

Breaking the Spell with *Credit-Easing*: Self-Confirming Credit Crises in Competitive Search Economies*

Gaetano Gaballo[†] and Ramon Marimon[‡]

February 15, 2014

Abstract

We analyse an economy where banks are uncertain about firms' investment opportunities and, as a result, credit tightness – due to banks' pessimistic beliefs – can result in excessive risk-taking. In the competitive credit market, banks announce credit contracts and firms apply to them, as in a directed search model. The Central Bank can affect banks' liquidity costs by changing its lending rate. We show that high-risk Self-Confirming Equilibria coexist with a low-risk Rational Expectations Equilibrium, in this competitive search economy. Misperceptions never disappear in a Self-Confirming Equilibrium. Lowering the CB policy rate reduces the set of parameters for which a high-risk SCE exists; nevertheless, within this set of economies, this interest rate policy is ineffective. However, a *credit-easing* policy can be an effective experiment, breaking the high-risk (low-credit) Self-Confirming Equilibrium. Since the latter does not arise from a *coordination failure*, the implications of the model differ from models of Self-Fulfilling credit freezes. In particular, we emphasise the social value of experimentation, often neglected in the recent literature that vindicates *robust decision making* as a form of good governance for central banks. (*JEL*: D53, D83, D84, D92, E44, E61, G01, G20, J64)

*Preliminary and incomplete draft. The views expressed in this paper do not necessarily reflect the ones of Banque de France.

[†]Banque de France, Monetary Policy Division. *Email*: gaetano.gaballo@banque-france.fr

[‡]European University Institute, UPF - Barcelona GSE, NBER and CEPR. *Email*: ramon.marimon@eui.eu

1 Introduction

This paper explores the macroeconomic consequences of individual uncertainty about others' agents opportunities and corresponding payoffs. In particular, we study a competitive economy where, in equilibrium, individual financial intermediaries (banks) – possibly, due to a collective ‘bad experience’ – can have pessimistic beliefs about firms' capacity or incentive to make low-risk investments. With such beliefs, banks charge a risk-premium to cover their expected losses. And, facing high interest loans, firms' optimal investment is a risky project. As a result, the behaviour of firms *self-confirms* banks' pessimistic beliefs and these beliefs persist in a high-risk equilibrium. It is in this context that there can be a rationale for *credit-easing* policies: breaking the high-risk (low-credit) Self-Confirming Equilibrium.

1.1 The model components

In our model, the competitive credit market consists of a continuum of firms and banks. The latter post credit contracts and the former apply for these contracts, in order to finance their investment projects. The interest rate on the loan defines the credit contract and, consequently, the type of project and size of the investment that maximises firm's profits. In the simplest version of our model, we assume that, at any point in time, firms can only invest in one project, which can be either a risk-free project, which involves a per-unit cost, or a risky project with zero per-unit cost. Firms' expected profits depend on the probability that their loan application is accepted and, if so, on the expected net return of their project. When a risky project fails firms only repay the principal of the loan (the *pledgeable* part) and, therefore, the lender-bank also bears part of the risk, which is compensated by the loan's interest risk-premium. Firms' have *rational expectations* in making their choice conditional on the existing menu of debt contracts.

Firms in our model can be non-banking financial intermediaries. In particular, intermediaries of Asset Backed Securities (ABS) that by incurring a per-unit monitoring cost can guarantee a ‘safe’ ABS package, while the latter becomes risky if they do not incur the cost. In contrast with models of *Self-fulfilling credit freezes* (e.g. Bebhuk and Goldstein 2011) a firm's project returns does not depend on other firms' financial conditions. As it will see, in our model *credit freezes* can arise even if there is no *coordination failure*.

Banks are financial intermediaries that borrow money from the Central Bank in order to provide loans to individual firms. There is free entry in this industry and banks cannot default on their Central Bank obligations, including CB lending-rate payments (i.e. default is too costly for banks). Banks know their costs with certainty but they are uncertain about their revenues, since they have to anticipate firm's reaction to their credit offers not knowing in which project firms will invest. We weaken the rational expectations hypothesis, with respect to banks, by assuming that their beliefs only need to be ‘locally-rational’ in equilibrium, meaning that they satisfy two conditions: first, as in directed

search-models, marginal variations of loans' interest rates are expected to be compensated by marginal variations on the number of applications; second, banks expect to loose if they offer non-equilibrium debt contracts and these beliefs are locally correct (i.e. for small, non-marginal, deviations); however, they may not be correct for large deviations that would result on a different choice of project. For example, banks may wrongly believe that offering a lower interest rate will not cover their expected losses since they miss-perceive the investments firms will undertake when borrowing at low interest rates. This equilibrium is a *Self-Confirming Equilibrium* and, only when banks' correctly perceive firms' reactions to all their offers, it is a *Rational Expectations Equilibrium*; i.e. only when the second rationality requirement is global. In our model, the *Rational Expectations Equilibrium* is unique.

In our model the Central Bank does not have superior information than private banks; nor has a better commitment technology than them, as it happens in the work of Karadi and Gertler (2011) on credit easing – it does. However, the CB maximise social welfare, not just banks' profits, and has access to society's resources (tax revenues). As we now will argue – in contrast with the work of Chari *et al.* (2010) on credit easing, where CB policy is ineffective – these two classical components of public policy provide a *rationale* for a *credit-easing* policy. Before characterising this policy it is useful to consider more conventional Central Bank interventions.

1.2 The Central Bank policy interventions

The Central Bank affects the banks' cost of liquidity by changing its lending rate. A conventional policy of lowering the CB policy rate reduces the set of economies for which an economy has a high-risk Self-Confirming Equilibrium. For economies within this set, the CB lower interest rate policy is ineffective. For economies outside this set, the same policy can have a radical effect by destroying the wrong perceptions of the banks: only the Rational Expectations low-interest rate equilibrium exists after the policy has been introduced.

In our economies neither for firms nor for banks there is an issue of a shortage of capital or liquidity: firms apply for loans without constraints at the interest rates specified by banks and, similarly, banks can borrow unlimited amounts from the Central Bank at its lending rate. Therefore, there is no role for an unconventional policy of capital injections from the CB to the banking system. If the latter is trapped in a high-risk Self-Confirming Equilibrium such a policy will be completely ineffective: it would only rebalance the respective balance sheets and we have abstracted from capital-requirements regulation issues.

Should the Central Bank pursue an unconventional policy of directly lending to firms? And, if so, how? Given that, as we have said, there are no capital shortages, the only reason could be if by doing so the CB could break a high-risk Self-Confirming equilibrium shifting the economy to the more efficient no-risk Rational Expectations Equilibrium. However, in our economies, the Central Bank does not have any informational advantage, wouldn't the CB have the same mis-perceptions than private banks? In our model the answer is yes; in

fact, it can even be more pessimistic than the private sector. However, as we have already emphasized the CB and private banks have different – social vs. private – objectives. The difference translates into a different assessment of the value of experimentation in the dynamic formulation of the model.

The dynamic model is basically a repeated version of the static model just described. In particular, in the simple version with two types of projects, one safe and one risky, the state of the economy has two components: the per-unit cost of the safe project and the probability of success of the risky project. Miss-perceptions by banks does not mean that they have degenerate beliefs assessing, for example, zero probability to the safe technology being available, but simply that their subjective beliefs are distorted with respect to the objective distribution generating the state of the economy. In our model once self-confirming mis-perceptions have been falsified, the economy remains in the unique rational expectations equilibrium and, therefore, the policy intervention only needs to be implemented temporarily. This was already true in the not-very effective interest rate policy and it justifies the more effective *credit-easing* policy. To understand this policy we first address the issue of *the value of experimentation* and then how the policy can be implemented in our economies.

Free entry in the banking industry implies that if banks (mistakenly) assign a low probability to firms investing in safe projects may have no incentive to ‘experiment’ with low interest rates, even if posting a credit line with very low interest rate will eventually dissipate the misperception. The reason being that the discovery would become immediately available to all firms and the zero-profit condition would be restored the following period. In contrast, a welfare maximising central bank may have a positive value for such experimentation – even if it attaches the same, or even lower, probability to the safe technology – since it values the move to a more efficient equilibrium against the cost of the temporary implementation of the *credit-easing* policy. In fact, this different can be very large if the CB is relatively patient, offsetting the pessimistic beliefs. Furthermore, the CB can be in the position to engage in a large scale experiment, which may reveal the true distribution of projects in one period. Notice that even if banks could coordinate to perform a large scale experiment, they will only anticipate the new zero-profit condition at the rational expectations equilibrium.

However, as we have also emphasized, the Central Bank does not have the technology to intermediate with firms and, therefore, cannot directly lend to them at a low interest rate. Banks must act as the ‘transmission’ for CB policies. The *credit-easing* policy in our economies is a policy of subsidizing banks’ risky loans, creating a *wedge* between the relatively high interest rate charged by banks and the relatively low interest rate faced by firms. This wedge modifies the behaviour of the banks as if they had revised their beliefs as give higher probability to firms investing in safe projects. Such policy, it is only costly while the test is being carried out and banks end up providing loans for risky projects. It is financed with tax revenues¹. For a relatively patient Central Bank it is worth to carry the experiment, even if she has the same

¹For simplicity, in this draft, we assume the large-scale test only takes one period.

misspecified pessimistic beliefs than private banks, provided that she is not a ‘robust decision-maker,’ for whom the test must be successful in the worst case scenario. However, the fact that the CB has to implement its policy through the banking system creates an interesting trade-off: the more aggressive is the policy (i.e. the higher the subsidy) the less costly – but, at the same time, the less potentially beneficial – it is for an individual bank to deviate from the Self-Confirming low-credit equilibrium, although there always exists an aggressive enough *credit-easing* policy that, by temporarily absorbing the risk, destroys the Self-Confirming low-credit equilibrium.

In sum, our theory provides a *rationale* for *credit-easing* policies or – one should better say, ‘experiments’ – since in our economies ‘experimentation’ is a public service. In particular, it can help to explain why these unconventional policies may result in either ineffectiveness or immediate success. A similar dichotomy is present in Bechuk and Goldstein’s (2011) model of credit freezes – to our knowledge, the closest paper to ours – and, therefore, it is useful to have a first contrast of both models.

1.3 The contribution of the paper

In sum, our theory provides a *rationale* for *credit-easing* policies or – one should better say, ‘experiments’ – since in our economies ‘experimentation’ is a public service. In particular, it can help to explain why these unconventional policies may result in either ineffectiveness or immediate success. It is a (complementary) alternative theory to the more standard theory of Self-fulfilling credit freezes. In our view, an example of this type of success has been the Term Asset Backed Securities Lending Facility (TALF) of the FED in the collapse of the ABS market in the 2008-2009 crisis. We discuss the comparison with Self-fulfilling credit freezes and the TALF experience in the next section.

Our work is also novel in bringing the concept of Self-Confirming Equilibrium to a competitive environment. This concept, was pioneered in game theory by Fudenberg and Levine (1993) and in macroeconomics by Sargent (1999). In both contexts, the beliefs of a non-negligible agent are misspecified out-of-equilibrium. This has two consequences. First, the agent’s deviation from the equilibrium path can disrupt the equilibrium and detect the miss-perception; something which is not possible in a frictionless competitive environment. Second, the individual and social value of experimentation is basically the same. We develop a model of *competitive economies with search frictions* where Self-Confirming Equilibria may exist and have predictable power (not everything goes); in particular, public policy experimentation can be a public service.

1.4 The roadmap

In Section 2 we compare our model with a model of Self-fulfilling credit freezes and briefly describe the TALF experiment. In Section 3 we describe, more formally, the components of the model. In section 4 we define and characterise

Self-confirming equilibria. In Section 5 a more detailed analysis of firms’ investment choices allows to sharpen the characterisation of Self-confirming equilibria. Section 6 discusses Central Bank policy interventions and, finally, Section 7 concludes (tbc).

2 Discussion

2.1 The contrast with Self-fulfilling credit freezes

In the B&G model, there are bad firms with useless projects and good firms whose projects can have positive returns if, and only if, the (parameterised) state of the economy is sufficiently good (high) and enough banks lend to firms. When the state of the economy is sufficiently good (bad) all banks lend (do not lend) to firms. For intermediate values, there is a coordination problem among banks and if they have precise information about the state of the economy there two rational expectations equilibria coexist: one with lending and one without. Following the *global-games* approach of Morris and Shin (2004) they consider the case where banks have imprecise information and, therefore, their decision to lend or not is given by a threshold value of the signal received. Correspondingly, there is a threshold value of the state of the economy determining a region characterised by inefficient credit freezes.

As we have seen, in our model there are no positive complementarities among firms and banks and the underlying rational expectations equilibrium is unique. Our (two-variables) state of the economy plays a similar role that the B&G state of the economy. Similarly, the analogy can be extended to consider that our Self-confirming high-risk (low credit) equilibrium correspond – even if they are qualitatively different – to their rational expectations Self-fulfilling credit freezes. Keeping this analogy, both models have a common property: policies that are effective in ‘destroying’ the low-credit equilibrium only need to be implemented temporarily: until they produce the desired ‘cleansing effect’. They also have a fundamental difference: in their model the inefficient Self-fulfilling credit-freeze equilibrium is a manifestation of a *coordination failure*, while in our model the inefficient Self-Confirming low-credit equilibrium is a manifestation of *persistent misperceptions*. As a result there is a ‘structural policy’ fundamental difference: Self-fulfilling credit-freezes can be avoided by concentrating the banking system in the B&G model, while banking concentration does not resolve the Self-confirming equilibrium problem (as long as expected profits are zero), while more competition in the banking system may weaken the *persistent misperceptions* problem, but worsen the *coordination failure*². Regarding the three mentioned policy interventions of the Central Bank a brief comparison of their effects is as follows:

Lowering the lending rate: it has a similar effect in both models: for certain states of the economy the policy destroys the low credit equilibrium, while

²This statement is a conjecture at the current stage of our work.

for other (not at the margin) states the policy is ineffective.

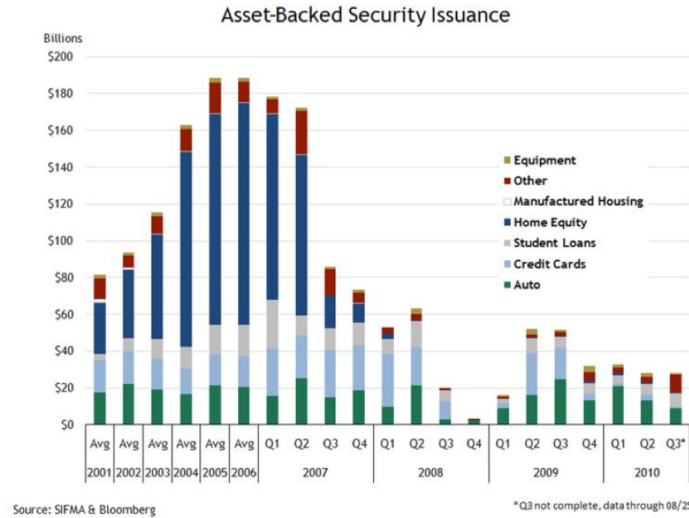
Capital injections into the banking system: in B&G Central Bank can be effective since they raise the amount of lending and, therefore, they also reduce the set of economies for which a Self-fulfilling credit freeze equilibrium exists. In contrast, in our model as we have already mentioned such a policy is ineffective.

Credit-easing: while both models vindicate this policy as the most effective in reducing the probability of a low-credit equilibrium, there is an important policy difference while, in terms of efficiency, *Credit-easing* can be dominated by a policy of capital injections into the banking system in B&G, as we have seen this can never happen in our economies. Furthermore, one expects that it should be ‘less costly’ to resolve a *persistent misperceptions* problem than a *coordination failure* with *credit-easing* this is an empirical (or properly simulated) question, for which modelling differences should be resolved. Some of this differences are: in B&G the Self-fulfilling credit freeze is associated with low risk in financial markets (in fact, only riskless government bonds are traded), while in our low-credit Self-confirming equilibrium is characterised by high risk investments; in B&G an economy falls into a Self-fulfilling credit freeze equilibrium as a result of an adverse shock to banks’ capital, in ours into a Self-confirming low-credit equilibrium because of an increase in risk in the financial sector (possibly more consistent with the evidence); in B&G the policy is implemented by directly lending to firms and the CB is prevented from substituting the private bank system by assuming that it cannot discriminate between bad and good firms as private banks do while, in our model as we have seen, the CB cannot substitute the private banks and, therefore, implements the policy through them (which seems closer to the evidence; see below the TALF example). Finally, in our model *Credit-easing* is an explicit policy experiment in which cost and benefits are assessed and the different public-private valuations compared as part of the policy design, while in their model the policy design problem is not addressed (a common feature of Self-fulfilling equilibrium models).

In sum, Self-Confirming and Self-Fulfilling inefficient low-credit equilibria are complementary theories for the same phenomena. With important differences – the main being the underlying mechanism – both vindicate the unconventional policy of *credit-easing*, but it is beyond the scope of this paper to provide a more definite comparison among the two. Nevertheless, the following recollection of the main *credit-easing* policy experiment can help to see strengths and weakness of both theories.

2.2 The collapse of the ABS market and the TALF policy

In the 2008-2009 crisis of the Asset Backed Securities (ABS) market in US. In the second half of 2007 the ASB market has experienced a sudden contraction



after a constant increase in volumes since early 2000.³ The crash was mostly driven by lower-than-expected returns in the housing markets which depressed the value of subprime home equities. The dramatic increase in perceived risk and the lack of confidence in rating agencies did even not spare an abrupt freeze of the AAA-rated ABS segment whose interest-rate rose at exceptionally high levels reflecting unusually high risk premiums⁴. As a consequence private liquidity collapsed rapidly, and investors directed available resources to quality assets like treasury bills which almost doubled their daily volumes of trade from 40 to 80 USD billions during 2008-2009.

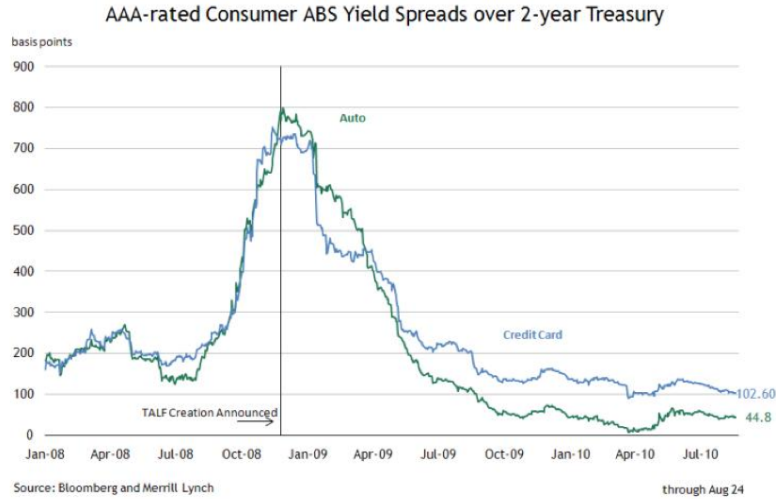
Within this context the FED stepped in with the launch of the Term Asset Backed Securities Lending Facility (TALF) which supplied about 71 billions of non-recourse loans⁵ at lower interest rates, to any U.S. company provided of highly rated (AAA and AAA-) collateral. This intervention was made primarily to sustain the credit market in a period of high perceived counterpart risk. More precisely, the FED acted as a borrower of last resort taking the risk of experimenting contractual conditions which were perceived as too risky by the private sector.

Nevertheless, despite malign prophecies welcoming the birth of the programs, on the 30th of September 2010, the FED announced that more than 60 percent

³New issuances of consumer ABS plunged from \$50 billion per quarter of new originations in 2007 to only \$4 in the last quarter of 2008.

⁴"This can be seen in the fact that AAA-rated student loan tranches, with underlying loans 97% guaranteed by the federal government, climbed to yield levels as much as 400 basis points over LIBOR." Speech by William C. Dudley, President and Chief Executive Officer on the 4th. of June 2009.

⁵Meaning that if the economy performs very badly and the securities fall sharply in value, an investor can put the collateral that secures its TALF loan back to the Fed, only losing the collateral haircut.



of the TALF loans have been repaid in full, with interest, ahead of their legal maturity dates. In other words, the more favorable conditions offered by the FED, instead of triggering an adverse selection process, have been the prelude of a remarkable business performance and expansion of consumption credit.

Should we conclude then that the market was mispricing highly rated ABS? Not necessarily in the sense that probably at such high interest rates failures to repay the loans could have been much more likely. However, a real risk psychosis was preventing the private sector from experimenting and so revealing the existence of profit opportunities in the high quality segment of the ABS market. In the absence of the FED intervention such bias could have not been corrected, self-sustaining a suboptimal outcome with major consequences for the already tatty American economy.

3 General framework

3.1 Environment: Actions and Timing

A continuum of firms of mass one look for credit to implement investment opportunities. Atomistic banks borrows money in the interbank market, or equivalently at a rate determined by the central bank (CB), and lend to firms. Both are risk neutral. Firms liability is limited to principal component whereas banks cannot default on the interbank (or CB) lending. Therefore when a project fails, which occurs with an exogenous probability, a bank loses interest repayments whereas it is enforced to repay the interests on its own loan. The return on a project can be secured adopting a possibly more costly type of investment. A firm adopts the type of investment in its own interest according

to its payoff structure which depends on the state of the world. Nevertheless the type of project adopted is not observable, hence banks cannot screen firms for project quality. The timing in the market is the following:

1. Nature selects the state of the world ω drawing from a set of possible states Ω : each state of the word characterizes a different payoff structure of firms;
2. a bank can borrow at a rate R_{CB} , controlled by the central bank, in the interbank market;
3. a bank pays a cost c to post a credit contract by which it commits to lend any amount at a fix chosen rate R ;
4. a firm chooses to which posted R to apply for credit;
5. once a match is formed a firm chooses the investment policy depending on the interest rate of the contract R and which state has realised ;
6. if the project is successful a firm pays back interest and principal to the bank, and only principal otherwise;
7. banks pay back their loan irrespective of the success of the project financed.

Banks bear two kind of risks: one is associated with the probability that a vacancy is not filled, one other originating from the partial enforceability of the posted contracts. An entrepreneur (firm) instead does not incur any cost if she does not match or if her project fail. Nevertheless, the exposure to risk of banks depends on firms' choices. In particular, to optimally solve thier problem banks need to anticipate the probability that a posted contract finds a match and firms' reaction once matched, that is they need to identify in which state of the word they act. Let us now describe the matching technology and then the payoffs structure of firms and banks.

3.2 Directed search

The matching framework presented here closely follows the competitive setting introduced by Moen (1997) along a simplified variant described by Shi (2006). The probability that a match is formed is described by a matching function which is a map $\mathfrak{R}_+^2 \rightarrow \mathfrak{R}_+$ from a couple (u, v) - being respectively the measure of illiquid firms and the measure of vacant credit lines - to $x(u, v)$ being a flow of new firm-bank matches. The matching function $x(u, v)$ encapsulates a search friction assumed in the competitive credit market. Following standard assumptions, let x be concave and homogeneous of degree one in (u, v) with continuous derivatives. Let $p = x(u, v)/u = x(1, \theta) = p(\theta)$ denote the transition rate from illiquid to liquid for an illiquid firm, and $q = x(u, v)/v = q(\theta)$ the arrival rate of firms for an open credit line, where θ is the credit market tightness u/v .

Let $\lim_{\theta \rightarrow 0} p(\theta) = \lim_{\theta \rightarrow \infty} q(\theta) = 1$ and $\lim_{\theta \rightarrow \infty} p(\theta) = \lim_{\theta \rightarrow 0} q(\theta) = 0$. For the sake of simplicity and without loss of generality, we will assume that the matching function has a Cobb-Douglas form

$$x(u, v) = Au^\gamma v^{1-\gamma}$$

so that $p(\theta) = A\theta^{\gamma-1}$ and $q(\theta) = A\theta^\gamma$. This assumption, which is standard in the literature, ensures a constant elasticity to the fraction of vacancies and illiquid firms.

The search is *directed*, meaning that at a certain interest rate R there is a subset of illiquid firms and banks with open credit lines looking for a match at that specific R . The number of matches in the submarket R is $x(u(R), v(R))$, where $u(R)$ is the measure of firms and $v(R)$ the measure of vacant credit lines searching for a counterpart in a credit contract at an interest rate R . The arrival rates of trading partners for firms and banks in this market are thus $p(\theta(R))$ and $q(\theta(R))$, respectively, where $\theta(R) = u(R)/v(R)$ is the specific tightness associated to the submarket R . Both firms and banks are free to move between submarkets. Once the match is formed any amount of credit is provided at a rate R . We will say that a submarket is active if there is at least a vacancy posted.

3.3 Firms and Banks

Firms choose to which posted contract $R \in H$ to send its application for funds. Once matched at the targeted R a firm implement its investment policy $f(R, \omega)$, namely, a vector of optimal choices in response to a couple (R, ω) . The objective of a firm is to maximize ex-ante profit

$$J(R) \equiv p(R) \Pi(f(R, \omega)),$$

where $p(R)$ is the probability of having an application accepted at a rate R and $\Pi(f(R, \omega))$ is the expected profit from the investment which depend on $f(R, \omega)$.

Banks are first movers in the search: they choose whether to enter in the market and eventually post a credit line in a submarket. A credit line is a contract by which banks commit to lend any amount of liquidity to successful applicants at a fix rate R . The ex-ante value of a credit line for a bank is given by

$$V(R, \omega) \equiv q(R) Y(R, f(R, \omega)), \quad (1)$$

where $p(R)$ is the probability of having an application accepted at a rate R and $Y(R, f(R, \omega))$ is the expected return on the credit which depend on R and the investment policy of the firm. Therefore

$$\max_R E^\beta [V(R, \omega) - c], \quad (2)$$

where β is the system of subjective beliefs held by banks on the realization of the state of the word, and $c > 0$ an exogenous cost associated with the opening

of a credit line. In particular, let define the E^β operator as follows

$$E^\beta[(\cdot)] \equiv \int_{\Omega} (\cdot) \beta(\tilde{\omega}) d\tilde{\omega},$$

where $\beta(\tilde{\omega})$ is the banks'subjective probability density function on the random variable $\tilde{\omega}$ representing⁶ the state of the world that Nature selected. Notice that for a bank to solve (2) it has to anticipate the reaction of firms $f(R, \omega)$ to the posted R . In equilibrium free entry requires $E^\beta[V(R, \omega) - c] = 0$.

3.4 Determination of the tightness

The tightness is a ratio representing the number of firms looking for a credit line *per-unit of vacant open lines*. This means that the tightness is independent of the absolute number of vacancies open in a certain market. The matching function is just a one-to-one map between p and q through a ratio θ . In other words, suppose an equilibrium is associated with a particular probability to obtain credit $\bar{p}$, then the matching function gives a $\theta = p^{-1}(\bar{p})$ and so a $\bar{q} = q(p^{-1}(\bar{p}))$ that is a probability of filling a vacant line in that submarket. The latter argument is independent on how many vacancies are open in that particular submarket. With a single vacancy open (resp. a measure ε of vacancies open), θ (resp. $\varepsilon\theta$) will be the expected number of firms searching in that submarket.

What then determines the tightness? The tightness of the market is determined by the rational behavior of the firms which act after banks publicly announce their contracts. In particular, consider the case where two different offers R' and R'' are publicly announced. In case $J(R') > J(R'')$ then firms will be more willing to send applications to get the R'' contract rather than the R' . As a consequence of a larger number of applications in the submarket R' , the probability of matching $p(R')$ must decrease lowering $J(R')$. Symmetrically, as a consequence of a smaller number of applications in the submarket R' , the probability of matching $p(R'')$ must increase enhancing $J(R'')$. Therefore rationality from the side of firms implies that, for a given set of posted contracts $H = \{R_1, R_2, R_3, \dots\}$, the tightness is determined by

$$\bar{J} = p(\theta(R, \bar{J})) \Pi(f(R, \omega)) \quad (3)$$

where $J(R) = \bar{J}$ is constant for each $R \in H$. In particular, $\bar{J}$ is the ex-ante utility that firms expect from participating to the market. Therefore there exists a unique tightness associated to each interest rate R which is conditional to a given level of ex-ante utility guaranteed by the participation of firms to the market.

⁶We adopt the convention that a tilde is used to denote a random variable as opposed to a realization.

4 Equilibria

4.1 Competitive SCE and REE

In this section we provide and discuss a definition self-fulfilling equilibrium (SCE) that goes beyond the specific payoff structure that we will analyse later. Then we will relate SCE to rational expectation equilibria (REE).

Definition 1 *For a given $\omega \in \Omega$, a self-confirming equilibrium (SCE) is a set H^* of interest rates such that, for each $R^* \in H^*$:*

(i) *firms maximize expected profits*

$$R^* = \arg \sup_{R \in H^*} p(\theta(R, J^*)) \Pi(f(R, \omega)) \quad (4)$$

where $J^* = p(\theta(R^*, J^*)) \Pi(f(R^*, \omega))$;

(ii) *banks maximize expected profits*

$$R^* = \arg \sup_{R \in \mathfrak{R}} E^\beta [q(\theta(R, \bar{J})) Y(R, f(R, \omega))] \quad (5a)$$

$$\text{s.t. } \bar{J} = p(\theta(R, \bar{J})) E^\beta [\Pi(f(R, \omega))] \quad (5b)$$

where $J(R) = \bar{J}$ is constant for each $R \in \mathfrak{R}$;

(iii) *banks correctly anticipates liquidity demand, and so the corresponding type of investment, at any local deviation from an equilibrium contract, that is*

$$E^\beta [f(R, \omega)] = f(R, \omega)$$

for any $R \in \mathfrak{S}(R^*)$ where $\mathfrak{S}(R^*) \subset \mathfrak{R}$ is a neighborhood of R^* .

The first requirement implies optimality from the side of the firms so that (17) defines the tightness of the submarkets. Notice that $\theta(R, J^*)$ depends on the ex-ante utility granted to firms at equilibrium conditions, and does not depend on individual choices.

The second condition requires that a bank posts a R^* that globally maximizes its expected value of a credit line. The relevant expectation is the one conditional to the bank subjective beliefs summarized by β .

The third condition restricts banks' beliefs about firms' actions to be correct in a neighborhood of an equilibrium R^* . This is a stronger beliefs' restriction of the one usually assumed in the directed search literature to get rid of trembling-hand imperfect equilibria, which involves only first-order deviations.

Importantly notice that the definition of a self-confirming equilibrium does not require banks to have correct beliefs about out-of-equilibrium behavior that is far away from the equilibrium considered. This leaves open the possibility that a self-confirming equilibrium banks are not actually maximizing in the *whole* domain of their actions. In fact banks' unbiased beliefs about firms' payoffs are only required limited to the neighborhood of the contract that is

implemented, and whose effects are therefore observed. Dominant contracts out of a neighborhood of the equilibrium could be wrongly believed by banks to be strictly dominated. Since such contracts will be never posted, then in equilibrium there do not exist counterfactual observations that could confute wrong beliefs.

A REE is a stronger notion than a SCE requiring that no agent holds wrong out-of-equilibrium beliefs. In the present model this equals to impose that banks' unbiased beliefs about firms' payoffs. In such a case the equilibrium contract is the one which objectively yields the highest reward with respect to every possible feasible contract.

Definition 2 *A rational expectation equilibrium (REE) is a self-confirming equilibrium H^* for which:*

iii-bis) banks correctly anticipates liquidity demand, and so the corresponding type of investment, at any feasible R , that is,

$$E^\beta [f(R, \omega)] = f(R, \omega),$$

for any $R \in \mathfrak{R}$.

A REE obtains from a tightening of condition (iii) in the definition of a SCE. This implies that every $R^* \in H^*$ is such that banks can exactly forecast their payoffs out of the equilibrium, as they can correctly anticipate firms' responses. Therefore condition ii) of the definition of a SCE becomes

$$R^* = \arg \sup_{R \in \mathfrak{R}} q(\theta(R, \bar{J})) Y(R, f(R, \omega)) \quad (6)$$

subject to $\bar{J} = p(\theta(R, \bar{J})) \Pi(R, \omega)$, in the case of a REE. That is, posting in the submarket R^* is a globally dominant strategy both from an objective and a subjective point of view.

4.2 Characterization of the Equilibria

The characterization of a self-confirming equilibrium of the model it is given by the following.

Proposition 1 *Consider two credit lines posted respectively at R_1 and R_2 . From the point of view of a single atomistic bank*

$$E^\beta [V(R_1, \omega)] \geq E^\beta [V(R_2, \omega)]$$

if and only if

$$E^\beta [\mu(R_1, \omega)] \geq E^\beta [\mu(R_2, \omega)] \quad (7)$$

with

$$\mu(R, \omega) \equiv \Pi(f(R, \omega))^{\frac{\gamma}{1-\gamma}} Y(R, f(R, \omega)),$$

for any profile of contracts offered by other banks.

Proof. Postponed to the appendix.

Corollary For a given ω , a set of contracts H^* is a SCE if for a given system of subjective beliefs β , any $R^* \in H^*$ is such that

$$\mu(R^*, \omega) \geq \mathbb{E}^\beta [\mu(R, \omega)], \quad (8)$$

for any $R \in \mathfrak{R}$, and

$$\mu(R^*, \omega) \geq \mu(R, \omega), \quad (9)$$

for any $R \in \mathfrak{S}(R^*)$. A SCE is a REE if and only if (8) and (9), both hold for any $R \in \mathfrak{R}$.

Notice one important feature of the condition above. With $\gamma = 0$ when all the surplus is extracted by banks (9) becomes $Y(R^*, f(R^*, \omega)) \geq Y(R, f(R, \omega))$, that is at the equilibrium only the interim payoff of banks is maximized as firms will always earn zero. With $\gamma = 1$ instead when the whole surplus is extracted by firms (9) becomes $\Pi(R^*, \omega) \geq \Pi(R, \omega)$, that is only the interim payoff of firms is maximized as banks will always earn nothing. Of course, (9) is satisfied locally by any interior SCE. In other words, (9) is a condition that implies the local holding of Hosios (1990) condition.

Proposition 2 An interior (i.e. when no participation constraints are binding) SCE is characterized by a single contract $H^* = \{R^*\}$ satisfying

$$\gamma \frac{\Pi'(f(R^*, \omega))}{\Pi(f(R^*, \omega))} + (1 - \gamma) \frac{Y'(R^*, f(R^*, \omega))}{Y(R^*, f(R^*, \omega))} = 0, \quad (10)$$

for a given parameterization of the functions $\Pi(f(R, \omega))$ and $Y(R, f(R, \omega))$.

Proof. Postponed to the appendix.

The proposition above states that a self-confirming equilibrium is such that all the matches occur at a unique R^* . Interior SCEs are *locally* optimal in the sense of the Hosios (1990) condition. In particular, at a SCE we have that

$$1 - \gamma = \frac{Y(R^*, f(R^*, \omega))}{Y(R^*, f(R^*, \omega)) - \frac{Y'(R^*, f(R^*, \omega))}{\Pi'(f(R^*, \omega))} \Pi(f(R^*, \omega))},$$

that is, the fraction of the surplus (properly evaluated) going to banks - the term on the right-hand side - reflects the elasticity of the matching function with respect to the fraction of illiquid firms in the market. This is exactly the condition for which banks internalize the social cost of opening a new vacancy.

5 A model of the credit market

In this section we explicitly model the payoff structure of firma and bank that can generate a SCE. Once in a match a firms chooses the size I and the type ς of investment. The expected profit of a firm is

$$\Pi(R, f(R, \omega)) \equiv \kappa(\varsigma, R, \omega) I - \frac{1}{2} I^2,$$

where $\kappa(\varsigma, R, \omega)$ is the expected per-unit gross return on the investment size I which is implemented paying a quadratic cost of $I^2/2$. A participation constraint imposes $\Pi \geq 0$. Firms can choose between two kinds of projects $\varsigma \in \{s, r\}$, respectively a safe and a risky one, which differ for the likelihood of success and per-unit adoption cost. Both types have the same gross per-unit return: in case of success is $1 + y$, whereas 1 in case of failure. Safe projects do not fail, but their adoption requires a fix per unit cost of k which sums up to the repayment of the bank's loan. Risky project do not have any fix per-unit additional cost, but they are successful only with a probability $\alpha \in (0, 1)$. Hence, a *risky* project gives a net per-unit expected return of

$$\kappa(r, R, \omega) = (y - R) \alpha,$$

whereas a *safe* project gives

$$\kappa(s, R, \omega) = y - k - R.$$

Let therefore complete our structure defining a state of the world as the couple of coefficients $\omega \equiv (\alpha, k)$, with $\Omega \equiv (0, 1) \times \mathfrak{R}_+$, shaping the payoffs of firms. In particular, the associated FOC for the demand of investment determines the investment policy being a mapping $f : \mathfrak{R}_+ \rightarrow \Theta \times \mathfrak{R}_+$ giving for each $R \in \mathfrak{R}_+$ a couple $f(R, \omega) = [\tau(R, \omega), I(R)]$ such that

$$I(R) = \tau(R, \omega) = \arg \max_{\varsigma \in \{s, r\}} \kappa(\varsigma, R, \omega), \quad (11)$$

as an optimal reply to an offer R . The choice of firms affect banks' expected profit

$$Y(R, f(R, \omega)) \equiv I(R) \pi(\tau(R, \omega), R, R_{CB}).$$

which depends on the size I of the loan and the net per-unit return π on the loan. The latter is determined by: i) τ , the optimal technology choice of firms in response to R , ii) R_{CB} , the cost of liquidity determined by the central bank, and, of course, iii) on the posted return R . The type of technology determines the degree of pledgeability of the investment, and so the expected repayment rate of the loan. In case of a risky type of project only a fraction α of firms will be able to repay back $I(1 + R)$, the whole loan, whereas $(1 - \alpha)$ will just repay the only pledgeable part, the principal. Therefore the expected per-unit return on the loan is given by

$$\pi(r, R, R_{CB}) = \alpha R - R_{CB},$$

and

$$\pi(s, R, R_{CB}) = R - R_{CB}$$

were notice that a bank cannot in any case default on the central bank loan, so that its repayment does not depend on α . Market participation of both a firm and a bank require $R \in (y, R_{CB}/\alpha)$.

Banks and firms act in a world in which ω is not random. Nevertheless depending on the set of contract is posted, firms' reaction will reveal either α or k . In particular, for a given R , the firm will adopt a safe technology if it will make higher profits out of it. For a given ω , this is the case when the return of the safe technology is sufficiently high, so that $\kappa(r, R, \omega) \leq \kappa(s, R, \omega)$, or equivalently, the offered contract is sufficiently low, that is

$$R \leq \bar{R} \equiv \max \left\{ y - \frac{k}{1 - \alpha}, R_{CB} \right\},$$

where for banks to make offers it has to be $R \geq 0$. In fact the cost of liquidity does not affect the choice of projects, but constraint the offer of credit lines. We will refer to $\bar{R}$ as the adoption frontier. As a consequence banks' uncertainty about one of this coefficient can survive. This feature makes SCE possible outcomes.

The following proposition describes the set of REE.

Proposition 3 *For given ω and R_{CB} , there exists a unique threshold value $\hat{\alpha}(k) \in (\underline{\alpha}, \bar{\alpha})$, with*

$$\underline{\alpha} = \frac{y - \hat{R}_s - k}{y - \hat{R}_s} \quad \text{and} \quad \bar{\alpha} = \frac{y - R_{CB} - k}{y - R_{CB}},$$

which is strictly decreasing in k such that:

(i) if $\alpha < \hat{\alpha}$ then there exists a unique "safe" REE characterized by $R_s^ \equiv \min(\bar{R}, \hat{R}_s)$ where*

$$\hat{R}_s \equiv \frac{1 - \gamma}{2} (y - k) + \frac{1 + \gamma}{2} R_{CB}; \quad (12)$$

(ii) if $\alpha \geq \hat{\alpha}$ and $\hat{R}_r > \bar{R}$, then there exists a "risky" REE characterized by $R_r^ \equiv \min(r, \hat{R}_r)$ where*

$$\hat{R}_r \equiv \frac{1 - \gamma}{2} y + \frac{1 + \gamma}{2\alpha} R_{CB}; \quad (13)$$

(iii) a "safe" REE and a "risky" REE both exist when $\alpha = \hat{\alpha}$ or

$$\frac{R_{CB}}{y} \geq \alpha \geq \frac{y - k - R_{CB}}{y - R_{CB}},$$

which requires $y \geq (y - R_{CB})^2/r$. In the latter both equilibria are degenerate, that is, $(R_s^*, R_r^*) = (R_{CB}, y)$.

Proof. Postponed to the appendix.

$\hat{R}^s$ and $\hat{R}^r$ represent interior solutions, namely they are the contracts which locally maximize banks' profits when no participation constraints are binding; R_s^* and R_r^* instead account for the possibility that participation constraints bind. In particular, $R_r^* > R_s^*$, that is, ceteris paribus, risky projects imply higher interest rates. Nevertheless, the profit of a firm and of a bank can be higher when a risky project is implemented depending on parameters. Notice that at the risky equilibrium banks' participation constraint $\alpha R_r^* - R_{CB} \geq 0$ is always satisfied whenever firms' participation constraint $r - R_r^* \geq 0$ is too: in fact $r \geq R_r^*$ implies $r \geq R_{CB}/\alpha$ which in turn yields $R_r^* \geq R_{CB}/\alpha$.

The proposition below states the possibility of a SCE

Proposition 4 *Given ω and R_{CB} there is a sufficiently high $E^\beta[k]$ such that for $\alpha < \hat{\alpha}$ a unique SCE that is not REE exists characterized by $R_r^* = \min(y, \hat{R}_r)$ with $\hat{R}_r > \bar{R}$. Otherwise only REE exist.*

Proof (sketch). By construction it is $\mu(R_r^*, \omega) \geq \mu(R)$ for any $R \in \mathfrak{S}(R_r^*)$ where $\mathfrak{S}(R_r^*)$ is a neighborhood of R_r^* . Since $\alpha < \hat{\alpha}$ then there exists at least a $R' \leq y - k/(1 - \alpha)$ such that $\mu(R_r^*, \omega) < \mu(R', \omega)$. Still this does not prevent $\mu(R_r^*, \omega) \geq E^\beta[\mu(R', \omega)]$ since for any $R \in \mathfrak{S}(R_r^*)$ with $R_r^* > 0$, $f(R, \omega)$ does not reveal the realization k . In particular, $\mu(R', \omega)$ is weakly decreasing in k from which the proposition.

On the other hand, there could not exist a safe SCE that is not REE. Suppose such an equilibrium exists, then it would arise as a corner solution posted at the frontier $\bar{R}$ because interior safe SCE are always REE. Nevertheless, by definition of a SCE, agents would have correct beliefs for marginal deviations from the equilibrium that in this case would provide information about the actual α . Therefore at a SCE posted along the frontier $\bar{R}$ agents would know the actual α . Hence banks can correctly forecast $f(R, \omega)$ at any R , and so they cannot sustain a safe SCE that is not a REE. A contradiction arises. ■

Figure 1 illustrates a baseline configuration of the economy. Let us firstly focus on panel A. The feasible range of equilibrium interest rates compatible with the adoption of a safe (risky) technology is the region below (above) the dotted curve representing the adoption frontier of firms. For any α value R_r and R_s are denoted by respectively the upper and lower solid/dashed lines. In particular, the solid line denotes the unique REE. For $\alpha < \hat{\alpha}$ a risky SCE coexists with a safe REE (the threshold $\hat{\alpha}$ is denoted by a vertical dotted line).

Panel B plots the corresponding levels of aggregate investment for the REE and SCE equilibria, measured in terms of cost-per-vacancy c .⁷

Panel C and D illustrate the individual maximization problem of a single bank for a specific value $\alpha = 0.7$ when all the others post equilibrium contracts. The REE is the safe equilibrium whereas the SCE is the risky one. This is evident from the inspection of panel C. The dotted line denotes the actual payoff that a bank would obtain conditional on all other banks posting at the risky equilibrium. The risky SCE equilibrium contract R_r corresponds to a local maximum of the dotted line where $V(R_r) - c$ takes value zero due to free entry. Posting R_r is a *locally*-optimal action since marginal deviations from that contract would produce ex-ante negative profits. Nevertheless, the risky equilibrium is not a REE because lowering the interest rate up to the point where firms will adopt safe projects would yield a strictly positive ex-ante profits. The solid line in panel C denotes the actual payoff that a bank would obtain conditional on all other banks posting at the safe equilibrium. The safe SCE equilibrium contract R_s corresponds to the absolute maximum of the solid line where $V(R_s) - c$ takes value zero as a consequence of free entry. Posting R_s is a *globally*-optimal action since any deviation from such contract would produce ex-ante negative profits.⁸

Nevertheless, the risky SCE equilibrium, and not the safe REE, can be sustained for sufficiently pessimistic beliefs on the value of k . For the sake of clarity, let me provide an extreme example. Suppose banks believe that with probability one $k = 0.007$ instead of $k = 0.0042$. Such a beliefs in fact is never confuted by observable produced at the risky SCE where no firm will implement the safe technology. The case with $k = 0.007$ is displayed, with the same convention of figure 1, in figure 2. Notice that the risky SCE in figure 1 is exactly the unique REE in figure 2. In fact as long as safe project are not adopted in equilibrium, the two economies are observationally equivalent: not only at the risky equilibrium but also for any marginal deviation from that equilibrium. The two pictures only differ for non-marginal out-of-equilibrium deviations which could trigger, in the first case but not in the second, a change in the type of investment.

⁷The aggregate investment is increasing in α (and so decreasing in R_r) when the economy is on a risky equilibrium, whereas it is decreasing whenever $\alpha \in (\underline{\alpha}, \bar{\alpha})$ for which the safe equilibrium arises as a corner contract constrained by the firms' profitability constraint. In the other regions where corner contract arises, namely when $R_r = y$ and $R_s = R_{CB}$, the economy displays no aggregate investment because, respectively firms and banks, are indifferent to participating or not to the market. For values of $\alpha < \underline{\alpha}$ instead, when the equilibrium contract is $\bar{R}_s$ the aggregate investment is insensitive to α .

⁸Concluding on the description of figure 1, Panel D plots the probabilities $q(R)$ (decreasing in R) and $p(R)$ (increasing in R) at the SCE and the REE. Their evolution reflect the effects of directed search. For a given equilibrium, the higher (lower) the interest rate posted by a bank, the lower (higher) the probability of filling that vacancy, and the higher (lower) the probability of firms obtaining funds at that rate.

6 Policy interventions (preliminary)

The aim of this section is to provide a preliminary discussion of policy interventions. In the first subsection we build up a simple example to clarify that a SCE does not exclude that banks put a positive probability on a true state of the world. It just requires that they expect that deviations would generate losses which discourage individual deviation. We will argue that, in such a situation, there does not exist any incentive-compatible form of cooperation that banks could undertake to fill their uncertainty.

6.1 The market and the problem of experimentation

Suppose for example, that, everything being the same, there exist two distinct states of the world characterized by two different values of k : low $k_l = 0.0042$ which corresponds to the economy in figure 1, and high $k_h = 0.007$ whose case is depicted, everything being equal, in figure 2. The best individual deviation from the risky SCE, namely offering the safe REE contract, yields an expected net profit of about $2.5c$ in the first case, whereas, a loss of about $-c$ in the second one. Consider the case where at the beginning of time nature selected k_l . The realization k_l is known to banks unless they post contracts in the safe adoption region. Let us investigate which beliefs can sustain a risky SCE that is not a REE.

The lines in the left panel of figure 5 display the expected payoff of a single bank, when all other banks post contracts at the risky SCE, for three cases which differ for the beliefs of banks about k . This difference shows up for interest rates lower than $\bar{R}$ the adoption frontier. The red dashed curve corresponds to the case when banks put zero probability on k_l which mimics the case of figure 2. The blue dashed curve denotes the case when banks put probability one on k_l in analogy to figure 1. Finally the solid blue is the one for which the expected profit from the best individual deviation in the safe territory equals the profit at the risky equilibrium. This curve obtains exactly when banks attach a probability 0.213 to k_l . More pessimistic banks would sustain a risky SCE that is not REE. For more optimistic beliefs instead only a safe REE exists.

To sum up, the low-fix-cost economy is the actual state of the world but banks believe it with a probability $\zeta < 0.213$: then conforming to the SCE prescriptions is a dominant action from the point of view of a subjective Bayesian agent. Nevertheless the absence of observable counterfactuals in equilibrium prevents learning about the true state. In this sense a SCE does not require banks ignoring the presence of safe investment opportunities; it instead requires banks wrongly believe that with high probability is too costly for firms implementing safe projects.

From the point of view of a bank, the evaluation of the benefits of a deviation concerns the expected return of one period only. In fact, independently from the

outcome of the deviation, which is publicly observable, free entry will guarantee zero expected profits on the best contract one period after. Free entry also deters cooperation among banks to experiment new markets since at the equilibrium banks have no resource to devote to subsidy eventual "explorers".

On the other hand, firms do not have incentive to reveal the truth as long as debt contracts cannot depend on past actions (i.e. banks cannot punish or reward firms for their past actions). More precisely, firms always have incentive to signal that they will implement a safe project - irrespective of the state of the world - to obtain a lower interest rate and so higher profits. Once the banks post a contract they cannot renegotiate the contract once the demand for funds has been realized. Unless banks do not cooperate to implement a collective punishment strategy - which however could well be not feasible⁹ or simply too costly (again because of free entry) - firms can exploit banks' confidence.

6.2 Conventional Policy: lowering the cost of money

Let us firstly consider the conventional policy of reduction of the cost of money. Figure 3 illustrates the effect of lowering the CB interest rate from 0.01 to 0.005 in the economy described by figure 1. A lower cost of money reflects in lower interest rates on equilibrium contracts. In particular, such a policy can reduce the set of economies where a risky SCE exists. This happens when the optimal risky contract became bounded by the profitability constraint of firms. In fact, in panel A and B, the dotted blue line exhibits a discontinuity in the range of α for which $\hat{R}^r \leq \bar{R}$. This is a case in which the best risky contract lies at the outside limit of the firms' profitability frontier (i.e. posting a $R' = \bar{R} + \varepsilon$ with $\varepsilon > 0$ infinitesimal), but posting a contract at the inside limit of the frontier (i.e. posting a $R' = \bar{R} - \varepsilon$ with $\varepsilon > 0$ infinitesimal) gives a higher payoff. This is plotted in panel C. Therefore the set of risky SCE is reduced by a decrease of the CB interest rate. At the same time, the aggregate investment at the REE increases.

Nevertheless, the effect of a decrease in the cost of money can be offset by a decrease in project return r . This scenario is illustrated by figure 4. r is lowered from 0.03 to 0.02, all the rates shift downwards but now the frontier $\bar{R}$ also follows the move. The scenario is qualitatively the same as in figure 1. Notice that in this particular example the risky REE arises as a corner equilibrium contract along the firms' profitability frontier.

6.3 Credit-Easing through banks' subsidies

The revelation of firms' true incentives through the opening of out-of-equilibrium credit markets breaks self-confirming coordination failures, correcting wrong beliefs. The result can be achieved though a credit-easing policy, that is direct lending from the CB to the firms. Nevertheless, it is realistic to assume that the

⁹For example if firms are active just one period.

central bank does not have lending facilities, so that it has to "use the market", that is, the policy should provide the right incentive for banks to maintain low interest rates. In particular we will explore the introduction of a subsidy s such that the payoff of banks which funds risky projects is

$$\pi(r, R, R_{CB}) = (\alpha + \text{sub}) R - R_{CB},$$

whereas the payoff of a bank funding safe project $\pi(r, R, R_{CB})$ is unaffected.

The right panel of figure 5 plots the expected payoff of a single bank, when all other bank post contracts at the risky SCE, in the case of a subsidy of $\text{sub} = 0.1$ for the same three cases and with the same conventions of the left panel. Notice that the subsidy has two contrasting effects. On one hand, a subsidy reduces the loss in the k_h scenario, so it increases the value of an individual deviation into the safe territory (that is, the red curve is higher). On the other hand, it reduces the distance between the equilibrium value of the risky interest rate - the risky option is less risky because of the subsidy - and the adoption frontier (the dashed blue curve decreases). This last effect implies a lower competitive advantage for a single "deviator", meaning less firms will apply to the out-of-equilibrium offer, and so value of an individual deviation lowers.

The net result of this trade-off is depicted in figure 6. The curve assigns for each value of the subsidy sub the minimal probability that banks put on k_l that makes them to sustain a risky SCE. A partial subsidy has initially a positive effect since it lowers the minimal belief that sustain a SCE. Nevertheless, as the subsidy approaches the totality of expected losses, in our calibration is 0.28 (as $\alpha = 0.72$), then the negative effect prevails and the subsidy makes more difficult to break a risky SCE. Nevertheless, at the limit of $\text{sub} = 0.28$ (and higher values) the banks are induced to post at (inside) the adoption frontier, for no matter which belief since the central bank is absorbing the whole risk. This way the authority can induce an experiment using the market.

6.4 *Credit-Easing* as a public good

The implementation of the policy does not require that the CB is more optimistic than private banks. The role of the policy maker is to act in the interest of the collectivity, banks *and* firms, investing a fraction of current available resources to provide a public good: the experimentation of new markets. In other words, the aim of the policy is to *produce* a piece of information which has a extremely high social value but cannot be internalized by current private transactions.

In an intertemporal perspective the objective of the CB is to maximize the overall welfare

$$W = \sum_{t=0}^{\infty} \delta^t W_t$$

with $W_t = J_t^* - T_t + V_t(R^*) - c$, where δ is a discount factor weighting the welfare of future generations, J_t^* and $V_t(R^*) - c$ denote the expected profit of respectively firms and banks at the equilibrium. T is a tax introduced to cover the expected cost of the subsidy experiment that will hit firms, since in equilibrium banks make zero profits.

Following our previous example suppose that the CB believe, as banks, that the probability of the low-fix-cost economy (the one of figure 1) is $\zeta < 0.213$. The resources for this incentive can be gathered through a taxation on actual firms T_t or through seigniorage. Imagine the experiment is conducted at time t . The expected social benefit at time t is given by $W_t + \Delta W_t(\text{sub}, \zeta)$ where $\Delta W_t(\text{sub}, \zeta)$ can be negative since the intervention distorts the equilibrium allocation. In particular, given the state of belief the optimal deviation R_d for a single bank is such that either

$$E^\beta [\mu'(R_d)] = 0 \quad \text{provided} \quad R_d < \bar{R}$$

or $R_d = \bar{R}$. Assuming that it takes one period for banks to update their beliefs, the expected welfare gain or loss relative to the period of experimentation is

$$\Delta W_t(\text{sub}, \zeta) = \frac{1}{2} \left(\zeta (y - k_l - R_d)^2 + (1 - \zeta) \alpha^2 (y - R_d)^2 \right) - \alpha \text{sub} (1 - \zeta) R_d$$

where $T = -\alpha s(1 - \zeta) R_d$.

Nevertheless although $\Delta W_t(\text{sub}, \zeta)$ can be negative, the authority can well decide that the experiment is socially valuable as it can provide with useful information for the future generation of traders. In fact, the change in welfare of all subsequent periods

$$\Delta W_{t+h}(\text{sub}, \zeta) = \zeta (J^s - J^r)$$

is a well defined positive value being equal to the expected increase in firms' REE profits (banks will earn zero due to free entry). Whether or not the gain in subsequent periods is a valuable social experimentation depends on the weights that future generations have in the preference of the policy maker. The social experiment is worth if and only if

$$\Delta W_t(\text{sub}, \zeta) < \frac{\zeta \delta}{1 - \delta} (J^s - J^r)$$

that is the current expected cost is smaller than the discounted stream of higher welfare in the future. In particular, provided $\Delta W_t(\text{sub}, \zeta^*)$ is bounded, there always exist a δ large enough such that the policy maker will make a social experiment as long as $\zeta \neq 0$. A more conservative (progressist) policy maker, that is one with a lower (higher) δ , will be less (more) incline to social experimentation. In this sense the policy maker must act at the opposite of what predicted

by a robust control criterion. With a $\delta \rightarrow 1$ in fact the policy maker has incentive to experiment every good state of the world which receives a strictly positive probability.

Two remarks to conclude. First, notice that the intertemporal perspective is not strictly necessary as there could be also cases in which $\Delta W_t(\text{sub}, \zeta)$ is a positive term. In such a case experimentation is worth for the current generation too. Finally, notice that by continuity, the argument works also with a policy maker that is more pessimist than the private sector.

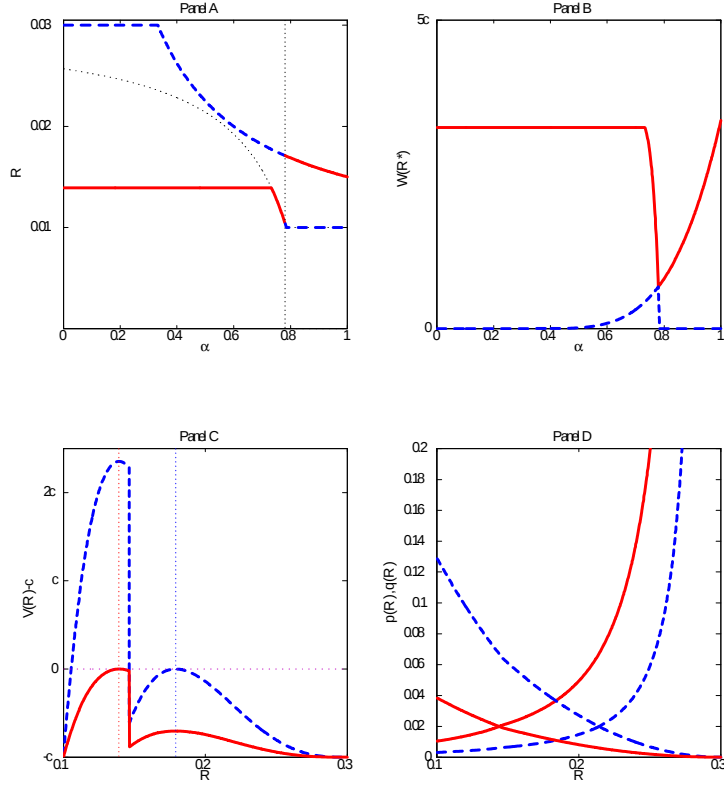


Figure 1: **Baseline parameterization.** Panel A: REE (solid) and SCE (dashed) in the (R, α) -space. Panel B: aggregate investment at the REE (solid) and the SCE (dashed). Panel C: payoff of a bank's individual deviation from the REE (solid) and the SCE (dashed) for $\alpha = 0.7$. Panel D: matching probabilities ($q(R)$ decreasing in R , $p(R)$ increasing in R) relative to a bank's individual deviation from the REE (solid) and SCE (dashed) for $\alpha = 0.72$. Other parameters taken fix across panels $\gamma = 0.3$, $R_{CB} = 0.01$, $k = 0.0042$, $A = 0.02$, $\gamma = 0.5$, $c = 10^{-6}$.

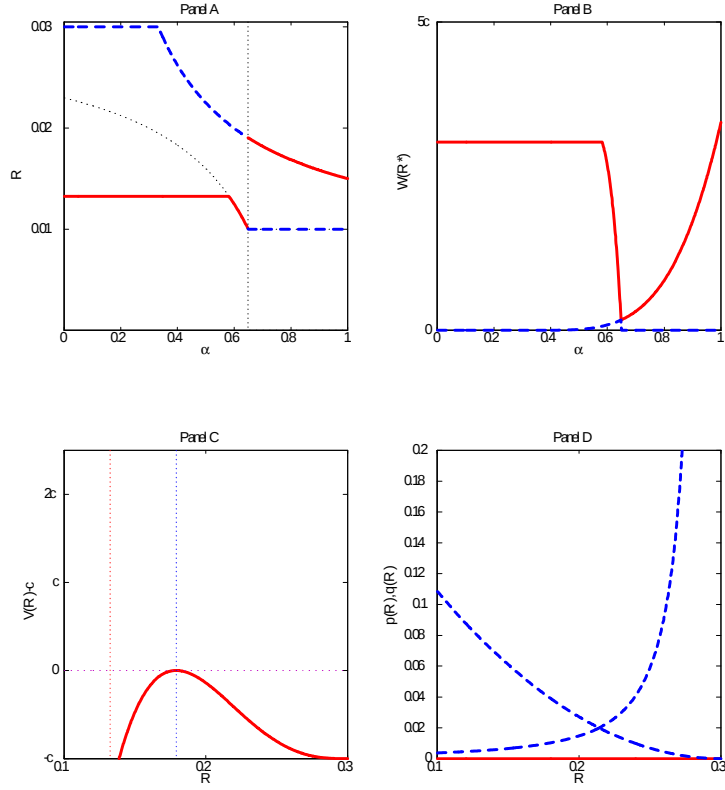


Figure 2: **Higher fix cost for safe projects** ($k = 0.007$). Panel A: REE (solid) and SCE (dashed) in the (R, α) -space. Panel B: aggregate investment at the REE (solid) and the SCE (dashed). Panel C: payoff of a bank's individual deviation from the REE (solid) and the SCE (dashed) for $\alpha = 0.7$. Panel D: matching probabilities ($q(R)$ increasing in R , $p(R)$ decreasing in R) relative to a bank's individual deviation from the REE (solid) and SCE (dashed) for $\alpha = 0.72$. Other parameters taken fix across panels $y = 0.3$, $R_{CB} = 0.01$, $A = 0.02$, $\gamma = 0.5$, $c = 10^{-6}$.

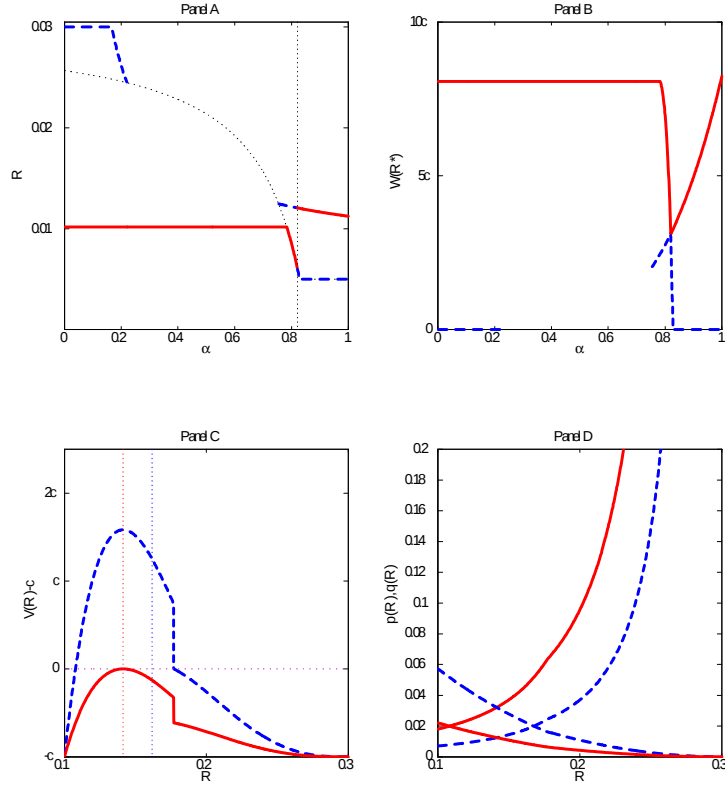


Figure 3: **Reduction in the cost of money** ($R_{CB} = 0.005$). Panel A: REE (solid) and SCE (dashed) in the (R, α) -space. Panel B: aggregate investment at the REE (solid) and the SCE (dashed). Panel C: payoff of a bank's individual deviation from the REE (solid) and the SCE (dashed) for $\alpha = 0.7$. Panel D: matching probabilities ($q(R)$ decreasing in R , $p(R)$ increasing in R) relative to a bank's individual deviation from the REE (solid) and SCE (dashed) for $\alpha = 0.72$. Other parameters taken fix across panels $y = 0.3$, $k = 0.0042$, $A = 0.02$, $\gamma = 0.5$, $c = 10^{-6}$.

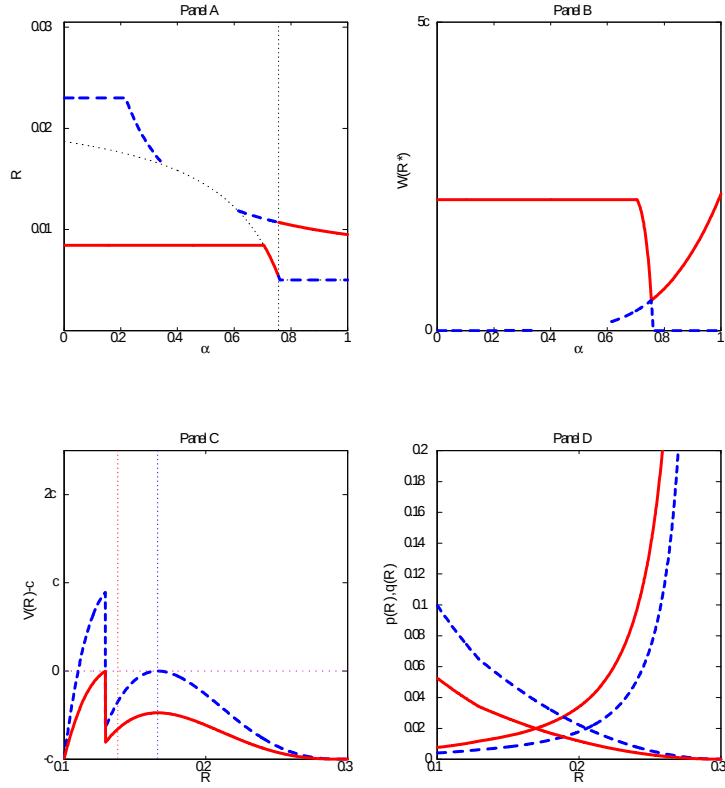


Figure 4: **Decrease in project returns** ($y = 0.02$). Panel A: REE (solid) and SCE (dashed) in the (R, α) -space. Panel B: aggregate investment at the REE (solid) and the SCE (dashed). Panel C: payoff of a bank's individual deviation from the REE (solid) and the SCE (dashed) for $\alpha = 0.7$. Panel D: matching probabilities ($q(R)$ decreasing in R , $p(R)$ increasing in R) relative to a bank's individual deviation from the REE (solid) and SCE (dashed) for $\alpha = 0.72$. Other parameters taken fix across panels $R_{CB} = 0.005$, $k = 0.0042$, $A = 0.02$, $\gamma = 0.5$, $c = 10^{-6}$.

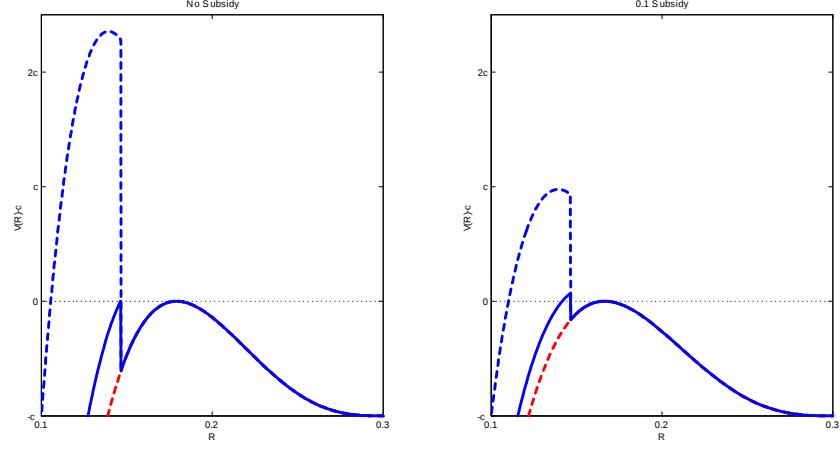


Figure 5: **Subsidy.** The blu line represents the expected payoff function when banks believe with probability 0.213 that $k = 0.0042$ and with probability 0.787 that $k = 0.007$. In the right panel the subsidy is set to zero $s = 0$ whereas in the left panel the subsidy accounts for 10% of banks margins. Other parameters taken fix across panels $\alpha = 0.72$, $y = 0.3$, $R_{CB} = 0.01$, $k = 0.0042$, $A = 0.02$, $\gamma = 0.5$, $c = 10^{-6}$.

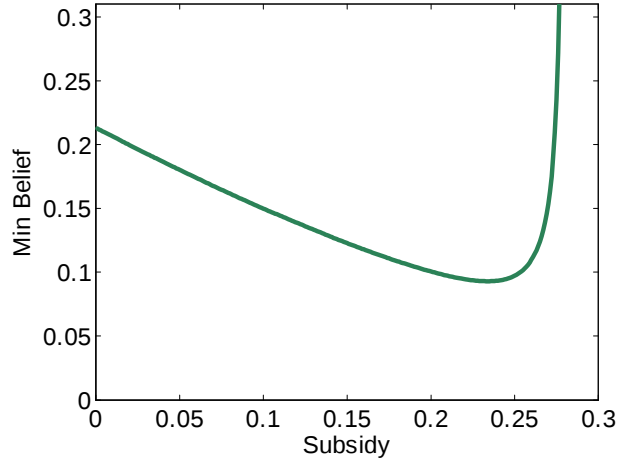


Figure 6: **Minimal pessimism for subsidy effectiveness.** The curve assigns for each value of the subsidy sub the minimal probability that banks put on k_l that makes them to sustain a risky SCE. The parameters is the same as in figure 5.

Appendix

Proof of proposition 1. Consider a single bank that evaluates an individual deviation R from an equilibrium contract R^* . A bank knows that firms act with perfect information, that is, banks know that firms enter the submarket that yield the highest expected profit.

Therefore, a bank anticipates that for any possible state of the world the tightness linked to the submarket R is such that

$$\bar{J} = p(\theta(R^*, \bar{J})) \Pi(R^*) = p(\theta(R, \bar{J})) \Pi(R, f) \quad (14)$$

where $\bar{J}$ is a constant ex-ante level of firm profits that a single bank cannot affect. For the sake of notational simplicity we denote $f(R, \omega)$ as just f (and $f(R^*, \omega)$ as f^*) when there is no ambiguity. Given the relation above the out-of-equilibrium function $p(\theta(R, \bar{J}))$ is obtained as

$$p(\theta(R, \bar{J})) = p(\theta(R^*, \bar{J})) \frac{\Pi(f)}{\Pi(f^*)}, \quad (15)$$

for any state of the world. Notice that if R decreases, $\Pi(f)$ increases and so $p(\theta(R, \bar{J}))$ decreases, that is the probability for a firm to match decreases in R .

Nevertheless a bank does not know with certainty which state of the world has realized. According to (1), the expected value of a credit line is

$$V(R, \theta(R)) = E^\beta [q(\theta(R)) Y(R, f)]$$

for the individual bank posting the contract R whereas all the others posting R^* . Given banks' knowledge of the tightness function (15), we can write

$$V(R, \theta(R, \bar{J})) = E^\beta \left[\theta(R, \bar{J}) p(\theta(R^*, \bar{J})) \frac{\Pi(f^*)}{\Pi(f)} Y(R, f) \right],$$

where we also used the fact $q(\theta) = \theta p(\theta)$. On the other hand, the bank who conforms to the equilibrium prescriptions expects

$$V(R^*, \theta(R^*, \bar{J})) = \theta(R^*, \bar{J}) p(\theta(R^*, \bar{J})) Y(R^*, f^*),$$

where notice we assume that there is no uncertainty at the equilibrium, that is $E^\beta [V(R^*, \theta(R^*, \bar{J}))] = V(R^*, \theta(R^*, \bar{J}))$.

The latter is a weakly dominant strategy if and only if

$$\begin{aligned} \theta(R^*, \bar{J}) p(\theta(R^*, \bar{J})) Y(R^*, f^*) &\geq \\ &\geq E^\beta \left[\theta(R, \bar{J}) p(\theta(R^*, \bar{J})) \frac{\Pi(f^*)}{\Pi(f)} Y(R, f^*) \right] \end{aligned}$$

or

$$E^\beta \left[\frac{\theta(R, \bar{J}) \Pi(f^*)}{\theta(R^*, \bar{J}) \Pi(f)} \frac{Y(R, f)}{Y(R^*, f^*)} \right] \leq 1 \quad (16)$$

To derive $\theta(R, J^*)$ we can use the definition of $\bar{J}$ in (14), so that

$$\bar{J} = A\theta(R, \bar{J})^{\gamma-1} \Pi(f)$$

implies

$$\theta(R, \bar{J}) = C\Pi(f)^{\frac{1}{1-\gamma}} \quad (17)$$

where $C \equiv (A/\bar{J})^{\frac{1}{1-\gamma}}$ is a constant for any single deviation R . Hence, we obtain

$$\frac{\theta(R, \bar{J})}{\theta(R^*, \bar{J})} = \left(\frac{\Pi(f^*)}{\Pi(f)} \right)^{-\frac{1}{1-\gamma}},$$

which is true at any state of the world. Hence (16) becomes

$$E^\beta \left[\left(\frac{\Pi(f^*)}{\Pi(f)} \right)^{-\frac{\gamma}{1-\gamma}} \frac{Y(R, f)}{Y(R^*, f^*)} \right] \leq 1$$

or

$$E^\beta \left[\Pi(f)^{\frac{\gamma}{1-\gamma}} Y(R, f) \right] \leq \Pi(f^*)^{\frac{\gamma}{1-\gamma}} Y(R^*, f^*)$$

which is (18) in the main text. Notice that the result does not depend on which $\bar{J}$ is believed by banks.

If instead we consider two arbitrary single deviations R_1 and R_2 from R^* we get that the latter weakly dominates the former when

$$\begin{aligned} E^\beta \left[\theta(R_2, \bar{J}) p(\theta(R^*, \bar{J})) \frac{\Pi(f(R^*, \omega))}{\Pi(f(R_2, \omega))} Y(R_2, f(R_2, \omega)) \right] &\geq \\ &\geq E^\beta \left[\theta(R_1, \bar{J}) p(\theta(R^*, \bar{J})) \frac{\Pi(f(R^*, \omega))}{\Pi(f(R_1, \omega))} Y(R_1, f(R_1, \omega)) \right] \end{aligned}$$

or

$$E^\beta \left[\theta(R_2, \bar{J}) \frac{Y(R_2, f(R_2, \omega))}{\Pi(f(R_2, \omega))} \right] \geq E^\beta \left[\theta(R_1, \bar{J}) \frac{Y(R_1, f(R_1, \omega))}{\Pi(f(R_1, \omega))} \right]$$

which becomes

$$E^\beta \left[C\Pi(R_2)^{\frac{1}{1-\gamma}} \frac{Y(R_2, f(R_2, \omega))}{\Pi(f(R_2, \omega))} \right] \geq E^\beta \left[C\Pi(R_1)^{\frac{1}{1-\gamma}} \frac{Y(R_1, f(R_1, \omega))}{\Pi(f(R_1, \omega))} \right]$$

after using (17). Finally we get

$$E^\beta [\mu(R_2)] \geq E^\beta [\mu(R_1)], \quad (18)$$

where

$$\mu(R) \equiv \Pi(f(R, \omega))^{\frac{\gamma}{1-\gamma}} Y(R, f(R, \omega)),$$

that is our operative criterion.

Therefore the best *interior* single deviation is the one that maximizes

$$E^\beta [\mu(R)]$$

whose first order condition is

$$E^\beta \left[\frac{\gamma}{1-\gamma} \Pi(f)^{\frac{\gamma}{1-\gamma}} \frac{\Pi'(f)}{\Pi(f)} Y(R, f) + \Pi(f)^{\frac{\gamma}{1-\gamma}} Y'(R, f) \right] = 0$$

at any state for the world. Hence

$$E^\beta \left[\left(\Pi(f)^{\frac{\gamma}{1-\gamma}} Y(R, f) \right) \left(\gamma \frac{\Pi'(f)}{\Pi(f)} + (1-\gamma) \frac{Y'(R, f)}{Y(R, f)} \right) \right] = 0$$

whose deterministic solution is

$$\gamma \frac{\Pi'(f)}{\Pi(f)} + (1-\gamma) \frac{Y'(R, f)}{Y(R, f)} = 0,$$

which reconciles with the solution under certainty we found using standard techniques.

Proof of proposition 2 An interior SCE R^* is a solution to

$$\begin{aligned} & \max_R q(R) Y(R, f) \\ & \text{s.t. } \bar{J} = p(R) \Pi(f) \end{aligned}$$

where $E^\beta [f(R, \omega)] = f(R, \omega)$ for any $R \in \mathfrak{S}(R^*)$.¹⁰ Therefore the following two FOCs

$$\begin{aligned} q'(R) Y(R, f) + q(R) Y'(R, f) &= 0 \\ p'(R) \Pi(f) + p(R) \Pi'(f) &= 0 \end{aligned}$$

have to be satisfied at the equilibrium R^* . Since $p(R) = A\theta(R)^{\gamma-1}$ and $q(R) = A\theta(R)^\gamma$ we can rewrite the latter as

$$(\gamma-1) \Pi(f) + \theta(R) \Pi(f)' = 0$$

and the former as

$$\gamma Y(R, f) + \theta(R) Y'(R, f) = 0$$

and finally use both to get rid of θ obtaining the relation

$$\gamma Y(R, f) + (1-\gamma) \frac{\Pi(f)}{\Pi'(f)} Y'(R, f) = 0,$$

or

$$\gamma \frac{\Pi'(f)}{\Pi(f)} + (1-\gamma) \frac{Y'(R, f)}{Y(R, f)} = 0$$

¹⁰To save on notation from here onwards we will omit dependence of functions from R .

After simple manipulations we have that, at the equilibrium,

$$1 - \gamma = \frac{Y(R^*, f^*)}{Y(R, f^*) - \frac{Y'(R^*, f^*)}{\Pi'(f^*)} \Pi(f^*)}$$

which demonstrates that the Hosios condition is met.

Proof of proposition 3 Here we deal with REE that is the case where banks' beliefs are correct also out of equilibrium.

FIRST STEP. The first step of the proof is to establish the set of best R contracts conditional to firms having incentives to adopt one specific type of technology. Call these locus best restricted contracts.

To obtain the best restricted contracts when *no participation constraints bind* it is enough to plug the explicit form of I and π into (10). In the case of a risky technology we have

$$\hat{R}_r = \frac{1 - \gamma}{2} r + \frac{1 + \gamma}{2\alpha} R_{CB}; \quad (19)$$

whereas in case of a safe technology

$$\hat{R}_s = \frac{1 - \gamma}{2} (r - k) + \frac{1 + \gamma}{2} R_{CB}. \quad (20)$$

(19) and (20) are the best interior contracts that a bank would provide if it was restricted to post a credit line respectively outside and inside the adoption frontier irrespective of any $\bar{J}$ value. This is intuitive since, as already noted, point ii) of definition of a SCE does not link $\bar{J}$ to the actual level granted in equilibrium J^* .

Nevertheless, best interior contracts could be unfeasable due to participation constraint. To find the best restricted contracts when *at least one participation constraint is binding* for at least one type of agent, we make use of a simple convexity argument. In the abstract case where only one technology is available, the maximization problem is nicely convex and has a single absolute maximum. Therefore, ceteris paribus, the closest the offered contract to $\hat{R}_r$ (resp. $\hat{R}_s$), the higher the expected profits $V(R)$ of a bank. This implies that whenever for a level of α it is $\hat{R}_r > r$ then the best risky restricted-contract is r . Moreover, whenever for a level of α it is $\hat{R}_r < \bar{R}$ then the best risky restricted-contract is $\bar{R}$. Finally whenever for a level of α it is $\hat{R}_s > \bar{R}$ then the best safe-restricted contract is $\bar{R}$.

SECOND STEP. For best restricted-contracts to be REE, it is necessary (but still not sufficient) that in a neighborhood of the equilibrium other contracts does not yield a strictly larger profit. For (19) and (20) this condition is always satisfied as both emerge as solution of a well-defined maximization problem in R . So, both (19) and (20) are candidate for the next step.

Also in the case where the best risky-restricted contract is r , then a bank cannot do better locally as any local deviation from r will remain strictly outside the adoption frontier. Therefore also r is a candidate for the next step.

When instead best restricted contracts lie along the adoption frontier $\bar{R}$ the argument is more involved. First consider the case of a best risky-restricted contract along $\bar{R}$. We need to use the criterion (18) to assess whether or not a single bank has incentive to deviate posting a contract inside the adoption frontier (i.e. posting a $\bar{R} - \varepsilon$ with $\varepsilon > 0$ infinitesimal) when all others post a contract outside the adoption frontier (i.e. posting a $\bar{R} + \varepsilon$ with $\varepsilon > 0$ infinitesimal): in the cases where this is true there do not exist REE risky contracts along $\bar{R}$, otherwise we proceed to the third step.

A deviation from the risky equilibrium (all banks post $\bar{R} + \varepsilon$) into safe territory (the deviant posts $\bar{R} - \varepsilon$) is worth if and only if,

$$\frac{\pi(s, \bar{R}, R_{CB})I(s, \bar{R})}{\pi(r, \bar{R}, R_{CB})I(r, \bar{R})} > \left(\frac{\pi(s, \bar{R}, R_{CB})I(r, \bar{R})}{\pi(r, \bar{R}, R_{CB})I(s, \bar{R})} \right)^\gamma \quad (21)$$

which is a rearrangement of (18). Hence, we have

$$\frac{(\bar{R} - R_{CB})(r - k - \bar{R})}{(\alpha\bar{R} - R_{CB})\alpha(r - \bar{R})} > \left(\frac{(\bar{R} - R_{CB})\alpha(r - \bar{R})}{(\alpha\bar{R} - R_{CB})(r - k - \bar{R})} \right)^\gamma \quad (22)$$

and after substituting for $\bar{R} = r - \frac{k}{1-\alpha} > R_{CB}$,

$$\frac{\bar{R} - R_{CB}}{\alpha\bar{R} - R_{CB}} > \left(\frac{\bar{R} - R_{CB}}{\alpha\bar{R} - R_{CB}} \right)^\gamma.$$

For $\alpha\bar{R} - R_{CB} < 0$ banks do not have incentive to enter in the market for risky credit so that a risky REE does not exist in this case. When instead $\alpha\bar{R} - R_{CB} \geq 0$ then the inequality always holds for whatever $\gamma < 1$. Therefore we can conclude that a risky REE does not exist along the $\bar{R}$ frontier.

Let us turn attention to the case of a best safe-restricted contract along $\bar{R}$. We apply the criterion (21) along $\bar{R} = r - \frac{k}{1-\alpha} > R_{CB}$ to assess whether a single bank has incentive to post a contract outside the adoption frontier (i.e. posting a $\bar{R} + \varepsilon$ with $\varepsilon > 0$ infinitesimal) when all others post a contract inside the adoption frontier (i.e. posting a $\bar{R} - \varepsilon$ with $\varepsilon > 0$ infinitesimal); in the cases where this is true there do not exist REE safe contracts along $\bar{R}$, otherwise we proceed to the third step. Applying (21) we have a relation which holds as a strict inequality whenever (22) is false. Therefore we can conclude that there could exist safe REE in the case $\hat{R}^s > \bar{R}$ along the $\bar{R}$ frontier.

THIRD STEP. This is the final step. Once selected the best restricted contracts (step 1) that are local maxima (step 2), we need to establish whether they are global maxima, that is, if they are REE. Now we apply (18) to the different cases, distinguishing between interior and corner contracts.

The relevant equation for an interior best restricted contract to be a REE when both type of interior best restricted contracts ($R^s = \hat{R}^s$ and $R^r = \hat{R}^r$) are feasible is

$$(1 - \gamma)(1 + \gamma)^{\frac{1+\gamma}{1-\gamma}} \left(\frac{r - k - R_{CB}}{2} \right)^{\frac{2}{1-\gamma}} > (1 - \gamma)(1 + \gamma)^{\frac{1+\gamma}{1-\gamma}} \left(\frac{r\alpha - R_{CB}}{2} \right)^{\frac{2}{1-\gamma}}$$

which holds is and only if $r - k/(1 - \alpha) > 0$. Since a necessary condition for the existence of an interior best safe-restricted contract is that the adoption region is not empty, $\bar{R} > R_{CB}$, this condition always holds: whenever an interior best safe-restricted contract exists then it is a REE. When instead we confront an interior best safe-restricted contract with a corner best risky-restricted contract posted at r , then the right-hand side of the disequality takes value zero so that the disequality is trivially satisfied. *We conclude that whenever an interior best safe-restricted contract exists then it is a REE.*

When instead the best safe-restricted contract is posted at R_{CB} , that is when

$$\alpha > \bar{\alpha} \equiv \frac{r - k - R_{CB}}{r - R_{CB}},$$

whereas the best risky-restricted contract is an interior ($r\alpha - R_{CB} > 0$), we have

$$0 > (1 - \gamma)(1 + \gamma)^{\frac{1+\gamma}{1-\gamma}} \left(\frac{r\alpha - R_{CB}}{2} \right)^{\frac{2}{1-\gamma}}, \quad (23)$$

which is always true. This implies that *whenever an interior best risky-restricted contract co-exists with a corner best safe-restricted contract being posted at R_{CB} , the former is always the unique REE.*

When instead the best safe-restricted contract is posted at $\bar{R} \neq R_{CB}$, whereas the best safe-restricted contract is an interior, we have

$$\left(r - \frac{k}{1 - \alpha} - R_{CB} \right) \left(\frac{\alpha k}{1 - \alpha} \right)^{\frac{1+\gamma}{1-\gamma}} > (1 - \gamma)(1 + \gamma)^{\frac{1+\gamma}{1-\gamma}} \left(\frac{r\alpha - R_{CB}}{2} \right)^{\frac{2}{1-\gamma}}. \quad (24)$$

The right-hand side is always monotonically increasing in α . The left-hand side instead is always monotonically decreasing in α in the relevant case

$$\alpha > \underline{\alpha} \equiv \frac{r - \hat{R}^s - k}{r - \hat{R}^s} = \frac{(\gamma + 1)(r - k - M)}{2k + (\gamma + 1)(r - k - M)},$$

for which the interior best safe-restricted contract is on the adoption frontier ($I(\hat{R}_s) = I(\hat{R}_s)$), given that

$$\frac{\partial \left(\left(r - \frac{k}{1 - \alpha} - R_{CB} \right) \left(\frac{\alpha k}{1 - \alpha} \right)^{\frac{1+\gamma}{1-\gamma}} \right)}{\partial \alpha} = \frac{(1 - \alpha)(1 + \gamma)(r - k - R_{CB}) - 2k\alpha}{\alpha(1 - \alpha)^2(1 - \gamma) \left(\frac{\alpha k}{1 - \alpha} \right)^{\frac{\gamma+1}{\gamma-1}}}.$$

Hence, we can conclude that

$$\left(r - \frac{k}{1 - \alpha} - R_{CB} \right) \left(\frac{\alpha k}{1 - \alpha} \right)^{\frac{1+\gamma}{1-\gamma}} = (1 - \gamma)(1 + \gamma)^{\frac{1+\gamma}{1-\gamma}} \left(\frac{r\alpha - R_{CB}}{2} \right)^{\frac{2}{1-\gamma}}$$

defines a threshold $\hat{\alpha}$, such that for $\alpha < \hat{\alpha}$ the corner best safe-restricted contract is the unique REE, whereas for $\alpha > \hat{\alpha}$ the interior best risky-restricted contract

is the unique a REE. The zero measure case $\alpha = \hat{\alpha}$ is the only one where two non-degenerate REE exist. In particular, notice that since

$$\frac{\partial \left(\left(r - \frac{k}{1-\alpha} - R_{CB} \right) \left(\frac{\alpha k}{1-\alpha} \right)^{\frac{1+\gamma}{1-\gamma}} \right)}{\partial k} = \frac{(1-\alpha)(1+\gamma)(r - R_{CB}) - 2k}{k(1-\alpha)(1-\gamma) \left(\frac{\alpha k}{1-\alpha} \right)^{\frac{\gamma+1}{\gamma-1}}} < 0$$

is true whenever $(1-\alpha)(1+\gamma)(r - k - R_{CB}) - 2k\alpha < 0$. This implies that $\hat{\alpha}$ has to be decreasing in k .

When instead the best safe-restricted contract is posted at $\bar{R} \neq R_{CB}$, whereas the best safe-restricted contract is posted at r ($r\alpha - R_{CB} < 0$), we have

$$\left(r - \frac{k}{1-\alpha} - R_{CB} \right) \left(\frac{\alpha k}{1-\alpha} \right)^{\frac{1+\gamma}{1-\gamma}} > 0, \quad (25)$$

so that $\hat{\alpha} = \bar{\alpha}$. This implies that *whenever a corner best risky-restricted contract co-exists with a corner best safe-restricted contract being posted at $\bar{R} > R_{CB}$, the latter is always the unique REE.*

Finally *whenever a corner best risky-restricted contract being posted at r (which requires $R_{CB}/r \leq \alpha$) co-exists with a corner best safe-restricted contract being posted at R_{CB} (which requires $r - k/(1-\alpha) \leq R_{CB}$), the two arise as two degenerate REE.*

References

- [1] Bebchuk L.A. and I. Goldstein (2011). Self-fulfilling Credit Market Freezes. *The Review of Financial Studies*, vol. 24, n.11, pp. 3519-3555.
- [2] Chari V.V. , A. Shourideh and A. Zetlin-Jones (2010). Adverse Selection, Reputation and Sudden Collapses in Secondary Loan Markets. NBER Working Papers 16080.
- [3] Fudenberg D. and D.K. Levine (1993). Self-Confirming Equilibrium. *Econometrica*, vol. 61, n.3, pp.523-545.
- [4] Gertler, M. and P. Karadi (2011). A model of unconventional monetary policy. *Journal of Monetary Economics*, Elsevier, vol. 58(1), pp.17-34.
- [5] Moen E. R. (1997). Competitive Search Equilibrium. *Journal of Political Economy*, vol. 105, n.2, pp. 385-411.
- [6] Morris, S. and H.S. Shin. (2004). Coordination Risk and the Price of Debt. *European Economic Review*, vol. 48, pp. 133-153.
- [7] Sargent T. J. (2001). *The Conquest of American Inflation*. Princeton University Press.
- [8] Shi S. (2006). Search Theory; Current Perspectives. *Palgrave Dictionary of Economics*.

TBC