



Modélisation du carnet d'ordres, Applications Market Making

Xiaofei Lu

► To cite this version:

Xiaofei Lu. Modélisation du carnet d'ordres, Applications Market Making. Autre. Université Paris Saclay (COmUE), 2018. Français. NNT: 2018SACLC069 . tel-01893379

HAL Id: tel-01893379

<https://theses.hal.science/tel-01893379>

Submitted on 11 Oct 2018

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Modélisation du carnet d'ordres, Applications Market Making

Limit order book modelling, Market Making applications

Thèse de doctorat de l'Université Paris-Saclay
préparée à CentraleSupélec

Ecole doctorale n°573 INTERFACES
Approches Interdisciplinaires : Fondements, Applications et Innovation

Spécialité de doctorat : Mathématiques appliquées

Thèse présentée et soutenue à Gif-sur-Yvette, le 04 Octobre 2018, par

XIAOFEI LU

Composition du Jury :

Mathieu ROSENBAUM Professeur, École Polytechnique	Président
Aurélien ALFONSI Professeur, École des Ponts ParisTech	Rapporteur
Enrico SCALAS Professeur, University of Sussex	Rapporteur (absent)
Frédéric ABERGEL Professeur, CentraleSupélec	Directeur de thèse
Nakahiro YOSHIDA Professeur, University of Tokyo	Examineur
Marouane ANANE Docteur, Ingénieur de recherche, BNP Paribas	Examineur
Ioane MUNI TOKE Maître de conférence, CentraleSupélec	Examineur
Sarah LEMLER Maître de conférence, CentraleSupélec	Examineur

Résumé

Cette thèse aborde différents aspects de la modélisation de la microstructure du marché et des problèmes de Market Making, avec un accent particulier du point de vue du praticien. Le carnet d'ordres, au cœur du marché financier, est un système de files d'attente complexe à haute dimension. Nous souhaitons améliorer la connaissance du LOB pour la communauté de la recherche, proposer de nouvelles idées de modélisation et développer des applications pour les Market Makers. Nous remercions en particulier l'équipe *Automated Market Making* d'avoir fourni la base de données haute-fréquence de très bonne qualité et une grille de calculs puissante, sans laquelle ces recherches n'auraient pas été possible.

Le Chapitre 1 présente la motivation de cette recherche et reprend les principaux résultats des différents travaux.

Le Chapitre 2 se concentre entièrement sur le LOB et vise à proposer un nouveau modèle qui reproduit mieux certains faits stylisés. Nous utilisons un processus de Hawkes de haute dimension pour modéliser les flux d'ordres. En introduisant des inhibitions entre les ordres, la performance est beaucoup améliorée. A travers cette recherche, non seulement nous confirmons l'influence des flux d'ordres historiques sur l'arrivée de nouveaux, mais un nouveau modèle est également fourni qui réplique beaucoup mieux la dynamique du LOB, notamment la volatilité réalisée en haute et basse fréquence.

Dans le Chapitre 3, l'objectif est d'étudier les stratégies de Market Making dans un contexte plus réaliste. À partir du modèle 'queue-reactive' proposé par [Huang et al. \[2015a\]](#), les améliorations apportées à deux aspects - la taille aléatoire des ordres et la forte influence des ordres qui consomment entièrement la première limite - améliorent largement le modèle initial. Les fonctions de valeur des stratégies optimales du Market Maker dans le nouveau modèle sont analysées en profondeur et comparées à celles du modèle initial. Les simulations de Monte Carlo et les backtests avec des données réelles montrent tous les deux que le nouveau modèle est plus réaliste. Cette recherche contribue à deux aspects: d'une part le nouveau modèle proposé est plus réaliste mais reste simple à appliquer pour la conception de stratégies, d'autre part la stratégie pratique de Market Making est beaucoup améliorée par rapport à une stratégie naive et est prometteuse pour l'application pratique.

La prédiction à haute fréquence avec la méthode d'apprentissage profond est étudiée dans le Chapitre 4. Premièrement, nous montrons que la prédiction du prix mid en 1-étape par les indicateurs LOB est non-linéaire, stationnaire et universelle dans une base de données européenne, similaire aux résultats de [Sirignano and Cont \[2018\]](#) pour les actions américaines. Cependant, l'algorithme d'optimisation parallèle disponible pour calibrer un modèle universel n'est pas assez puissant et se dégrade rapidement avec l'augmentation du nombre de processus. Un modèle universel de 25 processus parallèles est légèrement moins bien qu'un modèle stock par stock d'environ 100 stocks, ce qui pose le problème du compromis entre le temps de calibration et la précision de la prédiction. En outre, nous illustrons que le rendement en 1-étape est largement

influencé par la réversion artificielle du prix mid, de sorte que l'étude de la prédiction à plus long terme est nécessaire. Le modèle LSTM surpasse les modèles linéaires pour des prédictions en 10-étapes, mais se détériore rapidement pour un horizon plus long en raison du rapport signal par rapport au bruit plus faible. Une stratégie de market making simple basée sur des prédictions lissées par EMA montre la surperformance de LSTM par rapport aux modèles linéaires et le potentiel réel d'application de ce modèle d'apprentissage profond pour le Market Making.

Mots-clés: Trading haute fréquence, Microstructure du marché, Carnet d'ordres, Processus de Hawkes, Processus de décision Markovien, Apprentissage profond.

Abstract

This thesis addresses different aspects around the market microstructure modelling and market making problems, with a special accent from the practitioner’s viewpoint. The limit order book (LOB), at the heart of financial market, is a complex continuous high-dimensional queueing system. We wish to improve the knowledge of LOB for the research community, propose new modelling ideas and develop concrete applications to the interest of Market Makers. We would like to specifically thank the *Automated Market Making* team for providing a large high frequency database of very high quality as well as a powerful computational grid, without whom these researches would not have been possible.

The first chapter introduces the incentive of this research and resumes the main results of the different works.

Chapter 2 fully focuses on the LOB and aims to propose a new model that better reproduces some stylized facts. We use a high-dimensional Hawkes process to model the order flows, and find that by introducing inhibitions between orders the performance is much improved. Through this research, not only do we confirm the influence of historical order flows to the arrival of new ones, but a new model is also provided that captures much better the LOB dynamic, notably the realized volatility in high and low frequency.

In chapter 3, the objective is to study Market Making strategies in a more realistic context. Starting from the queue-reactive model proposed by [Huang et al. \[2015a\]](#), enhancements from two aspects – random order size and the strong influence of liquidity removal order – largely improve the initial model. Value functions of optimal market making strategies in the new model is analysed in-depth and compared to those in the initial model. Monte Carlo simulations and backtests with real data both show that the new model is more realistic. This research contributes in two aspects: from one hand the newly proposed model is more realistic but still simple enough to be applied for strategy design, on the other hand the practical Market Making strategy is of large improvement compared to the naive one and is promising for practical use.

High-frequency prediction with deep learning method is studied in chapter 4. Firstly, we show that the prediction of 1-step mid-price by LOB indicators is non-linear, stationary and universal in a European database, which is similar to the findings of [Sirignano and Cont \[2018\]](#) in US stocks. However, the available parallel optimization algorithm to calibrate a universal model is not powerful enough and exhibits a quick decrease in performances with the increase of number of processes. A 25-worker universal model is slightly worse than a stock-by-stock model of about 100 stocks, which raises the problem of trade-off between the calibration time and prediction accuracy. Moreover we illustrate that 1-step return is largely influenced by the artificial mean-reversion of mid-price so that the study of longer term prediction is necessary. LSTM model out-performs the linear models at 10-step predictions, but worsens fast for longer horizon due to the lower signal-to-noise ratio. A simple market making strategy based on EMA-smoothed predictions shows the out-performance of LSTM to linear models and the real

application potential of this deep learning model for Market Making.

Keywords: Market Microstructure, High-frequency trading, Limit order book, Hawkes process, Markov decision process, Deep learning.

Contents

Résumé	iii
Abstract	v
1 Contexte, Méthodes et Résultats (In French)	3
1.1 Contexte et objectifs	3
1.2 Modélisation du carnet d'ordres par un processus de Hawkes en grande dimension	4
1.2.1 Constatations empiriques: les interdépendances des événements du carnet d'ordres	5
1.2.2 Modélisation des dépendances par processus de Hawkes	7
1.2.3 Aspects numériques de la calibration du modèle	9
1.2.4 Conclusion	10
1.3 Modélisation du carnet d'ordre et stratégies de market making	10
1.3.1 Questionner le modèle "queue-reactive"	11
1.3.2 Market making au marché réel	13
1.3.3 Conclusion	16
1.4 Prédiction de prix avec des informations de microstructure boostée par deep learning	16
1.4.1 Prédiction en 1-étape	17
1.4.2 Régression par deep learning en plusieurs étapes	19
1.4.3 Conclusion	22
2 Limit order book modelling with high dimensional Hawkes processes	25
2.1 Introduction	26
2.2 Empirical findings: the interdependences of order book events	27
2.2.1 Data and Framework	27
2.2.2 Event definitions	27
2.2.3 Statistical dependencies between order book events	27
2.3 Modelling dependencies using Hawkes processes	30
2.3.1 Linear Hawkes process models	32
2.3.2 Performances of the linear Hawkes models	33
2.3.3 Nonlinear Hawkes process model	36
2.3.4 Performances of the nonlinear Hawkes models	38
2.4 Some numerical aspects of model calibration	41
2.4.1 Calibration with maximum likelihood estimation	42
2.4.2 Benchmarking the DE algorithm	44
2.4.3 Improvement in high dimensions	44
2.5 Conclusion	48
3 Order-book modelling and market making strategies	51
3.1 Introduction	52

3.2	Challenging the queue-reactive model	53
3.2.1	The queue-reactive model	53
3.2.2	The limitation of unit order size	54
3.2.3	The role of limit removal orders	57
3.2.4	Enriching the queue-reactive model	62
3.3	Market making in real markets	63
3.3.1	Optimal market making strategies	64
3.3.2	Backtesting the optimal strategies	68
3.4	Conclusion	71
4	Price prediction with microstructure information boosted by deep learning	73
4.1	Introduction	74
4.2	Deep learning experiment: one-step price change prediction	76
4.2.1	Indicators	76
4.2.2	Model construction	77
4.2.3	Empirical results	80
4.3	Multi-step deep learning regression	85
4.3.1	Model performance	87
4.3.2	Strategy building	89
4.4	Conclusion	92
	General conclusions and Outlooks	95
	Conclusion générale et perspectives	97
	Bibliography	104
	Appendix A Pseudocode of basic Differential Evolution	105
	Appendix B Markov decision processes and optimal strategies	107
B.1	Keep or cancel strategy for buying one unit	108
B.2	Market making (Making the spread)	109
	Appendix C Pseudocode of DC-ASGD algorithm	111

Chapter 1

Contexte, Méthodes et Résultats (In French)

1.1 Contexte et objectifs

Les marchés financiers sont des places de change permettant aux producteurs de se financer auprès des investisseurs ou de vendre leurs produits à des clients finaux. Ces marchés assurent donc l'apport de la liquidité à l'économie et garantissent à chaque investisseur de pouvoir changer d'investissement ou récupérer librement son argent sans impacter significativement l'activité économique globale.

Avec la révolution numérique qui a marqué la fin du vingtième siècle, les marchés financiers ont évolué progressivement vers des marchés électroniques. Cette évolution a amélioré considérablement le mécanisme de diffusion de l'information ainsi que la transparence de la formation du prix. Désormais, les acteurs des marchés affichent publiquement et anonymement leurs intérêts à l'achat et à la vente. Une transaction est effectuée chaque fois qu'un intérêt acheteur croise un intérêt vendeur. Les intérêts non exécutés (et non annulés) sont agrégés dans un carnet appelé le carnet d'ordre (*limit order book* en anglais) ([Abergel et al. \[2016\]](#) donne une étude élaborée du carnet d'ordres). Les enregistrements des variations de ces carnets d'ordres à la granularité la plus fine représentent une source très riche d'information qui permet aux mathématiciens de mieux modéliser et étudier la microstructure des marchés financiers.

Malgré ces avantages, le passage des marchés au monde électronique ne s'est pas fait sans effets secondaires indésirables. En particulier, la participation grandissante des automates dans les échanges financiers a accentué les risques de sélection adverse, de manipulation de prix et de crash boursier. Afin de faire face à ces risques, les principales bourses mondiales font désormais appel à des agents spécialisés appelés "teneur de marché" (*Market Maker* en anglais).

Pour définir le Market Making, nous pouvons se référer à la définition donnée par le *MiFID* (*Markets in Financial Instruments Directive*, voir [MiF](#)): "Un Market Maker est un membre ou un participant d'une ou plusieurs places de trading, qui utilise une stratégie soumettant simultanément des cotations de prix compétitifs et de quantités comparables à l'achat et à la vente pour un ou plusieurs actifs dans une seule place ou par plusieurs plateformes de trading différentes, avec pour résultat la fourniture de liquidité de manière régulière et fréquente pour l'ensemble du marché" .

Le Market Maker est donc par définition un professionnel averti qui a d'une part les connaissances nécessaires pour faire face aux risques du trading électronique et d'autre part la capacité financière pour faire face aux tentatives de manipulations de marché ou à l'absence ponctuelle de liquidité. La présence du Market Maker fournit une assurance aux autres investisseurs de trouver de la liquidité sur le marché à chaque instant. En particulier, ces derniers sont assurés de pouvoir acheter ou vendre à tout instant à un prix raisonnable, ce prix étant le résultat d'un consensus global du marché et non pas un artéfact d'un déséquilibre instantané.

Le carnet d'ordres étant au coeur du mécanisme de la formation des prix, la maîtrise de sa dynamique est donc essentielle pour tout Market Maker. La modélisation de cette dynamique est assez complexe à cause de la grande dimension du problème. En particulier, chaque ordre est caractérisé par son type, son prix, sa quantité et son temps d'arrivée. Les modélisations les plus basiques se basent sur une quantité unique, un prix discret avec un incrément unitaire, et un temps d'arrivée en processus de Poisson. Sans surprise, ces modèles simplistes ne permettent pas de retrouver les caractéristiques empiriques observées sur les marchés.

L'objectif de cette thèse est, d'une part, de contribuer aux efforts académiques, en place depuis plusieurs années, pour mieux modéliser les carnets d'ordres, et d'autre part de proposer des améliorations concrètes aux stratégies des vrais Market Makers. Les outils mathématiques que nous appliquons ont déjà été étudiés dans différents papiers académiques, notre contribution principale consiste à les appliquer aux vraies problématiques des Market Makers dans la pratique.

Dans la suite de ce chapitre, nous présentons les différentes parties de ce travail en récapitulant les méthodes et les résultats obtenus.

1.2 Modélisation du carnet d'ordres par un processus de Hawkes en grande dimension

Dans un marché dirigé par les ordres, les participants peuvent envoyer trois types d'ordres élémentaires: ordre limite, ordre marché et annulation:

- **Ordre limite:** Un ordre qui spécifie une limite de prix supérieure/inférieure à laquelle le trader est prêt à acheter / vendre un certain nombre d'actions.
- **Ordre marché:** Un ordre qui déclenche une transaction d'achat/vente immédiate pour un certain nombre d'actions au meilleur prix opposé disponible.
- **Annulation :** Un ordre qui supprime un ordre limite existante.

Le carnet d'ordres est l'agrégation de tous les ordres limites d'achat et de la vente qui ne sont ni exécutés ni annulés dans le marché. Un exemple typique avec explication détaillée est donné dans Figure 1.1.

Dans le chapitre 2, nous modélisons les processus d'ordres par un processus de Hawkes de grande dimension. Nous illustrons que le nouveau modèle arrive à reproduire des faits stylisés importants. Cette approche nous permet aussi de donner une interprétation financière pour les comportements à plusieurs échelles, ainsi que de montrer l'existence d'effet inhibiteur en même temps que l'effet d'excitation.

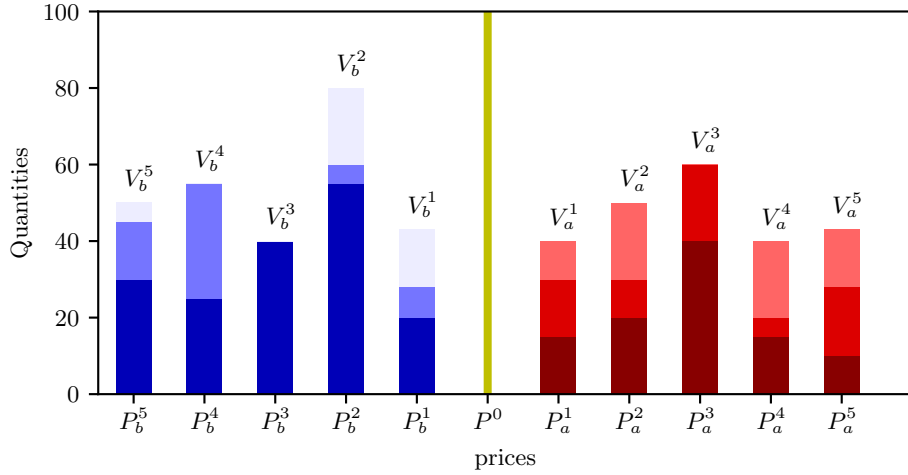


Figure 1.1: Illustration du carnet d'ordres. Les barres bleues à gauche de la figure représentent les ordres d'achat avec prix P_b^i et de quantités totales V_b^i . Ceux-ci correspondent au côté acheteur. Les participants du côté acheteur fournissent des liquidités avec des prix auxquels ils sont prêts à acheter certaines quantités de stock. Les barres de droite représentent le côté vente, où les participants désireux de vendre affichent leurs ordres avec les prix qu'ils sont prêts à vendre le stock. La ligne du milieu (P^0) correspond au niveau du prix moyen ("mid price") et est calculée comme la moyenne entre le meilleur (le plus élevé) prix d'achat et le meilleur (le plus bas) prix de vente. Une transaction se produit lorsqu'un ordre de vente et un ordre d'achat sont au moins partiellement appariés. Une file d'attente des ordres limites avec le même prix est appelée une limite. Des couleurs différentes dans la même limite représentent des ordres avec une priorité différente, les barres plus foncées ayant une priorité d'exécution plus élevée.

1.2.1 Constatations empiriques: les interdépendances des événements du carnet d'ordres

Nous considérons tous les ordres qui modifient la meilleure limite à l'achat ou à la vente comme un "événement", et nous les classons selon leurs types et leurs côtés du carnet (achat ou vente). Nous les étiquetons de plus selon qu'ils induisent ou pas un changement de prix. Tableau 1.1 résume les notations des événements considérés dans ce chapitre.

Pour étudier l'interdépendance entre les ordres, nous calculons le ratio de la probabilité d'un type d'événement, conditionnellement au dernier événement observé, divisée par sa probabilité non-conditionnelle. Le résultat est présenté dans Tableau 1.2. Pour chaque ligne i et colonne j du tableau, l'entrée représente la quantité $\frac{\mathbb{P}[O_j|O_i]}{\mathbb{P}[O_j]}$. Un ratio très élevé signifie que O_i incite l'apparition de O_j , alors que un ratio bas se traduit par l'effet inhibiteur de O_i sur O_j . Pour mieux visualiser les effets importants, les ratios supérieurs à 2 ou inférieurs à 0.2 sont représentés en gras dans le tableau.

Nous retrouvons plusieurs effets connus sur les marchés, telque les cross-excitations des ordres, par exemple suite à M_{buy}^0 , la probabilité d'avoir M_{buy}^0 et M_{buy}^1 augmente fortement, ceci est expliqué par le fractionnement des "meta-ordre" et par l'effet momentum. Nous observons également une augmentation modérée pour L_{buy}^1 et C_{sell}^1 suite au consensus du nouveau prix. L_{buy}^1 incite C_{buy}^1 et M_{sell}^1 , tous les deux viennent consommer la nouvelle liquidité fournie. Suite à C_{buy}^1 , Nous voyons que les probabilités de C_{buy}^0 et L_{sell}^1 augmentent. Pour le C_{buy}^0 le changement de probabilité est expliqué par un effet de momentum sur les annulations. Pour le L_{sell}^1 le

Table 1.1: Notations des différents types d'ordres

Notation	Définition
M, L, C, O	ordre marché, ordre limite, annulation, tout ordre.
M^0, L^0, C^0, O^0	ordre marché, ordre limite, annulation, tout ordre, qui ne change pas le prix.
M^1, L^1, C^1, O^1	ordre marché, ordre limite, annulation, tout ordre, qui change le prix.
M_{buy}, M_{sell}	ordre marché à l'achat/ à la vente.
M_{buy}^0, M_{sell}^0	ordre marché à l'achat/ à la vente qui ne change pas le prix
M_{buy}^1, M_{sell}^1	ordre marché à l'achat/ à la vente qui change le prix
L_{buy}, L_{sell}	ordre limite à l'achat/ à la vente.
L_{buy}^0, L_{sell}^0	ordre marché à l'achat/ à la vente qui ne change pas le prix
L_{buy}^1, L_{sell}^1	ordre limite à l'achat/ à la vente qui change le prix
C_{buy}, C_{sell}	annulation à l'achat/ à la vente.
C_{buy}^0, C_{sell}^0	annulation à l'achat/ à la vente qui ne change pas le prix:
C_{buy}^1, C_{sell}^1	annulation qui change le prix

changement vient de la recharge de la liquidité disparue. M_{buy}^1 augmente la probabilité de L_{buy}^1 , C_{sell}^1 et M_{buy}^1 , qui viennent tous de l'effet momentum court terme et du fractionnement des "meta-ordre".

Un autre effet très intéressant, et peu étudié dans la littérature, est la relation d'inhibition entre les ordres. Par exemple nous observons une baisse de la probabilité de C_{buy}^1 après L_{buy}^0 . Ceci est naturel car un L_{buy}^0 signifie qu'un nouvel ordre s'ajoute aux ordres déjà existants à la première limite. Comme deux ordres sont très rarement annulés en même temps, la probabilité de C_{buy}^1 chute mécaniquement. Nous observons aussi une baisse de la probabilité de C_{buy}^0 suite à un L_{buy}^1 , car l'annulation d'un seul ordre est très rarement partiel. Suite à C_{buy}^1 , plusieurs ordres ont une baisse de probabilité, dont les plus remarquables sont M_{sell}^0 et M_{sell}^1 . Une annulation d'un ordre qui change le prix signifie une augmentation de la différence du prix de vente et d'achat, qui baisse les intérêts de transactions. Et en plus il y a souvent une remise de liquidité sur le même prix (la probabilité élevée du L_{buy}^1), donc ça baisse autant plus l'intérêt des vendeurs pour une transaction.

Cette étude des probabilités empiriques montre qu'il y a une forte dépendance temporelle entre les événements du carnet d'ordres. Les dépendances les plus classiques sont les cross-excitations, qui engendrent le clustering des ordres, assez connu en microstructure. Il existe aussi des relations d'inhibition aussi importantes pour la dynamique du carnet, mais qui sont moins connues dans le monde académique.

Table 1.2: Probabilités conditionnelles relatives

	L_{buy}^0	L_{sell}^0	C_{buy}^0	C_{sell}^0	M_{buy}^0	M_{sell}^0	L_{buy}^1	L_{sell}^1	C_{buy}^1	C_{sell}^1	M_{buy}^1	M_{sell}^1
L_{buy}^0	1.4	0.5	1.3	1.0	1.4	0.7	1.2	0.4	0.0	1.2	1.6	0.3
L_{sell}^0	0.5	1.4	1.0	1.3	0.7	1.4	0.5	1.2	1.2	0.0	0.3	1.6
C_{buy}^0	1.1	1.2	1.4	0.4	0.8	0.8	0.7	1.1	1.7	0.4	0.6	0.7
C_{sell}^0	1.2	1.1	0.4	1.4	0.8	0.8	1.1	0.7	0.5	1.7	0.7	0.6
M_{buy}^0	1.0	0.4	0.3	0.5	11.5	0.8	2.4	0.2	0.4	2.3	15.0	0.6
M_{sell}^0	0.4	1.0	0.5	0.3	0.8	13.2	0.2	2.4	2.3	0.4	0.6	15.7
L_{buy}^1	1.6	0.5	0.1	1.3	1.2	1.3	1.1	0.4	3.0	1.7	1.5	2.9
L_{sell}^1	0.5	1.6	1.3	0.1	1.2	1.2	0.4	1.1	1.7	3.0	2.8	1.5
C_{buy}^1	0.7	0.6	2.5	0.2	0.3	0.1	2.1	1.0	1.0	0.3	0.3	0.1
C_{sell}^1	0.6	0.7	0.2	2.5	0.1	0.3	1.0	2.1	0.3	0.9	0.1	0.3
M_{buy}^1	0.6	0.3	0.2	1.5	1.3	0.9	6.2	1.1	0.4	2.9	2.1	1.6
M_{sell}^1	0.3	0.6	1.5	0.2	0.9	1.5	1.1	6.2	3.0	0.4	1.6	2.1

1.2.2 Modélisation des dépendances par processus de Hawkes

Nous modélisons les 12 types d'événements par un processus ponctuel multivarié $N(t) = (N_{L_{buy}^0}(t), \dots, N_{M_{sell}^1}(t))$ associé avec son processus d'intensité $(\lambda_{L_{buy}^0}(t), \dots, \lambda_{M_{sell}^1}(t))$. Deux différents modèles sont proposés.

D'abord nous considérons un processus de Hawkes linéaire, dont l'intensité s'écrit comme

$$\lambda_m(t) = \mu_m + \sum_{n=1}^M \int_0^t \phi_{mn}(t-s) dN_n(s) \quad (1.1)$$

En plus nous avons essayé deux types de noyaux exponentiels

- Modèle de Hawkes 1-exponentiel, $\phi_{ij}(t) = \alpha_{ij} \exp(-\beta_{ij}t)$
- Modèle de Hawkes 2-exponentiel, $\phi_{ij}(t) = \sum_{p=1}^2 \alpha_{ijp} \exp(-\beta_{ijp}t)$

Pour évaluer la performance de notre modèle, nous utilisons deux critères pour le comparer avec les données empiriques. Pour un processus ponctuel $((T_i, X_i))_{i \in \mathbb{N}_*}$, nous savons que les durées transformées $\tau_{m,i}$

$$\tau_{m,i} = \int_{T_i}^{T_{i+1}} \lambda_m(s) ds$$

sont i.i.d, et suit une loi d'exponentiel de paramètre 1. De ce fait, un test "Q-Q plot" qui compare les quantiles empiriques du modèle par rapport au quantile théorique de la loi exponentielle est appliqué. D'autre part, nous nous intéressons au "signature plot", qui illustre la relation entre la variance réalisée $RV(h)$ et la fréquence d'échantillonnage h , dont RV est défini comme

$$RV(h) = \frac{1}{T} \sum_{n=0}^{T/h} (X((n+1)h) - X(nh))^2.$$

Pour un processus stochastique $X_t, t \in [0, T]$.

Le modèle de processus de Hawkes est plus performant que le modèle de Poisson dans ces deux tests. Mais il y a un écart non-négligeable par rapport aux données empiriques. Plus

particulièrement la volatilité reproduite par le modèle est beaucoup plus élevée par rapport au vrai marché. Ceci vient du fait que les effets inhibiteur des ordres ne sont pas bien modélisés par un processus de Hawkes linéaire. Le minimum des noyaux est 0, donc l'intensité λ_m a une borne inférieure μ_m . Nous montrons dans Tableau 1.3 que certaines probabilités conditionnelles sont bien trop élevées dans le modèle. Par ailleurs, Nous montrons dans Tableau 1.4 les normes L_1 des noyaux calibrés pour la stimulation de C_{buy}^1 : il est évident que les couples d'ordres dans Tableau 1.3 ont 0 normes, M_{buy}^0 et C_{sell}^1 , mais aussi pour M_{sell}^0 , C_{buy}^1 et M_{sell}^1 qui sont moins explicites dans la matrice de dépendance.

Table 1.3: Comparaison de probabilité conditionnelle entre le flux d'ordres simulé et les données réelles

Stimulating pair	P_{simu}	P_{real}
$C_{buy}^1 L_{buy}^0$	0.402	0.048
$L_{buy}^1 L_{buy}^1$	1.628	0.141
$L_{sell}^1 L_{buy}^1$	1.288	0.171
$M_{sell}^0 C_{buy}^1$	0.545	0.068
$C_{buy}^1 C_{buy}^1$	0.548	0.072
$M_{sell}^1 C_{buy}^1$	0.854	0.037

Table 1.4: Médiane des noyaux L_1 dans le modèle du processus de Hawkes 2-exponentielle

	L_{buy}^0	L_{sell}^0	C_{buy}^0	C_{sell}^0	M_{buy}^0	M_{sell}^0	L_{buy}^1	L_{sell}^1	C_{buy}^1	C_{sell}^1	M_{buy}^1	M_{sell}^1
C_{buy}^1	0.1563	0.2357	0.9392	0.0914	0	0	0.3845	0.1607	0	0	0.0013	0

Pour intégrer les comportements inhibiteurs dans le modèle, nous proposons d'utiliser le processus de Hawkes non-linéaire, dont l'intensité est

$$\lambda(t) = (\mu + \Phi \star dN)_+, \quad (1.2)$$

avec des noyaux négatifs de forme

$$\phi_{mn} = \sum_{p=1}^2 -\alpha_{mnp} \exp(-\beta_{mnp} t)$$

De nombreux tests sont appliqués pour évaluer la performance de ce modèle. Nous observons une claire amélioration par l'introduction d'effets inhibiteurs dans tous ces tests. Notamment grâce à ce nouveau modèle, nous avons réussi à reproduire un signature plot (Figure 1.2) très proche des vraies données.

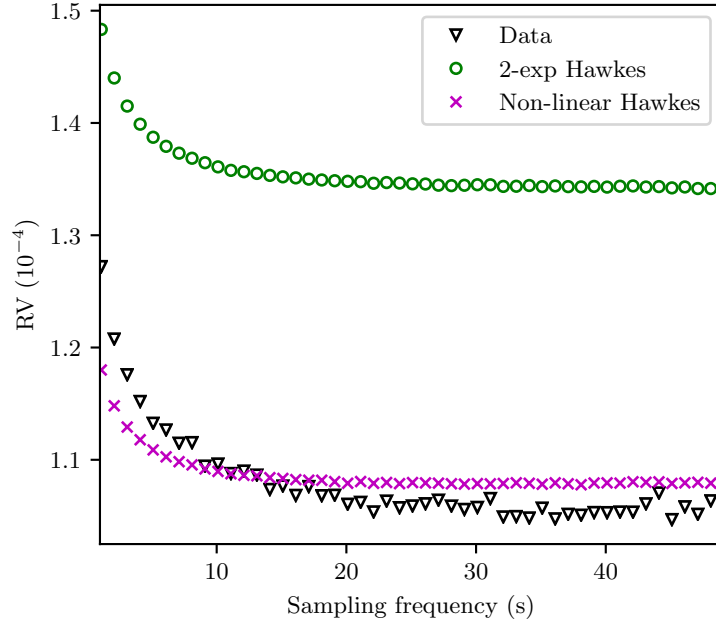


Figure 1.2: Comparaison de signature plot entre le modèle de processus de Hawkes 2-exponentielle linéaire, non-linéaire et les données réelles.

1.2.3 Aspects numériques de la calibration du modèle

La calibration du modèle de Hawkes se fait par la maximisation de la vraisemblance. Cette calibration n'est pas triviale en grande dimension car la fonction n'est pas concave, et elle a plusieurs optima locaux. Nous avons trouvé que l'algorithme classique de Nelder-Mead peut avoir de très mauvaises performances, selon le point de départ aléatoire, même dans un cas très simple. Nous avons donc utilisé un algorithme génétique *Differential Evolution (DE)* pour améliorer l'optimisation. Cet algorithme imite l'évolution génétique en partant d'une première population de points aléatoires, qui sont améliorés d'une génération à une autre par des mutations et des crossing-over. Cette procédure est répétée jusqu'à la convergence de l'algo ou l'atteinte du nombre maximale d'itérations. Les algorithmes Nelder-Mead et DE sont comparés dans Tableau 1.5 pour un processus de Hawkes à deux dimensions. Les paramètres calibrés sont considérés comme erronés s'ils ont plus de 10% d'écarts par rapport aux valeurs réelles.

Algorithm	T	μ_1	α_{11}	α_{12}	β_{11}	β_{12}	μ_2	α_{21}	α_{22}	β_{21}	β_{22}
DE	250	0.0	0.0	0.0	0.0	0.0	0.6	1.0	2.0	1.6	2.2
	2500	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.1	0.0	0.1
	25000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
NM random	250	2.1	2.1	2.0	2.0	2.1	24.5	27.6	30.7	24.8	30.5
	2500	1.4	1.4	1.4	1.4	1.4	12.1	14.8	16.9	14.6	18.8
	25000	1.0	1.0	0.8	1.0	1.0	9.9	12.6	14.0	11.4	16.7
NM perfect	250	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
	2500	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
	25000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0

Table 1.5: Taux d'erreur de calibration par algorithme Nelder-Mead et DE

Plusieurs adaptations de l’algorithme de base sont essayées et évaluées dans notre application. Nous avons choisi L-SHADE avec différentes méthodes d’amélioration. L’augmentation la taille de la population renforce largement la convergence de l’algorithme: elle empêche les points d’être piégés autour des maxima locaux. Et il est aussi inutile de garder tous les points surtout quand l’algorithme s’approche de la fin de ses itérations. Donc nous utilisons une réduction de population linéaire inspirée de [Li and Zhang \[2011\]](#) et [Yang et al. \[2013\]](#). Nous proposons aussi une nouvelle réduction de population.

Supposons que nous ayons trié les points par leurs valeurs de fonction objective par ordre décroissant. Pour chaque point $x_i = (\mu_i, \rho_i, \beta_i)$ de la population, si $\exists j \in \llbracket 1, b \rrbracket / \{i\}$ tel que toutes les conditions suivantes sont satisfaites,

$$\begin{aligned} |\mu_{mi} - \mu_{mj}| < e_r \mu_{mj} \quad \text{or} \quad |\mu_{mi} - \mu_{mj}| < e_a \\ |\rho_{mni} - \rho_{mnj}| < e_r \rho_{mnj} \quad \text{or} \quad |\rho_{mni} - \rho_{mnj}| < e_a \\ |\beta_{mni} - \beta_{mnj}| < e_r \beta_{mnj} \quad \text{or} \quad \rho_{mnj} < e_a \end{aligned}$$

alors x_i est supprimé de la population, parce que x_i et x_j ont convergés vers un seul point. Le b est au choix. En pratique, un petit b aide à supprimer des points de dupliqués et à explorer des espaces loin des meilleurs points actuels pour éviter les maxima locaux.

La combinaison de ces deux techniques de réduction de population permet d’augmenter la population initiale d’un facteur de 5 à 10 sans impact significatif sur le temps total de calcul, et la convergence est largement améliorée.

1.2.4 Conclusion

Dans cette section, nous avons étudié empiriquement la dynamique du carnet d’ordre, puis nous l’avons modélisée avec des processus de Hawkes multidimensionnels. Nous avons observé que les modèles classiques ne permettent pas de retrouver plusieurs effets réels importants telque le “Signature Plot” de la volatilité réalisée. Afin d’améliorer la modélisation, nous avons opté pour un processus de Hawkes non-linéaire capable de reporduir les effets d’inhibition entre les arrivée de certains événements du carnet d’ordres. Ce nouveau modèle permet une meilleure reproduction des observations empiriques et participe à enrichir la recherche théorique en microstructure. Nous avons aussi montré que face à la difficulté de l’optimisation numérique de la fonction de vraisemblance, il est essentiel de bien choisir son algorithme d’optimisation. Nous avons opté pour l’algorithme L-SHADE qui permet une amélioration significative par rapport à l’algorithme classique de DE, grâce à une meilleure initialisation et à un contrôle plus efficace de la population. Ces résultats ne sont pas spécifiques à la finance et peuvent être utiles dans d’autres problèmes d’optimisation globale.

1.3 Modélisation du carnet d’ordre et stratégies de market making

La problématique de stratégie de market making a été étudiée de manière approfondie. Pourtant la plupart des contributions concentrent sur la gestion de risque d’inventaire avec un modèle d’exécution simpliste pour faciliter la mise en place de la formulation du cadre de contrôle stochastique. Mais en pratique la discontinuité des prix et la dynamique de file d’attente intrinsèque du carnet d’ordre rendent ces simplifications trop simplistes par rapport aux marchés

réels, et il est évident pour le praticien que ces propriétés microstructurales réelles jouent un rôle fondamental dans l'évaluation de la rentabilité des stratégies de market making.

Il est par ces observations pratiques qui nous incitent à contribuer du point de vue de la modélisation du carnet d'ordres et de la conception de stratégie de market making. Nous utilisons les données de futur Eurostoxx 50 entre juin et juillet 2016 pour faire toutes les analyses, et faisons un backtest jusqu'à novembre pour la validation "out-of-sample". Il y a deux caractéristiques du futur Eurostoxx 50 tel qu'il est pertinent pour notre étude

1. c'est un instrument de "large tick" (la minimum incrémentation du prix est très grande), avec un "spread" (écart entre le meilleur prix de vente et d'achat) moyen très proche de 1 tick et des transactions à plusieurs limites extrêmement rares (inférieure à 0.5%);
2. la valeur d'un contrat est très élevée en euros, de sorte que nous le pensons en terme de nombre de contrats plutôt que de montant notionnel. Cela simplifie réellement le choix de l'unité.

Dans le chapitre 3, nous analysons et améliorons le modèle du carnet d'ordres "queue-réactive" proposé par Huang et al. [2015b] et puis étudie le placement optimal d'une paire d'ordres à l'achat et à la vente en tant que paradigme du market maker.

1.3.1 Questionner le modèle "queue-reactive"

Dans Huang et al. [2015b], les auteurs proposent un modèle de carnet d'ordres Markovien intéressant. Le carnet d'ordres est modélisé par un vecteur de dimension $2K$, avec $[Q_{-i} : i = 1, \dots, K]$ et $[Q_i : i = 1, \dots, K]$ représentent les limites d'achat et de vente respectivement pour les limites placées à $i-0.5$ ticks du prix de référence p_{ref} . Désignant les quantités correspondantes par q_i , le processus de $2K$ dimension $X(t) = (q_{-K}(t), \dots, q_{-1}(t), q_1(t), \dots, q_K(t))$ avec valeurs dans $\Omega = \mathbb{N}^{2K}$ est modélisé par une chaîne de Markov en temps continu avec le générateur infinitésimal de la forme

$$\begin{aligned} \mathcal{Q}_{q, q+e_i} &= f_i(q), \\ \mathcal{Q}_{q, q-e_i} &= g_i(q), \\ \mathcal{Q}_{q, q} &= - \sum_{q \in \Omega, p \neq q} \mathcal{Q}_{q, p}, \\ \mathcal{Q}_{q, p} &= 0 \text{ sinon.} \end{aligned}$$

Pourtant, quand ce modèle a été calibré aux données du futur Eurostoxx 50, des nouveaux phénomènes sont observés qui nous amène à enrichir le modèle dans deux directions.

D'abord nous nous intéressons à la quantité des ordres. Les modèles classiques du carnet suppose la taille d'ordre est unique et constante. Toutefois les tailles des ordres ont des propriétés statistiques remarquables, et que ces informations sont pertinentes pour la dynamique du carnet.

Nous présentons les tailles moyennes des ordres conditionnées par les états de la file d'attente dans Figure 1.3. La taille des ordres limite et d'annulation semble être assez stable pour différentes tailles de file d'attente (sauf pour les très petites files d'attente), tandis que la taille moyenne des ordres marchés augmente clairement avec la taille de la file d'attente.

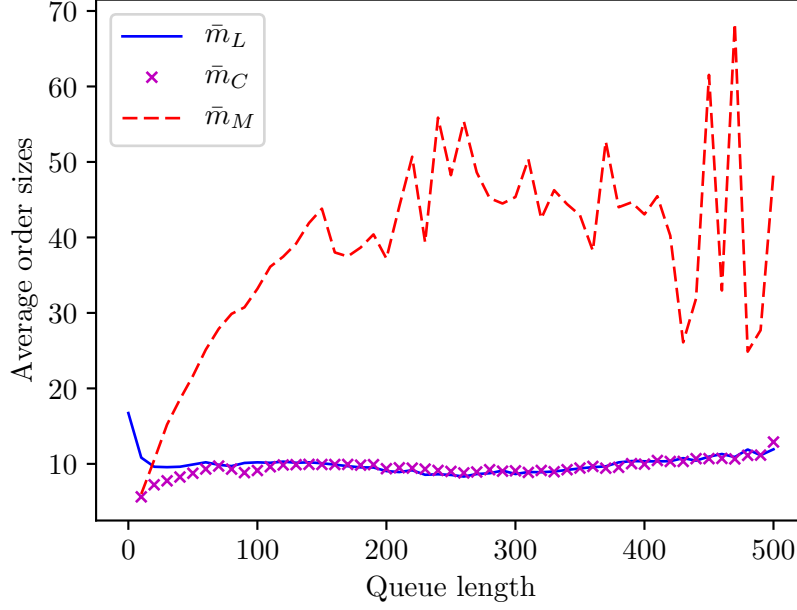


Figure 1.3: Queue reactive model mean order sizes

Une étude plus fine montre que les distributions empiriques pour les ordres limites et les annulations sont assez similaires et peuvent être modélisées avec une distribution géométrique. Au contraire, il existe des schémas intéressants dans les tailles des ordres de marché : la distribution ressemble à un mélange d'une distribution géométrique pour les petites tailles, et fonction Dirac correspondant à 50, 100 ... et la totalité de la file d'attente. Inspiré par ces observations, un modèle simple de tailles d'ordres peut être proposé :

- Tailles des ordres limites : lois géométriques $p_L(q; Q)$ de paramètres $p_0^L(Q)$;
- Tailles des annulations : lois géométriques tronquées de paramètres $p_C(q; Q)$;

$$p_C(q; Q) = \mathbb{P}[q|Q] = \frac{p_0^C(1 - p_0^C)^{q-1}}{1 - (1 - p_0^C)^Q} \mathbb{1}_{\{q \leq Q\}}$$

- Tailles des ordres marché : une mélange de lois géométriques et fonctions Dirac

$$p_M(q; Q) = \mathbb{P}[q|Q] = \theta_0 \frac{p_0^M(1 - p_0^M)^{q-1}}{1 - (1 - p_0^M)^Q} \mathbb{1}_{\{q \leq Q\}} + \sum_{k=1}^{\lfloor \frac{Q-1}{5} \rfloor} \theta_k \mathbb{1}_{\{q=5k+1\}} + \theta_\infty \mathbb{1}_{\{q=Q, Q \neq 5n+1\}}$$

dont les paramètres $\{p_0^M, \theta_0, \theta_k, \theta_\infty\}$ dépendent du Q .

Les paramètres peuvent être calibrés par la maximisation de la vraisemblance. Les résultats calibrés s'avèrent stables sur différentes longueurs de file d'attente

L'autre nouveauté de notre modèle est d'étudier l'influence du dernier événement consommant totalement la liquidité d'une limite sur l'évolution du carnet. Un tel ordre est appelé un *ordre d'élimination*. Il ne peut qu'être une annulation ou un ordre marché. De manière symétrique, un ordre qui crée une nouvelle limite sera appelé *ordre d'établissement*. Un tel ordre ne peut être qu'un ordre limite, mais il peut être placé du côté de l'achat ou de la vente car la limite est vide. Un couple d'ordre d'élimination et d'établissement peut résulter à un changement du prix

ou bien le prix revient pareil qu'avant. Pour le premier, nous l'appelons un suivi du prix, ou effet momentum, qui peut arriver par exemple quand il y a un ordre marché d'achat qui consomme totalement la meilleure limite de vente suivi par un ordre limite d'achat. Pour le dernier, nous l'appelons "reverte".

Les résultats principaux sont:

- Le type de l'ordre d'établissement dépend fortement du type et de la quantité du dernier ordre d'élimination. un ordre marché incite au prochain ordre limite de suivre ce changement du prix, alors qu'une annulation est souvent suivi par un ordre limite qui reverte. En plus un ordre marché de grande taille augmente d'autant plus cette probabilité d'effet momentum. Pour une annulation, la probabilité qu'elle soit grande est très faible, et l'influence de la taille sur le type de l'ordre d'établissement est relativement faible.
- La taille de l'ordre d'établissement est positivement corrélée à la taille de l'ordre d'élimination. Les ordres d'éliminations sont de petites tailles quand ils sont des annulations ou quand l'ordre de l'établissement reverte le prix. Donc il n'est intéressant que d'étudier quand l'ordre d'élimination est un ordre marché et l'ordre d'établissement suit le changement du prix. Dans cette situation, nous trouvons une très forte croissance de la quantité de l'ordre d'établissement quand l'ordre marché est plus grand.
- Le temps d'arrivé de l'ordre d'établissement en suivant l'ordre d'élimination a des propriétés particulières aussi. En suivant une annulation, la distribution du temps est indépendant de la taille et répartie de très haute fréquence vers quelques secondes. La situation est similaire pour un ordre marché de petite taille quand l'ordre d'établissement reverte. Un ordre marché grand suivi par un ordre d'établissement qui suit le changement du prix, a une distribution concentrée en dessous de $1ms$, et a une pique autour de $200\mu s$. Ça peut s'expliquer par la limitation de la latence aller-retour du marché.

Figure 1.4 présente les distributions des meilleures quantités limites échantillonnées à la fréquence 1s dans les différents modèles, ainsi que dans les données. Modèle 0 est le modèle queue-réactive initial, Modèle I modélise la taille d'ordres qui dépend de l'état du carnet, et Modèle II intègre l'influence du dernier ordre d'élimination dans le flux d'ordres. Le modèle 0 produit un LOB qui est très concentré autour de la taille limite d'équilibre constructive-destructrice (environ 300), avec une densité plus faible pour les files d'attente plus petites et presque aucune densité pour les longues files d'attente. L'ajout de la distribution de taille dans le Modèle I améliore déjà le modèle, avec une forme globale plus proche des données réelles. Le Modèle II améliore la densité trop élevée autour de la taille de la file d'attente d'équilibre et produit également des queues plus réalistes et plus grosses pour la distribution de la taille de la file d'attente.

1.3.2 Market making au marché réel

Suite aux précédentes analyses de la dynamique du carnet d'ordres, nous nous intéressons maintenant au problème de market making dans ce modèle de carnet d'ordres. Nous évaluons la qualité du modèle par simulation et backtest dans un vrai environnement du marché. Une 'stratégie optimale' perdante indique une mauvaise modélisation.

Nous reprenons l'approche dans [Hult and Kiessling \[2010\]](#) (voir aussi [Abergel et al. \[2017\]](#) [Bäuerle and Rieder \[2011\]](#) pour une approche similaire et les résultats). Bien que la plupart des travaux précédents assume qu'un market maker gagne de l'argent en moyenne systématiquement et ils

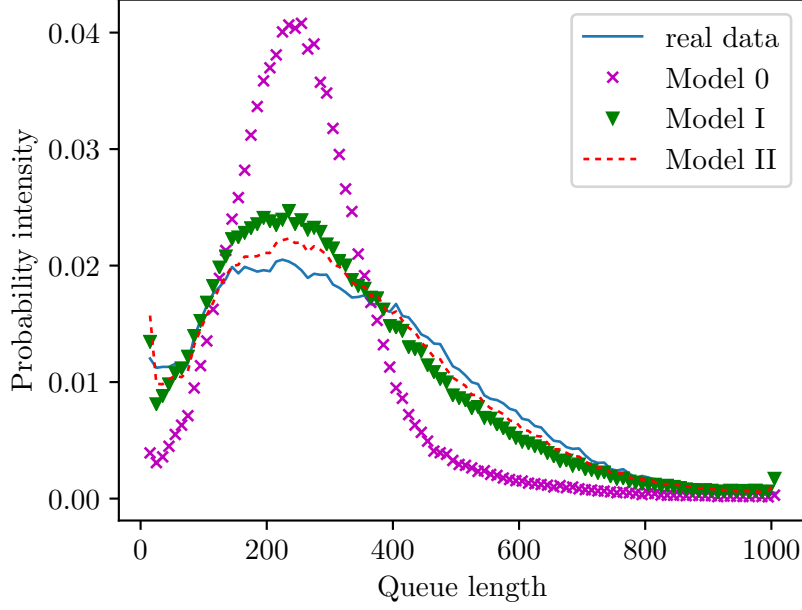


Figure 1.4: Best bid/ask quantities distributions.

ont pour objectif de diminuer la variance du profit. En pratique la stratégie de market making est très très difficile à faire du profit. Nous considérons un market maker *risque-neutre* qui a une fonction valeur le ‘profit & loss’ (P&L). Un pair d’ordre d’achat et de vente sera placé si la fonction valeur est strictement supérieure à 0, et aucun ordre est sur le marché sinon. D’ailleurs nous montrons aussi qu’une stratégie naïve de placer un ordre d’achat et un ordre de vente à tous les états du carnet a une espérance négative dans notre modèle.

L’état du carnet est décrit par un quadruple (x^B, x^A, y^B, y^A) représentant les quantités de la meilleur limite d’achat, de vente et les positions du pair d’ordres du market maker dans la file d’attente. Pour entrer dans le marché, le market maker place ses ordres à la fin de la limite.

$$x^B = y^B \quad x^A = y^A.$$

La fonction valeur des états initiaux (quand le market maker n’ont pas encore d’ordres dans le carnet) dans le Model 0 est illustré dans Figure 1.5. L’axe x et y sont les quantités de la limite d’achat et de vente respectivement. Il est clair que les valeurs sont tous positives et au moins supérieure à 0.5, c’est-à-dire les ordres du market maker sont très profitables peu importe l’état du carnet. Donc la stratégie optimale est la stratégie naïve de placer les ordres à tout moment.

Figure 1.6 montre les valeurs (à gauche) et décisions (à droite) en fonction de l’état initial dans le Model II. Clairement le résultat est très différent par rapport à ce de Model 0: dans la plupart des états initiaux, la décision optimale est de ne pas entrer dans le carnet. En effet, le market maker est pénalisé si le prix change avant que les deux ordres sont exécutés. L’existence de gros ordres marchés, qui font suivre le prix, augmente largement ce risque. La situation la plus favorable est quand les deux limites sont équilibrées et relativement longues. Dans ce cas-là les ordres gagnent de priorité progressivement avant que le carnet devient non-équilibré et le prix change.

Outre que les états initiaux, nous analysons aussi le Modèle II en fixant deux dimensions du quadruple (x^B, x^A, y^B, y^A) pour comprendre sa dépendance par rapport aux deux autres

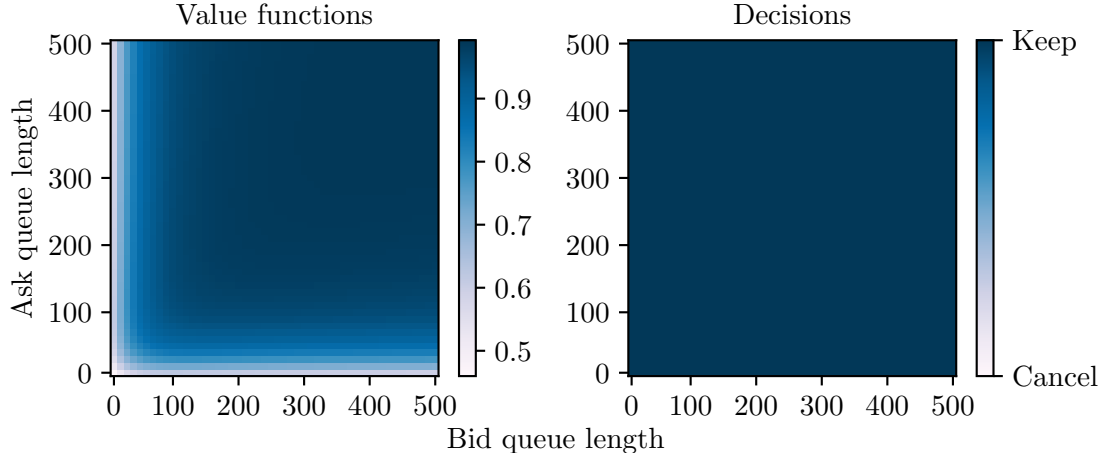


Figure 1.5: Valeurs des états initiaux dans Modèle 0

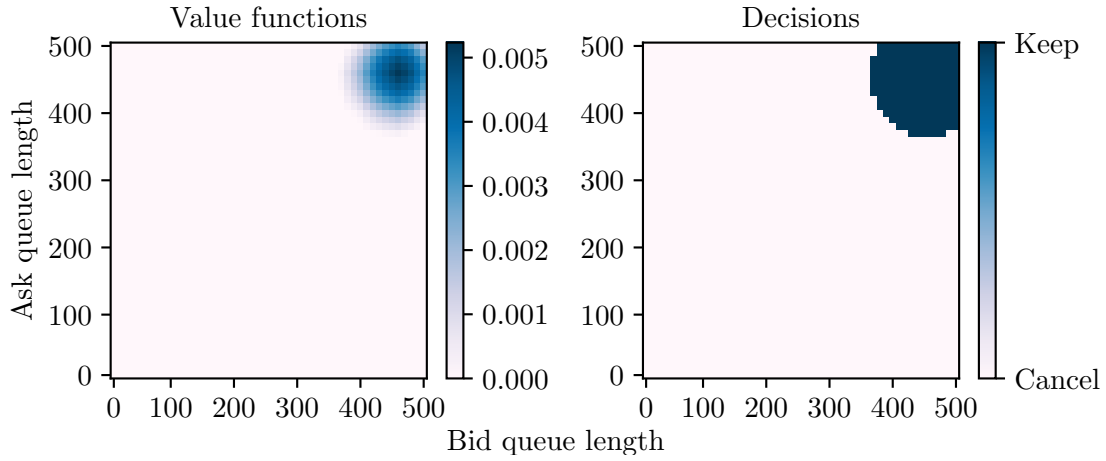


Figure 1.6: Valeurs des états initiaux dans Modèle II

dimensions. Les valeurs associées à chaque état nous permet de comprendre quantitativement l'espérance de placer un pair d'ordre de market making, et donne une décision pour la stratégie. En générale, un ordre qui a plus de priorité dans une file d'attente longue est plus avantageux.

Pour évaluer ces stratégies dans un vrai marché, nous faisons une petite relaxation de la contrainte d'inventaire 1 par rapport à la stratégie optimale. Nous soumettons continument les ordres d'achat et de vente une fois que l'ordre correspondant est exécuté et la fonction valeur est strictement positive. Cette stratégie, que nous appelons 'stratégie localement optimale', nous permet d'évaluer une stratégie continue dans la vraie condition du marché.

Tableaux 1.6 et 1.7 montrent respectivement les résultats de simulation Monte Carlo et backtest avec les données réelles. Le backtester a été développé avec la validation des stratégies qui sont en production dans le marché, donc très proche de la vraie condition avec les données réelles sans modèles. Les deux tests montrent une amélioration significative de la stratégie localement optimale par rapport à la stratégie naïve. Et nous voyons bien que la stratégie naïve est perdante dans ces deux scenarios. En plus nous trouvons que la profitabilité est assez faible même pour la stratégie optimale, qui est en accord avec les expériences pratique de la difficulté du market making.

Table 1.6: Simulation Monte Carlo pour la stratégie localement optimale et naïve

	Optimal	naïve
P&L	1.35 (0.47)	-7.12(0.71)
Inv	3.33 (0.02)	10.47(0.09)
turnover	19.01 (0.06)	113.46(0.22)

Table 1.7: Backtest avec données réelles de la stratégie localement optimale et naïve

	Optimal	Naïve		
		$q_{min} = 0$	$q_{min} = 250$	$q_{min} = 400$
P&L (ke)	0.54	-26.89	-3.43	-2.13
Turn over (Me)	45.20	3430.42	119.25	35.83
Profitability (bp)	0.12	-0.08	-0.29	-0.59

1.3.3 Conclusion

Dans ce chapitre, nous avons étudié un modèle du carnet d'ordre à la [Huang et al. \[2015b\]](#) et son application pour la stratégie de market making. Nous montrons la taille des ordres limites et annulations sont stables pour différentes longueurs de la file d'attente mais les ordres marchés ont des distributions spéciales. De plus, nous trouvons une forte dépendance de l'ordre d'établissement d'une limite par rapport à l'ordre d'élimination. Intégrer ces deux propriétés dans le modèle queue-réactive initial le renforce de manière significative. De plus, le modèle est appliqué pour concevoir la stratégie de market making comparée par rapport à une stratégie naïve, qui est le cas trouvé dans le modèle initial. La surperformance en simulation et en backtest montre que le nouveau modèle est bien plus réaliste et permet d'étudier de meilleures stratégies de market making.

1.4 Prédiction de prix avec des informations de microstructure boostée par deep learning

Market Makers encourent le risque d'inventaire, de l'exécution, et de l'adverse sélection comme indiqué par [Guilbaud and Pham \[2013b\]](#). Ce qui est le plus inquiétant est le risque du dernier qui est directement lié au P&L. Une perte est plus gênant que juste de la volatilité. Pour se protéger de ce risque, Market Makers a besoin de faire les prédictions sur les variations du prix futur selon les informations disponibles. Mais c'est très compliqué à cause de différentes stratégies et de différents horizons des investisseurs dans le marché. De ce fait, les Market Makers n'ont pas le choix que de réfléchir sur

- Utiliser quelles données?
- Prédire quel horizon?
- Quelle valeur à prédire?

Cette étude a pour objectif d'améliorer la prédiction du changement de prix à court terme, de l'étape-par-étape jusqu'à quelques minutes, en utilisant l'information du carnet d'ordres. Elle est particulièrement intéressante pour les Market Makers haute fréquence pour contrôler étroitement

le risque d'inventaire et de résoudre les problèmes de placement d'ordres. C'est un domaine spécialement intéressant et complémentaire pour des prédictions des horizons plus longs. Dans chapitre 4, nous analysons la prédiction du changement de prix à différents horizons avec LSTM. Nous étudions d'abord avec une base de données européennes les propriétés de la *non-linéarité*, la *stationnarité* et l'*universalité* que [Sirignano and Cont \[2018\]](#) rapportent. Ensuite la performance du modèle LSTM pour la prédiction aux horizons plus longs est présentée et discutée. Une stratégie de market making simple est construite à la fin pour illustrer l'apport du LSTM dans un contexte plus réaliste.

1.4.1 Prédiction en 1-étape

Nous nous intéressons uniquement au prix moyen ("mid price"), défini par la moyenne du meilleur prix d'achat et de la vente

$$P = 0.5(P^{b,1} + P^{a,1})$$

tout au long de ce chapitre. Et nous ne considérons que les temps événementiels où il y a eu des changements du P .

Trois types d'indicateurs sont définis:

- **Imbalance.** Ils résument l'information de l'équilibre d'achat et de vente des ordres passifs. Pour une profondeur de k ,

$$OBI_t^k := \frac{\sum_{i=1}^k Q_t^{b,i} - \sum_{i=1}^k Q_t^{a,i}}{\sum_{i=1}^k Q_t^{b,i} + \sum_{i=1}^k Q_t^{a,i}}$$

Pour garder le mémoire de l'indicateur, nous calculons aussi un *Exponential Moving Average (EMA)* avec un paramètre γ comme

$$OBI_t^{ema,\gamma} := \frac{Q_t^{b,ema,\gamma} - Q_t^{a,ema,\gamma}}{Q_t^{b,ema,\gamma} + Q_t^{a,ema,\gamma}}$$

où $Q_{t_i}^{b,ema,\gamma} = (1 - w)Q_{t_{i-1}}^{b,ema} + wQ_{t_i}^{b,1}$ et $Q_{t_i}^{a,ema} = (1 - w)Q_{t_{i-1}}^{a,ema,\gamma} + wQ_{t_i}^{a,1}$ avec $w = \min(1, \frac{t_i - t_{i-1}}{\gamma})$. Le mémoire augmente avec l'augmentation de γ

- **Log Return.** Le prix peut poursuivre son mouvement historique ou revenir vers la moyenne à court terme. Le rendement historique avec k étapes est

$$PR_t^k = \log P_t - \log P_{t-k}$$

- **Order Flow.** Si nous définissons le processus d'ordres marchés par Q_s , qui est un processus signé avec valeurs positives (négatives) représentant les transactions à la limite vente (à l'achat). L'indicateur avec un paramètre de poids β est

$$OF_t^\beta := \sum_{s \leq t} Q_s e^{-\beta(t-s)}$$

La mémoire augmente avec la baisse de β .

20 indicateurs sont utilisés: $OBI^1, OBI^2, OBI^3, OBI^{ema,1}, OBI^{ema,5}, OBI^{ema,10}, OBI^{ema,60}, OBI^{ema,300}, OBI^{ema,600}, PR^1, PR^5, PR^{10}, PR^{20}, PR^{50}, PR^{100}, OF^{0.1}, OF^1, OF^{10}, OF^{100}$,

OF^{1000} . L'idée est d'agréger l'information dans certains indicateurs au lieu de garder tous les états du LOB dans un processus de grande dimension. Pour cette prédiction d'une étape, nous prédisons si $\log P_{t+1} - \log P_t > 0$ ou $\log P_{t+1} - \log P_t < 0$ selon toutes les informations disponibles jusqu'à t .

Environ 100 actions des principaux indices de l'Europe continentale (CAC40, DAX, AEX) sont utilisées dans l'étude. Les données *tick-by-tick* sont acquises pour la période du février 2015 au avril 2018. La base est séparée par avant et après septembre 2017 pour engendrer l'échantillon d'entraînement et l'échantillon de test. Le réseau de LSTM est construit avec 25 neurones cachés, en utilisant les indicateurs de l'étape courante et les indicateurs décalés de jusqu'à 10 étapes. En même temps, 4 modèles linéaires sont proposés pour comparer avec le modèle LSTM: régression OLS, Ridge, LASSO et Logistic.

Le critère de surperformance est la **précision**, qui est la proportion de prédictions correctes divisée par rapport au nombre total d'observations,

$$\frac{\#\{y\hat{y} > 0\}}{\#\{\hat{y} \neq 0\}},$$

Nous observons d'abord qu'en résumant les états du carnet en indicateurs, la précision du modèle LSTM est plus élevée qu'utiliser directement les données brutes dans le modèle. Ceci confirme l'apport de la transformation avec expérience financière et justifie l'utilisation des indicateurs au lieu des données brutes dans le reste de l'étude.

Nos résultats sont en accord avec [Sirignano and Cont \[2018\]](#) sur cette base de données européenne sur deux points

- **Non-linéarité.** Comme présenté dans Figure 1.7, le modèle de LSTM est mieux que tous les modèles linéaires quand ils sont calibrés avec les indicateurs courants et avec les indicateurs décalés par environ 3%. En plus l'ajout des indicateurs décalés améliore plus le modèle de LSTM que des modèles linéaires.
- **Stationnarité.** Nous découpons l'échantillon d'entraînement par 10 périodes en temps, et calibrons un modèle de LSTM pour chaque période. Figure 1.8 présente les résultats de précisions ces modèles, ainsi que le modèle calibré sur l'ensemble des 10 périodes, évalués sur le même échantillon de test. La précision s'augmente très peu avec le rapprochement du temps vers l'échantillon de test, ce qui montre une stationnarité en temps de la relation entre les indicateurs et les variations de prix.

Quant à l'**universalité**, la calibration d'un modèle universelle prend environ deux semaines avec une seule machine. Donc ASGD, un algorithme d'optimisation parallèle, est utilisé pour calibrer le modèle universelle. Cependant, nous trouvons que l'algorithme parallèle dégrade significativement la performance du modèle. Le modèle universelle calibré avec 25 machines est légèrement moins bon que le modèle stock par stock calibré sur une seule machine, mais beaucoup meilleur que le modèle stock par stock calibré avec le même algorithme d'ASGD. Ce résultat montre non seulement l'existence de l'universalité dans notre base de données, mais aussi la nécessité d'un compromis entre le temps de calibration et l'amélioration de précision par le modèle universelle.

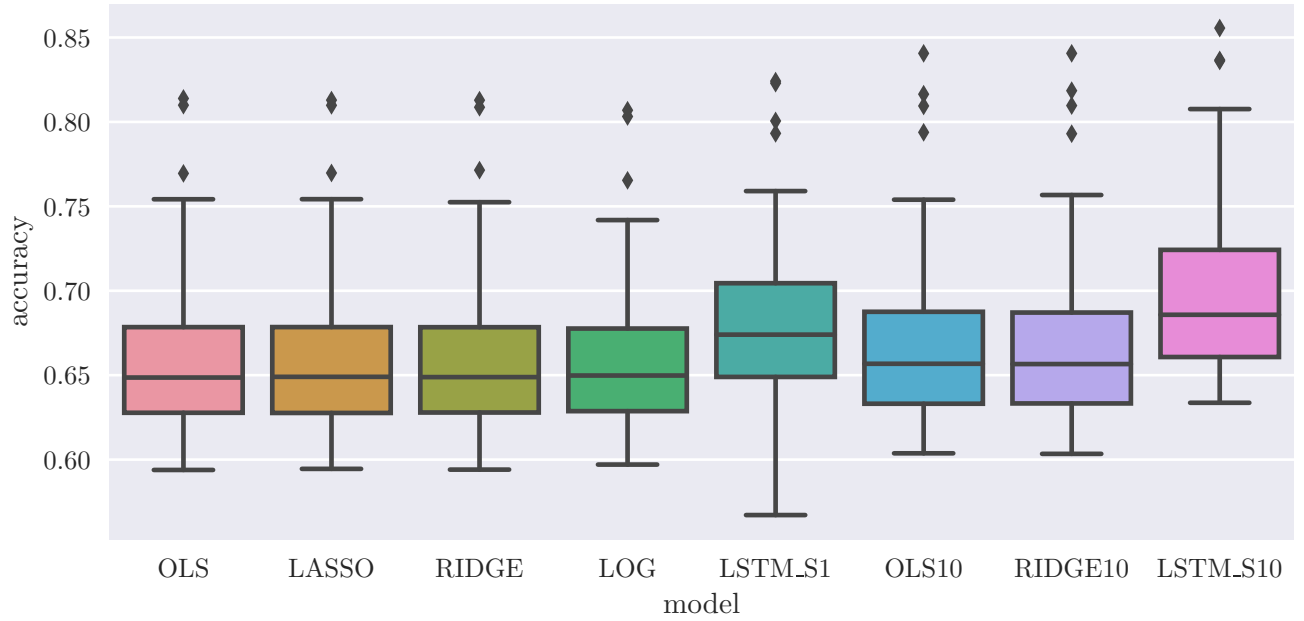


Figure 1.7: Comparaison des précisions de prédictions entre LSTM et modèles linéaires. OLS, LASSO, RIDGE, LOG et LSTM.S1 sont calibrés avec les indicateurs courants, alors qu'OLS10, RIDGE10 et LSTM.S10 sont calibrés avec les indicateurs courants et décalés jusqu'à 10 étapes.

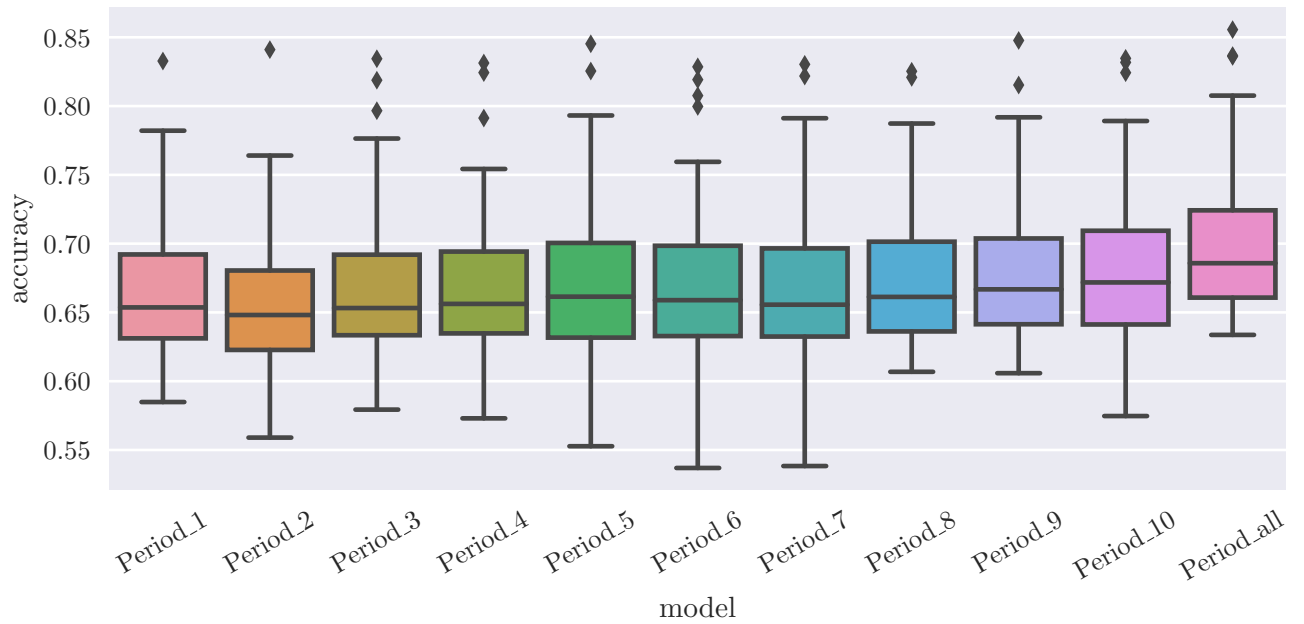


Figure 1.8: Précision de prédictions pour les modèles de LSTM calibré sur des périodes d'échantillon d'entraînement découpées. Period_10 représente le modèle calibré sur la période la plus récente et Period_1 est la plus ancienne.

1.4.2 Régression par deep learning en plusieurs étapes

Le changement du prix mid pour une étape est artificiellement facile à prévoir en raison des autocorrélations négatives élevées de la variation du prix mid. Ce phénomène de retour à la

moyenne résulte de la re-soumission de la liquidité lorsque les premières limites sont consommées ou annulées. Cependant, les market makers ne peuvent pas bénéficier directement de telles prédictions, car elles ne sont pas nécessairement associées aux transactions. Les ordres limites ne sont pas forcément toujours profitables à la présence de l’adverse sélection.

Figure 1.9 montre la profitabilité moyenne des ordres limites mesurée avec le prix mid k -étapes après son exécution. Elle approxime la profitabilité d’un market maker s’il est capable de recouvrir sa position k -étapes après au prix mid. La profitabilité est très élevée immédiatement après la transaction, cependant le market maker ne peut pas avoir une autre transaction opposée à l’horizon tellement court. Elle baisse très rapidement jusqu’à environ 10 étapes, et se stabilise à -0.2 bp. Si en plus nous prenons en compte un cout asymétrique de 0.5 bp pour le “liquidity taker” et 0 bp pour le “liquidity provider”, nous revenons à un cout de transaction à -0.3 bp et -0.2 bp respectivement pour les deux types de participants. Il existe dans le marché un équilibre entre les ordres limites et ordres marchés, avec un léger surcoût pour les ordres marchés car ils bénéficient la certitude d’exécution. Ce résultat illustre la nécessité de prédire pour un horizon plus long.

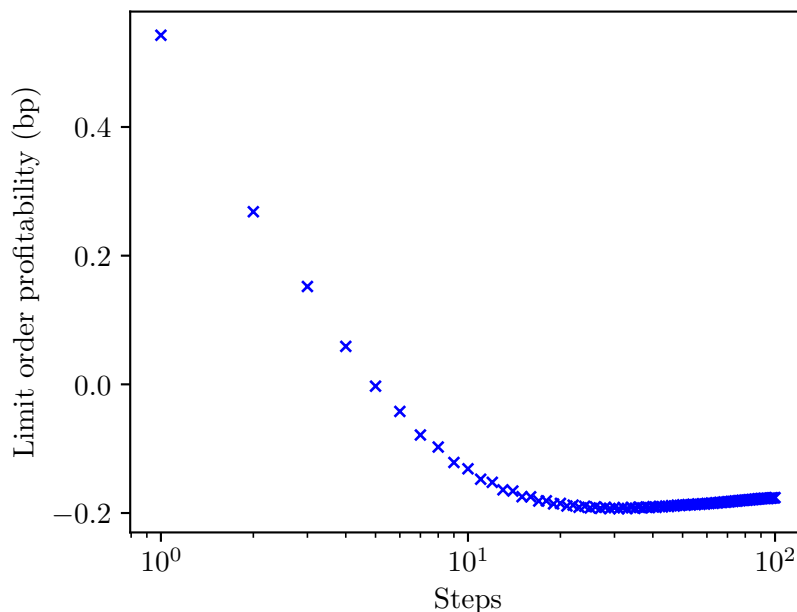


Figure 1.9: La profitabilité des ordres limites exécutées en fonction du prix mid k -étapes après.

Sans surprises, le modèle de LSTM et OLS dégradent beaucoup pour prédire un horizon plus long. Ils donnent une précision de 58% pour les rendements 10-étapes et 52% pour les rendements 100-étapes. La précision du modèle de LSTM devient similaire au modèle linéaire au modèle OLS. Nous illustrons la variation de la précision à différents horizons dans Figure 1.10. Nous voyons une plus grande déviation entre la performance in-sample et out-of-sample pour le modèle LSTM que le modèle OLS. En 100-étapes, le modèle LSTM est mieux in-sample que le modèle OLS alors que la performance out-of-sample est quasiment pareille. Le modèle LSTM subit autant plus la baisse du ratio signal ne bruit que le modèle OLS.

Pour conclure l’application du modèle LSTM pour la pratique, nous construisons une stratégie simple de market making. Nous prenons $\mathbf{sd}(\hat{y}_{train})$ comme un seuil significatif. Une position longue est prise (des stocks sont achetés) si $\hat{y}_t > \mathbf{sd}(\hat{y}_{train})$, et inversement une position short est prise si $\hat{y}_t < -\mathbf{sd}(\hat{y}_{train})$. Les prédictions à chaque étape sont très bruitées, donc nous calculons

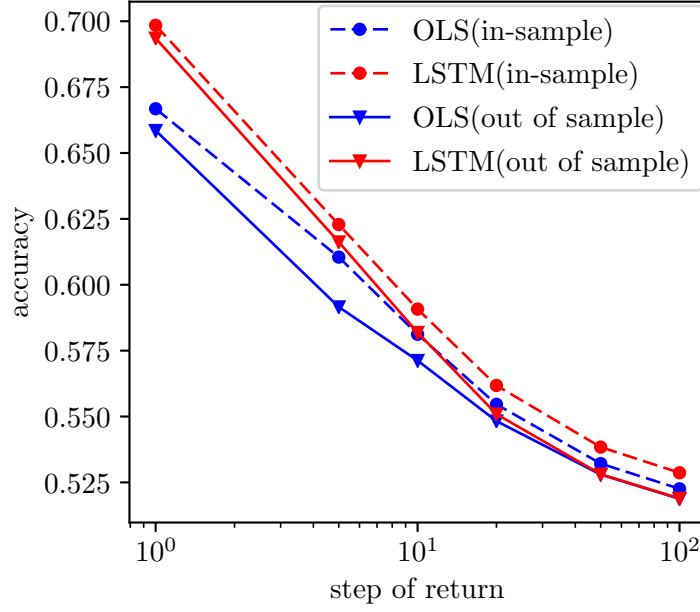


Figure 1.10: In-sample et out-of-sample précisions pour le modèle LSTM et OLS à différents horizons

aussi un signal lissé par EMA avec un paramètre de $\tau = 50$ étapes.

$$\widehat{EMA(y)}_k = \frac{1}{\tau} \hat{y}_k + (1 - \frac{1}{\tau}) \widehat{EMA(y)}_{k-1}$$

Table 1.8: Médiennes des mesures des stratégies

		OLS_1	OLS_10	OLS_100	LSTM_1	LSTM_10	LSTM_100
Signal brut	bp	0.43	0.42	0.40	0.42	0.45	0.36
	HP	1.3	2.0	1.8	1.1	1.5	1.7
	P&L sans coût	894	572	622	1195	824	575
	P&L avec coût	-218	-147	-168	-188	-123	-198
EMA	bp	0.67	0.58	0.60	0.70	0.71	0.57
	HP	17.6	23.9	26.2	10.0	17.7	28.7
	P&L sans coût	97	65	64	172	103	57
	P&L avec coût	17	6	7	33	24	6

Les statistiques des différentes stratégies sont résumées dans le Tableau 1.8. “Signal brut” signifie l’utilisation des prédictions étape par étape sorties directement des modèles, et “EMA” signifie les signaux lissés par la fonction EMA. Trois mesures de performances sont utilisées:

1. **Gain(P&L)**: Le profit ou la perte par jour par stock, un coût de 0.5 bp est appliqué pour le coût de transaction.
2. **Profitability (bp)**: Le P&L par unité de volumes de transactions.
3. **Holding period (HP)**: Durée de portage des positions entre le moment de l’acquisition et de la liquidation.

Le modèle de 100-étapes est moins bon que ceux de 1-étape et 10-étapes. La profitabilité et le P&L sont tous plus faibles. Cohérent avec notre observation pour l'analyse de précision, le modèle de LSTM se dégrade encore plus que le modèle d'OLS. Le bruit augmente de manière proportionnelle à $\sqrt{k}$, alors que le changement prédictible ne croît pas (ou très lentement) avec l'horizon. Comme LSTM performe moins bien dans un contexte de bruit fort, ce bruit perturbe plus le LSTM que l'OLS à horizon long.

Les stratégies générées par les signaux bruits ont des profitabilités faibles et des HPs très courts d'inférieur à 2 changements de prix moyen. Même si le P&L sans coût de transaction est élevé, ce tout réaliste tue totalement la performance à cause de la faible profitabilité. Les stratégies générées par les signaux lissés par EMA réduit largement les bruits des signaux. La profitabilité est renforcée à être capable de battre le coût de transaction de sorte que toutes les stratégies sont rentables après le coût. Et la durée de portage des positions est en moyenne des dizaines de changements du prix moyen. Ces résultats sont plus robuste pour résister l'artefact de réversion court-terme.

Si nous nous concentrons sur la comparaison des modèles LSTM et des modèles OLS avec des signaux lissés par l'EMA, nous trouvons que le premier est nettement supérieur au second pour les prédictions en 1-étape et 10-étapes. Les P&Ls avant et après le coût de transaction sont tous plus élevés pour le modèle LSTM. En particulier, le P&L après le coût de transaction est doublé de 17 à 33 pour 1-étape et quadruplé de 6 à 24 pour 10-étapes. Bien qu'il puisse encore y avoir une certaine influence des durées de portage, comme nous constatons que la durée de portage du modèle LSTM est plus courte que celle du modèle OLS. Mais même en comparant LSTM 10-étapes avec OLS 1-étape dont les durées de portage sont très proches, le P&L après le coût de transaction est toujours 50% supérieur. Nous pouvons conclure que le modèle LSTM surpasse le modèle linéaire pour construire une stratégie de market making aussi simple basée sur la prédiction du changement au prix mid à court terme.

1.4.3 Conclusion

Ce chapitre présente notre étude sur l'application du modèle LSTM pour la prédiction du prix selon les états du carnet d'ordres à court terme. Sur une base de données des stocks européens, nous retrouvons la non-linéarité, la stationnarité et l'universalité de la relation entre les informations du carnet d'ordres et le changement du prix en 1-étape rapporté par [Sirignano and Cont \[2018\]](#) pour des stocks américains. Un défi est aussi constaté que l'algorithme parallèle de l'optimisation pour faire le modèle universelle n'est pas suffisamment puissant, donc un compromis entre le temps de calculs et la précision est incontournable. Ensuite nous montrons l'insuffisance de juste prédire 1 étape de changement de prix, qui est fortement influencé par la réversion. Les précisions de prédiction par le modèle LSTM et le modèle OLS se dégradent avec la prolongation de l'horizon, mais la performance de LSTM décroît encore plus rapidement. Avec une analyse in-sample et out-of-sample, nous montrons que le LSTM a une plus grosse différence entre le résultat in-sample et out-of-sample du au problème d'overfitting dans un contexte de faible ratio signal par rapport au bruit. Des stratégies simples de market making ont été construites et testées pour illustrer l'apport du LSTM et le potentiel de l'application pour les market makers dans les vraies activités de market making.

Introduction to LOB

The core of our study is the limit order book, which is repeated in the three original papers. To avoid such repetition, we present the notation of LOB for this whole thesis here.

In an order-driven market, participants can submit orders of three basic types: limit order, market order and cancellation:

- **Limit order:** An order that specifies an upper/lower price limit (also called “quote”) at which one (commonly called “liquidity provider”) is willing to buy/sell a certain number of shares. The advantage of the limit order is that the transaction price is better than the instantaneous mid-price. However, there is no certainty that the limit order will be executed. Currently most markets adopt the “first in first out” rule, i.e. the priorities of limit orders are decided first according to price, and then to arrival time. A limit order can be entirely, partly or not executed.
- **Market order:** An order that triggers an immediate buy/sell transaction for a certain number of shares at the best available opposite quote(s). The advantage is to offer an immediate execution, however the price is worse than the mid-price. A market order can be executed with different limit orders as counterparties. The price is not necessarily the best limit price, if the quantity is big enough that the order eats up completely the first limit and hits the second or higher limits.
- **Cancellation :** An order that removes an existing limit order.

In addition to these three main types of orders, there exist various order services provided by the exchanges such as “stop loss”, “good til’ canceled”... Also note that some markets allow orders, such as “iceberg” orders, to provide hidden liquidity, making their presence difficult to infer from the order flow. Nevertheless, it is commonly agreed upon - and verified in practice - that the basic orders carry enough information for market microstructure studies.

An example is given in Figure 1.11.

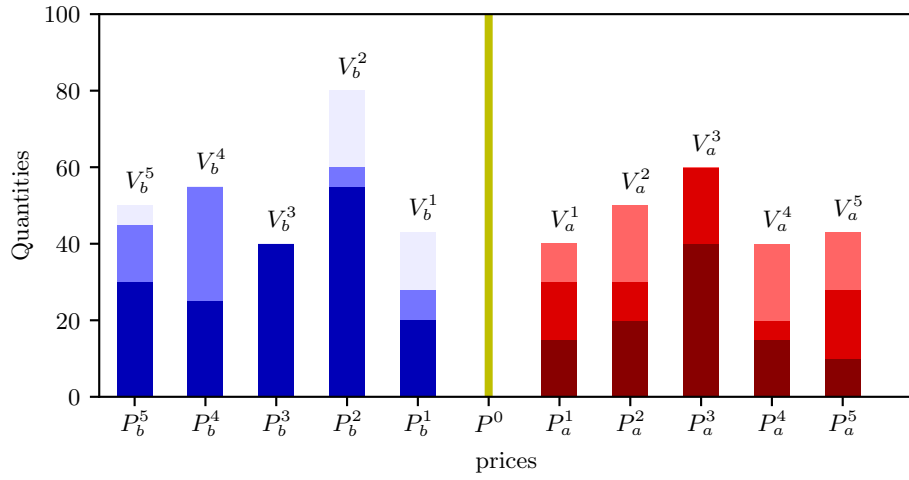


Figure 1.11: Illustrative order book example. Blue bars on the left half of the figure represent the available buy orders with prices P_b^i and total quantities V_b^i . These correspond to the buyer side, also called the bid side. Participants in the bid side are providing liquidity with prices at which they are ready to buy some quantities of the stock. The right hand bars represent the sell side, commonly called the ask side or offer side, where participants willing to sell post their orders with the prices they are ready to sell the stock. The line in the middle corresponds to the mid-price level and is computed as the average between the best (highest) bid price and the best (lowest) ask price. A transaction occurs when a sell order and a buy order are at least partially matched. A queue of limit orders with the same price is called a limit. Different colors in the same limit represent orders with different priority with darker bars having higher execution priority.

Chapter 2

Limit order book modelling with high dimensional Hawkes processes

Note:

- This chapter is published in “Quantitative Finance”.
- This chapter is presented in the conference “Market Microstructure: confronting many viewpoints”, Paris, December 2016.
- This chapter is presented in the workshop “Portfolio dynamics and limit order books”, CentraleSupélec, Châtenay-Malabry, December 2016.
- This chapter is presented in the workshop “ASC2017: Asymptotic Statistics and Computations”, University of Tokyo, Tokyo, January 2017.

Abstract

High-dimensional Hawkes processes with exponential kernels are used to describe limit order books in order-driven financial markets. The dependences between orders of various types are carefully studied and modelled, based on a thorough empirical analysis. The observation of inhibition effects is particularly interesting, and leads us to the use of non-linear Hawkes processes. A specific attention is devoted to the calibration problem, in order to account for the high dimensionality of the problem and the very poor convexity properties of the MLE. Our analyses show a good agreement between the statistical properties of order book data and those of the model.

Contents

2.1	Introduction	26
2.2	Empirical findings: the interdependencies of order book events	27
2.2.1	Data and Framework	27
2.2.2	Event definitions	27
2.2.3	Statistical dependencies between order book events	27
2.3	Modelling dependencies using Hawkes processes	30
2.3.1	Linear Hawkes process models	32
2.3.2	Performances of the linear Hawkes models	33
2.3.3	Nonlinear Hawkes process model	35
2.3.4	Performances of the nonlinear Hawkes models	37
2.4	Some numerical aspects of model calibration	42
2.4.1	Calibration with maximum likelihood estimation	42
2.4.2	Benchmarking the DE algorithm	44
2.4.3	Improvement in high dimensions	44
2.5	Conclusion	48

2.1 Introduction

Limit order books have attracted a considerable amount of attention since the electronification of financial markets in the early 90s. The historical quote-driven markets, where designated market makers used to provide liquidity to all participants, have largely evolved into order-driven markets, where buy and sell orders are matched continuously in a double auction queueing system.

Limit order books have been extensively studied, both from empirical and mathematical points of view, see e.g. [Abergel et al. \[2016\]](#) for a survey of their properties. In particular, the mathematical modelling of limit order books is itself an active research area that has many useful and practical applications, and this paper is a contribution to the field.

A particularly popular class of order book models is that of Markovian models, starting with the so-called *zero-intelligence* models as in [Smith et al. \[2003\]](#), then enriched with more complex and realistic contributions such as [Cont et al. \[2010\]](#); [Cont and de Larrard \[2013\]](#); [Radivojević et al. \[2014\]](#); [Scalas et al. \[2017\]](#); [Huang et al. \[2015b\]](#). In Markovian models the order flows are described by point processes with state-dependent conditional intensities.

More to the point, many empirical studies have identified some memory properties of financial markets. To name a few, [Gopikrishnan et al. \[2000\]](#); [Bouchaud et al. \[2009\]](#) underline a significant positive autocorrelation and slow decay of the trade flow. [Chakraborti et al. \[2011\]](#); [Scalas et al. \[2006\]](#) confirm that the Poisson hypothesis for the arrival of orders is not empirically satisfied, whereas [Eisler et al. \[2012\]](#) is an in-depth study of the correlation between, and price impact of, orders of all types. These findings advocate for a direct modelling of the temporal dependencies between order arrivals. As a consequence, Hawkes processes have come up as a natural modelling choice, and triggered a lot of interest in recent studies on market microstructure and limit order books. [Bacry and Muzy \[2014\]](#); [Bacry et al. \[2013, 2016\]](#) propose various models of price and order flow models. An extensive survey of the application of Hawkes process in finance can be found in [Bacry et al. \[2015\]](#). In the specific context of limit order books, [Large \[2007\]](#) is an early study of Hawkes processes applied to order book modelling, Hawkes-process-based limit order book models are introduced and mathematically investigated in [Abergel and Jedidi \[2015\]](#); [Zheng et al. \[2014\]](#) and, in a slightly different direction, [Rambaldi et al. \[2017\]](#) models the order volumes - in addition to their types - based on a multivariate Hawkes process.

This paper is a contribution to this latter strand of research.

In most papers involving Hawkes processes for order book modelling, the natural quantities of interest are the inter-event durations - or : inter-event forward recurrence times. They will be the main objects under scrutiny in the present work as well.

The quality of various Hawkes-process-based order book models will be assessed using some objective criteria: a model will be deemed satisfactory if it can reproduce as many as possible of the stylized facts of financial data. Our approach starts with a precise empirical analysis of the dependencies between order arrivals of various types. Then, models built from multivariate, possibly nonlinear, Hawkes processes with multiple exponential kernels are introduced. Once a model is designed, it is evaluated: the distribution of forward recurrence times, as well as the signature plot¹, are used as selection criteria. With this approach, we are able to discriminate between various Hawkes-process-based models, and provide a financial interpretation of the more successful ones in terms of their behaviour at various time scales, and the presence of inhibition

¹A characterization of the realized price volatility at various frequencies.

as well as excitation effects.

The paper is organized as follows: Section 2.2 presents our most relevant empirical findings, laying the ground for the modelling based on linear and nonlinear Hawkes processes discussed in Section 2.3. In Section 2.4, some numerical aspects of model calibration are discussed in detail. Finally, some concluding remarks are presented in Section 2.5.

2.2 Empirical findings: the interdependences of order book events

In this section, we present our main empirical findings on the dependencies between order arrivals. These findings pave the way for the modelling avenues followed in the next sections.

2.2.1 Data and Framework

This paper focuses on the DAX listed 30 stocks trading in XETRA - the electronic trading venue of the Frankfurt Stock Exchange. Three months (February to April 2016) of tick-by-tick data are used in this study. The data consist in the list of all trades and order book states any time a modification or a transaction occurs - with a resolution of $1\mu s$ ($10^{-6}s$). As is classical for high frequency financial data, see e.g. [Muni Toke \[2017\]](#) for a recent survey on order book reconstruction, some data cleaning was involved in order to identify limit orders, market orders and cancellations given the states of the order book and the list of trades.

Due to the large quantity of data, problems such as mismatches of quantities and lack of synchronization were expected. However, such anomalies represent less than 3% of the data, and our results are thus reliable.

2.2.2 Event definitions

In this study, any change that modifies the best limits of the order book is called an “event”². More precisely, an event can be a limit order, a market order, or a cancellation, and can affect the best bid or best ask. Moreover, events will be tagged according to whether they change the mid-price or not. Table 2.1 summarizes the definitions and notations for the various event types considered in this paper.

2.2.3 Statistical dependencies between order book events

Table 2.2 represents the empirical probabilities of occurrence of an event of type j (in column), **conditioned** on the fact that the last observed event is of type i (in row). The last row represents the unconditional probabilities of each type of events.

To simplify the interpretation of the results, Table 2.3 represents the ratio of conditional probabilities to unconditional probabilities, rounded to one decimal. It aims at revealing the mutual relationships between events, and ratios greater than two are highlighted.

²This simplifying choice essentially means that a level-1 order book is considered.

Table 2.1: Event types definitions.

Notation	Definition
M, L, C, O	market order, limit order, cancellation, any order.
M^0, L^0, C^0, O^0	market order, limit order, cancellation, any order, that does not change the mid-price.
M^1, L^1, C^1, O^1	market order, limit order, cancellation, any order, that changes the mid-price.
M_{buy}, M_{sell}	buy/sell market order.
M_{buy}^0, M_{sell}^0	buy/sell market order that does not change the mid-price: i.e. order quantity < best ask/bid available quantity.
M_{buy}^1, M_{sell}^1	buy/sell market order that changes the mid-price: ie. order quantity $\geq$ best ask/bid available quantity.
L_{buy}, L_{sell}	buy/sell limit order.
L_{buy}^0, L_{sell}^0	buy/sell limit order that does not change the mid-price: i.e. order price $\leq / \geq$ best bid/ask price.
L_{buy}^1, L_{sell}^1	buy/sell limit order that changes the mid-price: ie. order price $> / <$ best bid/ask price.
C_{buy}, C_{sell}	buy/sell cancellation.
C_{buy}^0, C_{sell}^0	buy/sell cancellation that does not change the mid-price: i.e. partial cancellation at best bid/ask limit or cancellation at another limit.
C_{buy}^1, C_{sell}^1	buy/sell cancellation that changes the mid-price: ie. total cancellation of best bid/ask limit order.

Table 2.2: Conditional probabilities of occurrences per event type (in %).

	L_{buy}^0	L_{sell}^0	C_{buy}^0	C_{sell}^0	M_{buy}^0	M_{sell}^0	L_{buy}^1	L_{sell}^1	C_{buy}^1	C_{sell}^1	M_{buy}^1	M_{sell}^1
L_{buy}^0	27.94	9.73	26.24	20.99	1.09	0.50	4.98	1.78	0.08	3.60	2.65	0.44
L_{sell}^0	9.63	28.43	20.36	26.50	0.58	1.01	1.83	4.94	3.55	0.08	0.43	2.68
C_{buy}^0	21.64	24.72	29.33	7.51	0.66	0.57	2.92	4.28	4.97	1.32	0.97	1.11
C_{sell}^0	24.04	21.97	7.45	29.82	0.64	0.58	4.30	2.85	1.32	5.00	1.08	0.95
M_{buy}^0	20.69	8.02	6.96	11.18	9.08	0.59	9.64	0.86	1.12	6.65	24.27	0.94
M_{sell}^0	7.19	20.72	10.18	6.90	0.64	9.63	0.72	9.39	6.67	1.10	0.93	25.92
L_{buy}^1	32.48	10.83	1.17	26.57	0.92	0.94	4.38	1.68	8.89	4.93	2.43	4.77
L_{sell}^1	10.24	33.61	26.27	1.14	0.94	0.90	1.71	4.32	4.88	8.97	4.54	2.47
C_{buy}^1	14.46	12.40	51.27	4.59	0.26	0.10	8.42	4.06	2.83	0.96	0.51	0.15
C_{sell}^1	11.85	14.60	4.32	52.25	0.10	0.22	3.96	8.41	0.91	2.77	0.14	0.48
M_{buy}^1	12.23	6.18	4.56	30.18	1.04	0.64	24.94	4.39	1.16	8.70	3.35	2.64
M_{sell}^1	5.93	12.67	29.80	4.63	0.71	1.09	4.36	24.68	8.88	1.13	2.59	3.52
O	19.93	20.42	20.23	20.74	0.79	0.73	4.02	3.99	2.93	2.95	1.62	1.65

Table 2.3: Conditional probability leverage.

	L_{buy}^0	L_{sell}^0	C_{buy}^0	C_{sell}^0	M_{buy}^0	M_{sell}^0	L_{buy}^1	L_{sell}^1	C_{buy}^1	C_{sell}^1	M_{buy}^1	M_{sell}^1
L_{buy}^0	1.4	0.5	1.3	1.0	1.4	0.7	1.2	0.4	0.0	1.2	1.6	0.3
L_{sell}^0	0.5	1.4	1.0	1.3	0.7	1.4	0.5	1.2	1.2	0.0	0.3	1.6
C_{buy}^0	1.1	1.2	1.4	0.4	0.8	0.8	0.7	1.1	1.7	0.4	0.6	0.7
C_{sell}^0	1.2	1.1	0.4	1.4	0.8	0.8	1.1	0.7	0.5	1.7	0.7	0.6
M_{buy}^0	1.0	0.4	0.3	0.5	11.5	0.8	2.4	0.2	0.4	2.3	15.0	0.6
M_{sell}^0	0.4	1.0	0.5	0.3	0.8	13.2	0.2	2.4	2.3	0.4	0.6	15.7
L_{buy}^1	1.6	0.5	0.1	1.3	1.2	1.3	1.1	0.4	3.0	1.7	1.5	2.9
L_{sell}^1	0.5	1.6	1.3	0.1	1.2	1.2	0.4	1.1	1.7	3.0	2.8	1.5
C_{buy}^1	0.7	0.6	2.5	0.2	0.3	0.1	2.1	1.0	1.0	0.3	0.3	0.1
C_{sell}^1	0.6	0.7	0.2	2.5	0.1	0.3	1.0	2.1	0.3	0.9	0.1	0.3
M_{buy}^1	0.6	0.3	0.2	1.5	1.3	0.9	6.2	1.1	0.4	2.9	2.1	1.6
M_{sell}^1	0.3	0.6	1.5	0.2	0.9	1.5	1.1	6.2	3.0	0.4	1.6	2.1

Results of Table 2.3 are quite symmetric and no significant differences are observed between the buy and the sell side. Therefore, only the buy side is interpreted in detail below:

- L_{buy}^0 : adds liquidity to the first limit, signalling an increase of market demand at the current price level. This stimulates L_{buy}^1 and M_{buy}^1 events, based on the new consensus for a higher price. The corresponding probabilities for orders of type 0 are also increased, based on a similar reasoning but in a less aggressive way. On the other hand, the selling activity decreases in general except for C_{sell}^1 , because some newly added limit orders may be cancelled shortly after. One notable thing is the sharp decrease of C_{buy}^1 , as the newly added limit order probably comes from another trader, making it very unlikely that the first limit should be cancelled.
- C_{buy}^0 : decreases liquidity on the buy side. It triggers successive cancellations C_{buy}^0 and C_{buy}^1 : cancellations tend to follow themselves. M_{buy}^0 , M_{sell}^0 , M_{buy}^1 and M_{sell}^1 become less likely, revealing the influence of low liquidity on the participants' willingness to generate executions.
- M_{buy}^0 : largely increases the probability of M_{buy}^0 and M_{buy}^1 . This is commonly attributed to order splitting and the momentum effect (other participants following the move). L_{buy}^1 and C_{sell}^1 are also stimulated as a new price consensus emerges.
- L_{buy}^1 : improves the offered price to buy. The first effect is a strong increase in the probability of M_{sell}^1 , i.e., participants entirely consume the new liquidity as the offered price has become higher. The second effect is an increase in the probability of C_{buy}^1 , i.e., the new liquidity is rapidly cancelled. This is consistent with a similar observation made for L_{buy}^0 orders, and might reflect some sort of market manipulation where agents are posting fake orders. Not surprisingly, the conditional probability of C_{buy}^0 is almost zero, because after one limit order is submitted, it cannot be partially cancelled. The fact that probability is not exactly 0 may be due to poor data synchronization, or the existence of hidden liquidity.
- C_{buy}^1 : a total cancellation of the best buy limit increases the probability of C_{buy}^0 - order cancellations come in succession as market makers lose interest to provide liquidity even at the new best limit - and that of L_{buy}^1 events, as traders may re-offer at the previous best price to gain priority. Events of other types become less frequent.
- M_{buy}^1 : consumes all the offered liquidity at the best ask. It stimulates L_{buy}^1 and C_{sell}^1 as a higher price consensus emerges among market participants. The probability of M_{buy}^1 increases, indicating a short term momentum effect, and order splitting.

As a conclusion to this empirical section, let us just say that strong temporal dependencies between events are identified. Some orders actually triggers other events, a fact that can be seen as self- or cross-excitation phenomena. There are also some inhibition effects, when incoming orders prevent other events to occur. These two important features will be the target of the modelling approach presented in the next section.

2.3 Modelling dependencies using Hawkes processes

It has now become widely accepted in the high frequency and market microstructure community that limit order books are worth modelling, and that the price dynamics can easily be extracted

from that of the order book. In fact, the complexity of inter-event dependencies is so rich that most significant features of the price dynamics : co-existence of time scales, leverage effect, signature plot, long term diffusivity... can be derived from advanced order book models.

In this section, point-process-based order book models are studied, building on the 12 event types previously introduced: $E = \{L_{buy}^0, L_{sell}^0, C_{buy}^0, C_{sell}^0, M_{buy}^0, M_{sell}^0, L_{buy}^1, L_{sell}^1, C_{buy}^1, C_{sell}^1, M_{buy}^1, M_{sell}^1\}$.

Recall that events with superscript 1 have an instantaneous price impact: in particular, events in $E_{up} = \{L_{buy}^1, C_{sell}^1, M_{buy}^1\}$ lead to a price increase, whereas those in $E_{down} = \{L_{sell}^1, C_{buy}^1, M_{sell}^1\}$ result in price decrease.

The arrival of order book events is modelled by a 12-variate simple point process $N(t) = (N_{L_{buy}^0}(t), \dots, N_{M_{sell}^1}(t))$. Of interest is the associated intensity process $(\lambda_{L_{buy}^0}(t), \dots, \lambda_{M_{sell}^1}(t))$.

Assuming that the process is simple means that two events cannot occur at the same time, a fairly realistic assumption due to the high time resolution of modern stock exchanges.

Since the focus in this paper is on temporal interdependencies, N is actually modelled as a 12-variate **counting process**, and the marks determining the price jump when an event of type 1 occurs are not modelled. Rather, a simplifying assumption is made, namely, that the jump of the best bid or ask price following an event of type 1 is always one tick. This approximation reduces the dimensionality of the point process, while being consistent with the real behaviour of the chosen data set, for which the average jump size of the best bid and ask prices is 1.08 ticks³. Under this assumption, the reconstructed mid-price dynamics easily obtains as a by-product of event arrivals:

$$S(t) = S(0) + \left(\sum_{e \in E_{up}} N_e(t) - \sum_{e' \in E_{down}} N_{e'}(t) \right) \times \frac{\eta}{2}, \quad t > 0$$

where $\eta > 0$ is the tick size.

This simplification will be taken into account when comparing the performances of the model with the behaviour of real data.

Events of type 0 do not directly influence the price, rather, their impact will come from their influence on the intensities of the type 1 event arrival process.

As already said in the introduction, it has long been recognized that the class of Hawkes processes is particularly well suited to the modelling of point processes interacting *via* their conditional intensities. Here, we build on the results of Section 2.2 and study two classes models respectively based on linear and nonlinear Hawkes processes that capture well the main characteristics of market dynamics.

The performances of the models are presented in this section, while some more technical aspects pertaining to their calibration are deferred until Section 2.4.

³Actually, for some large tick stocks, the average is even smaller than 1.01.

2.3.1 Linear Hawkes process models

In this short paragraph, we recall some essential definitions and results on the particularly interesting class of point processes introduced in [Hawkes and Oakes \[1974\]](#). We refer the interested readers to [Brémaud and Massoulié \[1996\]](#); [Massoulié \[1998\]](#) for a more in-depth presentation of these processes, and to [Zheng et al. \[2014\]](#); [Abergel and Jedidi \[2015\]](#) for their use in order book modelling.

A multivariate point process $((T_i, X_i))_{i \in \mathbb{N}_*}$, associated to a counting process $(N(t))_{t \in \mathbb{R}_+} = (N_1(t), \dots, N_M(t))_{t \in \mathbb{R}_+}$ with conditional intensity process $(\lambda(t))_{t \in \mathbb{R}_+} = (\lambda_1(t), \dots, \lambda_M(t))_{t \in \mathbb{R}_+}$, is called a (linear, multivariate) **Hawkes process** ° [Hawkes and Oakes \[1974\]](#); [Massoulié \[1998\]](#) if there holds for $m \in \{1, \dots, M\}$:

$$\lambda_m(t) = \mu_m + \sum_{n=1}^M \int_0^t \phi_{mn}(t-s) dN_n(s)$$

where μ_m are positive real numbers and ϕ_{mn} are nonnegative functions.

The μ_m are the *base intensities* and can be viewed as background intensities. Whenever an event occurs, the intensities increase, making subsequent events arrive at a higher frequency. Such effects are controlled by ϕ_{mn} . The functions ϕ_{mn} , the *kernel functions*, control the instantaneous increases and the relaxation speeds of the intensities in response to excitations.

For a multivariate Hawkes process, ϕ_{mm} describe the self-excitations, while ϕ_{mn} for $m \neq n$ measure the cross- (or: mutual) excitations, that is, the impact of an event of type n on the arrival of an event of type m .

A convenient, alternate way to express the intensity process is provided by the following equation:

$$\lambda(t) = \mu + \Phi \star dN \tag{2.1}$$

where $\Phi(t)$ is the $M \times M$ matrix whose entries are $\phi_{mn}(t)$, “ $\star$ ” denotes the “matrix convolution”

$$\Phi \star dN = \int_{\mathbb{R}} \phi(t-s) dN(s)$$

and $\phi(t-s)dN(s)$ stands for the standard matrix-vector product.

Hawkes processes are fully determined by their baseline intensity μ and the matrix Φ of kernel functions. In the following, we will concentrate on **exponential** kernels. This particular choice is classical, one of its main advantages being the Markovianity of the joint process (N, λ) , see e.g. [Massoulié \[1998\]](#). For the models considered in this work, the intensities follow Equation (2.1), where Φ is a 12×12 kernel function matrix describing the excitation between events of various types:

$$\Phi = (\phi_{ij})_{i,j \in E}.$$

What we call *1-exponential* and *2-exponential* Hawkes models differ by the number of exponential functions used to define each kernel, namely:

- For the 1-exponential Hawkes model, $\phi_{ij}(t) = \alpha_{ij} \exp(-\beta_{ij}t)$
- For the 2-exponential Hawkes model, $\phi_{ij}(t) = \sum_{p=1}^2 \alpha_{ijp} \exp(-\beta_{ijp}t)$

2.3.2 Performances of the linear Hawkes models

The adequacy of a linear Hawkes-process-based order book model is now evaluated, according to two criteria: a **goodness-of-fit** criterion for the distribution of forward recurrence times, and a criterion based on the **signature plot** generated by the model.

As a matter of fact, it is generally agreed upon that such statistical properties of the price process as the unconditional distribution of returns or the diffusive behaviour at large time scale, can easily be reproduced even with simpler models, whereas the signature plot and the inter-event durations offer a better challenge to discriminate among order book models.

2.3.2.1 Goodness of fit

It is well-known, see e.g. [Bowsher \[2007\]](#) that the transformed durations $\{\tau\}_i$ of a Hawkes process

$$\tau_i^m = \int_{T_i}^{T_{i+1}} \lambda^m(s) ds$$

are i.i.d. exponential random variables with parameter 1, where T_i are the event arrival times. This property is used to test the goodness-of-fit of the model to the data, by drawing Q-Q plots of the empirical quantiles with respect to the theoretical exponential distribution quantiles.

Though a global test can be conducted by concatenating all the transformed durations, plotting each dimension separately provides more information. This can be viewed as a marginal distribution fit test, i.e.: Given the law of other types of orders, how well can we fit the order under scrutiny?

The procedure is as follows: first, the parameters for several order book models (Poisson, 1-exponential linear Hawkes, 2-exponential linear Hawkes) are calibrated using maximum likelihood method (More details about the calibration are discussed in section 2.4), for each day in the study period. Then the transformed durations in the model are computed, and a Q-Q plot test is then performed. The results are shown in Figure 2.1.

As a first conclusion, one can easily see that a Poisson-process-based model globally fails to capture the distributional properties of recurrence times. The performances of the 1- and 2-exponential Hawkes models are similar, except for orders of type 0: the 2-exponential model significantly outperforms the 1-exponential model for L^0 and, to a lesser extent, for C^0 events. However, what is annoying is the behaviour for C^1 events: the distributions of the transformed durations in 1- and 2-exponential models are extremely close to one another, but neither is close to the theoretical exponential distribution.

This is an important, negative feature of the linear Hawkes models that will be revisited in the upcoming Subsection 2.3.3

2.3.2.2 Signature plots

The *signature plot* is a plot of the realized variance as a function of the sampling frequencies. It reveals some of the most important stylized facts about high frequency financial data.

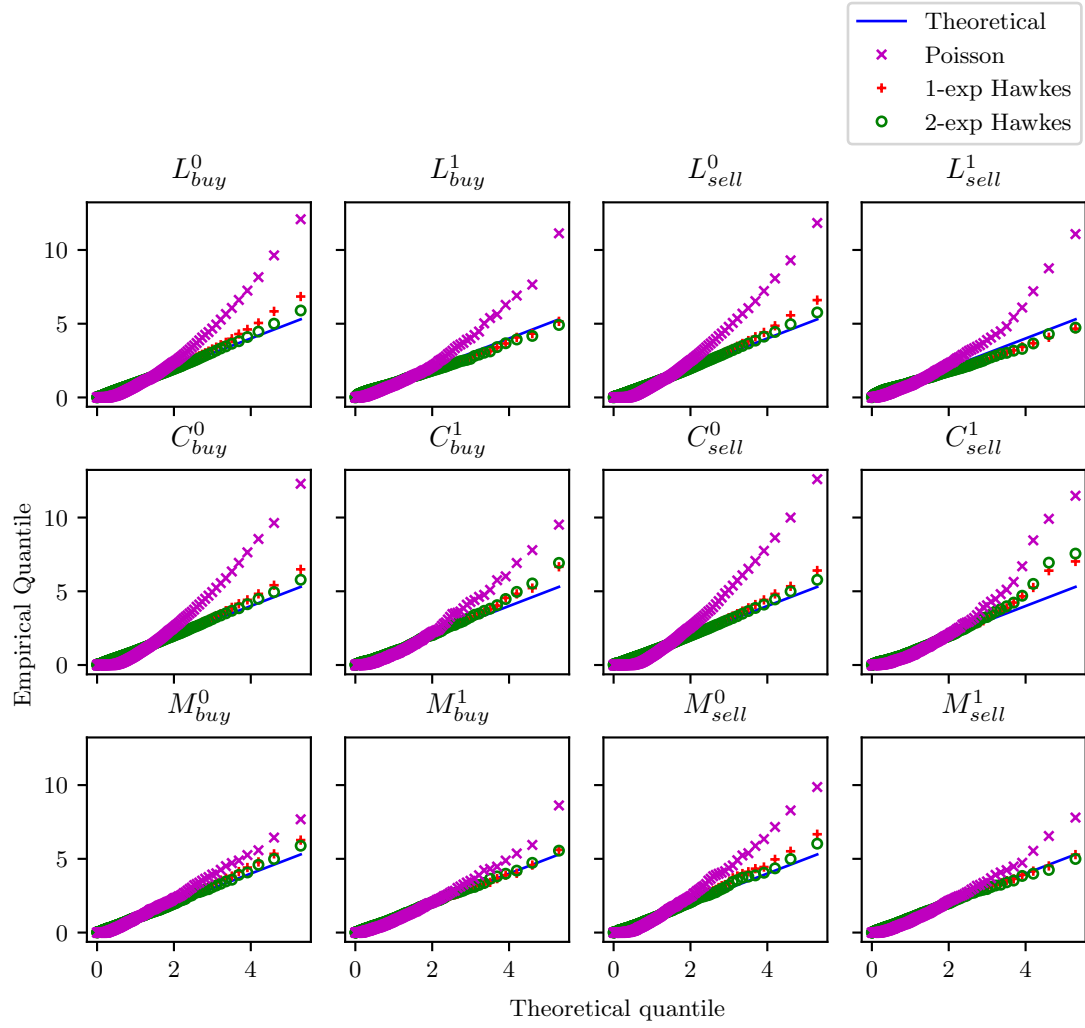


Figure 2.1: Q-Q plot goodness of fit tests of order book models.

The realized variance for a stochastic process X_t over a time period $[0, T]$ at a sampling frequency h is

$$RV(h) = \frac{1}{T} \sum_{n=0}^{T/h} (X((n+1)h) - X(nh))^2. \quad (2.2)$$

An important stylized fact of financial markets is that the quantity RV generally increases when h becomes small. This phenomenon is associated to the mean reverting behaviour of the price at short time scales. It has long been observed and was already reported in Andersen et al. [2000]. It is noteworthy that the signature plot becomes even steeper when computed on transaction prices rather than mid-prices because of the *bid-ask bounce*, which describes the phenomena that transactions can alternate between bid side and ask side even if both prices do not change. We will focus on mid-prices to avoid this spurious effect.

Once the model parameters are calibrated, the mid-price is easily simulated using Equation (2.3). Realized variances are calculated with sampling periods from 1 to 50 seconds, with a step of 1 second.

The results for the models and the real data are shown in Figure 2.2.

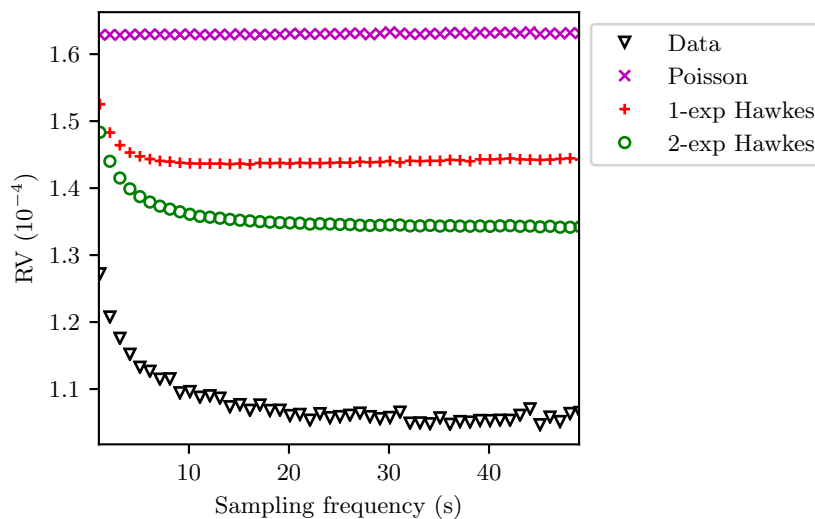


Figure 2.2: Mean signature plots of simulated price compared with real data.

Not surprisingly, the signature plot of the Poisson model is flat - this is expected, as the price dynamics in this model is that of a mid-price model with Poisson jumps, due to the mapping of orders that increase (resp. decrease) the price into upward (resp. downward) jumps.

The 1-exponential and 2-exponential Hawkes process models (more exponential kernels are similar to that of 2-exponential kernel, the improvement by increasing the number of exponentials in kernel functions is limited) behave similarly: the realized volatility decreases when the sampling interval increases, but the long-term volatility level is too high compared to the data. Though reproducing the overall shape of the signature plot, the linear Hawkes-process-based order book models are not satisfactory.

Table 2.4: Conditional probability comparison between simulated order flows and real data.

Pair	P_{simu}	P_{real}
$C_{buy}^1 L_{buy}^0$	0.402	0.048
$L_{buy}^1 L_{buy}^1$	1.628	0.141
$L_{sell}^1 L_{buy}^1$	1.288	0.171
$M_{sell}^0 C_{buy}^1$	0.545	0.068
$C_{buy}^1 C_{buy}^1$	0.548	0.072
$M_{sell}^1 C_{buy}^1$	0.854	0.037

2.3.3 Nonlinear Hawkes process model

This subsection addresses the shortcomings of linear Hawkes models in reproducing some characteristics of forward recurrence times and signature plots. Nonlinear Hawkes processes are introduced to overcome these difficulties, and their performances are studied.

2.3.3.1 Order dependencies: inconsistencies between real data and linear Hawkes models

The results presented in Paragraph 2.3.2.1 are now revisited in *event time*, temporarily ignoring the durations. When comparing the average conditional probability matrix of the 2-exponential Hawkes model with that of real data, one can check that most of the conditional probabilities are pretty close. However, for several pairs, there exist huge differences between the model and the real data, in particular for $C_{buy}^1|L_{buy}^0$, $M_{buy}^0|C_{buy}^1$ and $M_{sell}^0|C_{buy}^1$.

Table 2.4 below gives the list of all pairs (X, Y) for which the probability of an event of type X , conditioned on following an event of type Y , in the simulated order flow is either smaller than 50% or greater than 5 times the real conditional probability (only the buy side is shown, the sell side behaves similarly). The probabilities are calculated on 10000 independently simulated paths of 8 hours.

From a financial point of view, these discrepancies can easily be accounted for:

- A $C_{buy}^1|L_{buy}^0$ sequence almost never happens, because L_{buy}^0 is a limit order added to the current first limit and it is highly unlikely that two orders should be cancelled at the same microsecond.
- The low probabilities of $L_{buy}^1|L_{buy}^1$ and $L_{sell}^1|L_{buy}^1$ comes from the constraint of the bid-ask spread: an aggressive limit order decreases the spread, and when the spread becomes one tick wide, other price-changing limit orders are no longer possible.
- The remaining cases correspond to orders following a C_{buy}^1 order that increases the bid-ask spread. There is no physical constraint preventing the spread from being wide, but participants in the market are not seemingly ready to sell when a cancellation order has already decreased the best bid price.

From a mathematical point of view, this poor fit comes from an inherent shortcoming of the

Table 2.5: Medians of L_1 norm of kernels $\phi_{\cdot C_{buy}^1}$ in the 2-exponential model.

L_{buy}^0	L_{sell}^0	C_{buy}^0	C_{sell}^0	M_{buy}^0	M_{sell}^0	L_{buy}^1	L_{sell}^1	C_{buy}^1	C_{sell}^1	M_{buy}^1	M_{sell}^1
0.1563	0.2357	0.9392	0.0914	0	0	0.3845	0.1607	0	0	0.0013	0

linear Hawkes process model: the intensity for the arrival of an order of type e is written as

$$\lambda_e(t) = \mu_e + \sum_{e' \in E} \sum_{T_{e'} < t} \phi_{ee'}(t - T_{e'})$$

where $\phi_{e'} \geq 0$ and, by construction, $\lambda_e(t) < \mu_e$ cannot happen ! Consequently, **inhibition** effects, leading to a temporary decrease of certain short term conditional probabilities, are not modeled.

Note that, when calibrating the linear model (see Section 2.4 for details), the kernels corresponding to inhibitory behaviours are indeed forced to 0.

Below are the median values of the L_1 norms of the kernels stemming from the calibration results for C_{buy}^1 stimulations in Table 2.5: kernels corresponding to the event pairs listed in Table 2.4 have norms equal to 0.

Moreover, two other event pairs come out of the calibration with 0 kernel norms, $M_{buy}^0|C_{buy}^1$ and $C_{sell}^1|C_{buy}^1$. Although less obvious from the conditional probability matrix, this phenomenon is easy to interpret: a defensive cancellation on the bid side indicates a consensus of a fair price decrease in the market, therefore traders are less willing to buy at the previous ask price or cancel an existing ask order as it has already gained some queue priority with a profitable price.

2.3.3.2 Model definition

In order to incorporate inhibitory behaviours in the model, negative kernels are introduced in the Hawkes process. Then, a truncation is applied to avoid meaningless negative process intensities.

In the new model, the intensities satisfy the equation

$$\lambda(t) = (\mu + \Phi \star dN)_+, \quad (2.3)$$

where the entries of the matrix Φ are no longer constrained to take on positive values, and $()_+$ denotes the elementwise positive part function. With the function, intensities are no longer bounder by the base intensity, and can reach 0 if the inhibition effect is strong. Still, most of the time the intensities are positive and the intensity is the same as a linear Hawkes process. So we can benefit from the simpleness of process simulation and likelihood calculation.

When enriched with the nonlinearity, the 2-exponential Hawkes process model retains its Markovian nature, see e.g. Brémaud and Massoulié [1996]; Zhu [2015] for general results on nonlinear Hawkes processes. The negative kernels are chosen under the following form

$$\phi_{mn} = \sum_{p=1}^2 -\alpha_{mnp} \exp(-\beta_{mnp}t)$$

where the α s and β s are nonnegative real numbers. Note that we fix the same sign for the two exponentials, in order to avoid overfitting - it is actually unexpected for interdependencies

to have different time regimes, for example an inhibitory effect in the short term that would become an excitation in the long run.

2.3.4 Performances of the nonlinear Hawkes models

2.3.4.1 Goodness of fit

Similarly to the analysis presented in Paragraph 2.3.2.1, the Q-Q plots of the empirical quantiles with respect to those of the theoretical exponential distribution are shown on Figure 2.3.

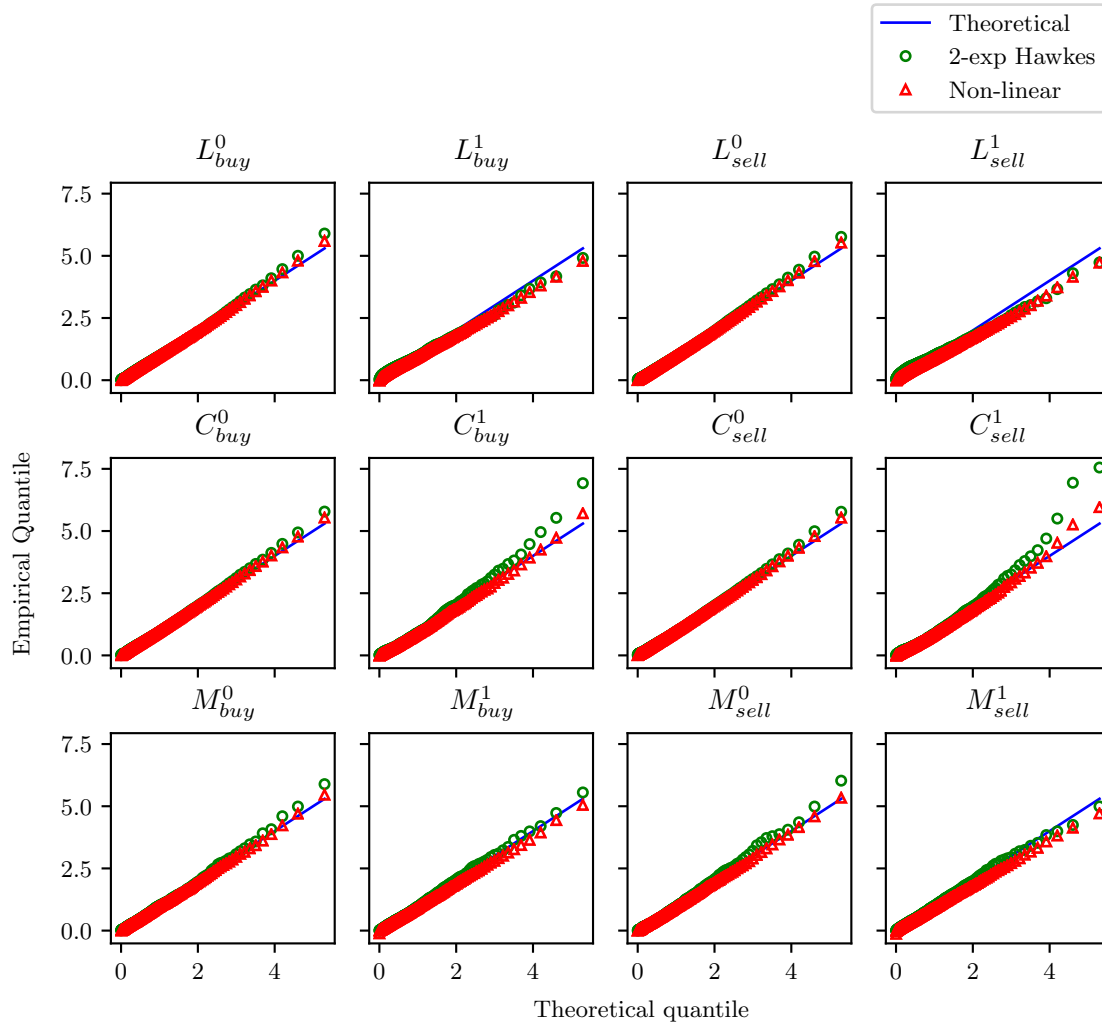


Figure 2.3: Q-Q plot goodness of fit tests of nonlinear Hawkes model.

It appears, simply by eyeballing the graphs, that the nonlinear Hawkes model leads to a statistically more satisfactory fit than the linear 2-exponential Hawkes model previously studied. This better performance will be confirmed by the analysis of the signature plots and forward recurrence times.

In Tables 2.6 and 2.7, the calibration results of the nonlinear model are compared to those of the 2-exponential linear model. The conditional probabilities are closer to real data, showing that a definite improvement is achieved by the nonlinear model. Also note that the kernels that

Table 2.6: Conditional probability amelioration of Non-linear Hawkes model.

Pair	P_{simu}	P_{simu}^{NL} ^a	P_{real}
$C_{buy}^1 L_{buy}^0$	0.402	0.105	0.048
$L_{buy}^1 L_{buy}^1$	1.628	0.254	0.141
$L_{sell}^1 L_{buy}^1$	1.288	0.235	0.171
$M_{sell}^0 C_{buy}^1$	0.545	0.135	0.068
$C_{buy}^1 C_{buy}^1$	0.548	0.163	0.072
$M_{sell}^1 C_{buy}^1$	0.854	0.108	0.037

^aConditional probabilities in nonlinear model.

Table 2.7: Medians of L_1 norm of kernels $\phi_{C_{buy}^1}$ in nonlinear model for those that were 0 in 2-exponential model.

M_{buy}^0	M_{sell}^0	C_{buy}^1	C_{sell}^1	M_{sell}^1
-0.0319	-0.1593	-0.0541	-0.1439	-0.1908

were formerly set to 0 now take on quite significant negative norms, a fact which confirms that the inhibition effect plays an important role in the order dynamics.

Another interesting insight is provided in Table 2.8, by comparing the optimal values of the log-likelihood functions for both models. One can actually see where inhibition effects become more pregnant: orders of type 1 are more influenced than orders of type 0, confirming the improvement already observed for the Q-Q plots in the goodness-of-fit test.

2.3.4.2 Signature plots

The signature plots of linear and nonlinear 2-exponential Hawkes models are shown in Figure 2.4, and compared to that of real data. The asymptotic volatility level significantly improves with the nonlinear model, and the resulting signature plot is overall a very good fit.

Table 2.8: Median of optimal likelihood functions for each type of order in 2-exponential and non-linear models.

Model	L_{buy}^0	L_{sell}^0	C_{buy}^0	C_{sell}^0	M_{buy}^0	M_{sell}^0
2-exp Hawkes	12862.2	14122.8	20018.4	21821.2	-2693.1	-1932.0
Non-linear Hawkes	12862.2	14122.8	20018.4	21821.2	-2584.9	-1784.6
	L_{buy}^1	L_{sell}^1	C_{buy}^1	C_{sell}^1	M_{buy}^1	M_{sell}^1
2-exp Hawkes	-1415.8	-1379.8	-1025.0	-1125.0	-1307.4	-1162.2
Non-linear Hawkes	-962.6	-999.4	-638.6	-803.0	-1099.0	-995.1

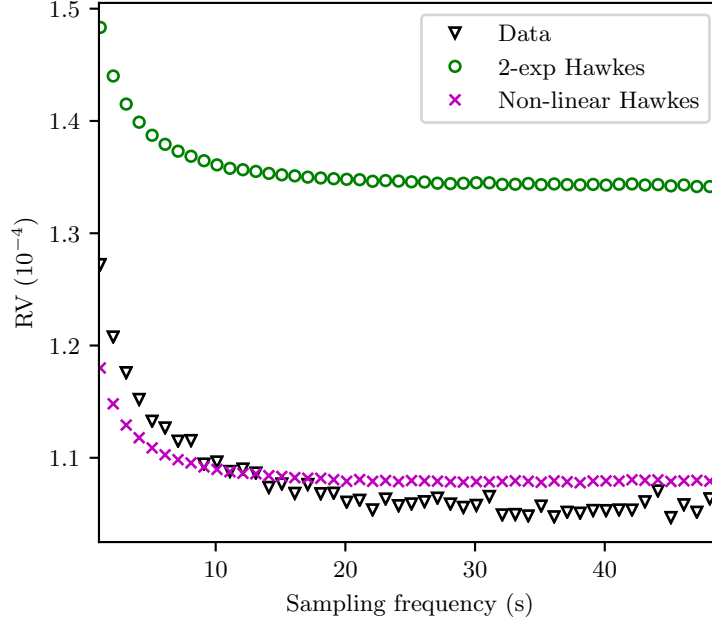


Figure 2.4: Mean signature plots of linear and nonlinear 2-exponential Hawkes process compared with real data.

2.3.4.3 Analysis of self- and cross-excitation recurrence times

The rationale behind the introduction of nonlinear Hawkes models was the empirically observed presence of inhibitory effects among events. As a consequence, one should hope that the inter-event recurrence times would behave in a more realistic way with these models.

Figure 2.5 and 2.6 show the cumulative distribution function (CDF) and the probability density of the (logarithm of) the forward recurrence times for all events of type 1 - that is, the forward recurrence times of (or: duration between) price jumps.

Specifically, define the inter-jump duration as

$$\Delta T_i = T_{i+1} - T_i$$

where T_i are the timestamps of the event arrivals.

According to the type of event causing the jump, these durations are furthermore separated into two subgroups: *self-excitation durations* $\Delta T^a \in \{\Delta T_i | X_i = X_{i+1}\}$ and *cross-excitation durations* $\Delta T^c \in \{\Delta T_i | X_i \neq X_{i+1}\}$.

Inter-jump durations predicted by the model are then computed, and compared to data: although the linear Hawkes model already performs well in reproducing the inter-jump duration distributions both for self- and cross-excitations, one can see that the nonlinear Hawkes process further improves the fit in the range between milliseconds and seconds ($\log_{10}(\Delta T) \in (-3, 1)$).

As a conclusion, one can say that the nonlinear Hawkes model provides a very satisfactory enhancement to the classical one, whether one uses Q-Q plots, signature plots or inter-jump recurrence times as benchmarks. This improvement is in fact quite natural, and is related to the empirical evidence presented in 2.3.3.1 on inhibition effects between events.

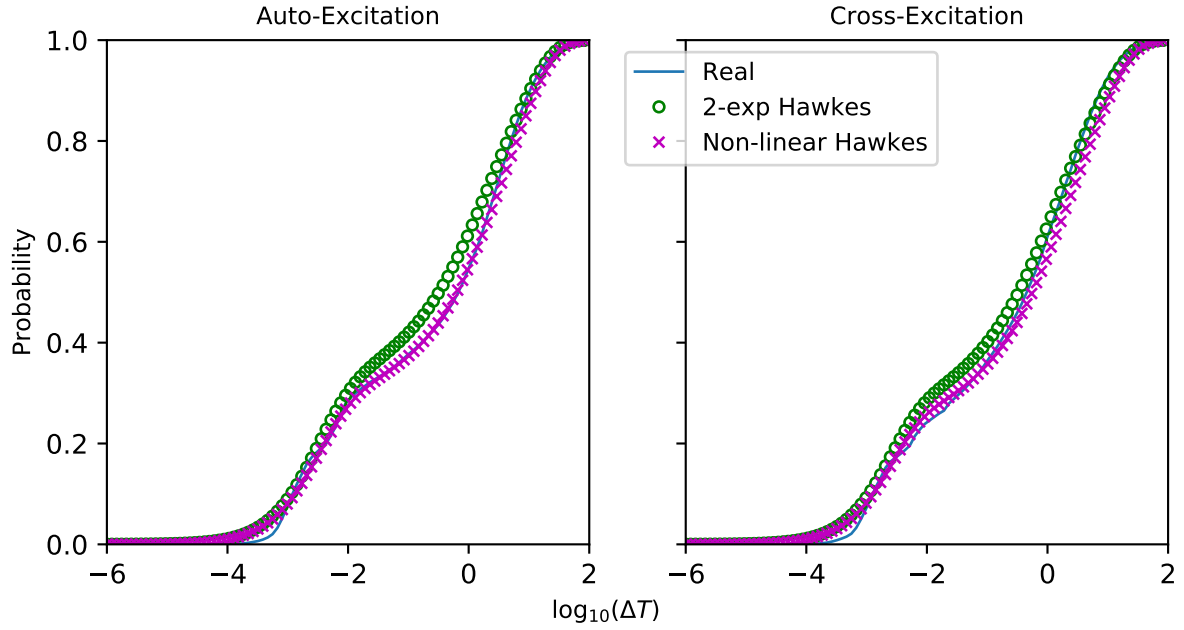


Figure 2.5: Cumulative distribution functions of log inter-jump durations for simulated price processes compared with real data.

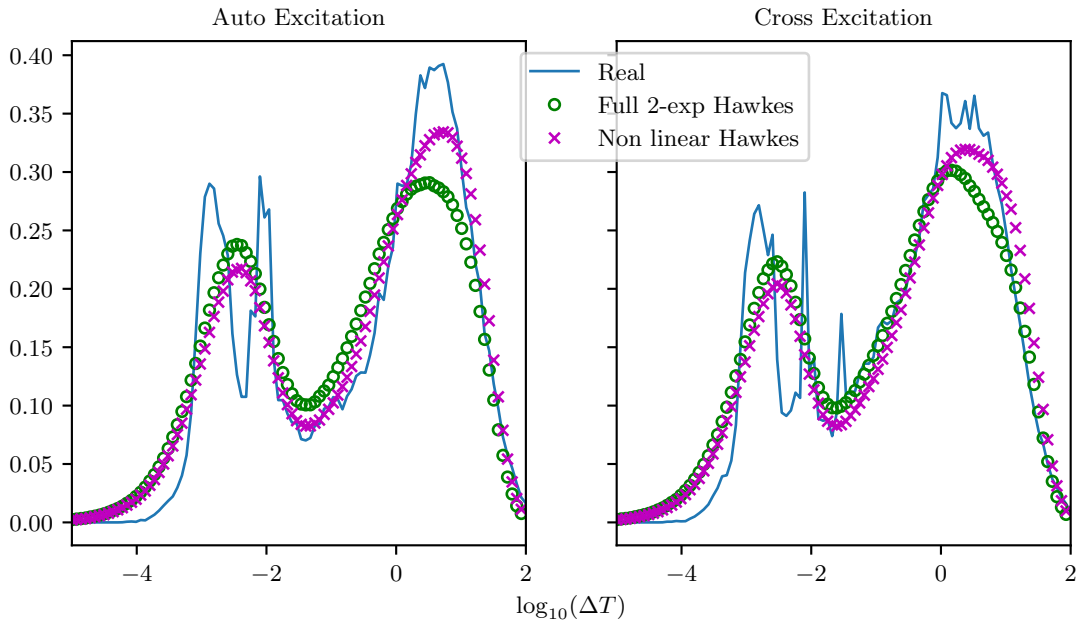


Figure 2.6: Probability density of log inter-jump durations.

2.4 Some numerical aspects of model calibration

This section is devoted to an analysis of the numerical algorithms used to calibrate the various models introduced in Section 2.3. Although rather technical, we think it is relevant - actually, very useful - for readers interested in calibrating high-dimensional Hawkes-processes to high

frequency financial data (or other types of data).

Several optimization procedures are discussed and compared, and the best performer among those we have tested is thoroughly investigated.

2.4.1 Calibration with maximum likelihood estimation

Let $((T_i, X_i))_{i \in \mathbb{N}^*}$ be a multivariate point process with associated counting process $(N_1(t), \dots, N_M(t))$, whose intensities are to be estimated.

The log-likelihood, see [Ozaki \[1979\]](#); [Rubin \[1972\]](#), of given intensities $(\lambda_1(t), \dots, \lambda_M(t))$, and a sample of observation $\{T_i, X_i\}_{i \in \{1, \dots, M\}}$, is defined by the sum of the log-likelihood of each component:

$$\begin{aligned} \ln L(\lambda, \{T_i, X_i\}_{i \in \{1, \dots, D\}}) &= \sum_m \ln L_m(\lambda_m, \{T_i, X_i\}_{i \leq D}) \\ &= \sum_{m=1}^M \left[\int_0^T \ln \lambda_m(s) dN_m(s) + \int_0^T (-\lambda_m(s)) ds \right]. \end{aligned}$$

In the case of a Hawkes process with exponential kernels, a straightforward computation gives:

$$\int_0^T \ln \lambda_m(s) dN_m(s) = \sum_{T_i^m} \ln \left[\mu_m + \sum_{n=1}^M \alpha_{mn} A_{mn}(i) \right]$$

and

$$\int_0^T \lambda_m(s) ds = \mu_m T - \sum_{n=1}^M \sum_{T_k^n} \frac{\alpha_{mn}}{\beta_{mn}} \left(e^{-\beta_{mn}(T - T_k^n)} - 1 \right),$$

where $A_{mn}(i) = \sum_{T_k^n < T_i^m} e^{-\beta_{mn}(T_i^m - T_k^n)}$ can be computed iteratively as

$$A_{mn}(i) = A_{mn}(i-1) e^{-\beta_{mn}(T_i^m - T_{i-1}^m)} + \sum_{T_{i-1}^m \leq T_k^n < T_i^m} e^{-\beta_{mn}(T_i^m - T_k^n)}$$

so that

$$\ln L_m(\lambda_m, \{T_i, X_i\}_{i \leq D}) = -\mu_m T + \sum_{n=1}^M \sum_{T_k^n} \frac{\alpha_{mn}}{\beta_{mn}} \left(e^{-\beta_{mn}(T - T_k^n)} - 1 \right) + \sum_{T_i^m} \ln \left[\mu_m + \sum_{n=1}^M \alpha_{mn} A_{mn}(i) \right].$$

It is however unfortunate that the likelihood function is not strictly concave. For example, in the 1-dimensional case, its expression simplifies to

$$\ln L(\lambda, \{T\}) = -\mu T + \sum_{T_i} \frac{\alpha}{\beta} \left(e^{-\beta(T_D - T_i)} - 1 \right) + \sum_{T_i} \ln \left[\mu + \alpha \sum_{T_j < T_i} e^{-\beta(T_i - T_j)} \right],$$

and, letting β tend to ∞ , there holds

$$\lim_{\beta \rightarrow +\infty} \ln L(\lambda, \{T\}) = -\mu T + N(T) \ln \mu,$$

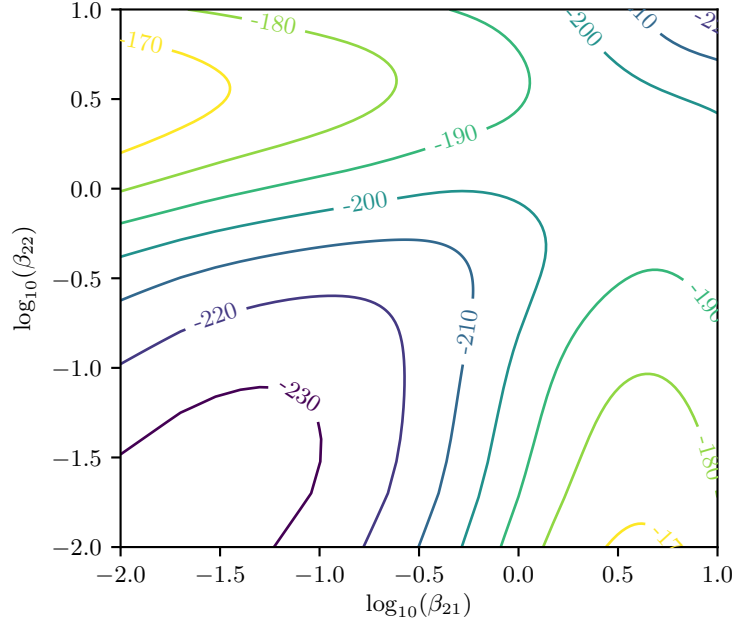


Figure 2.7: Example of local maxima in 2-d Hawkes process likelihood function $\ln L_2(\lambda_2)$

which is finite. However, a strictly concave continuous function having a local maximum cannot tend to a finite limit at infinity.

In fact, not only is the likelihood function not concave, but it actually has several local maxima. An illustrative example is given in Figure 2.7 where we draw the contour plot of the partial likelihood function $\ln L_2$ of a simulated 2-dimensional Hawkes process. The kernels are exponential functions with parameters specified in Equation (2.4). While μ_2 , α_{21} and α_{22} are kept fixed, the likelihood values are plotted as functions of β_{21} and β_{22} . The two axes are presented in logarithmic scale. We see at least two local minima in this example.

$$\mu = \begin{pmatrix} 0.1 \\ 0.2 \end{pmatrix} \quad \alpha = \begin{pmatrix} 5.0 & 10.0 \\ 1.0 & 2.0 \end{pmatrix} \quad \beta = \begin{pmatrix} 20.0 & 15.0 \\ 3.0 & 10.0 \end{pmatrix} \quad (2.4)$$

The existence of several local maxima make gradient-type algorithms less relevant for the maximum likelihood procedure and a global optimization algorithm appears necessary. The **Nelder-Mead** simplex algorithm (*NM*) has been widely used in previous works on the calibration of Hawkes processes; however we find it not stable enough when a good a priori guess is not available.

For these reasons, the **Differential Evolution** algorithm (*DE*) [Storn and Price \[1995\]](#) has been chosen to perform the optimization. *DE* is an efficient genetic evolutionary algorithm that has been adopted in various engineering domains such as electrical power systems, artificial neural networks, operation research, image processing etc. Starting from a population of randomly generated points, the algorithm performs a mutation-crossover-selection procedure, where the population is updated to have better objective function values and a large tentative space is scanned.

A pseudocode is given in Appendix A.

Table 2.9: Error rate (%) of calibration by Nelder-Mead algorithm and Differential Evolution algorithm.

Algorithm	T	μ_1	α_{11}	α_{12}	β_{11}	β_{12}	μ_2	α_{21}	α_{22}	β_{21}	β_{22}
DE	250	0.0	0.0	0.0	0.0	0.0	0.6	1.0	2.0	1.6	2.2
	2500	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.1	0.0	0.1
	25000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
NM random	250	2.1	2.1	2.0	2.0	2.1	24.5	27.6	30.7	24.8	30.5
	2500	1.4	1.4	1.4	1.4	1.4	12.1	14.8	16.9	14.6	18.8
	25000	1.0	1.0	0.8	1.0	1.0	9.9	12.6	14.0	11.4	16.7
NM perfect	250	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
	2500	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
	25000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0

2.4.2 Benchmarking the DE algorithm

Simulation-based numerical experiments are performed in order to compare the efficiency of the *NM* and *DE* algorithms. More specifically, we consider a 2-dimensional Hawkes process where the parameters are specified in (2.4). 100 process paths are simulated for each $T \in \{100, 250, 500, 1000, 2500, 5000, 10000, 25000\}$, and the parameters are calibrated from each simulated path with various algorithms.

NM is used with different initialization methods. For *NM random*, the initial reference point is drawn from uniform distributions. Denoting by ρ the L_1 norm of the kernel ($\rho = \frac{\alpha}{\beta}$), we choose

$$\mu \sim \mathcal{U}(0, 1) \quad \rho \sim \mathcal{U}(0, 1) \quad \beta \sim \mathcal{U}(0, 100) \quad (2.5)$$

and optimize with respect to ρ instead of α .

The algorithm *NM perfect* refers to *NM* where the true input parameters are used as reference point.

The empirical probability of error for each optimization algorithm is show in Table 2.9 for $T \in \{250, 2500, 25000\}$:

Clearly, with the possible exception of short time horizon, *DE* almost always finds the optimal point, getting very close to the *NM perfect* algorithm.

2.4.3 Improvement in high dimensions

The local maximum problem is more severe when dealing with higher dimension and real data instead of simulated data. In this section, we present some treatments designed to mitigate the numerical issues and boost the convergence towards a global maximum.

2.4.3.1 Some evolutions of the *DE* algorithm: a quick guided tour

Thanks to its wide variety of applications, *DE* has attracted a lot of interest, and the recent survey paper [Das et al. \[2016\]](#) documents a host of novel ideas to improve its classical form. Below is a brief summary of some of the proposed improvements (notations are those used in Algorithm 1):

- **Mutation strategy.** The donor vector $v_{i,g}$ in mutation can be generated with different strategies. The classical algorithm adopts a so-called “DE/rand/1” strategy

$$v_{i,g} = x_{r_1,g} + F(x_{r_2,g} - x_{r_3,g})$$

where r_1^g , r_2^g and r_3^g are mutually exclusive integers randomly chosen in $\llbracket 1, N \rrbracket \setminus \{i\}$. It could be preferable to approach the current best value

$$v_{i,g} = x_{best,g} + F(x_{r_2,g} - x_{r_3,g})$$

or use more points for deviation

$$v_{i,g} = x_{r_1,g} + F(x_{r_2,g} - x_{r_3,g}) + F(x_{r_4,g} - x_{r_5,g})$$

Combinations of these ideas are of course possible, which create vast candidate strategies.

- **Crossover.** Apart from the idea of the *binomial/uniform* crossover, another method called *exponential* crossover is also considered. The trial vector u takes the value of the donor vector v for adjacent coordinates. The benefit is limited to special structures of problems where neighboring variables are linked but relatively independent of other variables. As a result the *binomial* crossover is more frequently used.
- **Adaptation of control parameter (F and CR) and strategy.** It aims at adding learning performances to the offspring generation. Either the strategies are randomly chosen from fixed ensemble of strategies and parameters, which are designed to aid the algorithm to converge or explore larger space so that the combination can balance the two effects; or the mutation strategy is fixed, but the parameters can adapt to the evolution.
- **Population control.** The most natural idea is the reduction of population as they approach to each other and concentrate in a small region. Such reduction can be pre-scheduled or dynamically controlled based on the computational budget. On the other hand, varied population (instead of monotonically decreasing) is also introduced as a choice to adapt to the evolution of the algorithm.

Other extensions actually go beyond the classical framework, for example using new initialization techniques, adding clustering technique for the sub-population topology, and so on. Hybridization opens another branch of research: on the one hand *DE* is combined with other heuristic methods to explore the advantages of exploration strategies, and on the other hand, local search methods are injected into the *DE* algorithm to boost convergence and precision.

In the interest of tractability, we choose to concentrate on the non-hybrid extensions. In [Das et al. \[2016\]](#), the algorithm L-SHADE is reported to have the “best competitive performance among non-hybrid algorithms at the CEC 2014 competition on real parameter single-objective optimization”. Compared to the classical algorithm, L-SHADE combines adaptation in every respect - mutation, parameter control and population control:

- **Mutation** use the *current – to – pbest/1* strategy, where the new donor vectors are obtained by

$$v_{i,g} = x_{i,g} + F_i(x_{pbest,g} - x_{i,g}) + F_i(x_{r1,g} - x_{r2,g}),$$

where $x_{pbest,g}$ is randomly selected from the best $\lfloor pN \rfloor$ members in generation g , where $(p \in [0, 1])$. This strategy exhibits some greediness towards the current best points, but the existence of p leaves the flexibility for tradeoff between exploitation and exploration.

- **Parameter control** In order to dynamically adapt the parameters F and CR , a record of past candidates is maintained. Two lists of size H , M_{CR} and M_F , are kept. For each generation, F_i and CR_i are drawn randomly with certain distributions depending on randomly chosen means from the lists:

$$F_i = randc_i(M_{F,r_i}, 0.1), \quad CR_i = randn_i(M_{CR,r_i}, 0.1) \mathbb{1}_{\{M_{CR,r_i} \neq Null\}},$$

where $randn$ follows a normal distribution and $randc$, a Cauchy distribution. For each generation, the k th element ($k = g \bmod H$) of the list is updated, according to CR_i and F_i that succeed to find ameliorated points. Such mechanism introduces learning characteristics for the F and CR selection, in order to overcome the stagnation problem.

- **External archive introduction** To maintain diversity, an external archive is used so that parent vectors that are worse than the trial vectors are preserved in A . When generating donor vectors, $x_{r2,g}$ can be selected from $P \cup A$.
- **Linear population size reduction** The whole population N_g decreases according to the allowed total number of generations.

$$N_{g+1} = \text{round} \left(\left(\frac{N_{min} - N_{init}}{G} \right) * g + N_{init} \right),$$

where N_{init} is the classical initial population size, and N_{min} is the smallest possible population size for a mutation strategy.

The L-SHADE is a combination of interesting ideas. Roughly stated, the current-to-pbest/1 mutation helps approach the best candidates in the population, accelerating the convergence of the algorithm; the parameter control aims at learning the trade-off between exploration and exploitation; the external archive is to help keep diversity of the population so that exploration is partly internalized by the exploitation of the abandoned history; and the population size reduction saves computational cost to allow larger initial populations.

2.4.3.2 Calibrating high dimensional Hawkes order book models

Let us now turn towards the actual application of *L-shade* to the task at hand.

Starting from 100 different initial populations for each strategy with the same number of points and maximum generations, we plot the histograms of the final log-likelihood function values for one dimension of the 12-dimensional Hawkes model with real data in Figure 2.8, for different modifications of *DE*. The classical strategy, noted as “rand/1”, serves as a reference for the suggested “current-to-pbest/1”. The parameter adaptation is also combined with “rand/1” to provide better performances. We finally introduce a version with a refinement of the initial parameter intervals, noted as “better guess”. The right subplot is a zoom of the one on the left, to further show the improvement due to “better guess”.

Some comments are in order :

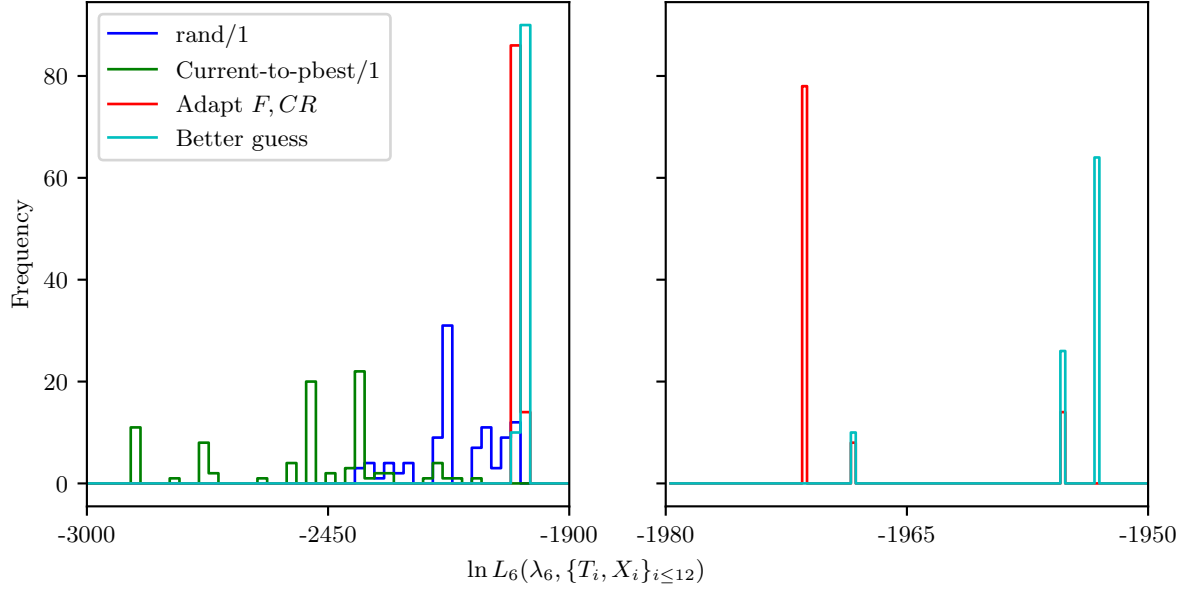


Figure 2.8: Distribution of optimal objective likelihood functions in different optimization strategy tests. The right one is zoomed at the optimal zone for further illustration.

- In the “current-to-pbest/1” strategy, the closer p is to 0, the greedier (converge rapidly to already found optimum) the algorithm is, and the more probable it is that the optimization gets trapped at a local maximum. The closer p is to 1, the more the algorithm favors exploration.
- The learning mechanism for F and CR in the adaptative version leads to some improvements. Parameters are initialized according to the following distributions:

$$\mu_m \sim \mathcal{U}(0, \frac{0.2N_m}{T}), \quad \rho_{mn} \sim (0, \min(\frac{0.2N_m}{N_n}, 0.5)), \quad \beta_{mn} = u_1 \mathbb{1}_{\{u_0=0\}} + u_2 \mathbb{1}_{\{u_0=1\}}$$

$$\text{for } u_0 \sim \mathcal{B}(1, 0.5), \quad u_1 \sim \mathcal{U}(0, 1) \quad u_2 \sim \mathcal{U}(0, 100)$$

derived from the physical interpretation of μ as the baseline intensity, of ρ as the integrated intensity of the influence from event arrival, and based on the relation

$$\mathbb{E}[\lambda_{m,\infty}]T = \mu_m T + \mathbb{E}[\sum_n \rho_{mn} N_n].$$

- Although different runs starting from different initial populations do not converge to the global maximum, some improvement may be gained from a “better guess” of the initial intervals.

An increase of the population size plays a major role in boosting the convergence: a larger population prevents points from getting trapped around the same local maximum. On the other hand, it is useless to keep all the population as the algorithm approaches the end of its iterations, since points tend to form clusters. As a consequence, it makes sense to consider effective population reduction techniques and use the saved computational budget to cover a larger search space.

Building on the linear population reduction method inspired by the combination of *DE* with clustering algorithms in Li and Zhang [2011] and the use of pairwise Euclidean distance for dynamic population control in Yang et al. [2013], we propose an additional reduction mechanism which allows, not only to decrease the function evaluation times, but also to avoid convergence to local maxima.

The algorithm is said to have converged if each coordinate of all the points in the population has converged. The convergence conditions of the coordinates are

$$\begin{aligned}\sigma(\mu_m) &< e_r \langle \mu_m \rangle \quad \text{or} \quad \max \mu_m - \min \mu_m < e_a, \\ \sigma(\rho_{mn}) &< e_r \langle \rho_{mn} \rangle \quad \text{or} \quad \max \rho_{mn} - \min \rho_{mn} < e_a, \\ \sigma(\beta_{mn}) &< e_r \langle \beta_{mn} \rangle \quad \text{or} \quad \langle \rho_{mn} \rangle < e_a,\end{aligned}$$

where $\sigma(\cdot)$ and $\langle \cdot \rangle$ are the standard deviation and the mean value respectively, and e_r and e_a are the relative and absolute error tolerance. At each generation, we eliminate points that are close to the current b best ones, using a criterion similar to the termination conditions for the population: suppose the points are sorted according to their objective function values by descending order. For a given point x_i , if $\exists j \in \llbracket 1, b \rrbracket / \{i\}$ such that all the following conditions are satisfied:

$$\begin{aligned}|\mu_{mi} - \mu_{mj}| &< e_r \mu_{mj} \quad \text{or} \quad |\mu_{mi} - \mu_{mj}| < e_a \\ |\rho_{mni} - \rho_{mnj}| &< e_r \rho_{mnj} \quad \text{or} \quad |\rho_{mni} - \rho_{mnj}| < e_a \\ |\beta_{mni} - \beta_{mnj}| &< e_r \beta_{mnj} \quad \text{or} \quad \rho_{mnj} < e_a,\end{aligned}$$

then x_i is eliminated from the population. In practice, it is convenient to select a small value for b . The decrease of population size saves some computational budget for the algorithm, which is very beneficial as the computation of the likelihood function is costly.

The combination of these population reduction techniques allows to increase the initial population by a factor of 5 to 10 with no significant impact on the total computation time, and the convergence is largely improved.

As a conclusion, one can say that the improved version *L-SHADE* of the *DE* algorithm drastically enhances the performances of the calibration, but despite all these efforts, we are still left with an average failure rate of approximately 5%.

2.5 Conclusion

This chapter is a study of Markovian Hawkes processes applied to high frequency limit order book data. Suitably designed nonlinear Hawkes processes that include inhibitory effects and a co-existence of time scales are shown to successfully model the dependencies between the arrival of order book events. Thanks to the particularly well-suited distinction between events that trigger, or do not trigger, an immediate change in the current price, the dynamics of the model fully reflect that of the price. Such a description helps cope with some shortcomings of order book models that were previously observed, particularly concerning the realized spot price volatility.

This chapter also gives a detailed analysis on a very important, albeit technical, topic: the choice of the optimization algorithm for the maximum likelihood estimation. The *L-SHADE* algorithm is a significant improvement over the classical *Differential Evolution* algorithm, thanks to better initializations and population control.

As a conclusion, one can say that nonlinear Hawkes processes satisfactorily capture such fundamental features of market dynamics as conditional probabilities, forward recurrence times, or the signature plot. They provide an accurate description of the order book in the high frequency realm, as well as a realistic behaviour of more macroscopic quantities. While leading to a better understanding of the mechanisms driving the markets, their use in the simulation of order driven markets can also lead to a host of potential applications.

Chapter 3

Order-book modelling and market making strategies

Note:

- This chapter is submitted to “Market Microstructure and Liquidity”.
- This chapter is presented in the conference “Financial Econometrics Conference: Market Microstructure, Limit Order Books and Derivative Markets”, Lancaster University, Lancaster, September 2018.

Abstract

Market making is one of the most important aspects of algorithmic trading, and it has been studied quite extensively from a theoretical point of view. The practical implementation of so-called “optimal strategies” however suffers from the failure of most order book models to faithfully reproduce the behaviour of real market participants.

The aim of this chapter is twofold. First, some important statistical properties of order driven markets are identified, advocating against the use of purely Markovian order book models. Then, market making strategies are designed and their performances are compared, based on simulation as well as backtesting. We find that incorporating some simple non-Markovian features in the limit order book greatly improves the performances of market making strategies in a realistic context.

Contents

3.1	Introduction	52
3.2	Challenging the queue-reactive model	53
3.2.1	The queue-reactive model	53
3.2.2	The limitation of unit order size	54
3.2.3	The role of limit removal orders	57
3.2.4	Enriching the queue-reactive model	62
3.3	Market making in real markets	63
3.3.1	Optimal market making strategies	64
3.3.2	Backtesting the optimal strategies	69
3.4	Conclusion	71

3.1 Introduction

Most modern financial markets are order-driven markets, in which all of the market participants display the price at which they wish to buy or sell a traded security, as well as the desired quantity. This model is widely adopted for stock, futures and option markets, due to its superior transparency. With the emerging of electronic markets, and the deregulation of financial markets, algorithmic trading strategies have become more and more important. In particular, market making - or: liquidity providing - strategies lay at the core of modern markets. Since there are no more *designated market makers*, every market participant can, and sometimes must, provide liquidity to the market, and the design of optimal market making strategies is a question of crucial practical relevance.

Originating with the seminal paper [Ho and Stoll \[1981\]](#), many researchers in quantitative finance have been interested in a theoretical solution to the market making problem. It has been formalised in [Avellaneda and Stoikov \[2008\]](#) using a stochastic control framework, and then extended in various contributions such as [Guéant et al. \[2013\]](#); [Cartea and Jaimungal \[2013a,b\]](#); [Cartea et al. \[2014\]](#); [Fodra and Pham \[2013, 2015\]](#); [Guilbaud and Pham \[2013a,b\]](#); [Bayraktar and Ludkovski \[2014\]](#); [Guéant et al. \[2012\]](#) or [Gueant and Lehalle \[2015\]](#). It is noteworthy that, in this series of papers, the limit order book is not modelled as such, and the limit orders are taken into account indirectly thanks to some probability of execution.

In practice, the price discontinuity and the intrinsic queueing dynamics of the LOB make such simplifications rather simplistic as opposed to real markets, and it is obvious to the practitioner that these actual microstructural properties of the order book play a fundamental role in assessing the profitability of market making strategies. There now exists an abundant literature on order book modelling, but, as regards market making strategies - or more general trading strategies, for that matter - only very recent papers such as [Abergel et al. \[2017\]](#) actually address the market making problem using a full order book model.

It is our aim in this chapter to contribute to the literature on the subject, both from the modeling and strategy design points of view, so that the chapter is twofold: it analyzes and enhances the queue-reactive order book model proposed by [Huang et al. \[2015b\]](#), and then study the optimal placement of a pair of bid-ask orders as the paradigm of market making.

A word on data: we use the Eurostoxx 50 futures data for June and July, 2016 for the entire analysis, and backtest until November for out-of-sample validation. Eurostoxx 50 futures offers two main advantages:

1. it is a very large tick instrument, with an average spread very close to 1 tick and extremely rare multiple-limit trades (less than 0.5%);
2. the value of a futures contract is very high in euros, so that one thinks in terms of number of contracts rather than notional. This actually simplifies the choice of the unit.

These two observations allow us to follow only the first (best) Bid and Ask limits, and focus on the question of interest to us, namely, the design of a model where the state of the order book as well as the type of the order that leads the book into its current state, are relevant. This approach, departing from the purely Markovian case, is based on empirical observations and will be shown to provide a more realistic and useful modelling framework. In a different mathematical setting, a similar reasoning is at the root of Hawkes-process-based order book models such as studied in [Lu and Abergel \[2018a\]](#); [Abergel and Jedidi \[2015\]](#).

The chapter is organized as follows: Section 3.2 presents the rationale and the calibration of the enriched queue reactive model that improve the performances of the initial model of [Huang et al. \[2015b\]](#). Section 3.3 addresses the optimal market making strategies in the context of this enhanced model, studying it both in a simulation framework, and in a backtesting engine using real data.

3.2 Challenging the queue-reactive model

This section presents empirical findings that lay the ground for two improvements to the queue-reactive model of [Huang et al. \[2015b\]](#). The first one is concerned with the distribution of order sizes, whereas the second, and maybe more original one, addresses the difference in the nature of the events leading to identical states of the order book.

These improvements will be incorporated in two order book models inspired by, but largely extending, the queue-reactive model. In Section 3.3, these models will be used in a simulation and backtesting framework to study optimal market making policies.

3.2.1 The queue-reactive model

In [Huang et al. \[2015b\]](#), the authors propose an interesting Markovian limit order book model. The limit order book (LOB in short) is seen as a $2K$ -dimensional vector of bid and ask limits $[Q_{-i} : i = 1, \dots, K]$ and $[Q_i : i = 1, \dots, K]$, the limits being placed $i - 0.5$ ticks away from a reference price p_{ref} . [Huang and Rosenbaum \[2017\]](#) further establish the ergodicity and asymptotic stability of more general Markov order book models.

Denoting the corresponding quantities by q_i , the $2K$ -dimensional process $X(t) = (q_{-K}(t), \dots, q_{-1}(t), q_1(t), \dots, q_K(t))$ with values in $\Omega = \mathbb{N}^{2K}$ is modeled as a continuous time Markov chain with infinitesimal generator $\mathcal{Q}$ of the form:

$$\begin{aligned} \mathcal{Q}_{q, q+e_i} &= f_i(q), \\ \mathcal{Q}_{q, q-e_i} &= g_i(q), \\ \mathcal{Q}_{q, q} &= - \sum_{q \in \Omega, p \neq q} \mathcal{Q}_{q, p}, \\ \mathcal{Q}_{q, p} &= 0 \text{ otherwise.} \end{aligned}$$

where e_i is a $2K$ -dimensional unit vector that is 1 at the i -th position and 0 elsewhere.

The authors study several choices for the function g : in the first and simplest one, queues are considered independent. The second one introduces some one-sided dependency, whereas the third one emphasizes the interaction between the bid and ask sides of the LOB. Some important statistical features of the limit order book can be reproduced within this model, such as the average shape of the LOB. However, when trying to calibrate the queue-reactive model on our dataset of EUROSTOXX50 future, we observe new phenomena that lead us to enrich the model in two directions.

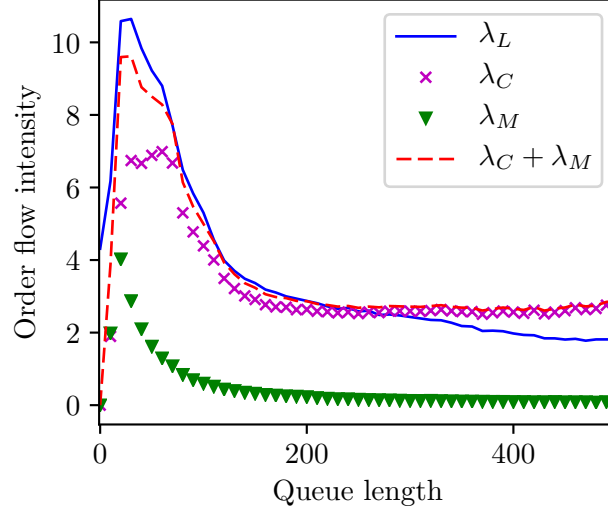


Figure 3.1: Queue reactive model order intensities

3.2.2 The limitation of unit order size

Following a procedure similar to that in [Huang et al. \[2015b\]](#), the conditional intensities of limit orders, cancellation and market orders are calibrated and presented in Figure 3.1. Note that the intensities of all order types are higher when the corresponding queue length is small. For each queue (bid and ask), the limit orders are liquidity constructive events and the other two are liquidity destructive.

There are three different regimes for the queue sizes:

- λ_L is slightly higher than $\lambda_C + \lambda_M$ when the queue size is smaller than (approximately) 70.
- They become comparable when the queue size lies between 70 and 300.
- When the queue size is above 300, λ_L decreases whereas $\lambda_C + \lambda_M$ stays stable.

Of special importance is the condition $\lambda_L < \lambda_C + \lambda_M$ when the queue size is large, a fact which guarantees that the system is ergodic.

Another interesting feature is that the intensity of market orders drastically decreases when the queue size increases, a fact that can be reformulated as the concentration of trades when the queue size is small. From a practical point of view, a small queue usually indicates a directional consensus, so that liquidity consumers race to take the liquidity before having to place limit orders and wait for execution at the same price.

In the work of [Huang et al. \[2015b\]](#), and many other related works, the order size is supposed to be constant. It is however clear from empirical analyses that the order sizes have remarkable statistical properties, and that such information is relevant to the LOB dynamics, see for instance [Abergel et al. \[2016\]](#); [Muni Toke \[2015\]](#); [Rambaldi et al. \[2017\]](#).

The mean order sizes conditional on the queue states are shown in Figure 3.2. The size of limit and cancellation orders appear to be quite stable across different queue sizes (except for very

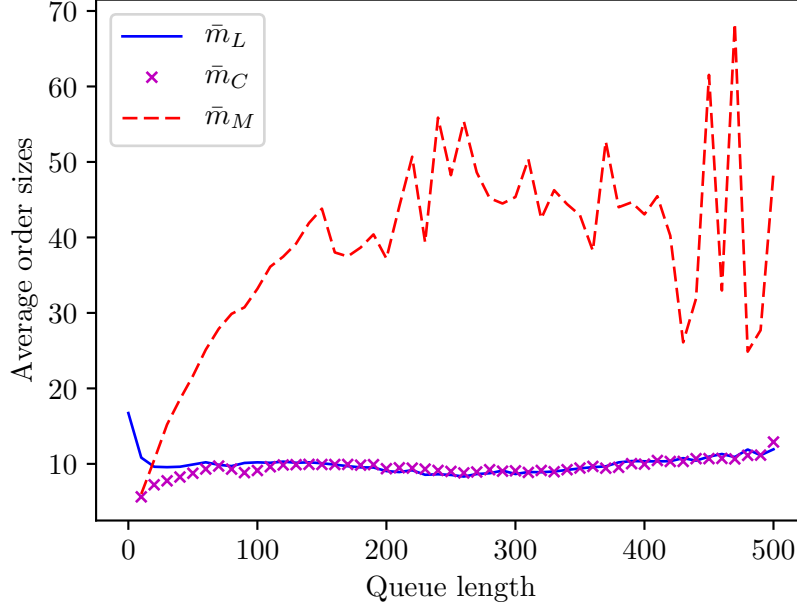


Figure 3.2: Queue reactive model mean order sizes

small queues), whereas the average size of market orders is clearly increasing with the queue size.

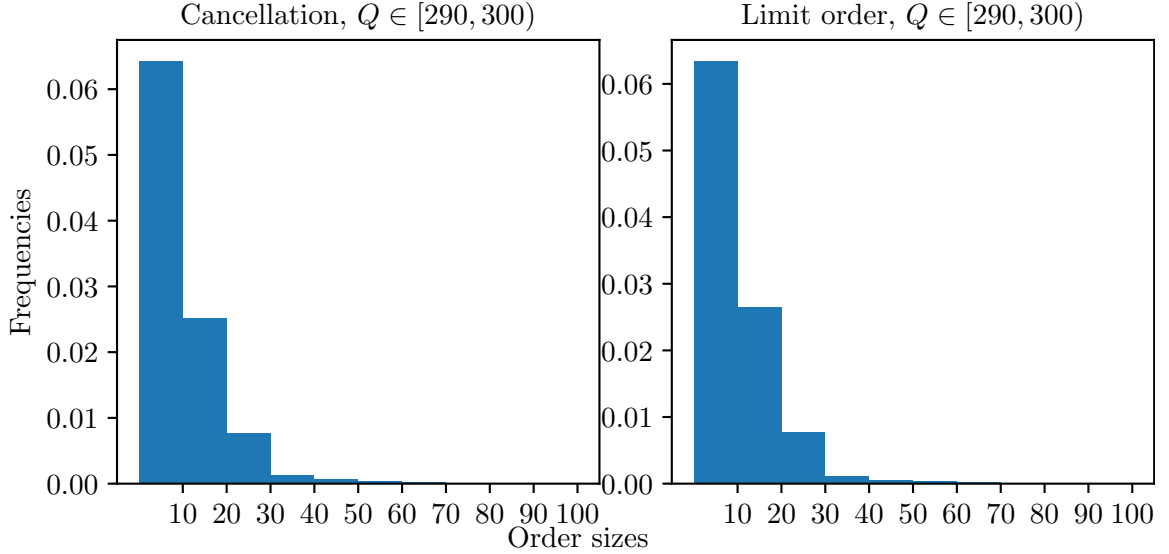


Figure 3.3: Limit order and cancellation size histogram

To further analyze the distributions of order sizes, histograms of order sizes with the queue length in the interval of $[290, 300)$ are shown in Figure 3.3 and 3.4. Bins of 10 futures on the x -axis are used to produce the histograms.

The empirical distributions for limit orders and cancellations are rather similar and can be modelled with a geometric distribution. On the contrary, there exist some interesting patterns in the sizes of market orders: the distribution looks like a mixture of a geometric distribution for small sizes, and Dirac functions at $\{[50, 60), [100, 110), [150, 160), \dots\}$ and $q = Q$. It is not surprising, because traders do not necessarily randomize their market orders, so that multiples of

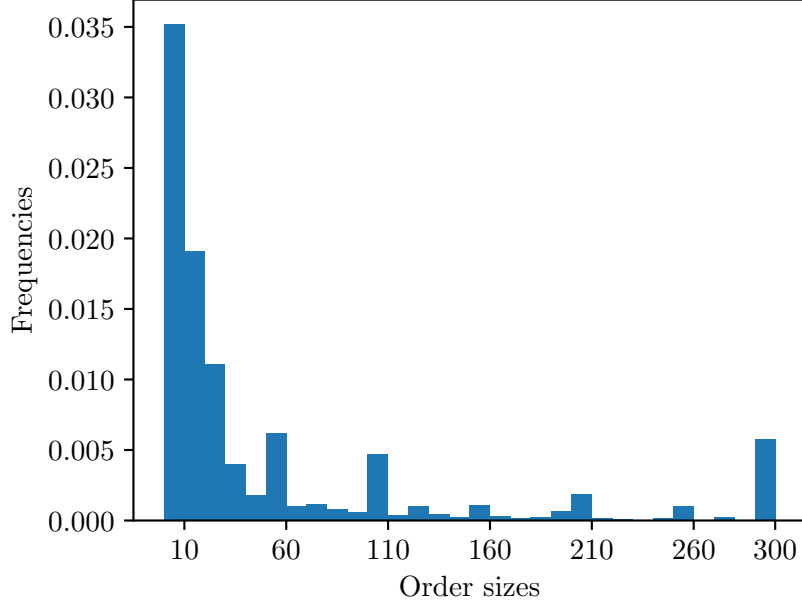


Figure 3.4: Market order size histogram when $Q \in [290, 299]$

50 occur quite frequently, and so do orders that completely eat up the first limit. For the rest of this subsection, we denote $q, Q \in \mathbb{N}^*$ to represent the bin of $[10(q-1), 10q)$ and $[10(Q-1), 10Q)$.

Inspired by such observations, a simple model of order sizes can be proposed:

- Limit order sizes follow geometric laws $p_L(q; Q)$ with parameters $p_0^L(Q)$ depending on the queue sizes;
- Cancellation sizes follow truncated geometric laws with parameters $p_C(q; Q)$

$$p_C(q; Q) = \mathbb{P}[q|Q] = \frac{p_0^C(1-p_0^C)^{q-1}}{1-(1-p_0^C)^Q} \mathbb{1}_{\{q \leq Q\}};$$

- Market order sizes follow a mixture of geometric laws and Dirac functions

$$p_M(q; Q) = \mathbb{P}[q|Q] = \theta_0 \frac{p_0^M(1-p_0^M)^{q-1}}{1-(1-p_0^M)^Q} \mathbb{1}_{\{q \leq Q\}} + \sum_{k=1}^{\lfloor \frac{Q-1}{5} \rfloor} \theta_k \mathbb{1}_{\{q=5k+1\}} + \theta_\infty \mathbb{1}_{\{q=Q, Q \neq 5n+1\}}$$

where the parameters $\{p_0^M, \theta_0, \theta_k, \theta_\infty\}$ depend on Q .

The parameters of limit order sizes are simply estimated. For the cancellation and market orders, a maximum likelihood method can be used. The market order log-likelihood is

$$\begin{aligned} \log(L) = & \sum_{q_i \neq 5n+1, q_i \neq Q} \log \theta_0 + (q_i - 1) \log(1 - p_0^M) + \log p_0^M - \log(1 - (1 - p_0^M)^Q) \\ & + \sum_{k=1}^{\lfloor \frac{Q-1}{5} \rfloor} \sum_{q_i=5k+1} \log(\theta_0(1 - p_0^M)^{q_i-1} p_0^M + \theta_k(1 - (1 - p_0^M)^Q)) - \log(1 - (1 - p_0^M)^Q) \\ & + \mathbb{1}_{\{Q \neq 5n+1\}} \sum_{q_i=Q} \log(\theta_0(1 - p_0^M)^{q_i-1} p_0^M + \theta_\infty(1 - (1 - p_0^M)^Q)) - \log(1 - (1 - p_0^M)^Q). \end{aligned}$$

The calibration results are presented in Table 3.1 and 3.2. As expected, the geometric distribution parameters for limit orders and cancellations are quite close to each other, and remain stable across different queue lengths. The geometric distribution parameter for market orders slightly decreases with the queue length. Dirac parameters θ_k are very stable across different queue lengths.

Table 3.1: Calibrated limit order and cancellation size parameters

Q	21	22	23	24	25	26	27	28	29	30
p_0^I	0.6421	0.6415	0.6458	0.6410	0.6443	0.6430	0.6418	0.6439	0.6404	0.6387
p_0^C	0.6578	0.6591	0.6600	0.6623	0.6598	0.6611	0.6557	0.6554	0.6538	0.6496

Table 3.2: Calibrated market order size parameters

Q	p_0^M	θ_0	θ_1	θ_2	θ_3	θ_4	θ_5^*	θ_6^*
21	0.3486	0.8357	0.0185	0.0338	0.0081	0.1038	-	-
22	0.3557	0.8311	0.0198	0.0338	0.0094	0.0215	0.0844	-
23	0.3383	0.8517	0.0148	0.0366	0.0084	0.0188	0.0697	-
24	0.3327	0.8475	0.0108	0.0373	0.0099	0.0192	0.0753	-
25	0.3333	0.8310	0.0234	0.0379	0.0084	0.0214	0.0779	-
26	0.3292	0.8391	0.0203	0.0408	0.0115	0.0167	0.0716	-
27	0.3250	0.8374	0.0188	0.0369	0.0114	0.0220	0.0116	0.0619
28	0.3134	0.8351	0.0191	0.0452	0.0086	0.0201	0.0164	0.0554
29	0.3090	0.8262	0.0192	0.0476	0.0086	0.0181	0.0135	0.0668
30	0.3050	0.8426	0.0205	0.0402	0.0096	0.0188	0.0103	0.0580

* For the cases of $Q \neq 5k + 1$, $\theta_{\lfloor \frac{Q-1}{5} \rfloor + 1}$ represents θ_∞

The enhanced LOB model is then driven by compound Poisson processes with intensities conditional on the queue size.

3.2.3 The role of limit removal orders

By construction, Markovian LOB models assume that the past has no influence on the future except through the present. Nevertheless, it has been established, see e.g. [Abergel et al. \[2016\]](#); [Lu and Abergel \[2018a\]](#) for some in-depth empirical studies, that the nature of past events actually influence the order flow and, therefore, the future states of the LOB.

In this section, the emphasis is set on the nature of the last event that totally removes the liquidity at one limit, and the subsequent evolution of the LOB.

An order that completely eats up a limit is termed a *limit removal order*. Such an order can only be a cancellation or market order. The liquidity removal process is denoted by $Y(t) = (q^r(t), O^r(t))$, where q^r stands for the order size and $O^r \in \{O^M, O^C\}$ for its type - O^M for a market order and O^C for a cancellation. Then, $Y(t)$ is a càdlàg process with jump times $\tau^r = \{\tau_i^r\}_{i=1,2,\dots}$.

In a symmetric way, an order that creates a new limit will be termed a *limit establishing order*. Such an order can only be a limit order, but it can be placed either on the bid

or ask side since the limit is empty. The liquidity establishing order process is denoted by $Z(t) = (q^e(t), O^e(t))$, where $O^e \in \{F, R\}$: F for ‘follow’, so that $P(\tau_i^e) \neq P(\tau_i^r-)$, and R for ‘revert’. The “removal-establishing” event results in twice the mid price changes. After the removal order, the mid price jumps a half tick to the price of the removal order. If the following establishing order is of the same nature as the limit before the removal event, the mid price will be the same as before, so the price “reverts”. Otherwise, the mid price “follows” the previous change and jumps another half tick. Again, $Z(t)$ is a càdlàg process with jump times as $\tau^e = \{\tau_i^e\}_{i=1,2,\dots}$.

3.2.3.1 Characterization of O^e

Table 3.3: O^e conditional probabilities

O^r	F	R	$\mathbb{P}[O^e = F O^r]$
O^M	7554	1409	84.3%
O^C	1043	2823	27.0%

We first investigate $O^e(\tau_i^e)$ conditional on $O^r(\tau_i^r)$. Table 3.3 presents the daily average number of events as well as their conditional probabilities.

A first important observation is that the price tends to move in the same direction if it is triggered by O^M , whereas mean-reversion is more likely in the case of an O^C -triggered price change. In other words, O^M is more informative than O^C , and is the main driver of price moves.

One can wonder whether this dependency structure could be simplified to one on the state of the LOB only, that is, whether the conditional distribution of $O^e|O^r, X$ could be explained by $O^e|X$. Figure 3.5 presents the empirical distributions of $O^e|(q^F, q^S)$ for $O^r \in \{O^M, O^C\}$ respectively, where $q^F = q_{-1}$ and $q^S = q_2$ at τ_i^e- , and $q^F = q_1$ and $q^S = q_{-2}$ for the corresponding bid side, which are LOB states just before the arrival of a limit establishing event.

Although the queue sizes may vary before the arrival of a limit establishing event, Figure 3.5 shows that their influence is negligible, and the conclusion is that O^e is highly dependent on the O^r but much less on (q^F, q^S) .

It seems also relevant to include the sizes $q^r(\tau_i^r)$ in the analysis. In Figure 3.6, the upper panel presents the conditional distribution of $O^e|O^r, q^r$ as a function of q^r , while the lower panel presents the cumulative distribution function of q^r . For both O^M and O^C , the probability of $O^e = F$ increases with the order size. It is however noteworthy that, for O^M , this probability reaches almost 1 when $q^r \geq 30$, a rather significant fact as there are over 20% of market orders that lay in the interval of $(30, +\infty)$.

For O^C , the situation is completely different: not only does the probability of $O^e = F$ remains close to 0.5, but the proportion of cancellation orders of size larger than 20 is very small and the decrease of $\mathbb{P}[O^e = F|O^r = O^C, q^r]$ when $q^r > 50$ is not statistically significant.

The interpretation of this phenomenon is direct: not only market orders are much more informative than cancellations but, the larger the market order is, the more likely it is to indicate a directional price movement that the market will follow. A market buy (sell) order of size larger than 30 that consumes the entire liquidity will almost always be followed by new liquidity providers placing bid (ask) limit orders at the previous trade price.

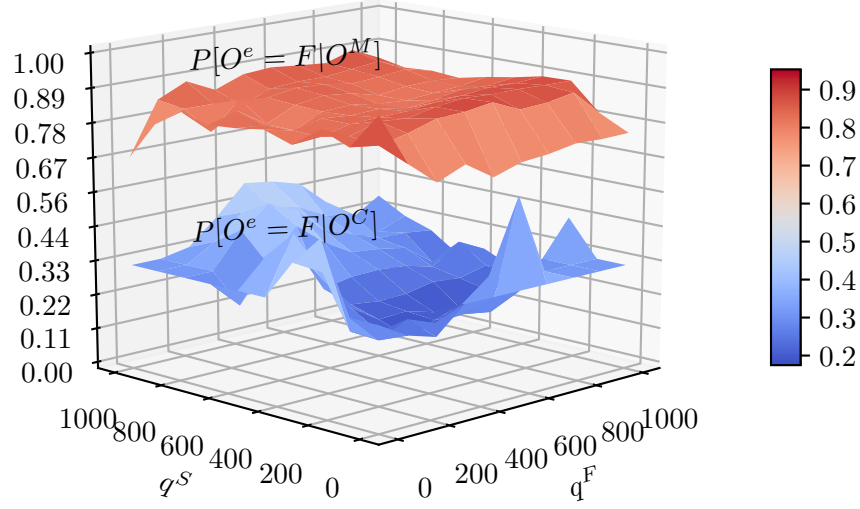


Figure 3.5: Conditional distribution O^e w.r.t the first limit queue sizes for different O^r

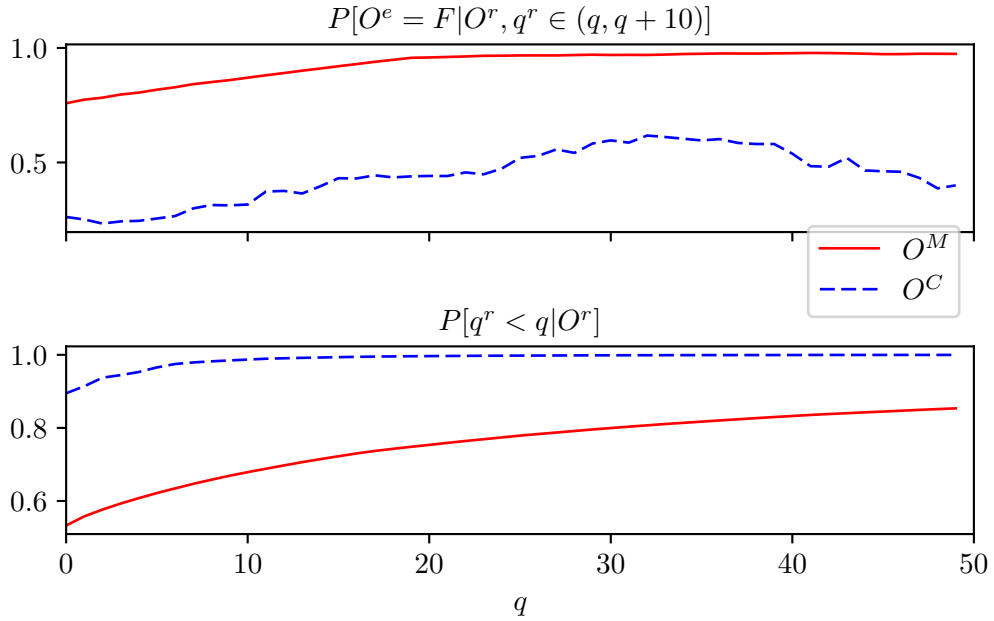


Figure 3.6: Conditional distribution O^e w.r.t q^r for different O^r

3.2.3.2 The size of liquidity establishing orders

Table 3.4 summarizes the average number of (O^r, O^e, q^r) events, as well as the average quantity q^e corresponding to each of these classes of events. As already observed in Section 3.2.3.1,

very few large cancellations result in an empty limit, most large orders are market orders that systematically lead to a follow event. Moreover, one can see a monotonically increasing relationship between q^e and q^r across all values of (O^r, O^e) .

Due to the low number of large orders in other classes, we will concentrate on (O^M, F) for an in-depth analysis.

Table 3.4: q^e conditional distribution on O^r , O^e and intervals of q^r

(O^r, O^e)	# of events per day			mean q^e size		
	$(0, 10)$	$[10, 20)$	$[20, \infty)$	$(0, 10)$	$[10, 20)$	$[20, \infty)$
(O^M, F)	3623	1128	2802	7.46,	10.18	27.93
(O^M, R)	1153	180	75	6.69,	6.75	27.98
(O^C, F)	906	112	23	7.22,	7.45	9.08
(O^C, R)	2552	243	26	5.60,	7.71	7.99

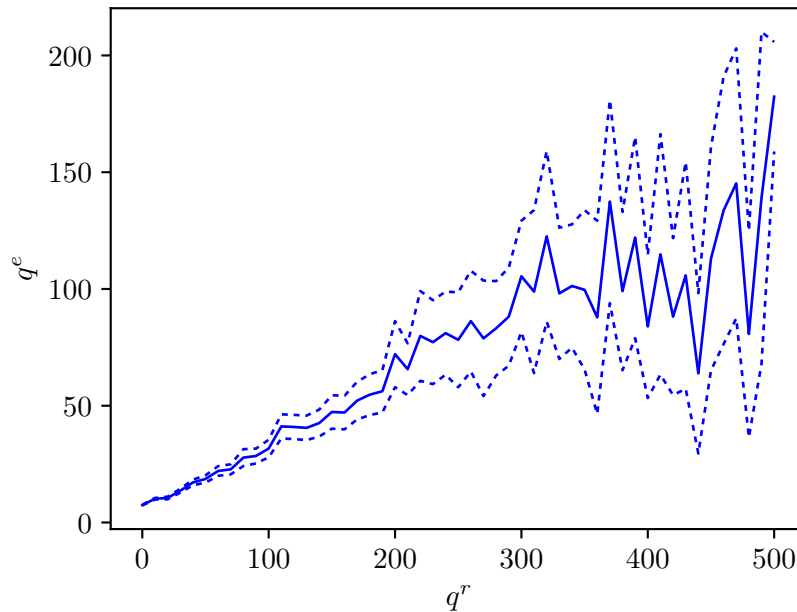


Figure 3.7: Mean q^e with respect to q^r . The full line presents the mean and the dashed line presents the 95% confidence interval of the estimation.

Figure 3.7 illustrates the variation of q^e with respect to q^r for $(O^r, O^e) = (O^M, F)$. There exists a definite monotonically increasing relationship between the two quantities, and the new limit can even reach a size of over 100 shortly after the old limit is consumed by a very large market order. The geometric distribution that we have previously advocated for the size of limit orders fails to represent such a phenomenon.

Figure 3.8 shows some empirical conditional distributions of q^e obtained by classifying q^r . When q^r is small, the distribution is close to geometric. However, as q^r increases, the density presents fatter tails and discrete peaks, and it is no longer appropriate to use a geometric distribution - one may rather consider using the empirical distribution to model $q^e|q^r$.

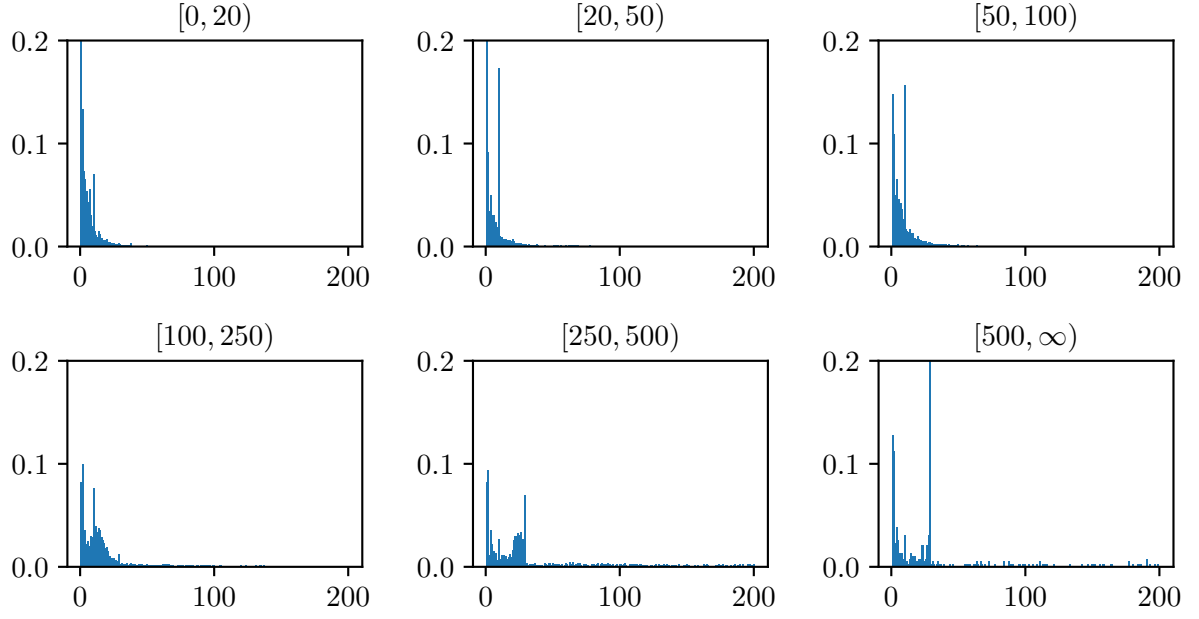


Figure 3.8: Histograms of q^e with respect to q^r in different intervals.

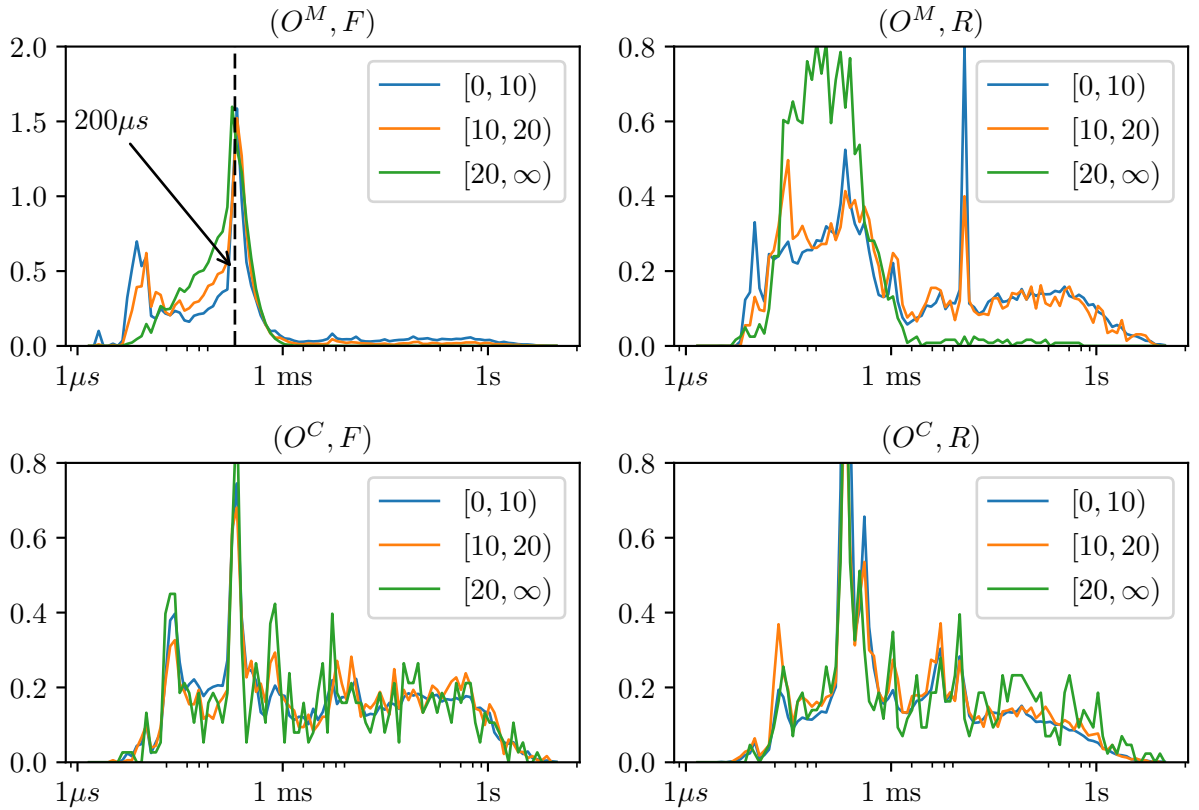


Figure 3.9: $\log_{10}(\tau^r - \tau^e)$ distribution with respect to O^r , O^e and q^r in different intervals.

3.2.3.3 Characterizing inter-event durations

Going further, we now analyze the empirical distributions of (the logarithm of) inter-arrival times $\log_{10}(\Delta\tau)$ in Figure 3.9. When the triggering event is O^C , the distribution of $\Delta\tau$ are very similar, no matter what O^e and the size of the order are. When the triggering event is O^M and the price reverts, $\Delta\tau$ distributions are quite similar when $q^r < 20$.

On the contrary, the distribution of $\Delta\tau|(O^M, F)$ is very different from the others, and is different for different intervals of q^r . First, as expected, except for very small q^r , the densities of $\Delta\tau > 1ms$ are close to 0. In addition, there is a sharp peak around $200\mu s$. One possible reason of such concentration is about the reaction to large trades. Though liquidity providers try to place limit orders immediately after a large trade, they are constrained by the round-trip latencies of the market and their own systems. As $200\mu s$ represents the typical market round-trip latency, a high density concentrates around this level. But there could have been a few market makers who react to exogenous information that is the same as the large trade but arrives late which result in the density below. And the above part comes simply from the higher latencies of some market makers.

3.2.4 Enriching the queue-reactive model

As a conclusion to this empirical study, one can see that the dynamics of the order book, and not only its state, must be used in order to gain a faithful representation of the market. In order to enhance the queue-reactive model, it is necessary to add a dependency of the order sizes and, more importantly, a dependency on the nature of the order that drove the book into its current state. One must therefore depart from the Markovian framework, but only slightly, and in the interest of a much more realistic modelling.

We then propose two LOB models that can be viewed as extensions of the queue-reactive model, but differ from it as regards price transitions:

- **Model I:** When either limit is empty, the next limit order that closes the spread depends only on whether the emptied limit was on the bid or ask side. The size of the order and the recurrence time are independent.
- **Model II:** The new limit order (O^e, q^e) depends not only on the side of the cleared limit, but also is a function of the last removal event (O^r, q^r) . The arrival time of the event is determined by $\tau^e = \tau^r + \Delta\tau$, where $\Delta\tau$ is dependent on (O^r, q^r) .

Each of these models is compared with the unit size model having the same transition rule as Model I, which we refer to as Model 0. Monte Carlo simulations are conducted for the different models, and the results are benchmarked against real data.

As an example, Figure 3.10 presents the distributions of the best limit quantities sampled at 1s frequency in the various models, as well as in the data. Model 0 produces a LOB that is very concentrated around the constructive-destructive equilibrium limit size (around 300), with lower density for smaller queues and almost no density for long queues. Adding the size distribution in Model I already improves the model, with a global shape closer to the real data. Model II improves the too high density around the equilibrium queue size, and also produces more realistic, fatter tails for the queue size distribution.

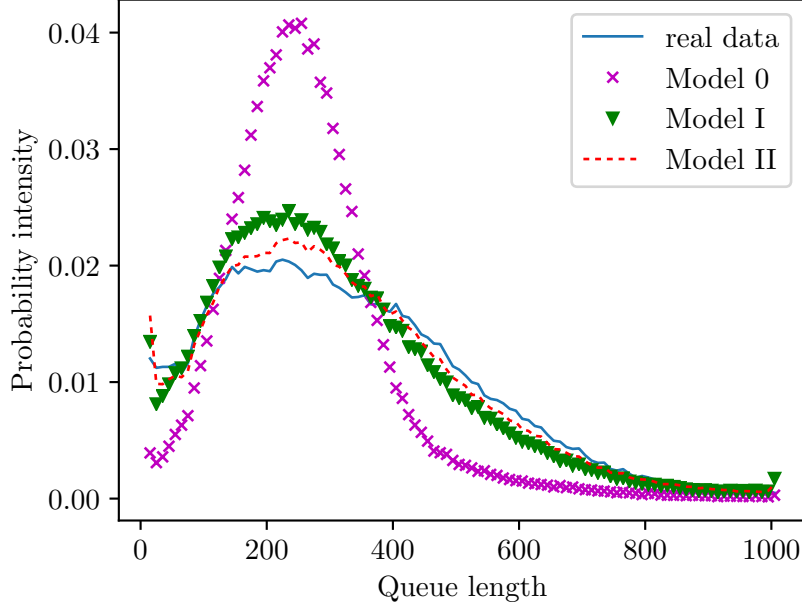


Figure 3.10: Best bid/ask quantities distributions.

3.3 Market making in real markets

In this section, we address the problem of defining an optimal market making (or: liquidity providing) policy in an order-driven financial market. Our goal is to mimick as well as possible the situation of an electronic market maker, and a realistic order book model incorporating the various empirical properties we have just shown must therefore be designed. Once a model is set up, stochastic control is used in order to derive the optimal strategy. The characteristics of the optimal strategy are quite useful in determining whether a model makes sense or not, as its performances in a real market environment can be measured, either through backtesting or direct experimentation. An 'optimal' strategy that would lose money in the market would be an indication of poor modelling !

Based on the work in [Hult and Kiessling \[2010\]](#) (see also [Abergel et al. \[2017\]](#); [Bäuerle and Rieder \[2011\]](#) for related approach and results), optimal market making strategies are numerically and empirically studied under the various hypotheses corresponding to the models labelled 0 and II introduced in Section 3.2.4. In the interest of readability, the theoretical framework and main results are recalled in Appendix B, while the current section is devoted to the presentation and discussion of the results.

As usual in stochastic control, the quantity of interest is the *value function* in each state of the LOB, that is, the expected future profit and loss (P&L) of a pair of bid and ask orders: the risk-neutral market maker will place an order if the value is strictly greater than 0, and stay out of the market otherwise. The associated optimal strategy is thus defined according to the value functions.

After a simulation-based study and analysis of optimal market making strategies, with or without inventory control, the optimal strategies for each models are backtested against realistic market conditions.

Note that, for technical reasons, the maximum queue size is set equal to 500 contracts and

the bid-ask spread is supposed to be always 1 tick - as a matter of fact, for liquid, large-tick instruments, trades almost never occur when the spread is larger than 1 tick, so that this hypothesis stands.

3.3.1 Optimal market making strategies

This section is devoted to a comparison of the optimal strategies for Model 0 and Model II. The results for Model I are only slightly different from Model 0 and will not be presented here.

The state of the LOB is described by a quadruple (x^B, x^A, y^B, y^A) representing the best bid and ask quantities in number of lots, and the position of the Market Maker (MM)'s orders in the queue. The key question in market making is the time to enter the order book.

Initially, the MM places his order at the end of the queue, so that

$$x^B = y^B \quad x^A = y^A.$$

3.3.1.1 Value function and optimal strategies for Model 0

We illustrate the value as a function of the initial state in Figure 3.11. The x and y axis are respectively the bid and ask queue lengths in the figure. Once the state values are known, the optimal strategy is straightforward: if the value is positive, the optimal action is to stay in the order book, whereas if the value is 0, the optimal action is to cancel the orders.

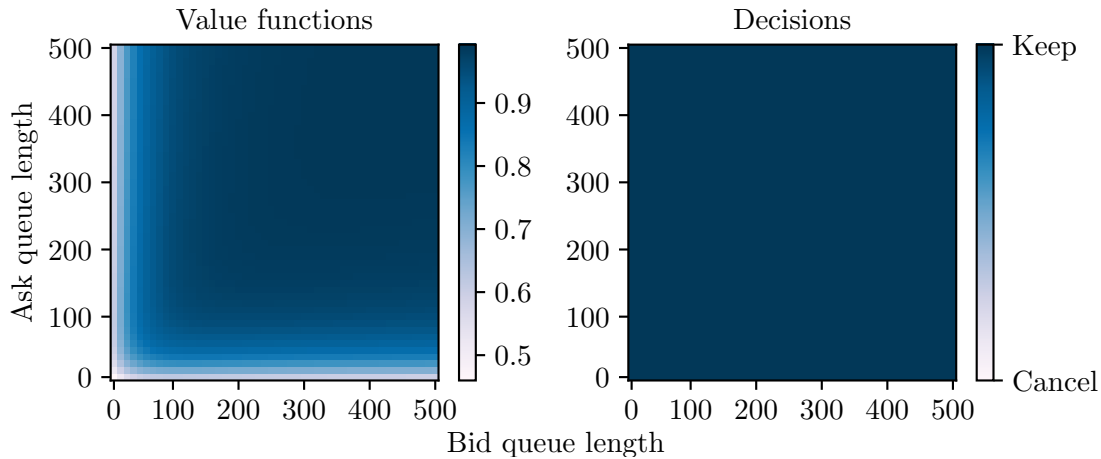


Figure 3.11: Initial state values in Model 0

One observes that the values are mostly positive. When the order book is highly imbalanced, the order value is lower. The lowest values appear when both of the limit queues are short. However, the minimum value is still higher than 0.5. The result can be interpreted as follows: when the queue length is short, the probability that the queue will become empty increases before other limit orders arrive behind the MM's orders; when his order is executed on one side of the order book but not on the opposite side, the MM will have to close the position at a loss with a market order. When the queue is long, the probability that both MM's orders get executed before the price changes is higher, so the value is closer to 1.

However, such a result is hardly a reflection of reality. From practical experience we know that

a simple market making strategy is unprofitable or poorly profitable in the real market. The value functions calibrated with Model 0 are at least 0.5 tick, indicating that we should follow a naive strategy, always placing an order on both the bid and ask sides. This unrealistic behaviour comes from the fact that Model 0 has very stable limits, so that the MM's bid and ask orders tend to be both executed before either limit is cleared.

3.3.1.2 Value function and optimal strategies for Model II

Figure 3.12 shows the values (left) and the decisions (right) depending on the initial state. The state is non-profitable when the state value is 0, so that the MM cancels orders on both sides and wait for a further transition, otherwise, he stays in the book and wait for his orders to get executed.

Clearly, the results are very different from those obtained previously: in most of the initial states, the optimal policy is to cancel. In fact, the MM is penalized when there is a price change before both his orders get executed, and the existence of large market orders increases such a risk. The most favourable situation is when the two queues are balanced and relatively long, so that his orders can gradually gain priority before the order book becomes imbalanced again and the price changes.

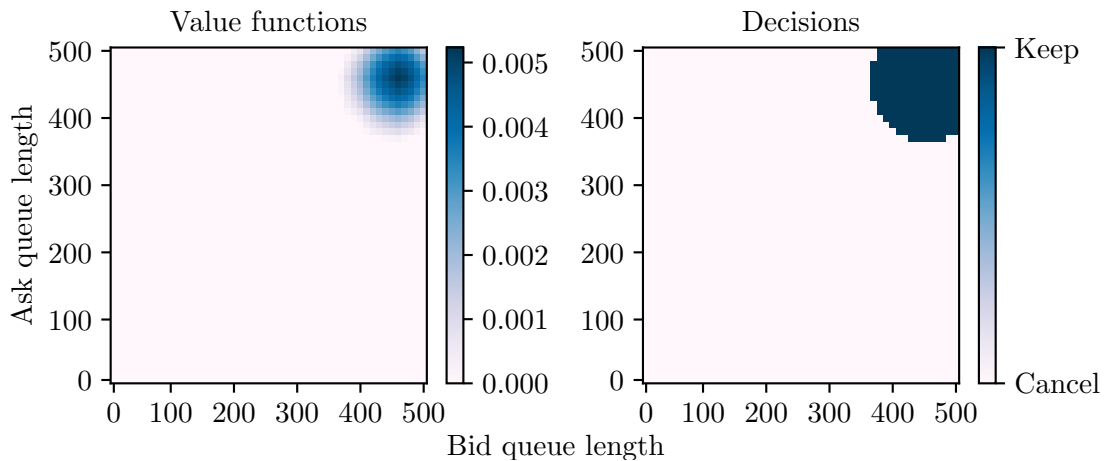


Figure 3.12: Initial state values in enriched Queue-reactive model

The relative values of different bid and ask positions are also of great importance. Figure 3.13 shows the example of the state $x = (200, 300)$. As expected, the states become more valuable as the MM's orders get closer to the top of the queues. In addition, the values are not symmetric: the value of $(200, 300, 50, 10)$ is larger than that of $(200, 300, 10, 50)$ because, the longer the queue, the more valuable the priority.

To further illustrate the differences in values and strategies for different states, let us now fix the positions in the queues to be $y^B = 300$ and $y^A = 200$. Figure 3.14 is a plot of the value function and optimal decision according to queue lengths.

As expected, the priorities become more valuable when the queues are longer, and the asymmetry also exists. For instance, in this case, and regardless of the bid queue length, the MM should not stay in the order book when the ask queue length is 220 or less.

Figure 3.15 shows the value function according to the MM's priority for small queue sizes.

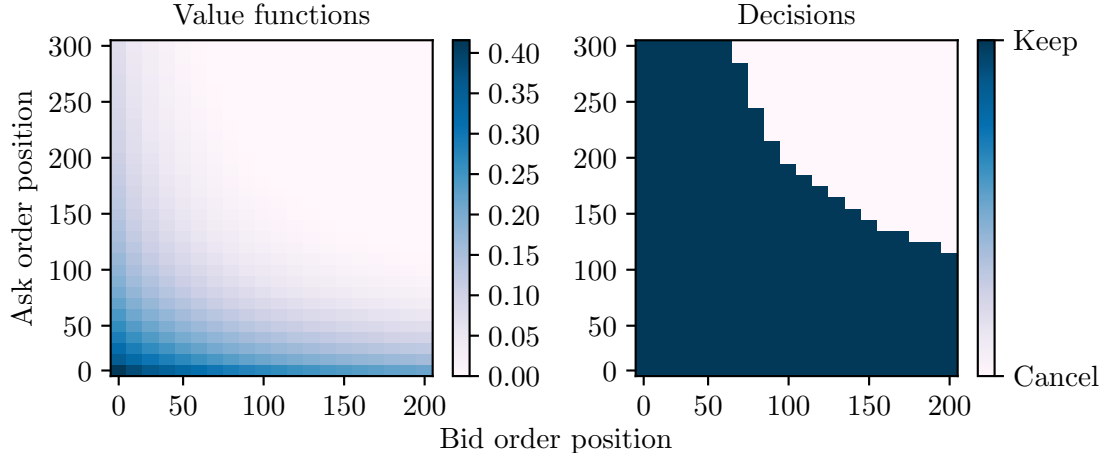


Figure 3.13: State values and decisions when the order book state in $x = (200, 300)$

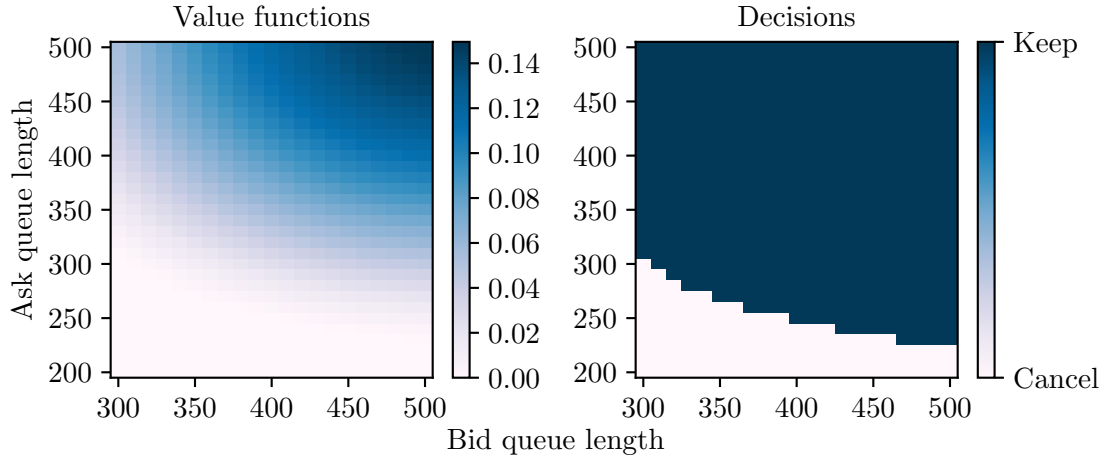


Figure 3.14: State values and decisions when the market maker's order priorities are $(y^B, y^A) = (300, 200)$

When the queue lengths become small, the proportion of market orders compared to that of limit orders increase and the risk of a price move becomes higher. As a consequence, it becomes uninteresting for the MM to stay in the market when the order book is in the state $(50, 100)$, except when he is at the top of both queues.

When one order is executed, the strategy becomes a “buy-one-lot” or “sell-one-lot” problem. The optimal buy-one-lot strategy after the MM's ask order has been executed is described in Figure 3.16 and Figure 3.17. Actually, an ask queue length of $x^A = 140$ acts as a threshold, beyond which the market maker should always stay in the bid queue and hope for an execution as a limit order. Otherwise, if the ask queue is too short (or the bid queue too long), the price tends to go upwards, and the MM may be better off using a market buy to close his inventory.

3.3.1.3 Relaxing the inventory constraints

The previously defined market making strategies allowed for an inventory of at most 1, a severe restriction in practice. Relaxing this constraint, one may decide to continuously submit limit bid or ask orders when the previous ones are executed and the state immediately after the

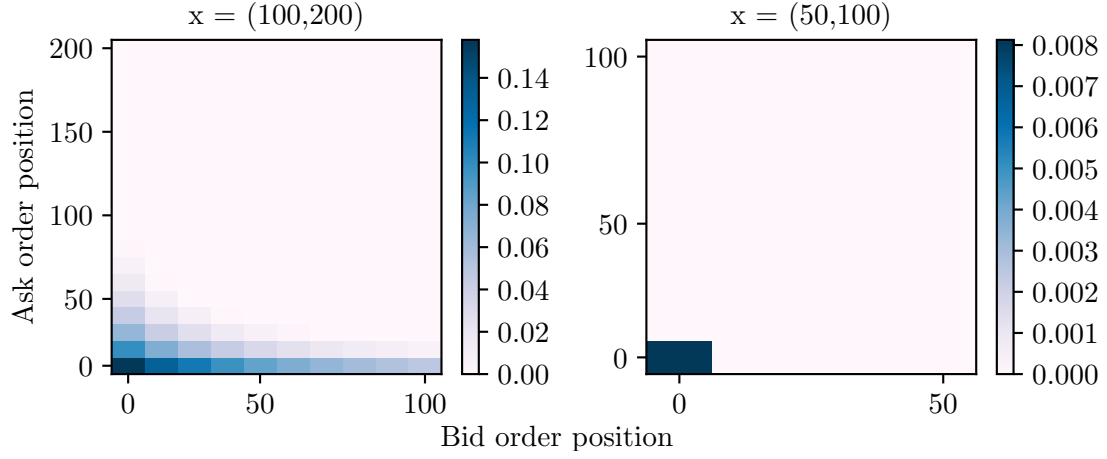


Figure 3.15: State values for short queues

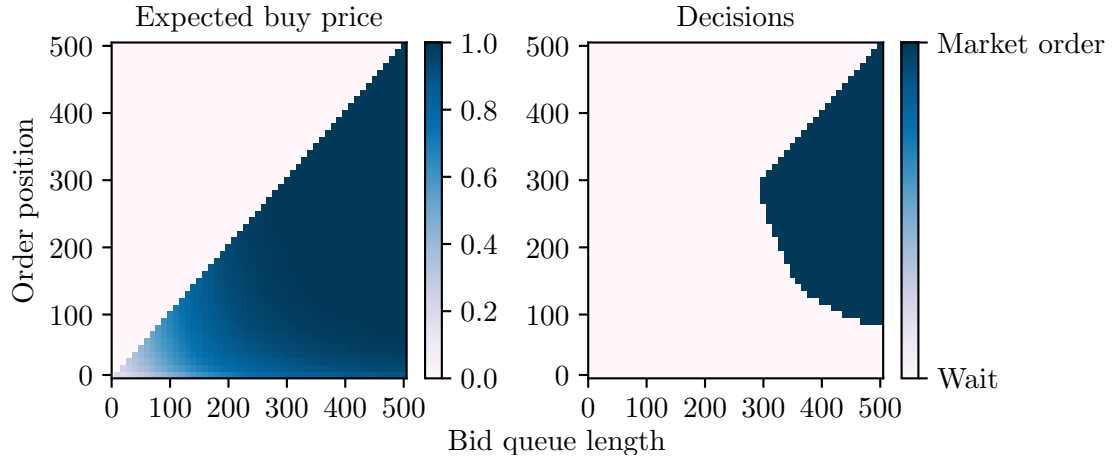


Figure 3.16: Buy-one-lot strategy when $x^A = 100$

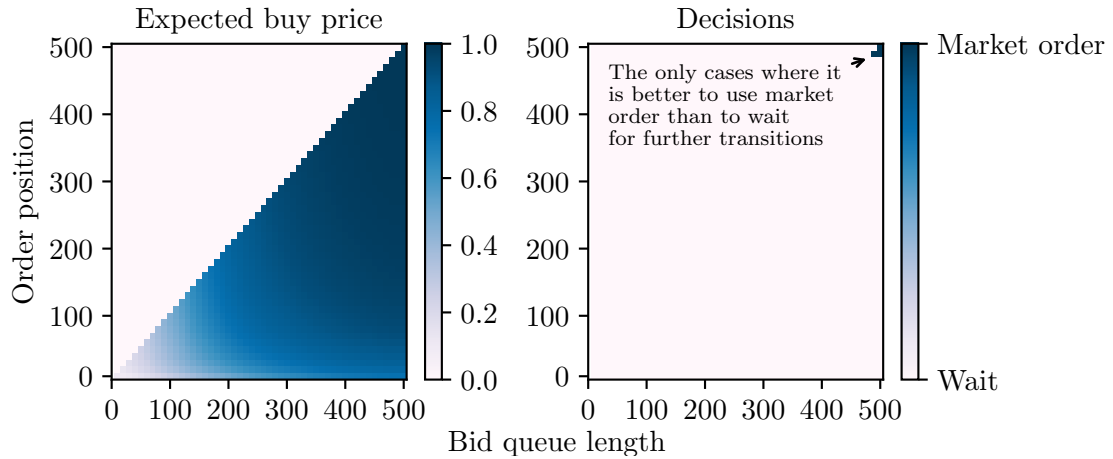


Figure 3.17: Buy one unit strategy when $x^A = 140$

execution has value greater than 0, regardless of the inventory. This new strategy violates the initial constraint but however provides some useful information as to whether the profits made by making the spread can cover the inventory risk. This new strategy will be referred to as a “locally optimal strategy”. As previously seen, the value function for Model 0 is positive for

all states, which implies that the MM should follow a naive strategy, continuously submitting orders whatever the state of the LOB. This naive strategy will be used as a benchmark.

Under the assumptions of Model II, Monte Carlo simulation is used to compare the performance of the locally optimal and naive strategies.

For the locally optimal strategy, the simulation runs for an hour of market activity, whereas for the naive strategy, it does for only 10 minutes as the turnover of the naive strategy is much higher. The final inventory is supposed to be closed at the end of the simulation using market orders.

The daily average P&L, absolute inventory and turnover - measured in number of contracts - of both strategies are summarized in Table 3.5, the standard deviations of the means being given between parentheses. The P&L of the continuous strategy is much better than that of naive strategy. Its P&L is significantly positive, with a Z-score of 2.87, whereas the P&L of the naive strategy is fairly negative, because of the adverse selection effects embedded in Model II.

Table 3.5: Continuous and naive strategy Monte Carlo simulation

	continuous	naive
P&L	1.35 (0.47)	-7.12(0.71)
Inv	3.33 (0.02)	10.47(0.09)
turnover	19.01 (0.06)	113.46(0.22)

3.3.2 Backtesting the optimal strategies

To better assess the performance of Model II in reproducing a real trading environment, the optimal strategy is backtested using tick-by-tick order book data.

Since building a backtester that can replay historical data in a realistic way is a notably difficult task, we first present our methodology before showing the results.

3.3.2.1 The backtester

The purpose of a backtesting engine is to reproduce as well as possible the performances of a trading strategy, were it to really be traded in the market. It can produce results that are quite different from those obtained by simulation.

One of the first issues is latency. As pointed out in Figure 3.9, the round-trip latency is reflected in market data. In practice, the incoming data feed and outgoing orders can have different latencies, as they do not necessarily share the same venue (public broadcast reception and private sending). When one reacts to exogeneous information, or information from markets that do not share the same location, these two latencies have to be specified separately. In addition, the latency is not constant, and is typically higher during intense activities because the sequential processing of orders by the matching engine will take more time when orders are piling up. Unfortunately, this variable latency is not measurable, so that for our backtests, a fixed round-trip latency is considered.

Another, very important source of discrepancy between backtests and real market conditions, is of course market impact. Orders in the market are seen by other traders and become a source of information that influences the order flows. For example, a constantly monitored indicator is the market (LOB) imbalance, a clue to short term price movements: an order on the bid side posted by the market maker will actually tend to decrease the probability of execution of said order. However, the market impact of passive limit orders is smaller, and much harder to model, than that of market orders, and we have chosen to ignore it. Such a simplification is realistic when the order size is small compared to the typical queue size. In the case of SX5E futures, the lot size of 10 contracts is very small compared to the average quantity at the first limit, generally around 500 contracts.

When replaying the LOB and trades adding the MM's orders, the key issue is to set some priority rules these fictitious orders in the LOB.

The main question is that of cancellations. Since the exact order flow is not available, it is impossible to know which orders have been completely cancelled. Moreover, modification of an order is allowed with a loss of priority (it can be viewed as simultaneously cancelling an order and resubmitting a new one), so that a decrease in size is not distinguishable from a cancellation. For the sake of simplicity, we decide to randomly choose the position where a cancellation occurs. Since orders with low priorities are more likely to be canceled, a capped exponential law is used.

It is true that there are some ways to improve the identification of canceled orders, by registering the limit orders in a list and matching the quantities, but it is our practical experience that the improvement is marginal.

As for trading rule, a fictitious MM's order will be executed if its position in the queue is within the size of an incoming market order. Immediately after this trade, and before the MM submits another order, the queue length is set to be the same as that observed in the initial LOB data - as if we had increased the trade size to absorb the MM's order.

One important exception is the case when the market order clears one limit: even if the MM's order is at the bottom of the queue, it will be considered to be executed, in accordance with the markedly very high proportion of market orders that actually empty limit in real data. Comparisons with production results show that this choice can influence up to 20% of the turnover for any given strategy and moreover, that ignoring these executions result in overestimating the backtest performances.

As a conclusion, although the only way to validate a strategy is probably to run it in the market, backtesting engines are always useful, but a lot of care must be taken when designing them and interpreting the results they provide.

3.3.2.2 Backtesting market making strategies

The "locally optimal strategy" is backtested, as well as a naive strategy (possibly with a threshold).

The backtest is run on the period ranging from July to mid-November, 2016. The MM's order size is fixed to 1 lot (10 contracts). In the case of the locally optimal strategy, the value function calculated using market data is used to determine whether the MM orders should wait for execution or be cancelled. When one of the orders is executed, another order on the same side is submitted immediately if the value is positive, otherwise, a new order will not be submitted

and the existing order on the opposite side is cancelled.

For the naive strategy, bid and ask orders are continuously submitted as soon as the previous one on the same side is executed. When a threshold q_{min} is set, the order stays in the LOB only if the corresponding queue is longer than the threshold.

Two different thresholds of respectively 250 and 400 contracts are studied.

Finally, for practical reasons, we also implement a simple inventory control: whenever the inventory reaches a certain level, new limit orders will no longer be submitted until the inventory falls below the level. In the examples shown here, the maximum inventory is 80, but other values have been tested and do not change the conclusions.

3.3.2.3 Results

The backtesting results are presented in Figure 3.18. Table 3.6 summarizes the daily average P&L, turnover and profitability of the different strategies. The locally optimal strategy is much more profitable than the naive strategies. Without a threshold, the naive strategy provides too much liquidity and the turnover is much higher than with the optimal strategy, making it difficult to compare the P&Ls. But even with a threshold of 400 (so as to match the turnover of the optimal strategy), the strategy has a decreasing trend and ends up with a negative P&L.

The locally optimal strategy is the only one that actually ends up positive.

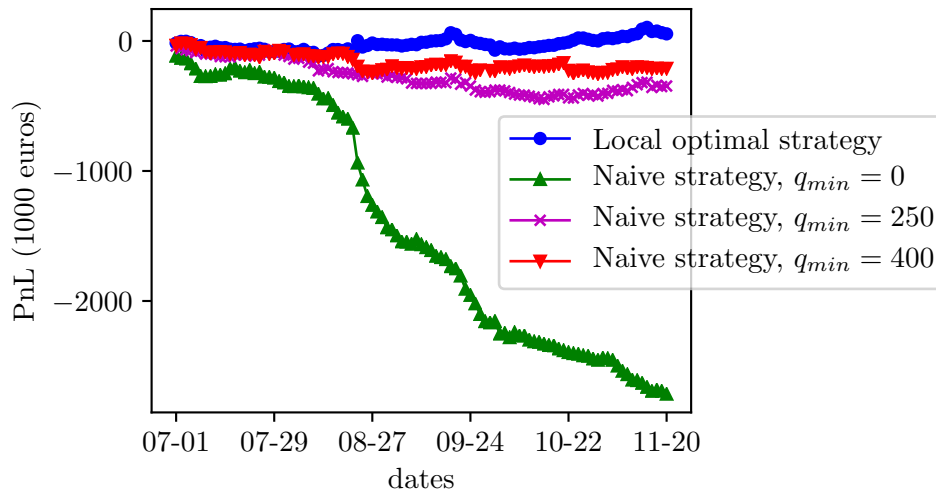


Figure 3.18: Strategies backtested with real data

Table 3.6: Average daily statistics for different strategies

	Optimal	Naive		
		$q_{min} = 0$	$q_{min} = 250$	$q_{min} = 400$
P&L (ke)	0.54	-26.89	-3.43	-2.13
Turn over (Me)	45.20	3430.42	119.25	35.83
Profitability (bp)	0.12	-0.08	-0.29	-0.59

3.4 Conclusion

This paper addresses the calibration of Markovian limit order book models *à la* [Huang et al. \[2015b\]](#), and their practical usefulness for market making strategies.

First, we show that the size and arrival times of limit orders and cancellations are stable across different queue lengths, whereas market order behave completely differently. Moreover, our analysis shows a strong dependence of a limit establishing order on the nature of the limit removal order. Incorporating these two features allows us to enhance in a significant way the queue-reactive model.

Second, the model is used with the purpose of designing optimal market making strategies. The optimal strategy is compared with the naive strategy, first in a simulation framework where it obviously performs better, but also in a historical backtester, where we also find it to perform much better. This is an indirect proof that the enhanced queue-reactive model may be closer to describing the real market than most Markovian LOB models.

Chapter 4

Price prediction with microstructure information boosted by deep learning

Note:

- This chapter is finished recently and will be submitted to corresponding journals.

Abstract

*LSTM model has seen a big success on various domains such as speech recognition and time-series prediction, but is less applied in finance. With large tick-by-tick high-frequency data, we find one of the rare financial scenarios where an application of machine learning is less troubled by the small set problem. We first assess the LSTM model and compare it with the linear models for 1-step mid-price change prediction. It is found that our results are in line with the results calibrated on US data by [Sirignano and Cont \[2018\]](#), such that the relation between LOB states and the mid-price change is **non-linear**, **stationary** and **universal**. But the universal model is not as good as a stock-by-stock model due to the limitation of the parallel optimization algorithm (ASGD) used to calibrate the universal model. Furthermore, we show that a prediction of longer horizon is necessary to circumvent the impact of short-term reversion. LSTM can still over perform OLS model for 10-step prediction. However, it becomes equivalent or even worse for 100-step prediction due to the increase of noise. At last, a simple market making strategy based on the predicted signals is constructed to show that in a more realistic situation, the LSTM model allows building a more profitable trading algorithm.*

Contents

4.1	Introduction	74
4.2	Deep learning experiment: one-step price change prediction	76
4.2.1	Indicators	77
4.2.2	Model construction	78
4.2.3	Empirical results	80
4.3	Multi-step deep learning regression	86
4.3.1	Model performance	88
4.3.2	Strategy building	89
4.4	Conclusion	93

4.1 Introduction

Market makers provide liquidity in the financial markets by posting firm buying and selling quotes. They encounter various risks, such as inventory risk, execution risk and adverse selection risk as is pointed out in [Guilbaud and Pham \[2013b\]](#). Among them, which concerns the most is the adverse selection risk, as it relates directly to the profit and loss (P&L). Loss is much more troublesome than just the volatility.

Adverse selection comes naturally from an economic point of view. For an efficient market, the transaction price should always be the best reflection of the future price with all of the information up to the moment. There is no reason that market makers should always generate a positive expected profit. Regardless of the trading facilities (some low frequency asset managers don't have direct access to the market, but they send their trades to brokers who determine to execute with high frequency algorithms in an optimal way), the market price is decided in a manner of equilibrium of supply and demand. Admitting for simplification that all market participants can post limit orders and market orders (not very far from reality), the choice between these two types will purely be determined by economic benefit. There should not be an advantage of one to the other on average, see [Wyart et al. \[2008\]](#).

In order to protect from the adverse selection, market makers have to predict the forward price movements according to all of the available information. However this is a pretty complex problem because:

- There are different classes of strategies, which can be as simple as blind investment to hedging or complex cross-asset systematic strategies. With limited research resources, it is impossible to exploit all of the strategies for trading signals.
- Participants in the market are mixed of all horizons, from high frequency traders to very long term asset managers. Market makers have to determine how to project longer term predictions to the horizons of their interest.

So that the market makers have to consider several important aspects:

- What sources of data should be used. Whether it will concentrate on tick by tick microstructure data, minute-level price and volume or even daily data. Or they can also be external data, like company fundamentals, social media and sentiment analysis.
- What horizon should be predicted. Whether it should be extremely high frequency at each tick, or for several seconds, several minutes, or some hours and even some days.
- What values to predict. For low frequency prediction, it is classical to suppose a unique continuous price for each instrument. However, market makers have to consider high frequency scenarios, where it is not only the mid-price, but also the bid, ask, and even the queue priorities that influence the applicability of the prediction.

Our research aims to help improve the knowledge of how to predict the price change at short term horizon, from tick-by-tick variation to several minutes (multiples steps of variations) using the limit order book information. This is to the particular interest of intraday high frequency market makers to tightly control the inventory risk and solve order placement problems. Rarely other data are updated at this frequency, and the information of order books can hardly be enough to predict very long term price variation. However, as the limit order book as well as

the completed transactions contains all the instantaneous supply and demand information, it reflects the short term trading equilibrium that dictates the price change. This domain is thus specifically interesting as a complementary to other relatively longer term predictions.

Financial predictions have been largely studied with different approaches. For example, Stübinger and Bredthauer [2017]; Stübinger and Endres [2018]; Avellaneda and Lee [2010]; Brogaard et al. [2014]; Krauss [2017] study the problem from a statistical arbitrage viewpoint by directly modelling the prices of different stocks to detect abnormal price changes that tend to revert. Previous works like Cartea et al. [2018]; Lehalle and Mounjid [2017]; Cont et al. [2014]; Yang and Zhu [2016]; Stoikov [2017]; Gould and Bonart [2016] have pointed out the statistical properties of the LOB imbalance as well as the importance of incorporating this information into modelling and market making strategy design. Market impact, which refers to the price changes led by trades, is a hot topic that has been widely studied and particularly for the high-frequency context in Moro et al. [2009]; Toth et al. [2012]; Jaisson [2015]; Eisler et al. [2012].

On the other hand, deep learning methods have earned wide success in image and language processing. The well known universal approximation theorem established in Hornik et al. [1989] laid the theoretical foundation that neural networks are able to approximate all measurable functions. Though few theories have been developed to well understand the theoretical foundations of deep neural networks, they have been widely applied in many domains and are believed to be an even more powerful tool. It helps to build models with complex non-linear relations. However, it is also prone to overfitting, where the model has been too close to the sample data and the power of generalization becomes weak. Its success has largely relied on the size of available training data, which counts for millions to billions of realizations, and the stationarity of the characteristics. Nevertheless, these conditions are hardly satisfied in finance. Limit order book data is unique in that its abundance makes it an exceptional experiment field where the application of deep learning may make sense.

Among the few machine learning applications in finance, some point to more traditional areas of natural language processing and sentiment analysis in Antweiler and Frank [2004]; Tetlock [2007]; O'Connor et al. [2010]; Tetlock et al. [2008]. Kim [2003]; Ahmed et al. [2010]; Talih and Hengartner [2005]; Patel et al. [2015]; Fischer and Krauss [2017] exploit the possibility of time series predictions in financial markets. More recent works start to apply deep learning methods, for example Heaton et al. [2016]; Selvin et al. [2017]; Nelson et al. [2017]; Gu et al. [2018], but they use mostly traditionally low frequency data.

A LSTM network is a recurrent neural network (RNN) composed of LSTM (Long short-term memory) units is called. *long short-term* refers to that LSTM is a model for short term memory which can last for a long period of time. LSTM is proposed by Hochreiter and Schmidhuber [1997] and improved by Gers et al. [1999]. Nowadays, LSTM networks are widely used in recognition problems (for example, speech recognition and handwriting recognition). In particular, Sirignano and Cont [2018] use an LSTM with asynchronous stochastic gradient descent trained on 1000 US stocks to predict the one-step price change based on historical limit order book states. They report *non-linearity*, *universality*, *stationarity* and *path-dependence* of the relation. This is to our knowledge one of the first works that apply deep learning algorithms in finance, and the results are both encouraging and interesting (though similar algorithms were proposed in 1980s and 1990s, they encountered technical problems for prediction uses). Especially, they provide evidence that the long believed non-stationarity and uniqueness of each stock are not always true in high frequency limit order books.

Nonetheless, one-step price change is much easier to predict but very difficult to use for practitioners. The misleading concept comes from the confusion of the relation between price

change and P&L. Only transactions can result in P&L, but price change, especially the high frequency mid-price change is not necessarily associated to buys and sells. So a high accuracy of mid-price change will not necessarily provide higher P&L for traders. For it to be closer to the reality, one can either model the limit order book conditional on the trades, or to approximate it by predicting longer horizons (more steps) so that the transaction problem is less concerning.

On the other hand, some classical predictors, like imbalance, order flow and price reversion have long been proved to be efficient in practice. It is worth evaluating if synthesising limit order book states into indicators can help reduce the dimension of the input to limit the overfitting and improve the performance.

This chapter is organised as following: Section 4.2 deals with a similar framework as [Sirignano and Cont \[2018\]](#) to predict the one-step price changes in European market. Section 4.3 talks about longer term prediction, which is less biased by the artificial predictive power and is more useful for practitioners. Section 4.4 summarizes our findings and concludes the chapter with further discussions.

4.2 Deep learning experiment: one-step price change prediction

As the best bid (ask) prices are denoted as $P^{b,1}$ ($P^{a,1}$), we denote the mid-price P as the average of the best bid/ask price

$$P = 0.5(P^{b,1} + P^{a,1})$$

Order book dynamic is better expressed in “event time”. In the rest of the chapter an event is defined as a mid-price change.

4.2.1 Indicators

For the prediction we use some classical high frequency indicators of order book imbalance, order flow, and past returns as is mentioned in [Anane](#).

- **Order Book Imbalance.** Such indicators summarize the instantaneous buy-sell equilibrium. Though the limit orders are passive, they can still give an idea if the buy or sell force is more powerful in the market. We define the Order Book Imbalance at depth k as

$$OBI_t^k := \frac{\sum_{i=1}^k Q_t^{b,i} - \sum_{i=1}^k Q_t^{a,i}}{\sum_{i=1}^k Q_t^{b,i} + \sum_{i=1}^k Q_t^{a,i}}$$

To keep a memory of the indicator, we also compute the Exponential Moving Average (EMA) of the instantaneous indicator of delay γ as

$$OBI_t^{ema,\gamma} := \frac{Q_t^{b,ema,\gamma} - Q_t^{a,ema,\gamma}}{Q_t^{b,ema,\gamma} + Q_t^{a,ema,\gamma}}$$

where $Q_{t_i}^{b,ema,\gamma} = (1 - w)Q_{t_{i-1}}^{b,ema} + wQ_{t_i}^{b,1}$ and $Q_{t_i}^{a,ema} = (1 - w)Q_{t_{i-1}}^{a,ema,\gamma} + wQ_{t_i}^{a,1}$ with $w = \min(1, \frac{t_i - t_{i-1}}{\gamma})$. The smaller γ is, the shorter the memory is.

- **Log Return.** Short term prices may exhibit reverting or momentum effects. We define the past return of k steps as

$$PR_t^k = \log P_t - \log P_{t-k}$$

- **Order Flow.** As [Lu and Abergel \[2018b\]](#) show, market orders can largely influence future price change. On the one hand, the more intense the orders come from one side, the higher the pressure that price will be forced to move in this direction. On the other hand, one can consider it to be the market impact of successive orders that models should take into account. Denote the market order processes as Q_s , which is signed process with positive (negative) value represents trades on ask(bid) limits. We define the Order Flow with weight parameter β as

$$OF_t^\beta := \sum_{s \leq t} Q_s e^{-\beta(t-s)}$$

The bigger β is, the shorter the memory is.

We use a total of 20 indicators: $OBI^1, OBI^2, OBI^3, OBI^{ema,1}, OBI^{ema,5}, OBI^{ema,10}, OBI^{ema,60}, OBI^{ema,300}, OBI^{ema,600}, PR^1, PR^5, PR^{10}, PR^{20}, PR^{50}, PR^{100}, OF^{0.1}, OF^1, OF^{10}, OF^{100}, OF^{1000}$. The idea is to aggregate the information in some indicators instead of keeping all of the LOB states in a high-dimensional process.

In this section, we are specifically interested in predicting the 1-step return, that is to say if $\log P_{t+1} - \log P_t > 0$ or $\log P_{t+1} - \log P_t < 0$, with all the available information until t .

4.2.2 Model construction

Traditionally financial applications often use linear models because of their simplicity and the good performance in highly noisy environment. We first introduce some linear regression models as benchmarks and assess if deep neural network model can over-perform such classical models. Supervised learning problem is considered in all of our work. For a sample of n realisations, we denote X the input features of dimension p , and y denotes the target observations.

4.2.2.1 Linear models

In a linear regression model, it is assumed that

$$y = X\beta + \epsilon \tag{4.1}$$

where $y \in \mathbb{R}^n$ is the dependent variable, $x \in \mathbb{R}^{n \times p}$ is a p -vector of the regressors and $\beta \in \mathbb{R}^p$ is the vector of regression coefficients. $\epsilon \in \mathbb{R}^n$ denotes the error term that satisfies $\mathbb{E}[\epsilon|X] = 0$ and $Cov(\epsilon|X) = \sigma^2 I_n$.

The simplest way to estimate β is ordinary least squares (OLS) regression. The estimator is obtained by solving the optimization problem

$$\hat{\beta}^{OLS} = \arg \min_{\beta} \|y - X\beta\|_2^2$$

where $\|\cdot\|_2$ denotes the L_2 norm. This optimization has a closed form solution

$$\hat{\beta}^{OLS} = (X^T X)^{-1} X^T y$$

when $X^T X$ is invertible. And the obtained estimator is *BLUE* (best linear unbiased estimator) in the sense of variance minimization.

With the presence of multicollinearity between features, the OLS estimator will encounter the problem of high variance which is caused from the poor conditioning of the matrix $X^T X$. High correlations result in small eigenvalues of $X^T X$. When inverting that matrix, it can cause numerical instability and disproportionate variance of $\hat{\beta}$ due to finite sample. This troubles our application particularly because our three families of indicators are naturally correlated. For this end two classical penalisations of the variance of $\hat{\beta}$ are used, so we have two types of estimators: LASSO and Ridge, according to if we penalise the L_1 or L_2 norm.

$$\hat{\beta}^{LASSO} := \arg \min_{\beta} \frac{1}{2n} \|y - X\beta\|_2^2 + \alpha \|\beta\|_1$$

$$\hat{\beta}^{Ridge} := \arg \min_{\beta} \|y - X\beta\|_2^2 + \alpha \|\beta\|_2^2.$$

There is a closed-form solution to the Ridge estimator $\hat{\beta}^{Ridge} = (X^T X + \alpha I_p)^{-1} X^T y$. With regards to LASSO estimation, various methods exist to solve the problem numerically.

We also consider logistic regression a linear classification model that assumes the conditional probability of label $y \in \{-1, 1\}$ to be linearly dependent on X . [Gould and Bonart \[2016\]](#) also used this method to evaluate the relation between queue imbalance and the direction of the subsequent mid-price movement.

$$\mathbb{P}[Y = 1|X = x] = f(x^T w + c)$$

By taking f the sigmoid function $f(z) = \frac{1}{1+e^{-z}}$, it can be easily found that

$$\mathbb{P}(Y = y|X = x) = f(y(x^T w + c))$$

because $1 - f(z) = f(-z)$. For the observations $\{(x_i, y_i)\}_{i=1}^n$, w and c can be estimated via minimization of likelihood function

$$\begin{aligned} l(w, c) &= \sum_{i=1}^n \log f(y_i(x_i^T w + c)) \\ &= - \sum_{i=1}^n \log(1 + \exp(-y_i(x_i^T w + c))). \end{aligned}$$

And a similar L_1 or L_2 regularization can be applied to reduce the impact of multicollinearity of the parameters

$$\begin{aligned} (w, c) &= \arg \min \|w\|_1 + C \sum_{i=1}^n \log(1 + \exp(-y_i(x_i^T w + c))) \\ (w, c) &= \arg \min \frac{1}{2} \|w\|_2^2 + C \sum_{i=1}^n \log(1 + \exp(-y_i(x_i^T w + c))) \end{aligned}$$

4.2.2.2 LSTM model

A LSTM unit is composed of a cell(s_t), an input gate(i_t), a forget gate(f_t) and an output gate(o_t), where

- The memory cell stores the memory content of the LSTM unit,
- The forget gate is a filter to decide which old information is removed from the cell state,

- The input gate is used to specify the new information added to the cell state,
- The output gate controls which content from the cell state is used as output.

And the structure of a memory cell is illustrated in Figure 4.1. At timestep t , each of three gates is presented with the input x_t and the output h_{t-1} at the previous timestep $t-1$. More specifically, we have

$$\begin{aligned}
f_t &= \sigma(W_{f,x}x_t + W_{f,h}h_{t-1} + b_f), \\
i_t &= \sigma(W_{i,x}x_t + W_{i,h}h_{t-1} + b_i), \\
\tilde{s}_t &= \tanh(W_{\tilde{s},x}x_t + W_{\tilde{s},h}h_{t-1} + b_{\tilde{s}}), \\
s_t &= f_t \odot s_{t-1} + i_t \odot \tilde{s}_t, \\
o_t &= \sigma(W_{o,x}x_t + W_{o,h}h_{t-1} + b_o), \\
h_t &= o_t \odot \tanh(s_t),
\end{aligned}$$

where σ is the sigmoid function defined by $\sigma(z) = \frac{1}{1+e^{-z}}$; $\tanh$ is the hyperbolic tangent function defined by $\tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}}$; $\tilde{s}$ is the candidate values potentially added to the cell state s ; $\odot$ denotes the Hadamard product; $W_{\cdot}$ are the weight matrices and $b_{\cdot}$ are the bias vectors. The active functions are shown in Fig. 4.2.

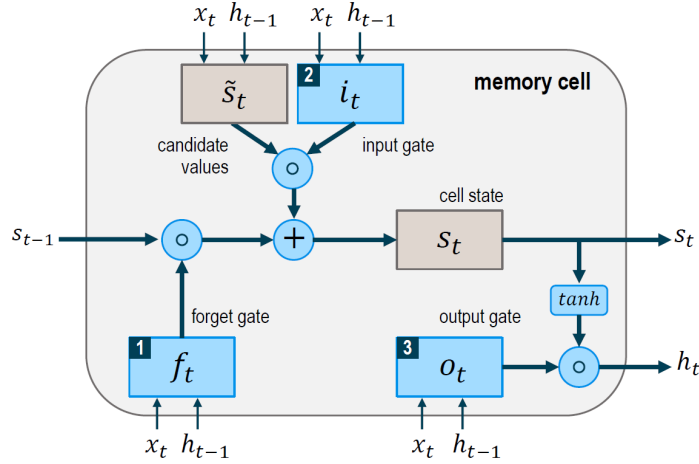


Figure 4.1: Structure of LSTM memory cell(Source: [Fischer and Krauss \[2017\]](#))

If there are h hidden units in the LSTM layer and i features in the input, then the number of parameters of the LSTM layer is:

$$\underbrace{4 \times hi}_{\text{weights } W_{\cdot,x}} + \underbrace{4 \times h^2}_{\text{weights } W_{\cdot,h}} + \underbrace{4 \times h}_{\text{bias vectors } b_{\cdot}} = 4h(i + h + 1).$$

We know that the dropout regularization is a method to prevent over-fitting in neural network. In the LSTM layer, as shown in [Gal and Ghahramani \[2016\]](#), we can replace (x_t, h_{t-1}) by $(x_t \odot z_x, h_{t-1} \odot z_h)$ with z_x, z_h random masks repeated at all time steps to realize the dropout regularization. “ $\odot$ ” denotes the Hadamard product, which is nothing but the element-wise product of two matrices. Moreover, early stopping is another regularization used to avoid over-fitting.

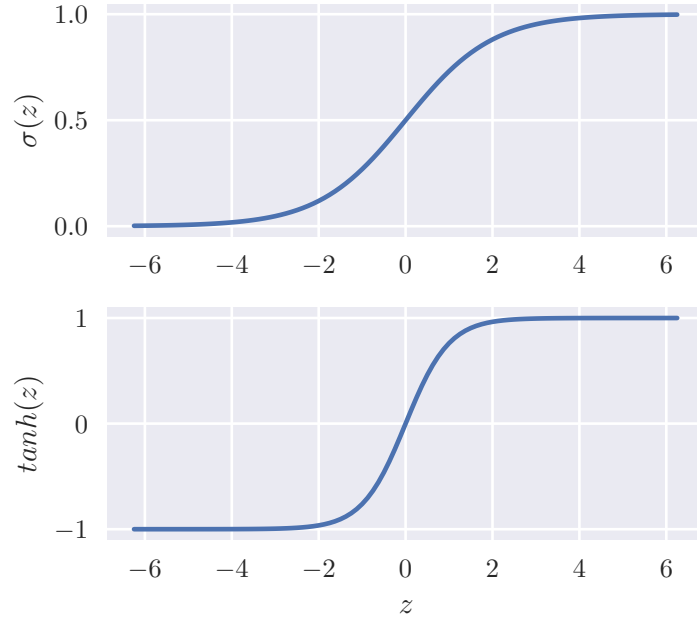


Figure 4.2: Active functions used in LSTM

4.2.3 Empirical results

We take the stocks of major continental European indices (CAC40, DAX, AEX) which count for about 100 stocks that are traded on two different exchanges (Euronext and Eurex). Data from a long period of February 2015 to April 2018 are acquired, and is split by September 2017 to be training sample and test sample. The indicators are calculated for each stock and the data are constructed by sampling on the mid-price change timestamps. The criterion for the following analysis is **accuracy**, that is the proportion of correct predictions to the number of total observations, i.e.

$$\frac{\#\{y\hat{y} > 0\}}{\#\{\hat{y} \neq 0\}}.$$

We only consider accuracy because the high frequency returns are balanced, that is to say we have about 50% $y > 0$ and 50% $y < 0$. The predictions are also verified to be balanced. In the following, without specification, we train one LSTM network of 25 hidden neurons for each stock on the training sample with cross-entropy loss function. The optimization algorithm is RMSprop as is described in [Tieleman and Hinton \[2012\]](#). All accuracies are evaluated on the test sample for out-of-sample performance.

4.2.3.1 Indicators vs Raw data

Though theoretically a neural network can approximate any function, with finite sample and with finite complexity of the network, it is possible that meaningful aggregation of the raw data into several indicators can reduce the noise without a loss of information and boost the prediction of the model. As for raw data, we refer to the limit quantities, one-step return and sum of signed order flow between price changes so that it is guaranteed the indicators are just the raw data transformed by some predetermined functions. We compare two models: a LSTM model with 10-steps of aggregated indicators and a LSTM model with 100-steps of raw data, because we roughly used 100 steps of raw data to build our indicators.

Figure 4.3 presents the boxplots of model trained on indicators and raw data as well as the histogram of the discrepancies for each stock. Though the accuracies are close, there is a clear, though small over performance of the former over the latter by around 1%. As we have already mentioned, in this large test sample, even 1% is significant. This test shows that constructed indicators can reduce the noise of the information and help improve the prediction model. However, it should not be ignored that the raw data LSTM model has a quite close accuracy and is still better than linear models. That is to say the network has been able to accomplish some of the indicator construction work implicitly.

Another advantage to use indicators instead of raw data is the decrease of computational time. The consumed time increases fast with the length of historical lengths. In our test, the 100-step raw data LSTM model takes about 7 times more time for the same conditions in the optimization.

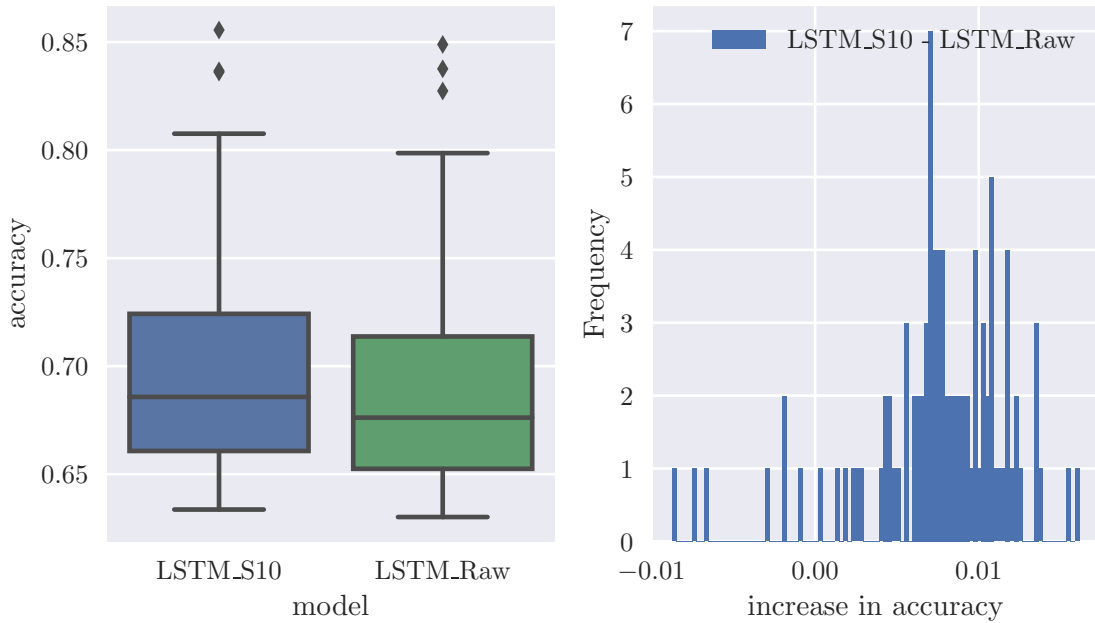


Figure 4.3: Comparison between indicators and raw data

4.2.3.2 Non-linearity

To compare the trained LSTM model to that of linear models, we calibrate 4 linear models: OLS, Ridge, LASSO and Logistic regressions. Except for logistic regression, all of the other methods provide a continuous predictor. So a threshold of 0 is used to decide whether the prediction was positive or negative.

We tested two model inputs: a model with only the current indicators (denoted OLS, LASSO, RIDGE, LOG and LSTM.S1) and with current and lagged indicators until 10-steps (denoted OLS10, RIDGE10 and LSTM_S10), in case that some short memory effects exist in the indicators to improve the accuracy. Figure 4.4 shows the box plots of accuracy of the models across different stocks. For each model, the central box shows the first and third quartiles by its bottom and top, with the horizontal line in the middle representing the median value. The bottom and top horizontal lines at the extremes of the vertical line are called whiskers. Noting the first and third quantile values as $Q1$ and $Q3$ respectively, the top whisker is the largest value that is smaller than $Q3 + 1.5 \times (Q3 - Q1)$ and the bottom whisker is the smallest value that is larger than

$Q1 - 1.5 \times (Q3 - Q1)$. Figure 4.5 shows the histogram of over performance of LSTM model to the OLS model.

We find that in both cases (current indicators and with lagged indicators) LSTM outperform largely OLS model in the accuracy by almost 2% in the current case and 3% in the 10-step lag case. The test sample is composed of 8 months of data, which is on average some millions of data points for each stock. In this case even a 1% improvement is significant and is very beneficial for trading. In addition, we plot a 3-d prediction function of the values before applying the final activation function in LSTM model conditional on OBI^1 and OF^{10} in Figure 4.6. We see a non-linear relation between the indicators and the predicted value functions.

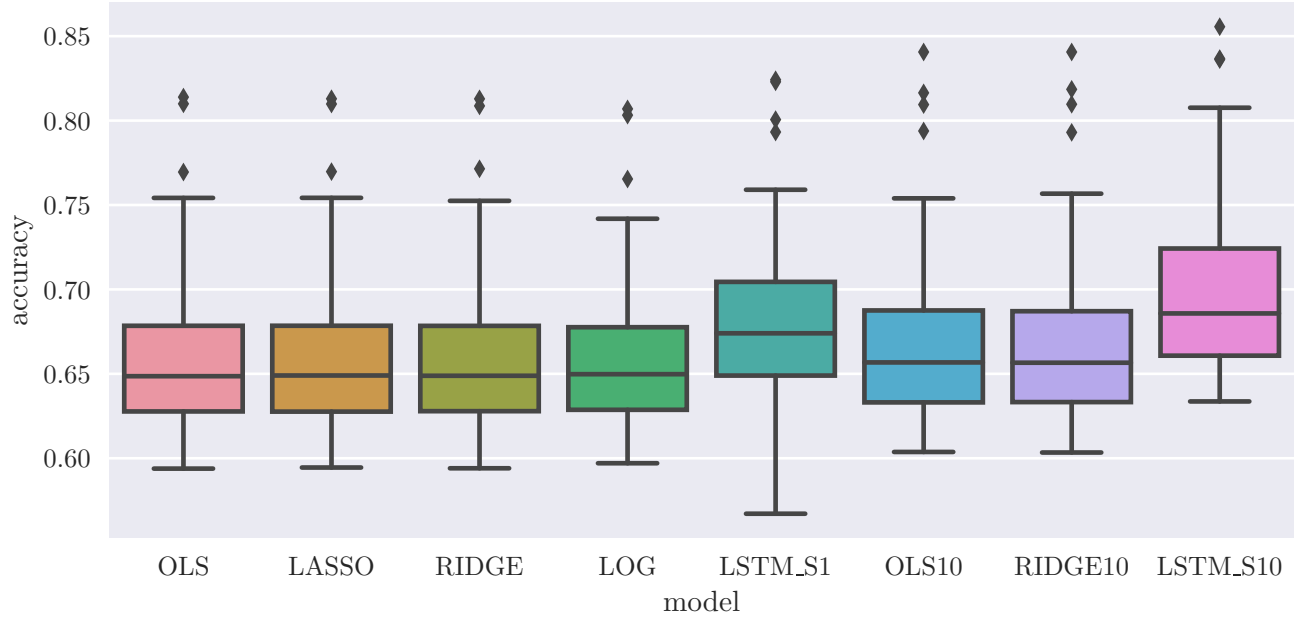


Figure 4.4: Comparison of accuracies for linear models and LSTM

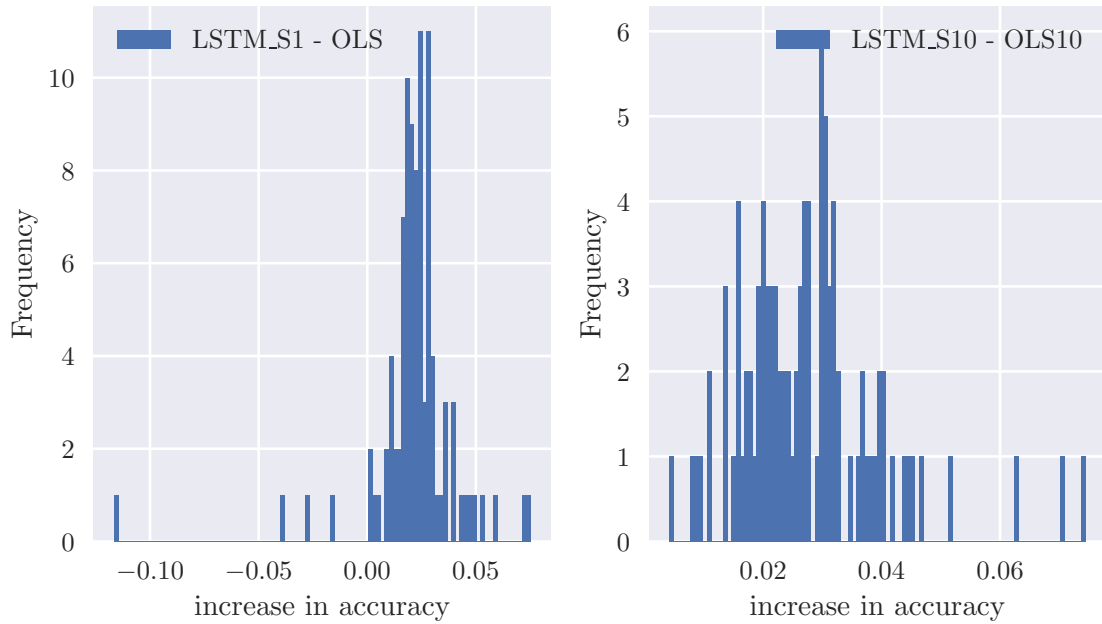


Figure 4.5: Histograms of increase in accuracy from OLS to LSTM

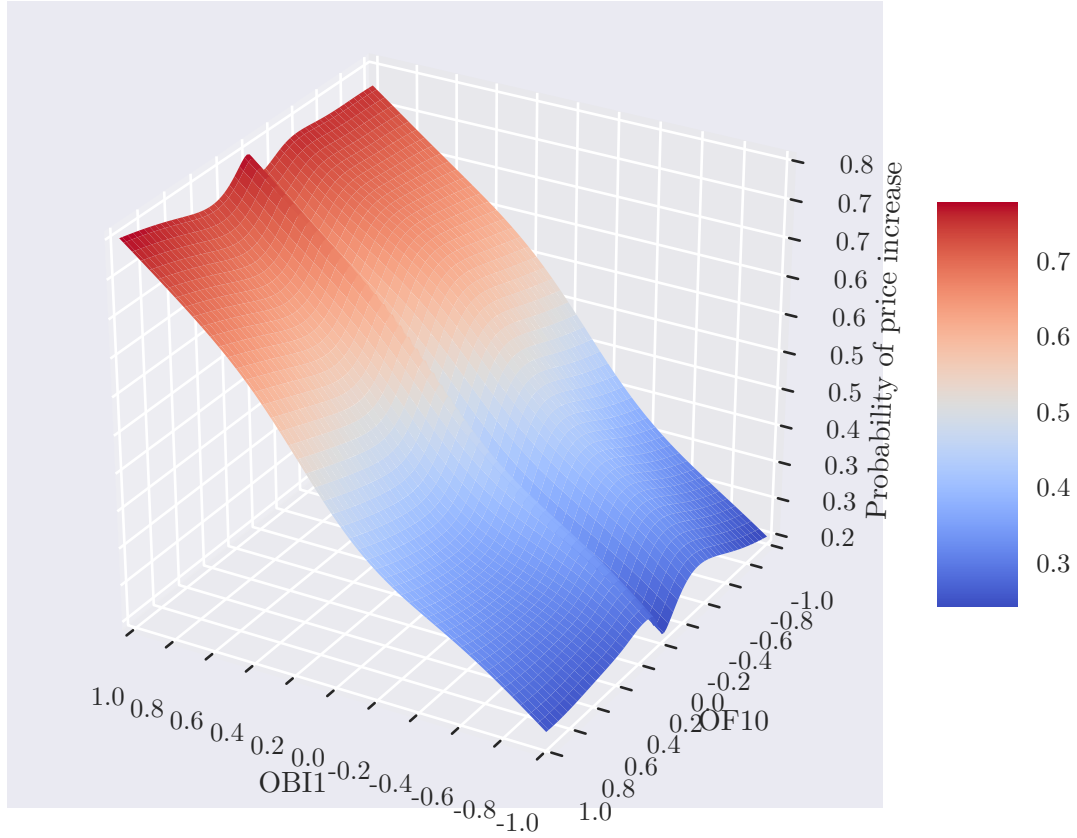


Figure 4.6: Projected value function of price increase probability with respect to OBI^1 and OF^{10} in LSTM model

4.2.3.3 Stationarity

The relation between the LOB states and the forward mid-price changes is said to be stationary if the accuracy of prediction is indifferent to the training periods. We separate our training sample of Feb 2015 to Aug 2017 into 10 periods to train a LSTM model on each period. Though the training samples are largely reduced, each still counts for 3 months which is reasonable in terms of size. The trained models are then applied on the whole test sample of Sep 2017 to Apr 2018. We illustrate the results of accuracy in Figure 4.7.

There is small tendency that the accuracy increases as the training period approaches the test period, where we see a 2% difference from the period 1 to the period 10. But as the size of the data in each training period is shorter, the variance of the model is also larger. It is not completely clear if such increase comes from the non-stationarity or the variance. The model trained on all of the periods outperforms even the model on the period 10. That shows the contribution of different periods to the prediction of the test periods. In the sense that old historical data are informative and help improve the performance of the model, the stationarity can be admitted.

4.2.3.4 Universality

All previous LSTM models are stock by stock, so it is possible to be trained on single machines. However if we turn to look for a single model for all of the stocks, even the raw data that

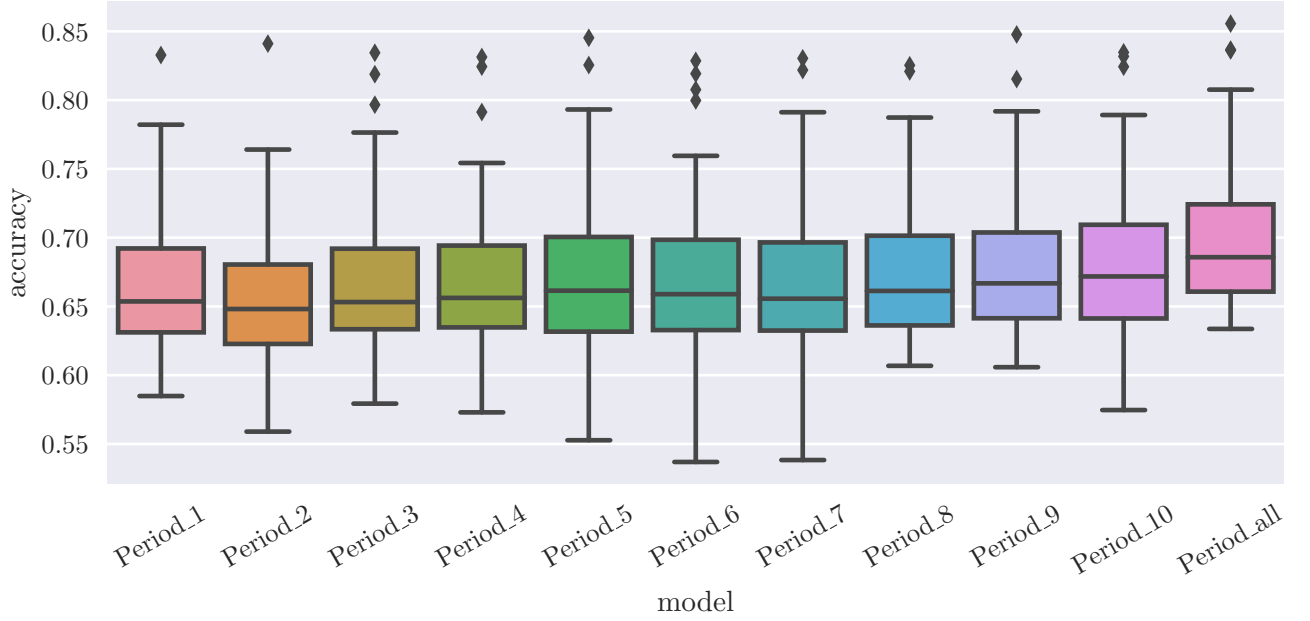


Figure 4.7: Boxplot of accuracy for different training periods

register the limit order book states, prices and transactions can account for more than 1 billion independent observations in total ($12,500$ per day $\times 800$ days $\times 100$ stocks), which makes it consume too long time even for only 1 epoch of gradient descent. So it is inevitable to use parallel algorithms for optimization.

We used the Delay Compensated Asynchronous Stochastic Gradient Descent (DC-ASGD) algorithm developed by Zheng et al. [2016] that improves the famous ASGD algorithm of Dean et al. [2012]. The ASGD algorithm is a host-worker system with each worker calculates the gradient of the network locally, reports the results to the host and the host decides how to update the network weights as well as give new parameters to the worker. Figure 4.8 illustrates the workflow schema of the ASGD method. However, it suffers from the *delayed gradient* problem, i.e. before a worker wants to update the global model; several other workers have already modified it with the gradients. DC-ASGD overcomes the problem by using Taylor expansion and Hessian approximation to compensate the delay. Pseudocode of this algorithm is provided in Appendix C.

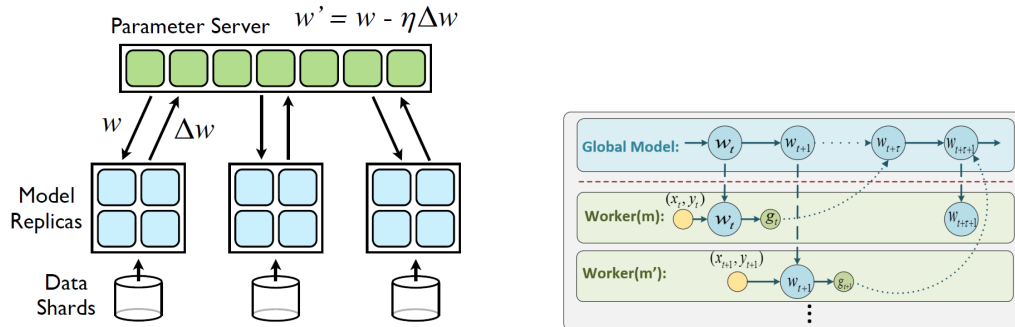


Figure 4.8: Diagram of Asynchronous Stochastic Gradient Descent. Source: Dean et al. [2012](left) and Zheng et al. [2016](right)

Surprisingly, we didn't find that such a universal model improves the accuracy to the stock specific model which is the main claim of [Sirignano and Cont \[2018\]](#). We thought about two possible reasons to explain these results:

- Market specificities. There are several different characteristics between US and European equity markets, such as the liquidity, trading models and types of participants. One issue that is particularly relevant is the tick size rules. In US stock exchanges, the minimal tick size is fixed to 0.01 dollar. However in European markets, the tick size is floating according to the price (and lately according to the liquidity as well based on MIFID II) so that the relative tick size is bounded in a certain range.
- More technically and concretely, it is reported in [Zheng et al. \[2016\]](#) that ASGD may decrease the training performance due to the parallelisation on different machines. Though [Sirignano and Cont \[2018\]](#) made a rigorous comparison to prove the superiority of the universal model using the same optimization algorithm, it was not compared to stock by stock model not using ASGD (it is not feasible to calibrate the universal model on a single machine). Such a stock-specic model can be even more useful in practice if it reaches similar or better performance.

To evaluate the deterioration of DC-ASGD with respect to the RMSprop that is originally used, we trained in the same single stock model with DC-ASGD algorithms of 1 and 25 workers and compare the respective accuracies. Figure 4.9 presents the results of the different models, which RMSprop, ASGD_1 and ASGD_25 are stock specific models and Univ_25 is a universal model trained with 25-worker ASGD. We observe a clear decrease of accuracy by increasing the number of workers in DC-ASGD, and using 25 workers lead to a loss of 5% of accuracy compared to 1 worker. If here we compare the stock by stock model trained with 25-worker DC-ASGD with the universal model, it is also worse. Such finding confirms the result of universality of the relation between LOB states and future price moves to some extent. However, training a universal model within a single machine consumes a very long time. So practically we have to consider the trade-off between the time-consumption and the accuracy improvement to decide if training a universal model is worthy, or a better parallel optimization algorithm has to be developed.

4.3 Multi-step deep learning regression

The mid-price change for one-step is artificially easy to predict due to the high auto-correlations of the mid-price change. This mean-reversion phenomenon results from the resubmission of liquidity when the first limits are consumed or cancelled. However, market makers cannot directly benefit from such predictions as it is not necessarily associated with trades. To better understand the difference, we presents an adverse selection example in Figure 4.10 where the mid-price is highly mean-reverting but the market maker may still lose money. The mid-price variation has an autocorrelation of -16.67%, and is principally stabilised between 100.5 and 102.5. However, as the market maker has sold at 101 and bought back at 102, they have lost 1 dollar with this pair of market making trades when finishing at 0 inventory at last.

Figure 4.11 presents the average forward return of executed limit orders with respect to logarithm of the number of mid-price changes. For each trade in the real data, if the trade is on the bid side, that is to say it is triggered by a market sell order, then the market maker has acquired a long position. The log-return of this position is $\log P_k^{mid} - \log P^{trade}$, with P_k^{mid} denoting the mid-price k changes after the trade. Accordingly, if the trade occurred on the ask side, the corresponding return for the market maker is $-(\log P_k^{mid} - \log P^{trade})$. It is not an accurate

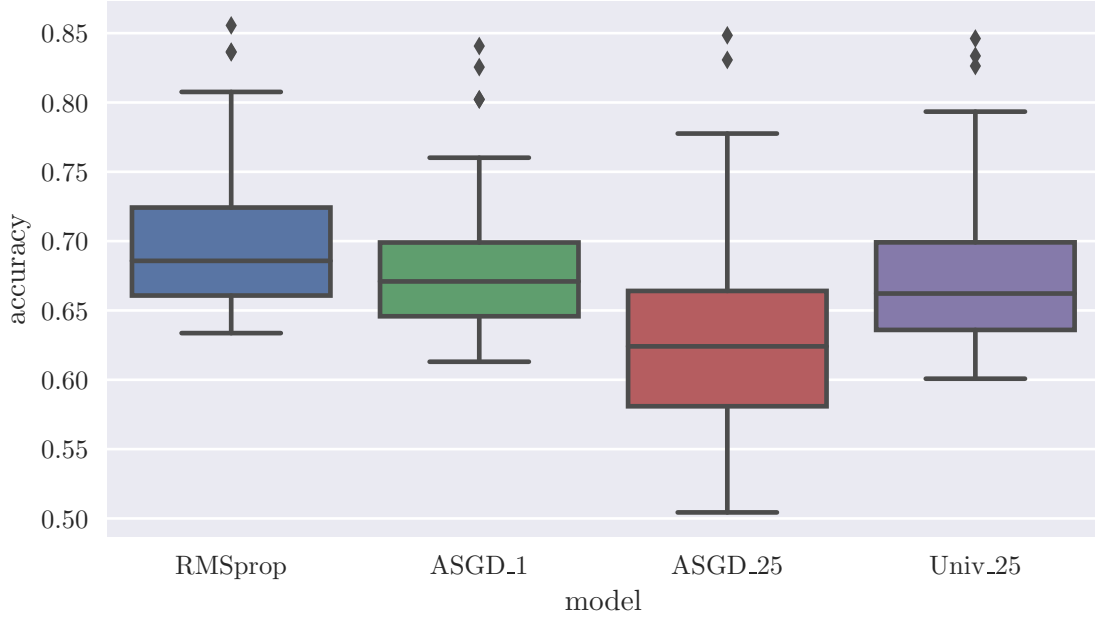


Figure 4.9: Comparison of ASGD and RMSprop

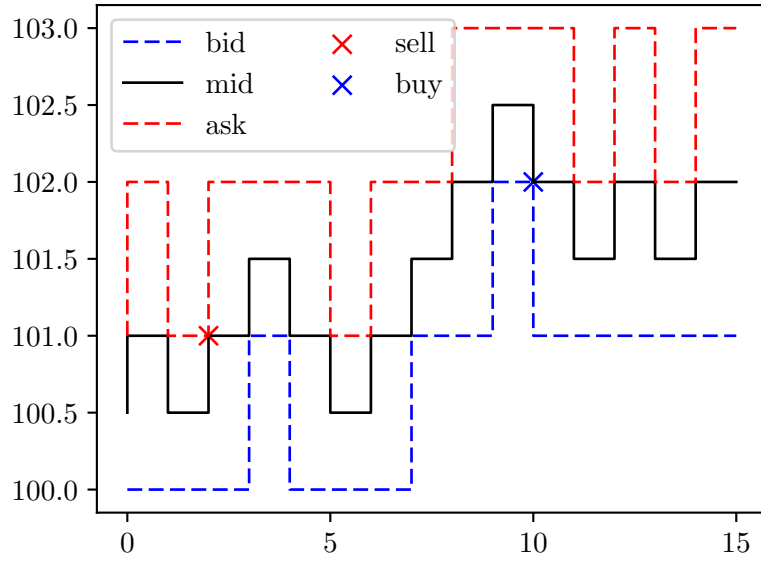


Figure 4.10: Illustrative adverse selection example.

measure of the P&L, but at least it approximates the profitability of the market makers on average if they are able to liquidate their position at mid after k price changes following the trade.

This return decreases very fast at short term, almost exponentially fast until 10 steps, and reverts a little at very long time horizon. There are fewer trades that totally consume the first limit than only partially the first limit. So immediately after a trade, the return basically represents the half-spread. With the absence of trades on the opposite side shortly after a transaction, the

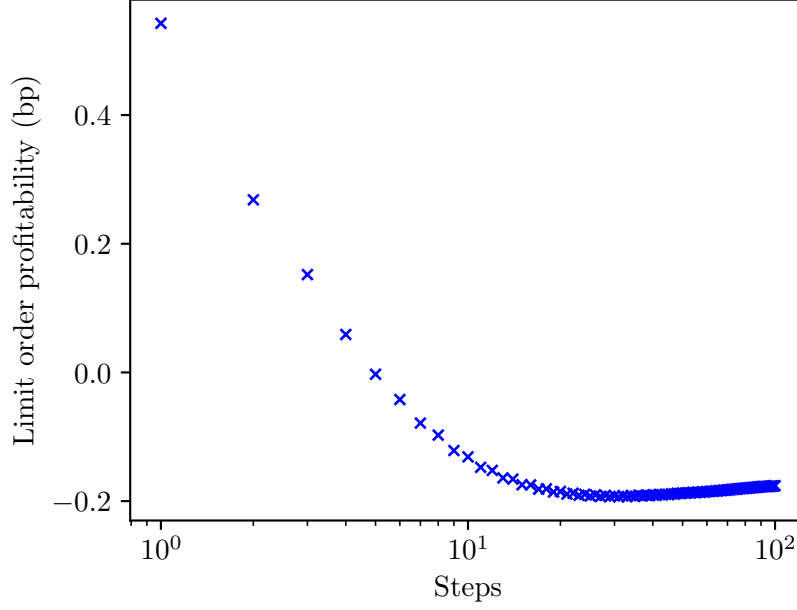


Figure 4.11: Mean signed limit order profitability after k -steps of mid-price change.

market makers are not able to really capture this profit. The long term profitability should be closer to the average performance in the market. Though the -0.2bp (basis point, 1% of 1%) seems surprising at first glance, it is much easier to understand when we take into account the asymmetric fees applied to liquidity providers and liquidity takers. The fee schedules are very complicated, which depends on market, types of orders, monthly volumes and etc. But a typical estimation is 0 bp for market makers and 0.5 bp for liquidity takers, which renders the final average profitability of each to be -0.2 bp and -0.3 bp respectively. Such close profitability also echoes the equilibrium as market participants have the choice to send limit orders or market orders. And as market orders guarantee an immediate execution, they are slightly less profitable as compensation.

Such practical observations incite us to search for longer horizon predictions. We first evaluate the prediction performance of linear models and LSTM models. A mid-to-mid strategy exploiting the predictions is then proposed in order to evaluate the practical profitability of each model.

4.3.1 Model performance

For a multi-step regression model, we replace the cross-entropy loss function used in the first two-class classification problem by the mean-squared error, which is equivalent to the objective function in OLS regression. We are first interested in the accuracy of the LSTM model compared to the linear model for 10-step and 100-step prediction. Figure 4.12 (boxplot of OLS10, LSTM10, OLS100, LSTM100) shows that the price changes are much harder to predict, with the median for 10 steps to be around 57% and 100 steps to be about 52%. In addition, the underperformance of OLS with respect to LSTM decreases dramatically when the horizon is longer. For 100-step prediction, the difference has almost disappeared.

Moreover, in practice we are more interested in the large values of predictions because we may regard them as “better opportunities” and keep a neutral point of view on the smaller values.

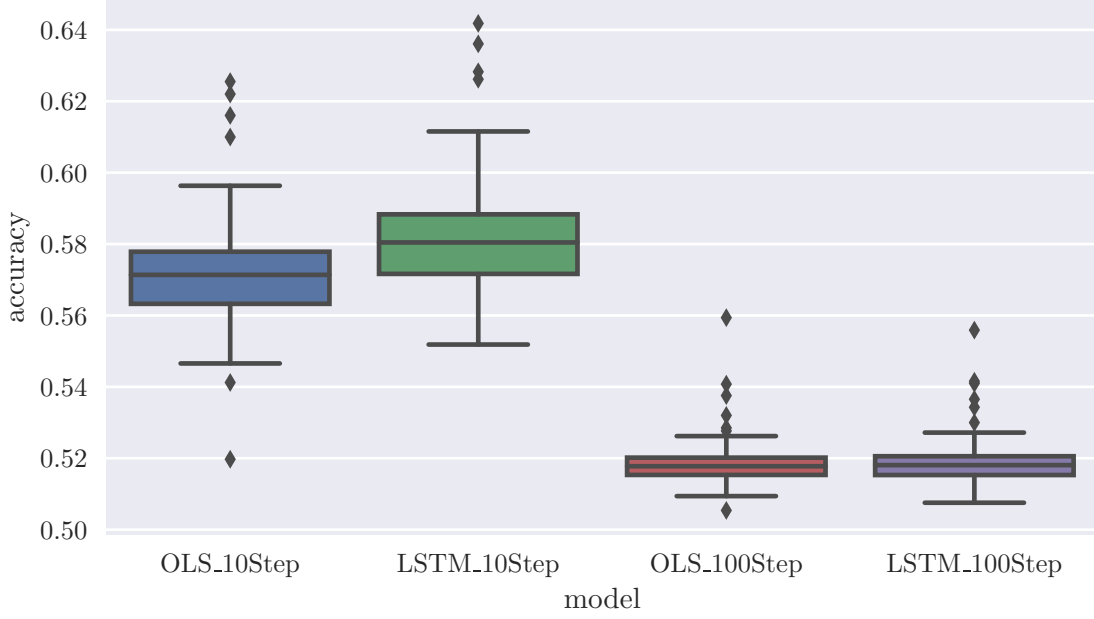


Figure 4.12: Accuracies of OLS and LSTM for 10- and 100-step return

To this end, we define for each stock the subset $|\hat{y}| > \mathbf{sd}(\hat{y}_{train})$ as the significant out-of-sample predictions, where $\mathbf{sd}(\cdot)$ is the sample standard deviation and $\hat{y}_{train}$ are the prediction values for the training set. The conditional accuracy is defined as

$$\frac{\#\{y\hat{y} > 0, |\hat{y}| \geq \mathbf{sd}(\hat{y}_{train}), y \neq 0\}}{\#\{|\hat{y}| \geq \mathbf{sd}(\hat{y}_{train}), y \neq 0\}}$$

Figure 4.13 illustrates the conditional accuracy of different models. All of the models have better conditional accuracies than their unconditional ones. The increase is higher for 10-step prediction than for 100-step prediction. It is interesting to see that at 10 steps, the LSTM is still around 2% better than the OLS model. Both results indicate the real application potential of the model in practice even at a more realistic horizon.

In addition, we define the conditional expectation

$$\mathbb{E}[y \cdot \text{sign}(\hat{y}) | |\hat{y}| \geq \mathbf{sd}(\hat{y}_{train})]$$

It is used to evaluate the expected profitability if each significant signal was used for trading. The results are shown in Figure 4.14. Consistent with previous findings, 100-step prediction is hard and provides a conditional expectation similar to that of OLS method at 10-step. At this horizon, the LSTM loses the edge compared to OLS method due to the high noise. However, LSTM has a conditional expectation 0.2 bp higher than OLS at 10-step horizon, which is very promising in practice.

To further investigate the decrease of accuracy with respect to the prediction horizon, we trained the OLS model and LSTM model in sample and out of sample for different steps from 1, 5, 10, 20, 50 and 100. “In-sample” means that the training and test sample are both the period of Sep 2017 to Apr 2018, which is the test sample in previous models. The evolution of accuracies are presented in Figure 4.15. All of the accuracies decrease with the prediction horizon, but LSTM model clearly deteriorates much faster out-of-sample than OLS model, whereas in-sample it is less the case. This result shows that the LSTM model are more prone to suffer from overfitting

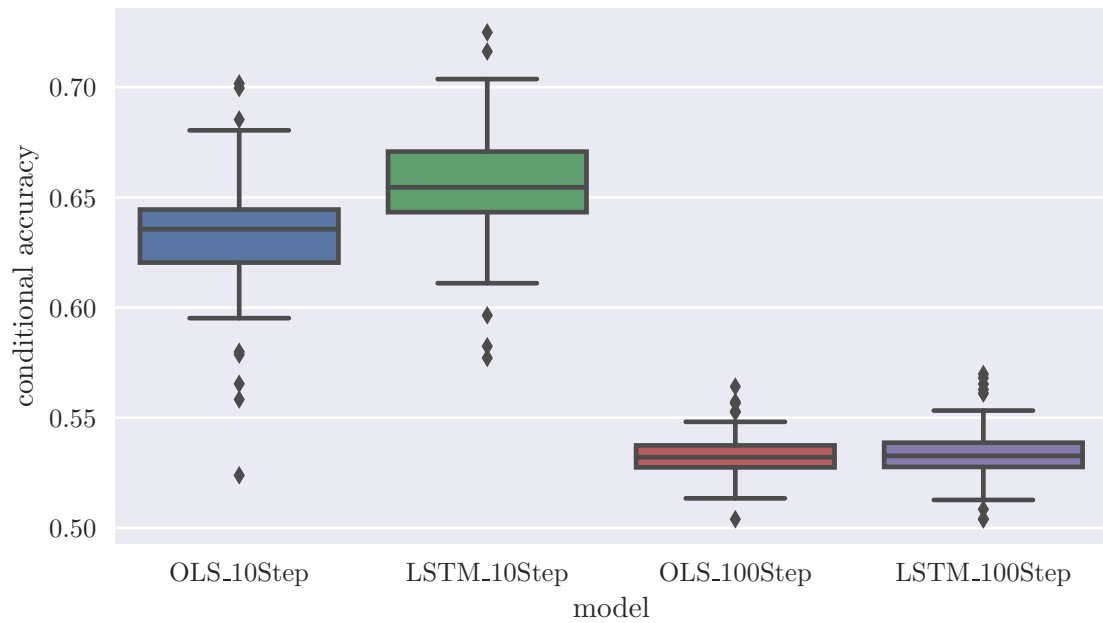


Figure 4.13: Conditional accuracy of OLS and LSTM for 10- and 100-step return

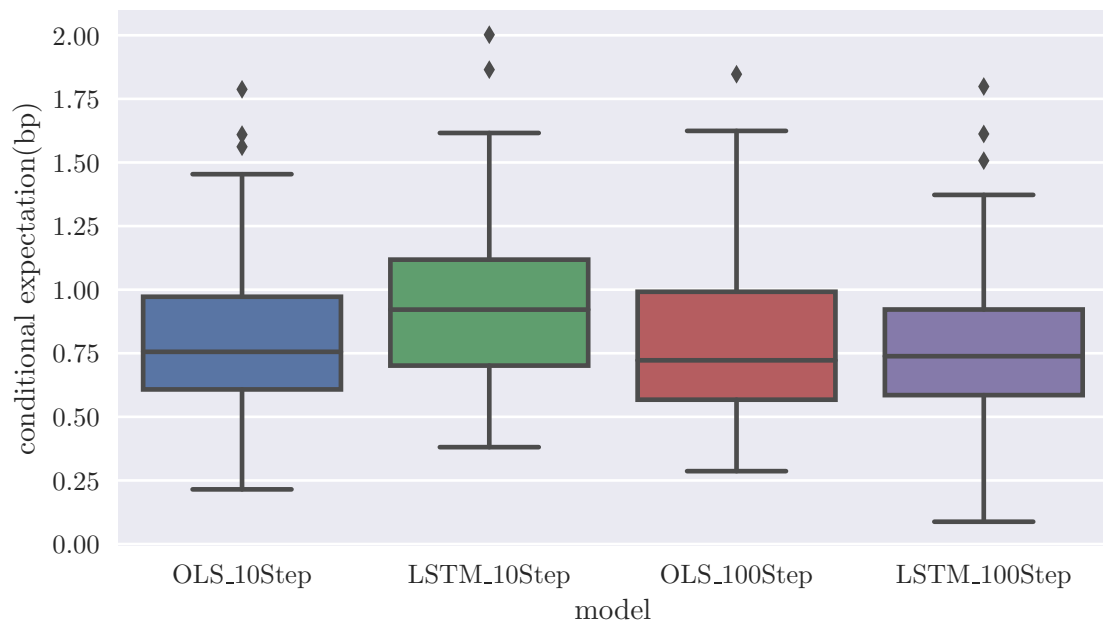


Figure 4.14: Boxplot of conditional expectation of OLS and LSTM for 10- and 100-step return

problems when the signal to noise ratio is too low, even though we are in the setting of tens of millions of training sample.

4.3.2 Strategy building

Although expected values and accuracies are interesting indicators for the predictive models, they are not so direct to be interpreted into potential trading performance. In order to better

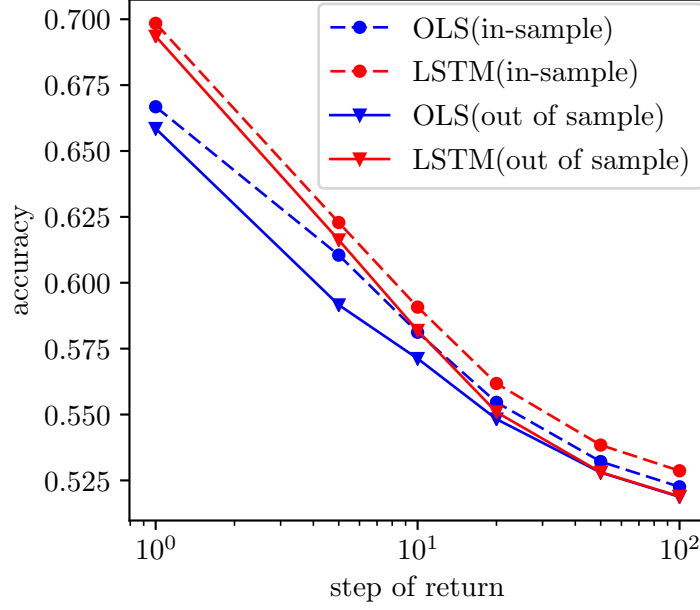


Figure 4.15: OLS and LSTM model in-sample and out of sample accuracies for different steps

evaluate the models in a practical context, we try to build a mid-to-mid trading strategy. Similar to conditional probability and expectation, we use $\mathbf{sd}(\hat{y}_{train})$ as a significance threshold, and we will enter into a long position (buy some stocks) if $\hat{y}_t > \mathbf{sd}(\hat{y}_{train})$ and oppositely into a short position (sell some stocks) if $\hat{y}_t < -\mathbf{sd}(\hat{y}_{train})$. We compute the following performance measures

1. **Gain(P&L)**: The total gain in Euros.
2. **Profitability (bp)**: The gain by unity of trading turnover.
3. **Holding period (HP)**: The average number mid price changes between a position is acquired and liquidated.

P&L is the natural measure for traders to estimate how much money a strategy can pocket in. It is also the most important indicator in the realisation of a strategy. Whereas as we are in a simulation environment, purely the simulation P&L is not enough to guarantee the realized P&L. In financial jargon, the cost between the simulation and realization of the P&L is called the slippage. Such cost comes from different ways, for example explicit cost like transaction fees and implicit cost like adverse selection, market impact, latencies and other discrepancies between the real market and simulation. The profitability is an indicator to evaluate if a strategy is robust facing the slippage. Besides, as we have pointed out very shortly after a trade the market maker earns a fake benefit when it is estimated with the mid-price, a too short average holding period will be largely influenced by this effect.

The predictions are highly noisy so that directly using the raw data results in very high trading turnover that decreases the profitability and shortens the holding period. To smooth the signal as well as reduce the impact from the short-term noise, we apply an *exponential moving average (EMA)* on the initial predictions. A strategy on EMA signals are constructed in a similar way as that of raw signal. With a significance threshold of $\mathbf{sd}(EMA(\hat{y}_{train}))$, a long or short position is taken if the EMA signal is greater than the threshold in absolute value.

Figure 4.16 illustrates the defined strategy for LSTM 10-step prediction raw signal and EMA signals for 15-mins trading between 16:45 and 17:00 of BNP PARIBAS on 1st September 2017. The first graph shows the mid-price change during that period. The second and third graphs present the raw and EMA signals as well as the corresponding significance thresholds respectively. Strategy positions are shown in the fourth graph. Graph 5 gives the trading turnovers of the two strategies. In this illustration, we observe a clear smoothing effect and the decrease of signal by applying an EMA on the signals.

We test 3 horizons of 1-step, 10-step and 100-step for both OLS model and LSTM model. The corresponding EMA signals are taken on event times, with a decay of $\tau = 50$ steps. That is to say, if we denote the raw signal as $\hat{y}_k$, the $\widehat{EMA}(y)_k$ is computed by

$$\widehat{EMA}(y)_k = \frac{1}{\tau} \hat{y}_k + (1 - \frac{1}{\tau}) \widehat{EMA}(y)_{k-1}$$

The medians of the measures for all of the strategies are presented in Table 4.1. A slippage of 0.5 bp is applied as an estimation of the trading cost, which should be rather conservative.

100-step models perform worse than 1-step and 10-step models. The profitability is lower and the P&L is lower as well. LSTM models deteriorate even more than OLS model, which is coherent with what we have found in accuracy analysis. The noise increases proportionally to the square root of the horizon, whereas the predictable price move does not grow (or grows very slowly) beyond a certain horizon. As the LSTM model performs worse in high noise context, the high noise at long horizon impairs it more than OLS model.

Strategies generated by raw signals all yield a low bp and very short holding periods of less than 2 mid-price changes, which is unrealistic in a real trading environment because it is hard to get opposite limit orders being executed at such short horizon to exit the position. Though the P&L before slippage is quite high, which can amount to nearly a thousand Euros per stock per day, the application of slippage totally kills it due to their low profitability. Strategies generated by EMA-smoothed signals largely reduce the noise of the signals. The profitability is boosted to beat the slippage so that all of the strategies are profitable after the cost, and the holding period is at least tens of mid-price changes. The results are much more robust to the short-term mid-price reversion artefact.

If we focus on the comparison of LSTM models and OLS models with EMA-smoothed signals, we find the former to be clearly superior to the latter for 1-step and 10-step predictions. Both the P&L before and after slippage are higher for LSTM models. Especially the P&L after slippage is doubled from 17 to 33 for 1-step and quadrupled from 6 to 24 for 10-step. Though there could still be some influence of holding periods, as it can be noticed that for corresponding horizon LSTM model is relatively shorter. But even though comparing LSTM_10 with OLS_1 whose holding periods are very close, the P&L after slippage is still 50 % higher. We can conclude the LSTM model outperforms linear model for building such a simple market making strategy based on short term mid-price change prediction.

We resume in Figure 4.17 the aggregated daily P&L performances of the strategies based on 10-step predictions by summing up all of the stock P&L for each day. The upper graph presents the raw signal OLS model and LSTM model before and after slippage while the lower one presents the same for EMA signals. Remarkably, we can finish up to almost half a million Euros by applying the LSTM strategy with maximum of 5,000 Euros position on each stock on half a year. In practice there are necessarily better way to profit from the improved predicting model to generate more P&L, by for example bearing more risk.

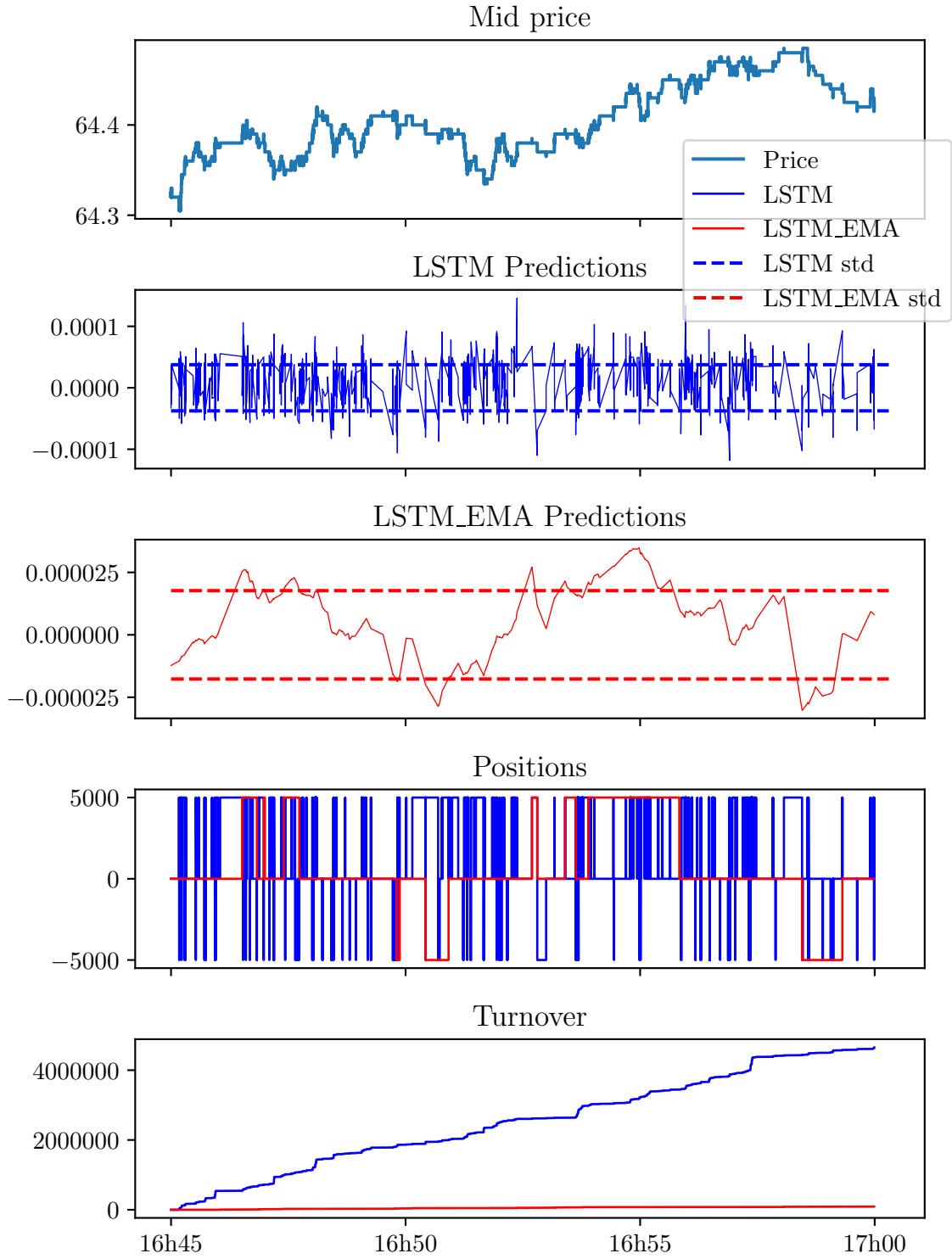


Figure 4.16: Illustration of market making strategy based on LSTM and LSTM.EMA 10-step prediction. Example of BNP PARIBAs on 1st September 2017

4.4 Conclusion

This chapter addresses the problem of short-term mid-price change prediction with LOB state information using LSTM and the practical applicability of the models.

Table 4.1: Medians of the strategy performance measures

		OLS_1	OLS_10	OLS_100	LSTM_1	LSTM_10	LSTM_100
Raw	bp	0.43	0.42	0.40	0.42	0.45	0.36
	HP	1.3	2.0	1.8	1.1	1.5	1.7
	P&L before slippage	894	572	622	1195	824	575
	P&L after slippage	-218	-147	-168	-188	-123	-198
EMA	bp	0.67	0.58	0.60	0.70	0.71	0.57
	HP	17.6	23.9	26.2	10.0	17.7	28.7
	P&L before slippage	97	65	64	172	103	57
	P&L after slippage	17	6	7	33	24	6

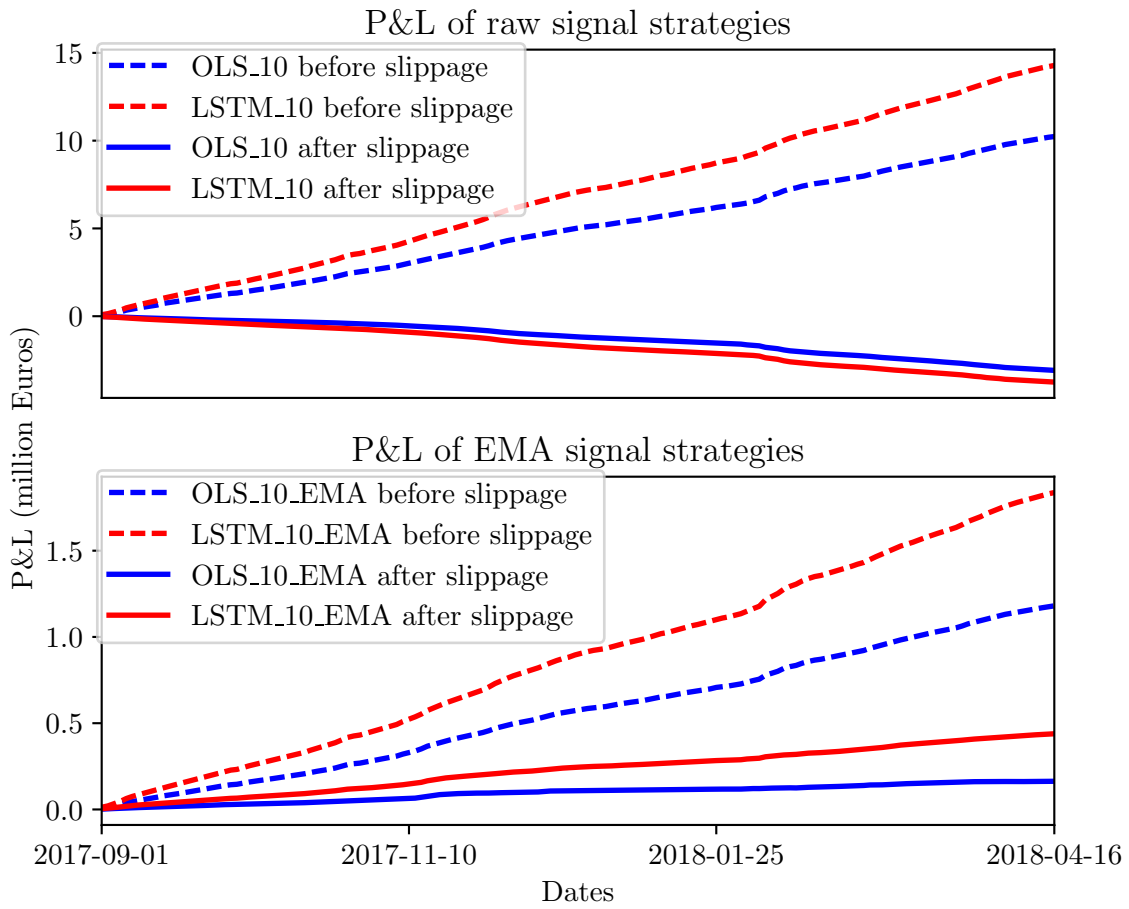


Figure 4.17: Cumulative P&L of 10-step strategies before and after slippage

The first part studies the problem of 1-step price move prediction on European continental stocks. We find that aggregating the LOB states into indicators improves the prediction accuracy and saves optimization time. Our study agrees with [Sirignano and Cont \[2018\]](#) that the LOB states and future price changes exhibit:

1. **Non-linearity** : Compared to linear models, the accuracy of LSTM is largely higher.
2. **Stationarity**: The models calibrated with old data give similar accuracies as recent data.

The prediction is time-stationary.

With respect to the universality, it is also found that the model calibrated on all of the stocks is better than a stock-by-stock model if the latter is calibrated using the same parallel optimization algorithm (ASGD). However, due to the limitation of this optimization algorithm, the universal model is not as good as a stock-by-stock model calibrated with a single-machine algorithm (RMSprop). Besides, the market specificities can also be a reason of less over-performance brought by universal model for European stocks.

In the second part, we show the insufficiency of 1-step return prediction, which is largely impacted by tick-level mean-reversion. Both LSTM and OLS models lose in accuracy for longer term predictions, but LSTM is more impacted by the increase of noise so that for 100-step prediction it is equivalent or worse than OLS model. A simple strategy that takes a position if the signal surpasses its historical standard deviation is used to evaluate the practical performance of Market Making strategies based on these predictive models. It is found that a smoothing by EMA is necessary to reduce the noise of the signals, and LSTM models out-perform OLS models for various measures.

Our study indicates that the LSTM model is better than linear model for short-term prediction with LOB data, but it is even more sensitive to longer horizon noise. A specifically interesting issue may be to learn how to choose the optimal horizon in terms of trade-off between the information content and noise.

General conclusions and Outlook

This thesis aims to enrich the research of market microstructure and market making strategies with empirical and applied studies.

In the first chapter, we proposed the modelling of order book events by high-dimensional non-linear Hawkes processes. Several statistical tests show the relevance of this model to reproduce the stylized facts of the financial data, notably the high- and low-frequency volatility. In addition, we have clarified the inhibitory effects between certain events besides excitation effects that are known much earlier. The numerical part of the optimization algorithm not only gives a concrete solution to calibrate a high-dimensional exponential kernel Hawkes process, but also proposes ways to improve ‘fine-tuning’ in a difficult global optimization problem.

The results in the second chapter suggest two simple but important enhancements to the *queue-reactive* model. The new model is more realistic in describing order book dynamics on the empirical distribution of first limit quantities and the price volatility. The simplicity of the model allows us to solve a problem of optimal market making strategy by the Markovian decision process theory. Simulations and backtests with market data show the relevance of the model as well as the real application potential of the strategy.

The results of the third chapter first confirmed the non-linearity, stationarity and universality of predicting 1-step price changes with order book information. We have also shown the need to predict the price at multiple steps, given the strong reversion artefact of the mid-price at 1-step. The LSTM model outperforms the 10-step OLS model as well, but becomes less well at 100-steps as it suffers the overfitting problem even more than OLS model when the signal-to-noise ratio is low. Simple market making strategies based on the prediction signals are tested and show that LSTM is always better than OLS at 1-step and 10-steps. The profitability of the LSTM strategy reaches 0.7 bp, which is above 0.5 bp of the approximate transaction cost, indicating its potential of exploitation in the practical activity of Market Making.

There are several research directions to pursue following these studies. For example, the model in Chapter 3 identified distributions of order quantities and the importance of limit removal orders, which typically correspond to the C^1 and M^1 orders of chapter 2. It is interesting to study how to model the limit order book with even larger Hawkes processes with a careful classification of order types and quantities. One of the difficulties is to be able to identify the most important types of events to be both more realistic and in a small dimension. Moreover, as the dimension becomes large, the events become necessarily sparser. It can also lead to technical problems of calibration.

In addition, the optimal strategy that we calibrated in chapter 3 was limited by rather tight boundary conditions of a pair of buy order and sell order. Ideally we would have considered a more general stochastic control modelling framework. The resolution of this control problem is rather complicated for its high dimension, with dimensions of inventory, time, quantities of both

first limits, and positions of the orders placed at the bid and ask limits. But it will be a great progress if this problem is solved.

In addition, we have only tested LSTM in chapter 4, but the machine learning algorithm family is much broader. We can try other algorithms in the same framework. Another question is what price to predict. The choice should be consistent with how it will be exploited in practice. The prediction of the mid-price is relevant if we consider the aggressive orders, but in this case an additional cost of a half bid-ask spread must be taken into account. For a Market Maker, they use mostly limit orders. So modelling the limit orders directly will be closer to reality.

I believe, and I hope, that this thesis brings novelties for both academic and industrial researchers. The academic community benefits from the empirically important observations, and some solutions using advanced mathematical tools for practical problems are provided. It also helped my team, *Automated Market Making* of *BNP Paribas*, to develop new Market Making models.

Finally, I sincerely thank all those who helped me to realize this work.

Conclusion générale et perspectives

Cette thèse a pour objectif d'enrichir l'étude de la microstructure du marché et des stratégies de Market Making d'un point de vue empirique et appliquée.

Dans le premier chapitre, nous avons proposé la modélisation des événements du carnet d'ordres par des processus de Hawkes non-linéaires en haute dimension. Plusieurs tests statistiques montrent la pertinence de ce modèle pour reproduire les faits stylisés des données financières, notamment la volatilité à haute et à basse fréquence. De plus nous avons éclairci les effets inhibiteurs entre certains événements à part des effets d'excitation qui sont plus connus avant. La partie numérique de l'algorithme d'optimisation non seulement donne une solution concrète pour calibrer un processus de Hawkes à noyaux exponentiel en haute dimension, mais aussi propose des pistes d'amélioration pour le "fine-tuning" dans un problème d'optimisation globale difficile.

Les résultats du deuxième chapitre proposent deux renforcements simples mais importants dans le modèle *queue-reactive*. Le nouveau modèle est plus réaliste pour décrire la dynamique du carnet d'ordres sur la distribution empirique des quantités de la première limite et la volatilité du prix. La simplicité du modèle nous permet de résoudre un problème de stratégie optimale de market making en utilisant la théorie du processus de décision Markovien. Les simulations et les backtests avec des données du marché montrent la pertinence du modèle ainsi que le potentiel d'applications réelles de la stratégie.

Les résultats du troisième papier ont d'abord confirmé la non-linéarité, la stationnarité et l'universalité de la prédiction des changements du prix en 1-étape avec les informations du carnet d'ordres. Nous avons aussi montré la nécessité de prédire le prix en plusieurs étapes, compte tenu l'artefact de réversion forte du prix mid en 1-étape. Le modèle de LSTM encore surperforme le modèle OLS en 10-étapes, mais deviens moins bien en 100-étapes car le modèle de LSTM subisse d'autant plus le problème d'overfitting que le modèle d'OLS quand le ratio signal-bruit est faible. Des stratégies simples de market making basées sur les signaux de prédictions sont testées et montrent que LSTM est toujours mieux qu'OLS pour la prédiction en 1-étape et 10-étapes. La rentabilité de la stratégie de LSTM atteint 0.7 bp, qui est au-dessus de 0.5 bp de coût de transaction approximative, indique le potentiel de l'exploitation dans l'activité pratique de Market Making.

Il y a plusieurs pistes de recherche à poursuivre à la suite de ces études. Par exemple, le modèle du Chapitre 3 a identifié les distributions des quantités d'ordres et l'importance des ordres d'éliminations, ces derniers correspondent typiquement aux ordres C^1 et M^1 du Chapitre 2. C'est intéressant d'étudier comment modéliser le carnet d'ordres par des processus de Hawkes de dimension d'encore plus élevée avec une classification attentive des types d'ordres et des quantités. Une des difficultés est de pouvoir identifier les types d'événements les plus importants pour à la fois être plus réaliste et avoir une dimension raisonnable. De plus, comme la dimension devient grand, les événements par classe deviennent forcément plus rares. Ça peut aussi poser

des questions techniques de calibration.

En outre, la stratégie optimale que nous avons calibrée dans le Chapitre 3 s'est limitée par des conditions aux bords assez restreintes d'une paire d'achat et de vente. Idéalement nous aurions considéré un cadre de modélisation plus général, du type de contrôle stochastique. La résolution de ce problème de contrôle est assez compliquée pour sa haute dimension, avec les dimensions de l'inventaire, du temps, des quantités des premières limites, et des positions des ordres placés à l'achat et à la vente. Mais ce serait un grand succès si le problème sera résolu.

De plus, nous n'avons que testé LSTM dans le Chapitre 4, mais la famille des algorithmes de machine learning est beaucoup plus large. Nous pouvons sans doute essayer d'autres algorithmes dans le même cadre. Une autre question est quel prix à prédire dans le modèle. Le choix devrait être cohérent avec la façon de l'exploiter en pratique. La prédiction du prix mid est pertinente si les ordres agressifs sont envisagés, mais dans ce cas-là il faut prendre en compte un coût supplémentaire conséquent d'un demi-*spread* entre la première limite à l'achat et à la vente. Pour un Market Maker, ils utilisent la plupart du temps les ordres limites. Donc une modélisation directement sur les ordres limites sera plus proche de la réalité.

J'estime, et j'espère que cette thèse apporte des nouveautés à la fois aux chercheurs académiques et industriels. La communauté académique bénéficie des observations empiriquement importantes, et quelques solutions à l'aide des outils mathématiques avancés sont proposées. Elle a aussi aidé mon équipe d'accueil, *Automated Market Making* du *BNP Paribas*, à développer des nouveaux modèles de Market Making.

Enfin, je remercie sincèrement à tous ceux qui m'ont aidé à réaliser ce travail.

Bibliography

- [MiF] *MiFID II, Article 17.4.*
- [Abergel et al., 2016] Abergel, F., Anane, M., Chakraborti, A., Jedidi, A., and Toke, I. M. (2016). *Limit order books*. Cambridge University Press.
- [Abergel et al., 2017] Abergel, F., Huré, C., and Pham, H. (2017). Algorithmic trading in a microstructural limit order book model. *arXiv preprint arXiv:1705.01446*.
- [Abergel and Jedidi, 2015] Abergel, F. and Jedidi, A. (2015). Long-time behavior of a Hawkes process-based limit order book. *SIAM Journal on Financial Mathematics*, 6:1026–1043.
- [Ahmed et al., 2010] Ahmed, N. K., Atiya, A. F., Gayar, N. E., and El-Shishiny, H. (2010). An empirical comparison of machine learning models for time series forecasting. *Econometric Reviews*, 29(5-6):594–621.
- [Anane] Anane, M. *Une approche mathématique de l’investissement boursier*. PhD thesis, CentraleSupélec.
- [Andersen et al., 2000] Andersen, T. G., Bollerslev, T., Diebold, F. X., and Labys, P. (2000). Great realizations. *Risk*, 13:105–108.
- [Antweiler and Frank, 2004] Antweiler, W. and Frank, M. Z. (2004). Is all that talk just noise? the information content of internet stock message boards. *The Journal of finance*, 59(3):1259–1294.
- [Avellaneda and Lee, 2010] Avellaneda, M. and Lee, J.-H. (2010). Statistical arbitrage in the us equities market. *Quantitative Finance*, 10(7):761–782.
- [Avellaneda and Stoikov, 2008] Avellaneda, M. and Stoikov, S. (2008). High frequency trading in a limit order book. *Quantitative Finance*, 8(3):217–224.
- [Bacry et al., 2013] Bacry, E., Delattre, S., Hoffmann, M., and Muzy, J.-F. (2013). Modelling microstructure noise with mutually exciting point processes. *Quantitative Finance*, 13(1):65–77.
- [Bacry et al., 2016] Bacry, E., Jaisson, T., and Muzy, J.-F. (2016). Estimation of slowly decreasing Hawkes kernels: application to high-frequency order book dynamics. *Quantitative Finance*, 16(8):1179–1201.
- [Bacry et al., 2015] Bacry, E., Mastromatteo, I., and Muzy, J.-F. (2015). Hawkes processes in finance. *Market Microstructure and Liquidity*, 1(01).
- [Bacry and Muzy, 2014] Bacry, E. and Muzy, J.-F. (2014). Hawkes model for price and trades high-frequency dynamics. *Quantitative Finance*, 14(7):1147–1166.
- [Bäuerle and Rieder, 2011] Bäuerle, N. and Rieder, U. (2011). *Markov decision processes with applications to finance*. Springer Science & Business Media.

- [Bayraktar and Ludkovski, 2014] Bayraktar, E. and Ludkovski, M. (2014). Liquidation in limit order books with controlled intensity. *Mathematical Finance*, 24(4):627–650.
- [Bouchaud et al., 2009] Bouchaud, J.-P., Farmer, J. D., and Lillo, F. (2009). How markets slowly digest changes in supply and demand. In Hens, T. and Schenk-Hoppé, K. R., editors, *Handbook of Financial Markets: Dynamics and Evolution*, pages 57–156. Elsevier, San Diego.
- [Bowsher, 2007] Bowsher, C. G. (2007). Modelling security market events in continuous time: Intensity based, multivariate point process models. *Journal of Econometrics*, 141(2):876–912.
- [Brémaud and Massoulié, 1996] Brémaud, P. and Massoulié, L. (1996). Stability of nonlinear Hawkes processes. *Ann. Probab.*, 24:1563–1588.
- [Brogaard et al., 2014] Brogaard, J., Hendershott, T., and Riordan, R. (2014). High-frequency trading and price discovery. *The Review of Financial Studies*, 27(8):2267–2306.
- [Cartea et al., 2018] Cartea, Á., Donnelly, R., and Jaimungal, S. (2018). Enhancing trading strategies with order book signals. *Applied Mathematical Finance*, pages 1–35.
- [Cartea and Jaimungal, 2013a] Cartea, Á. and Jaimungal, S. (2013a). Modeling asset prices for algorithmic and high frequency trading. *Mathematics and financial economics*, 20:512–547.
- [Cartea and Jaimungal, 2013b] Cartea, Á. and Jaimungal, S. (2013b). Risk metrics and fine tuning of high frequency trading strategies. *Mathematical Finance*, page doi: 10.1111/mafi.12023.
- [Cartea et al., 2014] Cartea, Á., Jaimungal, S., and Ricci, J. (2014). Buy low sell high: A high frequency trading perspective. *SIAM Journal of Financial Mathematics*, 5(1):415–444.
- [Chakraborti et al., 2011] Chakraborti, A., Muni Toke, I., Patriarca, M., and Abergel, F. (2011). Econophysics review: I. empirical facts. *Quantitative Finance*, 11(7):991–1012.
- [Cont and de Larrard, 2013] Cont, R. and de Larrard, A. (2013). Price dynamics in a Markovian limit order book market. *SIAM Journal on Financial Mathematics*, 4:1–25.
- [Cont et al., 2014] Cont, R., Kukanov, A., and Stoikov, S. (2014). The price impact of order book events. *Journal of financial econometrics*, 12(1):47–88.
- [Cont et al., 2010] Cont, R., Stoikov, S., and Talreja, R. (2010). A stochastic model for order book dynamics. *Operations Research*, 58:549–563.
- [Das et al., 2016] Das, S., Mullick, S. S., and Suganthan, P. (2016). Recent advances in differential evolution—an updated survey. *Swarm and Evolutionary Computation*, 27:1–30.
- [Dean et al., 2012] Dean, J., Corrado, G., Monga, R., Chen, K., Devin, M., Mao, M., Senior, A., Tucker, P., Yang, K., Le, Q. V., et al. (2012). Large scale distributed deep networks. In *Advances in neural information processing systems*, pages 1223–1231.
- [Eisler et al., 2012] Eisler, Z., Bouchaud, J.-P., and Kockelkoren, J. (2012). The price impact of order book events: market orders, limit orders and cancellations. *Quantitative Finance*, 12(9):1395–1419.
- [Fischer and Krauss, 2017] Fischer, T. and Krauss, C. (2017). Deep learning with long short-term memory networks for financial market predictions. *European Journal of Operational Research*.
- [Fodra and Pham, 2013] Fodra, P. and Pham, H. (2013). Semi markov model for market microstructure.
- [Fodra and Pham, 2015] Fodra, P. and Pham, H. (2015). High frequency trading and asymptotics for small risk aversion in a markov renewal model. Preprint ArXiv.

- [Gal and Ghahramani, 2016] Gal, Y. and Ghahramani, Z. (2016). A theoretically grounded application of dropout in recurrent neural networks. In *Advances in neural information processing systems*, pages 1019–1027.
- [Gers et al., 1999] Gers, F. A., Schmidhuber, J., and Cummins, F. (1999). Learning to forget: Continual prediction with lstm.
- [Gopikrishnan et al., 2000] Gopikrishnan, P., Plerou, V., Gabaix, X., and Stanley, H. E. (2000). Statistical properties of share volume traded in financial markets. *Physical Review E*, 62(4):R4493.
- [Gould and Bonart, 2016] Gould, M. D. and Bonart, J. (2016). Queue imbalance as a one-tick-ahead price predictor in a limit order book. *Market Microstructure and Liquidity*, 2(02):1650006.
- [Gu et al., 2018] Gu, S., Kelly, B. T., and Xiu, D. (2018). Empirical asset pricing via machine learning.
- [Gueant and Lehalle, 2015] Gueant, O. and Lehalle, C.-A. (2015). General intensity shapes in optimal liquidation. *Mathematical Finance*, 25(3):457–495.
- [Guilbaud and Pham, 2013a] Guilbaud, F. and Pham, H. (2013a). Optimal high frequency trading in a pro-rata microstructure with predictive information. *Mathematical Finance*, page doi: 10.1111/mafi.12042.
- [Guilbaud and Pham, 2013b] Guilbaud, F. and Pham, H. (2013b). Optimal high-frequency trading with limit and market orders. *Quantitative Finance*, 13:79–94.
- [Guéant et al., 2012] Guéant, O., Lehalle, C.-A., and Fernandez-Tapia, J. (2012). Optimal portfolio liquidation with limit orders. *SIAM Journal on Financial Mathematics*, 3(1):740–764.
- [Guéant et al., 2013] Guéant, O., Lehalle, C.-A., and Fernandez-Tapia, J. (2013). Dealing with the inventory risk: a solution to the market making problem. *Mathematics and financial economics*, 7:477–507.
- [Hawkes and Oakes, 1974] Hawkes, A. G. and Oakes, D. (1974). A cluster process representation of a Self-Exciting process. *Journal of Applied Probability*, 11:493–503.
- [Heaton et al., 2016] Heaton, J., Polson, N., and Witte, J. H. (2016). Deep learning in finance. *arXiv preprint arXiv:1602.06561*.
- [Ho and Stoll, 1981] Ho, T. and Stoll, H. R. (1981). Optimal dealer pricing under transactions and return uncertainty. *The Journal of Trading*, 9:47–73.
- [Hochreiter and Schmidhuber, 1997] Hochreiter, S. and Schmidhuber, J. (1997). Long short-term memory. *Neural computation*, 9(8):1735–1780.
- [Hornik et al., 1989] Hornik, K., Stinchcombe, M., and White, H. (1989). Multilayer feedforward networks are universal approximators. *Neural networks*, 2(5):359–366.
- [Huang et al., 2015a] Huang, W., Lehalle, C.-A., and Rosenbaum, M. (2015a). Simulating and analyzing order book data: The queue-reactive model. *Journal of the American Statistical Association*, 110:107–122.
- [Huang et al., 2015b] Huang, W., Lehalle, C.-A., and Rosenbaum, M. (2015b). Simulating and analyzing order book data: The queue-reactive model. *Journal of the American Statistical Association*, 110(509):107–122.
- [Huang and Rosenbaum, 2017] Huang, W. and Rosenbaum, M. (2017). Ergodicity and diffusivity of markovian order book models: a general framework. *SIAM Journal on Financial Mathematics*, 8(1):874–900.

- [Hult and Kiessling, 2010] Hult, H. and Kiessling, J. (2010). *Algorithmic trading with Markov chains*. PhD thesis, Doctoral thesis, Stockholm University, Sweden.
- [Jaisson, 2015] Jaisson, T. (2015). Market impact as anticipation of the order flow imbalance. *Quantitative Finance*, 15(7):1123–1135.
- [Kim, 2003] Kim, K.-j. (2003). Financial time series forecasting using support vector machines. *Neurocomputing*, 55(1-2):307–319.
- [Krauss, 2017] Krauss, C. (2017). Statistical arbitrage pairs trading strategies: Review and outlook. *Journal of Economic Surveys*, 31(2):513–545.
- [Large, 2007] Large, J. (2007). Measuring the resiliency of an electronic limit order book. *Journal of Financial Markets*, 10:1–25.
- [Lehalle and Mounjid, 2017] Lehalle, C.-A. and Mounjid, O. (2017). Limit order strategic placement with adverse selection risk and the role of latency. *Market Microstructure and Liquidity*, 3(01):1750009.
- [Li and Zhang, 2011] Li, Y.-l. and Zhang, J. (2011). A differential evolution algorithm with dynamic population partition and local restart. In *Proceedings of the 13th annual conference on Genetic and evolutionary computation*, pages 1085–1092. ACM.
- [Lu and Abergel, 2018a] Lu, X. and Abergel, F. (2018a). High-dimensional hawkes processes for limit order books: modelling, empirical analysis and numerical calibration. *Quantitative Finance*, 18(2):249–264.
- [Lu and Abergel, 2018b] Lu, X. and Abergel, F. (2018b). Order-book modelling and market making strategies. *arXiv preprint arXiv:1806.05101*.
- [Massoulié, 1998] Massoulié, L. (1998). Stability results for a general class of interacting point processes dynamics, and applications. *Stochastic Processes and their Applications*, 75:1–30.
- [Moro et al., 2009] Moro, E., Vicente, J., Moyano, L. G., Gerig, A., Farmer, J. D., Vaglica, G., Lillo, F., and Mantegna, R. N. (2009). Market impact and trading profile of large trading orders in stock markets.
- [Muni Toke, 2015] Muni Toke, I. (2015). The order book as a queueing system: average depth and influence of the size of limit orders. *Quantitative Finance*, 15(5):795–808.
- [Muni Toke, 2017] Muni Toke, I. (2017). Reconstruction of order flows using aggregated data. *Market microstructure and liquidity*.
- [Nelson et al., 2017] Nelson, D. M., Pereira, A. C., and de Oliveira, R. A. (2017). Stock market’s price movement prediction with lstm neural networks. In *Neural Networks (IJCNN), 2017 International Joint Conference on*, pages 1419–1426. IEEE.
- [O’Connor et al., 2010] O’Connor, B., Balasubramanyan, R., Routledge, B. R., Smith, N. A., et al. (2010). From tweets to polls: Linking text sentiment to public opinion time series. *Icwsn*, 11(122-129):1–2.
- [Ozaki, 1979] Ozaki, T. (1979). Maximum likelihood estimation of Hawkes’ self-exciting point processes. *Annals of the Institute of Statistical Mathematics*, 31:145–155.
- [Patel et al., 2015] Patel, J., Shah, S., Thakkar, P., and Kotecha, K. (2015). Predicting stock market index using fusion of machine learning techniques. *Expert Systems with Applications*, 42(4):2162–2172.

- [Radivojević et al., 2014] Radivojević, T., Anselmi, J., and Scalas, E. (2014). Ergodic transition in a simple model of the continuous double auction. *PloS one*, 9(2):e88095.
- [Rambaldi et al., 2017] Rambaldi, M., Bacry, E., and Lillo, F. (2017). The role of volume in order book dynamics: a multivariate hawkes process analysis. *Quantitative Finance*, 17(7):999–1020.
- [Rubin, 1972] Rubin, I. (1972). Regular point processes and their detection. *IEEE Transactions on Information Theory*.
- [Scalas et al., 2006] Scalas, E., Kaizoji, T., Kirchler, M., Huber, J., and Tedeschi, A. (2006). Waiting times between orders and trades in double-auction markets. *Physica A: Statistical Mechanics and its Applications*, 366:463–471.
- [Scalas et al., 2017] Scalas, E., Rapallo, F., and Radivojević, T. (2017). Low-traffic limit and first-passage times for a simple model of the continuous double auction. *Physica A: Statistical Mechanics and its Applications*, 485:61–72.
- [Selvin et al., 2017] Selvin, S., Vinayakumar, R., Gopalakrishnan, E., Menon, V. K., and Soman, K. (2017). Stock price prediction using lstm, rnn and cnn-sliding window model. In *Advances in Computing, Communications and Informatics (ICACCI), 2017 International Conference on*, pages 1643–1647. IEEE.
- [Sirignano and Cont, 2018] Sirignano, J. and Cont, R. (2018). Universal features of price formation in financial markets: perspectives from deep learning.
- [Smith et al., 2003] Smith, E., Farmer, J. D., Gillemot, L., and Krishnamurthy, S. (2003). Statistical theory of the continuous double auction. *Quantitative Finance*, 6:481–514.
- [Stoikov, 2017] Stoikov, S. (2017). The micro-price.
- [Storn and Price, 1995] Storn, R. and Price, K. (1995). *Differential evolution—a simple and efficient adaptive scheme for global optimization over continuous spaces*, volume 3. ICSI Berkeley.
- [Stübinger and Bredthauer, 2017] Stübinger, J. and Bredthauer, J. (2017). Statistical arbitrage pairs trading with high-frequency data. *International Journal of Economics and Financial Issues*, 7(4).
- [Stübinger and Endres, 2018] Stübinger, J. and Endres, S. (2018). Pairs trading with a mean-reverting jump–diffusion model on high-frequency data. *Quantitative Finance*, pages 1–17.
- [Talih and Hengartner, 2005] Talih, M. and Hengartner, N. (2005). Structural learning with time-varying components: tracking the cross-section of financial time series. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 67(3):321–341.
- [Tetlock, 2007] Tetlock, P. C. (2007). Giving content to investor sentiment: The role of media in the stock market. *The Journal of finance*, 62(3):1139–1168.
- [Tetlock et al., 2008] Tetlock, P. C., Saar-Tsechansky, M., and Macskassy, S. (2008). More than words: Quantifying language to measure firms’ fundamentals. *The Journal of Finance*, 63(3):1437–1467.
- [Tieleman and Hinton, 2012] Tieleman, T. and Hinton, G. (2012). Lecture 6.5—RmsProp: Divide the gradient by a running average of its recent magnitude. COURSERA: Neural Networks for Machine Learning.
- [Toth et al., 2012] Toth, B., Eisler, Z., Lillo, F., Kockelkoren, J., Bouchaud, J.-P., and Farmer, J. D. (2012). How does the market react to your order flow? *Quantitative Finance*, 12(7):1015–1024.

- [Wyart et al., 2008] Wyart, M., Bouchaud, J.-P., Kockelkoren, J., Potters, M., and Vettorazzo, M. (2008). Relation between bid–ask spread, impact and volatility in order-driven markets. *Quantitative Finance*, 8(1):41–57.
- [Yang et al., 2013] Yang, M., Cai, Z., Li, C., and Guan, J. (2013). An improved adaptive differential evolution algorithm with population adaptation. In *Proceedings of the 15th annual conference on Genetic and evolutionary computation*, pages 145–152. ACM.
- [Yang and Zhu, 2016] Yang, T.-W. and Zhu, L. (2016). A reduced-form model for level-1 limit order books. *Market Microstructure and Liquidity*, 2(02):1650008.
- [Zheng et al., 2014] Zheng, B., Roueff, F., and Abergel, F. (2014). Modelling bid and ask prices using constrained Hawkes processes: Ergodicity and scaling limit. *SIAM Journal on Financial Mathematics*, 5:99–136.
- [Zheng et al., 2016] Zheng, S., Meng, Q., Wang, T., Chen, W., Yu, N., Ma, Z.-M., and Liu, T.-Y. (2016). Asynchronous stochastic gradient descent with delay compensation for distributed deep learning. *arXiv preprint arXiv:1609.08326*.
- [Zhu, 2015] Zhu, L. (2015). Large deviations for Markovian nonlinear Hawkes processes. *The Annals of Applied Probability*, 25(2):548–581.

Appendix A

Pseudocode of basic Differential Evolution

Algorithm 1 Differential Evolution algorithm

```
1: Input. Maximum total generation  $G$ , population size  $N \geq 4$ , mutation factor  $F \in (0, 2)$ ,  
   crossover rate  $CR \in (0, 1)$ , parameter domain  $\Omega$ , termination criteria.  
2: Output. optimal point (optimal function value, termination generation etc.).  
3: // Initialization phase  
4:  $g=1$ ; Initialize the initial population  $(x_{1,1}, \dots, x_{N,1})$  randomly such that  $x_{i,1} \in \Omega$ ;  
5: while  $g \leq G$  and termination criteria not met do  
6:   for  $i \leftarrow 1, N$  do  
7:     // Mutation  
8:     Choose randomly  $r_1, r_2$  and  $r_3$  in  $\llbracket 1, N \rrbracket$  such that  $i, r_1, r_2$  and  $r_3$  are distinct;  
9:     Construct donor  $v_{i,g+1} \leftarrow x_{r_1,g} + F(x_{r_2,g} - x_{r_3,g})$ ;  
10:    // Crossover. Construct trial element  $u_{i,g+1}$   
11:     $I_{rand}$  is a random integer from  $\llbracket 1, D \rrbracket$ ;  
12:    for  $j \leftarrow 1, D$  do  
13:       $rand_{j,i} \sim \mathcal{U}(0, 1)$ ;  
14:      if  $rand_{j,i} \leq CR$  or  $j = I_{rand}$  then  
15:         $u_{j,i,g+1} \leftarrow v_{j,i,g+1}$ ;  
16:      else  
17:         $u_{j,i,g+1} \leftarrow x_{j,i,g}$ ;  
18:      end if  
19:    end for  
20:    //  $I_{rand}$  ensures that  $u_{i,g+1} \neq x_{i,g}$   
21:    // Selection  
22:    if  $f(u_{i,g+1}) \leq f(x_{i,g})$  then  
23:       $x_{i,g+1} \leftarrow u_{i,g+1}$ ;  
24:    else  
25:       $x_{i,g+1} \leftarrow x_{i,g}$ ;  
26:    end if  
27:  end for  
28:   $g \leftarrow g + 1$ ;  
29: end while
```

Appendix B

Markov decision processes and optimal strategies

Let $(X_n)_{n=0}^\infty$ be a Markov chain in discrete time on a countable state space $\mathbb{S}$ with transition matrix P . Let $\mathbf{A}$ be a finite set of possible actions. Every action can be classified as either continuation action or termination action. The set of continuation actions is denoted $\mathbf{C}$ and the set of termination actions $\mathbf{T}$.

$$\mathbf{C} \cup \mathbf{T} = \mathbf{A} \quad \mathbf{C} \cap \mathbf{T} = \emptyset$$

The Markov chain is terminated when a termination action is selected.

Every action is not available in every state of the chain. Let $A : \mathbb{S} \mapsto 2^{\mathbf{A}}$ be a function associating a non-empty set of actions $A(s)$ to each state $s \in \mathbb{S}$. $2^{\mathbf{A}}$ is the power set consisting of all subsets of $\mathbf{A}$. The set of continuation actions available in state s is denoted $\mathbf{C}(s) = A(s) \cap \mathbf{C}$ and the set of termination $\mathbf{T}(s) = A(s) \cap \mathbf{T}$. For each $s, s' \in \mathbb{S}$ and $a \in A(s)$ the transition probability from s to s' when selecting action a is denoted $P_{ss'}(a)$.

For every action there are associated values. The value of continuation is denoted $v_C(s, a)$, which can be non-zero only when $a \in \mathbf{C}(s)$. The value of termination is denoted $v_T(s, a)$, it can be non-zero only when $a \in \mathbf{T}(s)$. It is assumed that both v_C and v_T are non-negative and bounded.

A policy $\alpha = (\alpha_0, \alpha_1, \dots)$ is a sequence of functions: $\alpha_n : \mathbb{S}^{n+1} \mapsto \mathbf{A}$ such that $\alpha_n(s_0, \dots, s_n) \in A(s_n)$ for each $n \geq 0$ and $(s_0, \dots, s_n) \in \mathbb{S}^{n+1}$.

The expected total value starting in $X_0 = s$ and following a policy α until termination is denoted by $V(s, \alpha)$. It can be interpreted as the expected payoff of a strategy starting from state s by take the policy α . The purpose of Markov decision theory is to analyse optimal policies and optimal expected values. A policy α_* is called optimal if, for all states $s \in \mathbb{S}$ and policies α ,

$$V(s, \alpha_*) \geq V(s, \alpha)$$

The optimal expected value V_* is defined by $V_*(s) = \sup_{\alpha} V(s, \alpha)$

If an optimal policy α_* exists, then $V_*(s) = V(s, \alpha_*)$. It is proved in [Hult and Kiessling \[2010\]](#) that, if all policies terminate in finite time with probability 1, an optimal policy α_* exists and the optimal expected value is the unique solution to a Bellman equation. Furthermore, the optimal policy α_* is stationary, that is to say the policy does not change with time. The optimal values

as well as the associated stationary optimal policies can be approached by a recursive algorithm Algorithm 1.

Algorithm 2 Optimal strategy value approximation

Input. Tolerance TOL, transition matrix P , state space $\mathbb{S}$, continuation actions $\mathbf{C}$, termination strategies $\mathbf{T}$, continuation value v_C , termination value v_T .

Output. Lower bound of optimal value V_n and almost optimal policy α_n .

while $d > \text{TOL}$ **do**

 Put

$$V_n(s) = \max\left(\max_{a \in \mathbf{C}(s)} v_C(s, a) + \sum_{s' \in \mathbb{S}} P_{ss'}(a) V_{n-1}(s'), \max_{a \in \mathbf{T}(s)} v_T(s, a)\right)$$

and

$$d = \max_{s \in \mathbb{S}} V_n(s) - V_{n-1}(s), \text{ for } s \in \mathbb{S}$$

$$n = n + 1$$

end while

Define $\alpha : \mathbb{S} \mapsto \mathbf{C} \cup \mathbf{T}$ as a maximizer to

$$\max\left(\max_{a \in \mathbf{C}(s)} v_C(s, a) + \sum_{s' \in \mathbb{S}} P_{ss'}(a) V_{n-1}(s'), \max_{a \in \mathbf{T}(s)} v_T(s, a)\right)$$

B.1 Keep or cancel strategy for buying one unit

Denote X_n the order book state after n transitions and X_0 the initial state. An agent wants to buy one unit at price $j_0 < j^A(X_0)$. After each market transition, the agent can choose between keeping the limit order or cancelling it and submit a market buy order at the best ask level $j^A(X_n)$ if the price is lower than the predetermined stop loss price $J > j^A(X_0)$. If j^A reaches J before the agent's order is executed, he cancels the bid order and places a market order at J to fulfil the trade. It is assumed that there are always sufficient limit orders at level J .

Denote Y_n the position of the limit order of the agent, and (O_n^r, q_n^r) the last limit removal order type and quantity, where n is the number of event orders from time 0. $S_n = (X_n, Y_n, O_n^r, q_n^r)$ is still a Markov chain in $\mathbb{S} \subset \mathbb{N}^{2K} \times \{0, 1, 2, \dots\} \times \{O^M, O^C\} \times \mathbb{N}$ where $Y_n \leq X_n^{j_0}$.

The generator matrix of S is denoted $W = (W_{ss'})$. The jump chain associated with the process is denoted as $S = (S_n)_{n=0}^\infty$. The jump chain is of greater importance in the model. Without ambiguity, we will use S to represent both the continuous process and the jump chain. The transition matrix of the jump chain is denoted $P = P_{ss'}$.

Let $s = (x, y) \in \mathbb{S}$. There are three possible cases:

- $y < 0$ and $j^A(x) < J$. Then the possible continuation action is $\mathbf{C}(s) = \{0\}$, representing waiting for next market transition. And the possible termination action is $\mathbf{T}(s) = \{-1\}$, representing cancellation of the limit order and submission of market order at $j^A(x)$.
- $y < 0$ and $j^A(x) = J$. The process terminates as the ask price reaches stop loss price. The limit order is cancelled and a market order is submitted at $j^A(x) = J$, represented by $\mathbf{T}(s) = \{-1\}$. And $\mathbf{C}(s) = \emptyset$.

- $y = 0$. The process terminates with the execution of the limit order, represented by $\mathbf{T}(s) = \{-2\}$. And $\mathbf{C}(s) = \emptyset$.

The 1 tick hypothesis is important here for the boundary conditions. Once $J^B = J - 2$ and $J^A = J - 1$, and the ask is cleared, without the hypothesis we should have terminated the process, except that we have ignored the possibility of the price reversion by a new ask order, so that the execution probability of the market maker's bid order is underestimated.

The expected value (cost), interpreted as the expected saving with respect to stop loss price, is given by

$$V_\infty(s, \alpha) = \begin{cases} \sum_{s' \in \mathbb{S}} P_{ss'} V_\infty(s', \alpha) & , \quad \alpha(s) = 0 \\ \pi^J - \pi^{j^A(x)} & , \quad \alpha(s) = -1 \\ \pi^J - \pi^{j_0} & , \quad \alpha(s) = -2 \end{cases} \quad (\text{B.1})$$

The waiting value is zero. The value function could then be approximated by Algorithm with the iteration

$$\begin{aligned} V_{n+1}(s) &= \max\left(\max_{a \in \mathbf{C}(s)} \sum_{s' \in \mathbb{S}} P_{ss'}(a) V_n(s'), \max_{a \in \mathbf{T}(s)} v_T(s, a)\right) \\ &= \begin{cases} \max(\sum_{s' \in \mathbb{S}} P_{ss'} V_n(s'), \pi^{j^A}(s)) & , \quad \text{for } y > 0, j^A < J \\ \pi^J - \pi^{j^A} = 0 & , \quad \text{for } y > 0, j^A = J \\ \pi^J - \pi^{j_0} & , \quad \text{for } y = 0 \end{cases} \end{aligned}$$

B.2 Market making (Making the spread)

The extended Markov chain here is defined as $(X_n, Y_n^0, Y_n^1, j_n^0, j_n^1, O_n^r, q_n^r)$, where $Y_n^0(Y_n^1)$ is the positions of the market maker's bid(ask) order in the bid(ask) price level $j_n^0(j_n^1)$. We have both Y_n^0 and Y_n^1 are non-increasing, and

$$X_n^{j_n^0} \geq Y_n^0 \geq 0 \quad X_n^{j_n^1} \geq Y_n^1 \geq 0$$

The market maker predetermines a best buy level $J^{B^0} < j^A(X_0)$, a worst level $J^{B^1} > j^A(X_0)$, a best sell level $J^{A^1} > j^B(X_0)$ and a worst sell level $J^{A^0} < j^B(X_0)$. As in the buy one strategy, the order is cancelled and executed at the stop loss price if the corresponding best limit price reaches the worst price level. And it is assumed that the execution at the stop loss price is always available. The state space is defined on $\mathbb{S} \subset \mathbb{N}^d \times \{0, 1, 2, \dots\} \times \{0, 1, 2, \dots\} \times \{J^{B^0}, \dots, J^{B^1} - 1\} \times \{J^{A^0} + 1, \dots, J^{A^1}\} \times \{O^M, O^C\} \times \mathbb{N}$.

The possible actions in this strategy are:

- The market maker can choose to wait for next market transition and cancel both orders before any of the orders is executed
- When one of the orders has been executed, the market maker has one order on the opposite side waiting for execution. The market maker follows a buy(sell)-one-unit strategy.

Let $V_{\infty}^B(x, y, j, o^r, q^r)$ denote the optimal(minimal) expected buy price (cost rather than value) in state (x, y, j, o^r, q^r) for buying one unit, with best buy level J^{B^0} and worst level J^{B^1} . Similarly, $V_{\infty}^A(x, y, j)$ denotes the optimal (maximal) expected sell price in state (x, y, j, o^r, q^r) for selling one unit, with best sell level J^{A^1} and worst sell level J^{A^0} . The optimal expected value is then given by

$$V_{\infty}(s) \begin{cases} \max(\sum_{s' \in \mathbb{S}} P_{ss'} V_{\infty}(s'), 0) & , \text{ for } y^0 > 0, y^1 > 0 \\ \pi^{j^1} - V_{\infty}^B(x, y^0, j^0, o^r, q^r) & , \text{ for } y^0 > 0, y^1 = 0 \\ V_{\infty}^A(x, y^1, j^1) - \pi^{j^0, o^r, q^r} & , \text{ for } y^0 = 0, y^1 > 0 \end{cases}$$

An extended version is also possible, to take the change of limit price into consideration. Under this strategy, the available actions for the market maker are:

- Before any of the orders is executed, the market maker can choose from waiting for next transition, cancel both orders or cancel either order and resubmit at new levels k^0 et k^1 .
- When one of the orders have been processed, the outstanding limit order is proceeded according to the buy(sell)-one-unit strategy. And the price level is also renewable after each market transition.

In this strategy, the optimal expected value is determined by

$$V_{\infty}(s) \begin{cases} \max(\sum_{s' \in \mathbb{S}} P_{ss'} V_{\infty}(s'), \max V_{\infty}(s_{k^0 k^1}), 0) & , \text{ for } y^0 > 0, y^1 > 0 \\ \pi^{j^1} - V_{\infty}^B(x, y^0, j^0, o^r, q^r) & , \text{ for } y^0 > 0, y^1 = 0 \\ V_{\infty}^A(x, y^1, j^1, o^r, q^r) - \pi^{j^0} & , \text{ for } y^0 = 0, y^1 > 0 \end{cases}$$

where $s_{k^0 k^1}$ describes the cancel and resubmit of one limit order.

Appendix C

Pseudocode of DC-ASGD algorithm

Algorithm 3 DC-ASGD: worker m

Input. learning rate η
while *True* **do**
 Pull $\mathbf{w}_t$ from the parameter server
 Compute delta weight $\Delta\mathbf{w}_m = -\eta g_m = -\eta \nabla f_m(\mathbf{w}_t)$
 Send $\Delta\mathbf{w}_m$ to the parameter server
end while

Algorithm 4 DC-ASGD: parameter server

Input. learning rate η , control parameter λ , total iteration number N
Initialization. $\mathbf{w}_0$ is initialised randomly, $\mathbf{w}_{bak}(m) = \mathbf{w}_0, m \in \{1, 2, \dots, M\}$
while $t < N$ **do**
 if recieve “ $\Delta\mathbf{w}_m$ ” **then**
 $\mathbf{w}_{t+1} \leftarrow \mathbf{w}_t - \eta(g_m + \lambda g_m \odot g_m \odot (\mathbf{w}_t - \mathbf{w}_{bak}(m)))$
 $\leftarrow \mathbf{w}_t + \Delta\mathbf{w}_m - \frac{\lambda}{\eta} \Delta\mathbf{w}_m \odot \Delta\mathbf{w}_m \odot (\mathbf{w}_t - \mathbf{w}_{bak}(m))$
 $t \leftarrow t + 1$
 else
 if receive “pull request” from worker m **then**
 $\mathbf{w}_{bak}(m) \leftarrow \mathbf{w}_t$
 Send $\mathbf{w}_t$ back to worker m
 end if
 end if
end while

Titre : Modélisation du carnet d'ordres, Applications Market Making

Mots clés : Trading haute fréquence, Microstructure du marché, Carnet d'ordres, Processus de Hawkes, Processus de décision Markovien, Apprentissage profond.

Résumé : Cette thèse aborde différents aspects de la modélisation de la microstructure du marché et des problèmes de Market Making, avec un accent particulier du point de vue du praticien. Le carnet d'ordres, au cœur du marché financier, est un système de files d'attente complexe à haute dimension. Nous souhaitons améliorer la connaissance du LOB pour la communauté de la recherche, proposer de nouvelles idées de modélisation et développer des applications pour les Market Makers. Nous remercions en particulier l'équipe *Automated Market Making* d'avoir fourni la base de données haute-fréquence de très bonne qualité et une grille de calculs puissante, sans laquelle ces recherches n'auraient pas été possible.

Le Chapitre 1 présente la motivation de cette recherche et reprend les principaux résultats des différents travaux. Le Chapitre 2 se concentre entièrement sur le LOB et vise à proposer un nouveau modèle qui reproduit mieux certains faits stylisés. A travers cette recherche, non seulement nous

confirmons l'influence des flux d'ordres historiques sur l'arrivée de nouveaux, mais un nouveau modèle est également fourni qui réplique beaucoup mieux la dynamique du LOB, notamment la volatilité réalisée en haute et basse fréquence. Dans le Chapitre 3, l'objectif est d'étudier les stratégies de Market Making dans un contexte plus réaliste. Cette recherche contribue à deux aspects : d'une part le nouveau modèle proposé est plus réaliste mais reste simple à appliquer pour la conception de stratégies, d'autre part la stratégie pratique de Market Making est beaucoup améliorée par rapport à une stratégie naive et est prometteuse pour l'application pratique. La prédiction à haute fréquence avec la méthode d'apprentissage profond est étudiée dans le Chapitre 4. De nombreux résultats de la prédiction en 1-étape et en plusieurs étapes ont retrouvé la non-linéarité, stationarité et universalité de la relation entre les indicateurs microstructure et le changement du prix, ainsi que la limitation de cette approche en pratique.

Title : Limit order book modelling, Market Making Applications

Keywords : High-frequency trading, Market microstructure, Limit order books, Hawkes process, Markov decision process, Deep learning

Abstract : This thesis addresses different aspects around the market microstructure modelling and market making problems, with a special accent from the practitioner's viewpoint. The limit order book (LOB), at the heart of financial market, is a complex continuous high-dimensional queueing system. We wish to improve the knowledge of LOB for the research community, propose new modelling ideas and develop concrete applications to the interest of Market Makers. We would like to specifically thank the *Automated Market Making* team for providing a large high frequency database of very high quality as well as a powerful computational grid, without whom these researches would not have been possible.

The first chapter introduces the incentive of this research and resumes the main results of the different works. Chapter 2 fully focuses on the LOB and aims to propose a new model that better reproduces some stylized facts. Through this research, not only do we

confirm the influence of historical order flows to the arrival of new ones, but a new model is also provided that captures much better the LOB dynamic, notably the realized volatility in high and low frequency. In chapter 3, the objective is to study Market Making strategies in a more realistic context. This research contributes in two aspects : from one hand the newly proposed model is more realistic but still simple enough to be applied for strategy design, on the other hand the practical Market Making strategy is of large improvement compared to the naive one and is promising for practical use. High-frequency prediction with deep learning method is studied in chapter 4. Many results of the 1-step and multi-step prediction have found the non-linearity, stationarity and universality of the relationship between microstructural indicators and price change, as well as the limitation of this approach in practice.

