

Negative Momentum for Improved Game Dynamics

Gauthier Gidel*, Reyhane Askari Hemmat*, Mohammad Pezeshki,
Gabriel Huang, Remi Lepriol, Simon Lacoste-Julien, Ioannis Mitliagkas

**equal contribution*

Simple Min-max smooth game:

Gradient dynamic:

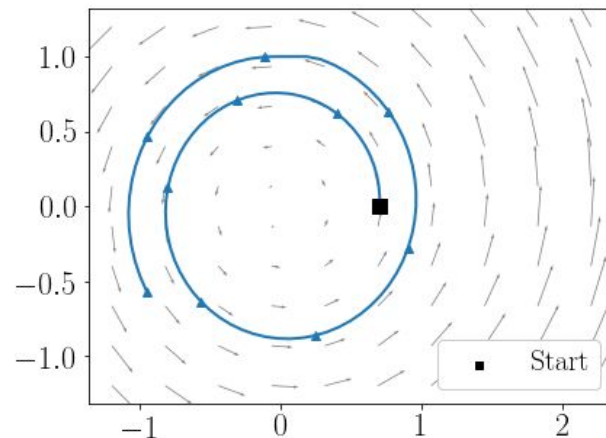
$$\min_{\theta \in \mathbb{R}} \max_{\phi \in \mathbb{R}} \theta \cdot \phi$$

$$\begin{aligned}\theta_{t+1} &= \theta_t - \eta \phi_t \\ \phi_{t+1} &= \phi_t + \eta \theta_t\end{aligned}$$

Simple Min-max smooth game:

$$\min_{\theta \in \mathbb{R}} \max_{\phi \in \mathbb{R}} \theta \cdot \phi$$

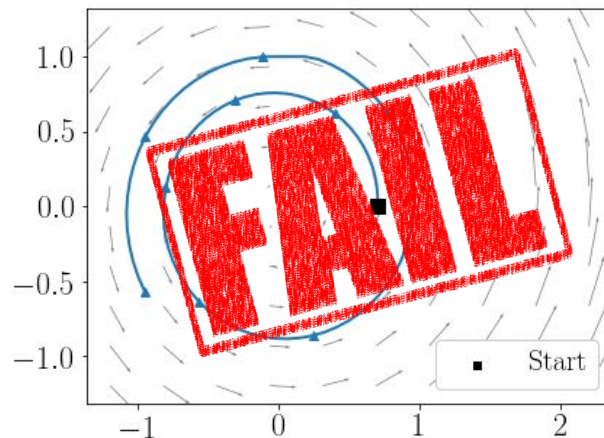
Gradient dynamic:



Simple Min-max smooth game:

$$\min_{\theta \in \mathbb{R}} \max_{\phi \in \mathbb{R}} \theta \cdot \phi$$

Gradient dynamic:



Way to optimize bilinear games

Method	update rule	Convergence (iterates)
Simultaneous Gradient	$\theta_{t+1} = \theta_t - \eta \phi_t$ $\phi_{t+1} = \phi_t + \eta \theta_t$	X

Way to optimize bilinear games

Method	update rule	Convergence (iterates)
Simultaneous Gradient	$\theta_{t+1} = \theta_t - \eta \phi_t$ $\phi_{t+1} = \phi_t + \eta \theta_t$	$\times$
Alternated Gradient	$\theta_{t+1} = \theta_t - \eta_1 \phi_t$ $\phi_{t+1} = \phi_t + \eta_2 \theta_{t+1}$	$\times$

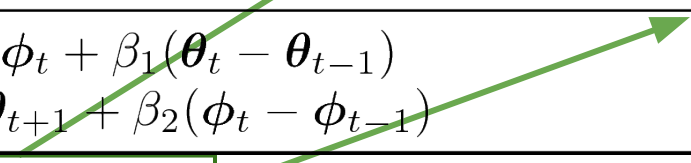
Way to optimize bilinear games

Method	update rule	Convergence (iterates)
Simultaneous Gradient	$\theta_{t+1} = \theta_t - \eta \phi_t$ $\phi_{t+1} = \phi_t + \eta \theta_t$	✗
Alternated Gradient	$\theta_{t+1} = \theta_t - \eta_1 \phi_t$ $\phi_{t+1} = \phi_t + \eta_2 \theta_{t+1}$	✗
Alternated Gradient + negative momentum	$\theta_{t+1} = \theta_t - \eta_1 \phi_t + \beta_1(\theta_t - \theta_{t-1})$ $\phi_{t+1} = \phi_t + \eta_2 \theta_{t+1} + \beta_2(\phi_t - \phi_{t-1})$	✓

Way to optimize bilinear games

Method	update rule	Convergence (iterates)
Simultaneous Gradient	$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta M \boldsymbol{\phi}_t$ $\boldsymbol{\phi}_{t+1} = \boldsymbol{\phi}_t + \eta M^\top \boldsymbol{\theta}_t$	✗
Simultaneous Gradient + negative momentum	$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta M \boldsymbol{\phi}_t + \beta(\boldsymbol{\theta}_t - \boldsymbol{\theta}_{t-1})$ $\boldsymbol{\phi}_{t+1} = \boldsymbol{\phi}_t + \eta M^\top \boldsymbol{\theta}_t + \beta(\boldsymbol{\phi}_t - \boldsymbol{\phi}_{t-1})$	✗ (but less worse)
Alternated Gradient	$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta_1 M \boldsymbol{\phi}_t$ $\boldsymbol{\phi}_{t+1} = \boldsymbol{\phi}_t + \eta_2 M^\top \boldsymbol{\theta}_{t+1}$	✗
Alternated Gradient + negative momentum	$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta_1 M \boldsymbol{\phi}_t + \beta_1(\boldsymbol{\theta}_t - \boldsymbol{\theta}_{t-1})$ $\boldsymbol{\phi}_{t+1} = \boldsymbol{\phi}_t + \eta_2 M^\top \boldsymbol{\theta}_{t+1} + \beta_2(\boldsymbol{\phi}_t - \boldsymbol{\phi}_{t-1})$	✓

Way to optimize bilinear games

Method	update rule	Convergence (iterates)
Simultaneous Gradient	$\theta_{t+1} = \theta_t - \eta \phi_t$ $\phi_{t+1} = \phi_t + \eta \theta_t$	
Simultaneous Gradient + negative momentum	$\theta_{t+1} = \theta_t - \eta M \phi_t + \beta(\theta_t - \theta_{t-1})$ $\phi_{t+1} = \phi_t + \eta M^\top \theta_t + \beta(\phi_t - \phi_{t-1})$	
Alternated Gradient	$\theta_{t+1} = \theta_t - \eta_1 \phi_t$ $\phi_{t+1} = \phi_t + \eta_2 \theta_{t+1}$	
Alternated Gradient + negative momentum	$\theta_{t+1} = \theta_t - \eta_1 \phi_t + \beta_1(\theta_t - \theta_{t-1})$ $\phi_{t+1} = \phi_t + \eta_2 \theta_{t+1} + \beta_2(\phi_t - \phi_{t-1})$	
This talk		

General 2 player games:

Two players aim to minimize their respective cost functions:

$$\theta^* \in \arg \min_{\theta \in \theta} \mathcal{L}^{(\theta)}(\theta, \phi^*) \quad \text{and} \quad \phi^* \in \arg \min_{\phi \in \phi} \mathcal{L}^{(\phi)}(\theta^*, \phi)$$

General 2 player games:

Two players aim to minimize their respective cost functions:

$$\boldsymbol{\theta}^* \in \arg \min_{\boldsymbol{\theta} \in \boldsymbol{\theta}} \mathcal{L}^{(\boldsymbol{\theta})}(\boldsymbol{\theta}, \boldsymbol{\phi}^*) \quad \text{and} \quad \boldsymbol{\phi}^* \in \arg \min_{\boldsymbol{\phi} \in \boldsymbol{\phi}} \mathcal{L}^{(\boldsymbol{\phi})}(\boldsymbol{\theta}^*, \boldsymbol{\phi})$$

Examples:

- Simple class of zero-sum games: ($\mathcal{L}^{(\boldsymbol{\theta})} = -\mathcal{L}^{(\boldsymbol{\phi})}$)

$$\min_{\boldsymbol{\theta} \in \mathbb{R}^d} \max_{\boldsymbol{\phi} \in \mathbb{R}^p} \alpha \|\boldsymbol{\theta}\|_2^2 + (1 - \alpha) \boldsymbol{\theta}^\top \boldsymbol{M} \boldsymbol{\phi} - \alpha \|\boldsymbol{\phi}\|_2^2, \quad \alpha \in [0, 1], \quad \boldsymbol{M} \in \mathbb{R}^{d \times p}$$

General 2 player games:

Two players aim to minimize their respective cost functions:

$$\boldsymbol{\theta}^* \in \arg \min_{\boldsymbol{\theta} \in \boldsymbol{\theta}} \mathcal{L}^{(\boldsymbol{\theta})}(\boldsymbol{\theta}, \boldsymbol{\phi}^*) \quad \text{and} \quad \boldsymbol{\phi}^* \in \arg \min_{\boldsymbol{\phi} \in \boldsymbol{\phi}} \mathcal{L}^{(\boldsymbol{\phi})}(\boldsymbol{\theta}^*, \boldsymbol{\phi})$$

Examples:

- Simple class of zero-sum games: ($\mathcal{L}^{(\boldsymbol{\theta})} = -\mathcal{L}^{(\boldsymbol{\phi})}$)

$$\min_{\boldsymbol{\theta} \in \mathbb{R}^d} \max_{\boldsymbol{\phi} \in \mathbb{R}^p} \alpha \|\boldsymbol{\theta}\|_2^2 + (1 - \alpha) \boldsymbol{\theta}^\top \boldsymbol{M} \boldsymbol{\phi} - \alpha \|\boldsymbol{\phi}\|_2^2, \quad \alpha \in [0, 1], \quad \boldsymbol{M} \in \mathbb{R}^{d \times p}$$

- Generative Adversarial Networks:

$$\mathcal{L}^{(\boldsymbol{\theta})}(\boldsymbol{\theta}, \boldsymbol{\phi}) \stackrel{\text{def}}{=} -\mathbb{E}_{\boldsymbol{x}' \sim q_{\boldsymbol{\theta}}} \log f_{\boldsymbol{\phi}}(\boldsymbol{x}') \quad \text{and} \quad \mathcal{L}^{(\boldsymbol{\phi})}(\boldsymbol{\theta}, \boldsymbol{\phi}) \stackrel{\text{def}}{=} -\mathbb{E}_{\boldsymbol{x} \sim p} \log f_{\boldsymbol{\phi}}(\boldsymbol{x}) - \mathbb{E}_{\boldsymbol{x}' \sim q_{\boldsymbol{\theta}}} \log(1 - f_{\boldsymbol{\phi}}(\boldsymbol{x}'))$$

(non-saturating GAN from Goodfellow et al. 2014)

General 2 player games:

Two players aim to minimize their respective cost functions:

$$\theta^* \in \arg \min_{\theta \in \theta} \mathcal{L}^{(\theta)}(\theta, \phi^*) \quad \text{and} \quad \phi^* \in \arg \min_{\phi \in \phi} \mathcal{L}^{(\phi)}(\theta^*, \phi)$$

Dynamics of gradient based method depends on the gradient vector fields:

$$v(\phi, \theta) := \begin{bmatrix} \nabla_{\phi} \mathcal{L}^{(\phi)}(\phi, \theta) \\ \nabla_{\theta} \mathcal{L}^{(\theta)}(\phi, \theta) \end{bmatrix}$$

General 2 player games:

Two players aim to minimize their respective cost functions:

$$\theta^* \in \arg \min_{\theta \in \Theta} \mathcal{L}^{(\theta)}(\theta, \phi^*) \quad \text{and} \quad \phi^* \in \arg \min_{\phi \in \Phi} \mathcal{L}^{(\phi)}(\theta^*, \phi)$$

Dynamics of gradient based method depends on the gradient vector fields:

$$v(\phi, \theta) := \begin{bmatrix} \nabla_{\phi} \mathcal{L}^{(\phi)}(\phi, \theta) \\ \nabla_{\theta} \mathcal{L}^{(\theta)}(\phi, \theta) \end{bmatrix}$$

And its associated Jacobian,

$$\nabla v(\phi, \theta) := \begin{bmatrix} \nabla_{\phi}^2 \mathcal{L}^{(\phi)}(\phi, \theta) & \nabla_{\phi} \nabla_{\theta} \mathcal{L}^{(\phi)}(\phi, \theta) \\ \nabla_{\phi} \nabla_{\theta} \mathcal{L}^{(\theta)}(\phi, \theta)^T & \nabla_{\theta}^2 \mathcal{L}^{(\theta)}(\phi, \theta) \end{bmatrix}$$

Fixed point dynamics

Gradient method is defined as the repetition of the operator:

$$F_{\eta}(\phi, \theta) = [\phi \quad \theta]^{\top} - \eta v(\phi, \theta).$$

Thus, the sequence computed is

$$(\phi_t, \theta_t) = \underbrace{F_{\eta} \circ \dots \circ F_{\eta}}_t(\phi_0, \theta_0) = F_{\eta}^{(t)}(\phi_0, \theta_0)$$

Fixed point dynamics

Gradient method is defined as the repetition of the operator:

$$F_{\eta}(\phi, \theta) = [\phi \quad \theta]^{\top} - \eta v(\phi, \theta).$$

Thus, the sequence computed is

$$(\phi_t, \theta_t) = \underbrace{F_{\eta} \circ \dots \circ F_{\eta}}_t(\phi_0, \theta_0) = F_{\eta}^{(t)}(\phi_0, \theta_0)$$

We aim to converge to a *Nash Equilibrium*:

Tuning the step size

Jacobian of our fixed point operator:

$$\nabla F_{\eta}(\boldsymbol{\theta}^*, \boldsymbol{\phi}^*) = I_n - \eta \nabla \boldsymbol{v}(\boldsymbol{\theta}^*, \boldsymbol{\phi}^*)$$

- To have fixed point we need $\nabla \boldsymbol{v}(\boldsymbol{\phi}^*, \boldsymbol{\theta}^*)$ to be definite positive.
- Thus, small enough step-size $\implies$ Eigenvalues in the unit disk.
- Want to find optimal step-size.

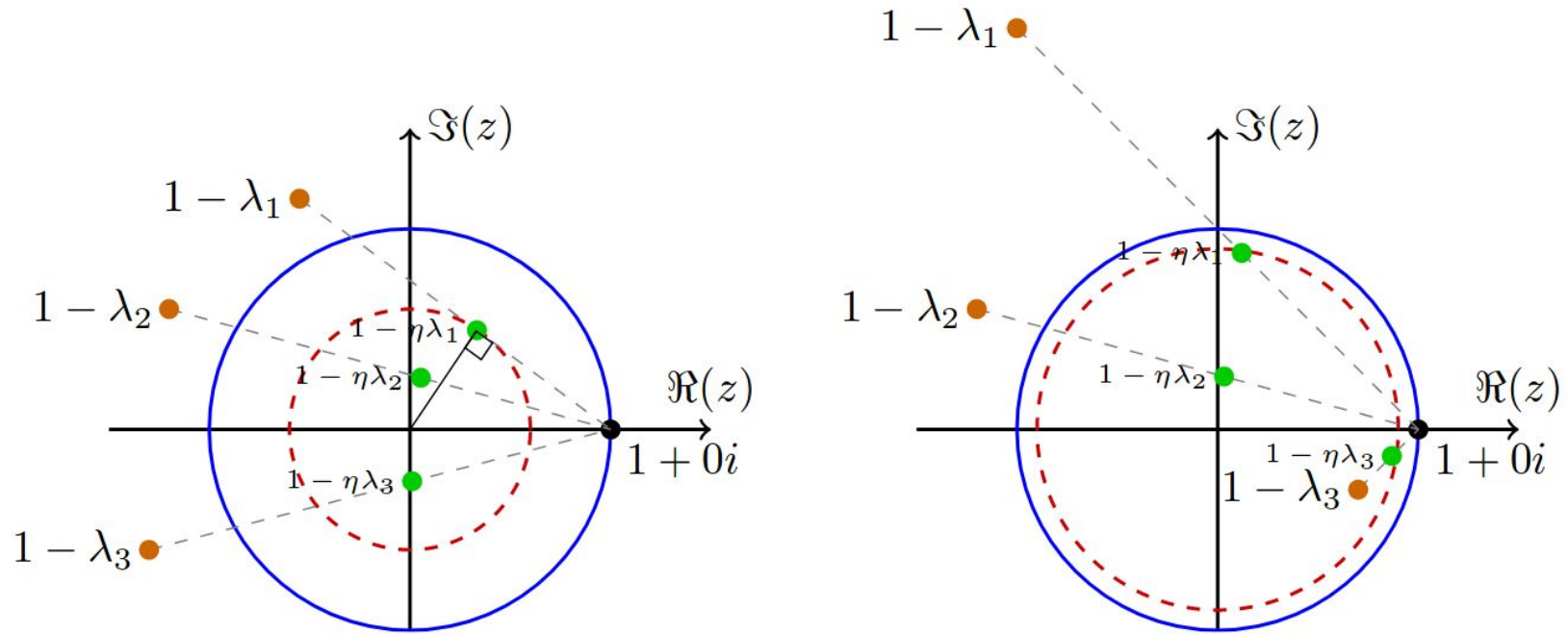
Fixed point dynamics

Jacobian of our fixed point operator: $\nabla F_\eta(\boldsymbol{\theta}^*, \boldsymbol{\phi}^*) = I_n - \eta \nabla \boldsymbol{v}(\boldsymbol{\theta}^*, \boldsymbol{\phi}^*)$

Theorem 1 (Proposition 4.4.1 Bertsekas (1999)). *Let $\mu_{\max}$ be the largest eigenvalue (in magnitude) of $\nabla F_\eta(\boldsymbol{\phi}^*, \boldsymbol{\theta}^*)$. If $\rho_{\max} := |\mu_{\max}| < 1$ then, for $(\boldsymbol{\phi}_0, \boldsymbol{\theta}_0)$ in a neighborhood of $(\boldsymbol{\phi}^*, \boldsymbol{\theta}^*)$, the sequence $(\boldsymbol{\phi}_t, \boldsymbol{\theta}_t)$ converges to the local Nash equilibrium at a rate of $\mathcal{O}((\rho_{\max} + \epsilon)^t)$, $\forall \epsilon > 0$.*

- Local convergence.
- Stationary point may not be a Nash equilibrium. (See Adolphs et al. 2018)
- But any Nash equilibrium is an stationary point.
- In this talk: local results on stationary points.

Tuning the step size



Negative Momentum

Recall Polyak's momentum :

$$\mathbf{x}_{t+1} = \mathbf{x}_t - \eta \mathbf{v}(\mathbf{x}_t) + \beta(\mathbf{x}_t - \mathbf{x}_{t-1}), \quad \mathbf{x}_t = (\boldsymbol{\theta}_t, \phi_t)$$

Fixed point operator requires a **state augmentation** : (because need previous iterates)

$$F_{\eta, \beta}(\mathbf{x}_t, \mathbf{x}_{t-1}) := \begin{bmatrix} \mathbf{I}_n & \mathbf{0}_n \\ \mathbf{I}_n & \mathbf{0}_n \end{bmatrix} \begin{bmatrix} \mathbf{x}_t \\ \mathbf{x}_{t-1} \end{bmatrix} - \eta \begin{bmatrix} \mathbf{v}(\mathbf{x}_t) \\ \mathbf{0}_n \end{bmatrix} + \beta \begin{bmatrix} \mathbf{I}_n & -\mathbf{I}_n \\ \mathbf{0}_n & \mathbf{0}_n \end{bmatrix} \begin{bmatrix} \mathbf{x}_t \\ \mathbf{x}_{t-1} \end{bmatrix}$$

Negative Momentum

Theorem 3. *The eigenvalues of $\nabla F_{\eta,\beta}(\phi^*, \theta^*)$ are*

$$\mu_{\pm}(\beta, \eta, \lambda) := (1 - \eta\lambda + \beta) \frac{1 \pm \Delta^{\frac{1}{2}}}{2}, \quad (9)$$

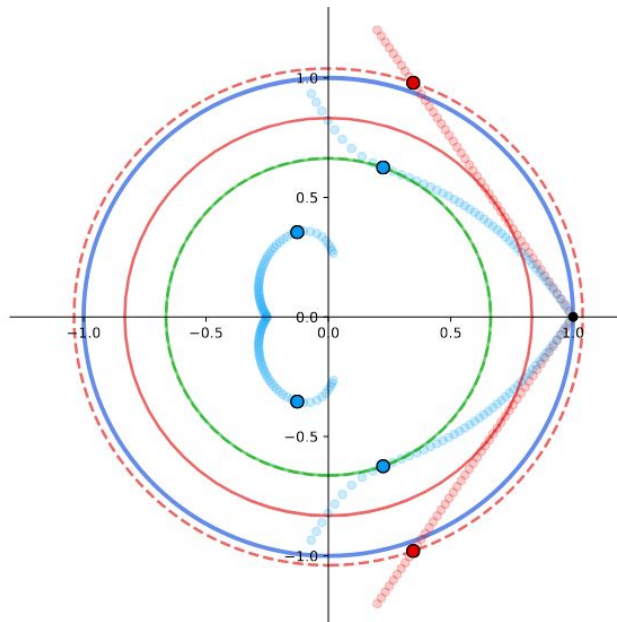
where $\Delta := 1 - \frac{4\beta}{(1-\eta\lambda+\beta)^2}$, $\lambda \in Sp(\nabla v(\phi^*, \theta^*))$ and $\Delta^{\frac{1}{2}}$ is the complex square root of Δ with positive real part³. Moreover we have the following Taylor approximation,

$$\mu_+(\beta, \eta, \lambda) = 1 - \eta\lambda - \beta \frac{\eta\lambda}{1 - \eta\lambda} + O(\beta^2) \quad \text{and} \quad \mu_-(\beta, \eta, \lambda) = \frac{\beta}{1 - \eta\lambda} + O(\beta^2) \quad (10)$$

³ If Δ is a negative real number we set $\Delta^{\frac{1}{2}} := i\sqrt{-\Delta}$

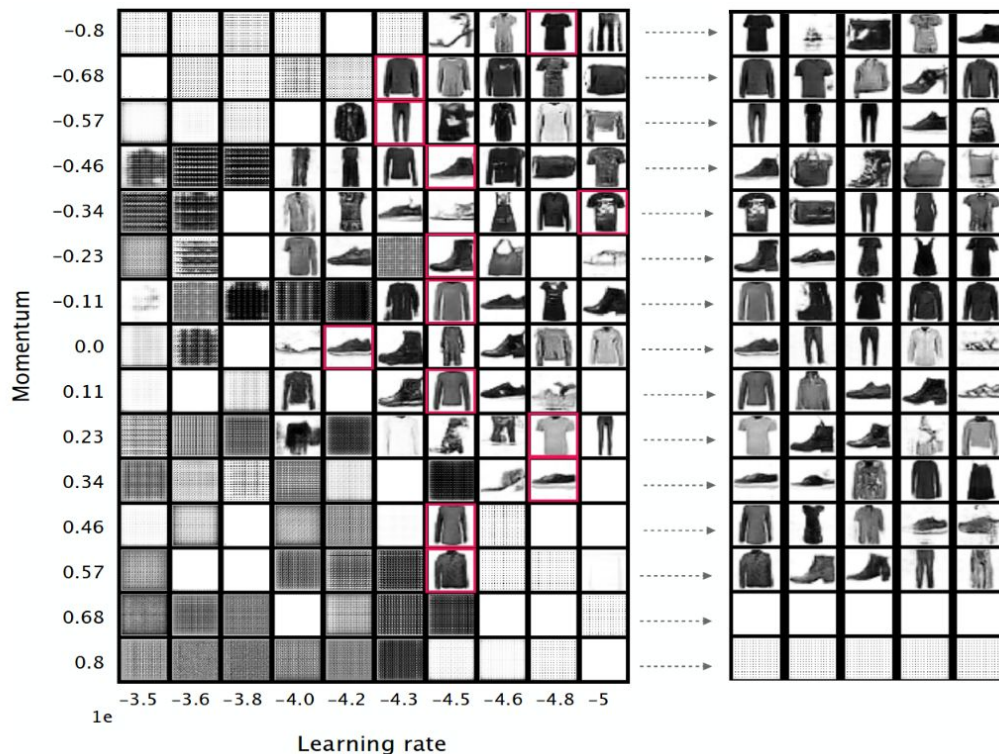
Negative Momentum

- **Fixed** momentum.
(- 0.25)
- Step-size is **not** fixed.
- Helps when the eigenvalue has large imaginary part.



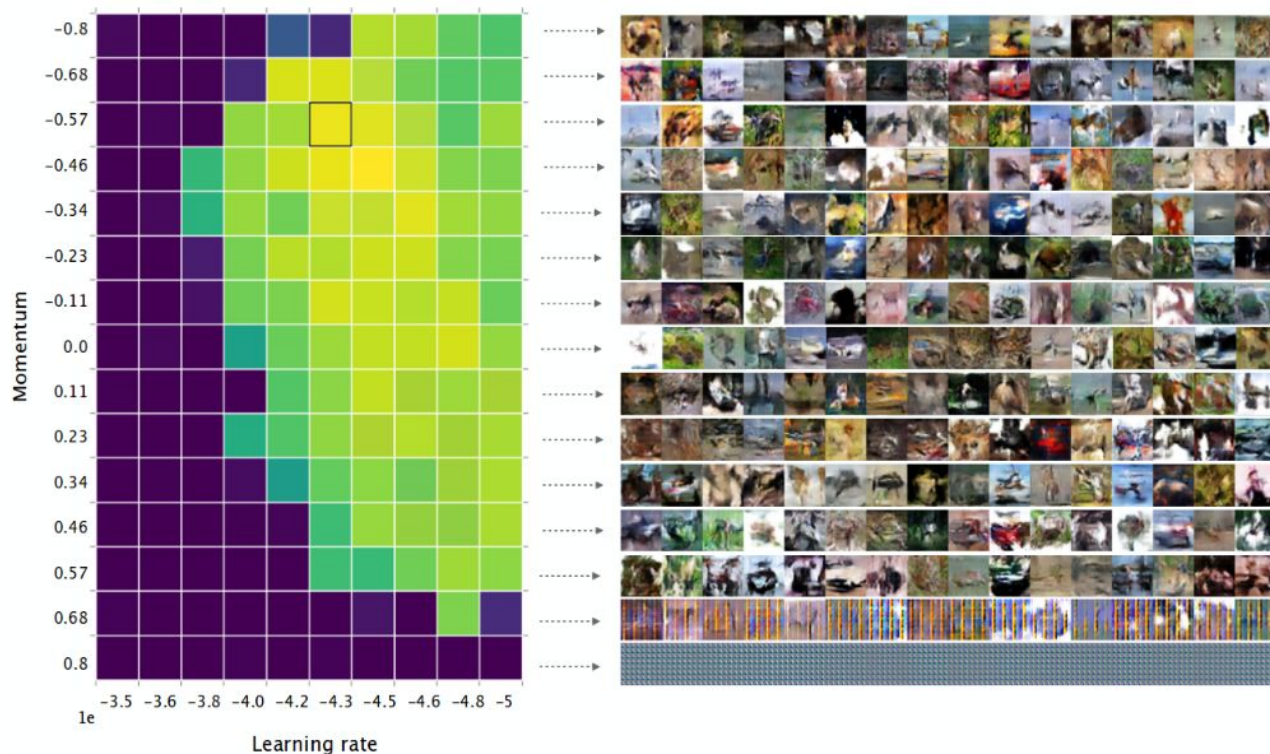
What happens in practice ?

Fashion MNIST:



What happen in practice ?

CIFAR-10:



Negative Momentum

To sum up:

- Negative momentum seems to improve the behaviour of the “bad” eigenvalues.
- If small enough seems to always help.
- It also allows larger step-size.

Thank you !

If you are interested in that topic:

- **NIPS Workshop : Smooth Games Optimization and Machine Learning**

Co-organized with:

Simon Lacoste-Julien · Ioannis Mitliagkas · Vasilis Syrgkanis · Eva Tardos

· Leon Bottou · Sebastian Nowozin

Soon : Call for contributions !!!