the end of this century – while some of them predict the extinction of Spanish. Such dynamics are already hinted at in the predictions towards 2030 (7).

Again, the preferred strategy for language coexistence would be to widen the perception of bilingualism in the society. In this case, that option would favor the preservation of Spanish. Interestingly, the best strategies for Catalan persistence alone (hence disregarding the survival of Spanish) would be opposite in the Barcelona area (where coexistence is the safest way) and the rest of Catalonia (where betting for extinction dynamics has a large chance of favoring Catalan over Spanish).

S5 Appendix. Meaning of the model and model parameters.

The current model has been extensively discussed in the literature and by different authors, with some papers offering a broad overview (often discussing also other models) and others focusing on specific parameters (Castelló et al., 2013; Juane et al., 2019; Mira and Paredes, 2005; Mira et al., 2011; Otero-Espinar et al., 2013; Seoane and Mira, 2017). A brief discussion follows aimed at clarifying what the equations mean in the real world, and exactly how each of the parameters of the model acts. This appendix has been written keeping in mind those readers that might be completely alien to the mathematical modeling of these or other systems (not necessarily social ones), with the hope

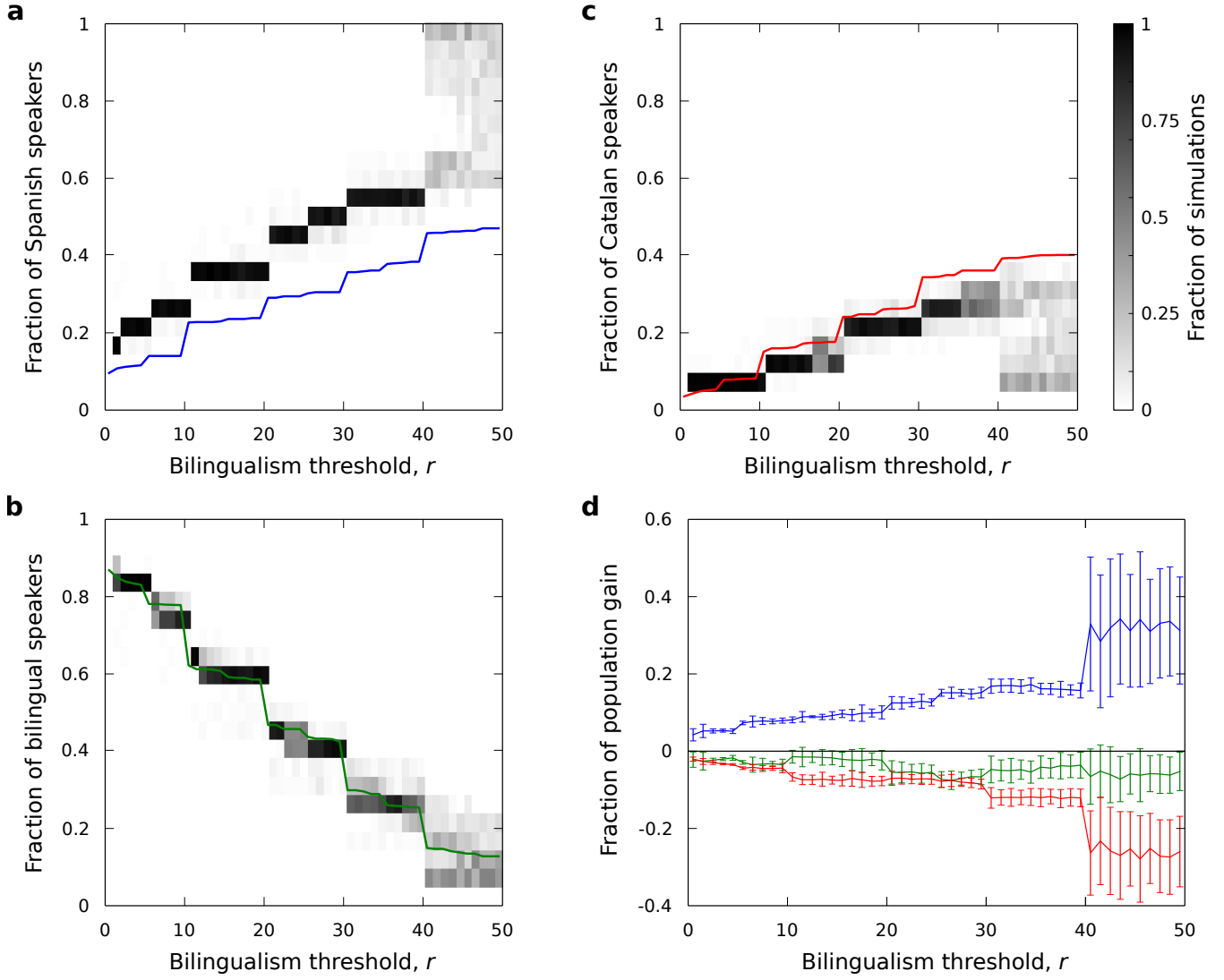


FIG. 6 **Year-2030 model projections, Barcelona and metropolitan area.** Projections have been weighted and binned as for the aggregated data in Fig 3 (main text) for Spanish (a), bilingual (b), and Catalan (c) speakers, with blue, green, and red lines indicating as a reference the fractions of speakers in 1990, the last data point in the EULP' series. d Gains and losses of speakers for each of the groups.

of building bridges between fields that are often far apart. However, we have found that it is also often enlightening to dissect the assumptions behind the model with the detail that we do here, often pointing as at the hypotheses that might more readily fail if our equations fail to reproduce the data.

The model builds upon the assumption that the whole population can be separated into three groups: X , Y , and B (Supporting Fig 8). We further assume that the total population is constant all the time. The ratio of people of each group to the total number of people across all groups is noted by the variables x , y , and b ; and these numbers must add up to 1 ($x + y + b = 1$). Departing from here, we can write down a very general model of how people in our population move from one group to another. Let us call P_{XY} to the probability that, within an amount of time $c\Delta t$, each person in X leaves this group and moves to group Y for whatever reasons. Then, the relative amount of people moving from X to Y is the relative amount of people in X multiplied by the probability that each of these people makes this change: xP_{XY} . In other words, this is the flux of people between groups X and Y . We can similarly write down fluxes starting (and finishing) in each other group. Thus, the increase, within an amount of time $c\Delta t$, of the relative number of people in group X will be the sum of all people that enter this group minus all the people that

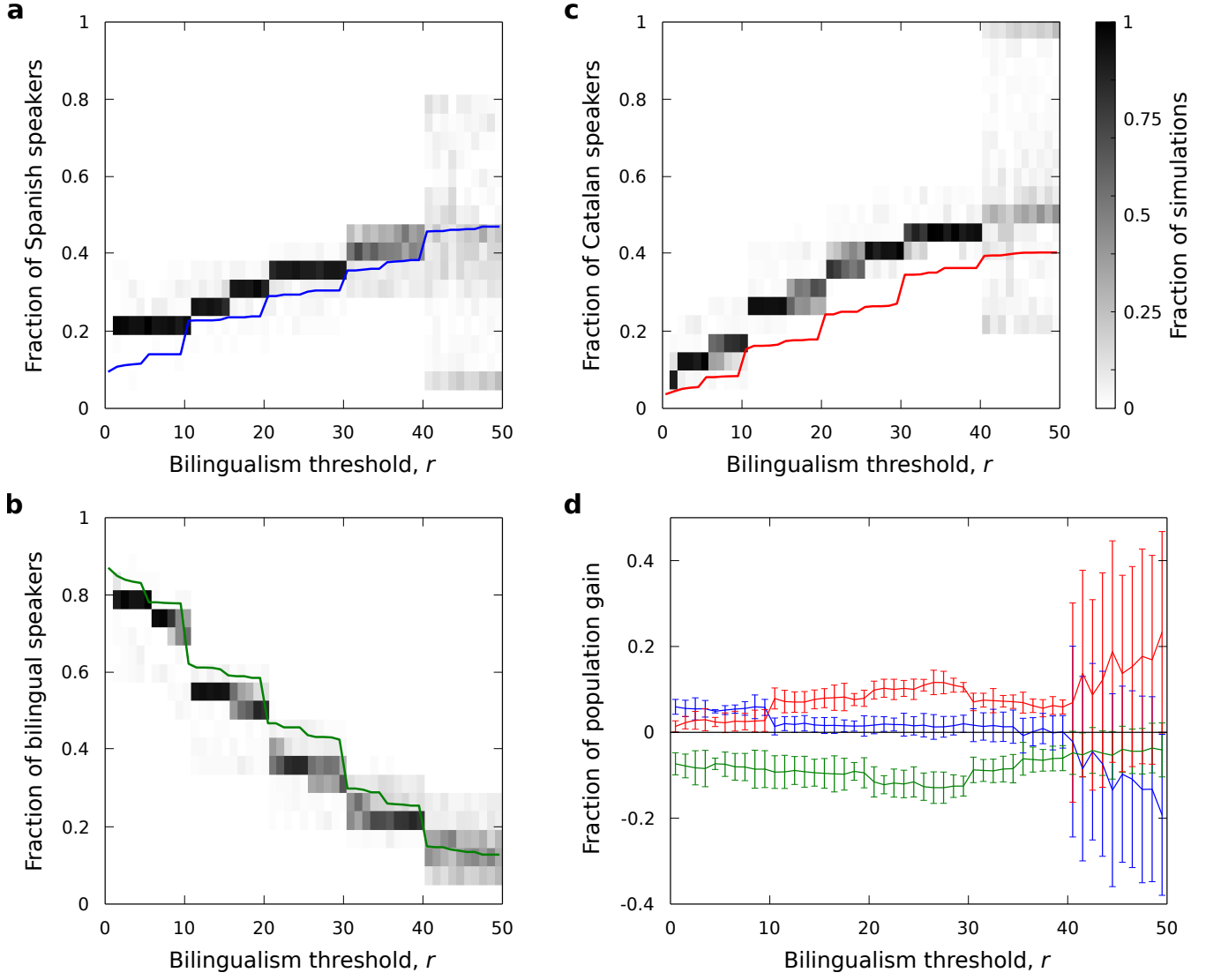


FIG. 7 **Year-2030 model projections, outside Barcelona and its metropolitan area.** Projections have been weighted and binned as for the aggregated data in Fig 3 (main text) for Spanish (a), bilingual (b), and Catalan (c) speakers, with blue, green, and red lines indicating as a reference the fractions of speakers in 1990, the last data point in the EULP' series. d Gains and losses of speakers for each of the groups.

leave this group:

$$\begin{aligned}\Delta x &= (yP_{YX} + bP_{BX} - xP_{XY} - xP_{XB}) \cdot c\Delta t = \\ &= [yP_{YX} + bP_{BX} - x(P_{XY} + P_{XB})] \cdot c\Delta t.\end{aligned}\quad (4)$$

We can write down similar equations for the increase of people in each of the other two groups:

$$\begin{aligned}\Delta y &= (xP_{XY} + bP_{BY} - yP_{YX} - yP_{YB}) \cdot c\Delta t = \\ &= [xP_{XY} + bP_{BY} - y(P_{YX} + P_{YB})] \cdot c\Delta t.\end{aligned}\quad (5)$$

$$\begin{aligned}\Delta b &= (yP_{YB} + xP_{XB} - bP_{BX} - bP_{BY}) \cdot c\Delta t = \\ &= [yP_{YB} + xP_{XB} - b(P_{BX} + P_{BY})] \cdot c\Delta t.\end{aligned}\quad (6)$$

Note that these increases can be negative (i.e. decreases) depending on a number of factors.

Note also that in our simple model no person leaves the population. Whoever leaves one of the three groups has to enter some other. Vice versa: nobody enters the population, so whoever comes into one group must come from

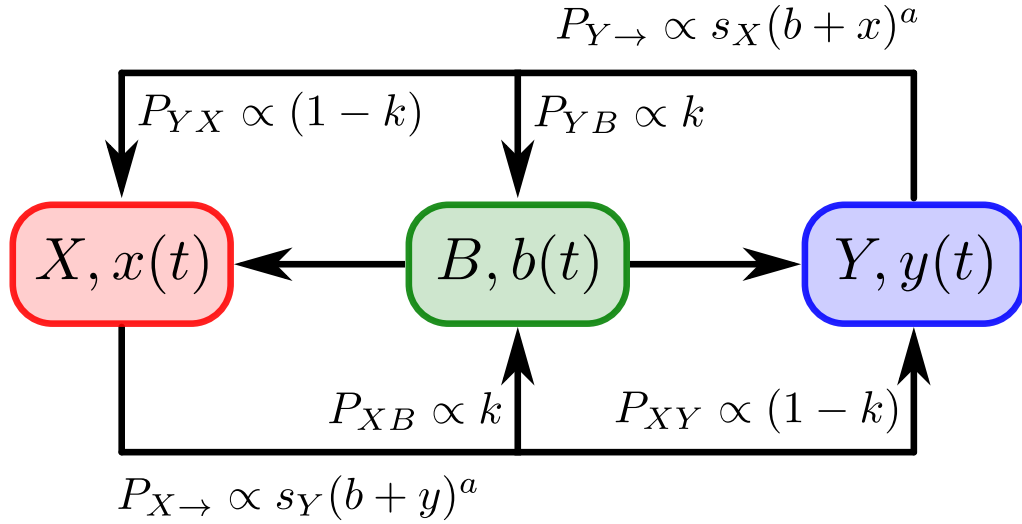


FIG. 8 **Fluxes of speakers between three sociolinguistic groups.** If we think about people as categorically divided within three abstract groups (X , Y , and Z), we can come up with a very general model of how people move from one group to another assuming that they do so with certain probabilities. Such probabilities can be anything, but we make them to obey some *tentative* and intuitive sociolinguistic hypothesis – e.g. the likelihood that you decide to learn a new language may be proportional to the number of people speaking that language and to the prestige that you might perceive that that language contributes – disregarding of whether that prestige is real or imagined.

one of the other two. This means that the change in the number of people in each group have to balance each other: $\Delta x + \Delta y + \Delta b = 0$. Thanks to this, if we know how much people entered (or left) the groups X and Y , we know right away how many people entered (or left) the group B . Hence we do not need to solve the three equations above and we can focus in just two. We choose to solve the equations for Δx and Δy (but we could have chosen any other combination).

These variables (Δx and Δy) tell us the relative number of people that enters groups X and Y within a time span $c\Delta t$. To solve equations like the ones above, it is often convenient to write not the number of people that enters or leaves a group, but the rates ($\Delta x/\Delta t \equiv dx/dt$ and $\Delta y/\Delta t \equiv dy/dt$) at which people enter or leave the group. Hence the relevant equations become:

$$\frac{dx}{dt} = c \cdot [yP_{YX} + bP_{BX} - x(P_{XY} + P_{XB})], \quad (7)$$

$$\frac{dy}{dt} = c \cdot [xP_{XY} + bP_{BY} - y(P_{YX} + P_{YB})]. \quad (8)$$

$$(9)$$

The notation dx/dt and dy/dt are just fancy ways to write these rates when the time span becomes infinitely small. Note how a parameter of the model, c , multiplies all fluxes in both equations. Everything being equal, increasing this parameter entails that all fluxes in the system are increased accordingly. In other words: c sets an overall velocity at which fluxes flow. Such velocity is defined with respect to an external time scale (e.g. does the process take place over years or seconds?), but it has no telling in how fast fluxes are with respect to each other. To answer that next question we need to look at the elements inside the brackets. A recent, lengthy discussion of the parameter c with empirical applications can be found in (Juane et al., 2019).

Back to that next question (how internal fluxes of people relate to each other, and not to an overall external scale): Note that this model is very general about “people” moving from one group to another. Actually, we could have said that these people were “objects” moving between three boxes, and we would have arrived to the same equation. This is: whatever we do next, our problem shares all this much with any other problem of people or objects moving around three options. Let us now make some sociolinguistic assumption about what might entice people to change groups. Therefore, groups X , Y , and B become less abstract, and from now on they represent respectively the group of monolingual speakers of languages X and Y , and the group of bilingual speakers.

If you are a monolingual speaker of language X , a good reason for you to learn language Y are all the people that you are going to be able to speak with with this new language. This is, potentially, all the people in groups B and Y . The more people in either of these groups, the higher a pressure you might feel to learn Y . Hence, the probability that you leave group X will be somehow directly proportional to b and y . Of course you can already speak with all

people in B , since they are bilinguals, thus they might not exert a great pressure for you to learn a new language. Such possibility was considered by (Heinsalu et al., 2014) in their discussion about the *attracting population*. To make things easier here, we assume that people in both B and Y entice you similarly to learn that new language. (If this were a crucial and wrong choice, our models would sooner or later fail to reproduce the data.) With this, the probability that we leave group X will be somehow proportional to $b + y$. Another good reason to learn language Y for monolinguals of language X is that the former might have associated some prestige s_Y as perceived by that person. We do not know what this prestige is or where might it originate. It might be a political pressure, a socio-economic one, or an aggregate of these and other factors. We do not care, but we assume that such prestige exists, and that the bigger it is, the more likely it is that monolingual speakers of X will be enticed to learn Y . Thus the probability that a person leaves group X is somehow proportional to both the attracting population $b + y$ and the prestige s_Y of language Y .

But both these contributions are different in nature so a question arises as to how to bring them together in an equation. The most fair way to do this is to introduce yet another parameter, a , that balances the relative importance of these two factors. Thus, the probability that monolingual speakers of X choose to learn language Y must be proportional to $s_Y(b + y)^a$. Note how if a becomes close to 0, the prestige is the only relevant factor to learn language Y . If a becomes very large, then the attracting population becomes much more relevant. More subtly, this parameter has the ability of ‘dissolving’ an existing group when a is too low. This is so because when only prestige matters you might choose to leave your own group even if it agglutinates most of the people. For this reason, the parameter has been termed *volatility* (Castelló et al., 2013) (while it actually measures the inverse of a volatility).

Finally, of all the monolingual speakers of X that chose to learn language Y , a fraction k of them will factually retain both languages (use it in its every day), while a fraction $1 - k$ will stop using X all together. This is, an amount $P_{XB} \equiv k s_Y(b + y)^a$ will leave group X to join group B , and an amount $P_{XY} \equiv (1 - k) s_Y(b + y)^a$ will leave group X to join group Y . We can follow a similar reasoning to reach $P_{YB} \equiv k s_X(b + x)^a$ and $P_{YX} \equiv (1 - k) s_X(b + x)^a$. Finally, we assume that bilingual speakers retain the chance to stop using each of the languages, and that they do so with probabilities: $P_{BX} \equiv (1 - k) s_X(b + x)^a$ and $P_{BY} \equiv (1 - k) s_Y(b + y)^a$. This is, they are subjected to the same attractions as monolinguals. As before, if this were a critical and wrong assumption, sooner or later our model would fail to reproduce the data. As long as the model reproduces the data within acceptable accuracy, we cannot disprove (nor prove!) our assumptions.

Back to the parameter k , since it measures the amount of people that keep using both languages, it is fair to assume that this number would be larger for sets of tongues which their speakers perceive as similar. Based on this, k was termed *interlinguistic similarity*. There is a very interesting and fundamental question: after which value of k can we consider that two languages are the same? Put otherwise, how much diversity (leading to decrease in k) can a same language tolerate? These are matters for other papers. The fact that this parameter is termed interlinguistic similarity is not crucial to interpret the results: what matters is that it has a causal effect on speakers retaining both languages for whatever reasons they might have chosen.

Note also that a couple of languages might be retained for speakers in a region while it might not be retained by speakers in another region (as discussed, e.g., by (Fishman, 2013)). This means that whatever parameters we derive, they are not inherent to the languages considered. Instead, they emerge from the idiosyncrasies and historical dependencies encapsulated by the available data. This is: different people might perceive the same set of languages as having different interlinguistic similarity, prestige, or volatility.

We can now plug all the numbers into the equations above to get:

$$\begin{aligned} \frac{dx}{dt} &= c \cdot [y(1 - k) s_X(b + x)^a + b(1 - k) s_X(b + x)^a - x((1 - k) s_Y(b + y)^a + k s_Y(b + y)^a)], \\ \frac{dy}{dt} &= c \cdot [x(1 - k) s_Y(b + y)^a + b(1 - k) s_Y(b + y)^a - y((1 - k) s_X(b + x)^a + k s_X(b + x)^a)]. \end{aligned} \quad (10)$$

This, with a little bit of algebra, can be written as Eq (1) from the main text, which much more compact and makes working with these objects easier.

This is a system of differential equations, and we know methods to solve them and, even when we cannot solve them, to extract properties about their solutions. A solution to these equations means *an evolution of the proportions of people in groups X , Y , and B over time*, and these are numbers that we can compare to real proportions of speakers of a couple of coexisting tongues. A choice of values for the parameters s_X , s_Y , k , a , and c results in a specific time evolution of the variables x , y , and b . If, for a given choice of the parameters, we find solutions that are very different from the actual number of speakers of two coexisting languages, then we know at least that those parameters are not good descriptors of that system of coexisting languages. By finding the values of parameters that better describe some given data, we are looking for the best description of that data in the terms dictated by the model – and that is the intent of our present work.

As said above, we can also extract general properties of the solutions even if we could not solve these equations. For example, (Otero-Espinar et al., 2013) have proven mathematically that there are not periodic nor oscillating solutions to these equations. This means that, according to these equations, it cannot happen that the proportions of speakers fall and raise and fall and raise several times consecutively. If we would observe such a behavior *very clearly* in our data, then we would know not only that some parameters do not work, but the the model is wrong all together (i.e. some of our assumptions is crucial and incorrect). In our data sets there are cases in which the number of speakers appears to raise, then fall, then raise, etc. Such fluctuation seem rather attributable to inherent noise in the system (e.g. different people are asked for different polls, so that results in the polls only approximate the actual number of speakers). All in all, these oscillating behaviors of the data are not *very clear*, as needed to disprove the model. Note, again, that this does not prove either that the model is valid at all.

References

- Castelló X, Loureiro-Porto L, San Miguel M (2013) Agent-based models of language competition. *International Journal of the Sociology of Language* 2013(221):21-51 *Europhys Lett* 69:1031-1034
- Fishman JA (2013) Language Maintenance, Language Shift, and Reversing Language Shift. In *The handbook of bilingualism and multilingualism*, Bhatia T K, Ritchie W C, Wiley J (Eds.), p.466.
- Generalitat de Catalunya (2008a) Enquesta d'usos lingüístics de la població 2008: anàlisi. Volum I. Les llengües a Catalunya: coneixements, usos, transmissió i actituds lingüístics. (Survey of language use of the population 2008: analysis. Volume I. Languages in Catalonia: knowledge, usage, transmission, and attitudes towards language) Generalit de Catalunya
- Generalitat de Catalunya (2008b) Enquesta d'usos lingüístics dela població 2008: anàlisi. Volum II. Els factors clau de les llengües a Catalunya: edadt, ocupació, classe social, lloc de naixement, territori i aprenentatge de català. (Survey of language use of the population 2008: analysis. Volume II. Key aspects of the languages in Catalonia: age, occupation, social class, place of birth, territory, and learning of Catalan) Generalit de Catalunya
- Generalitat de Catalunya, Direcció General de Política Lingüística (2013) Anàlisi de l'Enquesta d'usos lingüístics de la població 2013. Resum dels factors clau. (Analysis of the survey of language use of the population 2013. Summary of key aspects) Generalit de Catalunya
- Heinsalu E, Patriarca M, Léonard J (2014) The role of bilinguals in language competition. *Advances in Complex Systems* 17(1):1450003
- Juane MM, Seoane LF, Muñuzuri AP, Mira J (2019) Urbanity and the dynamics of language shift in Galicia. *Nat Com* 10(1):1680.
- Mira J, Paredes Á (2005) Interlinguistic similarity and language death dynamics.
- Mira J, Seoane LF, Nieto JJ (2011) The importance of interlinguistic similarity and stable bilingualism when two languages compete. *New J Phys* 13:033007
- Otero-Espinar MV, Seoane LF, Nieto JJ, Mira J (2013) An analytic solution of a model of language competition with bilingualism and interlinguistic similarity. *Physica D* 264:17-26
- Seoane LF, MiraJ (2017) Modeling the life and death of competing languages from a physical and mathematical perspective. In: Lauchlan F, Parafita Couto MC, editors. *Bilingualism and Minority Languages in Europe: Current trends and development*. Cambridge Scholars Press:70-93
- Shalizi C (2013) *Advanced data analysis from an elementary point of view*. Cambridge University Press