



UNIVERSITÀ DI PISA

DIPARTIMENTO DI MATEMATICA
CORSO DI LAUREA MAGISTRALE IN MATEMATICA

MASTER OF SCIENCE FINAL DISSERTATION

**Mean field games:
prospects and intuitions**

CANDIDATE:

Davide Veronese

ADVISOR:

Professor Franco Flandoli

DISCUSSANT:

Professor Maurizio Pratelli

ACADEMIC YEAR 2013/2014

to Giuri and Adriana

Contents

Abstract	3
1 Introduction	5
1.1 A difficult problem	5
1.2 The mean field approach: an outline	6
2 “What time does the meeting start?”	11
2.1 The mean field resolution	12
2.2 Numerical simulation of the toy model	17
2.2.1 Cost function parameters	17
2.2.2 Agent population choice and quorum	19
2.2.3 Some results	20
2.3 A digression on the concept of “continuum of agents”	23
2.3.1 Extending the result to non identically distributed families	24
3 The full MFG scenario	27
3.1 The first step: the agent optimal control problem	29
3.2 The second step: aggregation	32
3.3 The mean field games PDE formulation	35
3.3.1 Existence proof	36
3.3.2 Uniqueness proof	42
3.4 MFG solution as an ε -Nash Equilibrium	44
4 A model of imitative behaviours in a population	47
4.1 Model set-up	47
4.2 PDE formulation	48
4.3 A stationary solution	50
5 An individual-based human capital growth model	53
5.1 Model set-up	54
5.2 PDE formulation	55
5.3 Comparison with common noise models	57
6 Prospects and conclusion	59
Appendix	61

Abstract

Mean field games is a branch of game theory that deals with large symmetrical games.

Born in 2006, it is now a lively research topic that poses tough mathematical challenges and already features a broad range of applications.

Among the challenges, a system of coupled PDE plays the major role, while current mean field games models deal with crowds behaviour, traffic flows, macroeconomics dynamics and cancer evolution.

Our work has a dual aim: to highlight the core intuitions of this approach and to present some of its prospects as an applied tool. To do this, we merge the Mean Field Games analytic backbone by P. Cardaliaguet with some of the applications by O. Guéant, J-M. Lasry and P-L. Lions.

The first part of our work introduces the main intuitions behind the mean field game solution, such as the limit game definition, the forward-backward reasoning and the fixed point nature of the mean field. We then solve a single-shot game via the mean field approach and present a numerical simulation of the results.

The second part shows how the same steps can be employed in the continuous time case to obtain the coupled PDE formulation. In this general problem setting, we prove an important existence and uniqueness theorem for the mean field solution and characterize this latter as an ε -Nash equilibrium.

In the last part we focus on the prospects for this recent theory, presenting some economics related applications. After a brief model of imitative preferences, we discuss an heterogeneous agents, human capital based growth model, highlighting how easily real-like scenarios can give birth to new unconventional PDE systems.

Chapter 1

Introduction

1.1 A difficult problem

Let us begin by considering the following framework:

N rational and homogeneous agents play in a common environment, during a time interval $[0, T]$. We allow them to have a continuous state, described by $X_t^i \in \mathbb{R}^d$, and let them freely adopt a control of the same dimension $\alpha_t^i \in \mathbb{R}^d$, that acts on their own state X_t^i according to the law ¹:

$$dX_t^i = \alpha_t^i dt + \sqrt{2} dB_t^i \quad (1.1)$$

Each agent plays his strategy in order to minimize his expected cost, that we suppose to have the same form for all the agents:

$$\begin{aligned} \mathcal{J}_i^N(\alpha^i, (\alpha^j)_{\substack{j \leq N \\ j \neq i}}) = \\ = \mathbb{E} \left[\int_0^T \left[\frac{1}{2} |\alpha_s^i|^2 + F \left(X_s^i, \frac{1}{N-1} \sum_{j \neq i} \delta_{X_s^j} \right) \right] ds + G \left(X_T^i, \frac{1}{N-1} \sum_{j \neq i} \delta_{X_T^j} \right) \right] \end{aligned} \quad (1.2)$$

where:

- $|\alpha_s^i|^2$ is a simple choice for the cost density due to the employed control;
- $F(x, m)$ describes the cost density due to the interaction with other agents which occurs during game evolution;

¹We assume B_t^i to be a family of independent Brownian motions. If we allow the initial state X_0^i to be random and identically distributed with law m_0 , we also require that this law is independent from all the B_t^i .

- $G(x, m)$ describes the cost due to the final game status.

How can we deal with this game-type interaction?

It can be observed that the cost $\mathcal{J}$ depends on the other players' decisions in a particular way: an exhaustive summary for this dependence is given by the *mean distribution* of the states of the other players.

As a consequence, this remarkable property holds:

Property 1 (Symmetry). The problem setting is invariant through permutations among the agents, i.e. two agents who are in the same state and face the same external conditions undertake the same strategy.

Note that we did not include the random states $(X_t^1, \dots, X_t^N)$ as input variables for the cost functional $\mathcal{J}$, while we did put the corresponding controls. This is consistent with equation (1.1), since the expected value of the state X_t^i depends only on the adopted control α^i .

For this setting we may resort to classical game theory, but as N gets large the exact solution would become both analytically and numerically unpleasant.²

In fact, finite game theory considers a stochastic final distribution for the agents' states, which is something that is really hard to calculate and deal with.

A very elegant, innovative and subtle way to tackle this scenario was introduced by Jean-Michel Lasry and Pierre-Louis Lions in year 2006 [17, 18] under the evocative name of "Mean field games".

1.2 The mean field approach: an outline

In the framework described above, the mean field approach claims to provide an approximate solution of the problem for sufficiently large N by defining and solving a theoretical *infinite players* version of the same game.

This limit formulation relies on the introduction of a deterministic probability measure $m(\cdot, t)$ depending on time t over the state space R^d such that:

$$\begin{aligned} \lim_{N \rightarrow \infty} \mathcal{J}_i^N(\alpha^i, (\alpha^j)_{\substack{j \leq N \\ j \neq i}}) &= \mathcal{J}(\alpha^i, m) \\ &= \mathbb{E} \left[\int_0^T \left[\frac{1}{2} |\alpha_s^i|^2 + F(X_s^i, m(\cdot, s)) \right] ds + G(X_T^i, m(\cdot, T)) \right] \end{aligned} \quad (1.3)$$

²Research is still focusing on this topic, since to compute the exact solution has been proved to be a quite complex problem: the state of the art is well resumed in [11].

Remark 2. In this new formulation, given an agent i , its cost functional does not depend any more on the decisions of the other $(N - 1)$ players, as it was for (1.2), since it is evaluated upon his own strategy α_t^i and the newly introduced common factor m .

We could say that taking the limit for $N \rightarrow \infty$ has allowed us to replace $\mathcal{J}_i$ with $\mathcal{J}$, so that m appears in place of $\frac{1}{N-1} \sum_{j \neq i} \delta_{X_s^j}$ as resume of the other players' states.

This has a profound consequence:

Property 3 (i.i.d. states). In the above limit game formulation, due to the independence of the initial conditions X_0^i and the fact that the optimal strategy depends only on m , the agents states described by equation (1.1) constitute an *independent* infinite family of identically distributed random variables.

Let us now present the various steps that lead to build such a promising measure m according to the mean field approach.

Step 0: Before taking the operative steps, we think it is important to somewhat justify the above limit introduction and its result.

When N is small, each player has a certain influence on the others' decisions, since his state significantly affects the others' costs. Yet, if this contribution is averaged out all the other players' states, we expect that for large N each player becomes "small" and can hardly influence the overall game outcome.

Moreover, for finite N the final distribution of the player states (the game outcome) is *stochastic*, since it depends from the single realizations of the independent Brownian noises that appear in equation (1.1).

Given this framework, we see good hopes of simplifying our problem by letting N tend to $+\infty$.

In order to provide an intuitive view, we assume³:

$$\lim_{N \rightarrow \infty} \frac{1}{N-1} \sum_{j \neq i} \delta_{X_s^j} \quad \text{exists for some } i \text{ (implies for every } i \text{ by symmetry)} \quad (1.4)$$

³It is a reasonable assumption, but far from being obvious: X^i depends on the control α^i which in turn depends (for finite N) on all the other $(X^j)_{j \neq i}$, so that the family (δ_{X^i}) needs not to be independent. It turns out to be true, with respect to the distance we are going to define in chapter 3, due to the symmetry of the involved measures. The full proof can be found in [6].

then the most critical input variable to the cost functional (1.2) becomes:

$$\begin{aligned}
\lim_{N \rightarrow \infty} \frac{1}{N-1} \sum_{j \neq i} \delta_{X_s^j} &= \lim_{N \rightarrow \infty} \frac{N-1}{N} \left[\frac{1}{N-1} \sum_{j \neq i} \delta_{X_s^j} \right] \\
&= \lim_{N \rightarrow \infty} \left(\frac{N-1}{N} \left[\frac{1}{N-1} \sum_{j \neq i} \delta_{X_s^j} \right] + \frac{1}{N} \delta_{X_s^i} \right) \\
&= \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{j=1}^N \delta_{X_s^j} \\
&= \lim_{N \rightarrow \infty} m_N(s) := m(s) \quad (\text{independent from the initial considered } i)
\end{aligned}$$

Under reasonable regularity hypotheses on F and G the limit can move into⁴ $\mathcal{J}_i$ definition giving:

$$\begin{aligned}
\lim_{N \rightarrow \infty} \mathcal{J}_i^N(\alpha^i, (\alpha^j)_{j \leq N, j \neq i}) &= \mathbb{E} \left[\int_0^T \frac{1}{2} |\alpha_s^i|^2 + F \left(X_s^i, \lim_{N \rightarrow \infty} \frac{1}{N-1} \sum_{j \neq i} \delta_{X_s^j} \right) ds \right] + \\
&\quad + \mathbb{E} \left[G \left(X_{sT}^i, \lim_{N \rightarrow \infty} \frac{1}{N-1} \sum_{j \neq i} \delta_{X_T^j} \right) \right] \\
&= \mathbb{E} \left[\int_0^T \left[\frac{1}{2} |\alpha_s^i|^2 + F(X_s^i, m(s)) \right] ds + G(X_T^i, m(T)) \right] \\
&= \mathcal{J}(\alpha^i, m)
\end{aligned}$$

Another remarkable change occurs thanks to the infinite number of players.

As already noticed, if every agent undertakes his decisions minimizing the cost functional (1.3), the random variables family $(X^i)_{i \in \mathbb{N}}$ has the *i.i.d.* property.

We can then apply then apply the law of large numbers to the limit probability measure:

$$m(s) = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \delta_{X_s^i} \quad \text{for every time } s$$

thus proving that is indeed *deterministic* (yet still unknown).

We have just reached the first and most important milestone on our path: instead of analysing a very complicated game tree, where the final outcome is stochastic, we are now dealing with a single unknown *deterministic* distribution.

We can now move on to the operative steps.

A good starting point is to consider what happens if we assume that our rational agents initially know m : we highlight that this could not be done with a stochastic measure.

⁴It is a direct consequence of Hewitt-Savage theorem, see theorem A.6.

Step 1: If we assume the deterministic density m to be known, the optimal strategy adopted by every agent can be studied by minimizing the cost over a suitable class $\mathcal{A}$ of admissible controls:

$$\inf_{\alpha \in \mathcal{A}} \mathcal{J}(\alpha, m) \quad \text{if an optimal strategy } \tilde{\alpha} \text{ exists, it verifies:}$$

$$\tilde{\alpha} \quad \text{s.t.} \quad \mathcal{J}(\tilde{\alpha}) \leq \mathcal{J}(\alpha) \quad \forall \alpha \in \mathcal{A}$$

This optimization has to be carried out according to the cost formulation: in our problem, it is a matter of optimal stochastic control. The details will be analysed in section 3.1, where it will be shown that under reasonable hypotheses on F and G an optimal strategy uniquely exists.

The final aim of this step is to define a map that gives the optimal strategy $\tilde{\alpha}$ of any agent observing m :

$$m(\cdot, t) \longmapsto \tilde{\alpha}$$

Step 2: Since the agents states $(X_t^i)_{i \in \mathbb{N}}$ are i.i.d., it follows that the limit measure:

$$m^*(\cdot, t) = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{j=1}^N \delta_{X_s^j}$$

coincides with the law of (for example) X_t^1 on $\mathbb{R}^d$, so that we can denote with X_t the generic agent's state having law $m(t)$.

How this law can be determined depends on the problem settings. In the framework of this introduction, the calculation of the law of X_t requires to solve the SDE (1.1). We will perform this in chapter 3, by recalling the PDE–SDE relationship to get an equation satisfied by the density (denoted with the same letter) $m(x, t)$ on $\mathbb{R}^d$ of the measure $m(\cdot, t)$.

From a general perspective, the result of this step is to build the “aggregating” map:

$$\tilde{\alpha} \longmapsto m^*(\cdot, t)$$

Step 3: If we manage to prove that the composition of the steps above (the map that “updates” our initial guess on m) converges to a fixed point in a suitable space, we can claim that m is well defined and gets the expressive name of *mean field* of the problem.

Step 4: Finally, our *rationality* hypothesis assures that the m predicted by our model, which is completely known to agents, coincides with the real m . This happens since everyone plays the optimal strategy $\tilde{\alpha}$.

What happens if we come back to the real, finite case and equip the N players with the set of strategies $(\alpha^1, \dots, \alpha^N)$ associated with the mean field m , making they play like they would theoretically do in the infinite limit game?

The answer is indeed beautiful: this solution for the finite case can be characterized as an ε -Nash equilibrium for sufficiently large N (see section 3.4).

Before going deeper into the many details of this complex scenario, in the next chapter we deal with an interesting toy model that allows us to acquaint with the subtle innovations of the mean field approach.

Chapter 2

“What time does the meeting start?”

The situation we want to describe is quite a common one.

Largely inspired by [16], we think about a meeting with many participants that has been scheduled for a certain time, say 0: with rare exceptions, it is likely going to start some time after it.

We can imagine that the actual beginning time T (which thus coincides with the delay) is determined by the arrival times of the agents according to a *quorum rule* with required share θ , up to a limiting starting time $T_{\max}$.

This defines a game interaction between agents: nobody likes to wait for the others to arrive, but arriving with an excessive delay is doubtlessly embarrassing and makes it hard to understand what is going on at the meeting.

Let us say that each agent i controls its own target time τ_i while his actual arrival time t_i is subject to uncertainty according to:

$$t_i = \tau_i + \sigma_i \varepsilon_i$$

In our toy model we assume the agents to be a countable infinite and label them with the sequence $(i)_{i \in \mathbb{N}}$, so that $(\varepsilon_i)_{i \in \mathbb{N}}$ denotes a family of independent $\mathcal{N}(0, 1)$ random variables.

We allow for uncertainty to vary among agents. For simplicity¹, we focus on a finite set of h positive values for σ_i :

$$0 < \sigma_i \in \{\sigma_1, \dots, \sigma_h\}$$

¹Can we dare to assume a continuum of agents, each of those has an arbitrary σ_i ? See section 2.3 for a discussion.

This set can be interpreted, for example, as a way to account for the different transport means used by participants to reach the meeting location.

Considering an *infinite* number of agents leads us to assume the existence of these (finite and positive) limit shares:

$$0 < p_j = \lim_{N \rightarrow \infty} \frac{\#\{i \leq N : \sigma_i = \sigma_j\}}{N} \quad \text{for } j = 1, \dots, h$$

so that $\sum_{j=1}^h p_j = 1$.

Finally, we make a naive choice for the single agent cost function and think it as a sum of three components: ($[x]_+$ is the positive part of x)

- The reputation cost due to arriving after the scheduled time (which we have assumed to be 0), given by

$$a [t_i]_+$$

- The personal inconvenience due to arriving after the meeting has begun, given by

$$b [(t_i - T)]_+$$

- The waiting cost, simply given by

$$c [(T - t_i)]_+$$

Put together we have the following convex cost function: ($a, b, c > 0$)

$$\mathcal{J}(\tau_i, T) = \mathbb{E}[a [t_i]_+ + b [(t_i - T)]_+ + c [(T - t_i)]_+] \quad \text{where } t_i = \tau_i + \sigma_i \varepsilon_i$$

At this point, we have all the elements to investigate this infinite-players game interaction.

2.1 The mean field resolution

Let us build and discuss the mean field solution of this analytically simple, yet conceptually rich setting.

Remark 4. In order to better resemble the real circumstances, we allowed for different σ_i and therefore introduced a source of structural *eterogeneity*. In fact, the agent's cost function depends on σ_i and so the property 1 does not hold at the general level in our toy model.

Nevertheless, it can be noticed that any two agents that share the same σ_i and the same

forecast for T choose the same τ_i as optimal target arrival time.

Hence, symmetry still holds within agents having the same uncertainty value and we are soon going to rely on this fact in order to apply the law of large numbers and define the deterministic mean field.

A good candidate to be the mean field of the problem is the limit cumulative distribution² of the agents' real arrival times (F), since it allows each player to evaluate his own cost without the need of further information on the other participants.

It can be defined by:

$$F(z) = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathbb{1}_{\{t_i \leq z\}} \quad (2.1)$$

A finer observation is that the meeting starting time T is an exhaustive summary of F for the cost evaluation of every agent and can be chosen as a simpler mean field.

Let us begin the resolution with a fundamental proposition:

Proposition 1. Given the existence of $(p_1, \dots, p_h)$ as defined before, F is a deterministic distribution on $\mathbb{R}$.

Proof. We can expand according to the different uncertainty values and apply the law of large numbers:

$$\begin{aligned} F(z) &= \lim_{N \rightarrow \infty} \sum_{i=1}^N \frac{1}{N} \mathbb{1}_{\{t_i \leq z\}} = \lim_{N \rightarrow \infty} \sum_{\substack{i=1 \\ \sigma_i = \sigma_1}}^N \frac{1}{N} \mathbb{1}_{\{t_i \leq z\}} + \dots + \sum_{\substack{i=1 \\ \sigma_i = \sigma_h}}^N \frac{1}{N} \mathbb{1}_{\{t_i \leq z\}} \\ &= \lim_{N \rightarrow \infty} \sum_{\substack{i=1 \\ \sigma_i = \sigma_1}}^N \frac{\#\{i : \sigma_i = \sigma_1\}}{N} \frac{\mathbb{1}_{\{t_i \leq z\}}}{\#\{i : \sigma_i = \sigma_1\}} + \dots + \sum_{\substack{i=1 \\ \sigma_i = \sigma_h}}^N \frac{\#\{i : \sigma_i = \sigma_h\}}{N} \frac{\mathbb{1}_{\{t_i \leq z\}}}{\#\{i : \sigma_i = \sigma_h\}} \\ &= \lim_{N \rightarrow \infty} \frac{\#\{i : \sigma_i = \sigma_1\}}{N} \frac{\sum_{i=1}^N \mathbb{1}_{\{t_i \leq z\}}}{\#\{i : \sigma_i = \sigma_1\}} + \dots + \lim_{N \rightarrow \infty} \frac{\#\{i : \sigma_i = \sigma_h\}}{N} \frac{\sum_{i=1}^N \mathbb{1}_{\{t_i \leq z\}}}{\#\{i : \sigma_i = \sigma_h\}} \\ &= p_1 \lim_{N \rightarrow \infty} \frac{\overbrace{\sum_{i=1}^N \mathbb{1}_{\{t_i \leq z\}}}^{i.i.d}}{\#\{i : \sigma_i = \sigma_1\}} + \dots + p_h \lim_{N \rightarrow \infty} \frac{\overbrace{\sum_{i=1}^N \mathbb{1}_{\{t_i \leq z\}}}^{i.i.d}}{\#\{i : \sigma_i = \sigma_h\}} \end{aligned}$$

²Our setting is unidimensional, so we can easily employ the cumulative distribution to carry out the calculations relative to the probability measure m .

Hence we have:

$$\begin{aligned}
F(z) &= p_1 \mathbb{E} [\mathbb{1}_{\{t_i \leq z \wedge \sigma_i = \sigma_1\}}] + \dots + p_h \mathbb{E} [\mathbb{1}_{\{t_i \leq z \wedge \sigma_i = \sigma_h\}}] \\
&= p_1 \mathbb{P} \left\{ \mathcal{N}(0, 1) \leq \frac{z - \tau_i}{\sigma_1} \right\} + \dots + p_h \mathbb{P} \left\{ \mathcal{N}(0, 1) \leq \frac{z - \tau_i}{\sigma_h} \right\} \\
&= p_1 \Phi \left(\frac{z - \tau_i}{\sigma_1} \right) + \dots + p_h \Phi \left(\frac{z - \tau_i}{\sigma_h} \right)
\end{aligned} \tag{2.2}$$

■

It immediately follows that $T(F)$ is also deterministic and the following steps are well defined.

Step 1: Assuming T to be known and deterministic allows each agent to easily minimize his cost:

$$\tilde{\tau}_i = \arg \min_{z \in \mathbb{R}} \mathbb{E}[\mathcal{J}_i(z + \sigma_i \varepsilon_i, T)]$$

This can be achieved by verifying this implicit first order condition, sufficient for the convexity of the cost function:

$$a \Phi \left(\frac{\tilde{\tau}_i}{\sigma_i} \right) + (b + c) \Phi \left(\frac{\tilde{\tau}_i - T}{\sigma_i} \right) = c \tag{2.3}$$

where $\Phi(z)$ is the cumulative distribution function of a standard Gaussian variable.

Proof of equation (2.3).

To begin, we write $\mathcal{J}$ as:

$$\mathcal{J} = \mathbb{E}[a [t_i]_+ + (b + c) [(t_i - T)]_+ - c (t_i - T)]$$

For a stationary point (with respect to the target time τ_i) it holds:

$$\frac{\partial}{\partial \tau_i} \mathcal{J} = a \mathbb{P}(\tau_i + \sigma_i \varepsilon_i > 0) + (b + c) \mathbb{P}(\tau_i - T + \sigma_i \varepsilon_i > 0) - c = 0$$

$$a \Phi \left(\frac{\tau_i}{\sigma_i} \right) + (b + c) \Phi \left(\frac{\tau_i - T}{\sigma_i} \right) = c$$

■

We observe that the optimal $\tilde{\tau}_i$ is unique, so the former relation allows us to well define the map $T \mapsto (\tilde{\tau}_i)_{i \in \mathbb{N}}$ we are looking for in this step. Considering remark 4, this map can be expressed in the following form:

$$T \mapsto (\tilde{\tau}_1, \dots, \tilde{\tau}_h) \quad \text{thus accounting for the } h \text{ different types of agents}$$

Step 2: We can now move on to the aggregation problem and calculate the updated distribution F^* as built in equation (2.1):

$$\begin{aligned} F^*(z, \tilde{\tau}_1, \dots, \tilde{\tau}_h) &= \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathbb{1}_{\{\tau_i \leq z\}} \\ &= p_1 \Phi\left(\frac{z - \tilde{\tau}_1}{\sigma_1}\right) + \dots + p_h \Phi\left(\frac{z - \tilde{\tau}_h}{\sigma_h}\right) \end{aligned}$$

Then, for every fixed $\theta \in (0, 1)$, T^* is defined by:

$$T^*(F^*) = [F^*]^{-1}(\theta) \quad (2.4)$$

This concludes our update on T given the strategies $(\tilde{\tau}_1, \dots, \tilde{\tau}_h)$.

Step 3: Finally, we would like to prove that the following map is a contraction:

$$\begin{array}{c} \widehat{T} : [0, T_{\max}] \longrightarrow [0, T_{\max}] \\ T \xrightarrow{\text{agent optimizing}} (\tilde{\tau}_i)_{i \leq h} \xrightarrow{\text{LLN}} F^* \xrightarrow{\text{quorum rule}} T^* \end{array}$$

Proposition 2. If $a, b, c > 0$ then $\widehat{T}$ is a contraction mapping of $[0, T_{\max}]$ with a unique fixed point T .

Proof.

Firstly we want to find an estimate for $\frac{d\tilde{\tau}_i}{dT}$, where $\tilde{\tau}_i$ is the optimal target time for agent i . By totally differentiating (clearly $\sigma_i > 0$) we get:

$$\begin{aligned} &\overbrace{a \Phi'\left(\frac{\tilde{\tau}_i}{\sigma_i}\right)}^U d\tilde{\tau}_i + (b+c) \overbrace{\Phi'\left(\frac{\tilde{\tau}_i - T}{\sigma_i}\right)}^V (d\tilde{\tau}_i - dT) = 0 \\ &[aU + (b+c)V] d\tilde{\tau}_i = (b+c)V dT \\ &\frac{d\tilde{\tau}_i}{dT} = \frac{(b+c)V}{aU + (b+c)V} \end{aligned}$$

And the following bounds hold:

$$\begin{aligned} U &\geq \Phi'\left(\frac{T_{\max}}{\sigma_i}\right) \geq \Phi'\left(\frac{T_{\max}}{\min_i \sigma_i}\right) \\ V &\leq \frac{1}{\sqrt{2\pi}} \end{aligned}$$

so for every positive (a, b, c) it exists $K < 1$ such that:

$$\frac{d\tilde{\tau}_i}{dT} \leq K < 1$$

We can now move on to verify the contraction mapping condition.

Denoting with $(\tilde{\tau}_i)_T$ the h optimal arrival times for a given T , we introduce the following map:

$$\mathcal{F}^*(z, T) = F^*(z, (\tilde{\tau}_i)_T)$$

and note that from the previous point it holds $\forall T$, and $y > 0$:

$$(\tilde{\tau}_i)_{T+y} \leq (\tilde{\tau}_i)_T + Ky \quad \forall i \leq h$$

Therefore an opposite inequality holds for the cumulative distributions:

$$\begin{aligned} F^*(z, (\tilde{\tau}_i)_{T+y}) &\geq F^*(z, (\tilde{\tau}_i)_T + Ky) = F^*(z - Ky, (\tilde{\tau}_i)_T) \\ \mathcal{F}^*(z, T+y) &\geq \mathcal{F}^*(z - Ky, T) \end{aligned} \quad (2.5)$$

Recalling the (2.4) definition for T^* , we observe that for every $y > 0$ and cumulative distribution function $F(\cdot)$ it holds:

$$T^*(F(\cdot - y)) = T^*(F(\cdot)) + y$$

and T^* is lower for cumulative functions that rise earlier: hence, applying T^* to both sides of (2.5), the inequality is again reversed and we finally get:

$$\begin{aligned} \hat{T}(T+y) &= T^*(\mathcal{F}^*(z, T+y)) \leq T^*(\mathcal{F}^*(z - Ky, T)) = T^*(\mathcal{F}^*(z, T)) + Ky \\ \hat{T}(T+y) &\leq \hat{T}(T) + Ky \\ \hat{T}(T+y) - \hat{T}(T) &\leq Ky \quad \text{where } K < 1 \end{aligned}$$

■

Step 4: As in the general framework, from rationality we derive that the expected T will coincide with the realized one.

Therefore our analysis, which is exact due to the infinite number of agents considered, gives $(T, (\tilde{\tau}_i)_{i \in \mathbb{N}})$ as a consistent solution.

How far have we moved from the section 1.1 framework?

Our toy model is not a differential game with continuous time t , since agents take just a single strategic decision (a so called one-shot game). Its state space has dimension 1; traces of the Brownian motions can be found in the normal realizations ε_i .

Moreover, the mean field T is not a probability measure on $\mathbb{R}^d$ but reduces to a scalar: this latter fact happens since T acts as an exhaustive summary for each agent of the deterministic distribution F .

We can easily imagine a context where F would be the problem mean field: for example, we can assume (according to the equation (1.3) form) that the agent i reputation cost depends on the share of participants that have arrived before t_i , and not just on the distance from T .

2.2 Numerical simulation of the toy model

The straightforward structure of our model allows a full numerical simulation.

We consider h agents' categories, corresponding to the different agents uncertainties.

From the previous analysis we know that for every fixed T there are exactly h different strategies $(\tau_j)_{j=1,\dots,h}$, which can be easily computed by numerically solving equation (2.3). This corresponds to step 1, as we are computing the response of every agent to a given mean field T .

From equation (2.2) we know the exact expression of $F^*(z)$ as function of the τ_i , so it can be numerically inverted to obtain the unique deterministic T^* such that:

$$F(T^*) = \theta$$

In this way we completely built the “mean field updating” map $\hat{T}$. Proposition 2 grants that it is a contraction from $[0, T_{\max}]$ to itself, so we can numerically find its unique fixed point.

Let us now discuss in further detail some aspects relative to the parameter choices.

2.2.1 Cost function parameters

Denoting with τ the agent's target arrival time, t the actual arrival time and T the meeting starting time (while it was scheduled for $T=0$), the agent cost is given by:

$$\mathcal{J}_{(a,b,c)}(t, T) = \mathbb{E} \begin{cases} c(T - t) & \text{for } t \leq 0 \quad (\text{implies } t \leq T) \\ c(T - t) + at & \text{for } 0 \leq t \leq T \\ at + b(t - T) & \text{for } t \geq T \end{cases}$$

It can be noticed that given an ordered triple (a, b, c) it holds:

$$\mathcal{J}_{(\lambda(a,b,c))}(t, T) = \lambda \mathcal{J}_{(a,b,c)}(t, T)$$

and they have the same minimal conditions, hence our model is *scale invariant*. Therefore, we can assign freely one of the three, and we choose to set:

$$b = 1$$

Are there any value sets for (a, c) that produce a foreseeable outcome for the fixed point T ?

Proposition 3. Let us consider a particular choice for the agent population and the required quorum: we assume that $h = 1$, so all the agents show the same uncertainty level, and $\theta = 0.5$. With these parameters, for every given T all the agents adopt the same $\tilde{\tau}_i = \tau$ and moreover it is going to be $T = \tau$ since by the time τ exactly a half of them will have arrived.

We claim that some values of the cost parameters (a, c) can a priori determine the starting time T , in the following way:

- For every triple such that $c < a + b$, the starting time T is necessarily 0.
- For every triple such that: $c > 2a + b$, the starting time T is necessarily $T_{\max}$.

Proof. Let us consider the minimization problem of an agent that knows the starting time T and is subject to an uncertainty σ on his arrival time. We want to show that for any T , his choice of $\tilde{\tau}$ will be such that $\tilde{\tau} < T$: this proves that the only possible fixed point for T is zero.

Let us introduce $w = t - T$ and write $\mathcal{J}$ as:

$$\mathcal{J} = \mathbb{E} \begin{cases} -cw & t \leq 0 (\leq T) \\ -(c-a)w + aT & 0 \leq t \leq T \\ (a+b)w + aT & t \geq T \end{cases}$$

By hypothesis $a + b > c > c - a$, so it is clear that the cost relative to a positive w is always higher than the one given by $-w$.

Since $w \sim \mathcal{N}(\tau - T, \sigma^2)$ (a symmetrical distribution), this implies that the optimal τ lies strictly before T .

For the other statement, the proof is very similar: for the same $\mathcal{J}$, being by hypothesis $c > 2a + b \rightarrow c > c - a > a + b$, the above symmetry argument leads to the conclusion that for every T , $\tilde{\tau}$ lies after it and the only possible fixed point is $T_{\max}$. ■

According to the above remarks, we run the model in this very particular case by limiting the set of adopted parameters $(a, 1, c)$ to the region:

$$(a, c) : \begin{cases} c < 2a + 1 \\ c > a + 1 \end{cases} \quad \text{parametrized by} \quad \begin{cases} c = \frac{1}{\sqrt{2}}(u + v) \\ a = 1 + \frac{1}{\sqrt{3}}(u + 2v) \end{cases} \quad (u, v) \in \mathbb{R}_+ \times \mathbb{R}_+$$

so that our equilibrium time $T(u, v)$ is not known immediately known from our cost choice.

A graphical output of the result is:

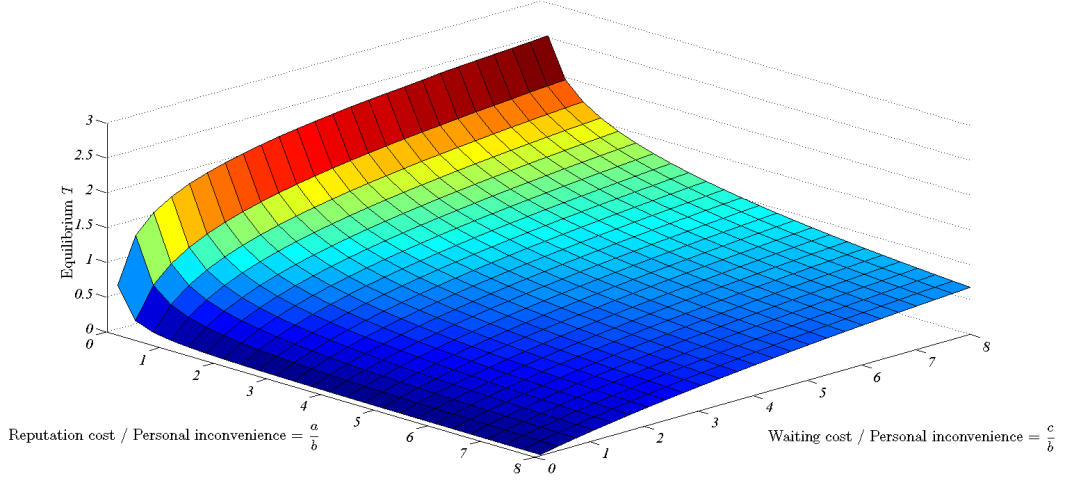


Figure 2.1: Equilibrium T for $u, v \in [0.05, 8]$, uniform $\sigma = 1$, $\theta = 0.5$

2.2.2 Agent population choice and quorum

After having shown that cost parameters can heavily influence the output significance, we now try to adopt a more realistic set-up for the agent population, described in general by $(\sigma_1, p_1), \dots, (\sigma_h, p_h)$.

Considering a time frame of 100 units, we arbitrarily divide the meeting participants in three groups:

- a) *The organizers:* they account for the 10% of total and have the lowest σ , set to 2.
- b) *The ones who come from afar:* they account for 15% of total and have the largest uncertainty, which we fix to 20.
- c) *The others:* they are the remaining ones and have a medium value for σ , let us say 6.

Coming to the quorum, we will consider a reasonable range, say: $\theta \in [0.41, 0.50]$.

2.2.3 Some results

Here below we present some resuming graphs for the equilibrium value of T at various quorum levels.

Note that the delay time increases as individuals care less about their formal late (reputation cost) and more about the time they could wait (waiting cost), both referred to a fixed level of personal lateness inconvenience (recall $b = 1$).

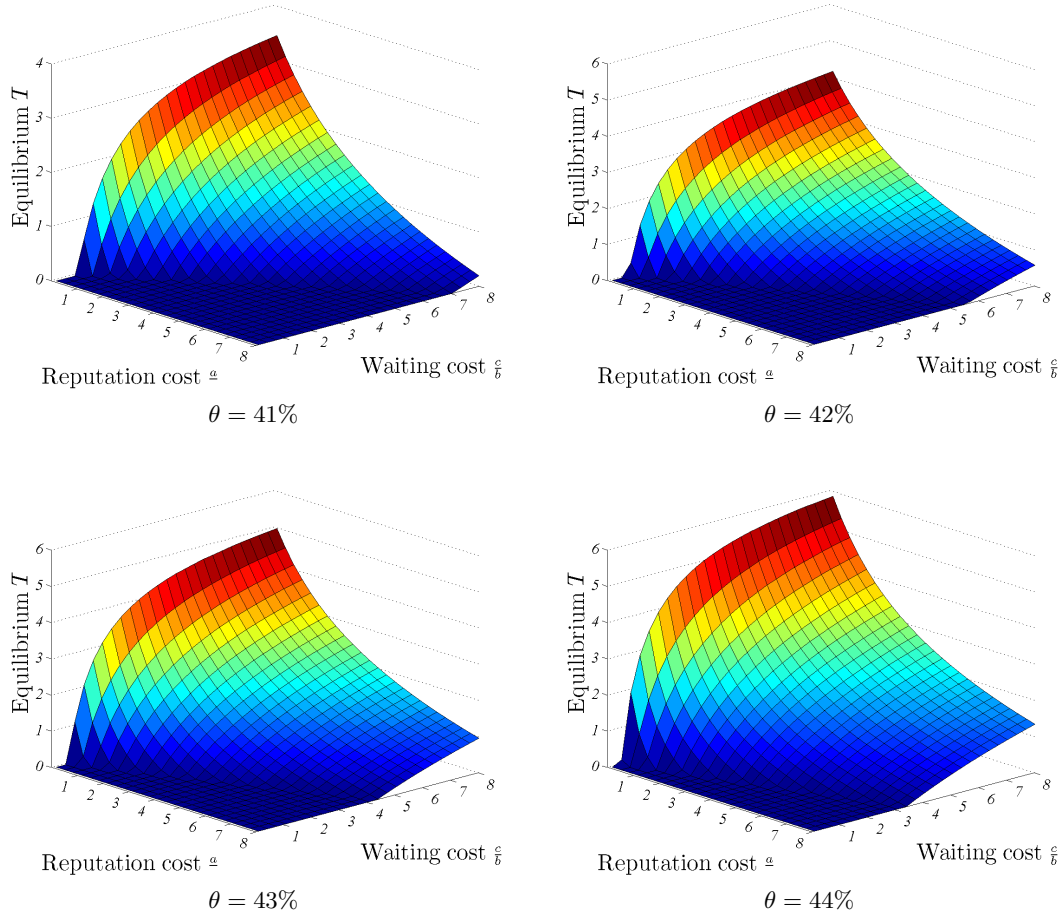


Figure 2.2: Equilibrium T for $u, v \in [0.1, 8]$, various θ , heterogeneous population

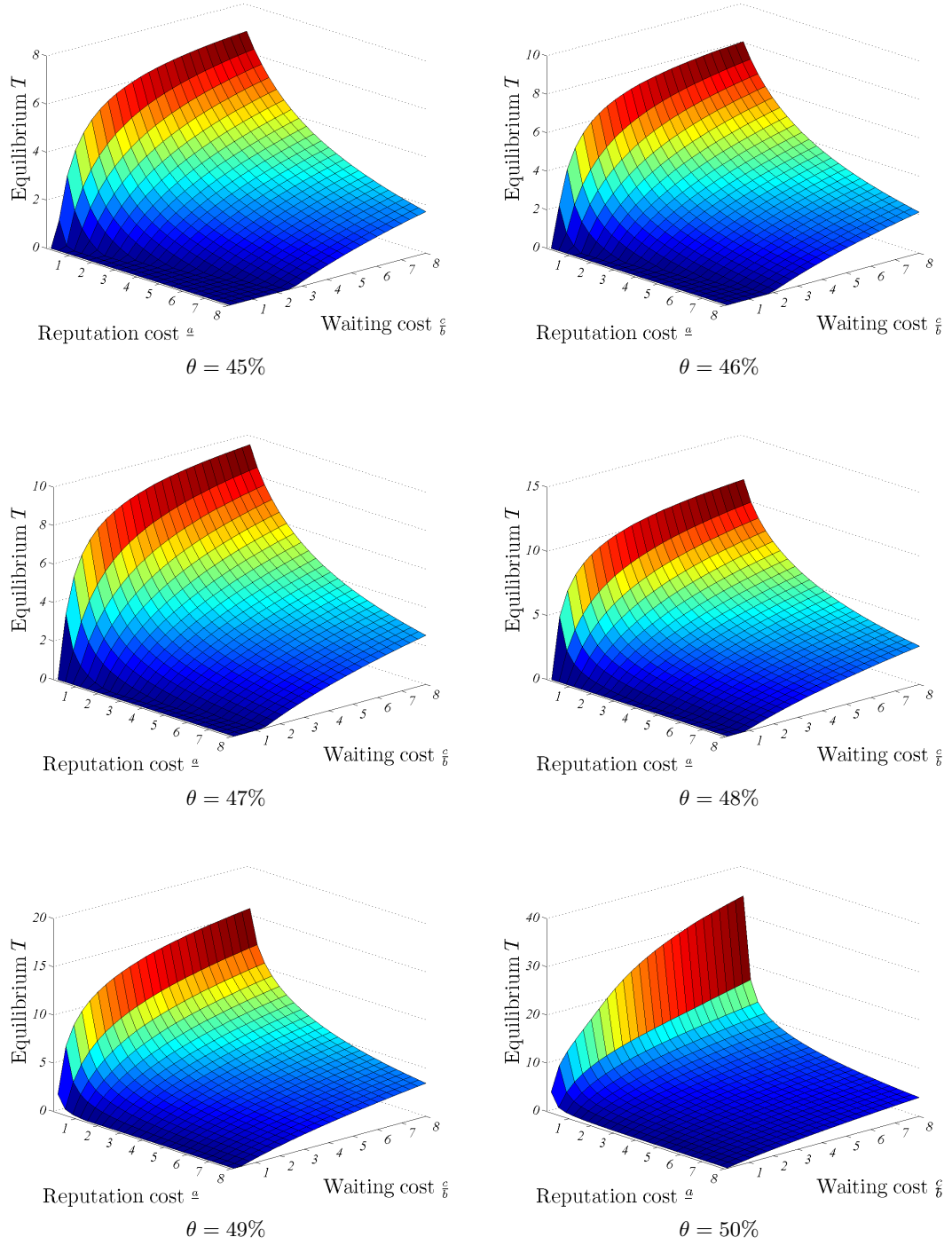


Figure 2.3: Equilibrium T for $u, v \in [0.1, 8]$, various θ , heterogeneous population

We can run a single meeting simulation and calculate the actual realized T from the quorum rule:

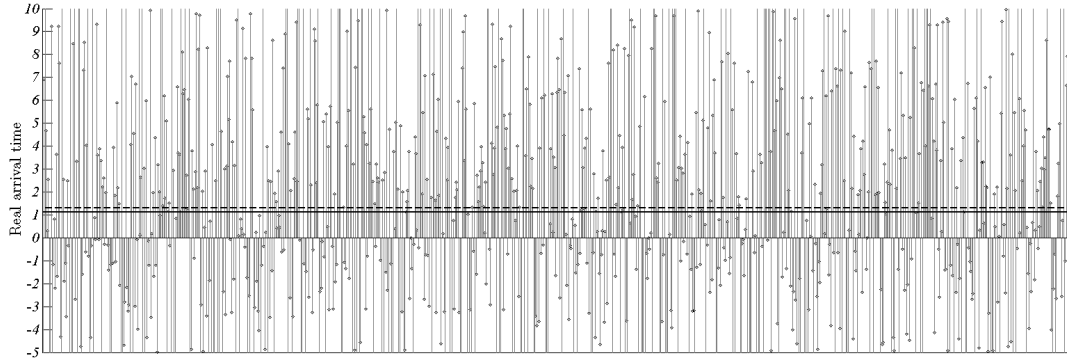


Figure 2.4: Arrival times of 1025 agents, $a = 1.5$, $b = 1$, $c = 3$, $\theta = 0.48$ and heterogeneous population. The solid line is the actual T , dashed one is the mean field computed T .

In the following graph we can observe the convergence of the realized T to the MFG computed one due to the law of large numbers, by using an increasing sample of participants and several realizations:

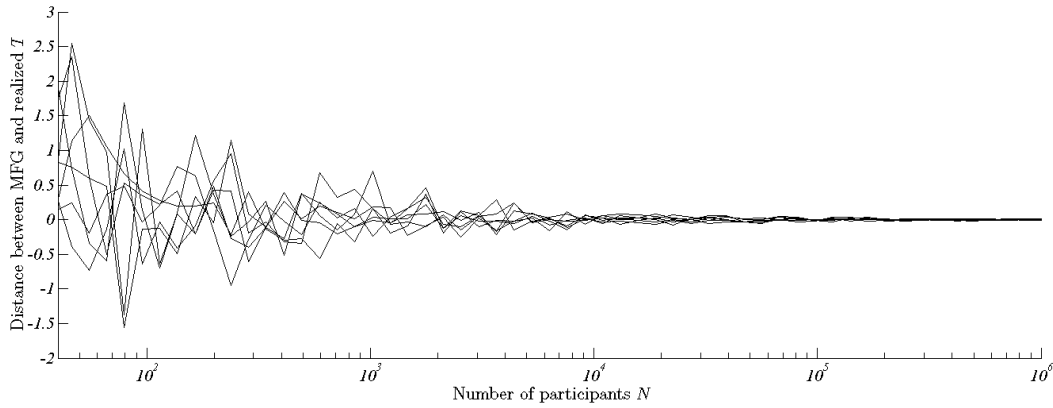


Figure 2.5: Same parameters as above, 7 different realizations for increasing values of N .

2.3 A digression on the concept of “continuum of agents”

This digression arises from the toy model presented in [16], where a continuum of agents is considered.

This particular assumption can be frequently found in the literature, especially in Mathematical Economics: a good overview can be found in [2].

The most commonly employed framework considers every agent as subject to an *idiosyncratic* standard Gaussian noise ε_i (meaning that $(\varepsilon_i)_{i \in [0,1]}$ is an i.i.d. $\mathcal{N}(0, 1)$ family) and then it is claimed that the law of large numbers applies.

Some conceptual and technical issues are immediately encountered when trying to formalize the *aggregate* noise distribution F .

First of all, we would like to speak of a continuous family of independent (normal) random variables on $\mathbb{R}$. Does such an object exists?

The answer is yes, and the construction relies on Kolmogorov’s Extension Theorem, helped by the fact that $\mathbb{R}$ is a σ -compact space: by asking for every finite dimensional distribution to be a spherical Gaussian vector we have the existence of a process:

$$Z : (\Omega, \mathcal{F}, \mathbb{P}) \rightarrow \left(\mathbb{R}^{[0,1]}, \mathcal{B} \left(\mathbb{R}^{[0,1]} \right) \right)$$

for some suitable probability space.

The real troubles come when we try to perform the aggregation and express the final cumulative distribution. Let us fix the independent family:

$$(Z_x)_{x \in [0,1]} \quad Z_x \sim \mathcal{N}(0, 1) \quad \forall x$$

A first aggregation formula can be proposed by analogy with the empirical distribution in the finite case:

$$F_N(t) = \frac{1}{N} \sum_{n=1}^N \mathbb{1}_{Z_n \leq t} \xrightarrow{\text{moving to the continuum}} F(t) = \int_0^1 \mathbb{1}_{Z_x \leq t} dx \quad (2.6)$$

where the Lebesgue measure on $[0,1]$ has been used as weight to average the single realizations.

The quite nasty fact is that the map $x \mapsto Z_x$ needs not to be Lebesgue-measurable and thus the integral above has in general no sense. This is strictly related to the fact that the process Z_x is not jointly measurable in $(\Omega \times [0, 1], \mathcal{F} \otimes \mathcal{B}([0, 1]))$.

Remark 5. Note that $\mathbb{1}_{Z_x \leq t}$ is a Bernoulli with parameter $p = \mathbb{P}\{Z \leq t\}$, so the integral expression above for a certain (t) can be expressed as:

$$W = \int_0^1 B_x dx \quad \text{where } (B_x)_{x \in [0,1]} \text{ is a family of i.i.d. Bernoulli of parameter } p \quad (2.7)$$

Our intuition would state that the above W should be constantly equal to p , since from law of large numbers we expect that, given an infinite number of B_x , a deterministic portion p of them takes value 1.

Surprisingly enough, this proves not to be true, since the set of variables which take value 1 could be not measurable, or (even worse) its measure could be arbitrary and not equal to p .³

This problem was already noticed by Doob [9]; since then, many different workarounds have been proposed in order to back the intuitive use of LLN and independence with rigorous foundations.

A brilliant comparison is presented in [2]; for our purpose we just mention the *Fubini extension* approach, which was introduced with nonstandard arguments by Y. Sun [21] and then recently expressed in the standard framework by K. Podczeck [20].

2.3.1 Extending the result to non identically distributed families

Now, supposing we managed to give full meaning to equation (2.6), this LLN application can be generalized to non identically distributed continuous families.

For example, the meeting participants in [16] have real times determined by independent Gaussian variables of varying mean and variance:

$$t_i \sim \mathcal{N}(\tau_i, \sigma_i) = \tau_i + \sigma_i \varepsilon_i$$

Choosing $[0, 1]$ as continuum indexing set, the resulting aggregate distribution is then given by:

$$F(w) = \int_0^1 \mathbb{1}_{\{t_x \leq w\}} dx = \int_0^1 \mathbb{1}_{\{\varepsilon_x \leq \frac{w - \tau_x}{\sigma_x}\}} dx = \int_0^1 \mathbb{1}_{\{\varepsilon_x \leq g(x)\}} dx \quad \text{for a measurable map } g(x)$$

Let us now assume that $g(x) > -R$ for some R , then it exists a monotone succession (g_n) of simple functions in the form:

$$g_n = \sum_{k \leq N_n} a_k \mathbb{1}_{A_k}(x) \quad \text{for some measurable sets } (A_k) \text{ in } [0, 1]$$

such that $g_n(x) \uparrow g(x)$, hence:

$$\int_0^1 \mathbb{1}_{\{\varepsilon_x \leq g_n(x)\}} dx \uparrow \int_0^1 \mathbb{1}_{\{\varepsilon_x \leq g(x)\}} dx$$

³These negative results have been separately stated in 1985 by K.L. Judd and in the same year by M. Feldman and C. Gilles, both in the Journal of Mathematical Economics.

and we can now focus on the simpler integral:

$$\int_0^1 \mathbb{1}_{\{\varepsilon_x \leq \sum_{k=1}^{N_n} a_k \mathbb{1}_{A_k}(x)\}} dx = \sum_{k \leq N_n} \int_{A_k} \mathbb{1}_{\{\varepsilon_x \leq a_k\}} dx + \int_{A_k^c} \mathbb{1}_{\{\varepsilon_x \leq 0\}} dx$$

and observe that the two terms are indeed two Bernoulli variables of parameters $\mathbb{P}\{\varepsilon_x \leq a_k\}$ and $\mathbb{P}\{\varepsilon_x \leq 0\}$.

In this way we reduced the apparently hard problem of varying parameters (in this case mean and variance) to the one represented by equation (2.7).

In conclusion, the fundamental question is still the following:

Given a continuum of independent coin tosses, is it true that the portion p of heads is deterministic and coincides with the underlying probability of getting a head?

Chapter 3

The full MFG scenario

Let us briefly recall our initial difficult problem in its infinite players limit (“mean field”) formulation:

In the time interval $[0, T]$, an infinite number of homogeneous agents interact. Each agent i has starting condition X_0^i , a random variable independent from the others that has law m_0 , while his state $X_t^i \in \mathbb{R}^d$ evolves following a controlled drift:

$$dX_t^i = \alpha_t^i dt + \sqrt{2} dB_t^i$$

The cost for an agent with state X_t^i observing m and adopting a control α is given by:

$$\mathcal{J}(\alpha, m) = \mathbb{E} \left[\int_0^T \left[\frac{1}{2} |\alpha_s|^2 + F(X_s^i, m(\cdot, s)) \right] ds + G(X_T^i, m(\cdot, T)) \right]$$

What will our fixed-point reasoning involve this time? What are the mathematical tools that will prove helpful to define the mean field?

The game we are considering is built upon a stochastic differential dynamic, so we expect SDE and PDE theory to come in handy with our questions.

Let us begin by introducing some definitions.

Definition 4 ($\mathcal{P}_1$ space). We define $\mathcal{P}_1$ as the space of probability measures μ on $\mathcal{B}(\mathbb{R}^d)$ with finite first order moment, $\int_{\mathbb{R}^d} |x| d\mu(x) < +\infty$.

Hereby, it is assumed that all the measures in our exposition belong to $\mathcal{P}_1$.

In order to deal with convergence in $\mathcal{P}_1$, it seems very natural to introduce a distance:

Definition 5 (Kantorovitch - Rubinstein distance). We define the following distance on $\mathcal{P}_1$:

$$d_1(\mu, \nu) = \inf_{\gamma \in \Pi(\mu, \nu)} \left[\int_{\mathbb{R}^{2d}} |x - y| d\gamma(\mu, \nu) \right]$$

where $\Pi(\mu, \nu)$ is the following subset of probability measures on $\mathcal{B}(\mathbb{R}^{2d})$:

$$\Pi(\mu, \nu) = \left\{ \gamma \text{ s.t. } \forall A \in \mathcal{B}(\mathbb{R}^d), \quad \gamma(A \times \mathbb{R}^d) = \mu(A), \quad \gamma(\mathbb{R}^d \times A) = \nu(A) \right\}$$

Theorem 6 (Equivalent definition for d_1 - without proof). For any $m, m' \in \mathcal{P}_1$, it holds:

$$d_1(m, m') = \sup_{f \in 1-Lip} \left\{ \int_{\mathbb{R}^d} f(x) dm(x) - \int_{\mathbb{R}^d} f(x) dm'(x) \right\}$$

By using d_1 , we can introduce the following distance on $\mathcal{C}_0([0, T], \mathcal{P}_1)$:

$$\tilde{d}_1(\mu, \nu) = \sup_{t \in [0, T]} d_1(\mu(t), \nu(t)) \quad \text{for any continuous } \mu, \nu : [0, T] \rightarrow \mathcal{P}_1$$

It is also convenient to define the following space of real valued continuous functions:

Definition 7. We define $\mathcal{C}^{l+\alpha}$ for an integer $l \geq 0$ and $\alpha \in (0, 1)$ the set of continuous functions $u : \mathbb{R}^d \times [0, T] \rightarrow \mathbb{R}$ such that, for every non-negative integers $r, s : 2r + s \leq l$, the derivative $(\partial_t)^r (\nabla)^s$ exists, is bounded and both α -Hölder in space and $\alpha/2$ -Hölder in time.

This space is complete if equipped with the norm:

$$\|u\|^{(l+\alpha)} = \sum_{2r+s \leq l} (\|(\partial_t)^r (\nabla)^s u\|_\infty + \|(\partial_t)^r (\nabla)^s u\|_{(x, \alpha)} + \|(\partial_t)^r (\nabla)^s u\|_{(t, \alpha/2)})$$

being:

$$\|f(x, t)\|_{(x, \alpha)} = \sup_{\substack{(x, t) \in \mathbb{R}^d \times [0, T] \\ |x_1 - x_2| < k_0}} \frac{|f(x_1, t) - f(x_2, t)|}{|x_1 - x_2|^\alpha}$$

and similarly for $(t, \alpha/2)$. Choosing different k_0 given equivalent norms, so we neglect this dependence.

We define the following functions and assume them to be continuous:

$$\begin{array}{ll} m : [0, T] \longrightarrow \mathcal{P}_1 & F, G : \mathbb{R}^d \times \mathcal{P}_1 \longrightarrow \mathbb{R} \\ t \longmapsto m(t) & (x, m) \longmapsto F(x, m), G(x, m) \end{array}$$

It follows that the maps:

$$\begin{array}{l} (x, t) \longmapsto F(x, m(t)) \\ x \longmapsto G(x, m(T)) \end{array}$$

are continuous.

We also consider a probability space $(\Omega, \mathcal{F}, \mathbb{P})$ and a filtration $(\mathcal{F}_t)_{t \in [0, T]} \subset \mathcal{F}$ such that an infinite family $(B_t^i)_{i \in \mathbb{N}}$ of independent brownian motions, each of those is adapted to $(\mathcal{F}_t)$, is defined on it.

3.1 The first step: the agent optimal control problem

As in section 1.2, we begin by studying the agent's optimization problem under the assumption that $m(\cdot, t)$ is known:

$$\begin{aligned} \inf_{\alpha_t \in \mathcal{A}} \mathcal{J}(\alpha) = \inf_{\alpha_t \in \mathcal{A}} \mathbb{E} \left[\int_0^T \left[\frac{1}{2} |\alpha_s|^2 + F(X_s, m(\cdot, s)) \right] ds + G(X_T, m(\cdot, T)) \right] \\ \text{being } \begin{cases} dX_s = \alpha_s ds + \sqrt{2} dB_s & \text{for } s \in [0, T] \\ X_0 \sim m_0 \end{cases} \end{aligned} \quad (3.1)$$

Let us properly specify the above formulation. As $\mathcal{A}$ we consider:

$$A = \left\{ \alpha_t \in L^2([0, T] \times \Omega, \mathbb{R}^d); \quad \alpha_t \text{ progressively measurable} \right\}$$

and we assume that the maps:

$$\begin{aligned} (x, t) &\longmapsto F(x, m(t)) \\ x &\longmapsto G(x, m(T)) \end{aligned}$$

are continuous, bounded and $\frac{1}{2}$ -Hölder.

Under these hypotheses, dynamic programming theory (see [22]) proves the existence of the following optimal control in feedback form:

$$\tilde{\alpha}(X_t, t) = -\nabla u(X_t, t)$$

where the value function $u(x, t)$ is $\mathcal{C}^2$ in space, $\mathcal{C}^1$ in time and verifies in the classical sense the Hamilton-Jacobi-Bellman PDE for this optimal control problem.

In the following, we are going to partly justify the HJB origin and to rigorously prove that under our assumptions it is a sufficient description of the optimal control $\tilde{\alpha}_t$.

To begin, let us introduce the *value function* for the above cost functional, assuming that m is given:

$$\begin{aligned} u(x, t) = \inf_{\alpha \in \mathcal{A}} \mathbb{E} \left[\int_t^T \left[\frac{1}{2} |\alpha_s|^2 + F(X_s, m(\cdot, s)) \right] ds + G(X_T, m(\cdot, T)) \right] \\ \text{being } \begin{cases} dX_s = \alpha_s ds + \sqrt{2} dB_s & \text{for } t \leq s \leq T \\ X_t = x \text{ (the input variable is the initial condition)} \end{cases} \end{aligned}$$

The value of u represents the cost that any player would assign to the state x at time t .

Stochastic control theory states that if $u(x, t)$ is $\mathcal{C}^{2,1}$, it necessarily verifies a certain HJB (see Prop 3.5, chap. 4 of [22]) in the classical sense; however, u needs not to have such a regularity and from the above formulation it is not evident if this is the case.

A workaround for this obstacle is to adopt a verification technique: we can still write the HJB assuming $u \in \mathcal{C}^{2,1}$, find a solution $\tilde{u}(x, t)$ of class $\mathcal{C}^{2,1}$ and then apply some *verification* theorem to prove that this solution is the value function $u(x, t)$ for the initial optimal control problem.

The rigorous derivation of the HJB equation requires vast premises, and would lead us far from our aims.

Nevertheless, we find it interesting to present an heuristic derivation that holds if we restrict the set of possible controls to a.s. *continuous* processes.

We begin by considering an intermediate time $q \in [t, T]$, and introduce the following notation:

$$\mathcal{A}_t^q = \left\{ \alpha_s \in L^2([t, q] \times \Omega, \mathbb{R}^d); \quad \alpha_s \text{ a.s. continuous and progressively measurable} \right\}$$

Then, the *Bellman optimality principle* states that:

$$\begin{aligned} u(x, t) = \inf_{\alpha \in \mathcal{A}_t^q} \mathbb{E} \left[\int_t^q \left[\frac{1}{2} |\alpha_s|^2 + F(X_s, m(\cdot, s)) \right] ds + u(X_q, q) \right] \\ \text{being } \begin{cases} dX_s = \alpha_s ds + \sqrt{2} dB_s & \text{for } t \leq s \leq q \\ X_t = x \end{cases} \end{aligned} \quad (3.2)$$

We can move $u(x, t)$ to the right side and write $\mathbb{E}[u(X_q, q) - u(x, t)]$ in integral form using Ito's formula:

$$0 = \inf_{\alpha \in \mathcal{A}_t^q} \mathbb{E} \left[\int_t^q \left[\frac{1}{2} |\alpha_s|^2 + F(X_s, m(\cdot, s)) + \nabla u(X_s, s) \cdot \alpha_s + \partial_t u(X_s, s) + \Delta u(X_s, s) \right] ds \right]$$

Let us now consider $q = t + dt$ and let $dt \rightarrow 0$. At the limit, the *inf* involves only the control value at time t , denoted with α :

$$\inf_{\alpha \in \mathbb{R}^d} \left(\frac{1}{2} |\alpha|^2 + F(x, m(\cdot, t)) + \nabla u(x, t) \cdot \alpha + \partial_t u(x, t) + \Delta u(x, t) \right) dt = 0$$

The total cost contribution depending on α is then given by:

$$\nabla u(x, t) \cdot \alpha + \frac{1}{2} |\alpha|^2$$

Thus we get, for every t :

$$\begin{cases} -\partial_t u - \Delta u - \inf_{\alpha \in \mathcal{A}} \left(\nabla u \cdot \alpha_t + \frac{1}{2} |\alpha_t|^2 \right) = F(x, m(\cdot, t)) \\ u(x, T) = G(x, m(\cdot, T)) \end{cases} \quad (3.3)$$

It can be noticed that, for every $\alpha \in \mathbb{R}^d$:

$$\min_{\alpha \in \mathbb{R}^d} \left(\nabla u \cdot \alpha + \frac{1}{2} |\alpha|^2 \right) = -\frac{1}{2} |\nabla u|^2 \quad \text{for } \alpha = -\nabla u(x, t)$$

And the final form of the HJB partial differential equation is:

$$\begin{cases} -\partial_t u + \frac{1}{2} |\nabla u|^2 - \Delta u = F(x, m(\cdot, t)) \\ u(x, T) = G(x, m(\cdot, T)) \end{cases} \quad (3.4)$$

Thanks to the assumed regularity of $F(x, t)$, $G(x)$, theorem A.1 can be applied to the above PDE (when adequately transformed - see section 3.3.1), providing the unique existence of a classical solution u with bounded gradient ∇u . This assures the strong existence and uniqueness of a solution to the closed loop state equation:

$$\begin{cases} dX_t = -\nabla u(X_t, t) dt + \sqrt{2} dB_t, & t \in [0, T] \\ X_0 \sim m_0 \end{cases}$$

leading to the conclusion that $\tilde{\alpha}_t = -\nabla u(X_t, t)$ is an admissible control. To prove that it is optimal, we need a verification theorem.

Theorem 8 (Verification theorem for the agent optimal control problem).

Under the hypotheses at the beginning of the section, let u be the unique classical solution¹ to the following Hamilton-Jacobi-Bellman (HJB) equation:

$$\begin{cases} -\partial_t u - \Delta u + \frac{1}{2} |\nabla u|^2 = F(x, m(t)) & \text{in } \mathbb{R}^d \times (0, T) \\ u(x, T) = G(x, m(T)) & \text{in } \mathbb{R}^d \end{cases} \quad (3.5)$$

Then the strategy defined by:

$$\tilde{\alpha} = -\nabla u(x, t)|_{x=X_t^i} = -\nabla u(X_t^i, t) \quad (3.6)$$

is the optimal strategy for every agent i .

Its adoption causes the agents' state to satisfy:

$$\begin{cases} dX_t = -\nabla u(X_t, t) dt + \sqrt{2} dB_t, & t \in [0, T] \\ X_0 \sim m_0 \end{cases} \quad (3.7)$$

Proof. The last part is a straightforward consequence of the problem setting outlined in equation (3.1).

¹A unique solution exists for theorem A.1, which can be applied to the Hopf-Cole transformed version of (3.5).

We need to show that $\tilde{\alpha}_t$ actually minimizes the cost functional, i.e.:

$$\mathcal{J}(\tilde{\alpha}) \leq \mathcal{J}(\alpha) \quad \forall \alpha_t \in \mathcal{A} \quad (3.8)$$

To begin we observe that X_t is an Ito process and therefore, recalling $dX_s = \alpha_s dt + \sqrt{2} dB_s$ and (3.5), we get:

$$\begin{aligned} \mathbb{E}[G(X_T, m(T))] &= \mathbb{E}[u(X_T, T)] \\ &= \mathbb{E} \left[u(X_0, 0) + \int_0^T (\partial_t u(X_s, s) + \alpha_s \cdot \nabla u(X_s, s) + \Delta u(X_s, s)) ds \right] \\ &= \mathbb{E} \left[u(X_0, 0) + \int_0^T \left(\frac{1}{2} |\nabla u(X_s, s)|^2 + \alpha_s \cdot \nabla u(X_s, s) - F(X_s, m(s)) \right) ds \right] \end{aligned}$$

Since, for every $\alpha \in \mathbb{R}^d$, $\frac{1}{2} |\nabla u|^2 + \alpha \cdot \nabla u + \frac{1}{2} |\alpha|^2 \geq 0$, we have $\frac{1}{2} |\nabla u|^2 + \alpha \cdot \nabla u \geq -\frac{1}{2} |\alpha|^2$ and thus:

$$\begin{aligned} \mathbb{E}[G(X_T, m(T))] &\geq \mathbb{E} \left[u(X_0, 0) + \int_0^T \left(-\frac{1}{2} |\alpha_s|^2 - F(X_s, m(s)) \right) ds \right] \\ \mathbb{E}[u(X_0, 0)] &\leq \underbrace{\mathbb{E} \left[\int_0^T \left(\frac{1}{2} |\alpha_s|^2 + F(X_s, m(s)) \right) ds + G(X_T, m(T)) \right]}_{=\mathcal{J}(\alpha)} \end{aligned}$$

Finally, if we set $\alpha = \tilde{\alpha}_t$ in the above, all the estimates become equalities, so:

$$\mathcal{J}(\tilde{\alpha}) = \mathbb{E}[u(X_0, 0)] \leq \mathcal{J}(\alpha)$$

And this condition for an admissible control is sufficient for optimality. $\blacksquare$

In this way we have shown that the uniquely solvable Hamilton-Jacobi-Bellman PDE characterises the optimal agent strategy for a given $m(\cdot, t)$. The first step is completed.

3.2 The second step: aggregation

Let us now suppose to equip every agent with the optimal feedback strategy resumed by u .

What would we see at the aggregate level?

The state X_t^i of each agent is a process that verifies:

$$\begin{cases} dX_t^i = -\nabla u(X_t^i, t) dt + \sqrt{2} dB_t^i, & t \in [0, T] \\ X_0^i \sim m_0 \end{cases} \quad (3.9)$$

Since $(X_0^i)_{i \in \mathbb{N}}$ and $(B_t^i)_{i \in \mathbb{N}}$ are independent, the states X_t^i constitute an i.i.d. family. Let us henceforth write X_t to denote a process with the same law of every X_t^i .

Let us now consider the “limit measure” definition of $m(t)$, as initially introduced in section 1.2:

$$m(t) = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \delta_{X_t^i} \quad \text{in the weak convergence}$$

Naming $\mu_t^N = \frac{1}{N} \sum_{i=1}^N \delta_{X_t^i}$, for every continuous and bounded $\phi : \mathbb{R}^d \rightarrow \mathbb{R}$ we get:

$$\int_{\mathbb{R}^d} \phi(x) d\mu_t^N(x) = \frac{1}{N} \sum_{i=1}^N \phi(X_t^i) \xrightarrow{\text{LLN}} \mathbb{E}[\phi(X_t)] = \int_{\mathbb{R}^d} \phi(x) dm(t)(x)$$

This shows that $m(t)$ coincides with the law of X_t^i for every i .

To determine $m(t)$, the aggregate agent distribution at time t , it therefore suffices to solve the SDE verified by X_t . To do this, a nice prompt is offered by a fundamental property of the stochastic differential equations, the SDE–PDE correspondence.

To recall it, we need an important notion of weak solution for a particular parabolic PDE, the Fokker-Planck equation.

Definition 9. For any continuous $\mathbf{b} : \mathbb{R}^d \times [0, T] \rightarrow \mathbb{R}^d$, let us consider the Fokker-Planck PDE:

$$\begin{cases} \partial_t m - \Delta m + \operatorname{div}(m \mathbf{b}) = 0 & \text{in } \mathbb{R}^d \times (0, T) \\ m(x, 0) = m_0(x) \end{cases} \quad (3.10)$$

We say that a measure $\mu \in L^1([0, T], \mathcal{P}_1)$ is a *probability measure* solution of (3.10) if for every $\phi \in \mathcal{C}_c^\infty(\mathbb{R}^d \times [0, T])$ it verifies:

$$\int_{\mathbb{R}^d} \phi(x, 0) dm_0(x) + \int_0^t \int_{\mathbb{R}^d} [\partial_t \phi(x, s) + \nabla \phi(x, s) \cdot \mathbf{b}(x, s) + \Delta \phi(x, s)] dm(s)(x) ds = 0$$

and the following correspondence can be proved:

Theorem 10. (*Correspondence between SDE and PDE*)
On one hand, let X_t be the unique strong solution of:

$$\begin{cases} dX_t = \mathbf{b}(X_t, t) dt + \sqrt{2} dB_t, & t \in [0, T] \\ X_0 \sim m_0 \end{cases} \quad (3.11)$$

where we assume that the initial distribution m_0 has density, still denoted with $m_0(x)$, on $\mathbb{R}^d$ and the vector field $\mathbf{b} : \mathbb{R}^d \times [0, T] \rightarrow \mathbb{R}^d$ is continuous, uniformly Lipschitz in x

and bounded.

On the other hand, consider the following (Fokker-Planck) PDE:

$$\begin{cases} \partial_t m - \Delta m + \operatorname{div}(m \mathbf{b}) = 0 & \text{in } \mathbb{R}^d \times (0, T) \\ m(x, 0) = m_0(x) \end{cases} \quad (3.12)$$

Then the law of X_t has a density $m(x, t)$ on $\mathbb{R}^d$ and this density weakly solves the above Fokker-Planck partial differential equation.

Proof. Our proof relies on a direct use of Ito's Formula. If $\phi(x, t)$ is bounded, $\mathcal{C}^2$ in the spatial variable x and $\mathcal{C}^1$ in time, then for the Ito's process X_t that verifies equation (3.11):

$$\begin{aligned} \phi(X_t, t) &= \phi(X_0, 0) + \int_0^t [\partial_t \phi(X_s, s) + \nabla \phi(X_s, s) \cdot \mathbf{b}(X_s, s) + \Delta \phi(X_s, s)] \, ds + \\ &\quad + \int_0^t \nabla \phi(X_s, s) \cdot \, dB_s \end{aligned}$$

We can take the expected value of both sides:

$$\begin{aligned} \mathbb{E}[\phi(X_t, t)] &= \mathbb{E} \left[\phi(X_0, 0) + \int_0^t [\partial_t \phi(X_s, s) + \nabla \phi(X_s, s) \cdot \mathbf{b}(X_s, s) + \Delta \phi(X_s, s)] \, ds \right] + \\ &\quad + \underbrace{\mathbb{E} \left[\int_0^t \nabla \phi(X_s, s) \cdot \, dB_s \right]}_{=0} \end{aligned}$$

Writing explicitly the expected value in terms of $m(t)$, the law of X_t , we have:

$$\begin{aligned} \int_{\mathbb{R}^d} \phi(x, t) \, dm(t)(x) &= \int_{\mathbb{R}^d} \phi(x, 0) \, dm_0(x) + \\ &\quad + \int_0^t \int_{\mathbb{R}^d} [\partial_t \phi(x, s) + \nabla \phi(x, s) \cdot \mathbf{b}(x, s) + \Delta \phi(x, s)] \, dm(s)(x) \, ds \end{aligned}$$

In particular, substituting $t = T$ and recalling that ϕ has compact support in $\mathbb{R}^d \times [0, T)$, we obtain that m is a weak solution to equation (3.12).

It is also true that it has density with respect to Lebesgue measure if the initial distribution m_0 has density, but we omit this part of the proof since it is quite technical and not directly related to our objectives. $\blacksquare$

Remark 6. In 2007, an important uniqueness result [3] has been established for probability measure-valued solutions of parabolic equations. Thanks to it, we can state that the above Fokker-Planck admits a unique probability measure solution as in definition 9.

Coming to our question, we can set:

$$\mathbf{b}(x, t) = -\nabla u(x, t)$$

and find that if $u(x, t)$ is regular enough, the law of X_t admits a density $m(x, t)$ that verifies:

$$\begin{cases} \partial_t m - \Delta m - \operatorname{div}(m \nabla u) = 0 & \text{in } \mathbb{R}^d \times (0, T) \\ m(x, 0) = m_0(x) \end{cases} \quad (3.13)$$

What happens if we solve this PDE? Can we claim to have found the law of X_t ?

Proposition 11. Under the same hypotheses adopted in theorem 10, the unique probability measure solution of (3.13) is the law of X_t .

Proof. Let us consider the agent state SDE (3.9) and recall that its initial condition X_0 is supposed to have density m_0 on $\mathbb{R}^d$.

We know that it admits a unique strong solution X_t , whose law $m(x, t)$ verifies the PDE (3.13) for the above theorem. Hence, if we consider a weak probability measure solution $\mu(x, t)$ to the Fokker-Planck with initial density m_0 , by uniqueness it has to be $\mu(x, t) = m(x, t)$. $\blacksquare$

We can conclude that the Fokker-Planck PDE fully describes the function $m(\cdot, t)$ for a given value function $u(x, t)$ and the second step of the mean field approach is completed.

3.3 The mean field games PDE formulation

Recollecting the findings from our previous steps we obtain a system of fully coupled PDE:

$$\begin{cases} -\partial_t u - \Delta u + \frac{1}{2}|\nabla u|^2 = F(x, m) & \text{in } \mathbb{R}^d \times (0, T) \\ \partial_t m - \Delta m - \operatorname{div}(m \nabla u) = 0 & \text{in } \mathbb{R}^d \times (0, T) \\ u(x, T) = G(x, m(T)) \\ m(x, 0) = m_0(x) \end{cases} \quad (3.14)$$

Remark 7. Solving this system is equivalent to find the fixed point we have been looking for in the previous chapters.

In fact, a solution (u, m) resumes both the individual optimization and the state aggregation steps in a consistent manner: on one hand the optimal strategy described by u produces m as mean field and on the other hand the mean field m is the one observed by every agent when choosing u .

Following [6], an important existence and uniqueness result can be proved for this setting under some adequate assumptions.

We highlight that this coupled-type problem has been introduced less than ten years ago and since then many different research paths have been opened. The hypotheses set we are going to adopt offer a good compromise since they allow for a straightforward proof exposition while still conserving the conceptual challenges of the mean field approach, such as the necessity to deal with measure inputs and the strong coupling between the two PDEs.

Theorem 12 (Existence and uniqueness). *Let us assume that it exists some constant C_0 such that:*

1. *F and G are uniformly bounded by C_0 ;*
2. *They are uniformly Lipschitz with constant C_0 over their whole domain:*

$$\begin{aligned} |F(x_1, m_1) - F(x_2, m_2)| &\leq C_0[|x_1 - x_2| + d_1(m_1, m_2)] \\ |G(x_1, m_1) - G(x_2, m_2)| &\leq C_0[|x_1 - x_2| + d_1(m_1, m_2)] \end{aligned} \quad (3.15a)$$

Moreover, we require that $m_0 \ll L(\mathbb{R}^d)$, and the density $m_0(x) = \frac{dm_0}{dx}$ is Hölder continuous and satisfies:

$$\int_{\mathbb{R}^d} |x|^2 m_0(x) dx < +\infty$$

Finally, we assume the following (monotonicity) conditions on F and G :

$$\begin{aligned} \int_{\mathbb{R}^d} (F(x, m_1) - F(x, m_2)) d(m_1 - m_2)(x) &> 0 \quad \forall m_1 \neq m_2 \in \mathcal{P}_1 \\ \int_{\mathbb{R}^d} (G(x, m_1) - G(x, m_2)) d(m_1 - m_2)(x) &\geq 0 \quad \forall m_1, m_2 \in \mathcal{P}_1 \end{aligned} \quad (3.15b)$$

Under the above assumptions, there exists a unique pair (u, m) such that $u, m : \mathbb{R}^d \rightarrow \mathbb{R}$ are continuous, of class $\mathcal{C}^2$ in space and $\mathcal{C}^1$ in time and (u, m) satisfies (3.14) in the classical sense. Such a pair is said to be a classical solution to the mean field problem.

We divide the *proof* in its two parts.

3.3.1 Existence proof

This proof is built upon a classical functional analysis approach. A good starting point is the following set:

Proposition 13. For every positive constant $\mathcal{C}_1$, let us define:

$$\mathcal{K} = \left\{ \mu \in \mathcal{C}^0([0, T], \mathcal{P}_1) \quad \text{s.t.} \quad \sup_{s \neq t} \frac{d_1(\mu(s), \mu(t))}{|t - s|^{\frac{1}{2}}} \leq \mathcal{C}_1 \quad \text{and} \quad \sup_{t \in [0, T]} \int_{\mathbb{R}^d} |x|^2 d\mu(t)(x) \leq \mathcal{C}_1 \right\}$$

We state that $\mathcal{K}$ is a convex compact subset of $\mathcal{C}^0([0, T], \mathcal{P}_1)$.

Proof. Closeness follows from the fact that both conditions in the definition are closed.

For convexity, let $\nu(t) = h_1\mu_1(t) + h_2\mu_2(t)$ be a convex combination of measures in $\mathcal{K}$, then:

$$\begin{aligned} d_1(\nu(t), \nu(s)) &= \sup_{f \in 1-Lip} \left\{ h_1 \left[\int_{\mathbb{R}^d} f d\mu_1(t) - \int_{\mathbb{R}^d} f d\mu_1(s) \right] + h_2 \left[\int_{\mathbb{R}^d} f d\mu_2(t) - \int_{\mathbb{R}^d} f d\mu_2(s) \right] \right\} \\ &= h_1 d_1(\mu_1(t), \mu_1(s)) + h_2 d_1(\mu_2(t), \mu_2(s)) \\ &\leq \mathcal{C}_1 |t - s|^{\frac{1}{2}} (h_1 + h_2) = \mathcal{C}_1 |t - s|^{\frac{1}{2}} \end{aligned}$$

and the integral conditions similarly holds for $\nu(t)$.

Finally, to prove the compactness of $\mathcal{K}$ in $\mathcal{C}^0([0, T], \mathcal{P}_1)$ is a rather subtle fact, and we refer to [6]. ■

The core of our investigation, the “*mean field updating*” map $m = \Psi(\mu)$, can be defined from $\mathcal{K}$ to itself through the usual two steps process:

1. Firstly, we introduce u as the unique solution of class $\mathcal{C}^{2+\alpha}$ to the HJB backward equation with input μ :

$$\begin{cases} \partial_t u - \Delta u + \frac{1}{2} |\nabla u|^2 = F(x, \mu) & \text{in } \mathbb{R}^d \times (0, T) \\ u(x, T) = G(x, \mu(T)) & \text{in } \mathbb{R}^d \end{cases} \quad (3.16)$$

2. Then we define $m = \Psi(\mu)$ as the unique solution of class $\mathcal{C}^{2+\alpha}$ of the Fokker-Planck forward equation.

$$\begin{cases} \partial_t m - \Delta m - \operatorname{div}(m \nabla u) = 0 & \text{in } \mathbb{R}^d \times (0, T) \\ m(x, 0) = m_0(x) & \text{in } \mathbb{R}^d \end{cases} \quad (3.17)$$

This construction does not immediately imply that $\Psi(\mu) = m \in \mathcal{K}$ for some $\mathcal{C}_1 \in \mathbb{R}^+$ (recall that $\mathcal{K}$ definition features the fixed constant $\mathcal{C}_1$).

Most of our proof will deal with the following proposition, which provides good premises for the existence of a fixed point.

Proposition 14. It exists a real positive constant $\mathcal{C}_1$ such that Ψ is well defined and continuous from $\mathcal{K}$ to itself.

Proof. Let us prove that u exists and is unique in the class $\mathcal{C}^{2+\alpha}$ with the Hopf-Cole transform, given by $w = e^{\frac{u}{2}}$. Equation (3.16) transformed becomes:

$$\begin{cases} \partial_t w - \Delta w = wF(x, \mu) & \text{in } \mathbb{R}^d \times (0, T) \\ w(x, T) = e^{\frac{G(x, \mu(T))}{2}} & \text{in } \mathbb{R}^d \end{cases} \quad (3.18)$$

We observe here that the maps:

$$(x, t) \mapsto F(x, m(t)) \text{ and } x \mapsto e^{\frac{G(x, \mu(T))}{2}}$$

belong to $\mathcal{C}^{\frac{1}{2}}$, since $\mu \in \mathcal{K}$ and (3.15a) holds.

Therefore we can invoke some classical results on the heat equation (see theorem A.1 in Appendix) to conclude that it exists a unique u in the class $\mathcal{C}^{2+\frac{1}{2}}$.

Due to the assumptions (3.15a) on F and G some bounds for u can be derived using the comparison principle. In the following, $\mathcal{C}_0$ is the constant introduced in the assumptions of theorem 12, which is both an absolute bound and Lipschitz constant for F, G .

To derive an estimate, we observe that $v = (1 + t - T)\mathcal{C}_0$ verifies the equation:

$$\partial_t v - \Delta v = \mathcal{C}_0 \quad \text{where } v(x, T) = \mathcal{C}_0$$

Since $F(x, \mu(t)) - \frac{1}{2}|\nabla u(x, t)|^2 \leq \mathcal{C}_0$, it has to be:

$$u(x, t) \leq v(x, t) \leq (1 + T)\mathcal{C}_0 \quad \forall (x, t)$$

The same argument applied to the equation $-\partial_t u + \Delta u = F(x, \mu(t)) - \frac{1}{2}|\nabla u|^2$ verified by $-u$ bound leads to $|u| \leq (1 + T)\mathcal{C}_0$.

Under the aforementioned assumptions also ∇u can be estimated.

Let us consider $x, y \in \mathbb{R}^d$ and a fixed time $t \in [0, T]$. We denote with $^x\alpha$ the optimal control relative to the initial position x and introduce the following state laws on $\mathbb{R}^d$:

$$X_s^{(\alpha)} \sim \begin{cases} dX_s = \alpha_s ds + \sqrt{2} dW_s \\ X_t = x \end{cases} \quad Y_s^{(\beta)} \sim \begin{cases} dY_s = \beta_s ds + \sqrt{2} dW_s \\ Y_t = y \end{cases}$$

The HJB solution $u(x, t)$, as shown in section 3.1, is the value function of the agent optimal control problem. Hence it holds:

$$\begin{aligned} u(x, t) &= \inf_{\alpha \in \mathcal{A}} \mathbb{E} \left[\int_t^T \left[\frac{1}{2} |\alpha_s|^2 + F(X_s^{(\alpha)}, s) \right] ds + G(X_T^{(\alpha)}) \right] \\ &= \mathbb{E} \left[\int_t^T \left[\frac{1}{2} |^x\alpha_s|^2 + F(X_s^{(^x\alpha)}, s) \right] ds + G(X_T^{(^x\alpha)}) \right] \leq \mathbb{E} \left[\int_t^T \left[\frac{1}{2} |^y\alpha_s|^2 + F(X_s^{(^y\alpha)}, s) \right] ds + G(X_T^{(^y\alpha)}) \right] \end{aligned}$$

If we assign the same control $^y\alpha_s$ for both X_s, Y_s , their law verifies the same SDE with different initial conditions. Hence:

$$\begin{aligned}
u(x, t) - u(y, t) &\leq \\
&\leq \mathbb{E} \left[\int_t^T \left[\frac{1}{2} |^y\alpha_s|^2 - \frac{1}{2} |^y\alpha_s|^2 + F(X_s^{(y\alpha)}, s) - F(Y_s^{(y\alpha)}, s) \right] ds + G(X_T^{(y\alpha)}) - G(Y_T^{(y\alpha)}) \right] \\
&\leq \mathbb{E} \left[\int_t^T \mathcal{C}_0 |X_s^{(y\alpha)} - Y_s^{(y\alpha)}| ds + \mathcal{C}_0 |X_T^{(y\alpha)} - Y_T^{(y\alpha)}| \right] = \int_t^T \mathcal{C}_0 |x - y| ds + \mathcal{C}_0 |x - y| \\
&\leq (1 + T) \mathcal{C}_0 |x - y|
\end{aligned}$$

In the same way, assigning $^x\alpha_s$ we get:

$$\begin{aligned}
u(y, t) - u(x, t) &\leq \\
&\leq \mathbb{E} \left[\int_t^T \left[\frac{1}{2} |^x\alpha_s|^2 - \frac{1}{2} |^x\alpha_s|^2 + F(Y_s^{(x\alpha)}, s) - F(X_s^{(x\alpha)}, s) \right] ds + G(Y_T^{(x\alpha)}) - G(X_T^{(x\alpha)}) \right] \\
&\leq (1 + T) \mathcal{C}_0 |x - y|
\end{aligned}$$

Recollecting all our findings for u and ∇u , we have proved that:

$$\begin{aligned}
|u| &\leq (1 + T) \mathcal{C}_0 \\
|\nabla u| &\leq (1 + T) \mathcal{C}_0
\end{aligned} \tag{3.19}$$

Remark 8. Thanks to the uniform boundedness assumptions on F and G , these estimates do not depend on $\mathcal{C}_1$ (which indeed affects μ).

We can now move to the second step related to Ψ definition.

Our aim, like before, is to prove that equation (3.17) has a unique solution m .

We rewrite the Fokker-Planck equation in the form:

$$\begin{cases} \partial_t m - \Delta m - \nabla m \cdot \nabla u - m \Delta u = 0 & \text{in } \mathbb{R}^d \times (0, T) \\ m(x, 0) = m_0(x) & \text{in } \mathbb{R}^d \end{cases} \tag{3.20}$$

and notice that:

$$\mathbf{a}(x, t) = \nabla m, \quad b(x, t) = \Delta u$$

belong to $\mathcal{C}^{\frac{1}{2}}$, since $u \in \mathcal{C}^{2+\frac{1}{2}}$. Like before, the same heat equation general result can be applied (theorem A.1) to deduce the unique existence of a solution $m \in \mathcal{C}^{2+\frac{1}{2}}$.

Is this m in $\mathcal{K}$? To answer, we should find a value for $\mathcal{C}_1$ so that $\Psi(\mathcal{K}) \subset \mathcal{K}$.

From section 3.2 we recall that $m(t)$ is the law of X_t (a generic agent's state) and depends on ∇u through the (3.7) SDE.

Some estimates can be derived by using equation (3.19) to control ∇u .

Proposition 15. It exists a constant $l_0 = l_0(\mathcal{C}_0, T)$ such that:

- a) $d_1(m(t), m(s)) \leq l_0(1 + \mathcal{C}_0)|t - s|^{\frac{1}{2}} \quad \forall s, t \in [0, T]$
- b) $\int_{\mathbb{R}^d} |x|^2 dm(t)(x) \leq l_0 \left(\int_{\mathbb{R}^d} |x|^2 dm_0(x) + 1 + \mathcal{C}_0^2 \right)$

Proof.

- a) Recalling definition 5 let us consider for $s < t$ the pair (X_t, X_s) : its law γ belongs to $\Pi(m(t), m(s))$ and therefore it holds:

$$\begin{aligned}
 d_1(m(t), m(s)) &\leq \int_{\mathbb{R}^{2d}} |x - y| d\gamma(x, y) = \mathbb{E}[|X_t - X_s|] \\
 &\leq \mathbb{E} \left[\int_s^t |\nabla u((X_w, w))| dw \right] + \sqrt{2} \mathbb{E}[|B_t - B_s|] \\
 &\leq \|\nabla u\|_\infty (t - s) + \sqrt{2} \sqrt{t - s} \quad \text{since, for } Z \sim \mathcal{N}(0, \sigma^2), \mathbb{E}|Z| = \sigma \sqrt{\frac{2}{\pi}} \\
 &\leq \mathcal{C}_0(t - s) + \sqrt{2} \sqrt{t - s}
 \end{aligned}$$

- b) Coming to the integral estimate, it can be derived as follows:

$$\begin{aligned}
 \int_{\mathbb{R}^d} |x|^2 dm(t)(x) &= \mathbb{E}[|X_t|^2] \leq \mathbb{E} \left[\left| X_0 + \int_0^t |\nabla u((X_w, w))| dw + \sqrt{2} B_t \right|^2 \right] \\
 &\leq 3 \mathbb{E} \left[|X_0|^2 + \left| \int_0^t |\nabla u((X_w, w))| dw \right|^2 + 2|B_t|^2 \right] \\
 &\leq 3 \left(\int_{\mathbb{R}^d} |x|^2 dm_0(x) + t^2 \|\nabla u\|_\infty^2 + 2t \right) \\
 &\leq 3 \left(\int_{\mathbb{R}^d} |x|^2 dm_0(x) + T^2 \mathcal{C}_0^2 + 2T \right)
 \end{aligned}$$

It is easy to choose a suitable $l_0(\mathcal{C}_0, T)$ that complies with both estimates. ■

Returning to our question on m , we observe that we can choose:

$$\mathcal{C}_1 = \max \left\{ l_0(1 + \mathcal{C}_0), l_0 \left(\int_{\mathbb{R}^d} |x|^2 dm_0(x) + 1 + \mathcal{C}_0^2 \right) \right\}$$

and deduce that $m \in \mathcal{K}$ and $\Psi(\mu)$ is well defined.

Let us check that Ψ is continuous: if $\mu_n \rightarrow \mu$ in $(\mathcal{K}, \tilde{d}_1)$, we want to show that:

$$m_n = \Psi(\mu_n) \rightarrow \Psi(\mu) = m \quad \text{in } (\mathcal{K}, \tilde{d}_1)$$

We can denote with (u_n, m_n) the pair of solutions corresponding to the input μ_n and investigate the convergence of these solutions:

Definition 16 (Locally uniform convergence). A sequence of functions $f_n : E \supset D \rightarrow \mathbb{R}^d$, where D is a metric space, locally uniformly converges to f if for every element x in the domain there exists a neighbourhood V such that:

$$f_n|_{V \cap D} \rightarrow f|_{V \cap D} \quad \text{uniformly}$$

We can notice that

$$\begin{aligned} F(x, \mu_n(t)) &\rightarrow F(x, \mu(t)) \\ G(x, \mu_n(T)) &\rightarrow G(x, \mu(T)) \end{aligned} \quad \text{with locally uniform convergence}$$

since

$$\forall (x, t) \in \mathbb{R}^d \times [0, T] \quad |F(x, \mu_n(t)) - F(x, \mu(t))| \leq C_0 \tilde{d}_1(\mu_n, \mu)$$

Therefore, since all the terms of the HJB system tend to their limits with locally uniform convergence, also the value functions solutions do, due to the continuous dependence on parameters of the value functions (Prop. 4.1 of chap. 4 from [22]):

$$u_n \rightarrow u \quad \text{locally uniformly}$$

Also the gradient of u , which enters the defining equation for m (equation (3.17)), has good convergence properties:

Proposition 17. In the same framework, $\nabla u_n \rightarrow \nabla u$ with local uniform convergence.

Proof. Since ∇u_n is uniformly bounded as shown in (3.19), we introduce the *uniformly bounded* sequence:

$$f_n = \frac{1}{2} |\nabla u_n|^2 - F(x, \mu_n)$$

and we can apply an interior regularity result (theorem A.2) to the equation satisfied by (u_n) :

$$\partial_t u_n - \Delta u_n = f_n$$

thus finding that ∇u_n converges uniformly to ∇u . ■

We can move on to the last steps of our proof.

Let m_{n_k} be any subsequence of m_n ; belonging to a compact set in a metric space, it admits a converging sub-subsequence:

$$m_{n_k} \longrightarrow y \in \mathcal{K}$$

every m_{n_k} satisfies:

$$\int_{\mathbb{R}^d} \phi(x, 0) m_0 + \int_0^t \int_{\mathbb{R}^d} [\partial_t \phi(x, s) - \nabla \phi(x, s) \cdot \nabla u_{n_k}(x, s) + \Delta \phi(x, s)] m_{n_k}(s) dx ds = 0$$

and for the limit y it holds:

$$\begin{aligned} & \int_0^t \int_{\mathbb{R}^d} [\partial_t \phi(x, s) - \nabla \phi(x, s) \cdot \nabla u_{n_k}(x, s) + \Delta \phi(x, s)] (m_{n_k} - y)(s) dx ds = \\ & = \int_0^t \int_{\mathbb{R}^d} \Xi_{n_k}(x, s) (m_{n_k} - y)(s) dx ds \longrightarrow 0 \quad \nabla u_{n_k} \longrightarrow \nabla u \end{aligned}$$

because that Ξ_{n_k} are uniformly Lipschitz (ϕ has compact support) and we recall that $\sup_{t \in [0, T]} d_1(m_{n_k}(t), y(t)) \longrightarrow 0$. Therefore we proved that y verifies (3.17).

But there exists a unique probability measure solution m to the Fokker-Planck equation², hence $y = m$ and every subsequence m_{n_k} admits a converging sub-subsequence to the same limit m .

In a metric space this suffices to conclude that $m_n \rightarrow m$ and we have proved that $\Psi(\mu)$ is well defined and continuous. $\blacksquare$

Finally, we deduce the existence of a fixed point m by applying Schauder fixed point Theorem (A.4) to the map Ψ , which we now know to be continuous from the compact metric space $\mathcal{K} \subset \mathcal{C}^0([0, T], \mathcal{P}_1)$ to itself.

3.3.2 Uniqueness proof

Let us take any two solutions to problem (3.14), (u_1, m_1) and (u_2, m_2) and set: $\tilde{u} = u_1 - u_2$; $\tilde{m} = m_1 - m_2$.

From the Fokker-Planck equation we have, against any test function ϕ with compact support in $[0, T] \times \mathbb{R}^d$:

$$\begin{aligned} & \partial_t \tilde{m} - \Delta \tilde{m} - \operatorname{div} (m_1 \nabla u_1 - m_2 \nabla u_2) = 0 \\ & \left[\int_{\mathbb{R}^d} \tilde{m} \phi \right]_0^T - \int_0^T \int_{\mathbb{R}^d} (\partial_t \phi + \Delta \phi) \tilde{m} + \int_0^T \int_{\mathbb{R}^d} \nabla \phi \cdot (m_1 \nabla u_1 - m_2 \nabla u_2) = 0 \end{aligned} \quad (3.21)$$

²Recall section 3.2 and [3].

Remark 9 (A classical functional analysis result). We can use as test functions any map belonging to $\mathcal{C}^2$, since it can be approximated by a sequence ϕ_n of regularized truncates. Hence, we can use $\tilde{u}$ as test function ϕ .

The HJB equation instead yields, by multiplying by $\tilde{m}$ and integrating:

$$\begin{aligned} & \partial_t \tilde{u} + \Delta \tilde{u} - \frac{1}{2} (|\nabla u_1|^2 - |\nabla u_2|^2) + (F(x, m_1) - F(x, m_2)) = 0 \\ & \int_0^T \int_{\mathbb{R}^d} \left[\partial_t \tilde{u} + \Delta \tilde{u} - \frac{1}{2} (|\nabla u_1|^2 - |\nabla u_2|^2) + (F(x, m_1) - F(x, m_2)) \right] \tilde{m} = 0 \end{aligned} \quad (3.22)$$

Writing equation (3.21) with $\phi = \tilde{u}$ we find: (note that $\tilde{m}(0) = 0$)

$$\int_{\mathbb{R}^d} \tilde{m}(T) \tilde{u}(T) - \int_0^T \int_{\mathbb{R}^d} (\partial_t \tilde{u} + \Delta \tilde{u}) \tilde{m} + \int_0^T \int_{\mathbb{R}^d} \nabla \tilde{u} \cdot (m_1 \nabla u_1 - m_2 \nabla u_2) = 0 \quad (3.23)$$

Adding equations (3.22) and (3.23) we get:

$$\begin{aligned} & \int_{\mathbb{R}^d} \tilde{m}(T) [G(m_1(T)) - G(m_2(T))] + \int_0^T \int_{\mathbb{R}^d} \nabla \tilde{u} \cdot (m_1 \nabla u_1 - m_2 \nabla u_2) + \\ & - \frac{1}{2} \int_0^T \int_{\mathbb{R}^d} (|\nabla u_1|^2 - |\nabla u_2|^2) \tilde{m} + \int_0^T \int_{\mathbb{R}^d} (F(x, m_1) - F(x, m_2)) \tilde{m} = 0 \end{aligned}$$

We can notice that:

$$\begin{aligned} & \nabla \tilde{u} \cdot (m_1 \nabla u_1 - m_2 \nabla u_2) - \frac{1}{2} \tilde{m} (|\nabla u_1|^2 - |\nabla u_2|^2) = m_1 |\nabla u_1|^2 + m_2 |\nabla u_2|^2 + \\ & - (m_1 + m_2) \nabla u_1 \cdot \nabla u_2 - \frac{1}{2} (m_1 |\nabla u_1|^2 - m_1 |\nabla u_2|^2 + m_2 |\nabla u_2|^2 - m_2 |\nabla u_1|^2) = \\ & = \frac{1}{2} [m_1 |\nabla u_1|^2 + m_2 |\nabla u_2|^2 + m_1 |\nabla u_2|^2 + m_2 |\nabla u_1|^2] - (m_1 + m_2) \nabla u_1 \cdot \nabla u_2 = \\ & = \frac{1}{2} (m_1 + m_2) |\nabla u_1 - \nabla u_2|^2 = \frac{1}{2} (m_1 + m_2) |\nabla \tilde{u}|^2 \end{aligned}$$

So we end up with:

$$\begin{aligned} & \int_0^T \int_{\mathbb{R}^d} (F(x, m_1) - F(x, m_2)) \tilde{m} = \\ & = - \int_{\mathbb{R}^d} [G(m_1(T)) - G(m_2(T))] \tilde{m}(T) - \frac{1}{2} \int_0^T \int_{\mathbb{R}^d} (m_1 + m_2) |\nabla \tilde{u}|^2 \end{aligned}$$

Since the right side is either null or negative, it has to be:

$$\int_0^T \int_{\mathbb{R}^d} (F(x, m_1) - F(x, m_2)) d(m_1 - m_2) \leq 0$$

Due the assumptions 3.15b, we deduce $\tilde{m} = 0$ and therefore $u_1 = u_2$ for the uniqueness of the HJB solution.

3.4 MFG solution as an ε -Nash Equilibrium

In the previous sections we defined and built our mean field solution to the initial difficult problem (page 5).

We are now ready to present and prove a very interesting result, which in our opinion is a fair reward for the efforts made so far.

Let us consider the section 1.1 large N -player game, where costs are given by equation (1.2) (page 5) and do not include m (since N is finite). We suppose that the agents adopt the *mean field optimal* strategies:

$$\forall i \quad \tilde{\alpha}_t^i = -\nabla u(X_t^i, t) \implies dX_t^i = -\nabla u(X_t^i, t) dt + \sqrt{2} dB_t^i$$

We know these strategies to solve the infinite players limit game. Are they somehow close to the N -player game solution?

The following theorem clarifies in what sense the mean field solution turns out to be a “good approximation”, and its proof recollects many of the ideas we have developed so far.

Theorem 18 (ε -Nash equilibrium property).

For any $\varepsilon > 0$, there is some N_0 such that for every $N \geq N_0$ the symmetric strategy $(\tilde{\alpha}^1, \dots, \tilde{\alpha}^N)$ is an ε -Nash equilibrium in the game $(\mathcal{J}_1^N, \dots, \mathcal{J}_N^N)$:

$$\forall i \quad \mathcal{J}_i^N(\tilde{\alpha}^1, \dots, \tilde{\alpha}^N) \leq \mathcal{J}_i^N((\tilde{\alpha}^j)_{j \neq i}, \beta) + \varepsilon \quad (3.24)$$

for all controls $\beta \in \mathcal{A}$.

Proof. Let be $\varepsilon > 0$ fixed. Thanks to symmetry, it is enough to prove that for any $\beta \in \mathcal{A}$, definitely in N it holds:

$$\mathcal{J}_1^N(\tilde{\alpha}^1, \dots, \tilde{\alpha}^N) \leq \mathcal{J}_1^N((\tilde{\alpha}^j)_{j \geq 2}, \beta) + \varepsilon \quad (3.25)$$

For a given β , let us denote with X_t the state of agent 1:

$$X_t : \begin{cases} dX_t = \beta_t dt + \sqrt{2} dB_t \\ X_0 \sim m_0 \end{cases}$$

His cost is then provided by equation (1.2):

$$\begin{aligned} J_1^N((\tilde{\alpha}^j)_{j \geq 2}, \beta) &= \\ &= \mathbb{E} \left[\int_0^T \frac{1}{2} |\beta_s|^2 ds + \underbrace{\int_0^T F \left(X_s, \frac{1}{N-1} \sum_{j \geq 2} \delta_{X_s^j} \right) ds + G \left(X_T, \frac{1}{N-1} \sum_{j \geq 2} \delta_{X_T^j} \right)}_{\mathcal{W}_{\beta, N}} \right] \\ &= \mathbb{E} \left[\int_0^T \frac{1}{2} |\beta_s|^2 ds + \mathcal{W}_{\beta, N} \right] \end{aligned}$$

A first observation is that if β is such that (an “expensive control”):

$$\mathbb{E} \left[\int_0^T \frac{1}{2} |\beta_s|^2 ds \right] \geq 2(T\|F\|_\infty + \|G\|_\infty) + \frac{1}{2} \mathbb{E} \left[\int_0^T |\tilde{\alpha}_s^1|^2 ds \right] := R$$

Then we can immediately conclude, since:

$$J_1^N((\tilde{\alpha}^j)_{j \neq i}, \beta) \geq \mathbb{E} \left[\int_0^T \frac{1}{2} |\beta_s|^2 ds \right] \geq R \geq \mathcal{J}_1^N(\tilde{\alpha}^1, \dots, \tilde{\alpha}^N)$$

Hence, in the following we can assume that $\mathbb{E} \left[\int_0^T \frac{1}{2} |\beta_s|^2 ds \right] \leq R$.

It could be useful to find a bound of some kind for $\sup_{t \in [0, T]} |X_t|$; we proceed as follows:

$$\mathbb{E} \left[\sup_{t \in [0, T]} |X_t|^2 \right] \leq 3\mathbb{E} \left[|X_0|^2 + \frac{1}{2} \int_0^T |\beta_s|^2 ds + 2 \sup_{t \in [0, T]} |B_t|^2 \right] \leq \underbrace{3\mathbb{E} [|X_0|^2] + 3R + 24T}_{R_1}$$

where from Doob’s inequality we derived:

$$\mathbb{E} \left[\sup_{t \in [0, T]} |B_t|^2 \right] \leq 4\mathbb{E}[|B_T|^2] = 4T$$

Then, naming R_1 the final right side written above, Chebishev’s inequality yields:

$$\mathbb{P} \left[\sup_{t \in [0, T]} |X_t| \geq \frac{1}{\sqrt{\varepsilon}} \right] \leq R_1 \varepsilon \quad (3.26)$$

In the following (large) disc set we can prove that the mean field approximation works fine:

$$D = \left\{ \sup_{t \in [0, T]} |X_t| < \frac{1}{\sqrt{\varepsilon}} \right\}$$

Since the other agents follow the mean field strategy, their states $(X_t^j)_{j \geq 2}$ form an i.i.d. family of random variables³ with law $m(t)$.

Therefore, we can apply a Hewitt-Savage corollary (theorem A.6 and corollary A.7) to find an estimate for the cost difference between the finite N formulation and the mean field one:

$$\begin{aligned} & \forall t \in [0, T], \quad \exists N_0 : \forall N \geq N_0 \\ & \mathbb{E} \left[\sup_{|y| \leq 1/\sqrt{\varepsilon}} \left| F \left(y, \frac{1}{N-1} \sum_{j \geq 2} \delta_{X_t^j} \right) - F(y, m(t)) \right| \right] \leq \varepsilon \\ & \mathbb{E} \left[\sup_{|y| \leq 1/\sqrt{\varepsilon}} \left| G \left(y, \frac{1}{N-1} \sum_{j \geq 2} \delta_{X_T^j} \right) - G(y, m(T)) \right| \right] \leq \varepsilon \end{aligned} \quad (3.27)$$

³See section 3.2, page 32

In the first estimate, we can also choose N_0 independently of t . In fact, thanks to the Lipschitz conditions (3.15a), for any (s, t) it holds:

$$\begin{aligned} & \mathbb{E} \left[\sup_{|y| \leq 1/\sqrt{\varepsilon}} \left| F \left(y, \frac{1}{N-1} \sum_{j \geq 2} \delta_{X_t^j} \right) - F \left(y, \frac{1}{N-1} \sum_{j \geq 2} \delta_{X_s^j} \right) \right| \right] \\ & \leq \mathbb{E} \left[\mathcal{C}_0 d_1 \left(\frac{1}{N-1} \sum_{j \geq 2} \delta_{X_t^j}, \frac{1}{N-1} \sum_{j \geq 2} \delta_{X_s^j} \right) \right] \end{aligned}$$

and recalling definition 5 and proposition 15:

$$\begin{aligned} \mathcal{C}_0 \mathbb{E} \left[d_1 \left(\frac{1}{N-1} \sum_{j \geq 2} \delta_{X_t^j}, \frac{1}{N-1} \sum_{j \geq 2} \delta_{X_s^j} \right) \right] & \leq \mathcal{C}_0 \frac{1}{N-1} \sum_{j \geq 2} \mathbb{E} [|X_t^j - X_s^j|] \\ & \leq \mathcal{C}_0 l_0 (1 + \|\nabla u\|_\infty) (t - s)^{\frac{1}{2}} \end{aligned}$$

Since $[0, T]$ is compact, this last estimate allows to choose N_0 uniformly in t .

In fact, we can adopt it for both $F \left(y, \frac{1}{N-1} \sum_{j \geq 2} \delta_{X_t^j} \right)$ and $F(y, m(t))$ and observe that, for every $\varepsilon > 0$, condition (3.27) uniformly holds in every small time interval $[t - \delta_\varepsilon, t + \delta_\varepsilon]$. For example, we can set:

$$\delta_\varepsilon < \frac{\varepsilon^2}{18\mathcal{C}_0 l_0 (1 + \|\nabla u\|_\infty)}$$

Since it exists a finite cover of $[0, T]$ made of such intervals, it suffices to define N_0 as the maximum in this cover to get a globally uniform bound.

We can now replace the $\mathcal{W}_{\beta, N}$ term with the mean field version, since on the set D this substitution is controlled by the above estimates, and D^c has an arbitrarily small probability.

For any $N \geq N_0$ we find:

$$\begin{aligned} J_1^N ((\tilde{\alpha}^j)_{j \geq 2}, \beta) &= \mathbb{E} \left[\int_0^T \frac{1}{2} |\beta_s|^2 ds \right] + \mathbb{E} [\mathcal{W}_{\beta, N} | D] \mathbb{P}(D) + \mathbb{E} [\mathcal{W}_{\beta, N} | D^c] \mathbb{P}(D^c) \geq \\ &\geq \mathbb{E} \left[\int_0^T \left[\frac{1}{2} |\beta_s|^2 + F(X_s, m(s)) - \varepsilon \right] ds + G(X_T, m(T)) - \varepsilon \right] - \mathbb{P}(D^c) 2(T\|F\|_\infty + \|G\|_\infty) \\ &\geq \mathbb{E} \left[\int_0^T \left[\frac{1}{2} |\beta_s|^2 + F(X_s, m(s)) \right] ds + G(X_T, m(T)) \right] - \varepsilon(1 + T) - \varepsilon R_1 2(T\|F\|_\infty + \|G\|_\infty) \\ &\geq J_1^N ((\tilde{\alpha}^j)_{j=1, \dots, N}) - L\varepsilon \quad \text{since } \tilde{\alpha}^1 \text{ is optimal} \end{aligned}$$

for some constant L independent from N, β . The inequality (3.24) is proved. $\blacksquare$

Chapter 4

A model of imitative behaviours in a population

The formulation discussed in chapter 3 can be easily employed to study the evolution of a large population driven by imitative behaviours.

We consider a great number of agents having the same preferences; their individual characteristics are resumed by a finite dimensional vector $X \in \mathbb{R}^d$.

Every dimension represents a particular social aspect of the individual, such as geographic position, wealth, education or fitness, and we assume that everyone tries to resemble the others over time in order to feel more comfortable.

In order to do so, an agent can apply a drift α to his state, paying an increasing cost for larger efforts, while at the same time his condition naturally evolves following a random path.

Assuming to know the initial distribution m_0 , can we predict the overall behaviour of the agents?

4.1 Model set-up

As in the previous chapter, we consider $(\Omega, \mathcal{F}, \mathbb{P})$ and a filtration $(\mathcal{F}_t)_{t \in [0, T]} \subset \mathcal{F}$ such that an infinite family $(B_t^i)_{i \in \mathbb{N}}$ of independent brownian motions adapted to $(\mathcal{F}_t)$ is defined on it.

Our first abstraction is to model the individual evolution with the following stochastic differential equation:

$$dX_t^i = \alpha_t dt + \sigma dB_t^i \quad \text{in } \mathbb{R}^d$$

where the additive noise B_t^i represents the natural state evolution that every agent experiences.

We then choose the following payoff function to account for imitative preferences:

$$\mathcal{J} = \mathbb{E} \left[\int_0^T \left(\ln(m(X_t^i, t)) - \frac{1}{2} |\alpha_t|^2 \right) e^{-\varrho t} dt \right] \quad (4.1)$$

where $m(x, t)$ is the agents' distribution density in the "social space" $\mathbb{R}^d$ and $\varrho > 0$ acts as a discount factor.

4.2 PDE formulation

Let us derive the HJB equation for this model. To begin, we introduce the value function $u(x, t)$:

$$\begin{aligned} u(x, t) &= \sup_{\alpha_t \in \mathcal{A}} \mathbb{E} \left[\int_t^T \left(\ln(m(X_s, s)) - \frac{1}{2} |\alpha_s|^2 \right) e^{-\varrho(s-t)} ds \right] \\ &= \sup_{\alpha_t \in \mathcal{A}} e^{\varrho t} \mathbb{E} \left[\int_t^T \left(\ln(m(X_s, s)) - \frac{1}{2} |\alpha_s|^2 \right) e^{-\varrho s} ds \right] \end{aligned}$$

where:

$$\mathcal{A} = \left\{ \alpha_t \in L^2 \left([0, T] \times \Omega, \mathbb{R}^d \right); \alpha_t \text{ progressively measurable} \right\}$$

As in section 3.1, we can assume that $u \in \mathcal{C}^{2,1}$ and write the Bellman principle for an intermediate time q :

$$u(x, t) = \sup_{\alpha_t \in \mathcal{A}_t^q} e^{\varrho t} \mathbb{E} \left[\int_t^q \left(\ln(m(X_s, s)) - \frac{1}{2} |\alpha_s|^2 \right) e^{-\varrho s} ds + e^{-\varrho q} u(X_q, q) \right]$$

Using Ito's formula to express $\mathbb{E}[e^{-\varrho(q-t)} u(X_q, q) - u(x, t)]$ in integral form we find:

$$\begin{aligned} \mathbb{E} \left[e^{-\varrho(q-t)} u(X_q, q) - u(x, t) \right] &= \int_t^q \frac{d}{ds} \mathbb{E} \left[e^{-\varrho(s-t)} u(X_s, s) \right] ds \\ &= e^{\varrho t} \int_t^q \mathbb{E} \left[\nabla u(X_s, s) \cdot \alpha_s - \varrho u(X_s, s) + \partial_t u(X_s, s) + \frac{\sigma^2}{2} \Delta u(X_s, s) \right] e^{-\varrho s} ds \end{aligned}$$

Substituting in the Bellman principle we get:

$$\begin{aligned} \sup_{\alpha_t \in \mathcal{A}_t^q} e^{\varrho t} \mathbb{E} \left[\int_t^q \left(\ln(m(X_s, s)) - \frac{1}{2} |\alpha_s|^2 + \frac{\sigma^2}{2} \Delta u(X_s, s) + \right. \right. \\ \left. \left. + \nabla u(X_s, s) \cdot \alpha_s - \varrho u(X_s, s) + \partial_t u(X_s, s) \right) e^{-\varrho s} ds \right] = 0 \end{aligned}$$

We can now let $q \rightarrow t$ and assume to deal with a continuous control, obtaining:

$$\sup_{\alpha \in \mathbb{R}^d} \left(\ln(m(x, t)) - \frac{1}{2}|\alpha|^2 + \frac{\sigma^2}{2}\Delta u(x, t) + \nabla u(x, t) \cdot \alpha - \varrho u(x, t) + \partial_t u(x, t) \right) = 0$$

$$\partial_t u + \frac{\sigma^2}{2}\Delta u + \ln(m(x, t)) - \varrho u + \sup_{\alpha \in \mathbb{R}^d} \left(\nabla u \cdot \alpha - \frac{1}{2}|\alpha|^2 \right) = 0$$

The *supremum* is actually realized by:

$$\max_{\alpha \in \mathbb{R}^d} \left(\nabla u \cdot \alpha - \frac{1}{2}|\alpha|^2 \right) = \frac{1}{2}|\nabla u|^2 \quad \text{for } \alpha = \nabla u$$

The final form for the Hamilton-Jacobi-Bellman equation is then:

$$\begin{cases} \partial_t u + \frac{\sigma^2}{2}\Delta u + \ln(m(x, t)) - \varrho u + \frac{1}{2}|\nabla u|^2 = 0 \\ u(x, T) = 0 \end{cases} \quad (\text{a conventional value}) \quad (4.2)$$

Remark 10. Unfortunately, there is no immediate way to apply theorem A.1 to the above PDE and deduce the existence of a unique classical solution.

Therefore, until we prove the existence of a classical u and the related verification theorem¹, we can not claim that a solution to this HJB represents an optimal control strategy.

Nevertheless, if we find a classical solution to the full MFG system, then we could a posteriori investigate the equivalence of our HJB equation to the original stochastic control problem. For comparison, the derivation made in section 3.1 has been allowed by the strong hypotheses on F, G , which provided some a priori estimates for the solution.

Let us move on to the Fokker-Planck equation to build the full MFG system.

The state equation with the feedback candidate optimal control $\nabla u(X_t, t)$ is:

$$\begin{cases} dX_t = \nabla u(X_t, t) dt + \sigma dB_t, & t \in [0, T] \\ X_0 \sim m_0 \end{cases} \quad (4.3)$$

If m_0 has density $m_0(x)$ on $\mathbb{R}^d$ and it exists a solution to (4.3), then its law verifies:

$$\begin{cases} \partial_t m - \frac{\sigma^2}{2}\Delta m + \operatorname{div}(m \nabla u(x, t)) = 0 & \text{in } \mathbb{R}^d \times (0, T) \\ m(x, 0) = m_0(x) \end{cases} \quad (4.4)$$

Like before, since we do not a priori know that u has a good regularity, we do not have sufficient conditions on ∇u to claim that the closed loop equation (4.3) has a unique

¹These are sufficient but not necessary conditions for the existence of an optimal control.

solution², so that the Fokker-Planck we derived describes the density of X_t . Again, we look forward to solve the full MFG system and to verify a posteriori that our equations are consistent with the initial problem.

Gathering the above results, we finally obtain this PDE system:

$$\begin{cases} \partial_t u + \frac{\sigma^2}{2} \Delta u + \frac{1}{2} |\nabla u|^2 - \varrho u = -\ln(m(x, t)) & \text{in } \mathbb{R}^d \times [0, T] \\ \partial_t m + \operatorname{div}(m \nabla u) = \frac{\sigma^2}{2} \Delta m & \text{in } \mathbb{R}^d \times [0, T] \\ m(x, 0) = m_0 \\ u(x, T) = 0 \end{cases} \quad (4.5)$$

We can recognize the backward/forward structure typical of the mean field approach: given a dynamic for m , every agent optimizes his movements adopting his best possible strategy (HJB backward equation). This consequently drives the aggregate evolution of m itself (FP forward equation). If a unique solution exists, it is consistent both with optimality and agent diffusion in the state space.

Then, the rationality assumption implies that this optimal solution will be effectively adopted by everyone, so that the predicted evolution of m will be actually realized by the agents' behaviour.

Note that the mean field existence and uniqueness result of the previous chapter can not be applied to the above formulation, since the HJB has unbounded coefficients and theorem A.1 does not apply.

Nevertheless, we can still try to find an explicit solution.

4.3 A stationary solution

If we renounce to solve the system for every assigned initial condition m_0 , we can turn things easier by looking for a stationary solution.

- By assuming a *stationary* distribution m , we immediately have a simplified diffu-

²This would be important also for the HJB consistency with the optimal control problem, since it would show that the feedback candidate is admissible.

sion equation:

$$\begin{aligned}\operatorname{div}(m \nabla u) &= \frac{\sigma^2}{2} \Delta m \\ \operatorname{div} \left(m \nabla u - \frac{\sigma^2}{2} \nabla m \right) &= 0\end{aligned}$$

Which is verified if:

$$m \nabla u = \frac{\sigma^2}{2} \nabla m$$

therefore a suitable form for m is:

$$m(x) = K e^{\frac{2}{\sigma^2} u(x)} \quad \text{for } K \text{ such that } m \text{ is a probability measure}$$

- Under the previous assignment and the stationarity requirement, the first HJB equation becomes:

$$\begin{aligned}\frac{\sigma^2}{2} \Delta u + \frac{1}{2} |\nabla u|^2 - \varrho u &= -\ln(K) - \frac{2}{\sigma^2} u \\ \frac{\sigma^2}{2} \Delta u + \frac{1}{2} |\nabla u|^2 + \left(\frac{2}{\sigma^2} - \varrho \right) u &= -\ln(K)\end{aligned}$$

Here we assume that the last coefficient has positive sign, i.e. $\varrho < \frac{2}{\sigma^2}$, and we look for a solution of the form:

$$u(x) = -\eta |x - \mu|^2 + \omega$$

The equation becomes³:

$$\begin{aligned}-\sigma^2 \eta d + 2\eta^2 |x - \mu|^2 + \varrho \eta |x - \mu|^2 + \left(\frac{2}{\sigma^2} - \varrho \right) \omega - \frac{2\eta}{\sigma^2} |x - \mu|^2 &= -\ln(K) \\ \left(2\eta^2 + \varrho \eta - \frac{2\eta}{\sigma^2} \right) |x - \mu|^2 + \sigma^2 \eta d - \left(\frac{2}{\sigma^2} - \varrho \right) \omega - \ln(K) &= 0\end{aligned}$$

- This yields the following relations for our parameters:

$$\begin{aligned}2\eta^2 + \varrho \eta - \frac{2\eta}{\sigma^2} &= 0 \implies \eta = \frac{1}{\sigma^2} - \frac{\varrho}{2} \\ \left\{ \begin{aligned} \sigma^2 \eta d - \left(\frac{2}{\sigma^2} - \varrho \right) \omega - \ln(K) &= 0 \\ K e^{\frac{2\omega}{\sigma^2}} \int_{\mathbb{R}^d} e^{-\frac{2\eta}{\sigma^2} |x - \mu|^2} &= K e^{\frac{2\omega}{\sigma^2}} \left(\frac{\pi \sigma^2}{2\eta} \right)^{\frac{d}{2}} = 1 \end{aligned} \right. \implies \ln(K) = -\frac{2\omega}{\sigma^2} - \frac{d}{2} \ln \left(\frac{\pi \sigma^2}{2\eta} \right)\end{aligned}$$

From the last equation we deduce:

$$\left(1 - \frac{\sigma^2}{2} \varrho \right) d + \varrho \omega = \frac{d}{2} \ln \left(\frac{2\eta}{\pi \sigma^2} \right)$$

and therefore ω can be determined.

³The Laplacian Δu gives: $\Delta u = \operatorname{div}(\nabla u) = -2\eta \operatorname{div}(x - \mu) = -2\eta d$

- In conclusion, we have found the following Gaussian distribution for m :

$$m(x) = K e^{-\frac{2\eta}{\sigma^2}|x-\mu|^2} \sim \mathcal{N}(\mu, s^2 I_d) \quad \text{where } s^2 = \frac{\sigma^4}{4 - 2\rho\sigma^2} \quad (4.6)$$

This solution shows that uniqueness does not hold in this framework, since our problem is invariant by translation (given by μ in the formula above). Even if we fix an origin to avoid this issue, uniqueness can not be granted (see [16], cap. 5).

Instead, as a positive result, we observe that ∇u is a linear function and therefore the equation (4.3) has a unique strong solution. The optimal control $\alpha_t = \nabla u(X_t, t)$ is then admissible and a verification theorem very similar to theorem 8 can be proved.

This model shows how easily the application of the MFG theory to real-like scenarios can lead to mathematically challenging problems, even when elementary payoff functions are involved.

Chapter 5

An individual-based human capital growth model

In this chapter we adopt the mean field game perspective to build an economic human capital growth model based on individual choices.

Our exposition is inspired by [16, 15], but differs in the adopted model structure since it consider independent noises for every agent. See section 5.3 for details.

Our ambition is to model the individual wealth growth in a stylized economy with a large numbers of interacting agents. We assume that their salaries are determined by the human capital distribution across the population and we denote with X the positive real amount of human capital every agents is endowed with.

A higher level of human capital allows for an increase in income due to the combination of two effects: the added value of being more competent and the decrease in the number of people that have a similar skill level.

Of course, enhancing their own human capital has a cost for individuals, and this cost sharply increases with their current level of competence, so that reaching the technological frontier becomes more and more expensive.

Each agents tries to maximize its lifetime wealth, knowing that the initial distribution of agents on the human capital line (and therefore income) is not uniform.

Let us now analyse the details of the above brief description with the elements introduced in chapter 3.

5.1 Model set-up

As always, we consider a probability space $(\Omega, \mathcal{F}, \mathbb{P})$ and a filtration $(\mathcal{F}_t)_{t \in [0, T]} \subset \mathcal{F}$ such that an infinite family $(B_t^i)_{i \in \mathbb{N}}$ of independent brownian motions adapted to $(\mathcal{F}_t)$ is defined on it.

For this model, the state space where agents' conditions evolve is the real positive half line, $X \in (0, +\infty)$.

We name $m(x, t)$ the agent's density over the possible levels of human capital and M its cumulative distribution, while as always $m(x, 0)$ is the initial density value.

We want to model the knowledge tendency to increase and not to assume negative values, so an additive Brownian noise does not fit in this particular scenario.

A possible choice for the (stochastic) human capital accumulation process of an agent i adopting a_t is given by:

$$dX_t^i = a_t dt + \sigma X_t dB_t^i \quad (5.1)$$

where we assume that $a_t \geq 0$.

Remark 11. This state equation is significantly different from equation (3.7), since the volatility term *varies* with the state X_t .

We can now look forward to formulate the agent's optimization problem. Let us define the income w as function of human capital level x :

$$w(x, m(x, t)) = \begin{cases} \mathcal{C} \frac{x^\alpha}{m(x, t)^\beta} & \text{if } x \text{ is in the support of } m(\cdot, t) \\ 0 & \text{otherwise} \end{cases} \quad (5.2)$$

where we introduced the real positive α and β such that $\alpha + \beta > 1$ and the positive constant $\mathcal{C}$.

Remark 12. Note that the above choice well models the opposite effects on salary of human capital level and abundance of similarly skilled workers.

Coming to the cost of increasing x at rate a , we set:

$$H(a, M(x, t)) = S \frac{a^{\alpha+\beta}}{(\alpha + \beta)[1 - M(x, t)]^\beta} \quad (5.3)$$

where S is a positive constant measuring inefficiency in human capital production and the exponents (α, β) are the same used above, so that the cost of human capital increasing is a convex function of the desired speed.

This choice for H keeps track of the dependence on both the increase speed and the proximity to the technological limit.

The next step is to specify an initial distribution m_0 of human capital (and therefore income), which we assume to have density $m(x, 0)$ on $\mathbb{R}^+$.

This choice should reflect the real income distribution in the population, therefore we adopt a Pareto law with coefficient k and conventionally set its starting point to 1.

$$m(x, 0) = k \frac{1}{x^{k+1}} \mathbb{1}_{x \geq 1}$$

In fact, this type of distribution can be widely found in income inequality studies, such as papers from T. Piketty or A.B. Atkinson. The coefficient $k \geq 1$ is an inverse indicator of inequality; the corresponding Gini index is given by $G = 1/(2k - 1)$.

The final step of our model set-up is to declare the agent payoff and formulate the related optimization problem:

$$\begin{aligned} \sup_{a_t \in \mathcal{A}} \mathcal{J}(a) = \sup_{a_t \in \mathcal{A}} \mathbb{E} \left[\int_0^T \left(c \frac{X_t^\alpha}{m(X_t, t)^\beta} - S \frac{a_t^{\alpha+\beta}}{(\alpha + \beta)[1 - M(X_t, t)]^\beta} \right) e^{-\varrho t} dt \right] \quad (5.4) \\ \text{being } \begin{cases} dX_t = a_t dt + \sigma X_t dB_t & \text{for } s \in [0, T] \\ X_0 \sim m(x, 0) \end{cases} \end{aligned}$$

where the set of admissible controls is:

$$A = \{a_t \in L^2([0, T] \times \Omega, \mathbb{R}^+); a_t \text{ progressively measurable}\}$$

and $\varrho > 0$ provides a discount factor.

5.2 PDE formulation

In order to express the differential formulation of the mean field problem, let us set for simplicity:

$$\alpha = \beta = 1$$

so that the payoff equation (5.4) becomes:

$$\mathcal{J}(a) = \mathbb{E} \left[\int_0^T \left(c \frac{X_t}{m(X_t, t)} - S \frac{a_t^2}{2[1 - M(X_t, t)]} \right) e^{-\varrho t} dt \right]$$

Let us introduce the value function u :

$$\begin{aligned} u(x, t) &= \sup_{a_t \in \mathcal{A}} \mathbb{E} \left[\int_t^T \left(c \frac{X_s}{m(X_s, s)} - S \frac{a_s^2}{2[1 - M(X_s, s)]} \right) e^{-\varrho(s-t)} ds \right] \\ &= \sup_{a_t \in \mathcal{A}} e^{\varrho t} \mathbb{E} \left[\int_t^T \left(c \frac{X_s}{m(X_s, s)} - S \frac{a_s^2}{2[1 - M(X_s, s)]} \right) e^{-\varrho s} ds \right] \end{aligned}$$

The corresponding HJB equation can be derived through the same heuristic steps adopted in sections 3.1 and 4.2, resulting in:

$$\partial_t u + \frac{\sigma^2}{2} x^2 \partial_{xx}^2 u - \varrho u + \sup_{a \in \mathbb{R}^+} \left(a \partial_x u - S \frac{a^2}{2[1 - M(x, t)]} \right) + \mathcal{C} \frac{x}{m(x, t)} = 0$$

Being, for every $a \in \mathbb{R}^+$,

$$\begin{aligned} \max_{a \in \mathbb{R}^+} \left(a \partial_x u - S \frac{a^2}{2[1 - M(x, t)]} \right) &= \frac{1 - M(x, t)}{2S} (\partial_x u)^2 \\ \text{for } a &= \frac{1 - M(x, t)}{S} \partial_x u \end{aligned}$$

the final form for the HJB equation is:

$$\begin{cases} \partial_t u + \frac{\sigma^2}{2} x^2 \partial_{xx}^2 u + \frac{1 - M(x, t)}{2S} (\partial_x u)^2 - \varrho u = -\mathcal{C} \frac{x}{m(x, t)} \\ u(x, T) = 0 \end{cases} \quad (5.5)$$

The argument of the maximum suggests this feedback form for the optimal control:

$$\tilde{a}_t = \tilde{a}(X_t, t) = \frac{1 - M(X_t, t)}{S} \partial_x u(X_t, t)$$

but we have no immediate argument to prove that $\tilde{a}_t$ is admissible. In fact, the closed loop equation:

$$\begin{cases} dX_s = \frac{1 - M(X_s, s)}{S} \partial_x u(X_s, s) dt + \sigma X_s dB_s & s \in [t, T] \\ X_t = x \end{cases} \quad (5.6)$$

needs not to admit a solution, since we lack of a priori information about u . Like in section 4.2, we can not prove a verification theorem for the HJB equation at this stage.

Coming to the Fokker-Planck that describes the aggregate density $m(x, t)$, we can apply the classical PDE-SDE relationship to equation (5.6) and get:

$$\begin{cases} \partial_t m = -\partial_x (\tilde{a}_t m) + \frac{\sigma^2}{2} \partial_{xx}^2 (x^2 m) \\ m(x, 0) = k \frac{1}{x^{k+1}} \mathbb{1}_{x \geq 1} \end{cases} \quad (5.7)$$

substituting for $\tilde{a}_t$ and expanding we obtain:

$$\begin{aligned} \partial_t m + \partial_x \left(m \frac{1 - M}{S} \partial_x u \right) - \frac{\sigma^2}{2} \partial_{xx}^2 (x^2 m) &= 0 \\ \partial_t m + \frac{1 - M}{S} (\partial_x m \partial_x u + m \partial_{xx}^2 u) - \frac{\sigma^2}{2} (x^2 \partial_{xx}^2 m + 4x \partial_x m + 2m) &= 0 \\ \partial_t m + \left(\frac{1 - M}{S} \partial_{xx}^2 u - \sigma^2 \right) m + \left(\frac{1 - M}{S} \partial_x u - 2\sigma^2 x \right) \partial_x m - \frac{\sigma^2 x^2}{2} \partial_{xx}^2 m &= 0 \end{aligned}$$

Collecting all the results, the final coupled PDE problem for our model is:

$$\begin{cases} \partial_t u + \frac{\sigma^2}{2} x^2 \partial_{xx}^2 u + \frac{1-M}{2S} (\partial_x u)^2 - \varrho u = -\mathcal{C} \frac{x}{m} \\ \partial_t m + \left(\frac{1-M}{S} \partial_{xx}^2 u - \sigma^2 \right) m + \left(\frac{1-M}{S} \partial_x u - 2\sigma^2 x \right) \partial_x m - \frac{\sigma^2 x^2}{2} \partial_{xx}^2 m = 0 \\ m(x, 0) = k \frac{1}{x^{k+1}} \mathbb{1}_{x \geq 1} \\ u(x, T) = 0 \end{cases} \quad (5.8)$$

These non-linear, fully coupled backward-forward PDE equations are the core of the mean field game problem.

These kinds of systems lie on the mathematical research frontier: for problem (5.8) we do not dispose of any existence theorem to apply, nor an explicit solution can be easily written.

Nevertheless, we find it very remarkable to have transformed our introductory model purpose into problem (5.8) thanks to the mean field approach versatility.

In doing this, we experienced the dual nature of mean field games theory: on one side, it constitutes a very interesting challenge that lies on the frontier of mathematics; on the other, it has the potential to become a powerful empirical tool once taken to reality.

5.3 Comparison with common noise models

The stochastic mean field model that is presented in [16, 15] uses a particular simplifying assumption to introduce randomness and hence admits an explicit solution.

Instead of considering N independent noises dB_t^i , a common Brownian motion dW_t acts on the state equation of all agents, describing evolving external conditions that equally affects the agents' states.

This leads the mean field m to depend on the particular realization of W_t and not to be deterministic; on the other hand, along a fixed trajectory $W_t(\omega)$ the function u results from a deterministic optimization and its HJB equation has no second order term.

We chose not to present this version of the model for it does not feature the aggregate uncertainty smoothing mechanism that we analysed in the previous chapters (see figure 2.5, page 22).

Some authors are currently performing research on general settings that feature both the idiosyncratic and the common noise terms, which is indeed a complicate scenario because it simultaneously deals with a non deterministic mean field and a stochastic optimization for the agents' strategies. An interesting preprint on the subject is [7].

Chapter 6

Prospects and conclusion

We are now concluding our travel through mean field games.

After the initial outline, we presented the core intuitions of the mean field approach through a thoroughly analysed toy model, discussing the role of rationality, symmetry, and the fixed-point reasoning keystone.

We then moved on to visit the theory frontier, proved an existence and uniqueness result and seized the ε -Nash recognition for our efforts.

Towards the end we deployed our tools in some economic flavoured scenarios, discovering their surprising versatility and emphasising the new open problems they pose.

Our travel has been driven by a strong belief in the future relevance of mean field games, both as a research topic in mathematics and empirical tool.

We indeed expect economists, statisticians, engineers and medical researchers to enjoy great benefits from mean field games development, since many difficult coupled PDE problems directly arise in these applications.

To conclude, we are glad to back these expectations by sharing some current prospects for mean field games.

- **Large heterogeneous agents models in Economics**

Lasry and Lions are now taking part to a promising research group at Princeton Economics. Thanks to mean field games theory, a broad category of heterogeneous agents macroeconomics models can now be investigated, both analytically and numerically. A draft paper that describes these frontier applications is [1].

Similar methods are being studied also at the European Central Bank: a working paper is [19].

- **Testing rationality**

Being an efficient way to describe an approximate Nash equilibrium, mean field games could provide a “benchmark aggregate behaviour” for many real interactions. By developing adequate statistical tools, it could be investigated how real (maybe irrational) aggregate outcome differs from the rational theoretical one, thus measuring the likelihood of the rationality assumption at the aggregate level.

- **Financial networks**

Themes such as interbank contagion, financial connectedness and systemic risk are widely studied after the recent world financial crisis. Some authors are currently adopting the mean field game approach to model the banks’ network: a preprint is [8].

- **Cancer evolution models in Medicine**

Mathematical oncology has greatly expanded in the last decade, yet many of its initial objectives have been redefined during its evolution.

Current research attempts to describe the evolution mechanisms that characterize the various phases of a tumour. Such a model could be employed by diagnosticians to determine the phase a tumour is in given a few time-lapsed observations.

The mean field perspective can effectively help to deal with the complex microscopic-macroscopic interactions that occur in tumour evolution, leading to macroscopic PDE-ODE models. Related literature is flourishing; a representative example is [13]. In this framework, the game theory modelling of cell behaviour constitutes a novelty (see [5]) that could pave the way for future developments.

Appendix

This section is mainly addressed to recollect some important Functional Analysis results. General references are [4, 14, 10].

Theorem A.1 (Heat equation existence and uniqueness). *This theorem is taken from [14], and we recall definition 7.*

Let us consider the following heat equation:

$$\begin{cases} \partial_t w - \Delta w - \mathbf{a}(x, t) \cdot \nabla w - b(x, t)w = f(x, t) & \text{in } \mathbb{R}^d \times (0, T) \\ w(x, 0) = g(x) & \text{in } \mathbb{R}^d \end{cases} \quad (\text{A.1})$$

If, for some $\alpha \in (0, 1)$, $a_i \forall i$, b , f , g belong to $\mathcal{C}^\alpha$ and are bounded, then the equation above has a unique classical solution $w \in \mathcal{C}^{2+\alpha}$.

Remark A.1. Note that this result holds also for the following backward formulation, where v_T is known in place of v_0 :

$$\begin{cases} \partial_t v + \Delta v + \mathbf{a}(x, t) \cdot \nabla v + b(x, t)v = -f(x, t) & \text{in } \mathbb{R}^d \times (0, T) \\ v(x, T) = g(x) & \text{in } \mathbb{R}^d \end{cases} \quad (\text{A.2})$$

In fact, it suffices to set $v(x, t) = w(x, T - t)$ and apply the above theorem.

Theorem A.2 (Interior regularity for a class of PDE solutions). *Let us consider the following form of heat equation:*

$$\begin{cases} \partial_t w - \Delta w = f(x, t) & \text{in } \mathbb{R}^d \times (0, T) \\ w(x, 0) = w_0(x) & \text{in } \mathbb{R}^d \end{cases} \quad (\text{A.3})$$

If we require f to be continuous and bounded, any classical, bounded solution of this equation verifies, for any compact set $K \subset \mathbb{R}^d \times (0, T)$:

$$\sup_{(x,t),(y,s) \in K} \frac{|\nabla w(x, t) - \nabla w(y, s)|}{|x - y|^\beta + |t - s|^{\beta/2}} \leq L(K, \|w\|_\infty) \|f\|_\infty \quad (\text{A.4})$$

where $\beta \in (0, 1)$ depends on d and L depends more broadly on K , $\|w\|_\infty$ and d .

Theorem A.3 (A compactness criterion for sets in $\mathcal{P}_1$). *Let $r < 1$ and $C \subset \mathcal{P}_1$ be such that:*

$$\sup_{\mu \in C} \int_{\mathbb{R}^d} |x|^r d\mu(x) < +\infty$$

Then the set C is relatively compact with respect to the d_1 distance.

Theorem A.4 (Schauder fixed point). *Let E be a Banach space, $\mathcal{C}$ be a nonempty closed convex set in E and $\mathcal{K}$ a compact subset of $\mathcal{C}$.*

Then every continuous map $\Psi : \mathcal{C} \rightarrow \mathcal{C}$ has a fixed point in $\mathcal{K}$.

We take some space to recall some important and useful measure theory findings.

Definition A.5 (Symmetric measure). A measure μ on the compact metric space $(Q)^n$ is *symmetric* if for every permutation $\sigma \in \mathcal{S}^n$ it holds:

$$\mu(x_{\sigma(1)}, \dots, x_{\sigma(n)}) = \mu$$

Theorem A.6 (Hewitt-Savage). *Let (m_n) be a sequence of symmetric probability measures on $(Q)^n$, where Q is a compact metric space, such that:*

$$\int_Q dm_{n+1}(x_{n+1}) = m_n \quad \forall n \geq 1$$

Then there is a probability measure μ on $\mathcal{P}(Q)$ which verifies, for any continuous map $f \in C^0(\mathcal{P}(Q))$:

$$\lim_{n \rightarrow \infty} \int_{Q^n} f\left(\frac{1}{n} \sum_{i=1}^n \delta_{x_i}\right) dm_n(x_1, \dots, x_n) = \int_{\mathcal{P}(Q)} f(m) d\mu(m)$$

Moreover,

$$m_n(A_1, \dots, A_n) = \int_{\mathcal{P}(Q)} m(A_1) \cdots m(A_n) d\mu(m) \quad \forall n, A_1, \dots, A_n \in \mathcal{B}(Q)$$

Corollary A.7. *Let $m_0 \in \mathcal{P}_1$ be a measure on $X = \mathbb{R}^d$ and let $m_N = \otimes_{i=1}^N m_0$ be the law on X^N of N i.i.d. random variables $(Y_1, \dots, Y_N)$ with law m_0 . Then, for any Lipschitz continuous map $f \in C^0(\mathcal{P}_1)$ it holds: (note that f accepts measures as input)*

$$\lim_{N \rightarrow \infty} \int_{X^N} f\left(\frac{1}{N} \sum_{i=1}^N \delta_{x_i}\right) dm_N(x_1, \dots, x_N) = \lim_{N \rightarrow \infty} \mathbb{E} \left[f\left(\frac{1}{N} \sum_{i=1}^N \delta_{Y_i}\right) \right] = f(m_0)$$

Bibliography

- [1] Achdou Y., Lasry J-M, Lions P-L, Moll B.: Heterogeneous agent models in continuous time. 2014 Working papers. Princeton, NJ: Princeton University.
- [2] Alós-Ferrer, C.: Individual Randomness in Economic Models with a Continuum of Agents. Working Paper, University of Vienna (2002)
- [3] Bogachev V.I., Da Prato G., Röckner M., Stannat W.: Uniqueness of solutions to weak parabolic equations for measures. Bull. Lond. Math. Soc. 39, 2007, no. 4, 631-640
- [4] Brezis H.: Functional Analysis, Sobolev Spaces and Partial Differential Equations
- [5] Bellomo N., Carbonaro B, Gramani L., On the kinetic and stochastic games theory for active particles: Some reasonings on open large living systems, Mathematical and Computer Modelling 48, 2008
- [6] Cardaliaguet P.: Notes on Mean Field Games (from P-L. Lions lectures at Collège de France), 2010
- [7] Carmona R., Delarue F., Lacker D.: Mean Field Games with common noise, 2014, preprint, <http://arxiv.org/abs/1407.6181v1>
- [8] Carmona R., Fouque, J-P. and Sun, L-H., Mean Field Games and Systemic Risk, 2013, preprint, <http://arxiv.org/abs/1308.2172v1>
- [9] Doob J.L.: Stochastic Processes, Wiley, New York, 1953.
- [10] Friedman A.: Partial differential equations of parabolic type, Prentice-Hall, 1964.
- [11] Goldberg P. W., Turchetta S., Query Complexity of Approximate Equilibria in Anonymous Games, preprint
- [12] Guéant O.: A reference case for mean field games models. J. Math. Pures Appl. 2009, 92, 276-294.
- [13] Hinow P., Gerlee P. et al., A spatial model of tumor-host interaction: application of chemotherapy, Math. Biosc. Engin., 2009.
- [14] Ladyzhenskaja O.A., Solonnikov V.A., Ural'ceva N.N., Linear and Quasi-linear Equations of Parabolic Type, American Mathematical Society, translations of mathematical monographs, 1968.

- [15] Lasry J-M, Lions P-L, Guéant, Application of Mean Field Games to Growth Theory, 2008. [jhal-00348376](#)
- [16] Lasry J-M, Lions P-L, Guéant: Mean field games and applications, Paris-Princeton Lectures on Mathematical Finance 2010, Springer, to appear
- [17] Lasry J-M, Lions P-L, Jeux à champ moyen. I. Le cas stationnaire. C. R. Acad. Sci. Paris 343 (2006), 619-625.
- [18] Lasry J-M, Lions P-L, Jeux à champ moyen. II. Horizon fini et contrôle optimal. C. R. Acad. Sci. Paris 343 (2006), 679-684
- [19] Nuño G.: Optimal control with heterogeneous agents in continuous time. 2013 Working Paper Series 1608. Frankfurt am Main, Germany: European Central Bank.
- [20] Podczeck K.: On the existence of rich Fubini extensions. Econ Theory 45, 2010
- [21] Sun Y.: A theory of hyperfinite processes: the complete removal of individual uncertainty via exact lln. Journal of Mathematical Economics, 29:419-503, 1998.
- [22] Yong J., Zhou X.Y.: Stochastic controls, Springer-Verlag New York, 1999.

Acknowledgements

I hereby want to express my gratitude to the ones who have encouraged and inspired me during this Mathematics master.

My first warm acknowledgement is for my advisor Franco, who entrusted me in every moment of this work and donated all his energies for its best completion: the least I can say is that I have been advised in the best possible way, and I really thank him for his all time availability, his sharp observations and his sincere, heartfelt encouragements.

I then want to acknowledge my parents for their trust and love, and for always allowing me to follow what I feel to be the best.

Among the people who inspired me during this master, I gratefully acknowledge Salvatore Rossi, Alan Kirman, Alessio Moneta, Valeria De Bonis, Francesca Campolongo and Nicoletta Cantarini for their motivating example.

Many friends and fellows have walked beside me during this journey, proving me with love, patient explanations and encouragements; I especially thank Federico, Ilaria, Cristina, Elena, Agnese, Ivan, Gianmichele, Luca, Stefano and Pierluigi for always being there.

The endless patience of my flatmates Luca, Antonio, Marco and Andrea deserves a special mention, as well as the support I always received from my home friends Damiano, Caterina, Matteo and Alessandro.

Finally, I am deeply grateful to Coralie and her respectful love, a smiley sun that has warmed me in every moment of this work.