

In the format provided by the authors and unedited.

A potential game approach to modelling evolution in a connected society

Jiabin Wu ¹ and Dai Zusai ^{2*}

¹Department of Economics, University of Oregon, Eugene, OR, USA. ²Department of Economics, Temple University, Philadelphia, PA, USA.

*e-mail: zusai@temple.edu

1 Supplementary Methods

1.1 Model

The base game and structured populations

In the base game, we let $\mathcal{A} := \{1, \dots, A\}$ be the action set and define the payoff function $\mathbf{F} : \mathbb{R}^A \rightarrow \mathbb{R}^A$ by $\mathbf{F}(\mathbf{x}) = \mathbf{\Pi}\mathbf{x}$.

The society is divided into P subpopulations $\mathcal{P} := \{1, \dots, P\}$. Denote by $x_a^p \in [0, 1]$ the *mass* of agents choosing action $a \in \mathcal{A}$ in population p ; let an A -dimensional (column) vector $\mathbf{x}^p := (x_a^p)_{a \in \mathcal{A}}$ be the action distribution in population p . We call $\mathbf{x}^{\mathcal{P}} = (\mathbf{x}^p)_{p \in \mathcal{P}}$ a social state:¹

$$\mathbf{x}^{\mathcal{P}} = (\mathbf{x}^1, \dots, \mathbf{x}^P) = (x_1^1, \dots, x_A^1, \dots, x_1^P, \dots, x_A^P) = (x_a^p)_{(a,p) \in \mathcal{A} \times \mathcal{P}} \in \Delta^{AP}.$$

Denote by $m^p := \sum_{a \in \mathcal{A}} x_a^p$ the mass of agents in population p ; let $\mathbf{m}^{\mathcal{P}} = (m^p)_{p \in \mathcal{P}} \in \Delta^P$. Let $\mathcal{X}^{\mathcal{P}}$ be the set of feasible social states, which depends on the freedom of choices for the agents in the society. $\mathcal{X}^{\mathcal{P}}$ is simply the whole space Δ^{AP} if the agents can freely choose their actions and population affiliations. In this case, $m^{\mathcal{P}}$ is not fixed. $\mathcal{X}^{\mathcal{P}}$ equals $m^{\mathcal{P}} \Delta^A := \times_{p \in \mathcal{P}} m^p \Delta^A$ if the agents can only choose their actions but not population affiliations. In this case, the population distribution $m^{\mathcal{P}}$ is fixed. Note that $m^{\mathcal{P}} \Delta^A \subsetneq \Delta^{AP}$.

The structured game

Consider a social state $\mathbf{x}^{\mathcal{P}}$. Entry g_{pq} in the connection matrix $\mathbf{G}$ is the intensity of interactions between populations p and q . We assume symmetry of $\mathbf{G}$, i.e., $g_{pq} = g_{qp}$. For agents in population p , the weighted mass of agents choosing action a in population q is given by $g_{pq}x_a^q$; the weighted mass of agents choosing action a in the whole society is given by $\tilde{x}_a^p := \sum_{q \in \mathcal{P}} g_{pq}x_a^q$. The distribution of weighted masses of agents choosing different actions for an agent from population p is given by $\tilde{\mathbf{x}}^p := \sum_{q \in \mathcal{P}} g_{pq}\mathbf{x}^q$. The payoff vector for agents from population p is given by

$$\tilde{\mathbf{F}}^p(\mathbf{x}^{\mathcal{P}}) := \mathbf{F}(\tilde{\mathbf{x}}^p) = \mathbf{\Pi}\tilde{\mathbf{x}}^p = \sum_{q \in \mathcal{P}} g_{pq}\mathbf{\Pi}\mathbf{x}^q.$$

The payoff function $\tilde{\mathbf{F}} : \mathcal{X}^{\mathcal{P}} \rightarrow \mathbb{R}^{AP}$ of the game now is given by²

$$\tilde{\mathbf{F}}(\mathbf{x}^{\mathcal{P}}) := \begin{pmatrix} \tilde{\mathbf{F}}^1(\mathbf{x}^{\mathcal{P}}) \\ \vdots \\ \tilde{\mathbf{F}}^P(\mathbf{x}^{\mathcal{P}}) \end{pmatrix} = \begin{pmatrix} \sum_{q \in \mathcal{P}} g_{1q}\mathbf{\Pi}\mathbf{x}^q \\ \vdots \\ \sum_{q \in \mathcal{P}} g_{Pq}\mathbf{\Pi}\mathbf{x}^q \end{pmatrix} = \underbrace{\begin{pmatrix} g_{11}\mathbf{\Pi} & \cdots & g_{1P}\mathbf{\Pi} \\ \vdots & \ddots & \vdots \\ g_{P1}\mathbf{\Pi} & \cdots & g_{PP}\mathbf{\Pi} \end{pmatrix}}_{AP \times AP} \underbrace{\mathbf{x}^{\mathcal{P}}}_{AP \times 1} = \underbrace{[\mathbf{G} \otimes \mathbf{\Pi}]}_{\tilde{\mathbf{\Pi}}} \mathbf{x}^{\mathcal{P}} \in \mathbb{R}^{AP}.$$

A structured game is defined by $\tilde{\mathbf{F}}$.

¹Strictly speaking, a bold letter like $\mathbf{v}$ denotes a column vector while an arrow like $\vec{v}$ denotes a row vector. Hence, we treat $\mathbf{x}^{\mathcal{P}}$ as an AP -dimensional column vector. Consider a $|\mathcal{U}|$ -dimensional real space. Each coordinate is labeled with an element of $\mathcal{U} = \{1, \dots, |\mathcal{U}|\}$. For a set $S \subset \mathcal{U}$, we define a $|S|$ -dimensional simplex $\Delta^{|\mathcal{U}|}(S)$ as $\Delta^{|\mathcal{U}|}(S) := \{\mathbf{x} \in \mathbb{R}_+^{|\mathcal{U}|} \mid \sum_{k \in S} x_k = 1 \text{ and } x_l = 0 \text{ for any } l \in \mathcal{U} \setminus S\}$. When S is the whole space $\mathcal{U}$ itself, we omit (S) and denote it by $\Delta^{|\mathcal{U}|}$.

²We denote by $\mathbf{A} \otimes \mathbf{B}$ a Kronecker product of matrices $\mathbf{A} = (A_{ij})_{i,j}$ and $\mathbf{B}$ such as $\mathbf{A} \otimes \mathbf{B} = (A_{ij}\mathbf{B})_{i,j}$.

Equilibria and the potential function

The formal definitions of the medium-run and long-run equilibria are given as below.

Definition 1.

- (1) Given population distribution $\mathbf{m}^{\mathcal{P}} = (m^p)_{p \in \mathcal{P}}$, social state (action distribution) $\mathbf{x}^{\mathcal{P}} \in m^{\mathcal{P}} \Delta^A$ is a **medium-run equilibrium**, if $\tilde{F}_a^p(\mathbf{x}^{\mathcal{P}}) \geq \tilde{F}_b^p(\mathbf{x}^{\mathcal{P}})$ for any $b \in \mathcal{A}$ whenever $x_a^p > 0$.
- (2) Social state $\mathbf{x}^{\mathcal{P}} \in \Delta^{AP}$ is a **long-run equilibrium**, if $\tilde{F}_a^p(\mathbf{x}^{\mathcal{P}}) \geq \tilde{F}_b^q(\mathbf{x}^{\mathcal{P}})$ for any $q \in \mathcal{P}, b \in \mathcal{A}$ whenever $x_a^p > 0$.

The most useful results for a potential game is that a local maximizer of the potential function in the feasible state space is a Nash equilibrium. In general, the set of Nash equilibria of a potential game is identical with the set of solutions of Karush-Kuhn-Tucker first-order conditions to maximize the potential function given the constraints to keep the maximizer in the feasible state space (Sandholm, 2001). In the settings considered in this paper, we have

- Given $m^{\mathcal{P}}$, a local maximizer of $\tilde{f}$ in $m^{\mathcal{P}} \Delta^A$ is a medium-run equilibrium.
- A local maximizer of $\tilde{f}$ in Δ^{AP} is a long-run equilibrium.

Construction of the connection matrix

From communities to populations Let $\mathcal{C}$ be the set of communities (taken as given in our model) and $c \in \mathcal{C}$ be the representative index of a community. An agent may belong to multiple communities. Define the set of populations $\mathcal{P}$ as the power set over $\mathcal{C}$ excluding the empty set. That is, a population p is defined as a subset of $\mathcal{C}$. For example, consider the case of three communities, $\mathcal{C} = \{c_1, c_2, c_3\}$. A subset $\{c_1, c_2\} \subset \mathcal{C}$ is a population. We interpret it as the set of agents who are affiliated with both communities c_1 and c_2 but not c_3 . Let p_{12} denote this population. With similar notations, we can write $\mathcal{P} = 2^{\mathcal{C}} \setminus \{\emptyset\} = \{p_1, p_2, p_3, p_{12}, p_{13}, p_{23}, p_{123}\}$.

Define a binary variable to represent agents' affiliations:

$$B_c^p = \begin{cases} 1, & \text{if } c \in p, \text{ i.e., agents in } p \text{ belong to community } c, \\ 0, & \text{if } c \notin p, \text{ i.e., agents in } p \text{ do not belong to community } c, \end{cases}$$

for each $p \in \mathcal{P}, c \in \mathcal{C}$.

Construction of the structure matrix Assume that each agent only interacts with other agents in the communities that the agent is affiliated with. In addition, the intensity of interaction between two agents is independent of the number of the communities that they share. In this case, g_{pq} take a binary value from $\{0, 1\}$. More specifically, $g_{pq} = 1$ if p and q have at least one shared community and $g_{pq} = 0$ if p and q has no shared community:

$$g_{pq} = \max\{B_c^p B_c^q \mid c \in \mathcal{C}\}.$$

Alternative construction of the structure matrix Alternatively, one can assume that an agent encounters another agent once in each community that both of them belong to. In this case, g_{pq} is proportional

to the number of communities shared over population p and population q :

$$g_{pq} = \sum_{c \in C} B_c^p B_c^q.$$

One can observe that if an agent from population p and another agent from population q have more shared affiliations, the social influence between them becomes stronger.

We can further extend the formulation to allow for *heterogeneous degrees of within community influences*. More specifically, agents can have different intensities of influences on one another within each community. In this case, we modify g_{pq} as

$$g_{pq} = \sum_{c \in C} B_c^p \delta_c B_c^q,$$

where $\delta_c \in [0, 1]$ captures the intensity of influence of agents affiliated with community c .

We can also allow the model to incorporate *heterogeneous degrees of affiliation to a community*. Agents may have different degrees of affiliations to a community depending on how intensive they engage in the activities offered by the community. To capture such a scenario, we can allow B_c^p to take a value from $[0, 1]$ instead of $\{0, 1\}$.

Canonical examples of connection structures

Different intensities of interactions within/across populations (Fig.1a) In Fig.1a, the society is divided into two populations $\mathcal{P} = \{1, 2\}$. Following Example 1, we introduce parameter $\beta \in \mathbb{R}$, let $g_{11} = g_{22} = 1$, $g_{12} = g_{21} = -\beta$; then, the structure matrix is given by

$$\mathbf{G} = \begin{pmatrix} 1 & -\beta \\ -\beta & 1 \end{pmatrix}.$$

The payoff matrix of the structured game is obtained as

$$\tilde{\mathbf{\Pi}} = \mathbf{G} \otimes \mathbf{\Pi} = \begin{pmatrix} \mathbf{\Pi} & -\beta \mathbf{\Pi} \\ -\beta \mathbf{\Pi} & \mathbf{\Pi} \end{pmatrix}.$$

In Example 1, we assume $\beta > 0$ (and thus $g_{12} = g_{21} = -\beta < 0$), so the game with agents in the other population has an opposite payoff matrix compared to the one within the same population.

The potential function is

$$\tilde{f}(\mathbf{x}^{\mathcal{P}}) = 0.5 \mathbf{x}^{\mathcal{P}} \cdot \tilde{\mathbf{\Pi}} \mathbf{x}^{\mathcal{P}} = 0.5 \sum_{p \in \mathcal{P}} \mathbf{x}^p \cdot \mathbf{\Pi} \mathbf{x}^p - \beta \mathbf{x}^1 \cdot \mathbf{\Pi} \mathbf{x}^2.$$

$\mathbf{G}$ is positive definite if $|\beta| < 1$, positive semidefinite if $|\beta| = 1$, and indefinite if $|\beta| > 1$. Note that when $|\beta| \leq 1$, the definiteness of the base game $\mathbf{\Pi}$ carries over to the structured game $\tilde{\mathbf{\Pi}}$.

Overlapping communities (Fig.1b) In Fig.1b, there are two communities L and R in the society. An agent can be affiliated with at least one community, possibly both.³ We categorize agents into subpopulations by affiliating communities. Agents in population Lo are affiliated with only community L ; agents in

³For simplicity, we exclude the possibility that some agents are not affiliated with any community.

population LR are affiliated with both community L and R ; agents in population oR are affiliated with only community c_R . Thus, $\mathcal{P} = \{Lo, LR, oR\}$.

We make two assumptions on the connection structure. First, assume that each agent only interacts with other agents in the communities that the agent is affiliated with. This implies that agents affiliated with both communities, i.e., those in the subpopulation LR , interact with more agents than others. Second, assume that the intensity of interaction between two agents is independent of the number of the communities that they share. The structure matrix $\mathbf{G}$ representing the connection structure is given by

$$\mathbf{G} = \begin{pmatrix} g_{Lo,Lo} & g_{Lo,LR} & g_{Lo,oR} \\ g_{LR,Lo} & g_{LR,LR} & g_{LR,oR} \\ g_{oR,Lo} & g_{oR,LR} & g_{oR,oR} \end{pmatrix} = \begin{pmatrix} 1 & 1 & 0 \\ 1 & 1 & 1 \\ 0 & 1 & 1 \end{pmatrix}. \quad (1)$$

The payoff matrix of the structured game is obtained as

$$\tilde{\mathbf{\Pi}} = \mathbf{G} \otimes \mathbf{\Pi} = \begin{pmatrix} \mathbf{\Pi} & \mathbf{\Pi} & 0 \\ \mathbf{\Pi} & \mathbf{\Pi} & \mathbf{\Pi} \\ 0 & \mathbf{\Pi} & \mathbf{\Pi} \end{pmatrix}.$$

The potential function is given by

$$\tilde{f}(\mathbf{x}^{\mathcal{P}}) = 0.5 \mathbf{x}^{\mathcal{P}} \cdot \tilde{\mathbf{\Pi}} \mathbf{x}^{\mathcal{P}} = 0.5 \left(\sum_{p \in \mathcal{P}} \mathbf{x}^p \cdot \mathbf{\Pi} \mathbf{x}^p + 2(\mathbf{x}^{Lo} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi} \mathbf{x}^{LR} \right),$$

where $\mathbf{x}^p = (x_I^p, x_O^p)$ for each $p \in \{Lo, LR, oR\}$. Note that the connection matrix $\mathbf{G}$ is indefinite. Hence, the payoff matrix $\tilde{\mathbf{\Pi}}$ of the structured game is also indefinite. Correspondingly, the potential function $\tilde{f}$ is neither concave nor convex.⁴

1.2 Dynamics

The two assumptions for admissible dynamics

The best response stationarity and the positive correlation are mathematically expressed as follows:⁵ 1) the strategy distribution should not change when the agents are taking optimal strategies given the current payoffs; 2) otherwise, they revise their strategies in a direction that is positively correlated with the relative payoffs so as to yield a positive payoff improvement on average.

Definition 2 (Best response stationarity of a subdynamic: BR stationarity). Subdynamic $\mathbf{v}^S : \Delta^{AP} \times$

⁴Alternatively, one can assume that the intensity of interaction between two agents increases proportionally with the number of their shared communities. In this case, $g_{LR,LR}$ in the structure matrix can be replaced by 2, because two agents in overpopulation LR share both communities L and R . This makes the connection matrix $\mathbf{G}$ positive semidefinite, implying that the payoff matrix $\tilde{\mathbf{\Pi}}$ of the structured game retains the same semidefiniteness as that of the base game $\mathbf{\Pi}$.

⁵The evolutionary dynamic we propose is not standard in the sense that we allow a mixture of different types of revision protocols to construct the dynamic. Therefore, instead of defining properties directly on the dynamic for the overall population as in the extant literature, it is appropriate to first define properties on each subdynamic and then combine them together to generate the desired properties on the aggregate level. The related contributions include Zusai (2018b), who extends equilibrium stationarity and positive correlation to heterogeneous dynamics, and Zusai (2018a), who provides a unified proof of stability of regular evolutionary stable states under general (nonimitative) evolutionary dynamics. In both papers, agents can adopt different protocols and they may not be able to observe other agents' types.

$\mathbb{R}^{AP} \rightarrow \mathbb{R}^{AP}$ satisfies **best response (BR) stationarity** if, for any $(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) \in \Delta^{AP} \times \mathbb{R}^{AP}$,

$$\mathbf{v}^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) = \mathbf{0} \iff \forall (a, p) \in \mathcal{A} \times \mathcal{P} [x_a^p > 0 \implies (a, p) \in \operatorname{argmax}_{(a', p') \in \mathcal{S}(a, p)} \pi_{a'}^{p'}]. \quad (2)$$

Definition 3 (Positive correlation of a subdynamic: PC). Subdynamic $\mathbf{v}^{\mathcal{S}} : \Delta^{AP} \times \mathbb{R}^{AP} \rightarrow \mathbb{R}^{AP}$ satisfies **positive correlation (PC)** if⁶

$$\boldsymbol{\pi}^{\mathcal{P}} \cdot \mathbf{v}^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) \begin{cases} \geq 0 & \text{for any } (\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) \in \Delta^{AP} \times \mathbb{R}^{AP}; \\ = 0 & \text{only if } \mathbf{v}^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) = \mathbf{0}. \end{cases} \quad (3)$$

Examples of admissible dynamics

Here we provide examples of well-established revision protocols and show how we generate subdynamics from them.

Example 1 (Best response dynamic). Consider an exact optimization protocol under which a revising agent switches only to a strategy that yields the greatest payoff among all available strategies:

$$\rho_{s, s'}^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) > 0 \implies s' \in \operatorname{argmax}_{s'' = (a'', p'') \in \mathcal{S}(s)} \pi_{s''} \quad \text{for any } s' \in \mathcal{S}(s) \setminus \{s\}.$$

Aggregation to a subdynamic defines a best response dynamic (BRD), proposed by Gilboa and Matsui (1991); Hofbauer (1995):

$$\mathbf{v}_{\text{BRD}}^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) = B^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) - \mathbf{x}^{\mathcal{P}},$$

where

$$B^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) = \sum_{s=(a,p) \in \mathcal{S}} x_a^p \Delta^{A \times \mathcal{P}} \left(\operatorname{argmax}_{s'=(a', p') \in \mathcal{S}(s)} \pi_{a'}^{p'} \right) = \sum_{s \in \mathcal{S}} \operatorname{argmax}_{\mathbf{y}_s \in x_a^p \Delta^{AP}(\mathcal{S}(s))} \mathbf{y}_s \cdot \boldsymbol{\pi}^{\mathcal{P}}$$

is the set of best response mixed strategies. Subdynamic $\mathbf{v}_{\text{BRD}}^{\mathcal{S}}$ satisfies both **BR** stationarity and **PC**.

Note that, in each type of revision opportunities, $B^{\mathcal{S}}$ reduces to

$$\begin{aligned} B^{\mathcal{A}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) &:= \prod_{p \in \mathcal{P}} \operatorname{argmax}_{\mathbf{y}^p \in m^p \Delta^A} \mathbf{y}^p \cdot \boldsymbol{\pi}^p, \\ B^{\mathcal{P}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) &:= \prod_{a \in \mathcal{A}} \operatorname{argmax}_{\mathbf{y}_a \in \sum_{p \in \mathcal{P}} x_a^p \Delta^{\mathcal{P}}} \mathbf{y}_a \cdot \boldsymbol{\pi}_a, \\ B^{AP}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) &:= \operatorname{argmax}_{\mathbf{y} \in \Delta^{AP}} \mathbf{y} \cdot \boldsymbol{\pi}^{\mathcal{P}}, \end{aligned}$$

where $\boldsymbol{\pi}^{\mathcal{P}} = (\pi_a^p)_{a \in \mathcal{A}}$ and $\boldsymbol{\pi}_a = (\pi_a^p)_{p \in \mathcal{P}}$ collect the payoffs of all available strategies $\mathcal{A}(a, p) = \mathcal{A} \times \{p\}$ or $\mathcal{P}(a, p) = \{a\} \times \mathcal{P}$ in revision opportunities $\mathcal{A}$ and $\mathcal{P}$, respectively. ■

Example 2 (Pairwise comparison dynamics). Under a pairwise comparison protocol, the probability that a revising agent switches to a new strategy is determined by the comparison of the payoffs associated with the agent's current strategy and the new one. The agent will not switch to the new strategy if the new strategy yields a lower payoff and the agent is more likely to switch to it if it gives the agent a higher payoff gain.

⁶If $\mathbf{v}^{\mathcal{S}}$ is a differential inclusion, i.e., it allows multiple transition vectors like the best response dynamic, then $\mathbf{v}^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) = \mathbf{0}$ should read as $\mathbf{0} \in \mathbf{v}^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}})$.

The protocol is formulated as:

$$\rho_{ss'}^S(\mathbf{x}^P, \boldsymbol{\pi}^P) = Q(\pi_{s'} - \pi_s) \quad \text{for any } s' \in \mathcal{S}(s) \setminus \{s\},$$

with an weakly increasing function $Q : \mathbb{R}_+ \rightarrow \mathbb{R}$ such as $Q(z) = 0$ if $z \leq 0$, $Q(z) > 0$ and it is strictly increasing in z when $z > 0$.

Aggregation to a subdynamic defines a pairwise comparison dynamic; in particular, if $Q(z) = [z]_+$, then the subdynamic reduces to a Smith dynamic (Smith, 1984).⁷ Under the above assumption on Q , subdynamic $\mathbf{v}_{\text{PCD}}^S$ satisfies both **BR** stationarity and **PC**. ■

Example 3 (Imitative dynamics). Under an imitative protocol, a revising agent randomly samples a candidate agent and decides whether to imitate the candidate's strategy based on the comparison of payoffs associated with the agent's current strategy and the candidate's strategy. The agent will not switch strategy if the candidate's strategy yields a lower payoff; and the agent is more likely to imitate the candidate's strategy if it gives the agent a higher payoff gain. This protocol is formulated as

$$\rho_{ss'}^S(\mathbf{x}^P, \boldsymbol{\pi}^P) = x_{s'} Q(\pi_{s'} - \pi_s) \quad \text{for any } s' \in \mathcal{S}(s) \setminus \{s\},$$

with a weakly increasing function $Q : \mathbb{R}_+ \rightarrow \mathbb{R}$ such as $Q(z) = 0$ if $z \leq 0$, $Q(z) > 0$ and it is strictly increasing in z when $z > 0$. $x_{s'}$ in the expression is the probability to sample a candidate using strategy s' in the entire society. It implies that the switching probability to strategy s' proportionally increases in the mass of agents who are currently taking s' .

Aggregation to a subdynamic defines an imitative dynamic; in particular, if $Q(z) = [z]_+$, then the subdynamic is a replicator dynamic (Schlag, 1998). Note that an imitative dynamic fails to guarantee **BR** stationarity when no agent is currently taking the best response and thus the best response cannot be sampled. Nevertheless it satisfies **PC**. ■

1.3 General analysis

By utilizing the two basic properties **BR** stationarity and **PC** defined on the subdynamics, we establish some fundamental theorems to link equilibrium states in a game with stationary points in a dynamic, as summarized in the main text. The theorems verify general equilibrium stationarity and equilibrium stability in potential games, while we need to tailor them to fit each of the two different settings—medium-run equilibria in the medium-run dynamics and the long-run equilibria in the long-run dynamics.

Equilibrium stationarity in a general population game: In any structured population game $\tilde{\mathbf{F}}$, every Nash equilibrium is a rest point of the dynamic and vice versa:

$$\dot{\mathbf{x}}^P = \mathbf{0} \quad \Longleftrightarrow \quad \mathbf{x}^P \text{ is a Nash equilibrium in } \tilde{\mathbf{F}}.$$

Equilibrium stability in a potential game: Suppose that a structured game $\tilde{\mathbf{F}}$ is a potential game. Then, i) the set of equilibria are globally attracting; ii) an isolated local maximizer of the potential function is asymptotically stable; iii) the potential function serves as a strictly increasing Lyapunov function.

⁷ $[z]_+ := \max\{0, z\}$.

The “Nash equilibrium” in the above definitions refers to a *medium-run equilibrium* for a medium-run dynamic in which agents can revise only actions, and to a *long-run equilibrium* for a long-run dynamic in which agents can revise both actions and affiliations. Similarly, a “maximizer” of a potential function refers to the social state $\mathbf{x}^{\mathcal{P}}$ that yields the greatest value of the potential in the set of *feasible* social states, which equals $\times_{p \in \mathcal{P}} m^p \Delta^A$ (i.e., the population distribution $m^{\mathcal{P}}$ is fixed) for a medium-run dynamic and Δ^{AP} for a long-run dynamic. Note that the set of local maximizers of a potential function is a subset of the Nash equilibrium set. Hence, the equilibrium stability defined above implies global convergence to the set of Nash equilibria.

In the theorems below, we will formally clarify why **BR** stationarity and **PC** are needed on the (sub)dynamic level to guarantee the two equilibrium properties defined above. On one hand, equilibrium stationarity can be naturally expected from **BR** stationarity. On the other hand, similar to Sandholm (2001)’s result for standard evolutionary dynamics, **PC** suggests that in a potential game, the value of the potential function ascends under the dynamic and the dynamic converges to those Nash equilibria corresponding to the local maximizers of the potential function.

Medium-run equilibrium stationarity and stability in the medium-run dynamics

Theorem 1 (Medium-run equilibrium stationarity and stability). *Consider a medium-run dynamic $\mathbf{v}^A : \Delta^{AP} \times \mathbb{R}^{AP} \rightarrow \mathbb{R}^{AP}$ under which agents receive only revision opportunities on $\mathcal{A}$ with rate $\lambda_A = 1$.*

- i) *If $\mathbf{v}^A$ satisfies **BR** stationarity, then the medium-run dynamic $\mathbf{v}^A$ satisfies the medium-run equilibrium stationarity in any structured game $\tilde{\mathbf{F}}$:*

$$\dot{\mathbf{x}}^{\mathcal{P}} = \mathbf{0} \quad \Longleftrightarrow \quad \mathbf{x}^{\mathcal{P}} \text{ is a medium-run equilibrium.}$$

- ii) *In addition, suppose that $\mathbf{v}^A$ satisfies **PC** as well, implying that the dynamic is admissible. Then, the medium-run dynamic $\mathbf{v}^A$ satisfies the medium-run equilibrium stability in any structured potential game.*

Part i) of Theorem 1 is immediate by the following observation. All agents are choosing their optimal strategies at a social state $\mathbf{x}^{\mathcal{P}}$ if and only if $\mathbf{x}^{\mathcal{P}}$ is a Nash equilibrium; then by **BR** stationarity, the strategy distribution would not change from $\mathbf{x}^{\mathcal{P}}$. Given that the dynamic is admissible, part ii) of Theorem 1 is an immediate application of the stability theorem in Sandholm (2001).

Long-run equilibrium stationarity and stability in the long-run dynamics

Unless $\lambda_A = \lambda_{\mathcal{P}} = 0$, a long-run dynamic cannot be analyzed in the same way as a standard evolutionary dynamic because a revising agent may be able to change either action or population affiliation, which makes the available strategy sets depend on the current strategy profile. Nevertheless, the following theorem shows that as long as all three subdynamics satisfy **PC** and $\mathbf{v}^{AP}$ satisfies **BR** stationarity, long-run equilibrium stationarity and positive correlation are guaranteed, and equilibrium stability of potential games holds in a long-run dynamic.

Theorem 2 (Long-run equilibrium stationarity and stability). *Consider a long-run dynamic $\mathbf{v} : \Delta^{AP} \times \mathbb{R}^{AP} \rightarrow \mathbb{R}^{AP}$ from subdynamics $\mathbf{v}^A$, $\mathbf{v}^{\mathcal{P}}$, and $\mathbf{v}^{AP}$. Assume $\lambda_{AP} > 0$.*

- i) *Suppose that all the three subdynamics $\mathbf{v}^A$, $\mathbf{v}^{\mathcal{P}}$, and $\mathbf{v}^{AP}$ satisfy **PC**. Then $\mathbf{v}$ satisfies **PC**.*

ii) In addition, suppose that (at least) $\mathbf{v}^{\mathcal{A}^P}$ satisfies **BR** stationarity. Then the long-run dynamic $\mathbf{v}$ satisfies the long-run equilibrium stationarity in any structured game $\tilde{\mathbf{F}}$:

$$\dot{\mathbf{x}}^P = \mathbf{0} \quad \Longleftrightarrow \quad \mathbf{x}^P \text{ is a long-run equilibrium.}$$

Therefore, the long-run dynamic $\mathbf{v}$ is admissible and thus satisfies the long-run equilibrium stability in any structured potential games.

The proof of Theorem 2 is relegated to Supplementary Methods 1.7. An immediate implication of Theorem 2 is that we can allow agents to adopt revision protocols that violate BR stationarity such as an imitative protocol for $\mathbf{v}^{\mathcal{A}}$ and $\mathbf{v}^P$ and still have a well-behaved long-run dynamic to guarantee long-run equilibrium stationarity and equilibrium stability in potential games. Hence, our long-run dynamic encompasses the dynamic considered in Boyd and Richerson (2009), who consider an imitative protocol for revision in actions. Theorem 2 shares some similarities with Sandholm (2010, Theorem 5.7.1), which suggests that when mixing two dynamics into a hybrid one by allowing agents to randomly use two different protocols, BR stationarity of either one dynamic is enough to guarantee PC and BR stationarity of the hybrid dynamic as long as both dynamics satisfy PC.

Strictly concave potential function

Theorem 1 and 2 establish the links between the evolutionary dynamics and the predictions of the traditional equilibrium analysis. In what follows, we provide methods for finding medium-run and long-run equilibria in two special cases: the potential function is either strictly concave or strictly convex. Note that the methods for finding the two types of equilibria are identical, except that additional constraints are imposed to keep the population distribution m^P when looking for the medium-run equilibria.

Theorem 3. Consider a structured game with a continuous and strictly concave potential function $\tilde{f}$. Then,

1. A local maximizer of $\tilde{f}$ must be the global maximizer.
2. There exists a unique global maximizer of the potential function.
3. If $\tilde{f}$ is continuously differentiable, the solution of Karush-Kuhn-Tucker condition is the unique global maximizer of $\tilde{f}$.

These are straightforward applications of well-known theorems on concave optimization: see Sundaram (1996, Sec.7.3).

Strictly convex potential function

We call a social state *pure strategy* if agents in each population choose the same action. Formally, in a pure strategy social state, in each population $p \in \mathcal{P}$, there is an action $a^p \in \mathcal{A}$ such that $x_{a^p}^p = m^p$ and $x_a^p = 0$ for any $a \neq a^p$. Denote by $\mathcal{E}^P \subset \Delta^{AP}$ the set of pure strategy social states. Let $E^A := \{\mathbf{e} \in \Delta^A \mid \exists i \in \mathcal{A} \ e_i = 1 \text{ and } e_j = 0 \text{ for any } j \neq i\}$. Then, $\mathcal{E}^P = \times_{p \in \mathcal{P}} m^p E^A$.

Theorem 4. Consider a structured game with a strictly convex potential function $\tilde{f}$. Local maximizers of the potential function exist and must be in $\mathcal{E}^P$.

The proof of Theorem 4 is provided in Supplementary Methods 1.7. One technical implication of the theorem is that, once the potential function $\tilde{f}$ is numerically defined, one can determine the asymptotically stable states by comparing the values of the potential at the pure strategy states (only finitely many: AP number of such states).

Theorems 3 and 4 provide us with useful methods for locating the medium-run and the long-run equilibria that are local maximizers of the potential function when the potential function is either strictly concave or strictly convex. Recall that the shape of the potential function is determined by the definiteness of $\tilde{\mathbf{\Pi}}$, which can be decomposed to those of $\mathbf{\Pi}$ and $\mathbf{G}$. For example, suppose that $\mathbf{\Pi}$ is negative definite, then we have

$\mathbf{G}$: positive definite $\Rightarrow \tilde{\mathbf{\Pi}}$: negative definite $\Rightarrow \tilde{f}$: strictly concave;

$\mathbf{G}$: negative definite $\Rightarrow \tilde{\mathbf{\Pi}}$: positive definite $\Rightarrow \tilde{f}$: strictly convex.

When $\tilde{\mathbf{\Pi}}$ is indefinite, however, we can no longer apply the above described methods for finding stable equilibria. In the following sections, we show how to find stable equilibria when $\tilde{\mathbf{\Pi}}$ is indefinite in three canonical examples.

1.4 Example 1

In Example 1, the base game is a binary coordination game with action set $\mathcal{A} = \{A, B\}$ and payoff matrix $\mathbf{\Pi}$ such as

$$\mathbf{\Pi} := \begin{pmatrix} \Pi_{AA} & \Pi_{AB} \\ \Pi_{BA} & \Pi_{BB} \end{pmatrix} = \begin{pmatrix} a & 0 \\ 0 & a \end{pmatrix},$$

where $a > 0$. The base game has two monomorphic Nash equilibria $\mathbf{x} = (1, 0)$, $\mathbf{x} = (0, 1)$ and one mixed Nash equilibria $\mathbf{x} = (0.5, 0.5)$. The game is played in the society consisting of two populations (groups) $\mathcal{P} = \{1, 2\}$ as in Fig.1a in the main text, with $g_{12} = g_{21} = -\beta < 0$.

Let us first consider the medium-run predictions of the game by fixing the population distribution $m^{\mathcal{P}} = (m^1, m^2)$. The medium-run equilibria are categorized as follows:

Type 1 The uniform replication of the base-game Nash equilibrium: $\mathbf{x}_A^{\mathcal{P}} = (0.5m^1, 0.5m^2)$.

Type 2 Border equilibria, which exist when $\beta \leq 1$: if $m^1 \geq \frac{\beta}{1+\beta}$, $\mathbf{x}_A^{\mathcal{P}} = (\frac{m^1+\beta m^2}{2}, m^2)$ and $\mathbf{x}_A^{\mathcal{P}} = (\frac{m^1-\beta m^2}{2}, 0)$ are equilibria; if $m^1 \leq \frac{1}{1+\beta}$, $\mathbf{x}_A^{\mathcal{P}} = (m^1, \frac{m^2+\beta m^1}{2})$ and $\mathbf{x}_A^{\mathcal{P}} = (0, \frac{m^2-\beta m^1}{2})$ are equilibria.

Type 3 Monomorphic equilibria, which exist when $\beta \leq 1$ and $m^1 \in [\frac{\beta}{1+\beta}, \frac{1}{1+\beta}]$: $x_A^{\mathcal{P}} = (0, 0)$; $x_A^{\mathcal{P}} = (m^1, m^2)$.

Type 4 Polarized equilibria: $x_A^{\mathcal{P}} = (m^1, 0)$; $x_A^{\mathcal{P}} = (0, m^2)$.

Under any admissible medium-run dynamic, Type-1 and Type-2 equilibria are unstable, Type 3 equilibria can be stable, and Type 4 equilibria are always stable:

Theorem 5. Consider a binary pure coordination game played in two homophilic populations with an arbitrary $\beta > 0$. With $\mathbf{m}^{\mathcal{P}}$ fixed, in any admissible medium-run dynamic,

- a) Type 1 and 2 equilibria are unstable;
- b) Type 3 equilibria are asymptotically stable when $\beta < 1$ and $m^1 \in (\frac{\beta}{1+\beta}, \frac{1}{1+\beta})$;
- c) Type 4 equilibria are always asymptotically stable.

Next, we consider the long-run predictions under the long-run dynamics.

Theorem 6. *Consider a binary pure coordination game played in two homophilic populations with $\beta > 0$. Under any admissible long-run dynamic, $\mathbf{x}_A^{\mathcal{P}} = (1, 0)$, $\mathbf{x}_A^{\mathcal{P}} = (0, 1)$, $\mathbf{x}_A^{\mathcal{P}} = (0, 0)$ with $\mathbf{m}^{\mathcal{P}} = (1, 0)$, and $\mathbf{x}_A^{\mathcal{P}} = (0, 0)$ with $\mathbf{m}^{\mathcal{P}} = (0, 1)$, are the only asymptotically stable long-run equilibria.*

In Supplementary Methods 1.8, we provide the verifications for these three types of equilibria as well as the proofs for Theorems 5 and 6.

1.5 Example 2

In Example 2, the base game is a binary anti-coordination game with action set $\mathcal{A} = \{I, O\}$ and payoff matrix $\mathbf{\Pi}$ such as

$$\mathbf{\Pi} := \begin{pmatrix} \Pi_{II} & \Pi_{IO} \\ \Pi_{OI} & \Pi_{OO} \end{pmatrix} = \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix},$$

where $a, b > 0$. We assume $a + b = 1$ without loss of generality. Note that $\Pi_{IO} = \Pi_{OI} = 0$ implies that $\mathbf{\Pi}$ is symmetric. The base game has a unique mixed strategy Nash equilibrium: $(x_I, x_O) = (b, a)$. The game is played in a society consisting of overlapping communities $\mathcal{P} = \{Lo, LR, oR\}$ as in Fig.1b in the main text. Let us first consider the medium-run predictions of the game by fixing the population distribution $\mathbf{m}^{\mathcal{P}} = (m^{Lo}, m^{LR}, m^{oR})$.

The medium-run equilibria are categorized as follows:

Type 1 The uniform replication of the base-game Nash equilibrium:

$$\mathbf{x}_I^{\mathcal{P}} = (bm^{Lo}, bm^{LR}, bm^{oR}).$$

Type 2 $x_I^{LR} = \min\{b, m^{LR}\}$, and $x_I^p = \max\{0, bm^p - am^{LR}\}$ for $p \in \{Lo, oR\}$.

Type 2' $x_O^{LR} = \min\{a, m^{LR}\}$, and $x_O^p = \max\{0, am^p - bm^{LR}\}$ for $p \in \{Lo, oR\}$.

In Type 2 and Type 2' equilibria, some agents in population LR choose a different action from the majority of the society. For example, when $m^{LR} \leq b$ and $m^{Lo}, m^{oR} \leq m^{LR}a/b$, the Type 2 equilibrium is $\mathbf{x}_I^{\mathcal{P}} = (0, m^{LR}, 0)$. In this equilibrium, all agents in population Lo and oR choose O , while all agents in population LR choose I .⁸

Under any admissible medium-run dynamics, the Type-1 equilibrium is unstable unless $m^{LR} = 0$, while Types 2 and 2' equilibria are always stable.

Theorem 7. *Consider a binary anti-coordination game played in overlapping communities. Under any admissible medium-run dynamic,*

a) *the Type 1 equilibrium corresponds to a saddle point of the potential function and thus it is unstable, unless $m^{LR} = 0$.*

b) *Type 2 and Type 2' equilibria are asymptotically stable.*

⁸Notice that the aggregate distribution of actions in this equilibrium is different from that in the Type-1 equilibrium: fewer players choose action I in the overall society because $\sum_p x_I^p = m^{LR} \leq b$. This implies that the best response for the agents in population LR is action I. At the same time, the proportion of agents choosing action O in community L at this equilibrium is smaller than that in the Type-1 equilibrium because $\hat{x}_O^{Lo}/(m^{Lo} + m^{LR}) = (1 + m^{LR}/m^{Lo})^{-1} \leq (1 + b/a)^{-1} = a$ by $a + b = 1$. Hence, the best response for the agents in population Lo is action O. Similarly, the best response for the agents in population oR is action O as well.

Note that if $m^{LR} = 0$ and thus no agent is affiliated with both the two communities, then all the three types of medium-run equilibria reduce to an identical social state $\mathbf{x}^P = (bm^{Lo}, am^{Lo}; 0, 0; bm^{oR}, am^{oR})$. That is, the two communities are completely segregated and the mixed equilibrium in the base game is simply played in each community.

When migration is allowed, there are only four long-run equilibria:

- a) Split replica equilibrium: $\mathbf{x}_I^P = (0.5b, 0, 0.5b)$ with $\mathbf{m}^P = (0.5, 0, 0.5)$;
- b) Center-integrated equilibrium: $x_I^P = (0, b, 0)$ with $\mathbf{m}^P = (0, 1, 0)$; and
- c) Population-wise diversified equilibria: $x_I^P = (0, b/(1+a), 0)$ with $\mathbf{m}^P = (a/(1+a), b/(1+a), a/(1+a))$, and $x_O^P = (0, a/(1+a), 0)$ with $\mathbf{m}^P = (b/(1+b), a/(1+b), b/(1+b))$;

The next theorem shows that the split replica equilibrium is the only stable long-run equilibrium under any admissible long-run dynamics. It implies that although the agents in Lo and oR could be connected through those in the intersection LR in the medium run, they eventually segregate themselves into two isolated communities of equal size in the long run when migration is allowed.

Theorem 8. *Consider a binary anti-coordination game played in overlapping communities. Under any admissible long-run dynamic, $\mathbf{x}^P = (0.5b, 0.5a; 0, 0; 0.5b, 0.5a)$ is the only local (and thus global) potential maximizer and thus globally asymptotically stable.*

In Supplementary Methods 1.9, we provide the verification for these three types of equilibria as well as the proofs for Theorem 7 and 8.

1.6 Example 3

In Example 3, the base game is a bilingual game with action set $\mathcal{A} = \{Ao, AB, oB\}$ and payoff matrix $\mathbf{\Pi}$ such as

$$\mathbf{\Pi} = \begin{pmatrix} a & a-c & d \\ a-c & a-c & b-c \\ d & b-c & b \end{pmatrix},$$

where $a > b > d+c$ (i.e., $b-c > d$) and $c, d > 0$. In the base game, each of the pure strategy social states like $x_{Ao} = 1$ constitutes a Nash equilibrium; in particular, those with $x_{Ao} = 1$ or $x_{oB} = 1$ are strict equilibria. Besides, there is a mixed-strategy Nash equilibrium $(x_{Ao}, x_{AB}, x_{oB}) = (0, c/(a-b+c), (a-b)/(a-b+c))$. The game is played in the overlapping populations such as in Fig.1b in the main text.

While the game has many medium-run equilibria, the stable ones are narrowed down as in the next theorem. They must be pure strategy social states: in each population, all agents take the same action (possibly different over populations). Below, we represent actions taken in the three populations in a pure strategy state by a triple. For example, (Ao, AB, oB) represents a pure strategy state with Ao in Lo, AB in LR and oB in oR. Similarly, among many long-run equilibria, we find that stable ones must be monomorphic, i.e., all agents in the society take the same action.

Theorem 9 (Medium-run stable equilibria). *Under any admissible medium-run dynamics, the following social states are stable.*

- i) *Monomorphic states where all agents take Ao regardless of their affiliations, under any $\mathbf{m}^P$.*
- ii) *Monomorphic states where all agents take oB regardless of their affiliations, under any $\mathbf{m}^P$.*

iii) Any pure strategy state in which agents in each population take the same action but possibly different across populations, if $m^{LR} = 0$.

iv) State (Ao, Ao, oB) , if $\frac{a-c-d}{c}m^{LR} < m^{oR} < \frac{c}{b-d}$; state (oB, Ao, Ao) , if $\frac{a-c-d}{c}m^{LR} < m^{Lo} < \frac{c}{b-d}$.

v) State (Ao, AB, oB) , if $\max\left\{\frac{c}{b-c-d}m^{Lo}, \frac{a-b}{c}m^{LR}\right\} < m^{oR} < \frac{a-c-d}{c}m^{Lo} + \frac{a-b}{c}m^{LR}$; state (oB, AB, Ao) , if $\max\left\{\frac{c}{b-c-d}m^{oR}, \frac{a-b}{c}m^{LR}\right\} < m^{Lo} < \frac{a-c-d}{c}m^{oR} + \frac{a-b}{c}m^{LR}$.

vi) State (AB, oB, oB) , if $cm^{LR} < (a-b)m^{Lo} < c(m^{LR} + m^{oR})$; State (oB, oB, AB) , if $cm^{LR} < (a-b)m^{oR} < c(m^{LR} + m^{Lo})$.

Theorem 10 (Long-run stable equilibria). *Under any admissible long-run dynamics, the following two connected sets of social states are locally stable.*

i) Monomorphic state where all agents take Ao regardless of their affiliations, with $m^{Lo} = 0$ or $m^{oR} = 0$.

ii) Monomorphic state where all agents take oB regardless of their affiliations, with $m^{Lo} = 0$ or $m^{oR} = 0$.

SI Supplementary Methods 1.10 provides the proofs of these theorems, as well as the complete identification of the medium-run and the long-run equilibria.

1.7 Proofs for General analysis

1.7.1 Proof of Theorem 2

i) First, at any $\mathbf{x}^P, \pi^P$, PC of each subdynamic $\mathcal{S}$ guarantees that $\mathbf{v}^S(\mathbf{x}^P, \pi^P) \cdot \pi^P \geq 0$. Taking a sum of these terms over $\mathcal{S} \in \{\mathcal{A}, \mathcal{P}, \mathcal{AP}\}$, we obtain

$$\mathbf{v}(\mathbf{x}^P, \pi^P) \cdot \pi^P = \sum_{\mathcal{S} \in \{\mathcal{A}, \mathcal{P}, \mathcal{AP}\}} \lambda_{\mathcal{S}} \mathbf{v}^S(\mathbf{x}^P, \pi^P) \cdot \pi^P \geq 0.$$

Nest, assume $\mathbf{v}(\mathbf{x}^P, \pi^P) \neq \mathbf{0}$. Since $\mathbf{v}(\mathbf{x}^P, \pi^P) = \sum_{\mathcal{S} \in \{\mathcal{A}, \mathcal{P}, \mathcal{AP}\}} \lambda_{\mathcal{S}} \mathbf{v}^S(\mathbf{x}^P, \pi^P)$, there always exists at least one of $\mathcal{S} \in \{\mathcal{A}, \mathcal{P}, \mathcal{AP}\}$ such that $\lambda_{\mathcal{S}} > 0$ and $\mathbf{v}^S(\mathbf{x}^P, \pi^P) \neq \mathbf{0}$. PC of this subdynamic $\mathcal{S}$ implies $\mathbf{v}^S(\mathbf{x}^P, \pi^P) \cdot \pi^P > 0$. With PC of the other subdynamics, it follows that

$$\mathbf{v}(\mathbf{x}^P, \pi^P) \cdot \pi^P \geq \lambda_{\mathcal{S}} \mathbf{v}^S(\mathbf{x}^P, \pi^P) \cdot \pi^P > 0.$$

ii) First of all, notice that PC implies *weak BR stationarity*:

$$\forall (a, p) \in \mathcal{A} \times \mathcal{P} [x_a^p > 0 \Rightarrow (a, p) \in \arg\max_{(a', p') \in \mathcal{S}(a, p)} \pi_{a'}^{p'}] \implies \mathbf{v}^S(\mathbf{x}^P, \pi^P) = \mathbf{0}.$$

For this, observe that the former condition implies $x_a^p(\pi_{a'}^{p'} - \pi_a^p) \leq 0$ for any $(a, p), (a', p') \in \mathcal{A} \times \mathcal{P}$. Therefore, we have

$$\mathbf{v}^S(\mathbf{x}^P, \pi^P) \cdot \pi^P = \sum_{(a, p) \in \mathcal{A} \times \mathcal{P}} x_a^p \sum_{(a', p') \in \mathcal{S}(a, p)} \rho_{(a, p), (a', p')}^S(\mathbf{x}^P, \pi^P) (\pi_{a'}^{p'} - \pi_a^p) \leq 0.$$

By PC, this implies $\mathbf{v}^S(\mathbf{x}^P, \pi^P) \cdot \pi^P = 0$ and thus $\mathbf{v}^S(\mathbf{x}^P, \pi^P) = \mathbf{0}$.

Next, we prove that the long-run equilibrium stationarity of the long-run dynamic $\mathbf{v}$. First, assume that $\mathbf{x}^{\mathcal{P}}$ is a long-run equilibrium. Then, we have

$$\forall (a, p) \in \mathcal{A} \times \mathcal{P} [x_a^p > 0 \Rightarrow (a, p) \in \operatorname{argmax}_{(a', p') \in \mathcal{S}(a, p)} \pi_{a'}^{p'}]$$

for each of the three subdynamics $\mathcal{S} \in \{\mathcal{A}, \mathcal{P}, \mathcal{AP}\}$. Since they satisfy PC and thus weak BR stationarity, it is followed by $\mathbf{v}^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) = \mathbf{0}$. Therefore,

$$\mathbf{v}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) = \sum_{\mathcal{S} \in \{\mathcal{A}, \mathcal{P}, \mathcal{AP}\}} \lambda_{\mathcal{S}} \mathbf{v}^{\mathcal{S}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) = \mathbf{0}.$$

Second, assume that $\mathbf{x}^{\mathcal{P}}$ is not a long-run equilibrium. Then, there exists a strategy $(a, p) \in \mathcal{A} \times \mathcal{P}$ such that

$$x_a^p > 0 \quad \text{and} \quad (a, p) \notin \operatorname{argmax}_{(a', p') \in \mathcal{A} \times \mathcal{P}} \pi_{a'}^{p'}.$$

This implies $\mathbf{v}^{\mathcal{AP}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) \neq \mathbf{0}$ by BR stationarity of subdynamic $\mathcal{AP}$, and further $\mathbf{v}^{\mathcal{AP}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) \cdot \boldsymbol{\pi}^{\mathcal{P}} > 0$ by PC of this subdynamic. With PC of the other subdynamics and $\lambda_{\mathcal{AP}} > 0$, we have

$$\mathbf{v}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) \cdot \boldsymbol{\pi}^{\mathcal{P}} \geq \lambda_{\mathcal{AP}} \mathbf{v}^{\mathcal{AP}}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}}) \cdot \boldsymbol{\pi}^{\mathcal{P}} > 0.$$

Therefore, $\mathbf{v}(\mathbf{x}^{\mathcal{P}}, \boldsymbol{\pi}^{\mathcal{P}})$ cannot be $\mathbf{0}$.

1.7.2 Proof of Theorem 4

We prove that any $\mathbf{x}^{\mathcal{P}} \in \Delta \setminus \mathcal{E}^{\mathcal{P}}$ cannot be a local maximum. Notice that $\mathbf{x}^{\mathcal{P}} \in \Delta \setminus \mathcal{E}^{\mathcal{P}}$ implies that $x_{a^1}^p, x_{a^2}^p \in (0, m^p)$ at some population $p \in \mathcal{P}$ and any two actions $a^1, a^2 \in \mathcal{A}$, say $a^1 = 1, a^2 = 2$ without loss of generality; note that p is fixed. Let $\mathbf{x}_{-1,2}^p = (0, 0, x_3^p, \dots, x_A^p) \in \Delta^A$, and $\mathbf{e}_a^p \in E^A$ be an action distribution (A -dimensional vector) in which only the a -th entry equals 1 while other entries are zero. Then we can write $\mathbf{x}^p = \mathbf{x}_{-1,2}^p + x_1^p \mathbf{e}_1^p + x_2^p \mathbf{e}_2^p$.

For an arbitrary $\varepsilon > 0$, define two distinct perturbed action distributions $\mathbf{x}_{1,\varepsilon}^p$ by

$$\mathbf{x}_{1,\varepsilon}^p = \mathbf{x}^p + \varepsilon(\mathbf{e}_1^p - \mathbf{e}_2^p) \quad \text{and} \quad \mathbf{x}_{2,\varepsilon}^p = \mathbf{x}^p + \varepsilon(\mathbf{e}_2^p - \mathbf{e}_1^p).$$

As long as ε is sufficiently small such as $\varepsilon < \min\{x_a^p, 1 - x_a^p : a = 1, 2\}$, they belong to Δ^A . Define $\mathbf{x}_{i,\varepsilon}^{\mathcal{P}} = (\mathbf{x}_{i,\varepsilon}^p, \mathbf{x}^{-p})$ by gathering this perturbed action distribution in population p and the original action distributions of $\mathbf{x}^{\mathcal{P}}$ in all the other populations.

Since $\mathbf{x}_{1,\varepsilon}^{\mathcal{P}} \neq \mathbf{x}_{2,\varepsilon}^{\mathcal{P}} = \mathbf{x}^{\mathcal{P}}$ and

$$0.5\mathbf{x}_{1,\varepsilon}^{\mathcal{P}} + 0.5\mathbf{x}_{2,\varepsilon}^{\mathcal{P}} = \mathbf{x}^{\mathcal{P}},$$

the strict convexity of $\tilde{f}$ implies

$$0.5\tilde{f}(\mathbf{x}_{1,\varepsilon}^{\mathcal{P}}) + 0.5\tilde{f}(\mathbf{x}_{2,\varepsilon}^{\mathcal{P}}) > \tilde{f}(\mathbf{x}^{\mathcal{P}}).$$

This strict inequality implies that at least one of $\tilde{f}(\mathbf{x}_{1,\varepsilon}^{\mathcal{P}})$ and $\tilde{f}(\mathbf{x}_{2,\varepsilon}^{\mathcal{P}})$ must attain a greater value than $\tilde{f}(\mathbf{x}^{\mathcal{P}})$. This holds for any $\varepsilon > 0$ and thus such a perturbed state $\mathbf{x}_{i,\varepsilon}^{\mathcal{P}}$ can be found anywhere close to $\mathbf{x}^{\mathcal{P}}$. Therefore, $\mathbf{x}^{\mathcal{P}} \in \Delta \setminus \mathcal{E}^{\mathcal{P}}$ cannot attain a local maximum of $\tilde{f}$. Note that the perturbed state $\mathbf{x}_{i,\varepsilon}^{\mathcal{P}}$ keeps the same population distribution m^p as $\mathbf{x}^{\mathcal{P}}$ and thus belongs to $m^p \Delta^A \subsetneq \Delta^{AP}$; the change from $\mathbf{x}^{\mathcal{P}}$ to $\mathbf{x}_{i,\varepsilon}^{\mathcal{P}}$ is

feasible both in the medium run and in the long run. Thus, this theorem applies to both the medium-run maximization of $\tilde{f}$ on $m^{\mathcal{P}}\Delta^A$ (with fixed $m^{\mathcal{P}}$) and the long-run maximization on Δ^{AP} .

1.8 Proofs for Example 1

First of all, we verify each type of the medium-run equilibria.

Type-1. To have a mixed equilibria in which both populations contain agents playing each of the actions, the following indifference conditions need to be satisfied:

$$\begin{aligned}\tilde{F}_A^1(\mathbf{x}^{\mathcal{P}}) &= ax_A^1 - \beta ax_A^2 = F_B^1(\mathbf{x}^{\mathcal{P}}) = ax_B^1 - \beta ax_B^2, \\ \tilde{F}_A^2(\mathbf{x}^{\mathcal{P}}) &= ax_A^2 - \beta ax_A^1 = F_B^2(\mathbf{x}^{\mathcal{P}}) = ax_B^2 - \beta ax_B^1.\end{aligned}$$

Together with the constraints that $x_A^1 + x_B^1 = m^1$ and $x_A^2 + x_B^2 = m^2$, we get $x_A^1 = 0.5m^1$ and $x_A^2 = 0.5m^2$.

Type-2. Consider a border equilibrium in the form of $x_A^1 \in (0, m^1), x_A^2 = m^2$. We should have

$$\begin{aligned}\tilde{F}_A^1(\mathbf{x}^{\mathcal{P}}) &= a(x_A^1 - \beta m^2) = \tilde{F}_B^1(\mathbf{x}^{\mathcal{P}}) = ax_B^1, \\ \tilde{F}_A^2(\mathbf{x}^{\mathcal{P}}) &= a(m^2 - \beta x_A^1) \geq \tilde{F}_B^2(\mathbf{x}^{\mathcal{P}}) = -a\beta x_B^1.\end{aligned}$$

These implies that

$$x_A^1 = \frac{m^1 + \beta m^2}{2} \leq \frac{m^2 + \beta m^1}{2\beta}.$$

The above inequality holds when $\beta \leq 1$. Also, since we have the constraint that $\frac{m^1 + \beta m^2}{2} = x_A^1 \leq m^1$, we also need $m^1 \geq \frac{\beta}{1+\beta}$. The verifications for the other three border equilibria follow the same logics.

Type-3. Consider the monomorphic equilibrium $\mathbf{x}_A^{\mathcal{P}} = (0, 0)$. It should satisfy

$$\begin{aligned}\tilde{F}_A^1(\mathbf{x}^{\mathcal{P}}) &= 0 \leq \tilde{F}_B^1(\mathbf{x}^{\mathcal{P}}) = am^1 - \beta am^2, \\ \tilde{F}_A^2(\mathbf{x}^{\mathcal{P}}) &= 0 \leq \tilde{F}_B^2(\mathbf{x}^{\mathcal{P}}) = am^2 - \beta am^1.\end{aligned}$$

Hence, we need $\beta \leq 1$ and $m^1 \in [\frac{\beta}{1+\beta}, \frac{1}{1+\beta}]$. The same conditions are needed for $\mathbf{x}_A^{\mathcal{P}} = (m^1, m^2)$ to be an equilibrium.

Type-4. Consider the polarized equilibrium $\mathbf{x}_A^{\mathcal{P}} = (m^1, 0)$. It should satisfy

$$\begin{aligned}\tilde{F}_A^1(\mathbf{x}^{\mathcal{P}}) &= am^1 \geq \tilde{F}_B^1(\mathbf{x}^{\mathcal{P}}) = -\beta am^2, \\ \tilde{F}_A^2(\mathbf{x}^{\mathcal{P}}) &= -\beta am^1 \leq \tilde{F}_B^2(\mathbf{x}^{\mathcal{P}}) = am^2.\end{aligned}$$

The two conditions always hold. Similar for $\mathbf{x}_A^{\mathcal{P}} = (0, m^2)$.

1.8.1 Proof of Theorem 5

The potential function can be written as follows:

$$\begin{aligned}\tilde{f}(\mathbf{x}^P) &= 0.5a((x_A^1)^2 - 2\beta x_A^1 x_A^2 + (x_A^2)^2 + (x_B^1)^2 - 2\beta x_B^1 x_B^2 + (x_B^2)^2) \\ &= 0.5a((x_A^1)^2 - 2\beta x_A^1 x_A^2 + (x_A^2)^2 + (m^1 - x_A^1)^2 - 2\beta(m^1 - x_A^1)(m^2 - x_A^2) + (m^2 - x_A^2)^2).\end{aligned}$$

With $x_B^P = m^P - x_A^P$, we can regard $\tilde{f}$ as a function of $\mathbf{x}_A^P = (x_A^1, x_A^2)$; we abuse notation to denote by $\tilde{f}(\mathbf{x}_A^P)$ the two-variable function given in the second line of the above equation.

When $\beta < 1$, both $\mathbf{\Pi}$ and $\mathbf{G}$ are positive definite, so $\tilde{\mathbf{\Pi}}$ is positive definite, implying that $\tilde{f}$ is strictly convex. According to Theorem 4, we only need to check the four vertices to find local maximizer of $\tilde{\mathbf{\Pi}}$.

We first check $\mathbf{x}_A^P = (m^1, 0)$. Since we have

$$\frac{\partial \tilde{f}}{\partial x_A^1}(m^1, 0) = a(m^1 + \beta m^2) > 0, \quad \frac{\partial \tilde{f}}{\partial x_A^2}(m^1, 0) = -a(\beta m^1 + m^2) < 0,$$

$\mathbf{x}_A^P = (m^1, 0)$ is an isolated local maximizer of $\tilde{f}$. Similarly, $\mathbf{x}_A^P = (0, m^2)$ is also an isolated local maximizer of $\tilde{f}$.

Next, we check $\mathbf{x}_A^P = (m^1, m^2)$. Since we have

$$\begin{aligned}\frac{\partial \tilde{f}}{\partial x_A^1}(m^1, m^2) &= a(m^1 - \beta m^2) > 0, \quad \text{if } m^1 > \beta m^2, \text{ i.e., } m^1 > \frac{\beta}{1+\beta}, \\ \frac{\partial \tilde{f}}{\partial x_A^2}(m^1, m^2) &= -a(\beta m^1 - m^2) > 0, \quad \text{if } m^2 > \beta m^1 \text{ i.e., } m^1 < \frac{1}{1+\beta},\end{aligned}$$

$\mathbf{x}_A^P = (m^1, m^2)$ is an isolated local maximizer if $m^1 \in (\frac{\beta}{1+\beta}, \frac{1}{1+\beta})$. Similarly, $\mathbf{x}_A^P = (0, 0)$ is an isolated local maximizer if $m^1 \in (\frac{\beta}{1+\beta}, \frac{1}{1+\beta})$.

Next, let us consider $\beta = 1$. In this case, the potential function can be simplified as

$$\tilde{f}(\mathbf{x}_A^P) = 0.5a((x_A^1 - x_A^2)^2 + (m^1 - m^2 - (x_A^1 - x_A^2))^2).$$

The only local maximizers are $\mathbf{x}_A^P = (m^1, 0)$ and $\mathbf{x}_A^P = (0, m^2)$.

Finally, consider $\beta > 1$; notice that $\mathbf{G}$ is indefinite in this case. $\mathbf{x}_A^P = (m^1, 0)$ and $\mathbf{x}_A^P = (0, m^2)$ are still isolated local maximizers by the same logic as in the case when $\beta < 1$. However, $m^1 > \beta m^2$ and $m^2 > \beta m^1$ cannot hold at the same time, implying that $\mathbf{x}_A^P = (m^1, m^2)$ and $\mathbf{x}_A^P = (0, 0)$ are not local maximizers. An immediate implication of indefiniteness of $\tilde{\mathbf{\Pi}}$ is that any interior Nash equilibrium of the game is a saddle point of $\tilde{f}$. Hence, we do not have an interior local maximizers. In addition, since border equilibria only exist when $\beta \leq 1$, we do not have border local maximizer in this case.

1.8.2 Proof of Theorem 6

First, observe that the asymptotically stable long-run equilibria must be a subset of the asymptotically stable medium-run equilibria. Hence, according to Theorem 5, Type 1 and Type 2 equilibria cannot be long-run asymptotically stable states. Second, let us look at Type 3 and Type 4 equilibria. Let $\mathbf{x}^P = (x_A^1, x_B^1; x_A^2, x_B^2)$.

Type 3 equilibria taking the form of $\mathbf{x}^P = (0, m^1; 0, m^2)$ constitute a continuous range of social states in $\mathcal{X}^P$. Hence, the asymptotically stable long-run equilibria must yield the greatest value of the potential

function among these equilibria whose m^P vary within the range. The potential reduces to

$$\tilde{f}(\mathbf{x}^P) = 0.5a((m^1)^2 - 2\beta m^1 m^2 + (m^2)^2),$$

which is maximized at $m^1 = 0$ or $m^2 = 0$. Similarly, the local maximum corresponding to the range of Type 3 equilibria $\mathbf{x}_A^P = (m^1, m^2)$ is attained at $m^1 = 0$ or $m^2 = 0$.

Type 4 equilibria $\mathbf{x}^P = (m^1, 0; 0, m^2)$ constitute a continuous range of social states in $\mathcal{X}^P$. Hence, a long-run asymptotically stable state must be the case-wise maximizer of the potential function on this range. The potential on this range is given by

$$\tilde{f}(\mathbf{x}^P) = 0.5a((m^1)^2 + (m^2)^2),$$

which is maximized at $m^1 = 0$ or $m^2 = 0$. Similarly, the local maximum corresponding to the range of Type 4 equilibria $\mathbf{x}^P = (0, m^1; m^2, 0)$ is attained at $m^1 = 0$ or $m^2 = 0$.

An asymptotically stable long-run equilibrium maximizes the potential in its neighborhood on Δ^{AP} . When finding the above long-run equilibria, we maximize the potential among each type of medium-run equilibria while varying m^P continuously in Δ^P . Since the asymptotically stable medium-run equilibria are local maximizers of the potential function in $m^P \Delta^A \subset \Delta^{AP}$ when m^P is fixed, each of the above long-run equilibria maximizes the potential function locally by allowing m^P to change locally. Thus, they are indeed the asymptotically stable long-run equilibria.

To conclude, the local maximizers we found above are asymptotically stable under the long-run BRD. They correspond to four long-run equilibria: $\mathbf{x}_A^P = (1, 0)$, $\mathbf{x}_A^P = (0, 1)$, $\mathbf{x}_A^P = (0, 0)$ with $m^P = (1, 0)$, $\mathbf{x}_A^P = (0, 0)$ with $m^P = (0, 1)$.

1.9 Proofs for Example 2

1.9.1 Proof of Theorem 7

First of all, we verify each type of the medium-run equilibria. Note that for an agent in population $p \in \{Lo, LR, oR\}$, I is a best response if and only if $-(1-b)\tilde{x}_I^p \geq -b\tilde{x}_O^p$, i.e., $\tilde{x}_I^p \leq b(\tilde{x}_I^p + \tilde{x}_O^p)$. Therefore, we obtain the following equilibrium condition for the action distribution in each population:

$$\text{for } p = Lo: \quad \begin{array}{c} x_I^{Lo} + x_I^{LR} \\ < \\ = b(m^{Lo} + m^{LR}) \\ > \end{array} \iff x_I^{Lo} \begin{cases} = m^{Lo} \\ \in [0, m^{Lo}] \\ = 0 \end{cases} \quad (4a)$$

$$\text{for } p = LR: \quad \begin{array}{c} x_I^{Lo} + x_I^{LR} + x_I^{oR} \\ < \\ = b \\ > \end{array} \iff x_I^{LR} \begin{cases} = m^{LR} \\ \in [0, m^{LR}] \\ = 0 \end{cases} \quad (4b)$$

$$\text{for } p = oR: \quad \begin{array}{c} x_I^{LR} + x_I^{oR} \\ < \\ = b(m^{LR} + m^{oR}) \\ > \end{array} \iff x_I^{oR} \begin{cases} = m^{oR} \\ \in [0, m^{oR}] \\ = 0 \end{cases} \quad (4c)$$

Type-1. A uniform replication of the base game equilibrium is identified by $x_I^p = bm^p \in [0, 1]$ for each $p \in \{Lo, LR, Ro\}$ and thus satisfies all these equilibrium conditions (4).

Type-2. Type 2 equilibria can be categorized into the following cases:

$$\mathbf{x}_I = (x_I^{Lo}, x_I^{LR}, x_I^{oR})$$

$$= \begin{cases} (0, b, 0) & \text{if (i) } m^{LR} \geq b, \\ (0, m^{LR}, 0) & \text{if (ii) } m^{LR} < b \text{ and } m^{Lo}, m^{oR} \leq m^{LR}a/b, \\ (0, m^{LR}, bm^{oR} - am^{LR}) & \text{if (iii) } m^{LR} < b \text{ and } m^{Lo} \leq m^{LR}a/b < m^{oR}, \\ (bm^{Lo} - am^{LR}, m^{LR}, 0) & \text{if (iii') } m^{LR} < b \text{ and } m^{oR} \leq m^{LR}a/b < m^{Lo}, \\ (bm^{Lo} - am^{LR}, m^{LR}, bm^{oR} - am^{LR}) & \text{if (iv) } m^{LR} < b \text{ and } m^{Lo}, m^{oR} > m^{LR}a/b. \end{cases}$$

Case i). Consider the case when $m^{LR} \geq b$. The equilibrium is identified by $x_I^{LR} = b \in [0, m^{LR}]$ and $x_I^{Lo} = x_I^{oR} = 0$, which satisfies all the equilibrium conditions (4).

Case ii). Consider the case when $m^{LR} < b$ and $m^{Lo}, m^{oR} \leq m^{LR}a/b$. The equilibrium is identified by $x_I^{LR} = m^{LR}$ and $x_I^{Lo} = x_I^{oR} = 0$. Since $m^{Lo} \leq m^{LR}a/b$ and $a + b = 1$, we have $m^{LR} = (a + b)m^{LR} \geq bm^{Lo} + bm^{LR}$ and thus

$$x_I^{Lo} + x_I^{LR} = m^{LR} \geq b(m^{Lo} + m^{LR})$$

and thus (4a) is confirmed. Similarly, we can confirm (4c) from $m^{oR} \leq m^{LR}a/b$ and $a + b = 1$. Since $m^{LR} < b$, we have

$$x_I^{Lo} + x_I^{LR} + x_I^{oR} = m^{LR} < b.$$

Hence, (4b) is confirmed.

Case iii). Consider the case when $m^{LR} < b$ and $m^{Lo} \leq m^{LR}a/b < m^{oR}$. The equilibrium is identified by $x_I^{Lo} = 0$, $x_I^{LR} = m^{LR}$ and $x_I^{oR} = bm^{oR} - am^{LR} \in (0, m^{oR})$. Similarly to case ii), (4a) is confirmed from $m^{Lo} \leq m^{LR}a/b$ and $a + b = 1$. Since $a + b = 1$, we have

$$x_I^{LR} + x_I^{oR} = m^{LR} + bm^{oR} - am^{LR} = b(m^{LR} + m^{oR})$$

and thus we verify that (4c) is satisfied.⁹ Further, we have

$$x_I^{Lo} + x_I^{LR} + x_I^{oR} = 0 + b(m^{LR} + m^{oR}) \leq b$$

and thus (4b) is confirmed.

Case iii'). The verification is similar to that of Case iii).

Case iv). Consider the case when $m^{LR} < b$ and $m^{LR}a/b < m^{Lo}, m^{oR}$. The equilibrium is identified by $x_I^{Lo} = bm^{Lo} - am^{LR} \in (0, m^{Lo})$, $x_I^{LR} = m^{LR}$ and $x_I^{oR} = bm^{oR} - am^{LR} \in (0, m^{oR})$. Similarly to the verification of (4c) in Case iii), we can verify that this action distribution satisfies (4a) and (4c) from $a + b = 1$. Since $m^{Lo} + m^{LR} + m^{oR} = 1$ and $a + b = 1$, we have

$$x_I^{Lo} + x_I^{LR} + x_I^{oR} = bm^{Lo} - am^{LR} + m^{LR} + bm^{oR} - am^{LR} = b - am^{LR} \leq b$$

⁹When $m^{LR}a/b < m^{oR}$, (4c) cannot be satisfied with $x_I^{oR} = 0$.

and thus (4b) is confirmed.

Type 2'. The verifications are similar to those for Type 2 equilibria.

Proof of Theorem 7. a) This is an immediate implication of indefiniteness of $\tilde{\mathbf{\Pi}}$, which equals $D^2\tilde{f}$. It implies that, if an interior point satisfies $D\tilde{f} = 0$, it is a saddle point of $\tilde{f}$. The Type-1 equilibrium indeed satisfies $D\tilde{f} = 0$.

b) For each case of Type 2 equilibria, we specify a feasible transition vector $\mathbf{z}^{\mathcal{P}} \in \mathbb{R}^{AP}$ of the action distribution and then confirm that any (sufficiently small) feasible transition only decreases the value of $\tilde{f}$ from the one attained at the equilibrium. Hence, the equilibrium is a locally strict maximizer and thus asymptotically stable under the medium-run dynamic. The proof of stability for Type 2' equilibria is similar and thus omitted.

To keep the population distribution fixed, any transition vector $\mathbf{z}^{\mathcal{P}}$ has to take a form of

$$\begin{aligned} \mathbf{z}^{\mathcal{P}} &:= (z_A^{Lo}, z_B^{Lo}; z_A^{LR}, z_B^{LR}; z_A^{oR}, z_B^{oR}) = (\varepsilon^{Lo}, -\varepsilon^{Lo}; \varepsilon^{LR}, -\varepsilon^{LR}; \varepsilon^{oR}, -\varepsilon^{oR}), \\ \text{i.e., } \mathbf{z}^{\mathcal{P}} &:= (z_A^p, z_B^p) = \varepsilon^p(1, -1) \text{ for each } p \in \{Lo, LR, oR\}. \end{aligned} \quad (5)$$

Each equilibrium has a different feasible range of $(\varepsilon^{Lo}, \varepsilon^{LR}, \varepsilon^{oR})$. Let $\mathbf{x}^{\mathcal{P}} = (x_A^{Lo}, x_B^{Lo}; x_A^{LR}, x_B^{LR}; x_A^{oR}, x_B^{oR})$.

Note that, if the social state moves from $\mathbf{x}^{\mathcal{P}}$ to $\mathbf{x}^{\mathcal{P}} + \mathbf{z}^{\mathcal{P}}$, then the value of the potential $\tilde{f}$ changes as

$$\tilde{f}(\mathbf{x}^{\mathcal{P}} + \mathbf{z}^{\mathcal{P}}) = \tilde{f}(\mathbf{x}^{\mathcal{P}}) + 0.5(2\mathbf{x}^{\mathcal{P}} + \mathbf{z}^{\mathcal{P}}) \cdot \tilde{\mathbf{\Pi}}\mathbf{z}^{\mathcal{P}}.$$

Therefore, $\tilde{f}(\mathbf{x}^{\mathcal{P}}) \geq \tilde{f}(\mathbf{x}^{\mathcal{P}} + \mathbf{z}^{\mathcal{P}})$ if and only if

$$\Delta\tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) := (2\mathbf{x}^{\mathcal{P}} + \mathbf{z}^{\mathcal{P}}) \cdot \tilde{\mathbf{\Pi}}\mathbf{z}^{\mathcal{P}} \leq 0.$$

This term can be decomposed as

$$\Delta\tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) = \begin{pmatrix} 2\mathbf{x}^{Lo} + \mathbf{z}^{Lo} \\ 2\mathbf{x}^{LR} + \mathbf{z}^{LR} \\ 2\mathbf{x}^{oR} + \mathbf{z}^{oR} \end{pmatrix} \cdot \begin{pmatrix} \mathbf{\Pi} & \mathbf{\Pi} & 0 \\ \mathbf{\Pi} & \mathbf{\Pi} & \mathbf{\Pi} \\ 0 & \mathbf{\Pi} & \mathbf{\Pi} \end{pmatrix} \begin{pmatrix} \mathbf{z}^{Lo} \\ \mathbf{z}^{LR} \\ \mathbf{z}^{oR} \end{pmatrix} = \Delta_1\tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) + \Delta_2\tilde{f}(\mathbf{z}^{\mathcal{P}}),$$

where

$$\begin{aligned} \Delta_1\tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) &= 2 \{ (\mathbf{x}^{Lo} + \mathbf{x}^{LR}) \cdot \mathbf{\Pi}\mathbf{z}^{Lo} + (\mathbf{x}^{Lo} + \mathbf{x}^{LR} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi}\mathbf{z}^{LR} + (\mathbf{x}^{LR} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi}\mathbf{z}^{oR} \}, \\ \Delta_2\tilde{f}(\mathbf{z}^{\mathcal{P}}) &= \mathbf{z}^{Lo} \cdot \mathbf{\Pi}\mathbf{z}^{Lo} + \mathbf{z}^{LR} \cdot \mathbf{\Pi}\mathbf{z}^{LR} + \mathbf{z}^{oR} \cdot \mathbf{\Pi}\mathbf{z}^{oR} + 2(\mathbf{z}^{Lo} + \mathbf{z}^{oR}) \cdot \mathbf{\Pi}\mathbf{z}^{LR}. \end{aligned}$$

With $\mathbf{z}^p = \varepsilon^p(1, -1)$, the latter reduces to

$$\begin{aligned} \Delta_2\tilde{f}(\mathbf{z}^{\mathcal{P}}) &= \sum_p \varepsilon^p \begin{pmatrix} 1 \\ -1 \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^p + 2(\varepsilon^{Lo} + \varepsilon^{oR}) \begin{pmatrix} 1 \\ -1 \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{LR} \\ &= -\sum_p \varepsilon^{p^2} - 2(\varepsilon^{Lo} + \varepsilon^{oR})\varepsilon^{LR} \quad (\because a + b = 1) \end{aligned}$$

Case i). From the equilibrium $\mathbf{x}^P = (0, m^{Lo}; b, m^{LR} - b; 0, m^{oR})$, a feasible transition vector in the form of (5) must satisfy

$$\varepsilon^{Lo} \in [0, m^{Lo}] \subset \mathbb{R}_+, \varepsilon^{LR} \in (-b, m^{LR} - b), \varepsilon^{oR} \in [0, m^{oR}] \subset \mathbb{R}_+.$$

Then, we have

$$\begin{aligned} (\mathbf{x}^{Lo} + \mathbf{x}^{LR}) \cdot \mathbf{\Pi z}^{Lo} &= \begin{pmatrix} b \\ m^{Lo} + m^{LR} - b \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{Lo} \\ &= \varepsilon^{Lo} b(-a + m^{Lo} + m^{LR} - b) \\ &= -\varepsilon^{Lo} b m^{oR} \quad (\because a + b = 1, \sum_p m^p = 1); \\ (\mathbf{x}^{Lo} + \mathbf{x}^{LR} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi z}^{LR} &= \begin{pmatrix} b \\ 1 - b \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{LR} \\ &= \varepsilon^{LR} b(-a + (1 - b)) = 0 \quad (\because a + b = 1); \\ (\mathbf{x}^{LR} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi z}^{oR} &= \begin{pmatrix} b \\ m^{LR} + m^{oR} - b \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{oR} \\ &= \varepsilon^{oR} b(-a + m^{LR} + m^{oR} - b) \\ &= -\varepsilon^{oR} b m^{Lo} \quad (\because a + b = 1, \sum_p m^p = 1). \end{aligned}$$

$$\therefore \Delta_1 \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) = -2(\varepsilon^{Lo} b m^{oR} + \varepsilon^{oR} b m^{Lo}).$$

Combining this with $\Delta_2 \tilde{f}$, we obtain

$$\Delta \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) = -\sum_p \varepsilon^{p2} - 2\varepsilon^{Lo}(\varepsilon^{LR} + b m^{oR}) - 2\varepsilon^{oR}(\varepsilon^{LR} + b m^{Lo}).$$

Since $\varepsilon^{Lo}, \varepsilon^{oR} \geq 0$ in a feasible transition vector, the last two terms are nonpositive if $\varepsilon^{LR} \geq -b \min\{m^{Lo}, m^{oR}\}$. Thus, as long as $m^{Lo}, m^{oR} > 0$, there is a sufficiently small neighborhood of $\mathbf{0}$ such that $\Delta \tilde{f} \leq -\sum_p \varepsilon^{p2} < 0$ is guaranteed for any non-zero feasible transition vector $\mathbf{z}^P$ in the neighborhood. If $m^{oR} = 0$, then ε^{oR} must be zero in any feasible transition vector. Thus,

$$\Delta \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) = -\varepsilon^{Lo2} - \varepsilon^{LR2} - 2\varepsilon^{Lo}\varepsilon^{LR} = -(\varepsilon^{Lo} + \varepsilon^{LR})^2 < 0$$

for any non-zero feasible transition vector $\mathbf{z}^P \neq \mathbf{0}$. Similarly, we can verify that $m^{Lo} = 0$ also implies $\Delta \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) = -(\varepsilon^{LR} + \varepsilon^{oR})^2 < 0$ for any non-zero feasible transition vector $\mathbf{z}^P \neq \mathbf{0}$. Therefore, in case i), the medium-run equilibrium attains a local strict maximum of $\tilde{f}$.

Case ii). From the equilibrium $\mathbf{x}^P = (0, m^{Lo}; m^{LR}, 0; 0, m^{oR})$, a feasible transition vector in the form of (5) must satisfy

$$\varepsilon^{Lo} \in [0, m^{Lo}] \subset \mathbb{R}_+, \varepsilon^{LR} \in [-m^{LR}, 0] \subset \mathbb{R}_-, \varepsilon^{oR} \in [0, m^{oR}] \subset \mathbb{R}_+.$$

Further, we have

$$\begin{aligned}
(\mathbf{x}^{Lo} + \mathbf{x}^{LR}) \cdot \mathbf{\Pi z}^{Lo} &= \begin{pmatrix} m^{LR} \\ m^{Lo} \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{Lo} \\
&= \varepsilon^{Lo}(-am^{LR} + bm^{Lo}) \\
&\leq 0 \quad (\because \varepsilon^{Lo} \geq 0, bm^{Lo} \leq am^{LR} \text{ in case ii}); \\
(\mathbf{x}^{Lo} + \mathbf{x}^{LR} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi z}^{LR} &= \begin{pmatrix} m^{LR} \\ 1 - m^{LR} \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{LR} \\
&= \varepsilon^{LR}(-am^{LR} + b(1 - m^{LR})) = \varepsilon^{LR}(b - m^{LR}) \quad (\because a + b = 1) \\
(\mathbf{x}^{LR} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi z}^{oR} &= \begin{pmatrix} m^{LR} \\ m^{oR} \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{oR} \\
&= \varepsilon^{oR}(-am^{LR} + bm^{oR}) \\
&\leq 0 \quad (\because \varepsilon^{oR} \geq 0, bm^{oR} \leq am^{LR} \text{ in case ii}). \\
\therefore \Delta_1 \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) &\leq 2\varepsilon^{LR}(b - m^{LR}).
\end{aligned}$$

Combining this with $\Delta_2 \tilde{f}$, we find

$$\Delta \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) \leq -(\varepsilon^{Lo} + \varepsilon^{LR})^2 - (\varepsilon^{LR} + \varepsilon^{oR})^2 + \varepsilon^{LR}(\varepsilon^{LR} + 2(b - m^{LR})).$$

Since $\varepsilon^{LR} \leq 0$, the last term is nonpositive if $\varepsilon^{LR} \geq -2(b - m^{LR})$. This is guaranteed for a sufficiently small nonpositive $\varepsilon^{LR} \in (-2(b - m^{LR}), 0]$. (Note that $b \geq m^{LR}$ in case ii.) If $\varepsilon^{LR} \in (-2(b - m^{LR}), 0)$, then we have $0 > \varepsilon^{LR}(\varepsilon^{LR} + 2(b - m^{LR})) \geq \Delta \tilde{f}(\mathbf{x}^P, \mathbf{z}^P)$. If $\varepsilon^{LR} = 0$, then $\Delta \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) \leq -\varepsilon^{Lo^2} - \varepsilon^{oR^2} < 0$ because ε^{Lo} or ε^{oR} cannot be zero at the same time given that $\varepsilon^{LR} = 0$ and $\mathbf{z}^P \neq 0$. Therefore, there is a sufficiently small neighborhood of $\mathbf{0}$ such that $\Delta \tilde{f} < 0$ is guaranteed for any non-zero feasible transition vector $\mathbf{z}^P$ in the neighborhood. This implies that in case ii), the medium-run equilibrium attains a local strict maximum of $\tilde{f}$.

Case iii). From the equilibrium $\mathbf{x}^P = (0, m^{Lo}; m^{LR}, 0; bm^{oR} - am^{LR}, a(m^{LR} + m^{oR}))$, a feasible transition vector in the form of (5) must satisfy

$$\varepsilon^{Lo} \in [0, m^{Lo}] \subset \mathbb{R}_+, \varepsilon^{LR} \in [-m^{LR}, 0] \subset \mathbb{R}_-, \varepsilon^{oR} \in [-(bm^{oR} - am^{LR}), a(m^{LR} + m^{oR})].$$

Further, we have

$$\begin{aligned}
(\mathbf{x}^{Lo} + \mathbf{x}^{LR}) \cdot \mathbf{\Pi z}^{Lo} &= \begin{pmatrix} m^{LR} \\ m^{Lo} \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{Lo} \\
&= \varepsilon^{Lo}(-am^{LR} + bm^{Lo}) \\
&\leq 0 \quad (\because \varepsilon^{Lo} \geq 0, bm^{Lo} \leq am^{LR} \text{ in case iii}); \\
(\mathbf{x}^{Lo} + \mathbf{x}^{LR} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi z}^{LR} &= \begin{pmatrix} b(m^{LR} + m^{oR}) \\ m^{LR} + a(m^{LR} + m^{oR}) \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{LR} \\
&= \varepsilon^{LR}\{-ab(m^{LR} + m^{oR}) + b(m^{oR} + a(m^{LR} + m^{oR}))\} = \varepsilon^{LR}bm^{Lo} \quad (\because a + b = 1);
\end{aligned}$$

$$\begin{aligned}
(\mathbf{x}^{LR} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi} \mathbf{z}^{oR} &= (m^{LR} + m^{oR}) \begin{pmatrix} b \\ a \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{oR} \\
&= \varepsilon^{oR} (m^{LR} + m^{oR}) (-ab + ab) = 0. \\
\therefore \Delta_1 \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) &\leq 2\varepsilon^{LR} b m^{Lo}.
\end{aligned}$$

Combining this with $\Delta_2 \tilde{f}$, we find

$$\Delta \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) \leq -\sum_p \varepsilon^{p2} + 2\varepsilon^{LR} (b m^{Lo} - \varepsilon^{oR} - \varepsilon^{Lo}).$$

Since $\varepsilon^{LR} \leq 0$, the last term is nonpositive if $\varepsilon^{oR} + \varepsilon^{Lo} \leq b m^{Lo}$. Thus, as long as $m^{Lo} > 0$, there is a sufficiently small neighborhood of $\mathbf{0}$ such that $\Delta \tilde{f} \leq -\sum_p \varepsilon^{p2} < 0$ is guaranteed for any non-zero feasible transition vector $\mathbf{z}^P$ in the neighborhood. If $m^{Lo} = 0$, then ε^{Lo} must be zero in any transition vector. Thus,

$$\Delta \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) = -\varepsilon^{LR2} - \varepsilon^{oR2} - 2\varepsilon^{LR} \varepsilon^{oR} = -(\varepsilon^{LR} + \varepsilon^{oR})^2 < 0$$

for any non-zero feasible transition vector $\mathbf{z}^P \neq \mathbf{0}$. Therefore, in case iii), the medium-run equilibrium attains a local strict maximum of $\tilde{f}$. The proof of Case iii') is similar to that of Case iii) and thus omitted.

Case iv). From the equilibrium $\mathbf{x}^P = (b m^{Lo} - a m^{LR}, a(m^{Lo} + m^{LR}); m^{LR}, 0; b m^{oR} - a m^{LR}, a(m^{LR} + m^{oR}))$, a feasible transition vector in the form of (5) must satisfy

$$\varepsilon^{Lo} \in [-(b m^{Lo} - a m^{LR}), a(m^{Lo} + m^{LR})], \varepsilon^{LR} \in [-m^{LR}, 0] \subset \mathbb{R}_-, \varepsilon^{oR} \in [-(b m^{oR} - a m^{LR}), a(m^{LR} + m^{oR})].$$

Further, we have

$$\begin{aligned}
(\mathbf{x}^{Lo} + \mathbf{x}^{LR}) \cdot \mathbf{\Pi} \mathbf{z}^{Lo} &= (m^{Lo} + m^{LR}) \begin{pmatrix} b \\ a \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{Lo} \\
&= \varepsilon^{Lo} (m^{Lo} + m^{LR}) (-ab + ab) = 0; \\
(\mathbf{x}^{Lo} + \mathbf{x}^{LR} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi} \mathbf{z}^{LR} &= \begin{pmatrix} b - a m^{LR} \\ a + a m^{LR} \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{LR} \\
&= \varepsilon^{LR} \{-a(b - a m^{LR}) + b(a + a m^{LR})\} = a \varepsilon^{LR} m^{LR} \quad (\because a + b = 1); \\
(\mathbf{x}^{LR} + \mathbf{x}^{oR}) \cdot \mathbf{\Pi} \mathbf{z}^{oR} &= (m^{LR} + m^{oR}) \begin{pmatrix} b \\ a \end{pmatrix} \cdot \begin{pmatrix} -a & 0 \\ 0 & -b \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \varepsilon^{oR} \\
&= \varepsilon^{oR} (m^{LR} + m^{oR}) (-ab + ab) = 0. \\
\therefore \Delta_1 \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) &= 2a \varepsilon^{LR} m^{LR}.
\end{aligned}$$

Combining this with $\Delta_2 \tilde{f}$, we find

$$\Delta \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) \leq -\sum_p \varepsilon^{p2} + 2\varepsilon^{LR} (a m^{LR} - \varepsilon^{Lo} - \varepsilon^{oR}).$$

Since $\varepsilon^{LR} \leq 0$, the last term is nonnegative if $\varepsilon^{Lo} + \varepsilon^{oR} \leq a m^{LR}$. Thus, as long as $m^{LR} > 0$, there is a sufficiently small neighborhood of $\mathbf{0}$ such that $\Delta \tilde{f} \leq -\sum_p \varepsilon^{p2} < 0$ is guaranteed for any non-zero feasible

transition vector $\mathbf{z}^P$ in the neighborhood. If $m^{LR} = 0$, then ε^{LR} must be zero in any transition vector. Thus,

$$\Delta \tilde{f}(\mathbf{x}^P, \mathbf{z}^P) = -\varepsilon^{Lo^2} - \varepsilon^{oR^2} < 0$$

for any non-zero feasible transition vector $\mathbf{z}^P$. Therefore, in case iv), the medium-run equilibrium attains a local strict maximum of $\tilde{f}$. $\square$

1.9.2 Proof of Theorem 8

First, we prove that the four equilibria mentioned in the main text are the only long-run equilibria. For this, we use the fact that a long-run equilibrium $\mathbf{x}^P$ must be a medium-run equilibrium with the population distribution fixed at the one implied by $\mathbf{x}^P$. Further, in each medium-run equilibrium, an agent in each population is indifference between used actions, i.e., the actions that are taken by a positive mass of agents. Therefore, by comparing agents' payoffs from a used action over populations at each type of the medium-run equilibria, we can identify the long-run equilibria.

Type 1. A type-1 medium-run equilibrium $\mathbf{x}^P$ takes a form of $\mathbf{x}^P = (m^{Lo}b, m^{Lo}a; m^{LR}b, m^{LR}a; m^{oR}b, m^{oR}a)$ with some $m^P \in \Delta^P$. Thus, the payoff from a used action in each population is

$$\begin{aligned}\tilde{F}_I^{Lo}(\mathbf{x}^P) &= \tilde{F}_O^{Lo}(\mathbf{x}^P) = -ab(m^{Lo} + m^{LR}), \\ \tilde{F}_I^{LR}(\mathbf{x}^P) &= \tilde{F}_O^{LR}(\mathbf{x}^P) = -ab(m^{Lo} + m^{LR} + m^{oR}) = -ab, \\ \tilde{F}_I^{oR}(\mathbf{x}^P) &= \tilde{F}_O^{oR}(\mathbf{x}^P) = -ab(m^{LR} + m^{oR}).\end{aligned}$$

First, we find a long-run equilibrium of this type with $m^{LR} > 0$. Then, the long-run equilibrium condition requires that

$$\tilde{F}_I^{LR}(\mathbf{x}^P) = -ab(m^{Lo} + m^{LR} + m^{oR}) \geq -ab(m^{Lo} + m^{LR}) = \tilde{F}_I^{Lo}(\mathbf{x}^P),$$

which is equivalent to $m^{oR} = 0$ by $a, b > 0$ and $m^P \geq 0$. Similarly, $m^{Lo} = 0$ is implied from $\tilde{F}_I^{LR}(\mathbf{x}^P) \geq \tilde{F}_I^{oR}(\mathbf{x}^P)$. Therefore, $\mathbf{x}^P = (0, 0; b, a; 0, 0)$ is the only candidate state for such a long-run equilibrium. By plugging this candidate $\mathbf{x}^P$ into the above payoff calculation, we can confirm that $\tilde{F}_a^p(\mathbf{x}^P) = -ab$ for all a, p and thus this $\mathbf{x}^P$ is indeed in a long-run equilibrium.

Now, we search for a type-1 long-run equilibrium with $m^{LR} = 0$. At least one of m^{Lo} and m^{oR} must be positive. Assume $m^{Lo} > 0$. Then, the long-run equilibrium condition requires that

$$\tilde{F}_I^{Lo}(\mathbf{x}^P) = -ab(m^{Lo} + m^{LR}) = -abm^{Lo} \geq -abm^{oR} = -ab(m^{LR} + m^{oR}) = \tilde{F}_I^{oR}(\mathbf{x}^P),$$

which is equivalent to $m^{oR} \geq m^{Lo} > 0$ by $a, b > 0$ and $m^P \geq 0$. Similarly, we can obtain $m^{Lo} > 0$ from $m^{oR} > 0$. Therefore, $m^{LR} = 0$ implies $m^{Lo}, m^{oR} > 0$. In a long-run equilibrium, this implies that the above inequality holds with equality; hence, $m^{oR} = m^{Lo}$. That is, $\mathbf{x}^P = (0.5b, 0.5a; 0, 0; 0.5b, 0.5a)$ is the only candidate state for such a long-run equilibrium. By plugging this candidate $\mathbf{x}^P$ into the above payoff calculation, we can confirm that $\tilde{F}_I^{Lo}(\mathbf{x}^P) = \tilde{F}_I^{oR}(\mathbf{x}^P) = -0.5ab \geq \tilde{F}_I^{LR}(\mathbf{x}^P) = -ab$ and thus this $\mathbf{x}^P$ is indeed in a long-run equilibrium.

Type 2, case i). A type-2-i) medium-run equilibrium $\mathbf{x}^P$ takes a form of $\mathbf{x}^P = (0, m^{Lo}; b, m^{LR} - b; 0, m^{oR})$ with some $m^P \in \Delta^P$. Thus, the payoff from a used action in each population is

$$\begin{aligned}\tilde{F}_O^{Lo}(\mathbf{x}^P) &= -b(m^{Lo} + m^{LR} - b), \\ \tilde{F}_I^{LR}(\mathbf{x}^P) &= \tilde{F}_O^{LR}(\mathbf{x}^P) = -b(m^{Lo} + m^{LR} + m^{oR} - b) = -ab, \\ \tilde{F}_O^{oR}(\mathbf{x}^P) &= -ab(m^{LR} + m^{oR} - b).\end{aligned}$$

The condition for case i) of type 2 implies $m^{LR} \geq b > 0$. Then, the long-run equilibrium condition requires that

$$\tilde{F}_I^{LR}(\mathbf{x}^P) = -b(m^{Lo} + m^{LR} + m^{oR} - b) \geq -b(m^{Lo} + m^{LR} - b) = \tilde{F}_O^{Lo}(\mathbf{x}^P),$$

which is equivalent to $m^{oR} = 0$ by $a, b > 0$ and $m^P \geq 0$. Similarly, $m^{Lo} = 0$ is implied from $\tilde{F}_I^{LR}(\mathbf{x}^P) \geq \tilde{F}_O^{oR}(\mathbf{x}^P)$. Therefore, $\mathbf{x}^P = (0, 0; b, 1 - b; 0, 0)$ is the only candidate state for such a long-run equilibrium. Since $a = 1 - b$, this is the one that we have found when considering the type-1 long-run equilibria with $m^{LR} > 0$; so it is indeed a long-run equilibrium.

Type 2, case ii). A type-2 case-ii) medium-run equilibrium $\mathbf{x}^P$ takes a form of $\mathbf{x}^P = (0, m^{Lo}; m^{LR}, 0; 0, m^{oR})$ with some $m^P \in \Delta^P$. Thus, the payoff from a used action in each population is

$$\begin{aligned}\tilde{F}_O^{Lo}(\mathbf{x}^P) &= -bm^{Lo}, \\ \tilde{F}_I^{LR}(\mathbf{x}^P) &= -am^{LR}, \\ \tilde{F}_O^{oR}(\mathbf{x}^P) &= -bm^{oR}.\end{aligned}$$

The condition for case ii) of type 2 implies $m^{LR} > 0$; otherwise, $m^{Lo} = m^{LR} = m^{oR} = 0$ by $m^{Lo}, m^{oR} \leq m^{LR}a/b$. Then, the long-run equilibrium condition requires that

$$\tilde{F}_I^{LR}(\mathbf{x}^P) = -am^{LR} \geq -b(m^{Lo} + m^{LR} - b) = \tilde{F}_O^{Lo}(\mathbf{x}^P),$$

which is equivalent to $m^{Lo} \geq m^{LR}a/b$ by $a, b > 0$ and $m^P \geq 0$. With the aforementioned case-ii) condition, this implies $m^{Lo} = m^{LR}a/b$. Similarly, we obtain $m^{oR} = m^{LR}a/b$ from $\tilde{F}_I^{LR}(\mathbf{x}^P) \geq \tilde{F}_O^{oR}(\mathbf{x}^P)$. These imply $m^P = (a/(1+a), b/(1+a), a/(1+a))$ by $2a+b=1+a$. Therefore, $\mathbf{x}^P = (0, a/(1+a); b/(1+a), 0; 0, a/(1+a))$ is the only candidate state for such a long-run equilibrium. By plugging this candidate $\mathbf{x}^P$ into the above payoff calculation, we can confirm that $\tilde{F}_O^{Lo}(\mathbf{x}^P) = \tilde{F}_I^{LR}(\mathbf{x}^P) = \tilde{F}_O^{oR}(\mathbf{x}^P) = -ab/(1+a)$ and thus this $\mathbf{x}^P$ is indeed in a long-run equilibrium.

Type 2, case iii), iii'). A type-2 case-iii) medium-run equilibrium $\mathbf{x}^P$ takes a form of $\mathbf{x}^P = (0, m^{Lo}; m^{LR}, 0; bm^{oR} - am^{LR}, a(m^{LR} + m^{oR}))$ with some $m^P \in \Delta^P$. Thus, the payoff from a used action in each population is

$$\begin{aligned}\tilde{F}_O^{Lo}(\mathbf{x}^P) &= -bm^{Lo}, \\ \tilde{F}_I^{LR}(\mathbf{x}^P) &= -a\{m^{LR} + (bm^{oR} - am^{LR})\} = -ab(m^{LR} + m^{oR}), \\ \tilde{F}_I^{oR}(\mathbf{x}^P) &= \tilde{F}_O^{oR}(\mathbf{x}^P) = -ab(m^{LR} + m^{oR}).\end{aligned}$$

The condition for case iii) of type 2 implies $m^{oR} > m^{LR}a/b \geq 0$. Then, the long-run equilibrium condition requires that

$$\tilde{F}^{oR}(\mathbf{x}^P) = -ab(m^{LR} + m^{oR}) \geq -bm^{Lo} = \tilde{F}_O^{Lo}(\mathbf{x}^P),$$

which is equivalent to $m^{Lo} \geq a(m^{LR} + m^{oR})$ by $a, b > 0$. With the condition for case iii) of type 2, $m^{Lo} \leq m^{LR}a/b < m^{oR}$, this implies $m^{Lo} \geq a(m^{LR} + m^{oR}) > am^{Lo} + bm^{Lo} = m^{Lo}$ by $a + b = 1$; it is a contradiction. Therefore, no long-run equilibrium exists in case iii) of type 2.

By the same token, we can conclude that no long-run equilibrium exists in case iii') of type 2.

Type 2, case iv). A type-2 case-iv) medium-run equilibrium $\mathbf{x}^P$ takes a form of $\mathbf{x}^P = (bm^{Lo} - am^{LR}, a(m^{Lo} + m^{LR}); m^{LR}, 0; bm^{oR} - am^{LR}, a(m^{LR} + m^{oR}))$ with some $m^P \in \Delta^P$. Thus, the payoff from a used action in each population is

$$\begin{aligned}\tilde{F}_I^{Lo}(\mathbf{x}^P) &= \tilde{F}_O^{Lo}(\mathbf{x}^P) = -ab(m^{Lo} + m^{LR}), \\ \tilde{F}_I^{LR}(\mathbf{x}^P) &= -a\{(bm^{Lo} - am^{LR}) + (bm^{oR} - am^{LR})\} = -a(b - am^{LR}), \\ \tilde{F}_I^{oR}(\mathbf{x}^P) &= \tilde{F}_O^{oR}(\mathbf{x}^P) = -ab(m^{LR} + m^{oR}).\end{aligned}$$

The condition for case iv) of type 2 implies $m^{Lo}, m^{oR} > m^{LR}a/b \geq 0$. Then, the long-run equilibrium condition requires that

$$\tilde{F}_I^{Lo}(\mathbf{x}^P) = -ab(m^{Lo} + m^{LR}) = -ab(m^{LR} + m^{oR}) = \tilde{F}_I^{oR}(\mathbf{x}^P),$$

which is equivalent to $m^{Lo} = m^{oR}$ by $a, b > 0, m^{LR} \geq 0$. If $m^{LR} = 0$, the population distribution must be $m^P = (0.5, 0, 0.5)$; the social state reduces to $\mathbf{x}^P = (0.5b, 0.5a; 0, 0; 0.5b, 0.5a)$. This is a long-run equilibrium that we have found when considering the type-1 long-run equilibria with $m^{LR} = 0$; so it is indeed a long-run equilibrium.

Next, consider equilibria with $m^{LR} > 0$. Then the long-run equilibrium condition also requires

$$\tilde{F}_I^{LR}(\mathbf{x}^P) = -a(b - am^{LR}) = -ab(m^{LR} + m^{oR}) = \tilde{F}_I^{oR}(\mathbf{x}^P),$$

which is equivalent to $m^{oR} = m^{LR}a/b \geq 0$. These imply $m^P = (a/(1+a), b/(1+a), a/(1+a))$ by $2a+b = 1+a$. Therefore, $\mathbf{x}^P = (0, a/(1+a); b/(1+a), 0; 0, a/(1+a))$ is the only candidate state for such a long-run equilibrium. This is the one that we have found when considering the type-2 case-ii) long-run equilibria with $m^{LR} > 0$; so it is indeed a long-run equilibrium.

Type 2'. By the same token as type 2, we can prove that $\mathbf{x}^P = (0.5b, 0.5a; 0, 0; 0.5b, 0.5a)$, $(0, 0; b, a; 0, 0)$, and $(b/(1+b), 0; 0, a/(1+b); b/(1+b), 0)$ are the only long-run equilibria found from the medium-run equilibria of type 2'.

Proof of Theorem 8. As in the proof of Theorem 7, we calculate the change in potential $\tilde{f}$ by a feasible transition $\mathbf{z}^P \in \mathbb{R}^{AP}$ from each long-run equilibrium $\mathbf{x}^P$ and check if the potential cannot increase as long as $\mathbf{z}^P$ is feasible and sufficiently small. In contrast to a medium-run equilibrium, feasibility of transition vector $\mathbf{z}^P$ only requires $\mathbf{x}^P + \mathbf{z}^P$ to stay in Δ^{AP} in a long-run equilibrium, which allows migration and changes in population distribution m^P . We keep the same notation as in the proof of Theorem 7.

Note that $\sum_{p \in \mathcal{P}} \mathbf{z}^p = \mathbf{0}$ implies

$$\begin{aligned}\Delta_2 \tilde{f}(\mathbf{z}^{\mathcal{P}}) &= \sum_{p \in \mathcal{P}} \mathbf{z}^p \cdot \mathbf{\Pi} \mathbf{z}^p + 2(\mathbf{z}^{Lo} + \mathbf{z}^{oR}) \cdot \mathbf{\Pi} \mathbf{z}^{LR} \\ &= \mathbf{z}^{Lo} \cdot \mathbf{\Pi} \mathbf{z}^{Lo} + \mathbf{z}^{oR} \cdot \mathbf{\Pi} \mathbf{z}^{oR} - \mathbf{z}^{LR} \cdot \mathbf{\Pi} \mathbf{z}^{LR}.\end{aligned}$$

a) Consider a long-run equilibrium $\mathbf{x}^{\mathcal{P}} = (0.5b, 0.5a; 0, 0; 0.5b, 0.5a)$. Transition vector $\mathbf{z}^{\mathcal{P}} \in \mathbb{R}^{AP}$ is feasible as long as $\sum_{p \in \mathcal{P}} \mathbf{z}^p = \mathbf{0}$ and $z_A^{LR}, z_B^{LR} \geq 0$. Then,

$$\begin{aligned}\Delta_1 \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) &= 2 \left\{ \left(0.5 \begin{pmatrix} b \\ a \end{pmatrix} + \mathbf{0} \right) \cdot \mathbf{\Pi} \mathbf{z}^{Lo} + \begin{pmatrix} b \\ a \end{pmatrix} \cdot \mathbf{\Pi} \mathbf{z}^{LR} + \left(0.5 \begin{pmatrix} b \\ a \end{pmatrix} + \mathbf{0} \right) \cdot \mathbf{\Pi} \mathbf{z}^{oR} \right\} \\ &= \begin{pmatrix} b \\ a \end{pmatrix} \cdot \mathbf{\Pi} (\mathbf{z}^{Lo} + 2\mathbf{z}^{LR} + \mathbf{z}^{oR}) = \begin{pmatrix} b \\ a \end{pmatrix} \cdot \mathbf{\Pi} \mathbf{z}^{LR}.\end{aligned}$$

The last equality comes from $\sum_{p \in \mathcal{P}} \mathbf{z}^p = \mathbf{0}$. Therefore, we obtain

$$\begin{aligned}\Delta \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) &= \Delta_1 \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) + \Delta_2 \tilde{f}(\mathbf{z}^{\mathcal{P}}) \\ &= \mathbf{z}^{Lo} \cdot \mathbf{\Pi} \mathbf{z}^{Lo} + \mathbf{z}^{oR} \cdot \mathbf{\Pi} \mathbf{z}^{oR} + \left(\begin{pmatrix} b \\ a \end{pmatrix} - \mathbf{z}^{LR} \right) \cdot \mathbf{\Pi} \mathbf{z}^{LR} \\ &\leq -a(b - z_A^{LR})z_A^{LR} - b(a - z_B^{LR})z_B^{LR}.\end{aligned}$$

The inequality comes from the fact that $\mathbf{\Pi}$ is negative semidefinite. Since $z_A^{LR}, z_B^{LR} \geq 0$, the last term is nonpositive as long as $b > z_A^{LR}$ and $a > z_B^{LR}$. By $a, b > 0$, this implies that $\Delta \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}})$ is nonpositive as long as $\mathbf{z}^{\mathcal{P}}$ is feasible and sufficiently close to $\mathbf{0}$; that is, $\tilde{f}$ attains a local maximum in Δ^{AP} at $\mathbf{x}^{\mathcal{P}} = (0.5b, 0.5a; 0, 0; 0.5b, 0.5a)$.

b) Consider a long-run equilibrium $\mathbf{x}^{\mathcal{P}} = (0, 0; b, a; 0, 0)$. Transition vector $\mathbf{z}^{\mathcal{P}} \in \mathbb{R}^{AP}$ is feasible as long as $\sum_{p \in \mathcal{P}} \mathbf{z}^p = \mathbf{0}$ and $z_A^{Lo}, z_B^{Lo}, z_A^{oR}, z_B^{oR} \geq 0$. Then,

$$\begin{aligned}\Delta_1 \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) &= 2 \left\{ \left(\mathbf{0} + \begin{pmatrix} b \\ a \end{pmatrix} \right) \cdot \mathbf{\Pi} \mathbf{z}^{Lo} + \begin{pmatrix} b \\ a \end{pmatrix} \cdot \mathbf{\Pi} \mathbf{z}^{LR} + \left(\mathbf{0} + \begin{pmatrix} b \\ a \end{pmatrix} \right) \cdot \mathbf{\Pi} \mathbf{z}^{oR} \right\} \\ &= 2 \left\{ \begin{pmatrix} b \\ a \end{pmatrix} \cdot \mathbf{\Pi} (\mathbf{z}^{Lo} + \mathbf{z}^{LR} + \mathbf{z}^{oR}) \right\} = 0.\end{aligned}$$

The last equality comes from $\sum_{p \in \mathcal{P}} \mathbf{z}^p = \mathbf{0}$. Therefore, we obtain

$$\begin{aligned}\Delta \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) &= \Delta_1 \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) + \Delta_2 \tilde{f}(\mathbf{z}^{\mathcal{P}}) \\ &= 0 + \Delta_2 \tilde{f}(\mathbf{z}^{\mathcal{P}}) = \mathbf{z}^{Lo} \cdot \mathbf{\Pi} \mathbf{z}^{Lo} + \mathbf{z}^{oR} \cdot \mathbf{\Pi} \mathbf{z}^{oR} - \mathbf{z}^{LR} \cdot \mathbf{\Pi} \mathbf{z}^{LR}.\end{aligned}$$

For an arbitrary $\varepsilon \in (-b, a)$, let $\mathbf{z}^{\mathcal{P}} = (0, 0; \varepsilon, -\varepsilon; 0, 0)$. Then,

$$\Delta \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) = 0 + 0 - (-a - b)\varepsilon^2 = \varepsilon^2 > 0.$$

Since $a, b > 0$, this implies that, in any sufficiently small neighborhood of $\mathbf{0}$, there exists a transition vector $\mathbf{z}^{\mathcal{P}}$ that increases the value of potential $\tilde{f}$ from the one at $\mathbf{x}^{\mathcal{P}} = (0, 0; b, a; 0, 0)$. That is, this equilibrium is not a local maximum of the potential.

c) Consider a long-run equilibrium $\mathbf{x}^{\mathcal{P}} = (0, a/(1+a); b/(1+a), 0; 0, a/(1+a))$. Transition vector $\mathbf{z}^{\mathcal{P}} \in \mathbb{R}^{4P}$ is feasible as long as $\sum_{p \in \mathcal{P}} \mathbf{z}^p = \mathbf{0}$ and $z_A^{Lo}, z_B^{LR}, z_A^{oR} \geq 0$. Then,

$$\begin{aligned} \Delta_1 \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) &= 2 \left\{ \frac{1}{1+a} \begin{pmatrix} b \\ a \end{pmatrix} \cdot \Pi \mathbf{z}^{Lo} + \frac{1}{1+a} \begin{pmatrix} b \\ 2a \end{pmatrix} \cdot \Pi \mathbf{z}^{LR} + \frac{1}{1+a} \begin{pmatrix} b \\ a \end{pmatrix} \cdot \Pi \mathbf{z}^{oR} \right\} \\ &= \frac{2}{1+a} \left\{ \begin{pmatrix} b \\ a \end{pmatrix} \cdot \Pi (\mathbf{z}^{Lo} + \mathbf{z}^{LR} + \mathbf{z}^{oR}) + \begin{pmatrix} 0 \\ a \end{pmatrix} \cdot \Pi \mathbf{z}^{LR} \right\} \\ &= -2abz_B^{LR}/(1+a). \end{aligned}$$

Therefore, we obtain

$$\begin{aligned} \Delta \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) &= \Delta_1 \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) + \Delta_2 \tilde{f}(\mathbf{z}^{\mathcal{P}}) \\ &= -2abz_B^{LR}/(1+a) + \mathbf{z}^{Lo} \cdot \Pi \mathbf{z}^{Lo} + \mathbf{z}^{oR} \cdot \Pi \mathbf{z}^{oR} - \mathbf{z}^{LR} \cdot \Pi \mathbf{z}^{LR}. \end{aligned}$$

For an arbitrary $\varepsilon \in (0, -b/(1+a))$, let $\mathbf{z}^{\mathcal{P}} = (\varepsilon/2, 0; -\varepsilon, 0; \varepsilon/2, 0)$. Then,

$$\Delta \tilde{f}(\mathbf{x}^{\mathcal{P}}, \mathbf{z}^{\mathcal{P}}) = 0 - a(\varepsilon/2)^2 - a(\varepsilon/2)^2 + a\varepsilon^2 = a\varepsilon^2/4 > 0.$$

Since $a, b > 0$, this implies that, in any sufficiently small neighborhood of $\mathbf{0}$, there exists a transition vector $\mathbf{z}^{\mathcal{P}}$ that increases the value of potential $\tilde{f}$ from the one at $\mathbf{x}^{\mathcal{P}} = (0, a/(1+a); b/(1+a), 0; 0, a/(1+a))$. That is, this equilibrium is not a local maximum of the potential. Similarly, we can negate a local maximum at $\mathbf{x}^{\mathcal{P}} = (b/(1+b), 0; 0, a/(1+b); b/(1+b), 0)$. $\square$

1.10 Proofs for Example 3

First of all, note that, if the state changes from $\mathbf{x}^{\mathcal{P}}$ to $\mathbf{x}^{\mathcal{P}} + \mathbf{z}^{\mathcal{P}}$, then the potential changes as

$$\tilde{f}(\mathbf{x}^{\mathcal{P}} + \mathbf{z}^{\mathcal{P}}) = \tilde{f}(\mathbf{x}^{\mathcal{P}}) + \mathbf{z}^{\mathcal{P}} \cdot \tilde{\Pi} \mathbf{x}^{\mathcal{P}} + 0.5 \mathbf{z}^{\mathcal{P}} \cdot \tilde{\Pi} \mathbf{z}^{\mathcal{P}}. \quad (6)$$

Note that $\tilde{\mathbf{F}}^{\mathcal{P}}(\mathbf{x}^{\mathcal{P}}) = \tilde{\Pi} \mathbf{x}^{\mathcal{P}}$. If the initial state $\mathbf{x}^{\mathcal{P}}$ is an equilibrium (either in the medium run or in the long run), any feasible change in the state should yield $\mathbf{z}^{\mathcal{P}} \cdot \tilde{\Pi} \mathbf{x}^{\mathcal{P}} \leq 0$. With a straightforward extension of the domain of $\tilde{\mathbf{F}}$ to $\mathbb{R}^{4P}$ such as $\tilde{\mathbf{F}}^{\mathcal{P}}(\mathbf{z}^{\mathcal{P}}) = \tilde{\Pi} \mathbf{z}^{\mathcal{P}}$, we can rewrite the third term as

$$\mathbf{z}^{\mathcal{P}} \cdot \tilde{\Pi} \mathbf{z}^{\mathcal{P}} = \mathbf{z}^{\mathcal{P}} \cdot \tilde{\mathbf{F}}^{\mathcal{P}}(\mathbf{z}^{\mathcal{P}}) = \sum_{p \in \mathcal{P}} \mathbf{z}^p \cdot \tilde{\mathbf{F}}^p(\mathbf{z}^p). \quad (7)$$

The best response actions

Given $\tilde{\mathbf{x}}$, the payoff for an agent from each action is

$$\begin{array}{llll} F_{Ao}(\tilde{\mathbf{x}}) = a & \tilde{x}_{Ao} + & (a-c)\tilde{x}_{AB} & + d\tilde{x}_{oB}, \\ F_{AB}(\tilde{\mathbf{x}}) = (a-c) & \tilde{x}_{Ao} + & (a-c)\tilde{x}_{AB} & + (a-c)\tilde{x}_{oB}, \end{array}$$

$$F_{oB}(\tilde{\mathbf{x}}) = d \quad \tilde{x}_{Ao} + \quad (b - c)\tilde{x}_{AB} \quad + b\tilde{x}_{oB}.$$

By comparing these, we can find that the best response action for the agent is

$$\begin{array}{llllll} Ao & \text{if} & c\tilde{x}_{Ao} & \geq & (b - c - d)\tilde{x}_{oB}; \\ AB & \text{if} & c\tilde{x}_{Ao} & \leq & (b - c - d)\tilde{x}_{oB} \\ & \text{and} & (a - c - d)\tilde{x}_{Ao} & + (a - b) & \tilde{x}_{AB} \geq & c\tilde{x}_{oB}; \\ oB & \text{if} & (a - c - d)\tilde{x}_{Ao} & + (a - b) & \tilde{x}_{AB} \leq & c\tilde{x}_{oB}. \end{array} \quad (8)$$

Note that each action is a best response if all agents in the observed action distribution take this action; it is the only best response if it is Ao or oB. If all observed agents take AB, both Ao and AB are best responses.

It is possible that two actions—Ao and AB, or AB and OB—are best response actions and thus a population may mix the two actions when two of the above inequality conditions are satisfied with equality. However, at stable equilibria, agents' best response must be robust to (sufficiently small) perturbation of $\tilde{\mathbf{x}}$ in any feasible direction. Notice that, to have two best response actions, the observed action distribution must be mixed. Thus, the equality is easily broken by perturbation and then the best response becomes either one action. So, in stable equilibria, agents in each population must be taking one single action.

In search for stable equilibria

Based on the finding that each population must concentrated on either one action in a stable equilibrium, here we check each possible combination of actions over the three populations and see if the combination can constitute a stable medium-run or long-run equilibrium and if so how the agents should be distributed among the populations.

Monomorphic states First, we consider a monomorphic state where all agents are taking the same action across the populations. In the medium run, such a state is an equilibrium; in particular, if they are taking either Ao or oB, it is a strict equilibrium and thus stable. We check if it can be stable in the long run.

Let $\mathbf{x}^{\mathcal{P}}$ be a state, such that, for either one action $a \in \mathcal{A}$, all the mass of each population take this action, i.e., $x_a^{\mathcal{P}} = m^{\mathcal{P}}$ for each $\mathcal{P}$. The payoff from this action in each population is

$$\begin{aligned} \tilde{F}_a^{Lo}(\mathbf{x}^{\mathcal{P}}) &= \pi_{aa}(m^{Lo} + m^{LR}), \\ \tilde{F}_a^{LR}(\mathbf{x}^{\mathcal{P}}) &= \pi_{aa}(m^{Lo} + m^{LR} + m^{oR}), \\ \tilde{F}_a^{oR}(\mathbf{x}^{\mathcal{P}}) &= \pi_{aa}(m^{LR} + m^{oR}). \end{aligned}$$

Since $\pi_{aa} > 0$, there is an incentive to agglomerate into LR for a positive scale merit. Actually, if $m^{oR} > 0$, then $\tilde{F}_a^{Lo}(\mathbf{x}^{\mathcal{P}}) > \tilde{F}_a^{LR}(\mathbf{x}^{\mathcal{P}}$ and thus $m^{oR} = 0$ in a long-run equilibrium; then, $\tilde{F}_a^{Lo}(\mathbf{x}^{\mathcal{P}}) = \tilde{F}_a^{LR}(\mathbf{x}^{\mathcal{P}}$ since their observed action distributions become the same, i.e., $\tilde{\mathbf{x}}^{Lo} = \tilde{\mathbf{x}}^{LR} = \mathbf{e}_a$, and thus agents are indifferent between Lo and LR. So any population distribution over Lo and LR constitutes a long-run equilibrium. Similarly, $m^{Lo} > 0$ implies $m^{Lo} = 0$ in a long-run equilibrium and any distribution between LR and oR is a long-run equilibrium. In sum, if a monomorphic state satisfies $m^{Lo} = 0$ or $m^{oR} = 0$, then it is a long run equilibrium; and vice versa.

Further, if agents concentrate on action Ao and the population distribution satisfies $m^{Lo} = 0$ or $m^{oR} = 0$, the potential is $\tilde{f} = a$. Since this is indeed the greatest value of the potential attainable in this game, such a

monomorphic state that agents locate either only in community L or community R and take the same action Ao attains the global maximum of $\tilde{f}^{\mathcal{P}}$ and thus the set of such states is stable both in the long run and in the medium run.

Now check a social state where all agents take AB: this is a medium-run equilibrium for any $\mathbf{m}^{\mathcal{P}}$ and it is a long-run equilibrium if $m^{Lo} = 0$ or $m^{oR} = 0$. Then, consider a state change $\mathbf{z}^{\mathcal{P}}$ such as $\mathbf{z}^p = \varepsilon m^p(1, -1, 0)$ for each $p \in \mathcal{P}$: a mass of agents switches from AB to Ao while staying in the same population. Then, we have

$$\begin{aligned}\tilde{\mathbf{F}}^{Lo}(\mathbf{z}^{\mathcal{P}}) &= \mathbf{\Pi}(\mathbf{z}^{Lo} + \mathbf{z}^{LR}) = \varepsilon(m^{Lo} + m^{LR})(c, 0, -(b - c - d))^T, \\ \tilde{\mathbf{F}}^{LR}(\mathbf{z}^{\mathcal{P}}) &= \mathbf{\Pi}(\mathbf{z}^{Lo} + \mathbf{z}^{LR} + \mathbf{z}^{oR}) = \varepsilon(c, 0, -(b - c - d))^T, \\ \tilde{\mathbf{F}}^{oR}(\mathbf{z}^{\mathcal{P}}) &= \mathbf{\Pi}(\mathbf{z}^{LR} + \mathbf{z}^{oR}) = \varepsilon(m^{LR} + m^{oR})(c, 0, -(b - c - d))^T.\end{aligned}$$

Therefore, by (7), we have

$$\mathbf{z}^{\mathcal{P}} \cdot \tilde{\mathbf{\Pi}}\mathbf{z}^{\mathcal{P}} = \varepsilon^2 \{m^{Lo}(m^{Lo} + m^{LR}) + m^{LR} + m^{oR}(m^{LR} + m^{oR})\} c.$$

This is positive for any $\varepsilon > 0$ with any $\mathbf{m}^{\mathcal{P}}$. Plugging this and $\mathbf{z}^{\mathcal{P}} \cdot \mathbf{F}(\mathbf{x}^{\mathcal{P}}) \leq 0$ into (6), we obtain $\tilde{f}^{\mathcal{P}}(\mathbf{x}^{\mathcal{P}} + \mathbf{z}^{\mathcal{P}}) > \tilde{f}^{\mathcal{P}}(\mathbf{x}^{\mathcal{P}})$: any small shift of the mass from AB to Ao in each population increases the potential. Hence, the monomorphic state concentrating on AB is not a local maximum of the potential and thus is not locally stable even in the medium run.

Now check a social state where all agents take oB: this is a medium-run equilibrium for any $\mathbf{m}^{\mathcal{P}}$ and it is a long-run equilibrium if $m^{Lo} = 0$ or $m^{oR} = 0$. Any feasible change $\mathbf{z}^{\mathcal{P}}$ in the state must satisfy

$$z_a^{oR}, z_{Ao}^p, z_{AB}^p \geq 0 \quad \text{for each } a \in \mathcal{A}, p \in \mathcal{P}; \quad \sum_{(a,p) \in \mathcal{A} \times \mathcal{P}} z_a^p = 0.$$

The payoff vector in each population is

$$\tilde{\mathbf{F}}^{Lo}(\mathbf{x}^{\mathcal{P}}) = \tilde{\mathbf{F}}^{LR}(\mathbf{x}^{\mathcal{P}}) = \begin{pmatrix} d \\ b - c \\ b \end{pmatrix}, \quad \tilde{\mathbf{F}}^{oR}(\mathbf{x}^{\mathcal{P}}) = m^{LR} \begin{pmatrix} d \\ b - c \\ b \end{pmatrix}.$$

Thus, we have

$$\begin{aligned}\mathbf{z}^{\mathcal{P}} \cdot \tilde{\mathbf{\Pi}}\mathbf{x}^{\mathcal{P}} &= \mathbf{z}^{\mathcal{P}} \cdot \tilde{\mathbf{F}}^{\mathcal{P}}(\mathbf{x}^{\mathcal{P}}) = d\hat{z}_{Ao} + (b - c)\hat{z}_{AB} + b\hat{z}_{oB} \\ &= b \sum_{(a,p) \in \mathcal{A} \times \mathcal{P}} z_a^p - (b - d)\hat{z}_{Ao} - c\hat{z}_{AB} - b(1 - m^{LR})z_{oB}^{oR},\end{aligned}$$

where $\hat{z}_a = z_a^{Lo} + z_a^{LR} + m^{LR}z_a^{oR}$ for each $a \in \mathcal{A}$; note $\hat{z}_{Ao}, \hat{z}_{AB} \geq 0$ for any feasible $\mathbf{z}^{\mathcal{P}}$. This is non-negative for any feasible $\mathbf{z}^{\mathcal{P}}$, since $b > 0, b - d > c > 0$; in particular, it is strictly negative if $z_{oB}^{oR} > 0$ or there exists a population p such that $z_{Ao}^p \neq 0$ or $z_{AB}^p \neq 0$. Thus, since $\tilde{\mathbf{F}}(\mathbf{x}^{\mathcal{P}})$ is the derivative vector of $\tilde{f}$ at $\mathbf{x}^{\mathcal{P}}$, this suggests that $\tilde{f}$ decreases locally when the state moves in the direction of $\mathbf{z}^{\mathcal{P}}$. Now consider the case that $\mathbf{z}^{\mathcal{P}} \cdot \tilde{\mathbf{\Pi}}\mathbf{x}^{\mathcal{P}} = 0$: that is, suppose that $z_{oB}^{oR} = 0$ and $z_{Ao}^p = z_{AB}^p = 0$ holds for any population $p \in \mathcal{P}$. Then, $\mathbf{z}^{\mathcal{P}}$ must take a form such as $\mathbf{z}^{Lo} = -\varepsilon \mathbf{e}_{oB}, \mathbf{z}^{LR} = \varepsilon \mathbf{e}_{oB}, \mathbf{z}^{oR} = \mathbf{0}$ with $\varepsilon \in \mathbb{R}$ such that $z \in [-m^{LR}, m^{Lo}]$. Then,

we have

$$\begin{aligned}\tilde{\mathbf{F}}^{Lo}(\mathbf{z}^{\mathcal{P}}) &= \tilde{\mathbf{F}}^{LR}(\mathbf{z}^{\mathcal{P}}) = \mathbf{\Pi}(\mathbf{z}^{Lo} + \mathbf{z}^{LR}) = (-\varepsilon + \varepsilon)\mathbf{\Pi}\mathbf{e}_{oB} = \mathbf{0}, \\ \tilde{\mathbf{F}}^{oR}(\mathbf{z}^{\mathcal{P}}) &= \mathbf{\Pi}\mathbf{z}^{LR} = \varepsilon\mathbf{\Pi}\mathbf{e}_{oB} = \varepsilon(d, b - c, b)^T.\end{aligned}$$

Therefore, by (7), we have

$$\mathbf{z}^{\mathcal{P}} \cdot \tilde{\mathbf{\Pi}}\mathbf{z}^{\mathcal{P}} = (\mathbf{z}^{Lo} + \mathbf{z}^{LR}) \cdot \mathbf{0} + \mathbf{0} \cdot \varepsilon(d, b - c, b)^T = 0.$$

Plugging this and $\mathbf{z}^{\mathcal{P}} \cdot \mathbf{F}(\mathbf{x}^{\mathcal{P}}) = 0$ into (6), we obtain $\tilde{f}(\mathbf{x}^{\mathcal{P}} + \mathbf{z}^{\mathcal{P}}) = \tilde{f}^{\mathcal{P}}(\mathbf{x}^{\mathcal{P}})$. Combining these, we find that $\tilde{f}$ attains a local maximum (though not strict) at such a monomorphic state with $m^{Lo} = 0$ or $m^{oR} = 0$. Therefore, the set of such states is locally stable.

Polimorphic states Now we consider polimorphic states where agents in different populations take different actions. Of course, we consider the states where two populations take the same action but the other one takes a different action.

We will find that, in many (though not all) of long-run equilibria, one population is abandoned (no mass of agents located there) and the other two take different actions. We start from negating long-run stability of such long-run equilibria.

LR is abandoned. Presume that there is a long-run equilibrium where all the agents in Lo take action a , all the agents in oR take action $b \neq a$ and no agents live in LR: $\mathbf{x}^{Lo} = m^{Lo}\mathbf{e}_a$, $\mathbf{x}^{oR} = m^{oR}\mathbf{e}_b$ and $\mathbf{x}^{oR} = \mathbf{0}$. In the medium run, any of such states is a strict equilibrium and thus (locally) stable. It is a long run equilibrium, if

$$\tilde{F}_a^{Lo}(\mathbf{x}^{\mathcal{P}}) = m^{Lo}\pi_{aa} = m^{oR}\pi_{bb} = \tilde{F}_b^{oR}(\mathbf{x}^{\mathcal{P}}),$$

i.e., the population distribution is such as $\mathbf{m}^{\mathcal{P}} = (\pi_{bb}/(\pi_{aa} + \pi_{bb}), \pi_{aa}/(\pi_{aa} + \pi_{bb}), 0)$.

To see stability in the long run, consider the shift of agents between populations Lo and oR, without changing actions taken in each population: the social state changes by $\mathbf{z}^{\mathcal{P}}$ such that $\mathbf{z}^{Lo} = \varepsilon\mathbf{e}_a$, $\mathbf{z}^{oR} = -\varepsilon\mathbf{e}_b$, $\mathbf{z}^{oR} = \mathbf{0}$ with $\varepsilon \in (-m^{Lo}, m^{oR})$. Then, we have

$$\tilde{\mathbf{F}}^{Lo}(\mathbf{z}^{\mathcal{P}}) = \mathbf{\Pi}\mathbf{z}^{Lo} = \varepsilon\mathbf{\Pi}\mathbf{e}_a, \quad \tilde{\mathbf{F}}^{oR}(\mathbf{z}^{\mathcal{P}}) = \mathbf{\Pi}\mathbf{z}^{LR} = -\varepsilon\mathbf{\Pi}\mathbf{e}_b.$$

Therefore, by (7), we have

$$\mathbf{z}^{\mathcal{P}} \cdot \tilde{\mathbf{\Pi}}\mathbf{z}^{\mathcal{P}} = \varepsilon^2(\mathbf{e}_a \cdot \mathbf{\Pi}\mathbf{e}_a - \mathbf{e}_b \cdot \mathbf{\Pi}\mathbf{e}_b) = \varepsilon^2(\pi_{aa} + \pi_{bb}).$$

By examining $\mathbf{\Pi}$ for each pair of two distinct actions, we can find both π_{aa} and π_{bb} are positive for any pair of two actions a, b . Therefore, $\mathbf{z}^{\mathcal{P}} \cdot \tilde{\mathbf{\Pi}}\mathbf{z}^{\mathcal{P}} > 0$ for any $\varepsilon > 0$. Since the initial state $\mathbf{x}^{\mathcal{P}}$ is presumed to be a long-run equilibrium, it must be the case that $\mathbf{z}^{\mathcal{P}} \cdot \mathbf{F}(\mathbf{x}^{\mathcal{P}}) = 0$. Plugging these two results into (6), we obtain $\tilde{f}^{\mathcal{P}}(\mathbf{x}^{\mathcal{P}} + \mathbf{z}^{\mathcal{P}}) > \tilde{f}^{\mathcal{P}}(\mathbf{x}^{\mathcal{P}})$: any small shift of the mass between Lo and oR (in any direction) increases the potential. Hence, the initial state $\mathbf{x}^{\mathcal{P}}$ is not a local maximum of the potential when the population distribution is not fixed, and thus it is not locally stable in the long run.

oR (or Lo) is abandoned. Presume that there is a long-run equilibrium where all the agents in Lo take action a , all the agents in LR take action $b \neq a$ and no agents live in oR: $\mathbf{x}^{Lo} = m^{Lo}\mathbf{e}_a$, $\mathbf{x}^{LR} = m^{LR}\mathbf{e}_b$ and $\mathbf{x}^{oR} = \mathbf{0}$. Note that $m^{oR} = 0$ implies $\tilde{\mathbf{x}}^{Lo} = \tilde{\mathbf{x}}^{LR} = \mathbf{x}^{Lo} + \mathbf{x}^{LR}$. Thus, agents in these two populations observe the same action distribution and thus have the same set of the best response actions. To have them choose different actions, $\tilde{\mathbf{x}}^{Lo}(= \tilde{\mathbf{x}}^{LR})$ must admit two actions and these actions must be in the support of the observed action distribution $\tilde{\mathbf{x}}^{Lo}$. According to our identification of the best response actions, it is the case only when the two actions are AB and oB and $\tilde{x}_{AB} = \tilde{x}_{oB} = (b-d)/(a-b)$. Hence, we have two long-run equilibria of this type:

$$\begin{aligned} x_{AB}^{Lo} = m^{Lo} = (b-d)/(a-d) & \quad \text{with } x_{oB}^{LR} = m^{LR} = (a-b)/(a-d); & \quad \text{and} \\ x_{oB}^{Lo} = m^{Lo} = (a-b)/(a-d) & \quad \text{with } x_{AB}^{LR} = m^{LR} = (b-d)/(a-d). \end{aligned}$$

Note that, to have such a social state in the medium-run equilibrium, the population distribution $\mathbf{m}^P$ must satisfy this. However, as we found when identifying the best response action, any small perturbation breaks this indifference between the two actions and benefits the one that receives more agents in the perturbation. Thus, this cannot be stable even in the medium run.

Henceforth, we denote a combination of actions such as (Ao,AB,oB), by which we mean that all agents in Lo take Ao, those in LR take AB and those in oR take oB. We focus on states with $\mathbf{m}^P \gg \mathbf{0}$; if any population is abandoned, it falls into the cases that we have checked so far.

(Ao,Ao,AB), or (AB,Ao,Ao). With action combination (Ao,Ao,AB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} m^{Lo} + m^{LR} \\ 0 \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} m^{Lo} + m^{LR} \\ m^{oR} \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} m^{LR} \\ m^{oR} \\ 0 \end{pmatrix}.$$

According to (8), Ao is indeed the only best response action in Lo and LR under any $\mathbf{m}^P \gg \mathbf{0}$, or as long as $m^{Lo} + m^{LR} > 0$. To make AB a best response in oR, we need

$$(a-c)m^{LR} + (a-c)m^{oR} \geq am^{LR} + (a-c)m^{oR},$$

which reduces to $cm^{LR} \leq 0$. Since $c > 0$ and $m^{LR} \geq 0$, this can be the case only if $m^{LR} = 0$. Therefore, a medium-run equilibrium with (Ao,Ao,AB) must be a duomorphic state with Ao in Lo, AB in oR and $\mathbf{m}^{LR} = 0$. Similarly, a medium-run equilibrium with (AB,Ao,Ao) must be a duomorphic state with AB in Lo, Ao in oR and $\mathbf{m}^{LR} = 0$. Thus, there is no medium-run equilibrium with (Ao,Ao,AB) or (AB,Ao,Ao) and $\mathbf{m}^P \gg \mathbf{0}$.

(Ao,Ao,oB) or (oB,Ao,Ao). With action combination (Ao,Ao,AB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} m^{Lo} + m^{LR} \\ 0 \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} m^{Lo} + m^{LR} \\ 0 \\ m^{oR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} m^{LR} \\ 0 \\ m^{oR} \end{pmatrix}.$$

According to (8), Ao is indeed the only best response action in Lo any $\mathbf{m}^P \gg \mathbf{0}$, or as long as $m^{Lo} + m^{LR} > 0$. To make Ao and AB best responses in LR and oR respectively, we need

$$\begin{aligned} \text{to make Ao better for LR than AB:} \quad & c(1 - m^{oR}) = c(m^{Lo} + m^{LR}) \geq (b - c - d)m^{oR}, \\ \text{to make oB better for oR than AB:} \quad & cm^{oR} \geq (a - c - d)m^{LR}. \end{aligned}$$

These are summarized as

$$\frac{a - c - d}{c} m^{LR} \leq m^{oR} \leq \frac{c}{b - d}.$$

The RHS is smaller than 1 by $b - d > c$. It is greater than the LHS if m^{LR} is sufficiently small (and note that the RHS); then there exists a continuous range of m^{oR} (and thus of $\mathbf{m}^P$) that meet this condition. With such an $\mathbf{m}^P$, (Ao,Ao,oB) constitutes a medium-run equilibrium. In particular, if m^{oR} satisfies the above equation with strict inequalities, the optimal actions are strictly better than any other actions for each population; thus, the equilibrium is strict and thus locally stable under admissible dynamics.

This action combination constitutes a long-run equilibrium with $\mathbf{m}^P \gg \mathbf{0}$, if and only if $\mathbf{m}^P$ satisfies the above condition for a medium-run equilibrium and

$$am^{Lo} + am^{LR} = am^{Lo} + am^{LR} + dm^{oR} = dm^{LR} + bm^{oR}$$

to make all populations indifferent and prevent migration. The first equality cannot hold with $m^{oR} > 0$ since $d > 0$. Thus, there is no such a long-run equilibrium with $\mathbf{m}^P \gg \mathbf{0}$.

Similarly, there are stable medium-run equilibria with (oB,Ao,Ao) and $\mathbf{m}^P \gg \mathbf{0}$ and but no long-run equilibrium with such a distribution.

(Ao,AB,Ao). With action combination (Ao,AB,Ao), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} m^{Lo} \\ m^{LR} \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} m^{Lo} + m^{oR} \\ m^{LR} \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} m^{oR} \\ m^{LR} \\ 0 \end{pmatrix}.$$

According to (8), Ao is the only best response in Lo and LR under any $\mathbf{m}^P \gg \mathbf{0}$, or as long as $m^{Lo} > 0$ and $m^{oR} > 0$ respectively for each population. To make AB a best response in LR, we need to

$$c(m^{Lo} + m^{oR}) \leq 0 = (b - c - d) \cdot 0.$$

Since $c > 0$, this can be the case only if $m^{Lo} + m^{LR} = 0$, i.e., $\mathbf{m}^P = (0, 1, 0)$. Then, it reduces to a monomorphic state with $\mathbf{m}^P = (0, 1, 0)$. Thus, there is no medium-run equilibrium with (Ao,AB,Ao) and $\mathbf{m}^P \gg \mathbf{0}$.

(Ao,AB,AB) or (AB,AB,Ao). With action combination (Ao,AB,AB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} m^{Lo} \\ m^{LR} \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} m^{Lo} \\ m^{LR} + m^{oR} \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} 0 \\ m^{LR} + m^{oR} \\ 0 \end{pmatrix}.$$

According to (8), Ao is indeed the only best response action in Lo under any $\mathbf{m}^P \gg \mathbf{0}$, or as long as $m^{Lo} > 0$; AB is the best response in oR as well as Ao under any $\mathbf{m}^P \gg \mathbf{0}$ (and not oB as long as $m^{LR} + m^{oR} > 0$). To make AB a best response in LR, we need to make AB better for LR than Ao:

$$(a - c)\{m^{Lo} + (m^{LR} + m^{oR})\} \geq am^{Lo} + (a - c)(m^{LR} + m^{oR}).$$

Since $c > 0$, this can be the case only if $m^{Lo} = 0$. Thus, there is no medium-run equilibrium with (Ao,AB,AB) and $\mathbf{m}^P \gg \mathbf{0}$. Similarly there is no medium-run equilibrium with (AB,AB,Ao) and $\mathbf{m}^P \gg \mathbf{0}$.

(Ao,AB,oB) or (oB,AB,Ao). With action combination (Ao,AB,AB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} m^{Lo} \\ m^{LR} \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} m^{Lo} \\ m^{LR} \\ m^{oR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} 0 \\ m^{LR} \\ m^{oR} \end{pmatrix}.$$

According to (8), Ao is indeed the only best response action in Lo under any $\mathbf{m}^P \gg \mathbf{0}$, or as long as $m^{Lo} > 0$. To make AB and oB best responses in LR and oR respectively, we need

$$\begin{aligned} \text{to make AB better for LR than Ao:} & \quad (b - c - d)m^{oR} \geq cm^{Lo}, \\ \text{to make AB better for LR than oB:} & \quad (a - c - d)m^{Lo} + (a - b)m^{LR} \geq cm^{oR}, \\ \text{to make oB better for oR than AB:} & \quad cm^{oR} \geq (a - b)m^{LR}. \end{aligned}$$

These are summarized as

$$\max \left\{ \frac{c}{b - c - d} m^{Lo}, \frac{a - b}{c} m^{LR} \right\} \leq m^{oR} \leq \frac{a - c - d}{c} m^{Lo} + \frac{a - b}{c} m^{LR}.$$

Note that, if m^{Lo} is positive but small enough, then the LHS is $m^{LR}(a - b)/c$ and the RHS is strictly greater than the LHS; then there exists a continuous range of m^{oR} (and thus of $\mathbf{m}^P$) that meet this condition. With such an $\mathbf{m}^P$, (Ao,AB,oB) constitutes a medium-run equilibrium. In particular, if m^{oR} satisfies the above equation with strict inequalities, the optimal actions are strictly better than any other actions for each population; thus, the equilibrium is strict and thus locally stable under admissible dynamics.

This action combination constitutes a long-run equilibrium with $\mathbf{m}^P \gg \mathbf{0}$, if and only if $\mathbf{m}^P$ satisfies the above condition for a medium-run equilibrium and

$$am^{Lo} + (a - c)m^{LR} = (a - c)m^{Lo} + (a - c)m^{LR} + (b - c)m^{oR} = (b - c)m^{LR} + bm^{oR}$$

to make all populations indifferent and prevent migration. The second equality and the second inequality in the medium-run equilibrium condition (or the condition to make AB better for LR than oB) imply $dm^{Lo} \leq 0$. It holds with $d > 0$ only if $m^{Lo} = 0$. Thus, there is no such a long-run equilibrium with $\mathbf{m}^P \gg \mathbf{0}$.

Similarly, there are stable medium-run equilibria with (oB,AB,Ao) and $\mathbf{m}^P \gg \mathbf{0}$ and but no long-run equilibrium with such a distribution.

(Ao,oB,Ao). With action combination (Ao,AB,AB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} m^{Lo} \\ 0 \\ m^{LR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} m^{Lo} + m^{oR} \\ 0 \\ m^{LR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} m^{oR} \\ 0 \\ m^{LR} \end{pmatrix}.$$

To make Ao the best response in Lo and oR and oB in LR, we need

$$\begin{aligned} \text{to make Ao better for Lo than oB:} & \quad am^{Lo} + dm^{LR} \geq dm^{Lo} + bm^{LR}, \\ \text{to make oB better for LR than Ao:} & \quad d(m^{Lo} + m^{oR}) + bm^{LR} \geq a(m^{Lo} + m^{oR}) + dm^{LR}, \\ \text{to make Ao better for oR than Ao:} & \quad dm^{LR} + am^{oR} \geq bm^{LR} + dm^{oR}. \end{aligned}$$

They imply

$$(a - d) \min\{m^{Lo}, m^{oR}\} \geq (b - d)m^{LR} \geq (a - d)(m^{Lo} + m^{oR}).$$

Since $a > d$, this can be the case only if $m^{Lo} = m^{oR} = m^{LR} = 0$. Thus, there is no medium-run equilibrium with (Ao,oB,Ao) (for any $\mathbf{m}^P$).

(Ao,oB,AB) or (AB,oB,Ao). With action combination (Ao,oB,AB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} m^{Lo} \\ 0 \\ m^{LR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} m^{Lo} \\ m^{oR} \\ m^{LR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} 0 \\ m^{oR} \\ m^{LR} \end{pmatrix}.$$

To make Ao the best response in Lo, oB in LR and AB in oR, we need

$$\begin{aligned} \text{to make Ao better for Lo than oB:} & \quad am^{Lo} + dm^{LR} \geq dm^{Lo} + bm^{LR}, \\ \text{to make oB better for LR than AB:} & \quad dm^{Lo} + bm^{LR} + (b - c)m^{oR} \geq (a - c)m^{Lo} + (b - c)m^{LR} + (a - c)m^{oR}, \\ \text{to make oB better for LR than Ao:} & \quad dm^{Lo} + bm^{LR} + (b - c)m^{oR} \geq am^{Lo} + dm^{LR} + (a - c)m^{oR}, \\ \text{to make AB better for oR than oB:} & \quad (b - c)m^{LR} + (a - c)m^{oR} \geq bm^{LR} + (b - c)m^{oR}. \end{aligned}$$

The first and third imply

$$(a - d)m^{Lo} \geq (b - d)m^{LR} \geq (a - d)m^{Lo} + (a - b)m^{oR}.$$

Since $a > b$, this can be the case only if $m^{oR} = 0$. Similarly, we can find the second and fourth hold only if $m^{Lo} = 0$. Thus, there is no medium-run equilibrium with (Ao,oB,AB) and $\mathbf{m}^P \gg \mathbf{0}$. Similarly there is no medium-run equilibrium with (AB,oB,Ao) and $\mathbf{m}^P \gg \mathbf{0}$.

(AB,Ao,AB). With action combination (AB,Ao,AB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} m^{LR} \\ m^{Lo} \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} m^{LR} \\ m^{Lo} + m^{oR} \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} m^{LR} \\ m^{oR} \\ 0 \end{pmatrix}.$$

According to (8), Ao is indeed the only best response action in LR under any $\mathbf{m}^P \gg \mathbf{0}$, or as long as $m^{Lo} > 0$. To make AB a best response in Lo and oR, we need

$$cm^{LR} \leq 0 = (b - c - d) \cdot 0.$$

Since $c > 0$, this can be the case only if $m^{LR} = 0$. Then, it reduces to a monomorphic state of AB with $\mathbf{m}^{LR} = 0$. Thus, there is no medium-run equilibrium with (AB,Ao,AB) and $\mathbf{m}^P \gg \mathbf{0}$.

(AB,Ao,oB) or (oB,Ao,AB). With action combination (AB,Ao,AB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} m^{LR} \\ m^{Lo} \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} m^{LR} \\ m^{Lo} \\ m^{oR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} m^{LR} \\ 0 \\ m^{oR} \end{pmatrix}.$$

According to (8), to make AB a best response in Lo, we need

$$cm^{LR} \leq 0 = (b - c - d) \cdot 0.$$

Since $c > 0$, this can be the case only if $m^{LR} = 0$. Then, it reduces to duomorphic state with AB in Lo, oB in oR and $\mathbf{m}^{LR} = 0$. $\mathbf{m}^{LR} = 0$. Similarly, a medium-run equilibrium with (oB,Ao,AB) must be a duomorphic state with oB in Lo, Ao in Lo and $\mathbf{m}^{LR} = 0$. Thus, there is no medium-run equilibrium with (AB,Ao,oB) or (oB,Ao,AB) and $\mathbf{m}^P \gg \mathbf{0}$.

(AB,AB,oB) or (oB,AB,AB). With action combination (AB,AB,oB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} 0 \\ m^{Lo} + m^{LR} \\ 0 \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} 0 \\ m^{Lo} + m^{LR} \\ m^{oR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} 0 \\ m^{LR} \\ m^{oR} \end{pmatrix}.$$

According to (8), AB is a best response in Lo as well as Ao under any $\mathbf{m}^P$. To make AB and oB best responses in LR and oR respectively, we need

$$\begin{aligned} \text{to make AB better for LR than Ao:} & \quad 0 = c \cdot 0 \geq (b - c - d)m^{oR}, \\ \text{to make AB better for LR than oB:} & \quad (a - b)(1 - m^{oR}) = (a - c - d) \cdot 0 + (a - b)(m^{Lo} + m^{LR}) \geq cm^{oR}, \\ \text{to make oB better for oR than AB:} & \quad (a - b)m^{LR} = (a - c - d) \cdot 0 + (a - b)m^{LR} \geq cm^{oR}. \end{aligned}$$

These are summarized as

$$\frac{a-b}{c}m^{LR} \leq m^{oR} \leq \frac{a-b}{a-b+c}.$$

There exists a continuous range of m^{oR} (and thus of $\mathbf{m}^P$) that meet this condition. With such an $\mathbf{m}^P$, (Ao,AB,oB) constitutes a medium-run equilibrium. However, since AB is indifferent in Lo with Ao, any shift in Lo from AB to Ao makes Ao the only best response. Thus, these medium-run equilibria cannot be stable.

This action combination constitutes a long-run equilibrium with $\mathbf{m}^P \gg \mathbf{0}$, if and only if $\mathbf{m}^P$ satisfies the above condition for optimal choices of actions and

$$(a-c)(m^{Lo} + m^{LR}) = (a-c)(m^{Lo} + m^{LR}) + (b-c)m^{oR} = (b-c)m^{LR} + bm^{oR}$$

to make all populations indifferent and prevent migration. Since $b > c$, the first equality can hold only if $m^{oR} = 0$. Thus, there is no such a long-run equilibrium with $\mathbf{m}^P \gg \mathbf{0}$.

(AB,oB,AB). With action combination (AB,oB,AB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} 0 \\ m^{Lo} \\ m^{LR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} 0 \\ m^{Lo} + m^{oR} \\ m^{LR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} 0 \\ m^{oR} \\ m^{LR} \end{pmatrix}.$$

According to (8), to make these actions best responses in each population, we need

$$\begin{aligned} \text{to make AB better for Lo than oB:} & \quad (a-b)m^{Lo} \geq cm^{LR}, \\ \text{to make oB better for LR than AB:} & \quad cm^{LR} \geq (a-b)m^{Lo} + (a-b)m^{oR}, \\ \text{to make AB better for oR than oB:} & \quad (a-b)m^{oR} \geq cm^{LR}. \end{aligned}$$

They hold together only if $m^{LR} = 0$. Then, this state reduces to a monomorphic state of AB with $m^{LR} = 0$. Thus, there is no such a medium-run equilibrium with $\mathbf{m}^P \gg \mathbf{0}$.

(AB,oB,oB) or (oB,oB,AB). With action combination (AB,oB,oB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} 0 \\ m^{Lo} \\ m^{LR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} 0 \\ m^{Lo} \\ m^{LR} + m^{oR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} 0 \\ 0 \\ m^{LR} + m^{oR} \end{pmatrix}.$$

According to (8), oB is the only best response in oR for any $\mathbf{m}^P \gg \mathbf{0}$. To make AB and oB best responses in LR and oR respectively, we need

$$\begin{aligned} \text{to make AB better in Lo than Ao:} & \quad 0 = c \cdot 0 \leq (b-c-d)m^{LR}, \\ \text{to make AB better in Lo than oB:} & \quad (a-b)m^{Lo} = (a-c-d) \cdot 0 + (a-b)m^{Lo} \geq cm^{LR}, \\ \text{to make oB better for LR than AB:} & \quad (a-b)m^{Lo} = (a-c-d) \cdot 0 + (a-b)m^{Lo} \leq c(m^{LR} + m^{oR}). \end{aligned}$$

These are summarized as

$$cm^{LR} \leq (a-b)m^{Lo} \leq c(m^{LR} + m^{oR}).$$

There exists a continuous range of m^{oR} (and thus of $\mathbf{m}^P$) that meet this condition. With such an $\mathbf{m}^P$, (Ao,AB,oB) constitutes a medium-run equilibrium. With such an $\mathbf{m}^P$, (AB,oB,oB) constitutes a medium-run equilibrium. In particular, if m^{Lo} satisfies the above equation with strict inequalities, the optimal actions are strictly better than any other actions for each population; thus, the equilibrium is strict and thus locally stable under admissible medium-run dynamics.

This action combination constitutes a long-run equilibrium with $\mathbf{m}^P \gg \mathbf{0}$, if and only if $\mathbf{m}^P$ satisfies the above condition for optimal choices of actions and

$$(a-c)m^{Lo} + (b-c)m^{LR} = (b-c)m^{Lo} + b(m^{LR} + m^{oR}) = b(m^{LR} + m^{oR})$$

to make all populations indifferent and prevent migration. Since $b-c > 0$, the second inequality holds only if $m^{Lo} = 0$. Thus, there is no such a long-run equilibrium with $\mathbf{m}^P \gg \mathbf{0}$.

(oB,Ao,oB). With action combination (AB,Ao,AB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} m^{LR} \\ 0 \\ m^{Lo} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} m^{LR} \\ 0 \\ m^{Lo} + m^{oR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} m^{LR} \\ 0 \\ m^{oR} \end{pmatrix}.$$

According to (8), to make these actions best responses in each population, we need

$$\begin{aligned} \text{to make oB the best in Lo:} & \quad (a-c-d)m^{LR} \leq cm^{Lo}, \\ \text{to make Ao the best in LR:} & \quad cm^{LR} \geq (b-c-d)(m^{Lo} + m^{oR}), \\ \text{to make oB the best in oR:} & \quad (a-c-d)m^{LR} \leq cm^{oR}. \end{aligned}$$

These are summarized as

$$\frac{b-c-d}{c}(m^{Lo} + m^{oR}) \leq m^{LR} \leq \frac{c}{2(a-c-d)}(m^{Lo} + m^{oR}).$$

Since $(b-c-d)/c > (b-d)/c > c/(a-d) > c(a-c-d)$, this can be the case only if $m^{Lo} + m^{oR} = m^{LR} = 0$. Thus, there is no medium-run equilibrium with (oB,Ao,oB) and $\mathbf{m}^P \gg \mathbf{0}$.

(oB,AB,oB). With action combination (oB,AB,oB), the observed action distribution for agents in each population is

$$\tilde{\mathbf{x}}^{Lo} = \begin{pmatrix} 0 \\ m^{LR} \\ m^{Lo} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{LR} = \begin{pmatrix} 0 \\ m^{LR} \\ m^{Lo} + m^{oR} \end{pmatrix}, \quad \tilde{\mathbf{x}}^{oR} = \begin{pmatrix} 0 \\ m^{LR} \\ m^{oR} \end{pmatrix}.$$

According to (8), to make these actions best responses in each population, we need

$$\begin{aligned} \text{to make oB the best in Lo:} & \quad (a-b)m^{LR} \leq cm^{Lo}, \\ \text{to make AB the best in LR:} & \quad (a-b)m^{LR} \geq c(m^{Lo} + m^{oR}), \end{aligned}$$

to make oB the best in oR:

$$(a - b)m^{LR} \leq cm^{oR}.$$

These are summarized as

$$c(m^{Lo} + m^{oR}) \leq (a - b)m^{LR} \leq 2c(m^{Lo} + m^{oR}).$$

Since $c > 0$, this can be the case only if $m^{Lo} + m^{oR} = m^{LR} = 0$. Thus, there is no medium-run equilibrium with (oB,AB,oB) and $\mathbf{m}^P \gg \mathbf{0}$.

2 Supplementary Discussion

2.1 Simulations of the bilingual game

Here we run simulations of the bilingual game in Example 3 with several variations. The results suggest robustness of the findings in this example. Furthermore, they show how we could utilize analytical results to conduct simulations.

2.1.1 The historical dependency of the limit states on the initial states

Analytically, we find two distinct sets of stable long-run equilibria, one with all agents taking action Ao and the other with all agents taking action oB. The multiplicity of stable equilibria suggests that the limit states of the dynamics depend on the initial states.

One would attempt to use simulations to show concretely which limit state the dynamic converges to from a given initial state. However, convergence is a notion on an infinite time horizon, while we can only run a simulation for a finite time period. Therefore, it would be difficult to determine the limit state within a finite time period especially if we expect nonmonotonic behavior of the dynamics.

Analytical results help us resolve this concern. First, we know that all agents should commonly take either Ao or oB as their best responses, regardless of their affiliations in the limit. Since both long-run equilibria are locally stable, the best responses would be unchanged as long as a social state is sufficiently close to one of the limit states.

In fact, Ao constitutes a strict equilibrium in the base game. The analysis of the best responses (see Supplementary Methods 1.10) suggests that, once Ao becomes the best response in a population, an increase in the share of Ao keeps action Ao the best response in this population. Therefore, once it becomes the best response in all populations, it is guaranteed that the dynamic converges to the stable equilibrium set where Ao is used. Similar for the stable equilibrium set where oB is used.

We utilize this observation in simulations. We run a long-run dynamic with $\rho^A = 0.0475$, $\rho^P = 0$ and $\rho^{AP} = 0.025$, so that an agent changes her action alone or her action-affiliation pair, but never her affiliation alone. All subdynamics are best response dynamics. The time is discretized so these revision rates are the probabilities for an agent to receive these revision opportunities in each period. In every 100 periods, it is checked if all populations agree to take Ao as their best responses or to take oB or neither (all AB or they do not agree on the best response actions). If they agree, the simulation stops and we conclude that the agreed best response action is the limit action. If they do not agree on Ao or oB, the simulation continues.

In simulations in Supplementary Fig.1, we set the payoffs of the base game as $a = 4, b = 3, c = 0.75$ and

$d = 0$, i.e.,

$$\mathbf{\Pi} = \begin{pmatrix} 4 & 3.25 & 0 \\ 3.25 & 3.25 & 2.25 \\ 0 & 2.25 & 3 \end{pmatrix}.$$

The initial state in each simulation is set to a medium-run stable equilibrium (Ao,AB,oB), i.e., the state where all agents in Lo take Ao, those in LR take AB and those in oR take oB; the limit action evolved from this initial state would be the least trivially predictable.

We vary the share of agents over the three populations. In the figure, the horizontal axis shows m^{Lo} in the initial state and the vertical one shows m^{oR} ; note that $m^{LR} = 1 - m^{Lo} - m^{oR}$. The color of each dot represents the limit action for the dynamic starting from a given initial state (indicated by the location of the point); red for Ao and blue for oB. This limit action is determined in the finite period simulation based on the above observation.

The figure suggests that i) the basin of attraction to Ao is much greater than that to oB, ii) the threshold between these two basins looks like a straight line, possibly represented by a linear equation, iii) on this threshold, an increase in m^{Lo} (accompanied with a decrease in m^{LR}) makes it more likely to converge to Ao and an increase in m^{oR} (accompanied with a decrease in m^{LR}) makes it more likely to converge to oB. This confirms our expectation of the dependency of the limit states on the initial states.

2.1.2 The population size effects

In the discussion section of the main text, we argue that the effects of the population sizes can be isolated to a certain extent by adding a common constant θ to all the cells in the base payoff matrix and investigating how the behavior/outcome of the dynamic changes with θ ; positive θ encourages agglomeration to a population that has a greater degree of interactions and negative θ encourages dispersion from such a population.¹⁰ If there is a population size effect, this comparative statics should show changes in the historical dependency of the limit states on the initial states and in the basins of attraction to each stable equilibrium. Further, from analytical investigation on the condition of maintaining local stability of each equilibrium, we can predict structural changes (bifurcation) in the set of possible limit states.

In Supplementary Fig.2, we conducted the above analysis by varying θ . The payoff matrix is changed as

$$\mathbf{\Pi}(\theta) = \begin{pmatrix} 4 + \theta & 3.25 + \theta & 0 + \theta \\ 3.25 + \theta & 3.25 + \theta & 2.25 + \theta \\ 0 + \theta & 2.25 + \theta & 3 + \theta \end{pmatrix}.$$

We run simulations with various values of θ . In Supplementary Fig.2a, θ is set to 100.¹¹ Compared to Supplementary Fig.1, it shows no difference in the limit actions from any initial states. Hence, we conclude that, when increasing the incentive for agglomeration, it does not change the basins of attraction to the two stable equilibria.

Careful examination of the proofs in Supplementary Methods 1.10 suggests that the assumption of $d(= \pi_{Ao,oB} = \pi_{oB,Ao}) > 0$ is used to eliminate some polymorphic action combinations, especially (Ao,oB,oB), from the long-run equilibria. Thus, we would expect changes in the long-run stable equilibrium set when

¹⁰More concretely, adding θ to the (original) base payoff matrix $\mathbf{\Pi}^0$ means that the payoff matrix changes to $\mathbf{\Pi}(\theta) = \mathbf{\Pi}^0 + \theta \mathbf{I}$ where $\mathbf{I}$ is a matrix with 1 in all cells. Fix a social state at $\mathbf{x}^p$; $\pi_a^p(\theta) := (\mathbf{\Pi} + \theta \mathbf{I})\tilde{\mathbf{x}}^p$ is the payoff in population p from action a at this social state with a given θ . Let $\tilde{y}^p = \sum_{a \in \mathcal{A}} x_a^p$ be the total sum of weighted masses (i.e., the total frequencies of interaction) for population p . Then, $\pi_\theta^p = \mathbf{\Pi}\tilde{\mathbf{x}}^p + \theta \mathbf{I}\tilde{\mathbf{x}}^p = \pi_0^p + \theta \tilde{y}^p$.

¹¹We also try $\theta = 1, 5, 10$ and confirm that these values do not change the basins of attraction of the two stable equilibria.

$\pi_{Ao,oB}^\theta = 0 + \theta < 0$. In Supplementary Fig.2b and 2c, we set θ to -0.05 and -1 , respectively. The small dots indicate that the best responses do not agree in 10000 periods; we found that, in these cases, the best responses are Ao in population Lo and oB in LR and oR at the terminal period 10000. Hence, it suggests that the dynamic converges to a new limit action combination (Ao,oB,oB). However, the figures also show that, from most of the initial states, the limit action is still either Ao for all populations or oB for all. The shapes of the basins of the attraction look quite similar to Supplementary Fig. 1. Therefore, the qualitative results mentioned at the end of the previous subsection should be robust. Comparison between Supplementary Fig.2b and 2c suggests that greater dispersion incentive (i.e., change from $\theta = -0.05$ to $\theta = -1$) shrinks the basin of attraction to Ao.

2.1.3 Asymmetric costs

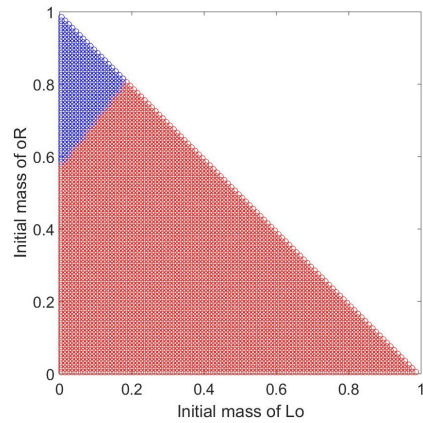
To see the robustness of the non-monotonic convergence in the simulation result in Example 3 in the main text, here we introduce an asymmetric cost $f > 0$. The cost only applies to agents who choose the bilingual action AB . Now the payoff matrix is given as

$$\mathbf{\Pi} = \begin{pmatrix} a & a - c & d \\ a - c - f & a - c - f & b - c - f \\ d & b - c & b \end{pmatrix}.$$

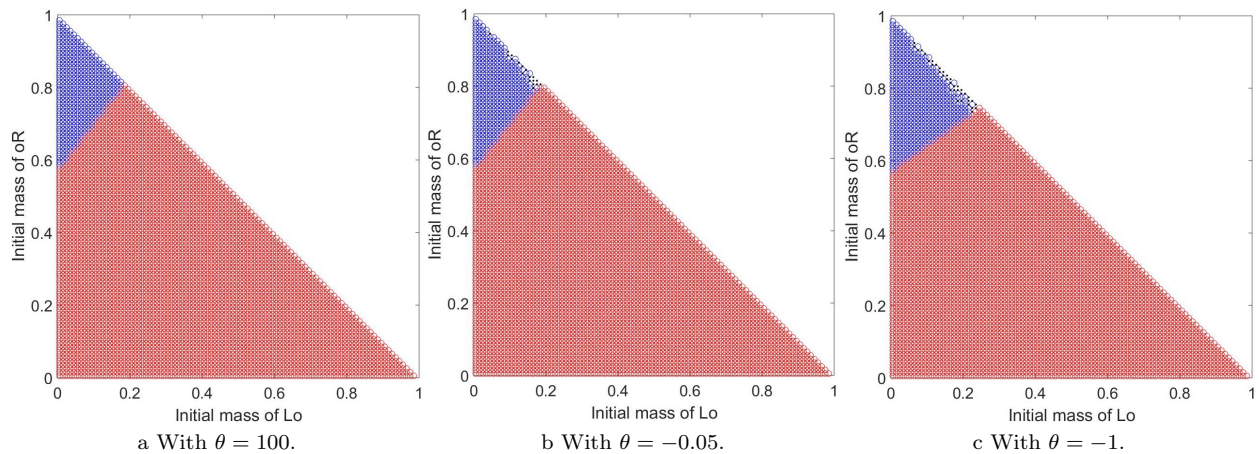
In Supplementary Fig.3, the asymmetric bilingual cost is only $f = -0.1$. All other parameters and the initial state are set to be the same as in the simulation in the main text. The evolution of the action distribution and the change in payoffs look almost the same as those in Fig.3 in the main text.

While we would expect a large asymmetric cost to make significant changes, the simulation in Supplementary Fig.4 shows that the dynamic eventually converges to the same long-run outcome as in the one in the main text.

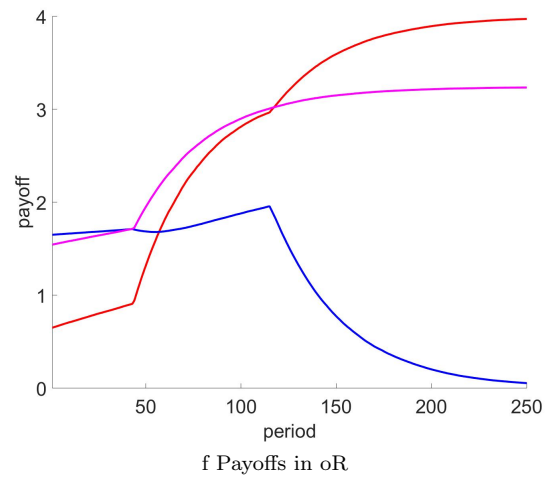
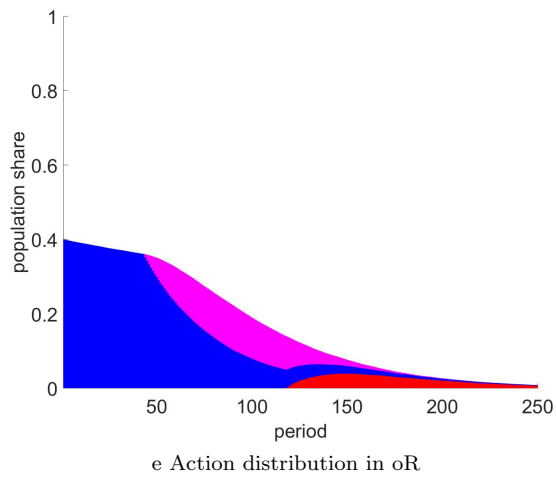
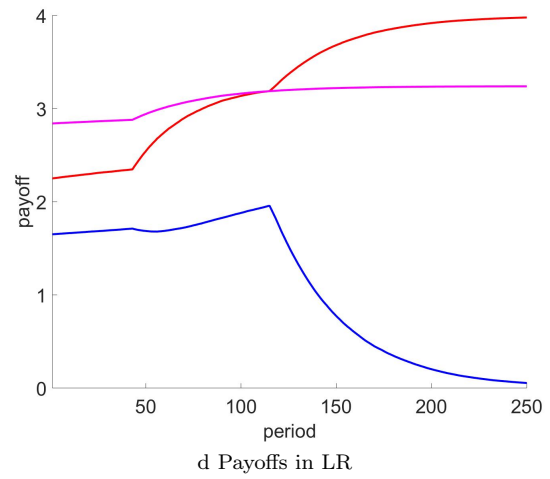
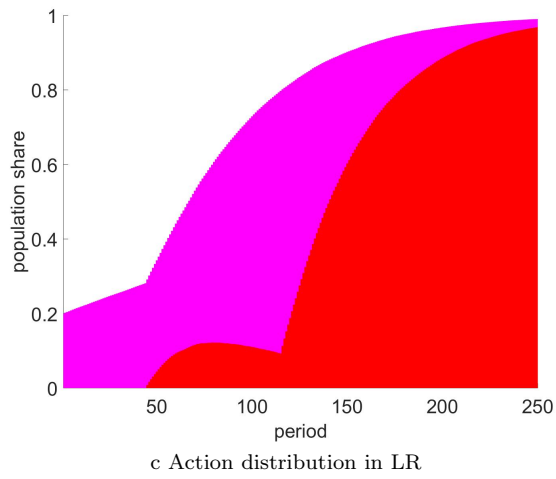
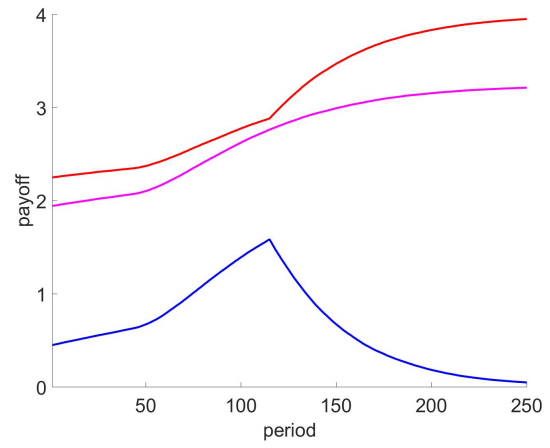
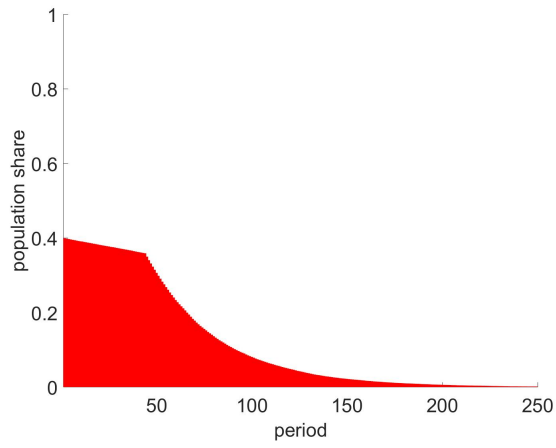
3 Supplementary Figures



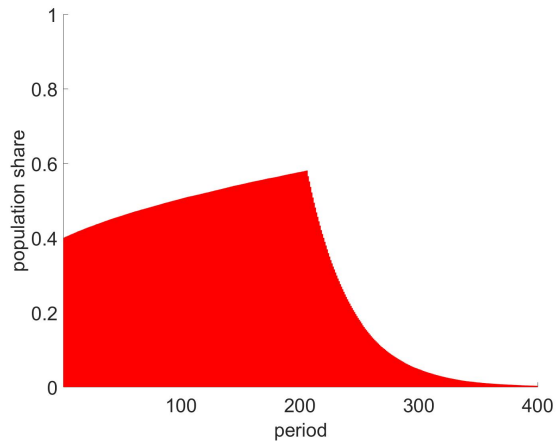
Supplementary Figure 1: The limit action from (Ao,AB,oB) with initial $\mathbf{m}^P$ varying over Δ^P . The horizontal and vertical axes show m^{Lo} and m^{oR} in the initial state. The color of each dot indicates the limit action, red for Ao and blue for oB, under the long-run best response dynamic from the initial state represented by the location of the point.



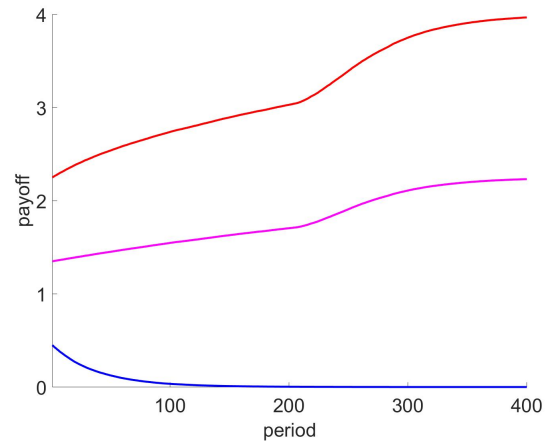
Supplementary Figure 2: The population size effect on the limit states in the bilingual game.



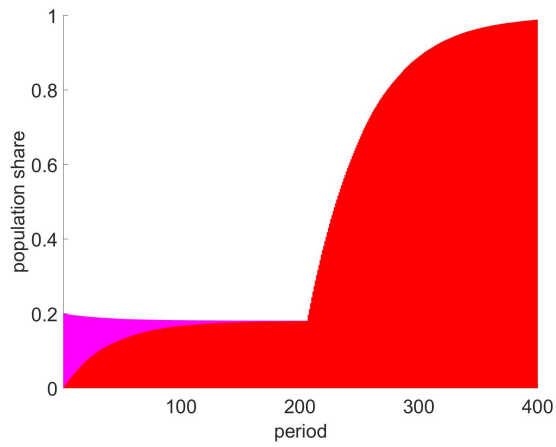
Supplementary Figure 3: Simulation of the long-run BRD in the bilingual game with a small asymmetric cost $f = 0.01$.



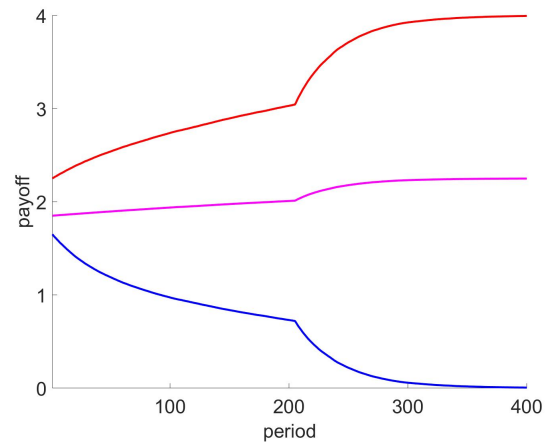
a Action distribution in Lo



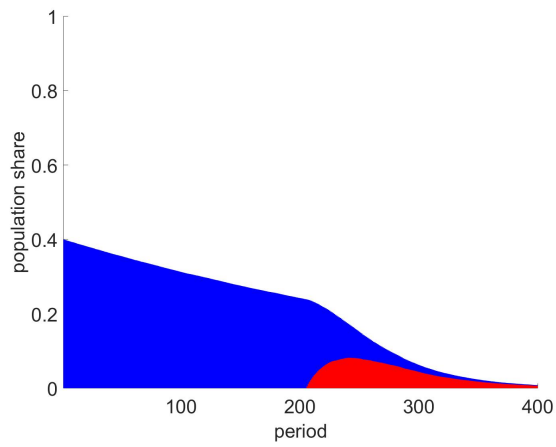
b Payoffs in Lo



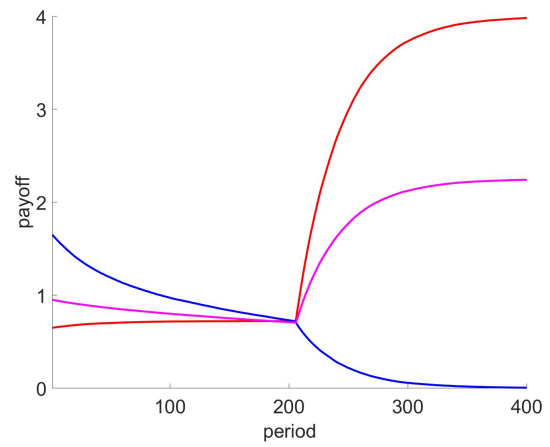
c Action distribution in LR



d Payoffs in LR



e Action distribution in oR



f Payoffs in oR

Supplementary Figure 4: Simulation of the long-run BRD in the bilingual game with a large asymmetric cost $f = 1$.

Supplementary references

- BOYD, R., AND P. J. RICHERSON (2009): “Voting with your feet: Payoff biased migration and the evolution of group beneficial behavior,” *Journal of Theoretical Biology*, 257(2), 331 – 339.
- GILBOA, I., AND A. MATSUI (1991): “Social Stability and Equilibrium,” *Econometrica*, 59(3), 859–867.
- HOFBAUER, J. (1995): “Stability for the best response dynamics,” mimeo, University of Vienna.
- SANDHOLM, W. H. (2001): “Potential Games with Continuous Player Sets,” *Journal of Economic Theory*, 97(1), 81–108.
- (2010): *Population games and evolutionary dynamics*. MIT Press.
- SCHLAG, K. H. (1998): “Why Imitate, and If So, How? A Boundedly Rational Approach to Multi-armed Bandits,” *Journal of Economic Theory*, 78, 130–56.
- SMITH, M. J. (1984): “The Stability of a Dynamic Model of Traffic Assignment: An Application of a Method of Lyapunov,” *Transportation Science*, 18(3), 245–252.
- SUNDARAM, R. K. (1996): *A First Course in Optimization Theory*. Cambridge University Press, Cambridge and New York.
- ZUSAI, D. (2018a): “Gains in evolutionary dynamics: a unifying approach to dynamic stability for contractive games and ESS,” mimeo, <http://sites.temple.edu/zusai/research/gain/>.
- (2018b): “Heterogeneity and aggregation in evolutionary dynamics: a general framework without aggregability,” mimeo, <https://sites.temple.edu/zusai/research/disaggevol/>.