

Interacting Particle Systems and Mean-field Games Abstracts

12–16 October 2026

September 28, 2026

TBA

Jose Antonio Carrillo
University of Oxford, UK

Pending

Computing Stationary Equilibria in Measure-Dependent Markov Systems

Jing Dong
Columbia University, USA

Many stochastic systems in operations and economics exhibit feedback between their long-run state distribution and the transition law governing their dynamics. In this paper, we develop a computational framework for stationary equilibria in such measure-dependent Markov systems when this feedback operates through a low-dimensional aggregate. We show that the original stationary-equilibrium problem can be reduced to a finite-dimensional self-consistency equation, separating steady-state analysis of the underlying Markov system from equilibrium computation. This reduction allows the structure of the self-consistency map to guide the choice among fixed-point, relaxed fixed-point, and residual-based methods. For residual-based computation, we develop finite-time infinitesimal perturbation analysis estimators for the required derivatives, with error bounds that separate Monte Carlo error from finite-time bias. We illustrate the framework through a strategic $G/G/c$ queue and an opinion-dynamics model, showing how different structural properties lead naturally to different equilibrium-computation methods. This is joint work with Bar Light and Xin T. Tong.

Learning to Rank under Strategic Corruption: A Nonatomic Approach

Yifan Feng
National University of Singapore, Singapore

We study a dynamic learning-to-rank problem in which N agents strategically inflate the reward signals that are used to rank them, with corruption budgets chosen endogenously. The learner observes noisy manipulated outcomes and seeks to maximize total true reward over T periods. Under a general supermodular ranking-reward model, we show that efficient learning remains achievable despite

pervasive strategic manipulation. Specifically, a simple policy attains regret bounded in T , namely $O(\min\{N \log T, N^3\})$, under a strategy profile that forms an $o(1)$ -equilibrium among agents as N and T grow. Our analysis proceeds in three parts. First, we show that agents' strategic interactions converge to a static nonatomic game. Second, we find a nonatomic equilibrium in this limit, where agents' manipulation intensity increases with their types. Third, we transport the limit equilibrium to the finite-agent, finite-horizon setting, which yields non-asymptotic approximate-equilibrium and regret guarantees.

This is joint work with Qinzheng Li from NUS and Hongfan (Kevin) Chen from CUHK.

Equilibrium in Multi-Agent Reinforcement Learning

Bar Light

National University of Singapore, Singapore

Standard solution concepts for stochastic games, such as Markov perfect equilibrium and Markov coarse correlated equilibrium, are computationally difficult, and thus, standard decentralized reinforcement-learning algorithms should not generally be expected to converge to them. In this paper, we study the equilibrium generated by such algorithms. In particular, we introduce a new solution concept for stochastic games, Markov Bayes coarse correlated equilibrium (MBCCE), defined as a distribution over states and stationary policy profiles such that, after observing the state but before observing her recommended action, no player can gain by choosing a different current action, with the sampled policy profile governing play thereafter. We discuss the parallels between MBCCE and coarse correlated equilibrium (CCE) in finite normal-form games and show that MBCCE retains several of its key properties. We then introduce a corresponding regret notion, adaptive Markov coarse regret (AMCR), and show that vanishing AMCR implies that every accumulation point of the empirical distribution of realized states and policy profiles is an MBCCE. Crucially, we show that achieving AMCR reduces to two standard learning tasks: minimizing external regret at each state and accurately evaluating the current joint policy. We then prove that under mild conditions these properties hold for two natural RL algorithmic designs: a decentralized asynchronous actor-critic algorithm through a new two-timescale stochastic-approximation analysis, and a standard episodic multi-agent projected policy-gradient method. Hence, both algorithms generate approximate MBCCEs, and we establish explicit finite-time convergence rates for both. Finally, we relate MBCCE to stationary MCCE and use this connection together with our algorithmic guarantees to identify a natural class of stochastic games in which these algorithms achieve approximate MCCE. Taken together, our results identify MBCCE as a new equilibrium notion for the long-run empirical behavior of standard

reinforcement-learning methods in general finite discounted stochastic games.

Learning Elliptic PDEs with Mean Field Langevin Dynamics

Yulong Lu

University of Minnesota, USA

In this talk, we study the theoretical foundations of feature learning within the frameworks of the deep Ritz method for solving a class of elliptic PDEs with the Neumann boundary condition. We analyze the optimization dynamics of the two-layer neural networks trained via the mean-field Langevin algorithm. By utilizing a quantitative propagation of chaos estimate for the mean-field Langevin algorithm, we provide quantitative estimates for the sample and running-time complexities of the algorithm needed to achieve given accuracy.

Designing Interacting Particle Systems for Sampling

Aimee Mauraïs

Massachusetts Institute of Technology, USA

Sampling from a target probability distribution is fundamental to modern computational science and machine learning. Sampling is the essence of Monte Carlo integration, enables uncertainty quantification in Bayesian inference, and underlies generative models that have the ability to synthesize convincing text, images, and far beyond. A powerful, emerging approach to sampling is to employ dynamics, in the form of an ordinary or stochastic differential equation, to evolve a collection of samples from a tractable reference distribution toward the desired target distribution. Such dynamic approaches are state-of-the-art in generative modeling, where they typically rely on velocity fields and drift functions which are pre-trained to achieve the desired evolution. While similar training-intensive approaches can be employed for sampling from a target density, an alternate approach is to rely on interaction: by devising dynamics which explicitly depend on the current distribution of the particle ensemble, we can evolve the ensemble towards a desired target without the need for training. In this talk I will discuss my work on enabling such interacting particle system approaches for sampling via (1) development of new, gradient-free particle systems for Bayesian sampling, (2) principled design of density paths via PDE-constrained optimization, and (3) scalability through exploitation of sparse conditional dependence structure.

This talk is based on joint work with Bamdad Hosseini (University of Washington), Shuigen Liu (Penn State), Youssef Marzouk (MIT), and Xin Tong (NUS).

Propagation of Chaos for Mean-Field Langevin Dynamics

Atsushi Nitanda

Nanyang Technological University, Singapore

Mean-field Langevin dynamics (MFLD) is an optimization method derived by taking the mean-field limit of noisy gradient descent for two-layer neural networks in the mean-field regime. Recently, the propagation of chaos (PoC) for MFLD has gained attention as it provides a quantitative characterization of the optimization complexity in terms of the number of particles and iterations. Remarkable progress in a recent study showed that the approximation error due to finite particles remains uniform in time and diminishes as the number of particles increases. By refining the defective log-Sobolev inequality a key result from earlier work an improved PoC result is established for MFLD, which removes the exponential dependence on the regularization coefficient from the particle approximation term of the optimization complexity.

Interacting Particle Formulations for Stochastic Optimal Control

Sebastian Reich

University of Potsdam, Germany

Pontryagin's maximum principle has been previously extended to stochastic optimal control leading to a coupled set of forward-backward SDEs. In this talk, an alternative mean-field approach will be presented which, when discretised, results in (deterministic) Hamiltonian interacting particle systems. The proposed approach is attractive since it naturally extends to mean-field control problems. We will also present initial results on applications to the control of partially observed diffusion processes, which can be rephrased as a mean-field control problem subject to a common noise driver.

TBA

Domènec Ruiz-Balet
Universitat de Barcelona, Spain

Pending

TBA

Bao Wang
University of Utah, USA

Pending

Learning-enhanced Structure Preserving Methods for Kinetic Plasma Models

Li Wang
University of Minnesota, USA

Plasma fusion holds great promise as a future source of clean energy. Despite decades of progress in developing reliable computational methods for plasma models, kinetic models—which provide a first-principles description of interacting charged particles—remain prohibitively expensive to solve. Their computational cost is far from meeting the millisecond-scale requirements of real-time burning-plasma control. In this talk, I will describe our recent efforts to integrate deep learning with conventional structure-preserving numerical methods to address some of the key challenges in simulating these kinetic models.

Incomplete Market Equilibrium under Social Comparisons

Marko Hans Weber
National University of Singapore, Singapore

We study how social comparisons shape asset prices in an incomplete-market economy with uninsurable income risk. Agents are heterogeneous in preferences, endowments, and in the weights they place on the consumption of others. Using a mean-field game framework with a directed network described by a finite-rank kernel,

we derive in closed form the equilibrium interest rate, stock price, and individual consumption and portfolio policies. A social multiplier matrix captures how network interactions reallocate risk across agents. Social comparisons reduce the market price of hedgeable risk by raising the effective risk tolerance of the economy, as consumption and social reference benchmarks co-move with common shocks. The network structure governing social comparisons determines whether the precautionary-savings motive induced by unhedgeable risk is amplified or mitigated. If all agents exert the same influence on others and are equally sensitive to the rest of the system, as in a doubly stochastic comparison network, the specific network topology does not affect equilibrium prices.

Continuous-time Mean Field Games: a Primal-dual Characterization

Jiacheng Zhang

The Chinese University of Hong Kong, Hong Kong SAR

This paper establishes a primal-dual formulation for continuous-time mean field games (MFGs) and provides a complete analytical characterization of the set of all Nash equilibria (NEs). We first show that for any given mean field flow, the representative player's control problem with *measurable coefficients* is equivalent to a linear program over the space of occupation measures. We then establish the dual formulation of this linear program as a maximization problem over smooth subsolutions of the associated Hamilton-Jacobi-Bellman (HJB) equation, which plays a fundamental role in characterizing NEs of MFGs. Finally, a complete characterization of *all NEs for MFGs* is established by the strong duality between the linear program and its dual problem. This strong duality is obtained by studying the solvability of the dual problem, and in particular through analyzing the regularity of the associated HJB equation.

Compared with existing approaches for MFGs, the primal-dual formulation and its NE characterization do not require the convexity of the associated Hamiltonian or the uniqueness of its optimizer, and remain applicable when the HJB equation lacks classical or even continuous solutions.

Learning Mean-Field Games through Mean-Field Actor-Critic Flow

Mo Zhou

University of California, Los Angeles, USA

Consensus-Based Optimization (CBO) is a class of derivative-free methods for global optimization based on interacting stochastic particle systems. Particles are attracted toward an objective-weighted consensus point that favors regions of low objective values. The resulting nonlinear and nonlocal interactions pose mathematical challenges for understanding both the mean-field approximation and the long-time behavior of these systems.

In this talk, we establish uniform-in-time propagation of chaos for first-order CBO and its kinetic (second-order) extension. Using stability estimates for the weighted consensus point and quantitative concentration inequalities, we show that the particle system tracks its mean-field limit at the classical Monte Carlo rate uniformly over infinite time horizons. For CBO with anisotropic noise, the error prefactor is independent of the problem dimension. We also discuss almost sure exponential convergence of the particle system and the additional challenges in establishing long-time stability when velocity dynamics are introduced.

We then turn to the long-time asymptotics of a dynamically regularized first-order CBO model. By combining Wasserstein convergence with Fisher-information estimates, we establish convergence of the rescaled mean-field density in L^1 , characterizing its asymptotic profile as the original distribution concentrates around a consensus point.

The propagation-of-chaos results are joint work with Nicolai Gerber, Franca Hoffmann, and Urbain Vaes (first order), and with Seung-Yeal Ha and Franca Hoffmann (second order). The long-time asymptotics results are joint work with Wuchen Li and Franca Hoffmann.

A New Framework for Mean-field Reinforcement Learning in the Physical World

Yuhua Zhu

University of California, Los Angeles, USA

This talk addresses mean-field reinforcement learning (RL) in the physical world, where the system dynamics evolve smoothly according to McKean–Vlasov equations. The central challenge is that the system evolves in continuous time, yet only

discrete-time data are available. We first show that classical RL algorithms incur substantial discretization error because they ignore the intrinsic smoothness of the underlying dynamics. We then introduce Optimal-PhiBE, a structure-aware framework that embeds discrete-time information into a continuous-time PDE. By exploiting the smooth structure of the dynamics, Optimal-PhiBE yields a provably more stable approximation to the optimal policy. In linear quadratic control, Optimal-PhiBE even recovers an accurate optimal policy in the population-level from discrete-time observations alone.

We further develop a model-free algorithm for solving Optimal-PhiBE and establish its convergence under model misspecification, along with finite-sample guarantees. Remarkably, a minimal change, a single line of code in standard RL algorithms, suffices to adapt them to physical-world RL problems and substantially reduces discretization error. Finally, we characterize a fundamental trade-off between discretization error and sample error that is inherent to RL in the physical world.
