

EEML Summer School - Day 1

01/07/19

Intro to Deep Learning - Razvan Pascanu (DeepMind)

- SGD intuition from Taylor expansion (1st order):

$$\arg\min_{\theta} L(\theta + \Delta\theta) \approx \arg\min_{\Delta\theta} \left[L(\theta) + \Delta\theta \frac{\partial L}{\partial \theta} \right] \text{ s.t. } \|\Delta\theta\| \leq \epsilon$$

$$\Leftrightarrow \arg\min_{\Delta\theta} L(\theta) + \Delta\theta \frac{\partial L}{\partial \theta} + \frac{1}{2} (\Delta\theta)^T I(\Delta\theta) \Delta\theta$$

Laplace H.
matrix
→ form of loss! rel!

- Black-Box ^{gen.} via Evolutionary algos $\Rightarrow$ Evolution $\oplus$ Next Generation
 → allows for potentially global optimizers $\Rightarrow$ problem of many proofs
 → More useful for hypothesis
- Population-based Training $\Rightarrow$ Train models in parallel / stop + evaluate
 afterwards only further develop best performing model
 → Between parametric and evolutionary optimization!

→ slow learning
→ compression / fast eval

- $h: X \rightarrow Y \Rightarrow$ Predictor $\Rightarrow$ Hypothesis
 (Classification / Regression / Structured Pred. (graphs))

$$\hookrightarrow h: \theta \times X \rightarrow Y \Rightarrow \text{fit } \theta \text{ via loss fct.}$$

→ ML / MXP perspective!

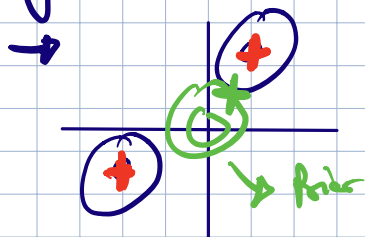
→ Parametric
 → Non-Parametric
 → bad scaling / slow time
 → Fast Learning

$$\arg\max_{\theta} P(D|\theta) = P(D|\theta) P(\theta) / P(D) \rightarrow \arg\max_{\theta} P(D|\theta) P(\theta)$$

- Uniform prior in MXP formulation results in Max. Likelihood!
- together with iid assumption $\rightarrow$ split data into $\frac{N}{n} \oplus$ by trials

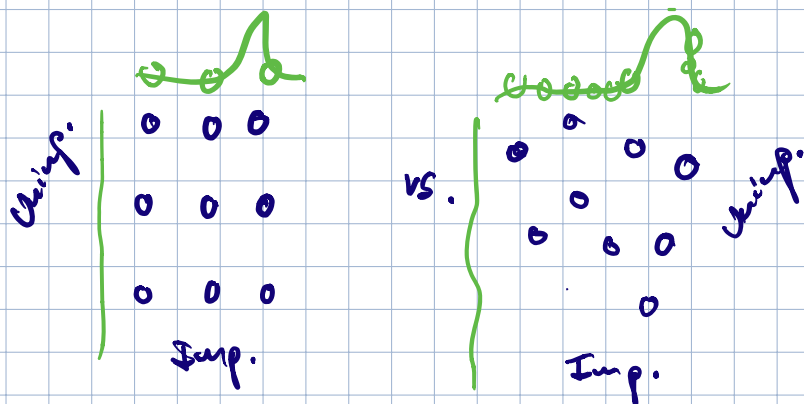
- Loss fct often times results from Bayesian / MXP perspective $\oplus$ assumptions on the data $\rightarrow$ E.g. MSE from MXP + Uniform Prior + Gaussian Data!
- Multi-Label Classification: Multinomial $\Leftrightarrow$ w. log Ltt.

- Regularize via not assuming $P(\theta)$ being uniform! [Together w. Outward approx.]



$\Rightarrow$ Reg. get rid of local min problem + by preferring one that is closer to 0 $\Rightarrow +$

- Hyperparameter Optimization: Random vs. Grid $\rightarrow$ Random covers more space!

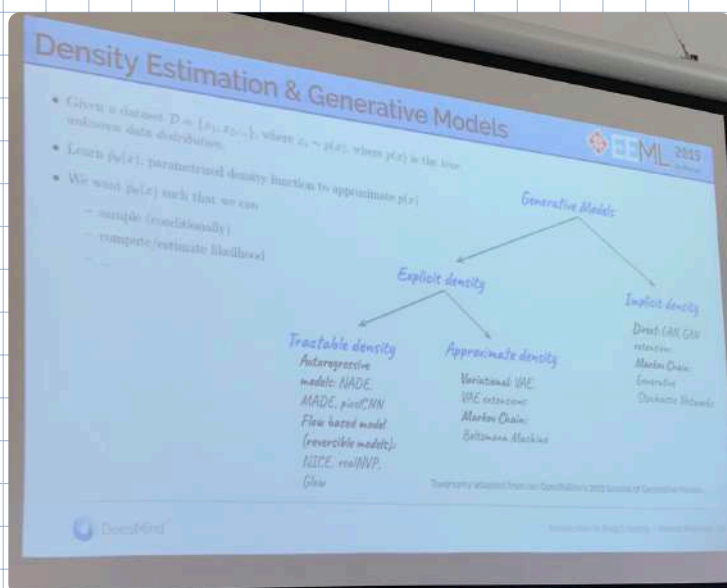


$\rightarrow$ Be careful not to overfit test data!
 $\rightarrow$ Unclear what global best configuration is!

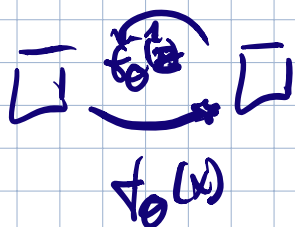
- Early-stopping as a prior $\rightarrow$ Form of L_2 $\Rightarrow$ See Bishop! Trust Region?
 $\rightarrow$ Duvenaud et al 2016 - Early Stopping as Hypothesis.

Generative Models

$\rightarrow$ Interpretable: Decompose joint as product of conditionals
 $\rightarrow$ Not necessarily clear how to do so if there is no temporal chain.
 Eg. where to start sampling image
 $\rightarrow$ Also: Slow Sampling!



$\rightarrow$ Flow-based / Reversible Models: Find $f_\theta^{-1}(z)$



$\rightarrow$ Change of variables $\rightarrow$ Gaussian in latent space
 $\rightarrow$ optimize jointly and invert

$$p_X(x) = p_Z(z) \left| \det \left(\frac{\partial x}{\partial z} \right) \right|^{-1}$$

$\rightarrow$ Not clear how to parametrize f_θ s.t. it is invertible!

$\rightarrow$ Easily sample based on sampling Gaussian z and then apply $f_\theta^{-1}(z)$

$\rightarrow$ VAE $\Rightarrow$ Restrict latent codes by having encoder output Gaussian parametrization $\rightarrow$ Optimize both Reconstruction as well as KL objective! $\rightarrow$ Restriction allows for efficient sampling

Variational Autoencoders (VAE)

$\ln p(x) = \ln \int p(x, z) dz$

$= \ln \int p(x, z) \frac{q(z|x)}{q(z|x)} dz$

$\geq E_{q(z|x)} \left[\ln \frac{p(x, z)}{q(z|x)} \right]$ (Jensen inequality)

$= E_{q(z|x)} \left[\ln \frac{p(x|z)p(z)}{q(z|x)} \right]$

$= E_{q(z|x)} [\ln p(x|z)] + E_{q(z|x)} \left[\ln \frac{p(z)}{q(z|x)} \right]$

$= E_{q(z|x)} [\ln p(x|z)] + \int q(z|x) \ln \frac{p(z)}{q(z|x)} dz$

$= E_{q(z|x)} [\ln p(x|z)] + D_{KL} [q(z|x) \| p(z)]$

$= \text{likelihood} + \text{KL to prior}$

$\log(z|x)$

Autoencoder

Original → Encoder → Compressed → Decoder → Reconstruction

Variational Autoencoder

Reparam. trick for differentiability

Encoder: $q(z|x)$

Decoder: $p(x|z)$

Loss: $D_{KL} [q(z|x) \| p(z)]$

DeepMind

↳ Gaussian cores are lucky for the computation as well as sample

↳ Gaussian + ReLU //

also new work on Gumbel ReLU all multivariate

↳ Problem to get sharp color maps!

↳ want VAE for latent interpolation!

→ Interpressive models vs. GMMs → Under GMM better, under BR

• ReLU activations: Splits space into linear regions → Can be a lot

↳ Montufar et al 2014 → NIPS

↳ Some of optimal efficient policy of the space!

↳ Optimization such that decision boundary becomes sharp!

↳ No anything some regions are restricted to close parameters!

→ Need high-dimensional data → Unfortunately both

↳ 1d mixture of Gaussians hard to fit!

• Loss Opt. Surfaces: Not necessarily crazy!

Not a lot known about restrictive parameters!

↳ Yann LeCun: Simple random heuristic of complex prob. of all dir. putting up → very interesting

↳ Sharp vs. flat minima! → What do they mean? LeCun

↳ Can switch eigenvalue profile without changing features!

• RMSprop → Running avg, Mean/Variance → larger step if small var

↳ Approx. of Natural Gradient

↳ Better when → Hessian approximates to identity → Fisher & NG!

- Convolution = restriction of net $\rightarrow$ Used were inductive biases!
- Graph Nets = Do Can not on locally defined pixels but on other local structure
- RNN $\rightarrow$ Exploding grad from product of Jacobians $\rightarrow$ Gradient clipping
 - $\hookrightarrow$ LSTM $\rightarrow$ varying gradient constant via linear cell state but need to learn gates
 - $\hookrightarrow$ Interp gate work as a form of low pass filters ^{reverses 0 at 1!}

Intro to RL - John Breck (DeepMind / McGill) $\rightarrow$ TD

- How to overcome sample efficiency problem with sparse rewards
 - $\rightarrow$ self-play $\rightarrow$ easy acquisition // intrinsic motivation $\rightarrow$ surrogate reward
- Key features:
 - * Trial-and-Error Search
 - * Stochastic Env
 - * Delayed reward $\rightarrow$ Temp. Credit Assignment
 - * $E-E$ Trade-off
- heavily with multiple output heads / auxiliary heads
 - $\hookrightarrow$ Relationship to Ego / Bounded Rationality ?!
- TD Ganner $\rightarrow$ (Tesauro et al, 1995) $\Rightarrow$ "Early Alpha-Go"
- Intrinsic learning / behavioral cloning $\rightarrow$ form of pretraining
 - $\rightarrow$ important to train in simulator with exploration to robustify!
- Different of interpretations: Survival prob., bias-variance tradeoff, self-heals
- ML methods $\approx$ Supervised $\rightarrow V(S_t)_{t+1} \leftarrow V(S_t)_{t+1} + \alpha [R_t + V(S_t)_{t+1} - V(S_t)_t]$
 - $\rightarrow$ Update for each transition $\Rightarrow$ not iid! Sequentially correlated
 - $\hookrightarrow$ Overcome via ER buffer for example!
- Bias-Variance Trade-Offs:
 - * Function approx. $\rightarrow$ Gradient updates
 - * Generalization across environments
 - * Discount $\rightarrow$ Estimation of future rewards
 - $\hookrightarrow$ schedule of discount params $\rightarrow$ relate to TD as in PER?!

- Bellman Self-Consistency Equations: Set of linear equations!

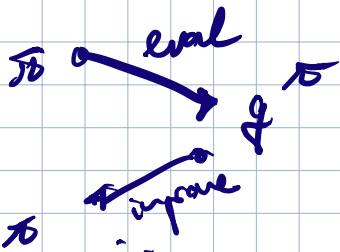
$$v_{\pi}(s) = \sum_a \pi(a|s) \sum_{s', r} p(s', r|s, a) [r + \gamma v_{\pi}(s')]$$

- ↳ But $v^{\pi}(s)$ is a non-linear set of equations due to $\arg\max$!
- ↳ Can only be done if full MDP is known! $\Rightarrow$ Estimate dynamics
- ↳ Importance $\gamma < 1$: Contraction $\Rightarrow$ Banach Fixed Point Theorem
- ↳ Problem also with large state spaces $\Rightarrow$ storage is expensive

- Contraction is what makes bootstrapping work!

↳ Would it make sense to start off with MC return backup and slowly move towards full bootstrap? k -step where $k=1 \rightarrow k=1$

- Policy Improvement: Assumption of no local optima otherwise exploitable heuristic



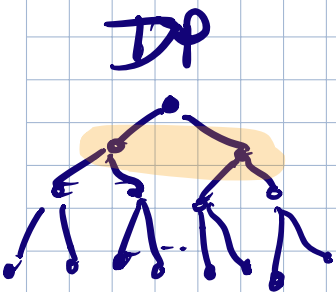
• Dyn. Pr.: Full-Width Backup

↳ TD: In between MC and DP

- Crucial assumption of Markovity for bootstrapping $\Rightarrow$ BXS

↳ Introduce bias to reduce variance

↳ MC not biased! $\Rightarrow$ k -step TD introduces!



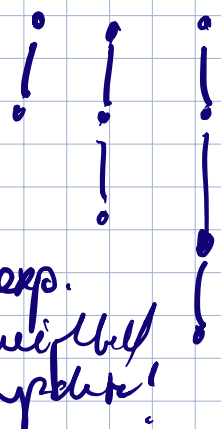
TD(1)



k -step TD



TD(2)



- Epilepsy project: Use SDS to simulate

↳ likelihood-free methods!!

exp. neighbor updates!

EEML Summer School - Day 1

Intro to RL II - Policy Search → Control

Stochastic Gradient Descent (SGD) is the idea behind most approximate learning

General SGD: $\theta \leftarrow \theta - \alpha \nabla_{\theta} \text{Error}_t^2$

For VFA: $\leftarrow \theta - \alpha \nabla_{\theta} [\text{Target}_t - \hat{v}(S_t, \theta)]^2$

Chain rule: $\leftarrow \theta - 2\alpha [\text{Target}_t - \hat{v}(S_t, \theta)] \nabla_{\theta} [\text{Target}_t - \hat{v}(S_t, \theta)]$

Semi-gradient: $\leftarrow \theta + \alpha [\text{Target}_t - \hat{v}(S_t, \theta)] \nabla_{\theta} \hat{v}(S_t, \theta)$

Linear case: $\leftarrow \theta + \alpha [\text{Target}_t - \hat{v}(S_t, \theta)] \phi(S_t)$

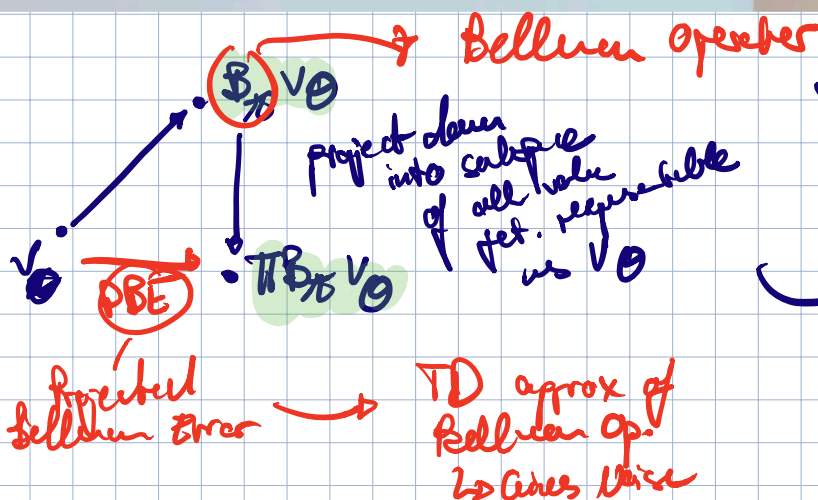
Action-value form: $\theta \leftarrow \theta + \alpha [\text{Target}_t - \hat{q}(S_t, A_t, \theta)] \phi(S_t, A_t)$

→ Transform into form of supervised learning

→ Assume that target does not depend on θ

→ Stability problem
→ Target networks

→ Would not be a problem if full MC



$$(B_{\pi} V) S := \sum_a \pi(s, a) [r(s, a) + \sum_{s'} p(s' | s, a) V(s')]$$

done until $PBE = 0$

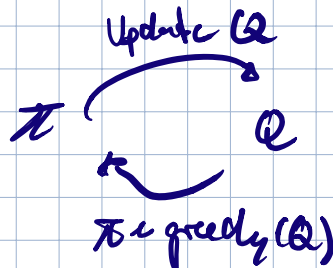


→ Policy shifts due policy shifts! → Need to constrain that shift
E.g. introducing trust region! → clipping / MC

→ Linear fct. approx: $MSVE(\theta_{lin}) \leq \frac{1}{1-\gamma} MSVE(\theta_{MC})$ → many give in con-linear!

→ u-step backups: Trade-Off bias-variance → Full MC has large variance!
→ of controls for this

• Monte Carlo Control:



→ Random exploring shifts!
low vs off-policy

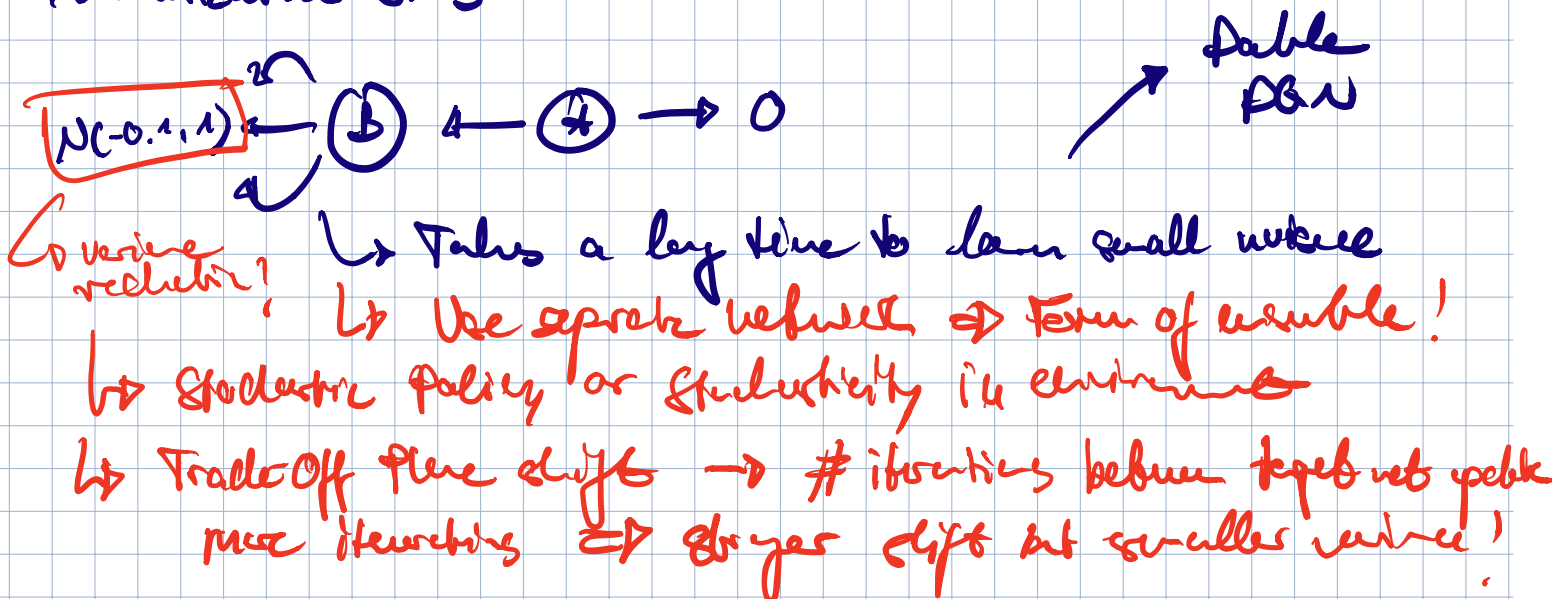
→ Exploration: ϵ -greedy → dithering → problem indep. of actual values
→ In practice: exploration schedules are hard to implement

- Targets: $y_t = r_{t+1} + \gamma Q(s_t, a_{t+1}; \theta_t)$ [SARSA]
 $y_t = r_{t+1} + \gamma \sum_a \pi(a|s_{t+1}) Q(s_{t+1}, a; \theta_t)$ [Exp. SARSA]
 $y_t = \sum_{s', r} p(s', r | s_t, a) [r + \gamma \sum_a \pi(a|s_{t+1}) Q(s_{t+1}, a; \theta_t)]$
 ↳ SARSA: on-policy → next action is chosen at t to update
 ↳ still explores!

- Q-learning: Exploring → since theoretically converge from completely random start!

↳ SARSA vs. Q-L: Safety vs. Optimality

- Maximization bias:



- Rainbow: Off-policy Parameters $\Rightarrow$ but maybe there is safety better! $\rightarrow$ META RL!
- Used for different set of optimizers in RL $\rightarrow$ different targets then $\rightarrow$ special learning
- Policy gradient Methods:

Objective: $J(\theta) = y_\theta(s_0)$ $\rightarrow$ initial state distr. $\rightarrow$ policy

$\theta_{t+1} := \theta_t + \alpha \nabla J(\theta_t)$

$\nabla J(\theta) = \sum_s d_\pi(s) \sum_a \underbrace{q_\pi(s, a)}_{\text{baseline estimator}} \nabla_\theta \pi(a|s, \theta)$

Deriving REINFORCE from the PGT

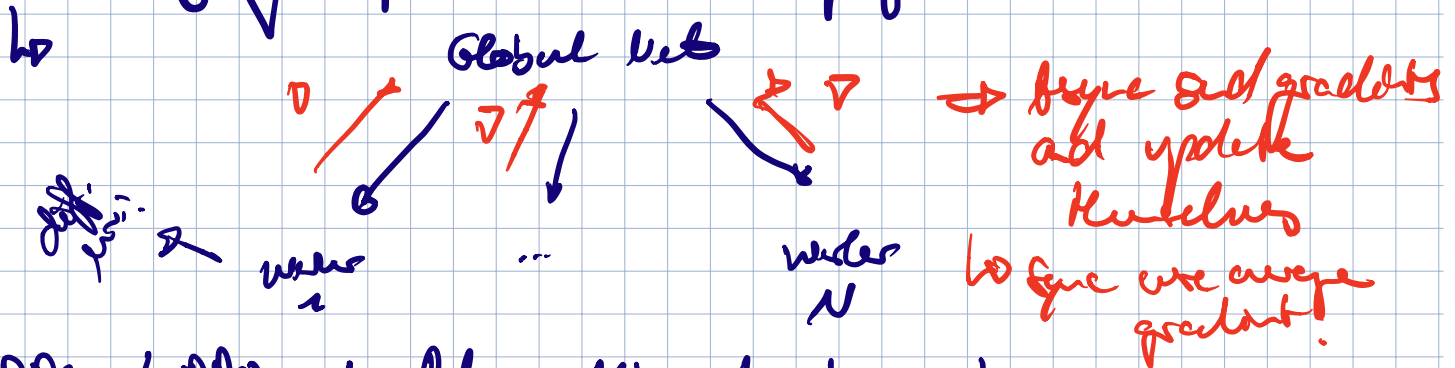
$$\begin{aligned}
 \nabla_{\theta} J(\theta) &= \sum_s d_{\pi}(s) \sum_a q_{\pi}(s, a) \nabla_{\theta} \pi(a|s, \theta) \\
 &= \mathbb{E}_{\pi} \left[\gamma^t \sum_a q_{\pi}(S_t, a) \nabla_{\theta} \pi(a|S_t, \theta) \right] \\
 &= \mathbb{E}_{\pi} \left[\gamma^t \sum_a \pi(a|S_t, \theta) q_{\pi}(S_t, a) \frac{\nabla_{\theta} \pi(a|S_t, \theta)}{\pi(a|S_t, \theta)} \right] \\
 &= \mathbb{E}_{\pi} \left[\gamma^t q_{\pi}(S_t, A_t) \frac{\nabla_{\theta} \pi(A_t|S_t, \theta)}{\pi(A_t|S_t, \theta)} \right] \quad (\text{replacing } a \text{ by the sample } A_t \sim \pi) \\
 &= \mathbb{E}_{\pi} \left[\gamma^t G_t \frac{\nabla_{\theta} \pi(A_t|S_t, \theta)}{\pi(A_t|S_t, \theta)} \right] \quad (\text{because } \mathbb{E}_{\pi}[G_t|S_t, A_t] = q_{\pi}(S_t, A_t))
 \end{aligned}$$

Thus

$$\theta_{t+1} \triangleq \theta_t + \alpha \widehat{\nabla_{\theta} J(\theta_t)} \triangleq \theta_t + \alpha \gamma^t G_t \frac{\nabla_{\theta} \pi(A_t|S_t, \theta)}{\pi(A_t|S_t, \theta)}$$

- No parameterization of value fct. needed
- Use returns
- MC estimation
- Agh noise!
- Need to stabilize!
- AC $\rightarrow$ parameterize value fct!

- *3C $\rightarrow$ reduce exploration by parallelization!
- Pooling of experiences $\rightarrow$ Mean of gradients



- TRPO / PPO $\rightarrow$ before KL constraint as large!

Multi-agent RL - Shimon Whiteson (Oxford)

- Motivation: Classical RL abstracts many aspects
↳ World is full of multi-agent systems

COOPERATIVE

- ↳ Shared reward
- ↳ Coordination

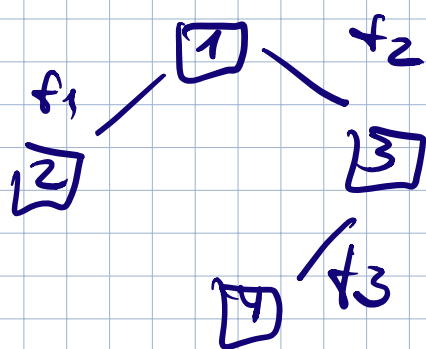
COMPETITIVE

- ↳ Zero-Sum Game
- ↳ Opposing rewards
- ↳ Minimax

MIXED

- ↳ General Sum
- ↳ Nash Eq.
- ↳ Q? Etc.

- Cooperative setting: Absence of social conflict $\Rightarrow$ risk of coordination
→ Exponential joint-action space. $\Rightarrow$ grows very fast
→ Question: what does $\Rightarrow$ Coordination Graphs
 $Q(u) = f_1(u^1, u^2) + f_2(u^1, u^3) + \dots$ → action agent 3



- ↳ Probabilistic graphical Model
 $\Rightarrow$ Solvers: MPP estimation!
↳ Encoding of cond. independence
 $\Rightarrow$ No separability!

↳ Variable elimination:

$$\max_u Q(u) = \max_{u_2, u_3, u_4} [f_3(u^3, u^4) + \underbrace{\max_{u_1} [f_1(u^1, u^2) + f_2(u^1, u^3)]}_{f_2(u^2, u^3)}]$$

↳ form of cooperative Model!

↳ Computationally expensive if we were connected the graph

- ↳ Computes only optimal value of joint action $\Rightarrow$ no balanced paths to actually derive that action!
- ↳ Need to know the graph / Cond. indep. conditions

- Max-Plus $\Rightarrow$ Vlassis et al (2004) $\Rightarrow$ Message Passing
↳ Two agent local payoff assumption

↳ Improving messages themselves

↳ Convergence guarantees in acyclic graphs

→ MDP (+ Synchronicity!) : Centralized Multi-Agent MDP

↳ All agents see state $\Rightarrow$ Not really multi-agent: simply huge action/state space $\rightarrow$ simply comb. optimization

→ Independent Q-learning : No attempt to model $Q(s, u)$

↳ each agent learns $Q(s_i, u_i^a)$! $\Rightarrow$ consistency if all agents learn $\rightarrow$ no convergence guarantee! $\rightarrow$ Tan et al (1988)

→ Coordinated Q-learning : Gordon et al (2002)

$$Q^{\text{tab}}(s, u) = \sum_{e \in \mathcal{E}} Q_e(s^e, u^e)$$

↳ extremely strong assumption! $\rightarrow$ subjects our agents / ad features

↳ Transitions might be coupled

↳ Update each factor $\Rightarrow$ graph is stationary!

• Partial observability $\rightarrow$ requires decentralized execution!

↳ learning should be centralized! $\rightarrow$ share params / gradients

↳ Minimal vs. right assumptions!

• Dec-POMDP

$Q(s, a): S \times \mathcal{A} \rightarrow \mathbb{R}$ $\rightarrow$ before the hist: $\tau^a \in T \subseteq (\mathbb{R} \times \mathcal{U})^T$

↳ Decentralized policies: $\delta^a(u^a) \gamma^a): T \times \mathcal{U} \rightarrow [0, 1]$

↳ dilemma: Explore private info vs. being predictable $\rightarrow$ ^{depends on} reward jet

• Policy Gradient Methods for MxRL:

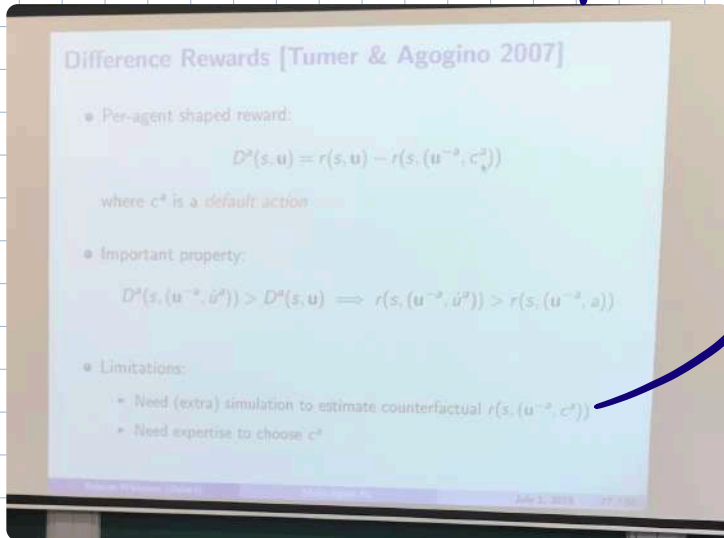
$\rightarrow$ PG useful if gradient is hard!

$\rightarrow$ MxRL independent: share parameters - still indep. inputs + joint value
↳ consistency, hard to have coordination, multi-agents can't coord.

• Counterfactual MDPs → Foster et al 2018 (COMA)

① Centralize Critic

② Wolpert & Tumer (2000) ⇒ Wondeful Life Utility
to reason about counterfactual reward if you would not have participated



→ had to go back into simulator!
to Counterfactual Baseline

$$A^a(s, u) = Q(s, u) - \sum_{u^a} \pi(u^a | \tau^a) Q(s, (u^{-a}, u^a))$$

→ problem: as π becomes deterministic we can't get the counterfactual aspect

→ COMA critic: output and gives all values

→ still nonstochastic joint value fun.

• Value Decomposition Networks - Sunehag et al 2017

→ per agent: $Q_{tot}(\tau, u) = \sum_{a=1}^N Q_a(\tau^a, u^a; \theta^a)$

to decentralizes the agents ⇒ no lower level gradification

• QMIX: Not only summation but mixing which

⇒ constraint to non-neg weights!

→ Can blow away after criticised learning!

SMAC framework → StarCraft Baseline!

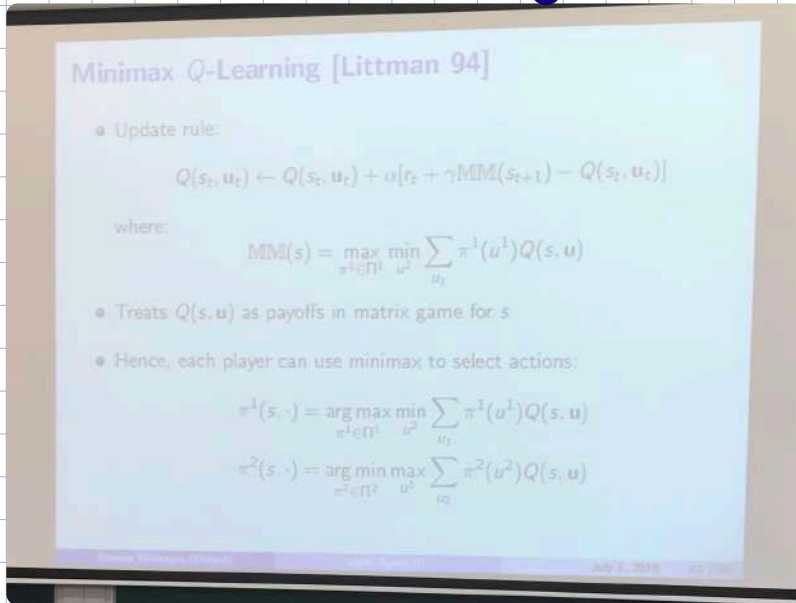
• Competitive setting: In competitive setting one optimal policy exists → how: Stochasticity

→ Mixed strategies ↔ Minimax Theorem

to Rational vs. Strategic agents → Symbolic ordering does not matter

→ inner player does not have to believe stochastically!

→ Minimax Q-Learning → Littman (1994)



→ Training approach: Learn about the opponent from data
↳ Fud. Q-Learning → implicit
↳ Explicit: Fictitious play
Rosenbluth
→ Opponent Modeling

→ Self-Play: Checkers → Arthur Samuel (1956)

↳ natural curriculum → as fast if game is not trivial
↳ leave yourself vulnerable to bad opponents while beating stronger ones → helps star to learn / kill of them!

• Fixed: Nash Equilibria → Existence!

↳ Nash Equilibrium Q-Learning → Hu & Wellman (1998)
→ almost never guaranteed to converge!

↳ "If multi-qit learning is the answer, what is the question?"

↳ Mechanism Design! → Economics

↳ Fictitious-Q-Learning