

TRPO and ACKTR

Olivier Sigaud

Sorbonne Université
<http://people.isir.upmc.fr/sigaud>



TRPO

Outline

- ▶ More PG with baselines: TRPO and ACKTR
- ▶ Three aspects distinguish TRPO:
 - ▶ Surrogate return objective
 - ▶ Natural policy gradient
 - ▶ Conjugate gradient approach
- ▶ Differences in ACKTR:
 - ▶ Approximate second order gradient descent (Hessian)
 - ▶ Using Kronecker Factored Approximated Curvature
- ▶ Then PPO (a quick overview of two versions)

Surrogate return objective

- ▶ The standard policy gradient algorithm for stochastic policies is:

$$\nabla_{\theta} J(\theta) = \mathbb{E}_t[\nabla_{\theta} \log \pi_{\theta}(\mathbf{a}_t | \mathbf{s}_t) \hat{A}_{\phi}^{\pi_{\theta}}]$$

- ▶ This gradient is obtained from differentiating

$$Loss^{PG}(\theta) = \mathbb{E}_t[\log \pi_{\theta}(\mathbf{a}_t | \mathbf{s}_t) \hat{A}_{\phi}^{\pi_{\theta}}]$$

- ▶ But we obtain the same gradient from differentiating

$$Loss^{IS}(\theta) = \mathbb{E}_t\left[\frac{\pi_{\theta}(\mathbf{a}_t | \mathbf{s}_t)}{\pi_{\theta_{old}}(\mathbf{a}_t | \mathbf{s}_t)} \hat{A}_{\phi}^{\pi_{\theta}}\right]$$

where $\pi_{\theta_{old}}$ is the policy at the previous iteration

- ▶ Because $\nabla_{\theta} \log f(\theta) |_{\theta_{old}} = \frac{\nabla_{\theta} f(\theta) |_{\theta_{old}}}{f(\theta_{old})} = \nabla_{\theta} \left(\frac{f(\theta)}{f(\theta_{old})} \right) |_{\theta_{old}}$

- ▶ Another view based on importance sampling
- ▶ See John Schulmann's Deep RL bootcamp lecture #5

<https://www.youtube.com/watch?v=xvRrgxcpaHY> (8')



Schulman, J., Levine, S., Moritz, P., Jordan, M. I., & Abbeel, P. (2015) Trust Region Policy Optimization. *CoRR*, abs/1502.05477

The policy gradient is on-policy

- ▶ The policy gradient calculation assumes that the training trajectories are obtained from the policy we are optimizing:
- ▶ Reminder: we want to find $\operatorname{argmax}_{\theta} \sum_{\tau} P(\tau, \theta) \psi(\tau)$
- ▶ We use

$$P(\tau^{(i)}, \theta_{\text{sample}}) = \prod_{t=1}^H p(s_{t+1}^{(i)} | s_t^{(i)}, a_t^{(i)}) \cdot \pi_{\theta_{\text{sample}}}(a_t^{(i)} | s_t^{(i)})$$

- ▶ Here, by definition, $\pi_{\theta_{\text{sample}}}(a_t^{(i)} | s_t^{(i)})$ is the policy which generated the trajectories
- ▶ Then we take the gradient and get the policy gradient formula
- ▶ If we want to optimize another policy $\pi_{\theta_{\text{other}}}(a_t^{(i)} | s_t^{(i)})$, the derivation is wrong

Importance sampling

- ▶ How can we estimate an expectation of a function over a distribution θ_1 if we know it from another distribution θ_2 ?

$$\begin{aligned}\mathbb{E}_{x \sim \theta_1}[f(x)] &= P(x|\theta_1)f(x) \\ &= \frac{P(x|\theta_1)}{P(x|\theta_2)} P(x|\theta_2)f(x) \\ &= \frac{P(x|\theta_1)}{P(x|\theta_2)} \mathbb{E}_{x \sim \theta_2}[f(x)]\end{aligned}$$

- ▶ $\frac{P(x|\theta_1)}{P(x|\theta_2)}$ is the importance sampling term
- ▶ In policy gradient methods, the distributions of interest are $\pi_{\theta_{smp}}$ and $\pi_{\theta_{other}}$.

Importance sampling: application to TRPO

- ▶ We sampled data from $\pi_{\theta_{samp}}$
- ▶ We want to optimize another policy $\pi_{\theta_{other}}$,
- ▶ We can rewrite

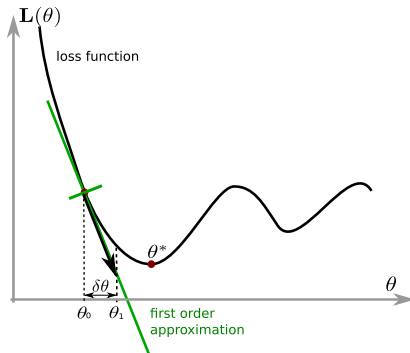
$$P(\tau^{(i)}, \theta_{other}) = \prod_{t=1}^H p(s_{t+1}^{(i)} | s_t^{(i)}, a_t^{(i)}) \cdot \pi_{\theta_{other}}(a_t^{(i)} | s_t^{(i)}) \cdot \frac{\pi_{\theta_{samp}}(a_t^{(i)} | s_t^{(i)})}{\pi_{\theta_{samp}}(a_t^{(i)} | s_t^{(i)})}$$

- ▶ Or

$$P(\tau^{(i)}, \theta_{other}) = \prod_{t=1}^H p(s_{t+1}^{(i)} | s_t^{(i)}, a_t^{(i)}) \cdot \frac{\pi_{\theta_{other}}(a_t^{(i)} | s_t^{(i)})}{\pi_{\theta_{samp}}(a_t^{(i)} | s_t^{(i)})} \cdot \pi_{\theta_{samp}}(a_t^{(i)} | s_t^{(i)})$$

- ▶ The term $\frac{\pi_{\theta_{other}}(a_t^{(i)} | s_t^{(i)})}{\pi_{\theta_{samp}}(a_t^{(i)} | s_t^{(i)})}$ is the importance sampling term
- ▶ In TRPO, $\pi_{\theta_{sample}} = \pi_{\theta_{old}}$, $\pi_{\theta_{other}} = \pi_{\theta}$
- ▶ We apply the same derivation as for the policy gradient...
- ▶ We get $Loss^{IS}(\theta) = \mathbb{E}_t[\frac{\pi_{\theta}(a_t | s_t)}{\pi_{\theta_{old}}(a_t | s_t)} \hat{A}_{\phi}^{\pi_{\theta}}]$

Trust region



- ▶ The gradient of a function is only accurate close to the point where it is calculated
- ▶ $\nabla_{\theta} J(\theta)$ is only accurate close to the current policy π_{θ}
- ▶ Thus, when updating, π_{θ} must not move too far away from a “trust region” around $\pi_{\theta_{old}}$



Kakade, S. & Langford, J. (2002) Approximately optimal approximate reinforcement learning. In *ICML*, volume 2, pages 267–274

Trust Region Policy Optimization

- ▶ Theory: monotonous improvement towards the optimal policy
(Assumptions do not hold in practice)
- ▶ To ensure small steps, TRPO uses a natural gradient update instead of standard gradient
- ▶ Minimize Kullback-Leibler divergence to previous policy



$$\max_{\theta} \mathbb{E}_t \left[\frac{\pi_{\theta}(\mathbf{a}_t | \mathbf{s}_t)}{\pi_{\theta_{old}}(\mathbf{a}_t | \mathbf{s}_t)} A_{\phi}^{\pi_{\theta_{old}}}(\mathbf{s}_t, \mathbf{a}_t) \right]$$

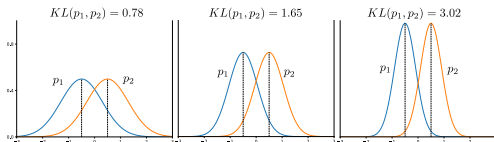
$$\text{subject to } \mathbb{E}_t [KL(\pi_{\theta_{old}}(\cdot | \mathbf{s}) || \pi_{\theta}(\mathbf{a}_t | \mathbf{s}_t))] \leq \epsilon$$

- ▶ In TRPO, optimization performed using a conjugate gradient method to avoid approximating the Fisher Information matrix



Schulman, J., Levine, S., Moritz, P., Jordan, M. I., & Abbeel, P. (2015) Trust Region Policy Optimization. *CoRR*, [abs/1502.05477](https://arxiv.org/abs/1502.05477)

Natural Policy Gradient



- ▶ One way to constrain two stochastic policies to stay close is constraining their KL divergence
- ▶ The KL divergence is smaller when the variance is larger
- ▶ Under fixed KL constraint, it is easier to move the mean further away when the variance is large
- ▶ Thus the mean policy converges first, then the variance is reduced
- ▶ Ensures a large enough amount of exploration noise
- ▶ Other properties presented in the Pierrot et al. (2018) paper



Sham M. Kakade. A natural policy gradient. In *Advances in neural information processing systems*, pp. 1531–1538, 2002



Pierrot, T., Perrin, N., & Sigaud, O. (2018) First-order and second-order variants of the gradient descent: a unified framework. *arXiv preprint arXiv:1810.08102*

Advantage estimation

- ▶ To get $\hat{A}_{\phi}^{\pi_{\theta}}$, an empirical estimate of $V^{\pi_{\theta}}(s)$ is needed
- ▶ TRPO uses a MC estimate approach through regression, but constrains it (as for the policy):

$$\min_{\phi} \sum_{n=0}^N \|V_{\phi}^{\pi_{\theta}}(s_n) - V^{\pi_{\theta}}(s_n)\|^2$$
$$\text{subject to } \frac{1}{N} \sum_{n=0}^N \frac{\|V_{\phi}^{\pi_{\theta}}(s_n) - V_{\phi_{old}}^{\pi_{\theta}}(s_n)\|^2}{2\sigma^2} \leq \epsilon$$

- ▶ Equivalent to a mean KL divergence constraint between $V_{\phi}^{\pi_{\theta}}$ and $V_{\phi_{old}}^{\pi_{\theta}}$
- ▶ Very similar to target critic in DQN, DDPG... Can be implemented in the same way

Properties

- ▶ Moves slowly away from current policy
- ▶ Key: use of line search to deal with the gradient step size
- ▶ More stable than DDPG, performs well in practice, but less sample efficient
- ▶ Conjugate gradient approach not provided in standard tensor gradient libraries, thus not much used
- ▶ Greater impact of PPO
- ▶ Related work: NAC, REPS



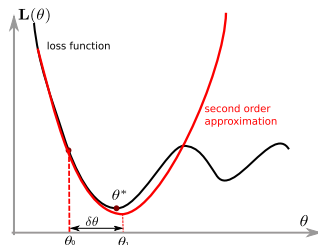
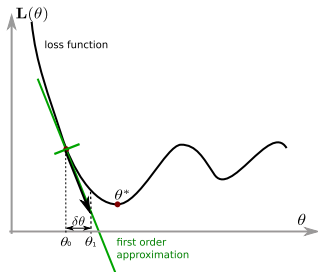
Jan Peters and Stefan Schaal. Natural actor-critic. *Neurocomputing*, 71 (7-9):1180–1190, 2008



Jan Peters, Katharina Mülling, and Yasemin Altun. Relative entropy policy search. In *AAAI*, pp. 1607–1612. Atlanta, 2010

ACKTR

First order versus second order derivative



- ▶ In first order methods, need to define a step size
- ▶ Second order methods provide a more accurate approximation
- ▶ They also provide a true minimum, when the Hessian matrix is symmetric positive-definite (SPD)
- ▶ In both cases, the derivative is very local
- ▶ The trust region constraint applies too

ACKTR

- ▶ K-FAC: Kronecker Factored Approximated Curvature: efficient estimate of the gradient
- ▶ Using block diagonal estimations of the Hessian matrix, to do better than first order
- ▶ ACKTR: TRPO with K-FAC natural gradient calculation
- ▶ But closer to actor-critic updates (see PPO)
- ▶ The per-update cost of ACKTR is only 10% to 25% higher than SGD
- ▶ Improves sample efficiency
- ▶ Not much excitement: less robust gradient approximation?



Yuhuai Wu, Elman Mansimov, Shun Liao, Roger Grosse, and Jimmy Ba (2017) Scalable trust-region method for deep reinforcement learning using Kronecker-factored approximation. *arXiv preprint arXiv:1708.05144*

Any question?



Send mail to: Olivier.Sigaud@upmc.fr



Kakade, S. and Langford, J. (2002).

Approximately optimal approximate reinforcement learning.
In *ICML*, volume 2, pages 267–274.



Kakade, S. M. (2001).

A natural policy gradient.

In Dietterich, T. G., Becker, S., and Ghahramani, Z., editors, *Advances in Neural Information Processing Systems 14 [Neural Information Processing Systems: Natural and Synthetic, NIPS 2001, December 3-8, 2001, Vancouver, British Columbia, Canada]*, pages 1531–1538. MIT Press.



Peters, J., Mülling, K., and Altun, Y. (2010).

Relative entropy policy search.

In *AAAI*, pages 1607–1612. Atlanta.



Peters, J. and Schaal, S. (2008).

Natural actor-critic.

Neurocomputing, 71(7-9):1180–1190.



Pierrot, T., Perrin, N., and Sigaud, O. (2018).

First-order and second-order variants of the gradient descent: a unified framework.
arXiv preprint arXiv:1810.08102.



Schulman, J., Levine, S., Abbeel, P., Jordan, M. I., and Moritz, P. (2015).

Trust region policy optimization.

In Bach, F. R. and Blei, D. M., editors, *Proceedings of the 32nd International Conference on Machine Learning, ICML 2015, Lille, France, 6-11 July 2015*, volume 37 of *JMLR Workshop and Conference Proceedings*, pages 1889–1897. JMLR.org.



Wu, Y., Mansimov, E., Grosse, R. B., Liao, S., and Ba, J. (2017).

Scalable trust-region method for deep reinforcement learning using kronecker-factored approximation.

In Guyon, I., von Luxburg, U., Bengio, S., Wallach, H. M., Fergus, R., Vishwanathan, S. V. N., and Garnett, R., editors, *Advances in Neural Information Processing Systems 30: Annual Conference on Neural Information Processing Systems 2017, December 4-9, 2017, Long Beach, CA, USA*, pages 5279–5288.

