

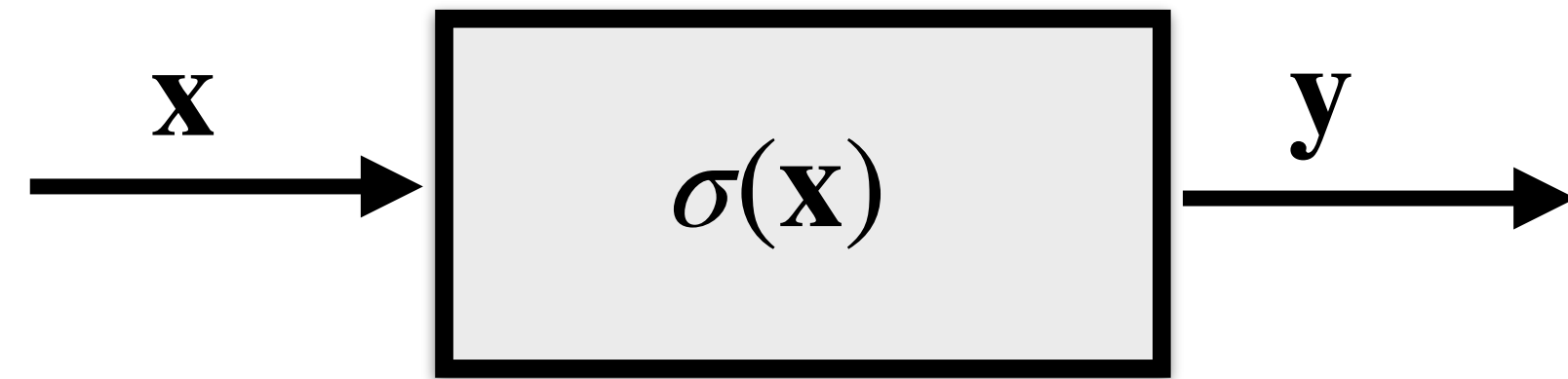
Implicit layers

Backpropagation through root finder, constrained and unconstrained optimization and ODE solvers

Karel Zimmermann

Explicit layers

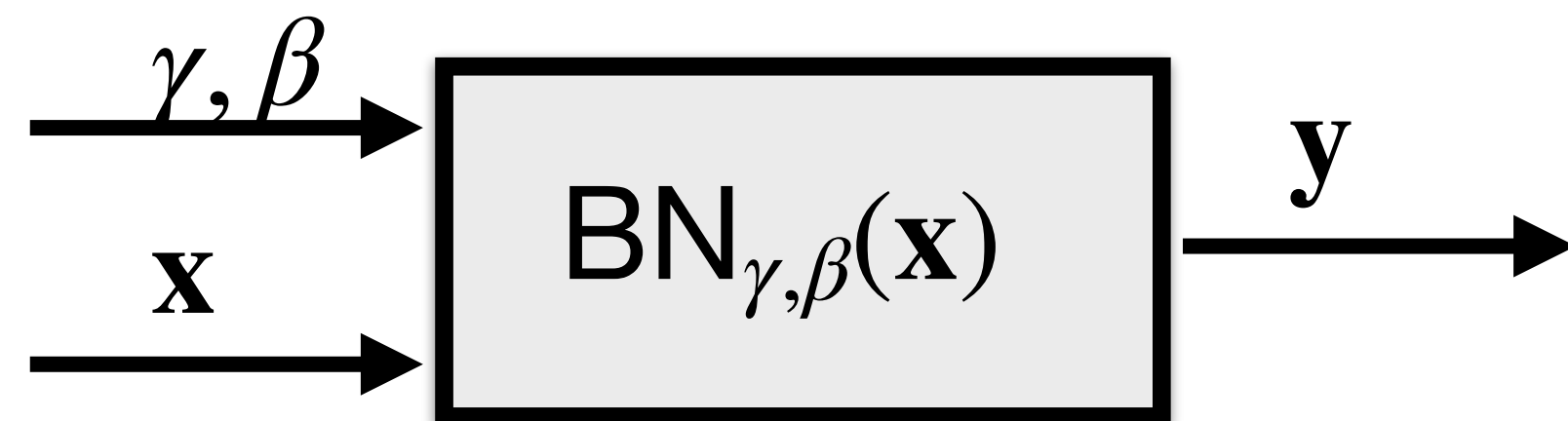
Sigmoid:



Convolution:



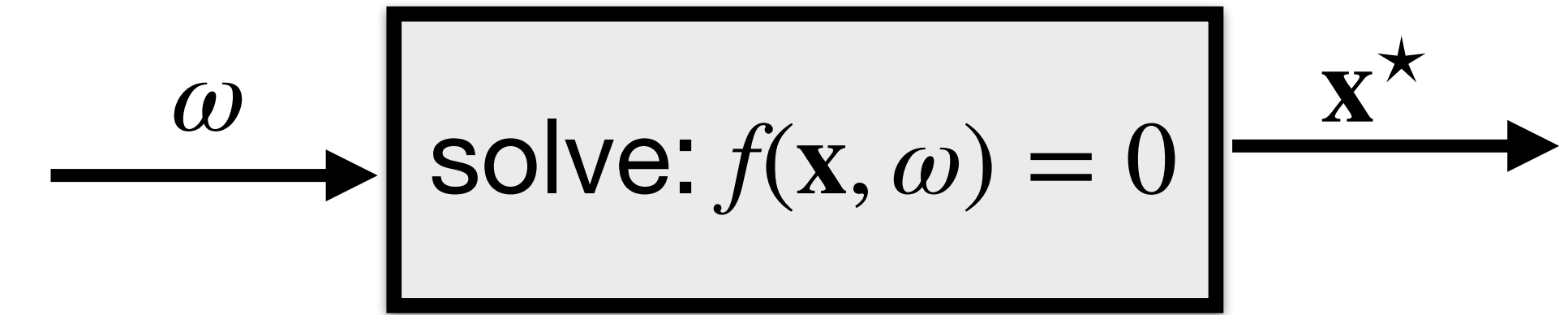
Batch-norm:



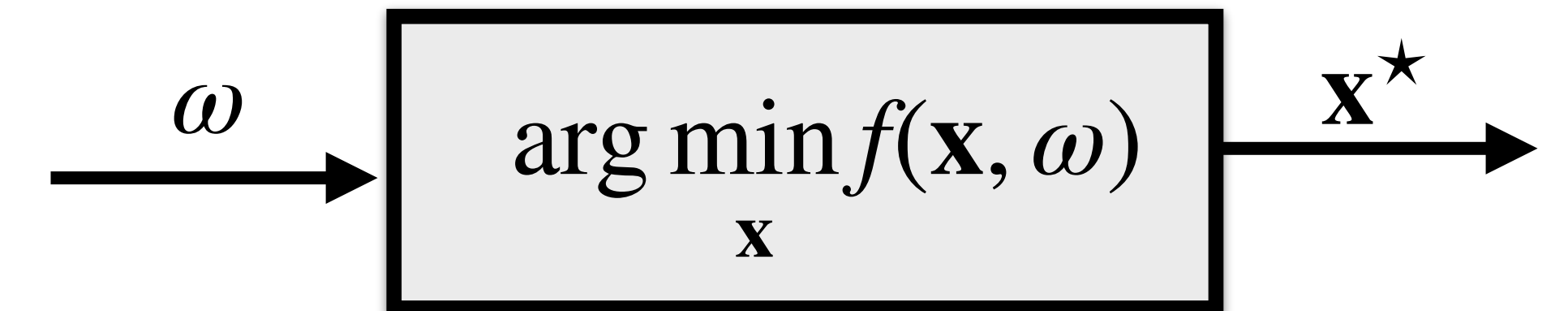
Can I use autodiff on the problem solver?

Implicit layers

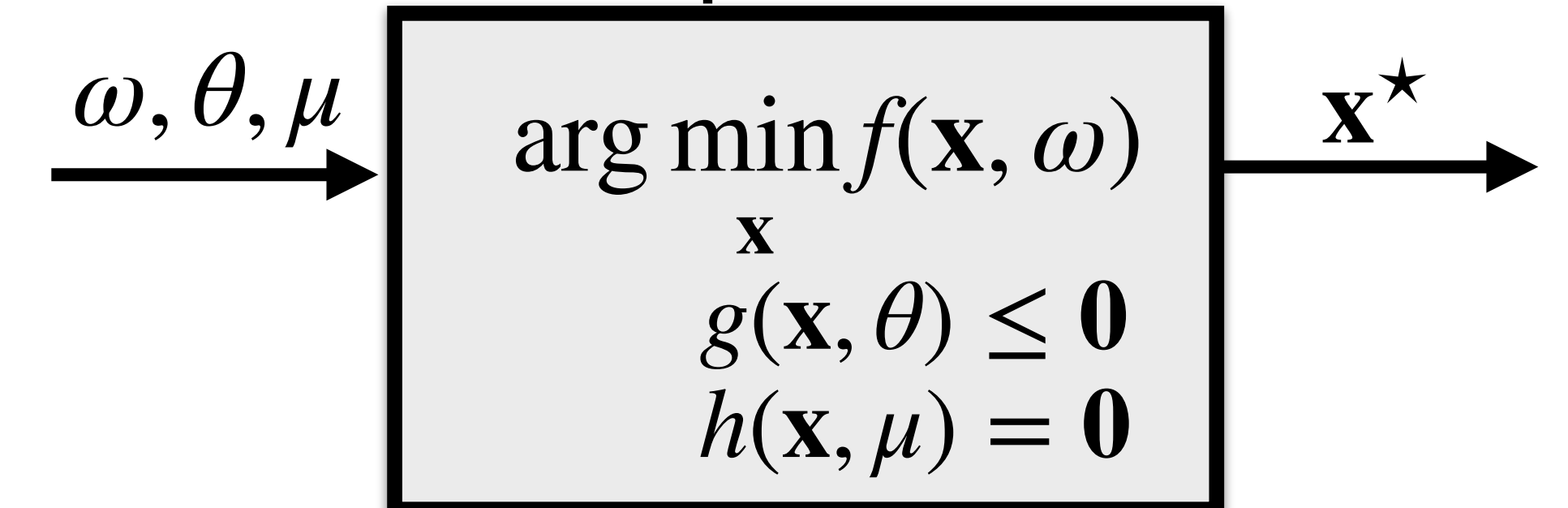
Root finder:



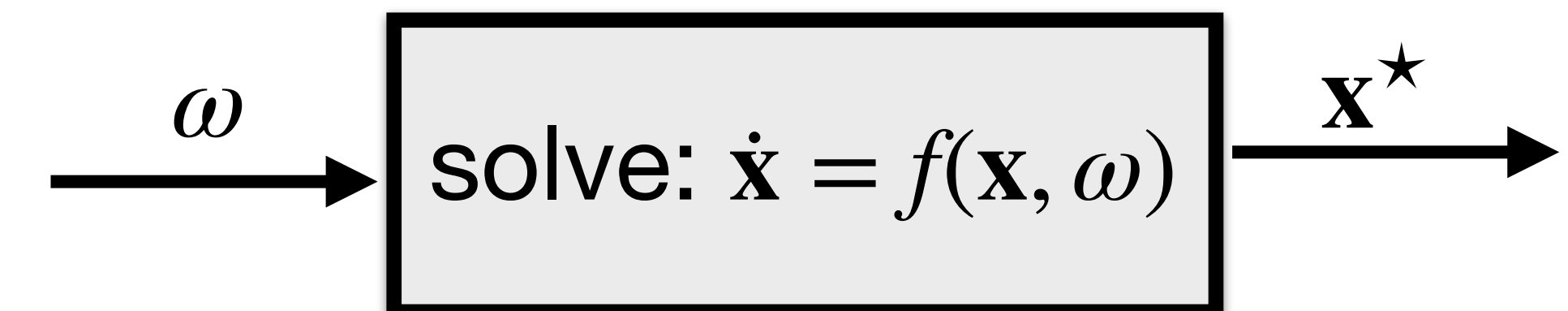
Unconstrained optimizer:



Constrained optimizer:

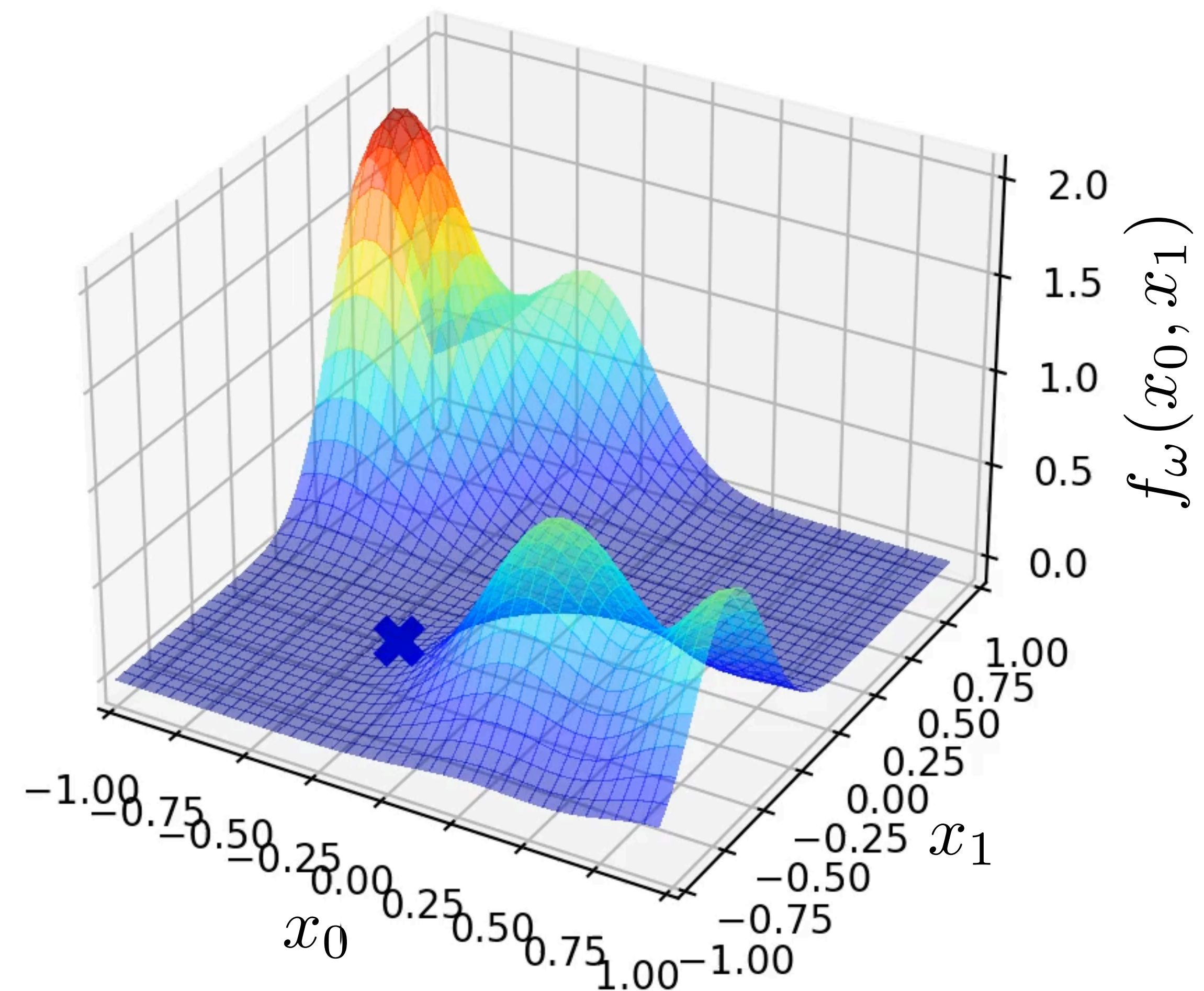


ODE solver:

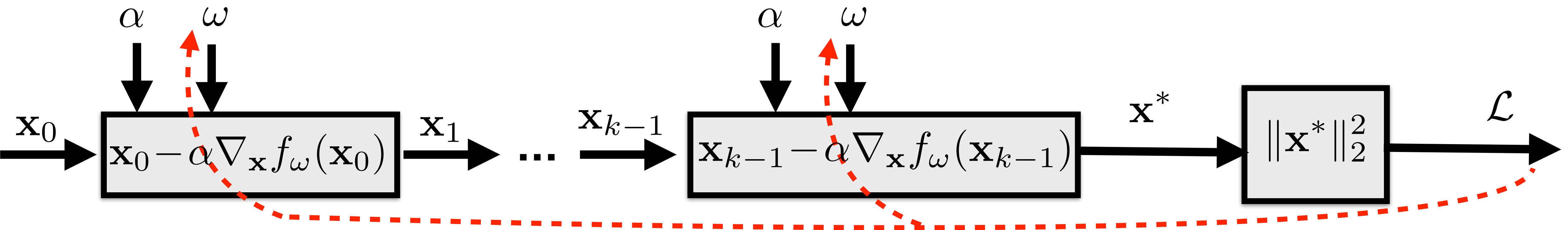


Backpropagation through the solver of an optimization problem

$$\omega^* = \arg \min_{\omega} \|\text{SGD: } \arg \min_{\mathbf{x}} f_{\omega}(\mathbf{x})\|^2$$

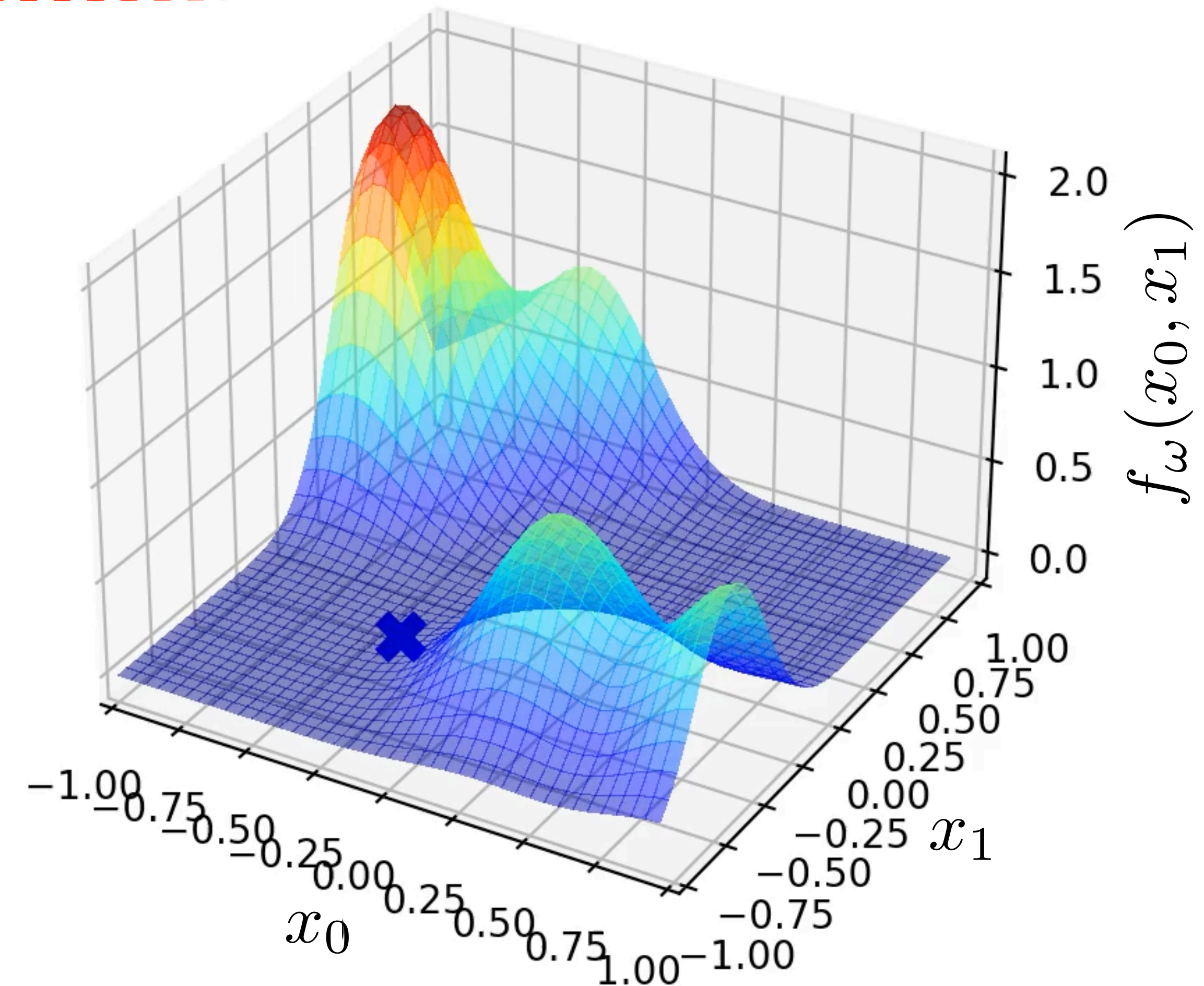


Backpropagation through the solver of an optimization problem



PyTorch:

```
for i in range(iter):  
    x = x - torch.autograd.grad(f, x, ...  
                                create_graph=True, ...  
                                retain_graph=True)  
  
    loss = x.norm()  
    torch.autograd.grad(loss, omega)
```



Naive straightforward solution

- **Feedforward** pass through the **implicit layer** (solver) is a **code**
=> Its computational graph can be **backpropagated via autograd**.
- **Disadvantage:**
 - Problems of autodiff in complex solvers (“if” has often zero or undefined grad)

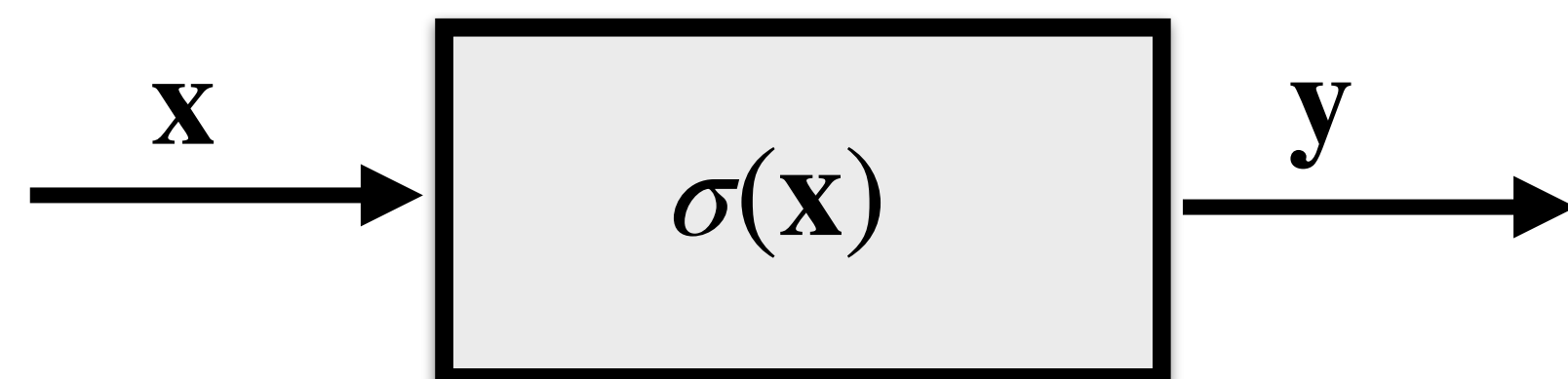
```
def solver(x, w):  
    if x>w:  
        y = 2*x  
    else:  
        y = sin(x)  
    return y
```

$$\frac{\partial \text{solver}(\mathbf{x}, \omega)}{\partial \omega} = \begin{cases} 0 & \mathbf{x} \neq \omega \\ \text{nan} & \mathbf{x} = \omega \end{cases}$$

- Memory and computationally inefficient (e.g. iterations of simplex method)
- **Advantage:**
 - It can backpropagate towards internal parameters of the solver and optimize the optimization itself.

Explicit layers

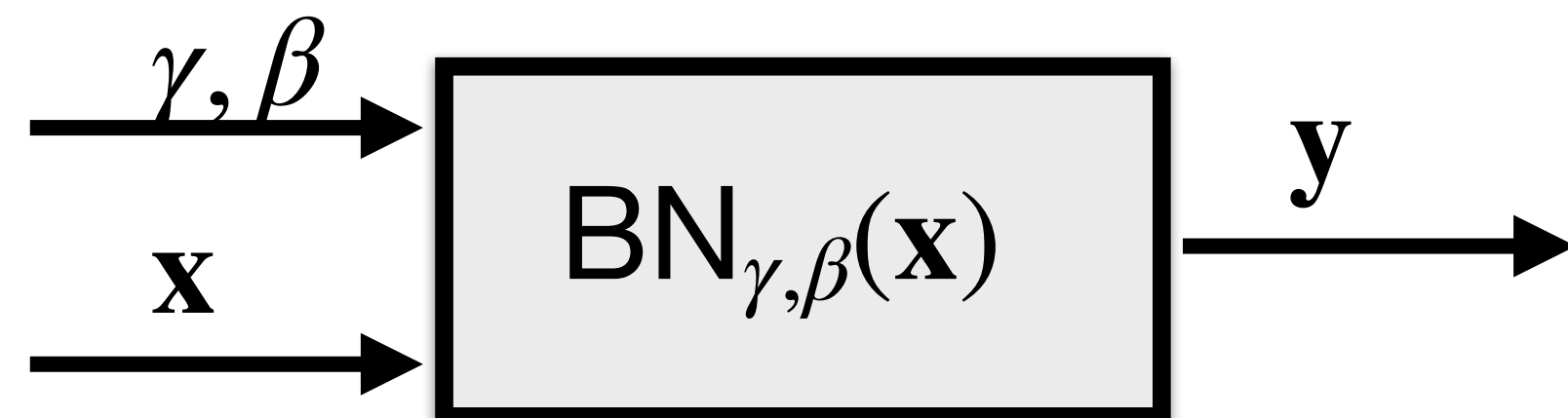
Sigmoid:



Convolution:

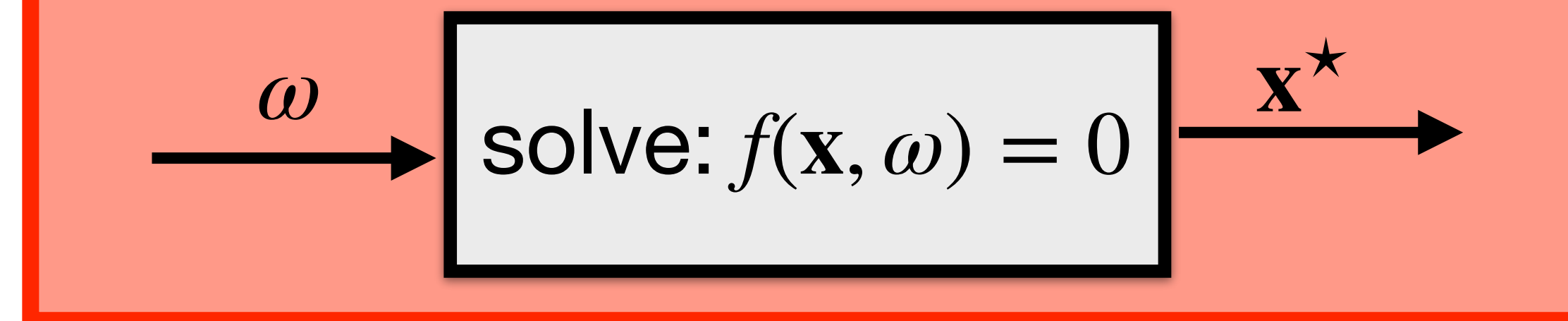


Batch-norm:

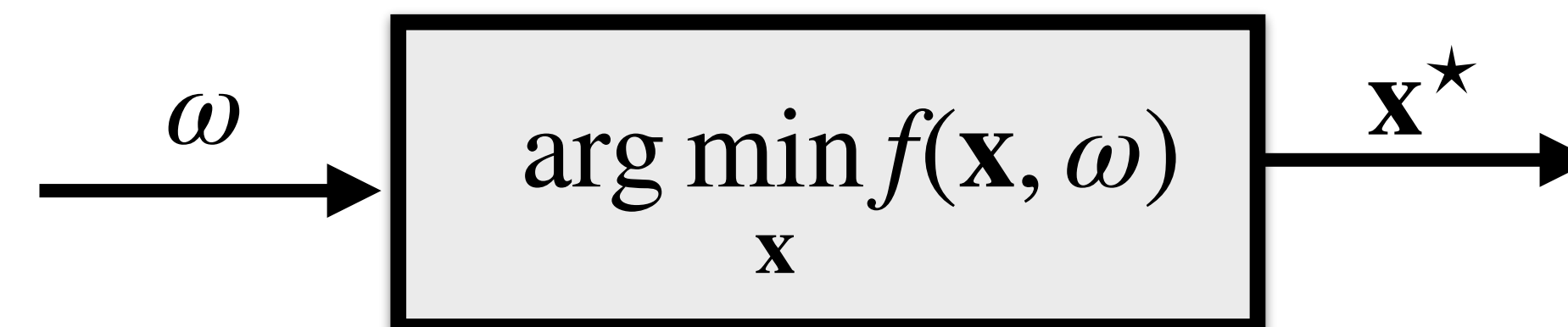


Implicit layers

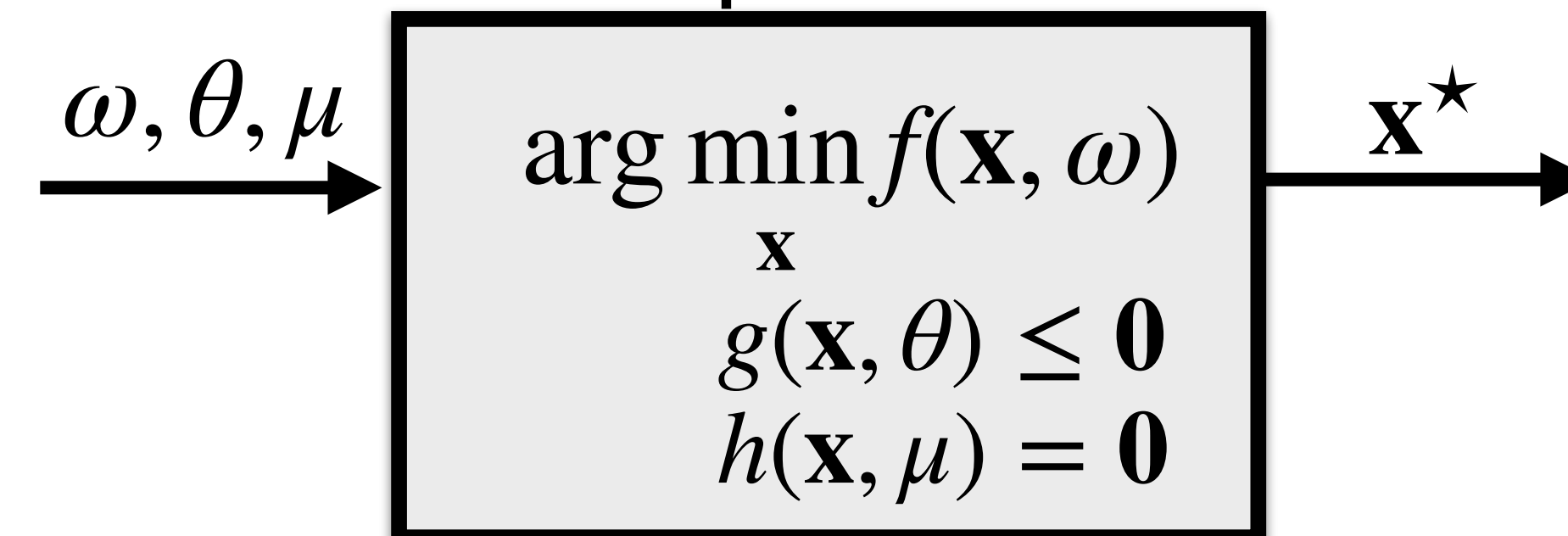
Root finder:



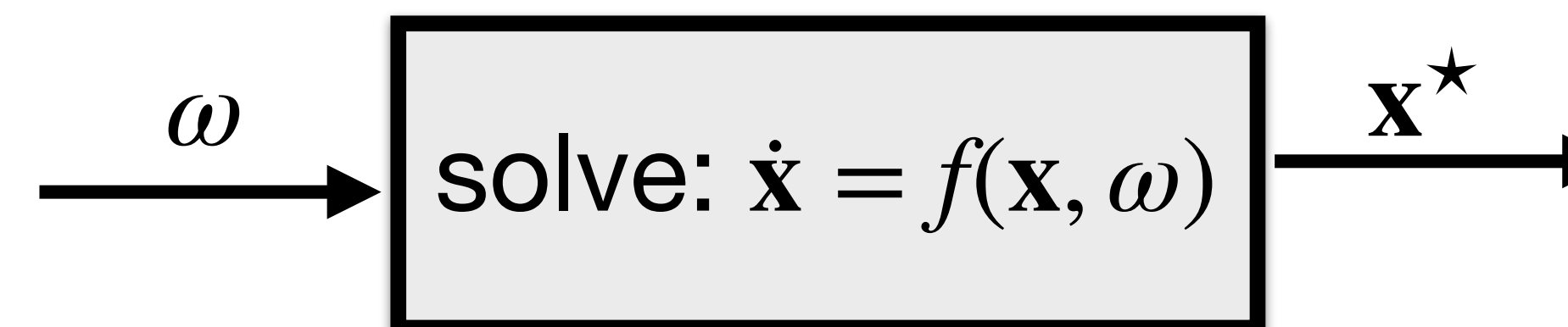
Unconstrained optimizer:



Constrained optimizer:



ODE solver:



Root finder

- Let's assume that root finder is a **compiled binary** w/o the access to its source
- Given parameters ω , the code delivers solution $\mathbf{x}^\star$ such that $f(\mathbf{x}^\star, \omega) = 0$
- This **solution depends on parameters**, we refer it $\mathbf{x}^\star(\omega) : \mathbb{R}^n \rightarrow \mathbb{R}^m$
- We want to find $\frac{\partial \mathbf{x}^\star(\omega)}{\partial \omega} = ?$ can we do better than a numerical differentiations?

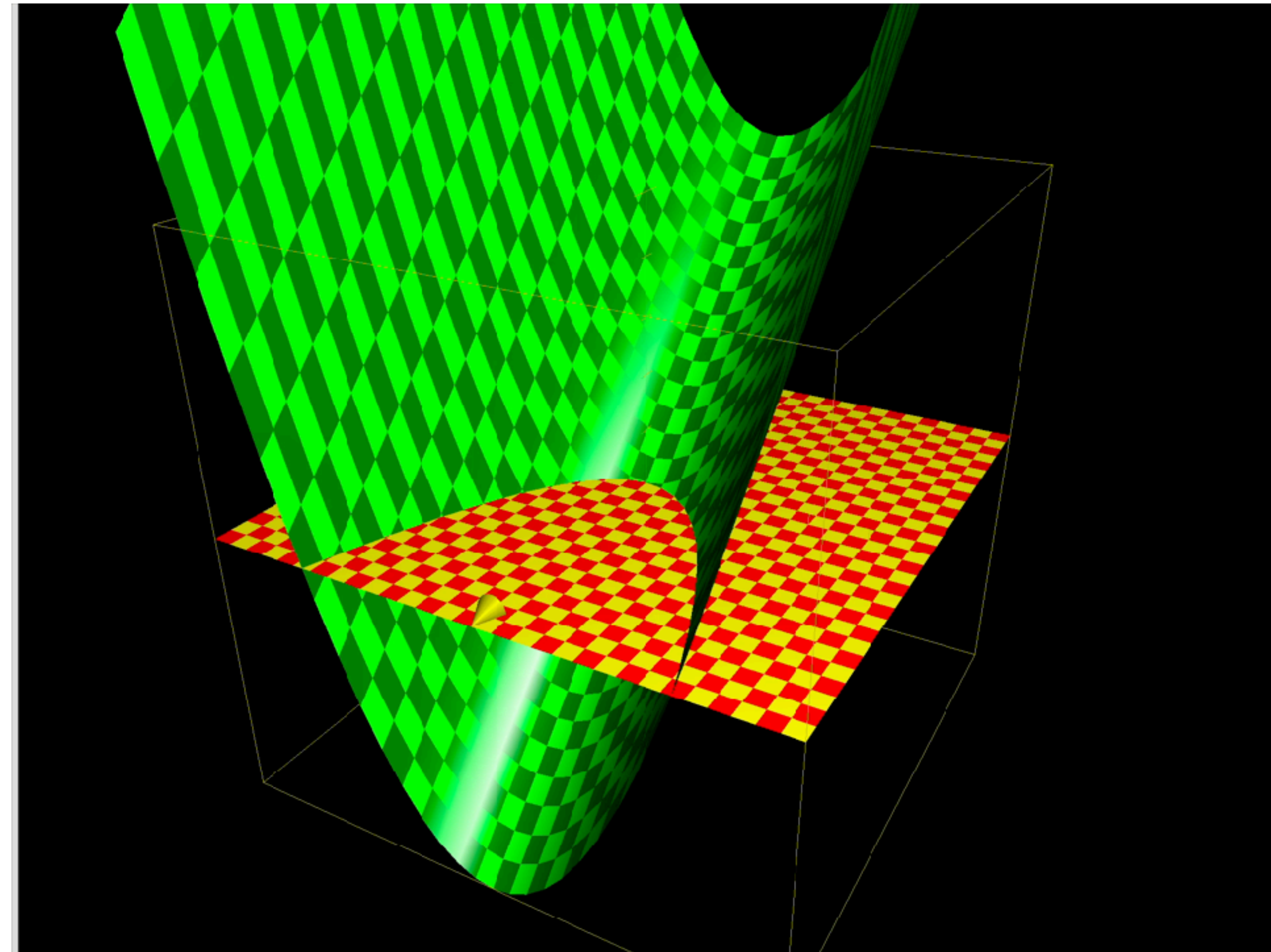
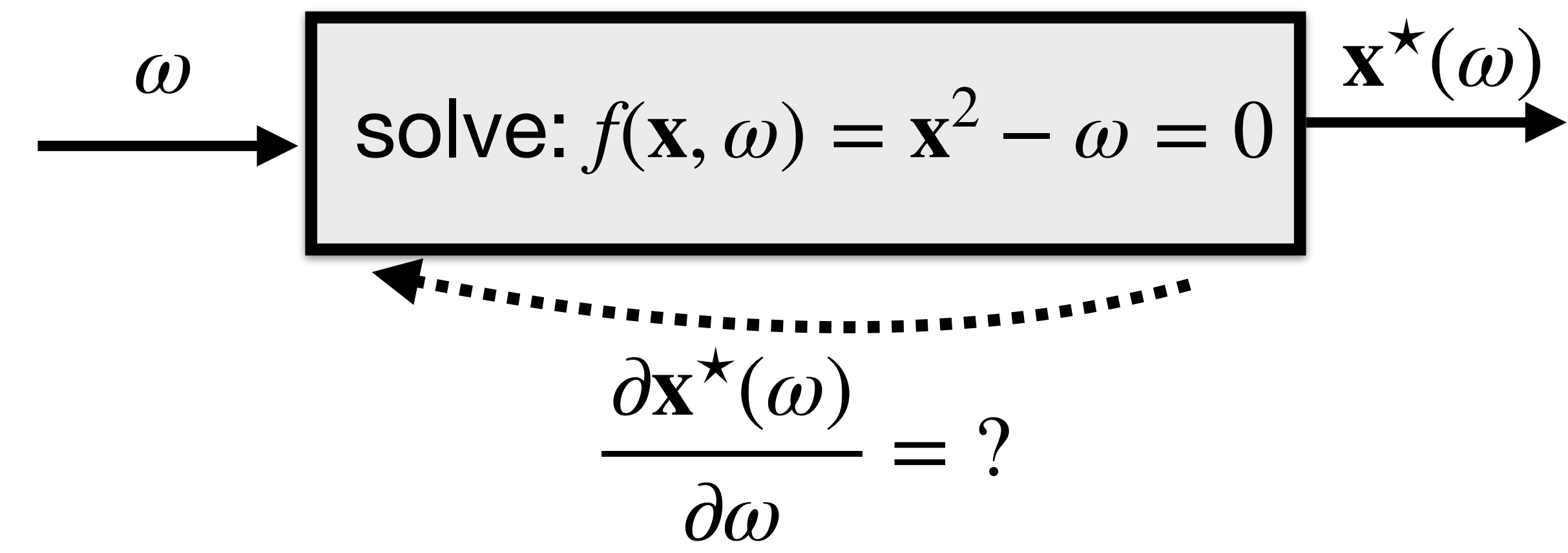
$$f(\mathbf{x}^\star(\omega), \omega) = 0$$

Given ω , the root function $\mathbf{x}^\star(\omega)$ satisfy this equation.

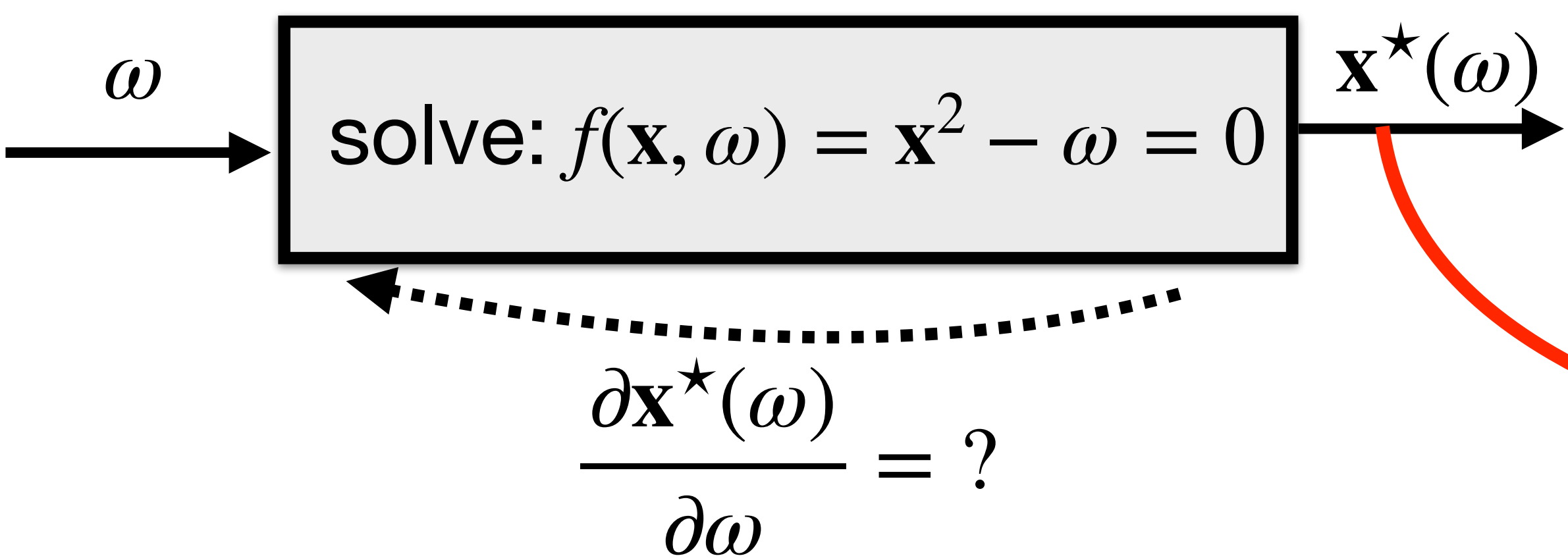
$$\frac{\partial f(\mathbf{x}^\star(\omega), \omega)}{\partial \omega} = 0$$

ω is allowed to **change only in directions** which **does not change** the value of $f(\mathbf{x}^\star(\omega), \omega)$ in order to stay within the solution manifold

Example:

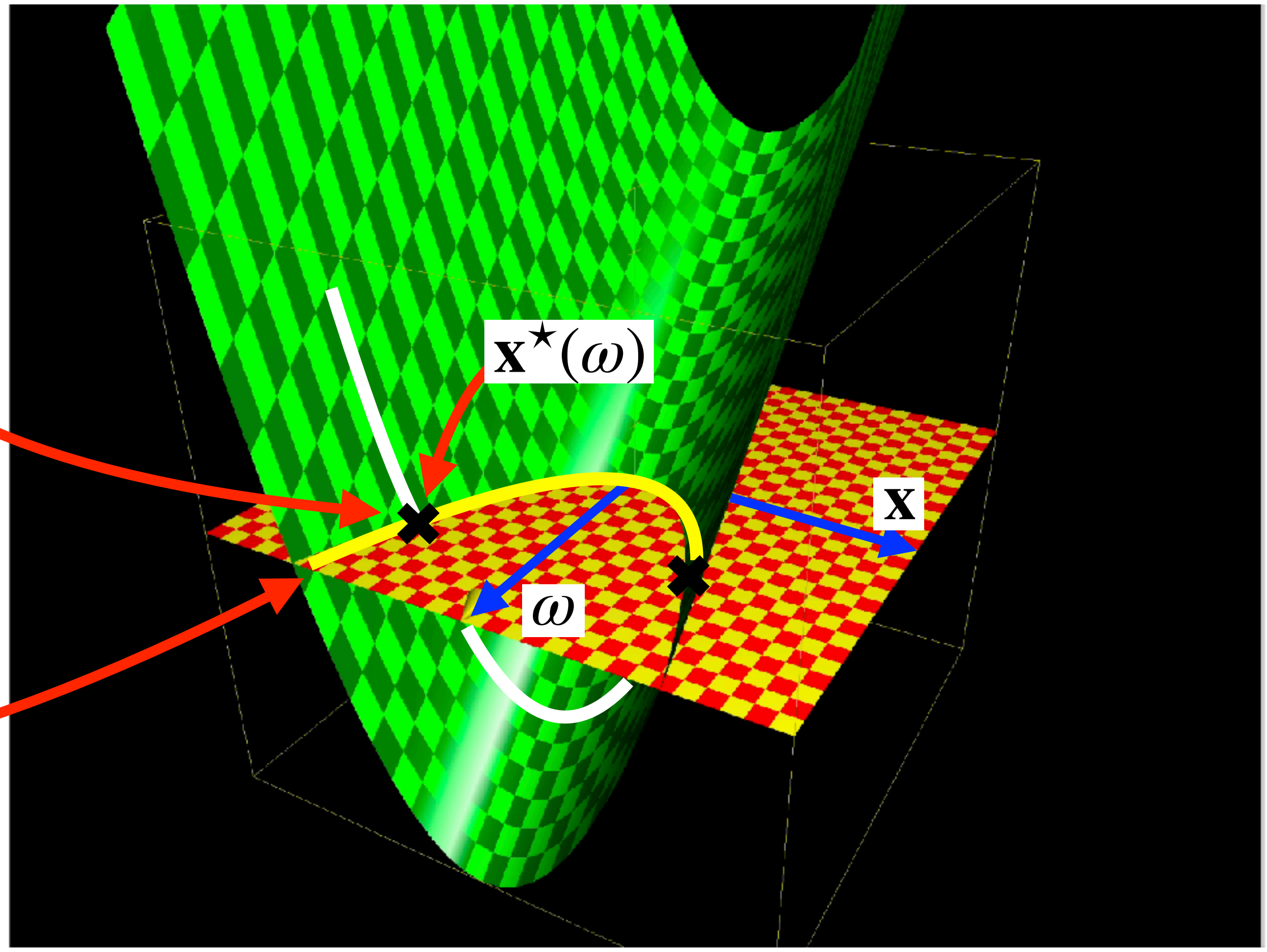


Example:

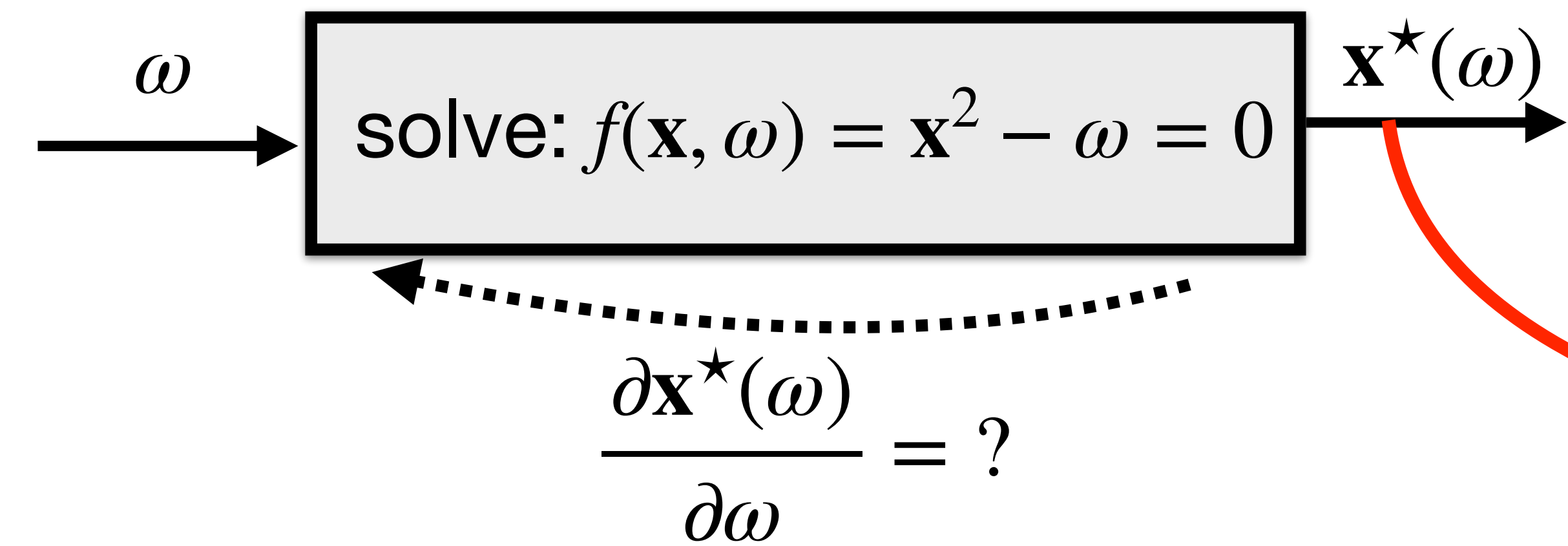


While changing ω , root function $\mathbf{x}^\star(\omega)$ follows the yellow curve

$f(\mathbf{x}^\star(\omega), \omega) = \mathbf{x}^\star(\omega)^2 - \omega = 0$

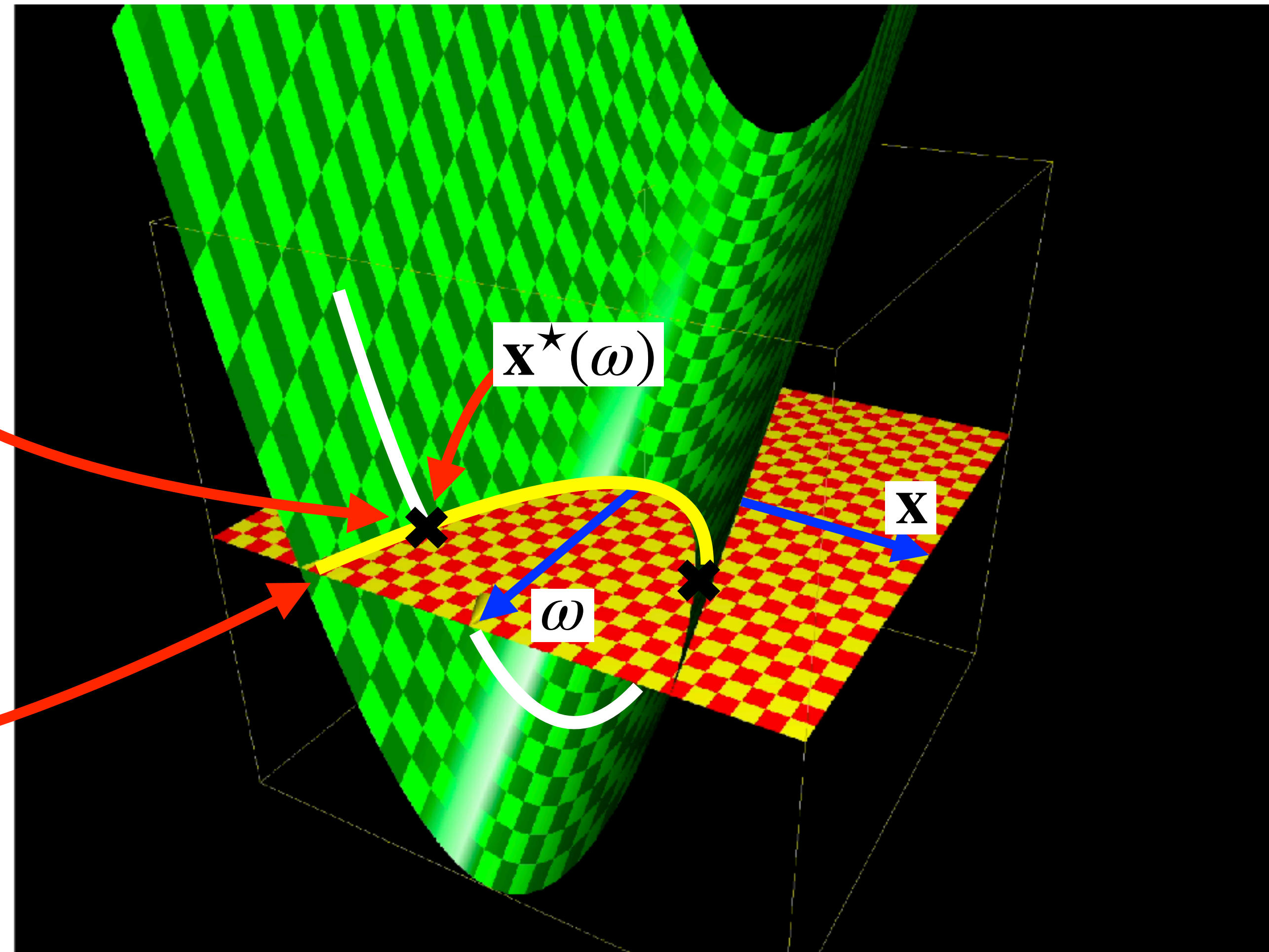


Example:

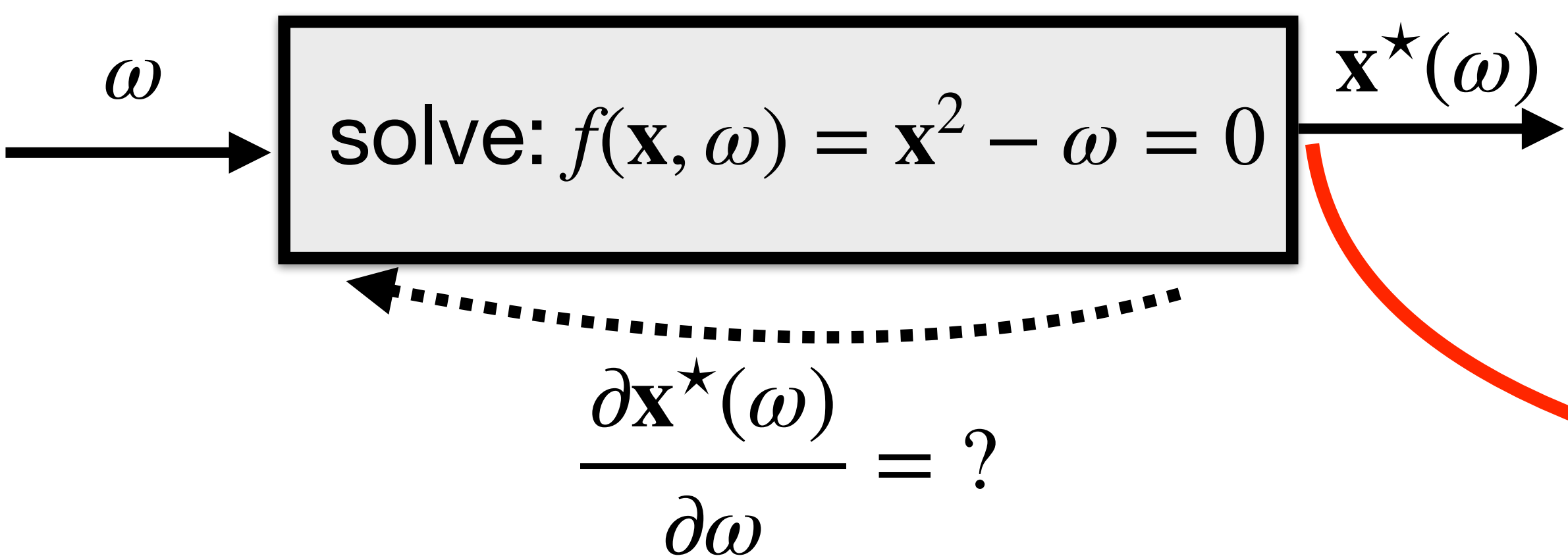


While changing ω , root function $\mathbf{x}^\star(\omega)$ follows the yellow curve

$$f(\mathbf{x}^\star(\omega), \omega) = \mathbf{x}^\star(\omega)^2 - \omega = 0$$

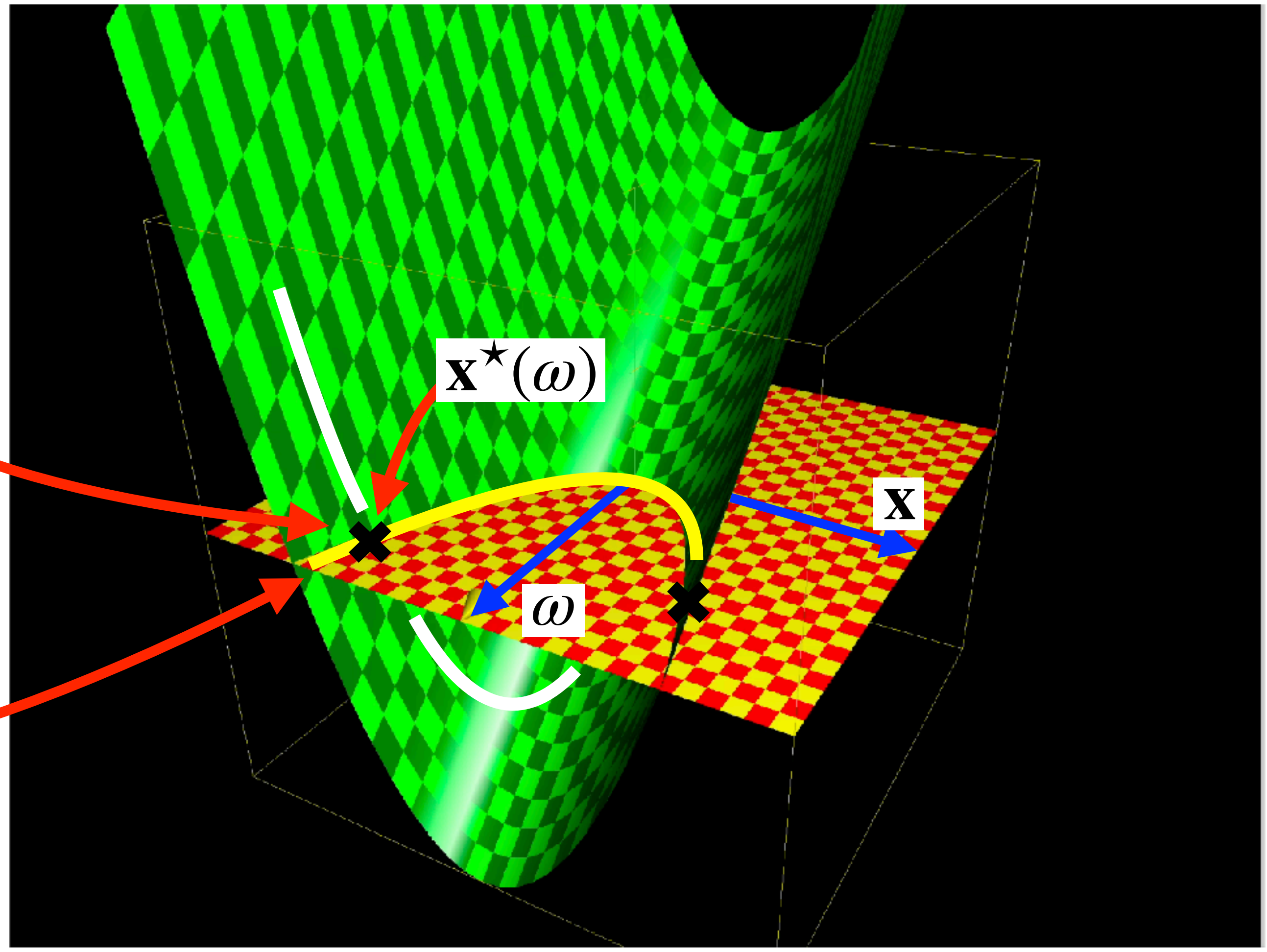


Example:

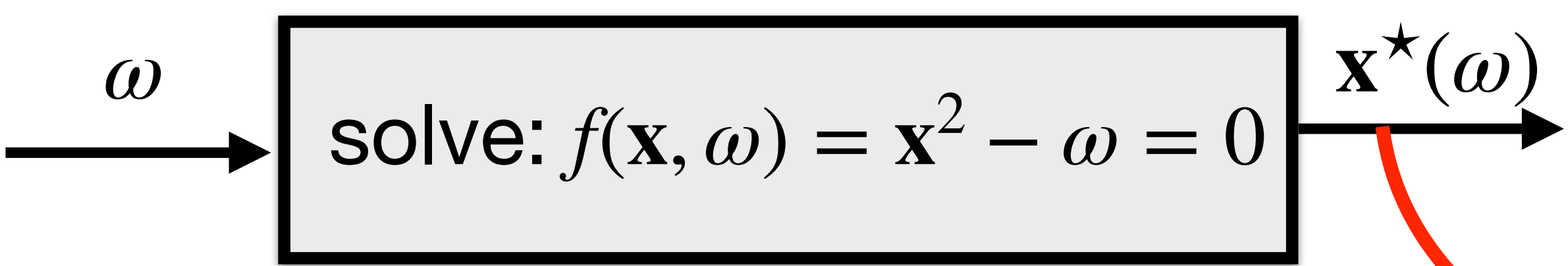


While changing ω , root function $\mathbf{x}^\star(\omega)$ follows the yellow curve

$$f(\mathbf{x}^\star(\omega), \omega) = \mathbf{x}^\star(\omega)^2 - \omega = 0$$



Example:



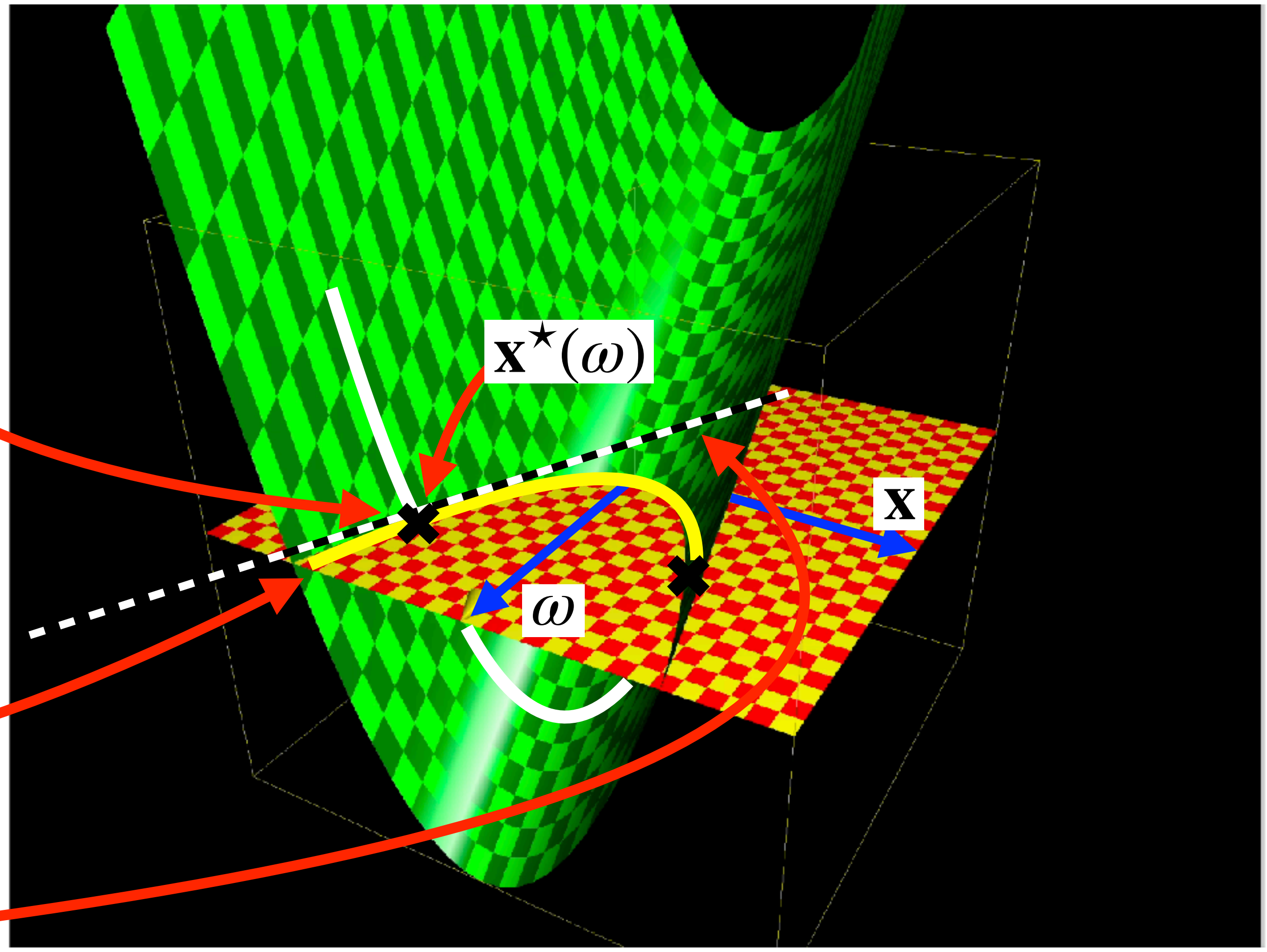
$$\frac{\partial \mathbf{x}^\star(\omega)}{\partial \omega} = ?$$

While changing ω , root function $\mathbf{x}^\star(\omega)$ follows the yellow curve

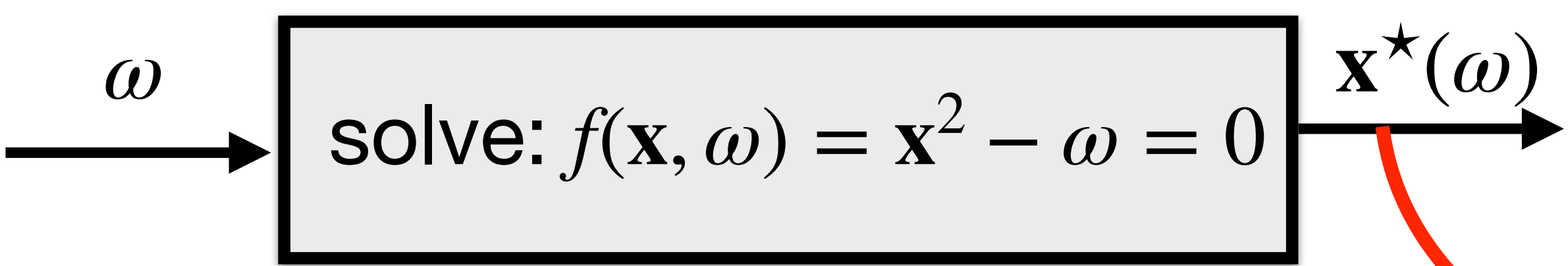
$$f(\mathbf{x}^\star(\omega), \omega) = \mathbf{x}^\star(\omega)^2 - \omega = 0$$

$$\frac{\partial f(\mathbf{x}^\star(\omega), \omega)}{\partial \omega} = 0$$

Yellow space is locally defined as the direction in ω -domain which locally preserves zero value of $f(\mathbf{x}^\star(\omega), \omega)$ function



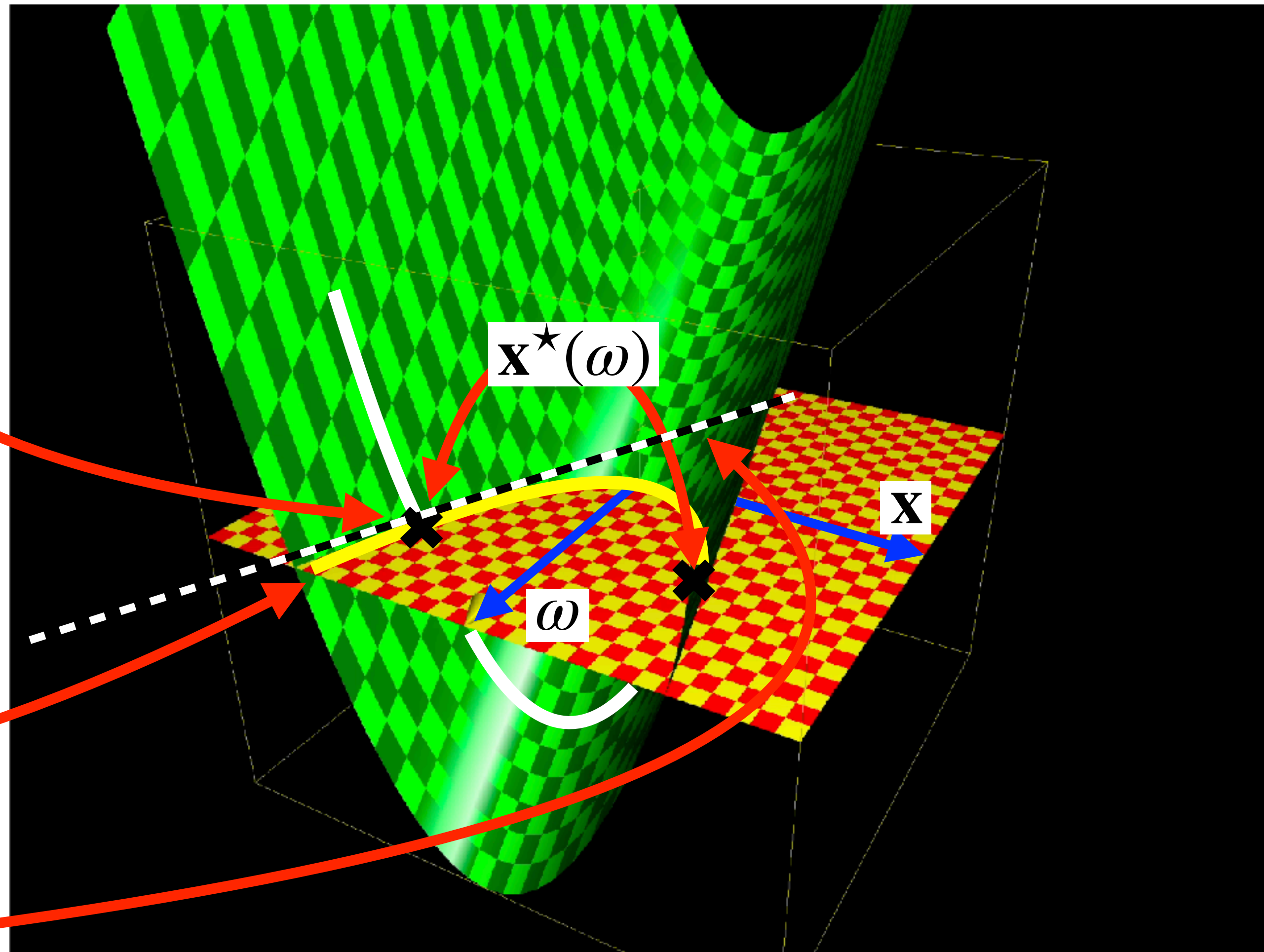
Example:



$$\frac{\partial \mathbf{x}^\star(\omega)}{\partial \omega} = ?$$

While changing ω , root function $\mathbf{x}^\star(\omega)$ follows the yellow curve

$$f(\mathbf{x}^\star(\omega), \omega) = \mathbf{x}^\star(\omega)^2 - \omega = 0$$



$$\frac{\partial f(\mathbf{x}^\star(\omega), \omega)}{\partial \omega} = 0$$

Yellow space is locally defined as the direction in ω -domain which locally preserves zero value of $f(\mathbf{x}^\star(\omega), \omega)$ function

$$2\mathbf{x}^\star(\omega) \cdot \frac{\partial \mathbf{x}^\star(\omega)}{\partial \omega} - 1 = 0 \Rightarrow \boxed{\frac{\partial \mathbf{x}^\star(\omega)}{\partial \omega} = \frac{1}{2\mathbf{x}^\star(\omega)}}$$

Can you compute gradient ???
How is it for $\omega = 0$???
How is it for other root ???

Root finder

$$\frac{\partial f(\mathbf{x}^*(\omega), \omega)}{\partial \omega} = 0$$

ω is allowed to **change only in directions** which **does not change** the value of $f(\mathbf{x}^*(\omega), \omega)$ in order to stay within the solution manifold

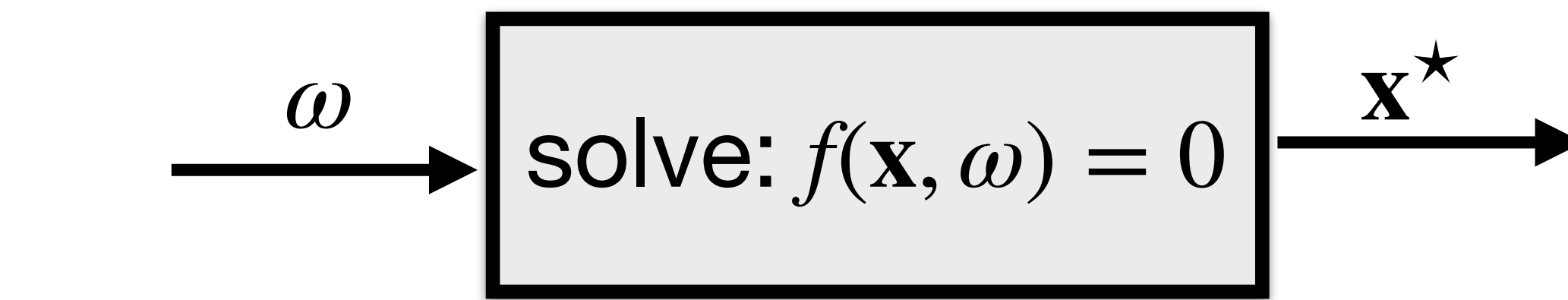
Introduce notation: $\frac{\partial f(a, b)}{\partial a} = \partial_1 f(a, b)$, $\frac{\partial f(a, b)}{\partial b} = \partial_2 f(a, b)$

Differentiation of compound function yields manifold eq. for searched $\frac{\partial \mathbf{x}^*(\omega)}{\partial \omega}$

$$\frac{\partial f(\mathbf{x}^*(\omega), \omega)}{\partial \omega} = \partial_1 f(\mathbf{x}^*(\omega), \omega) \frac{\partial \mathbf{x}^*(\omega)}{\partial \omega} + \partial_2 f(\mathbf{x}^*(\omega), \omega) = 0$$

$$\frac{\partial \mathbf{x}^*(\omega)}{\partial \omega} = - \left[\partial_1 f(\mathbf{x}^*(\omega), \omega) \right]^{-1} \partial_2 f(\mathbf{x}^*(\omega), \omega)$$

Root finder



$$\frac{\partial \mathbf{x}^\star(\omega)}{\partial \omega} = - \left[\partial_1 f(\mathbf{x}^\star(\omega), \omega) \right]^{-1} \partial_2 f(\mathbf{x}^\star(\omega), \omega)$$

Implicit gradient
(terms can be
obtained by autograd)

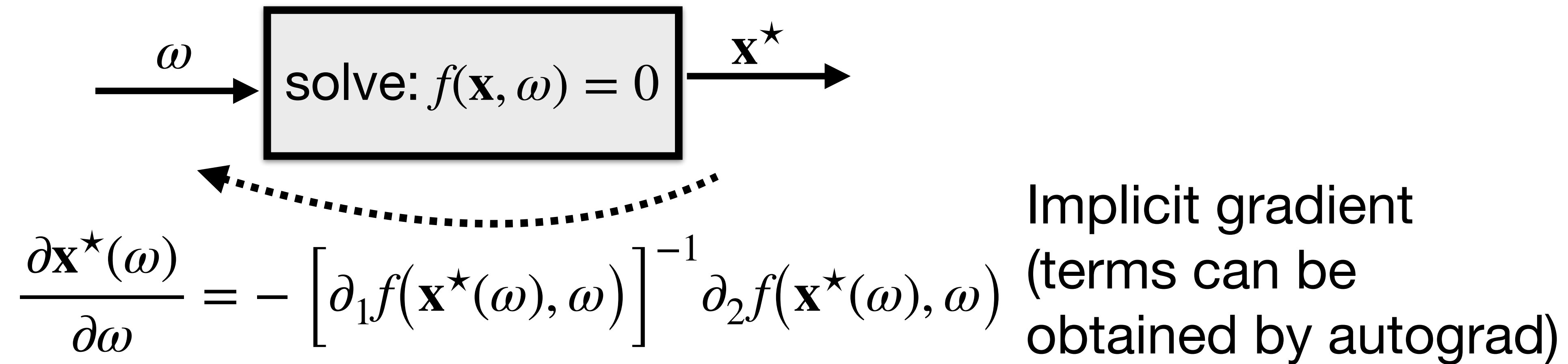
The implicit function theorem. Let $f : \mathbb{R}^p \times \mathbb{R}^n \rightarrow \mathbb{R}^n$ and $a_0 \in \mathbb{R}^p$, $z_0 \in \mathbb{R}^n$ be such that

1. $f(a_0, z_0) = 0$, and
2. f is continuously differentiable with non-singular Jacobian $\partial_1 f(a_0, z_0) \in \mathbb{R}^{n \times n}$.

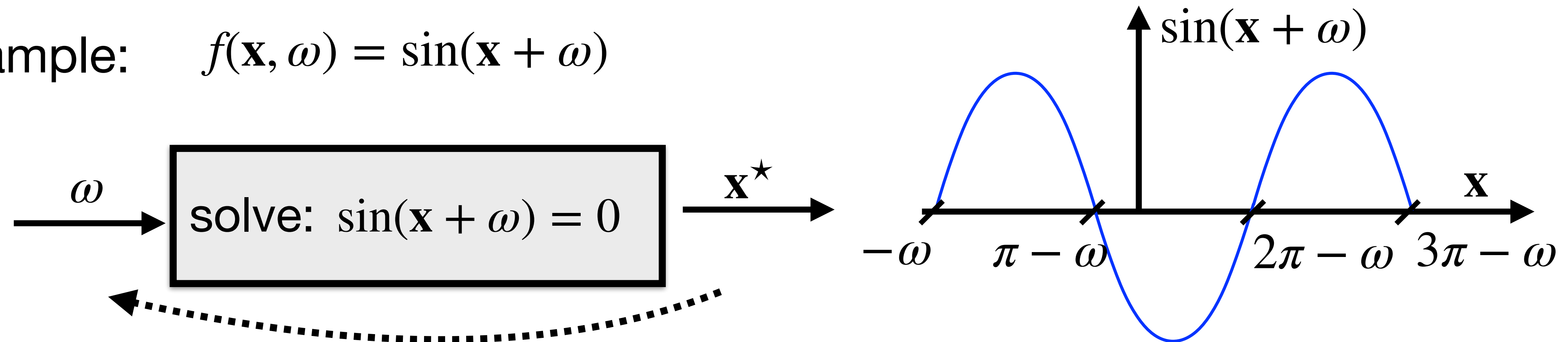
Then there exist open sets $S_{a_0} \subset \mathbb{R}^p$ and $S_{z_0} \subset \mathbb{R}^n$ containing a_0 and z_0 , respectively, and a unique continuous function $z^* : S_{a_0} \rightarrow S_{z_0}$ such that

1. $z_0 = z^*(a_0)$,
2. $f(a, z^*(a)) = 0 \quad \forall a \in S_{a_0}$, and
3. z^* is differentiable on S_{a_0} .

Root finder



Example: $f(\mathbf{x}, \omega) = \sin(\mathbf{x} + \omega)$

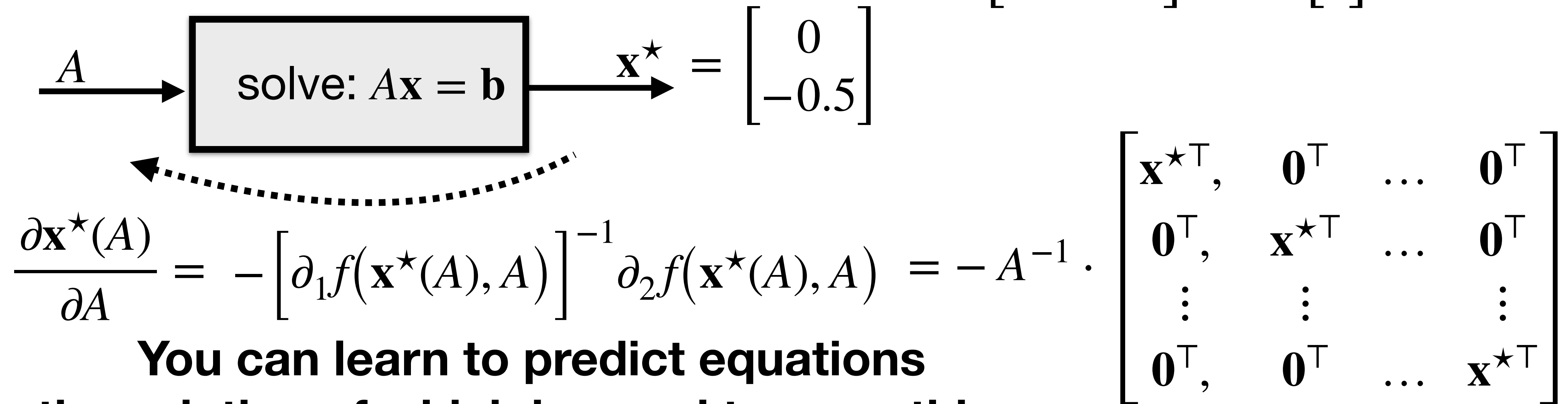


$$\frac{\partial \mathbf{x}^*(\omega)}{\partial \omega} = - \left[\partial_0 f(\mathbf{x}^*(\omega), \omega) \right]^{-1} \partial_1 f(\mathbf{x}^*(\omega), \omega) = - \frac{\cos(\mathbf{x}^* + \omega)}{\cos(\mathbf{x}^* + \omega)} = -1$$

Root finder (backprop through linear system / inverse matrix)

Example2: $f(\mathbf{x}, A) = A\mathbf{x} - \mathbf{b}$

$$A = \begin{bmatrix} 2 & -2 \\ -2 & 0 \end{bmatrix}, \mathbf{b} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$$



**You can learn to predict equations
the solution of which is equal to something**

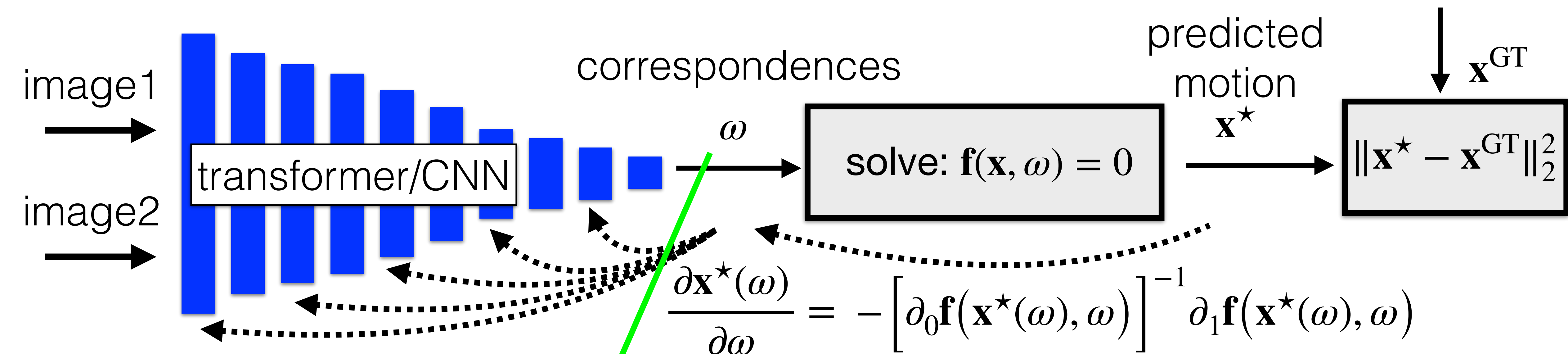
$$= - \begin{bmatrix} 2 & -2 \\ -2 & 0 \end{bmatrix}^{-1} \cdot \begin{bmatrix} x_1^* & x_2^* & 0 & 0 \\ 0 & 0 & x_1^* & x_2^* \end{bmatrix} = - \begin{bmatrix} 0 & 0 & 0 & 0.25 \\ 0 & 0.25 & 0 & 0.25 \end{bmatrix}$$

```
x_star = torch.linalg.inv(A) @ b
torch.autograd.grad(x_star[0], A)[0].reshape(4)
> tensor([-0.0000, -0.0000, -0.0000, -0.2500])
torch.autograd.grad(x_star[1], A)[0].reshape(4)
> tensor([-0.0000, -0.2500, -0.0000, -0.2500])
```

**Do not have to preserve
computational graph
of the matrix inverse
computation !!!**

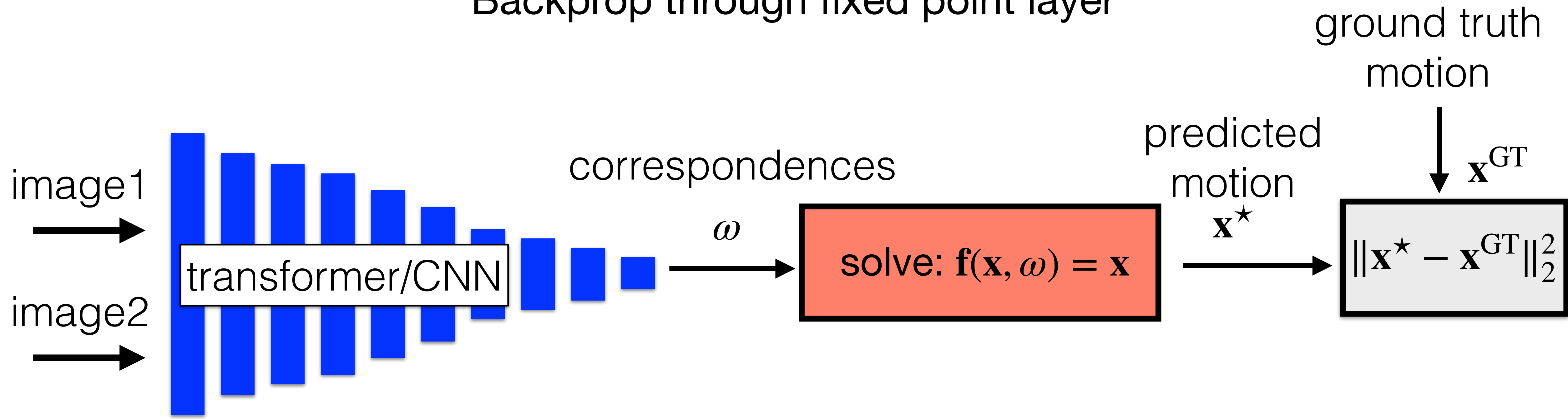
Root finder: backprop through set of equations typical application

ground truth
motion



$\mathbf{f}(\mathbf{x}, \omega) = 0$ is set of non-linear equations

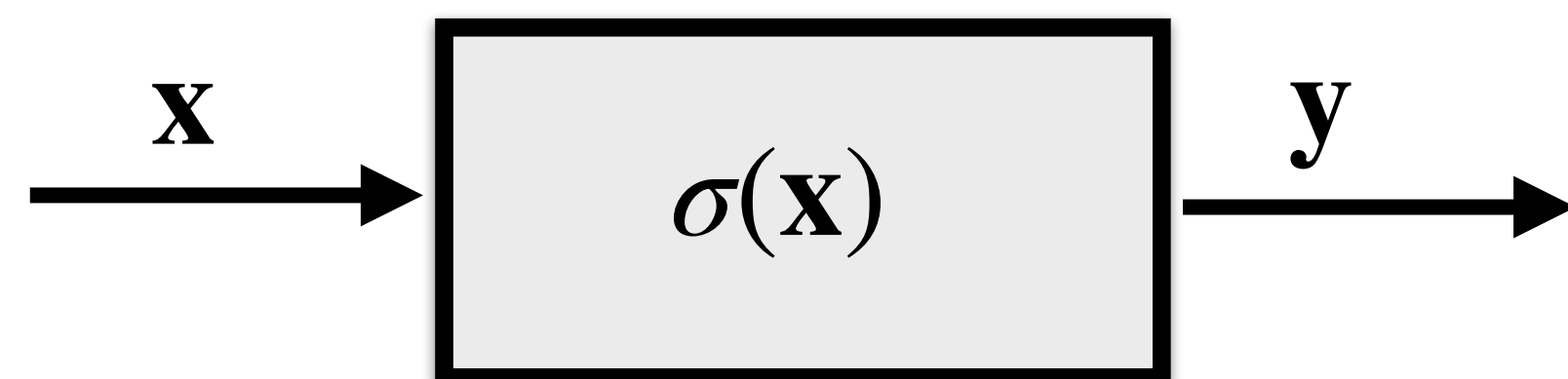
Backprop through fixed point layer



- Infinitely deep residual network layer
- Q-learning as a feed-forward pass
- Convergence assured for some function (contractivity or monotocinity)

Explicit layers

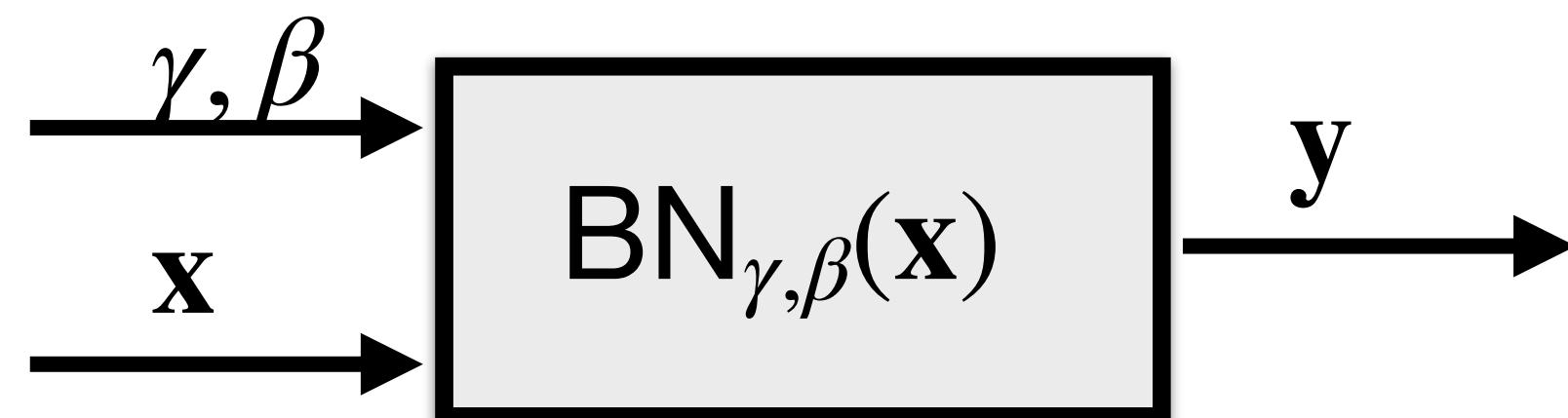
Sigmoid:



Convolution:

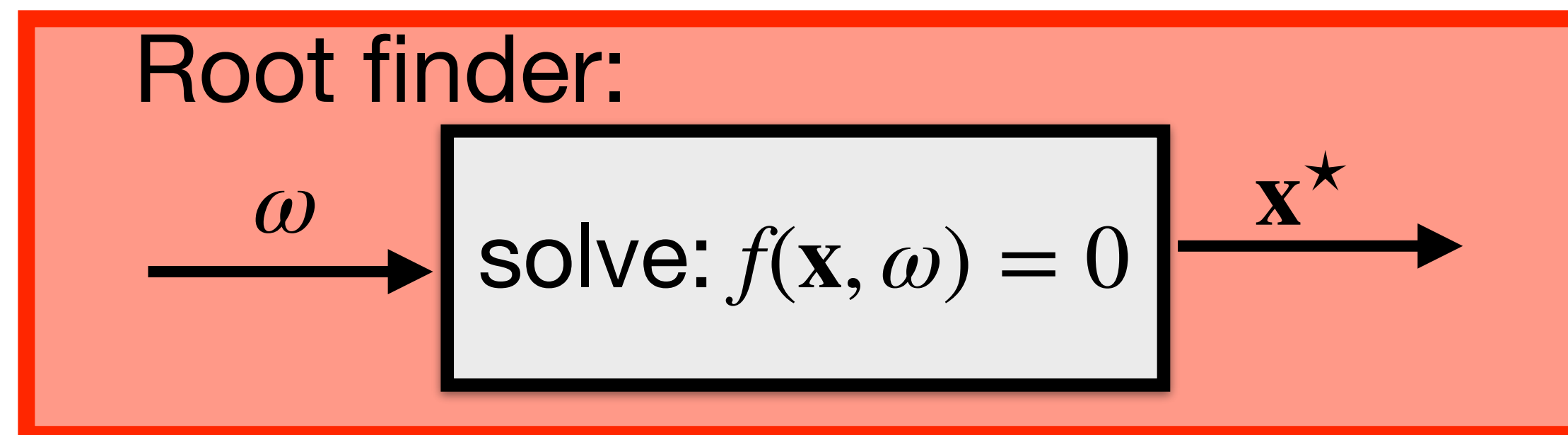


Batch-norm:

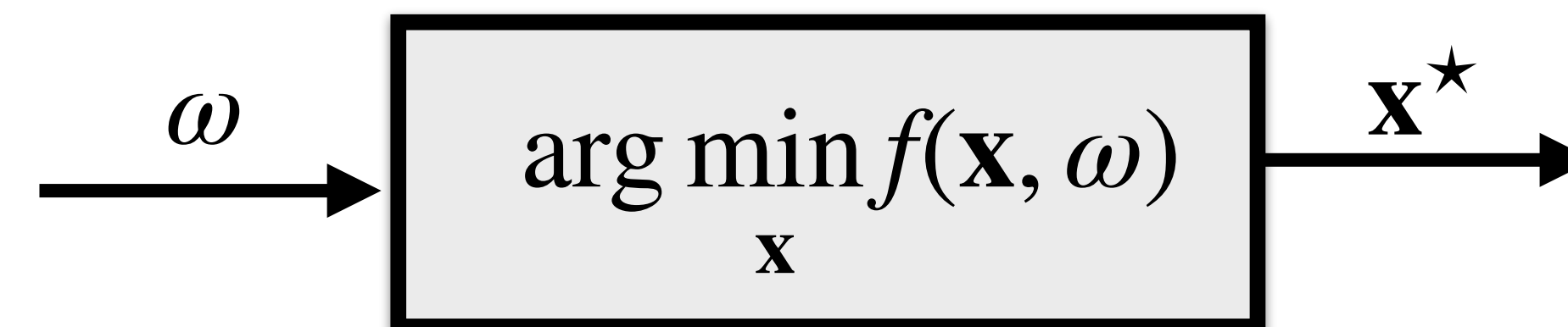


Implicit layers

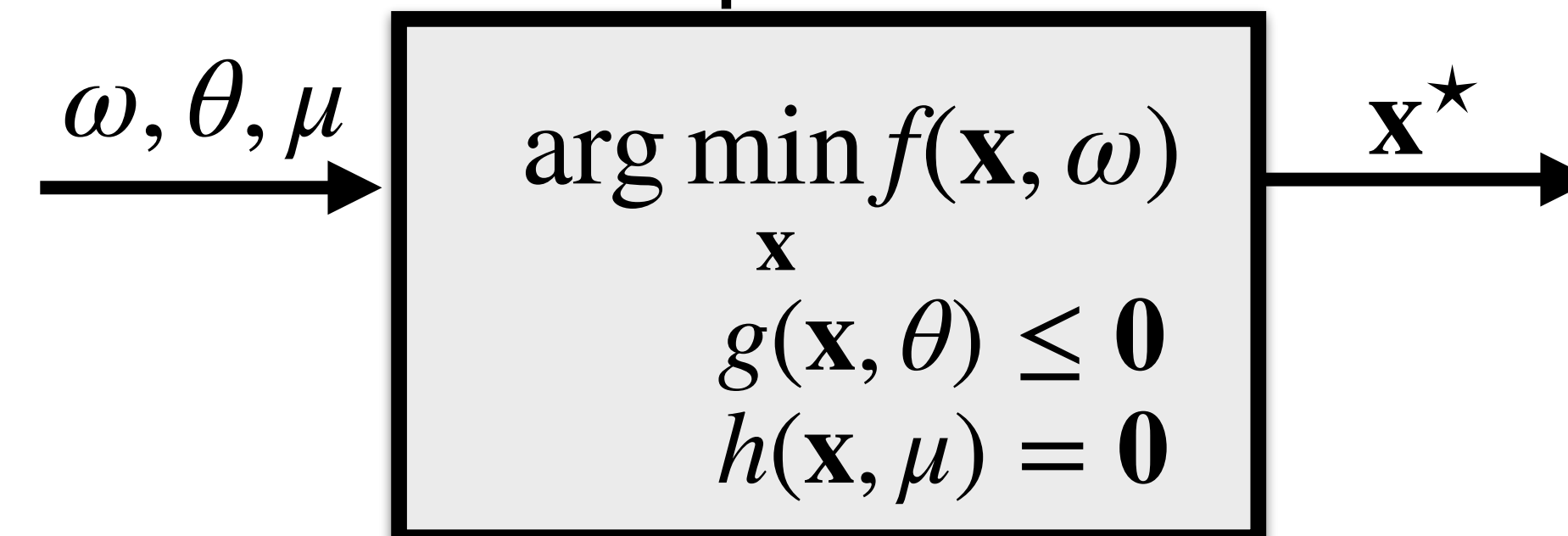
Root finder:



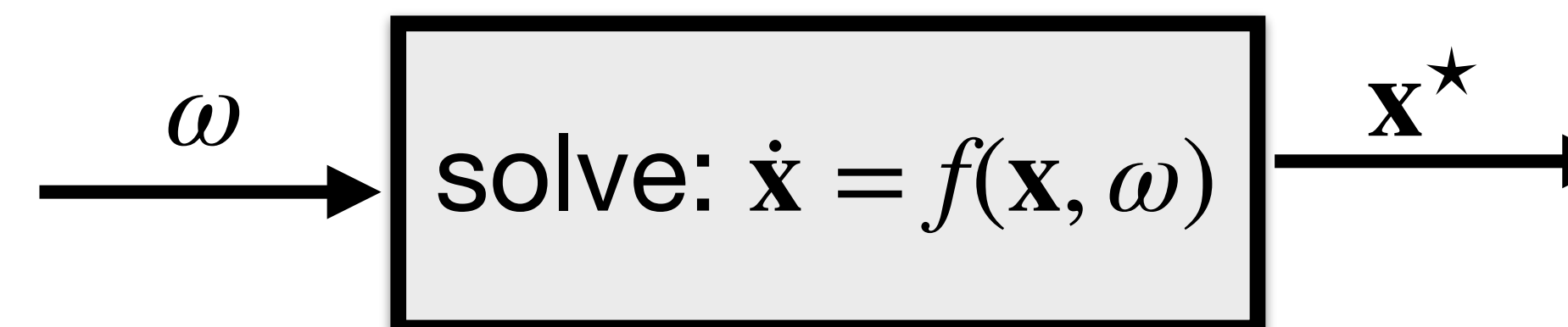
Unconstrained optimizer:



Constrained optimizer:



ODE solver:

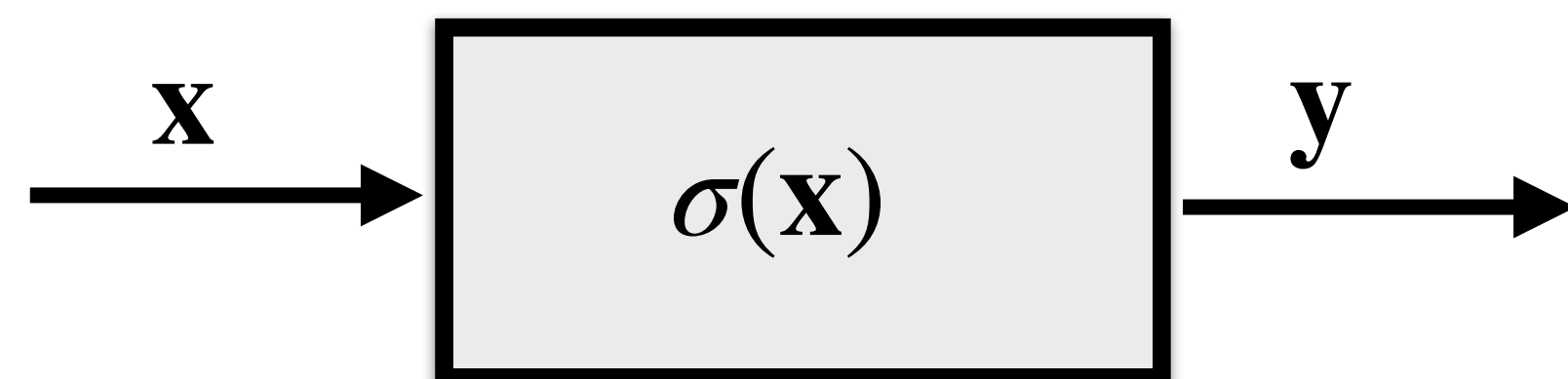


Explicit layers

The rest will not be
in the exam test

Implicit layers

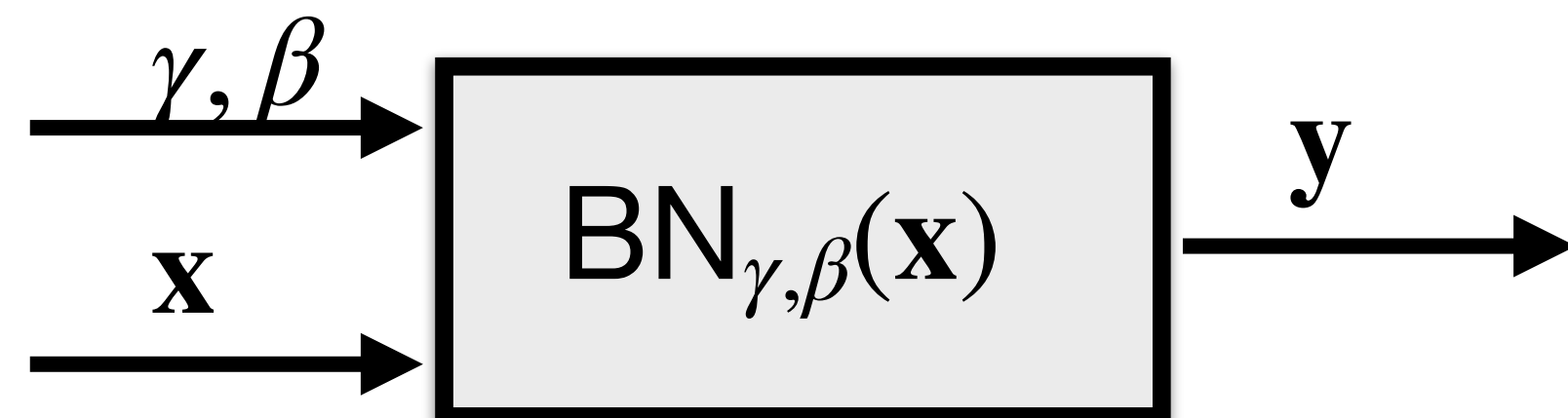
Sigmoid:



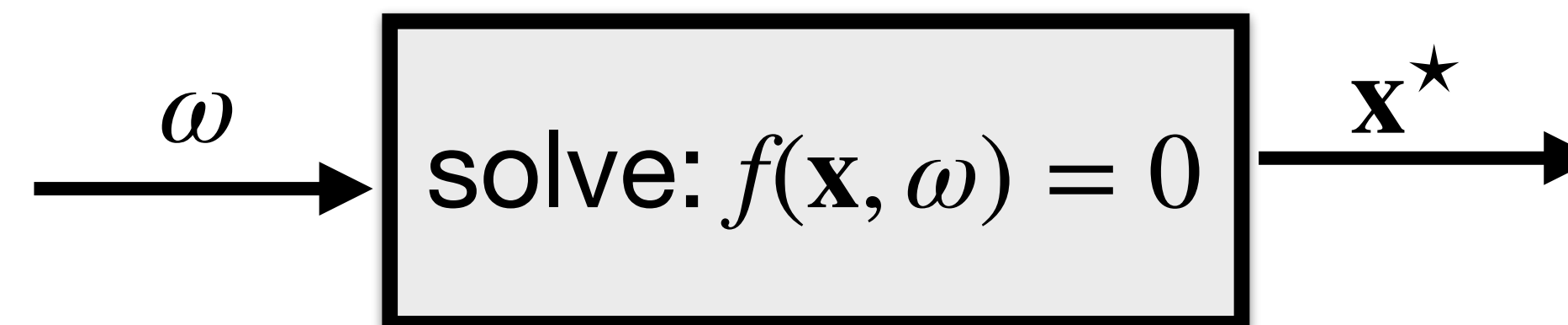
Convolution:



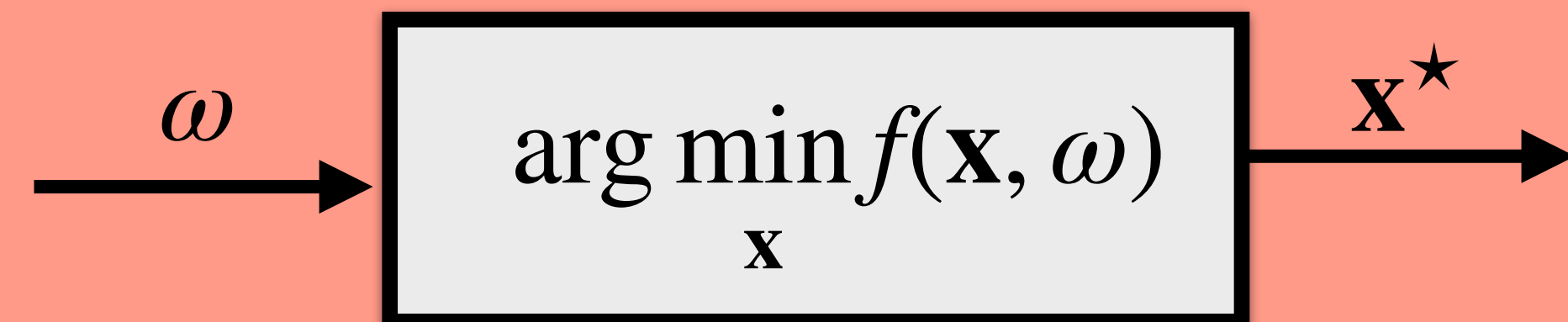
Batch-norm:



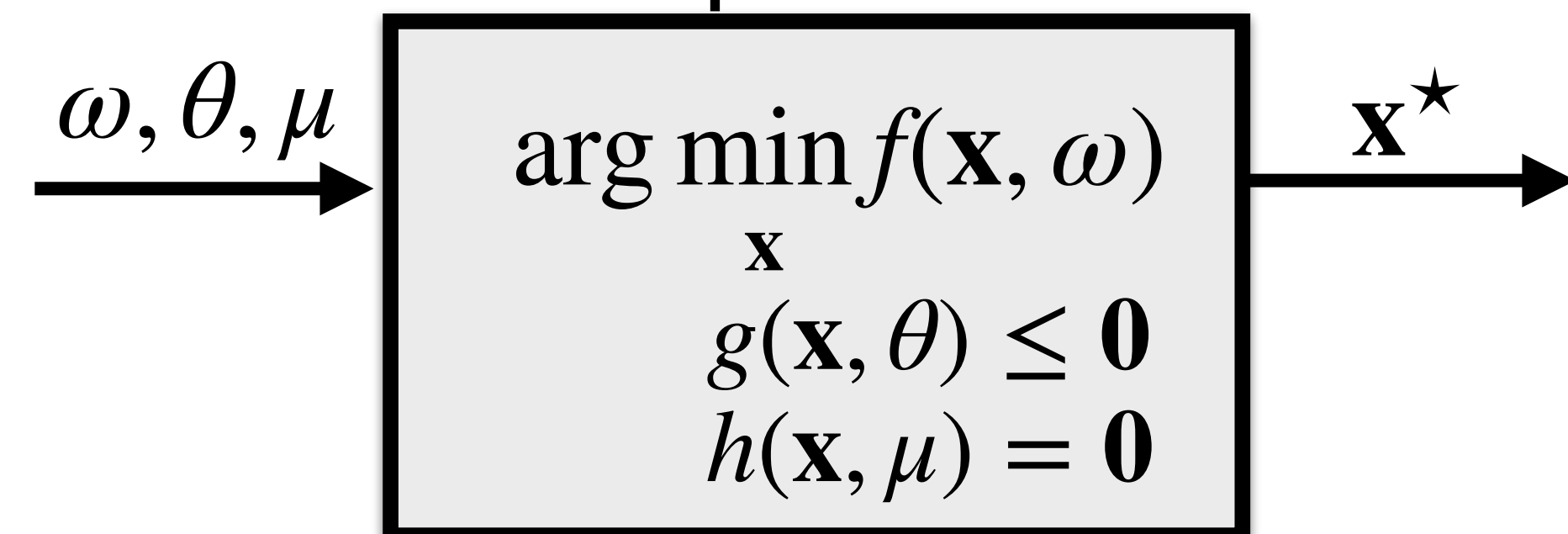
Root finder:



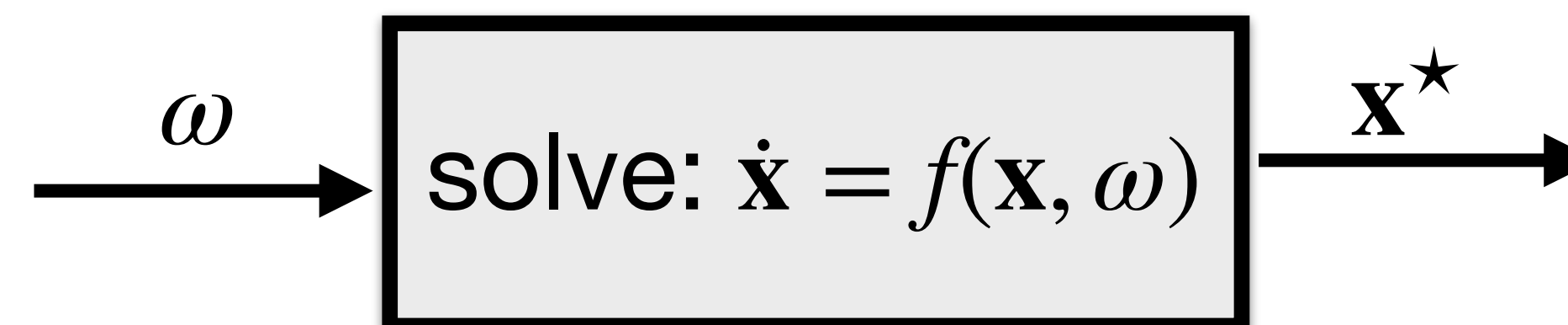
Unconstrained optimizer:



Constrained optimizer:



ODE solver:



Unconstrained optimizer

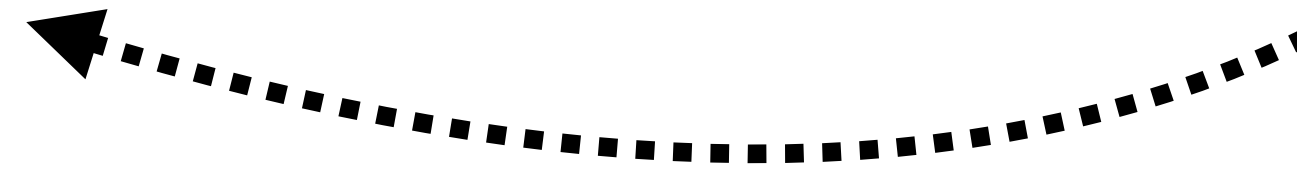
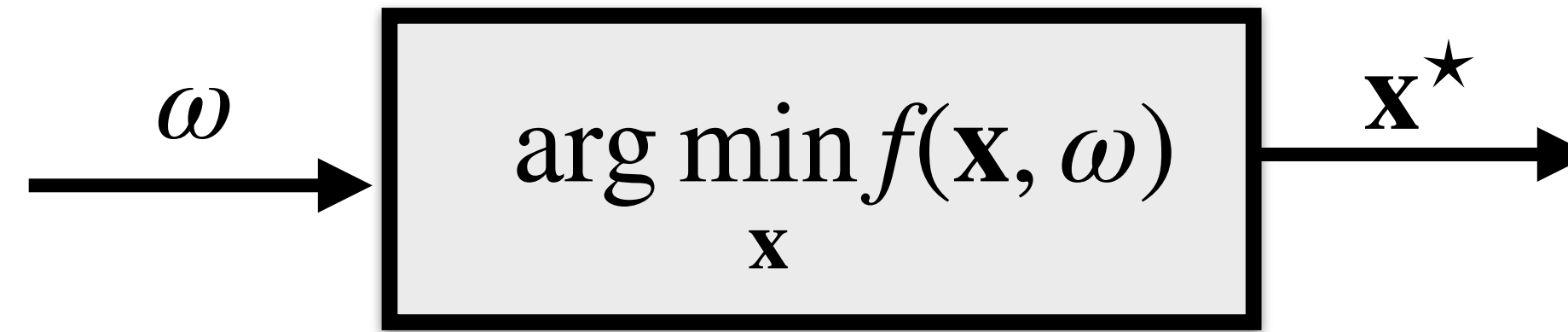
- Unconstrained local optimizer provides solution $\mathbf{x}^\star$ such that $\frac{\partial f(\mathbf{x}^\star, \omega)}{\partial \mathbf{x}^\star} = 0$
- Let's denote gradient $\frac{\partial f(\mathbf{x}^\star, \omega)}{\partial \mathbf{x}^\star} = g(\mathbf{x}^\star, \omega) = 0$
- This solution depends on parameters, we refer it $\mathbf{x}^\star(\omega) : \mathbb{R}^n \rightarrow \mathbb{R}^m$
- We treat unconstrained optimizer as **gradient root finder** => the rest is the same

$$\frac{\partial g(\mathbf{x}^\star(\omega), \omega)}{\partial \omega} = 0$$

ω is allowed to change only in directions which does not change the gradient $g(\mathbf{x}^\star(\omega), \omega)$ in order to stay within the solution manifold

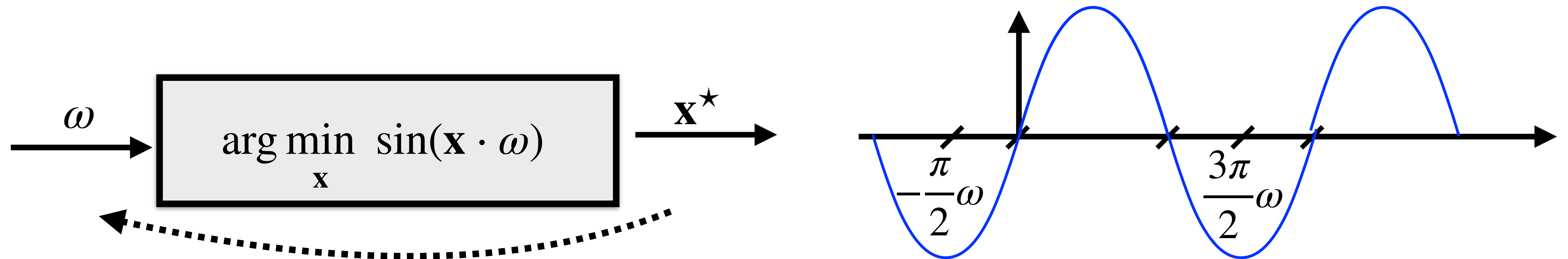
$$\begin{aligned} \frac{\partial g(\mathbf{x}^\star(\omega), \omega)}{\partial \omega} &= \partial_1 g(\mathbf{x}^\star(\omega), \omega) \frac{\partial \mathbf{x}^\star(\omega)}{\partial \omega} + \partial_2 g(\mathbf{x}^\star(\omega), \omega) = 0 \\ \frac{\partial \mathbf{x}^\star(\omega)}{\partial \omega} &= - \left[\partial_1 g(\mathbf{x}^\star(\omega), \omega) \right]^{-1} \partial_2 g(\mathbf{x}^\star(\omega), \omega) \end{aligned}$$

Unconstrained optimizer



$$\frac{\partial \mathbf{x}^*(\omega)}{\partial \omega} = - \left[\partial_1 g(\mathbf{x}^*(\omega), \omega) \right]^{-1} \partial_2 g(\mathbf{x}^*(\omega), \omega) \quad \text{where} \quad g(\mathbf{x}^*, \omega) = \frac{\partial f(\mathbf{x}^*, \omega)}{\partial \mathbf{x}^*}$$

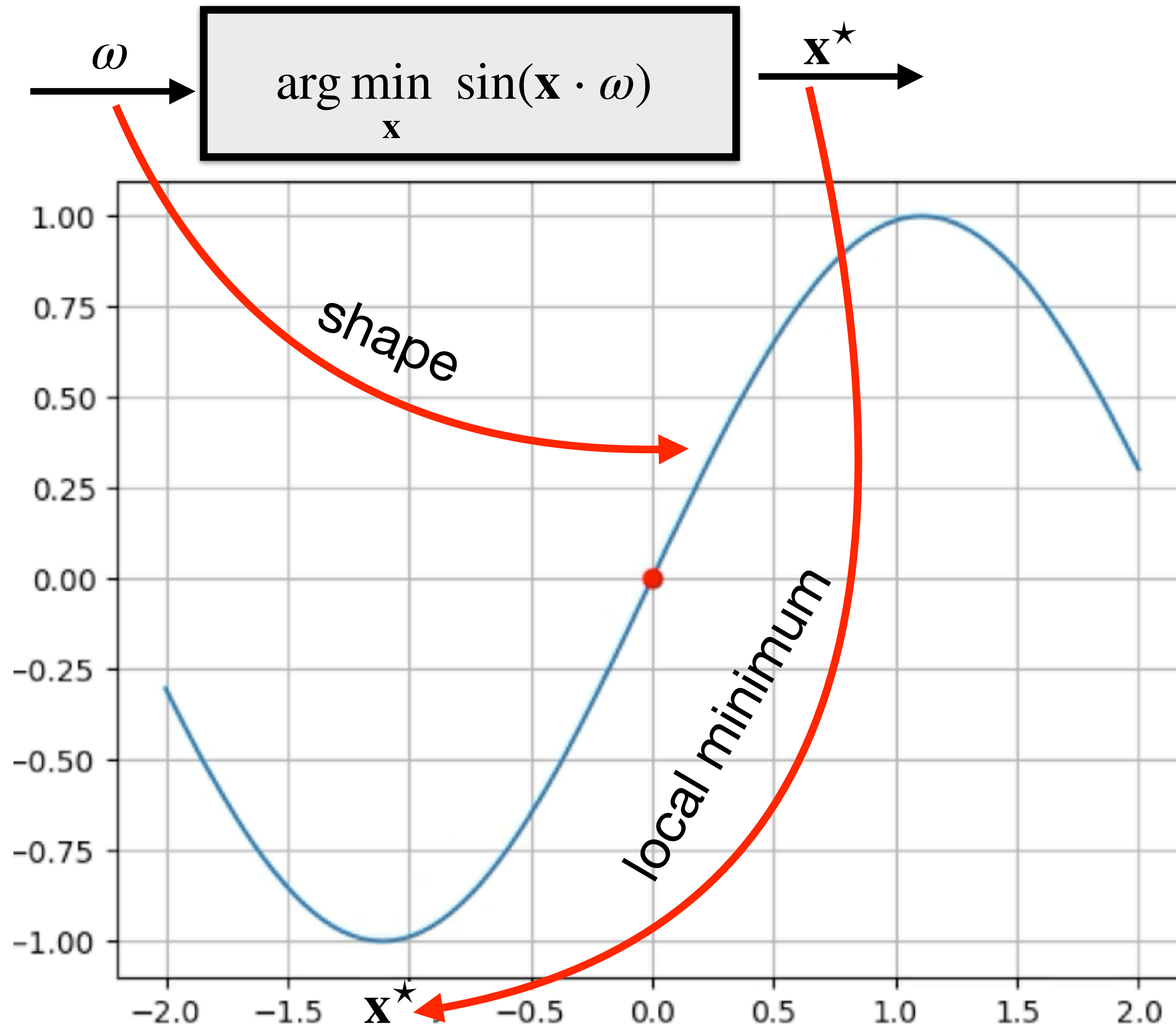
Example: $f(\mathbf{x}, \omega) = \sin(\mathbf{x} \cdot \omega) \quad \Rightarrow \quad g(\mathbf{x}, \omega) = \omega \cdot \cos(\mathbf{x} \cdot \omega)$



$$\frac{\partial \mathbf{x}^*(\omega)}{\partial \omega} = - \left[\underbrace{-\omega^2 \sin(\omega \mathbf{x})}_{\frac{\partial}{\partial \mathbf{x}}} \right]^{-1} \cdot \left(\underbrace{-\omega \mathbf{x} \sin(\omega \mathbf{x}) + \cos(\omega \mathbf{x})}_{\frac{\partial}{\partial \omega}} \right)$$

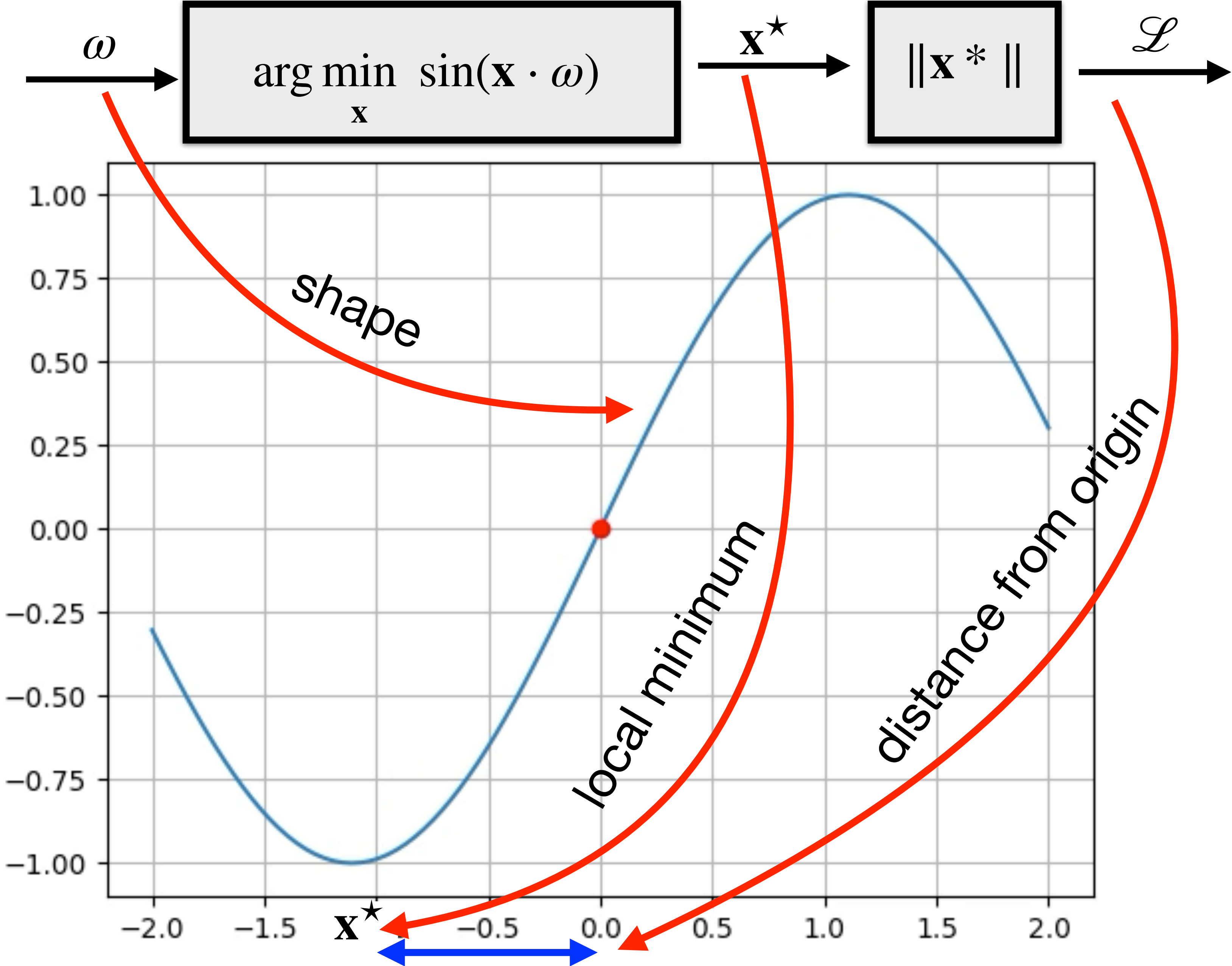
Unconstrained optimizer

Example:



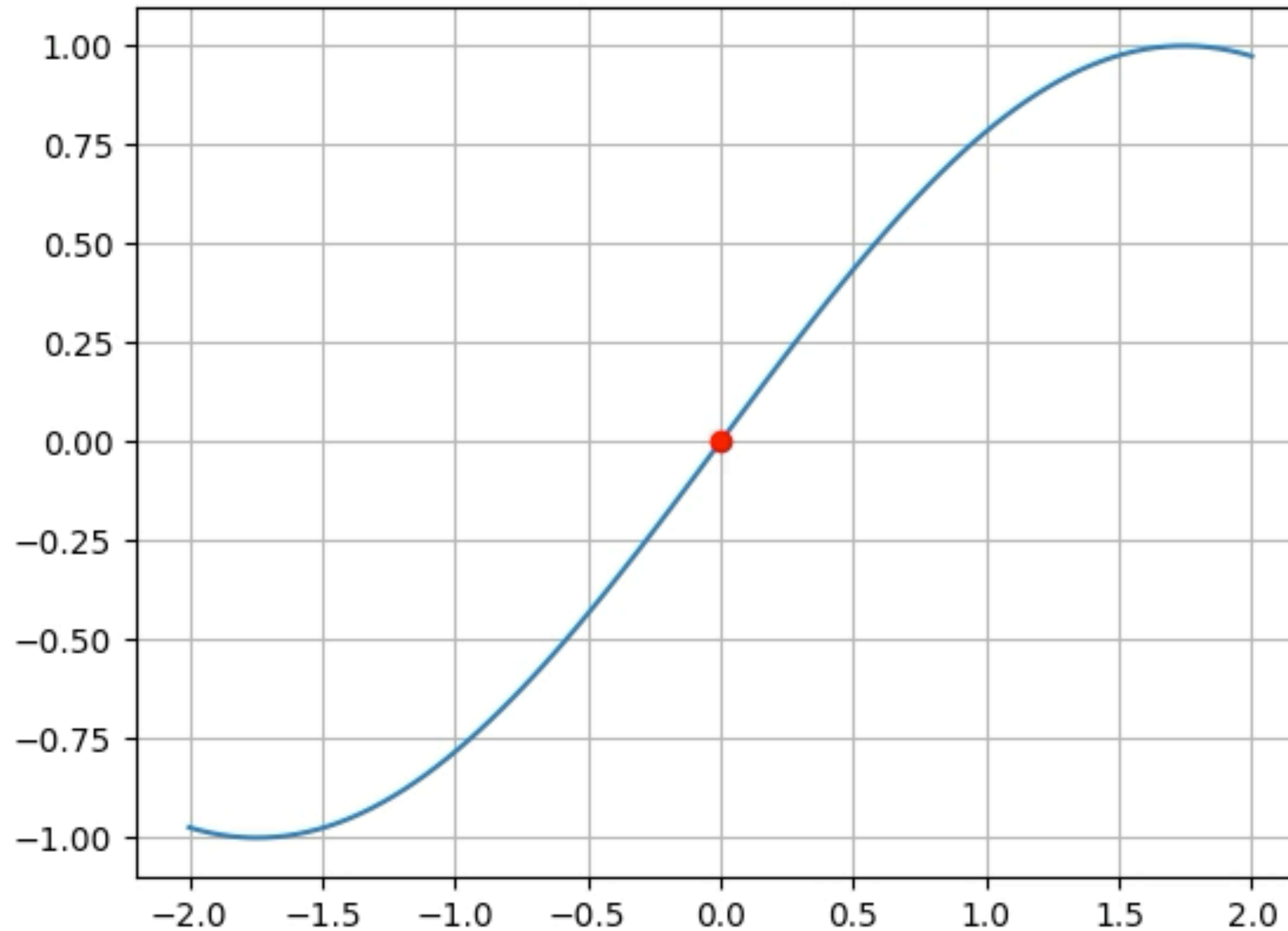
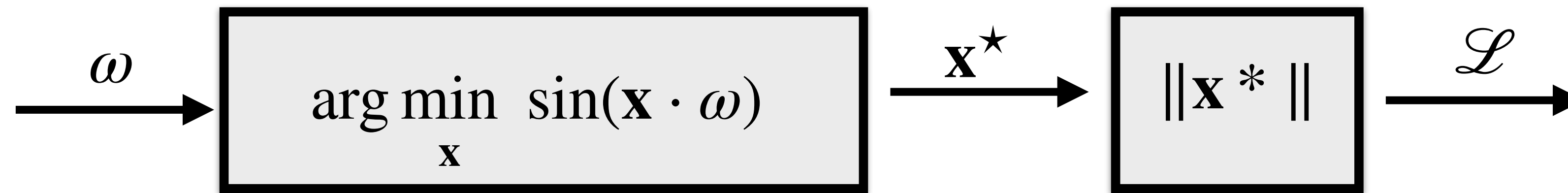
Unconstrained optimizer

Example:



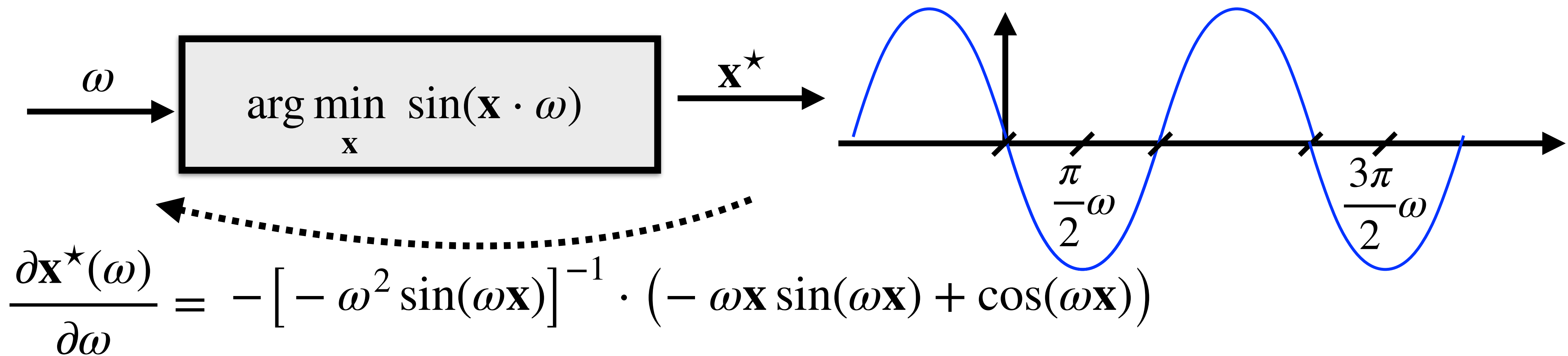
Unconstrained optimizer

Example:



Unconstrained optimizer

Example:



```
def f(x, omega):  
    return torch.sin(x*omega)  
  
omega = torch.tensor(2.0, requires_grad=True)  
x = torch.tensor(0.0, requires_grad=True)  
for j in range(25):  
    x = x - 0.1*torch.autograd.grad(f(x, omega), x, retain_graph=True, create_graph=True)[0]
```

explicit gradient

```
torch.autograd.grad(x, omega)[0]
```

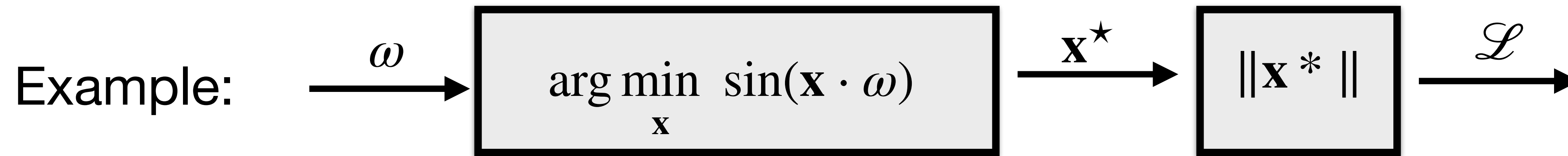
```
> tensor(0.3927)
```

implicit gradient

```
-(-torch.sin(omega*x)*omega*x + torch.cos(omega*x))/(-torch.sin(omega*x)*(omega**2))
```

```
> tensor(0.3927)
```

Unconstrained optimizer



What happens to **explicit gradient** if number of **iterations** are **decreased**?

```
for j in range(25):
    x = x - 0.1*torch.autograd.grad(f(x, omega), x, retain_graph=True, create_graph=True)[0]

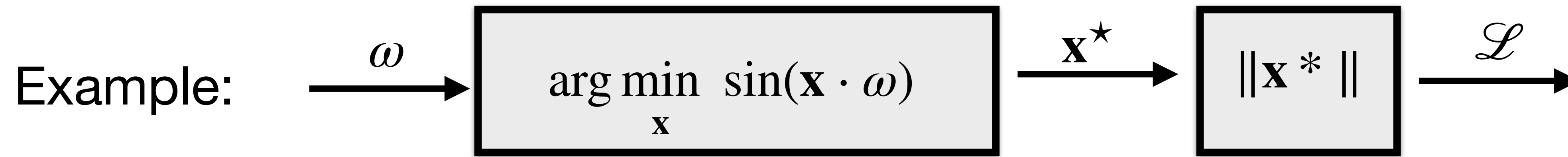
torch.autograd.grad(x.norm(), omega)[0]
> tensor(-0.3927)
```

```
for j in range(5):
    x = x - 0.1*torch.autograd.grad(f(x, omega), x, retain_graph=True, create_graph=True)[0]

torch.autograd.grad(x.norm(), omega)[0]
> tensor(0.4527)
```

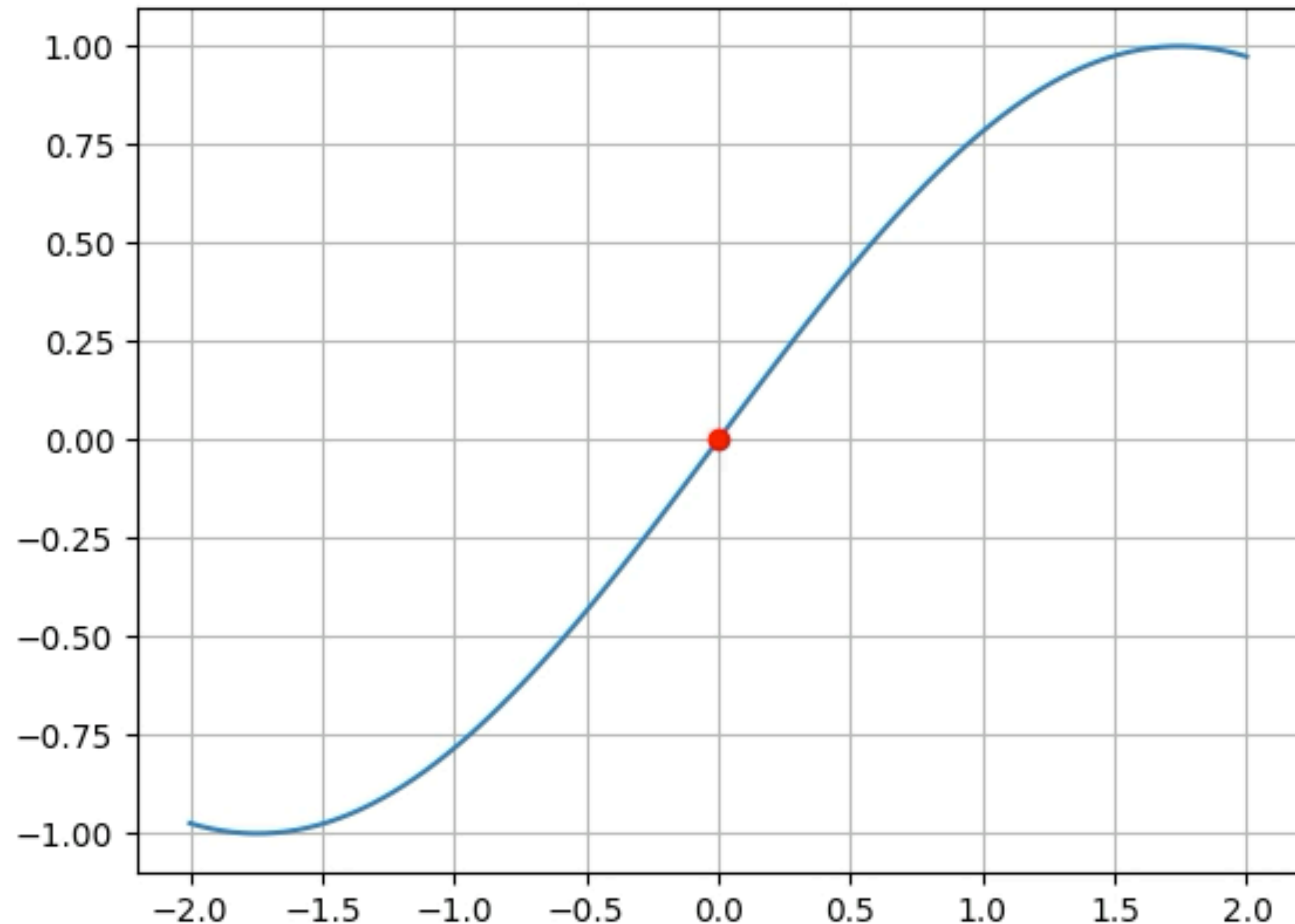
Is it correct?

Unconstrained optimizer

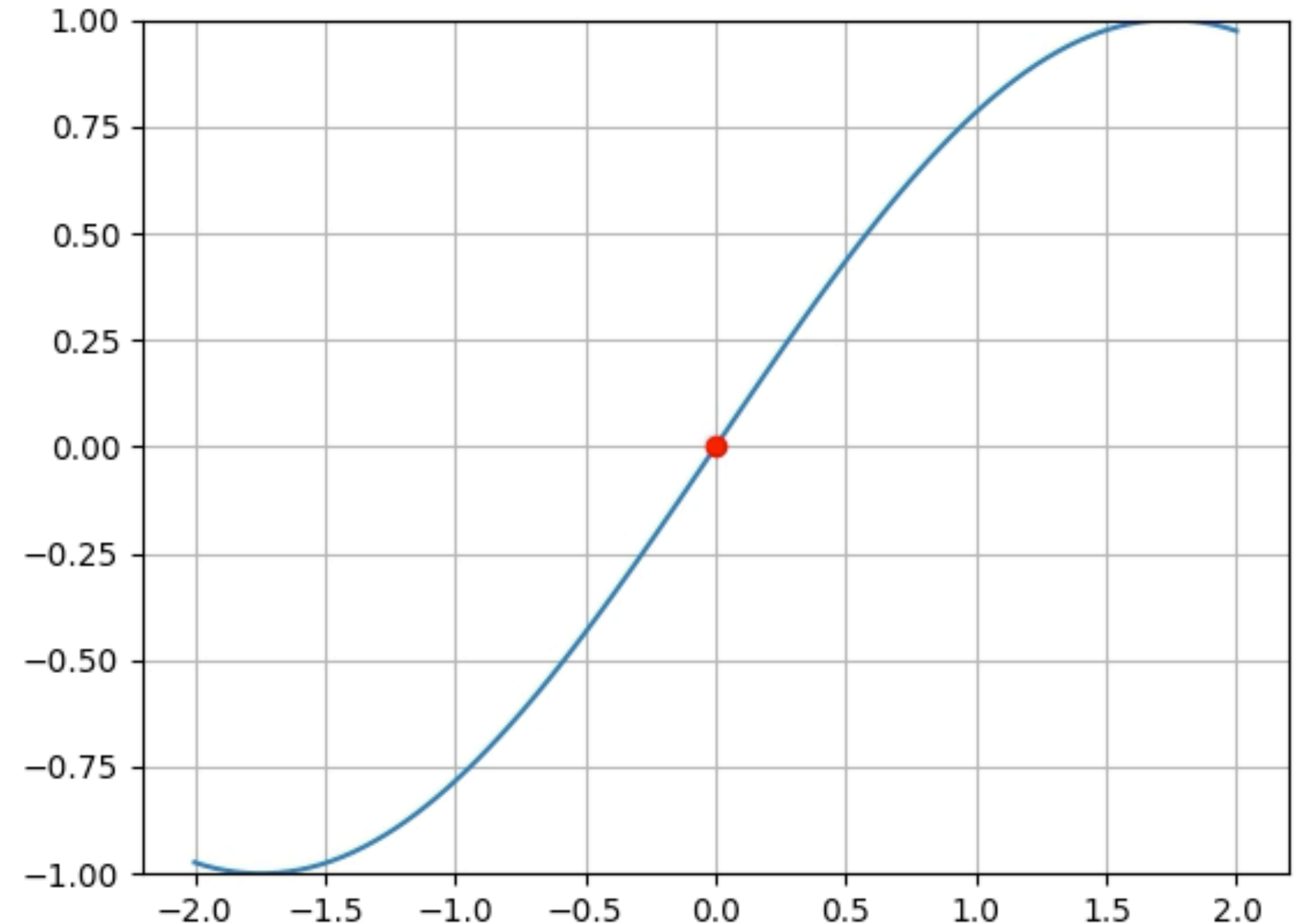


What happens to explicit-gradient if number of iterations are decreased?

```
for j in range(25):
```

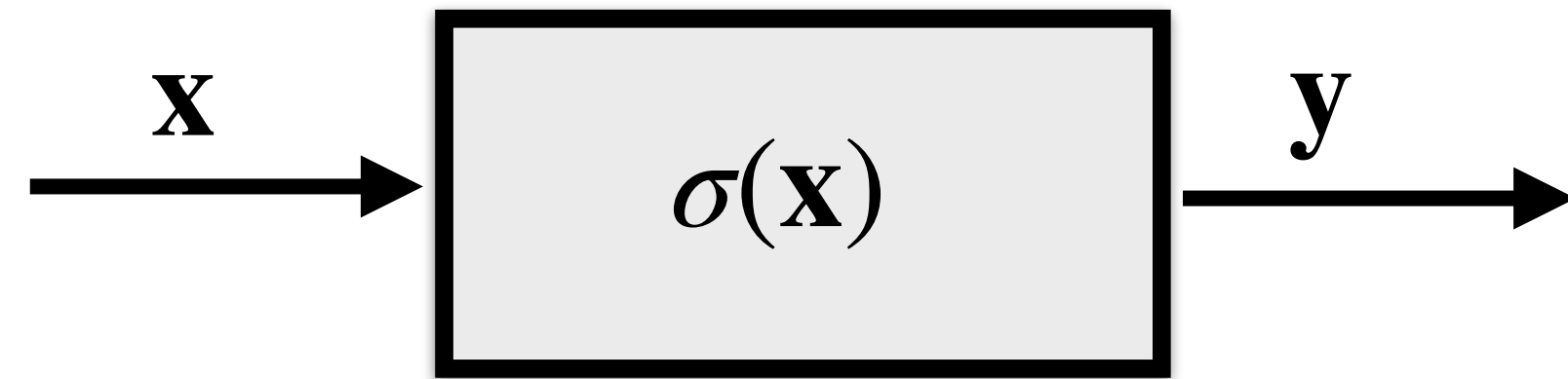


```
for j in range(5):
```



Explicit layers

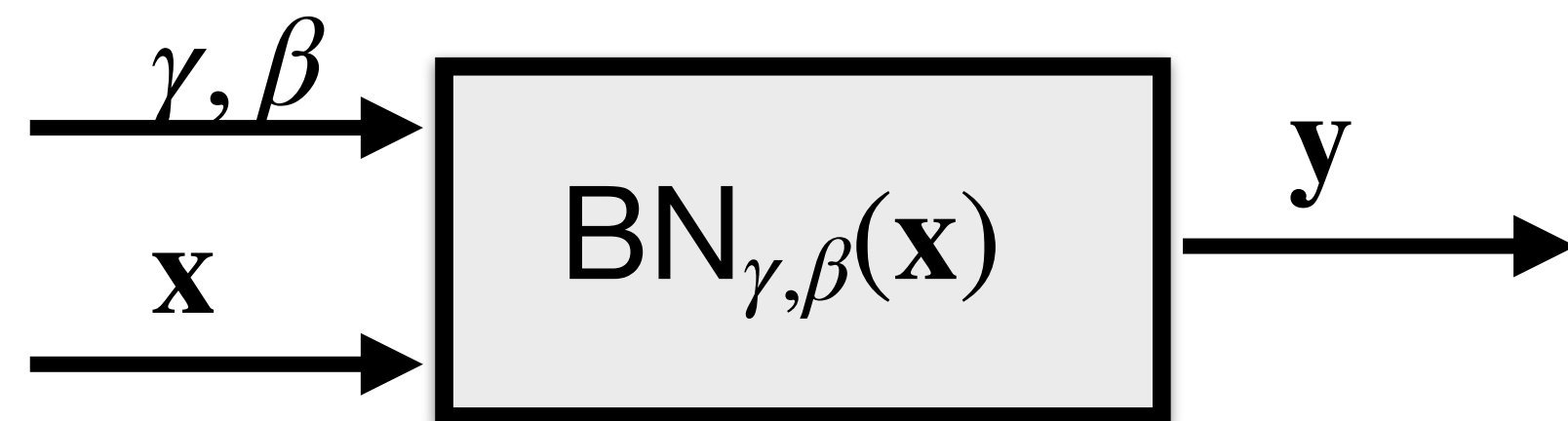
Sigmoid:



Convolution:

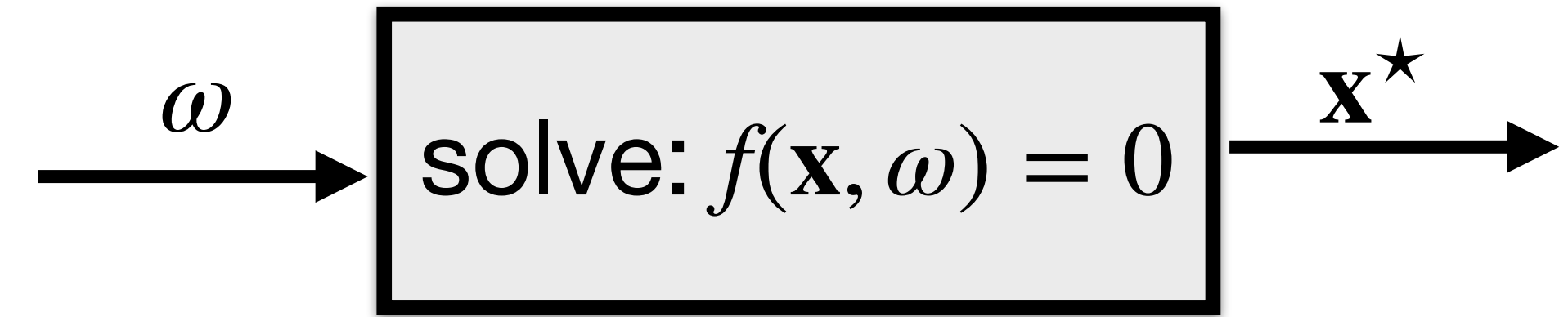


Batch-norm:

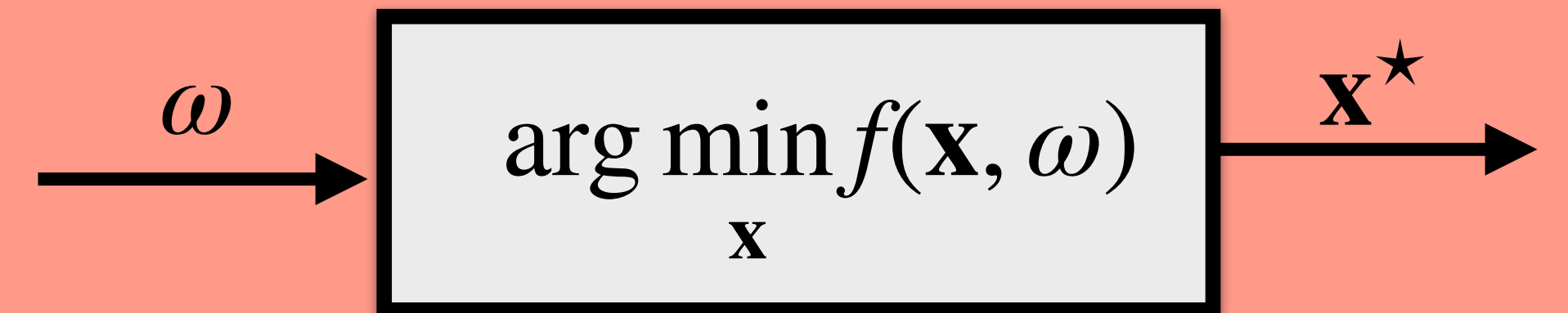


Implicit layers

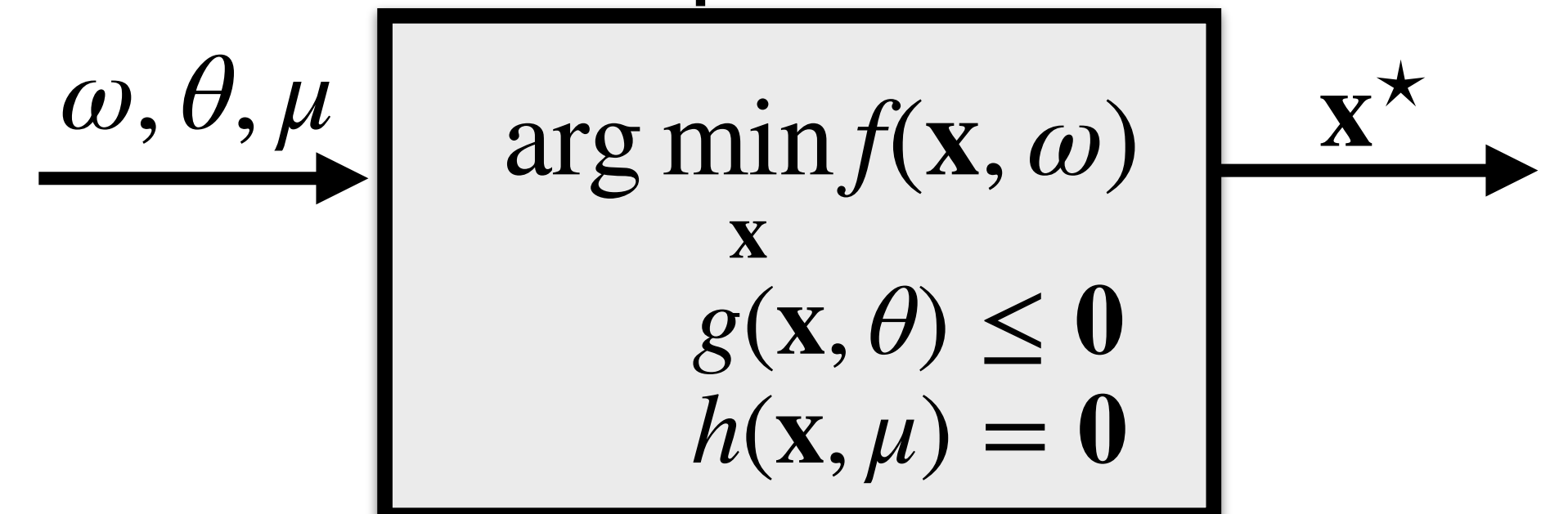
Root finder:



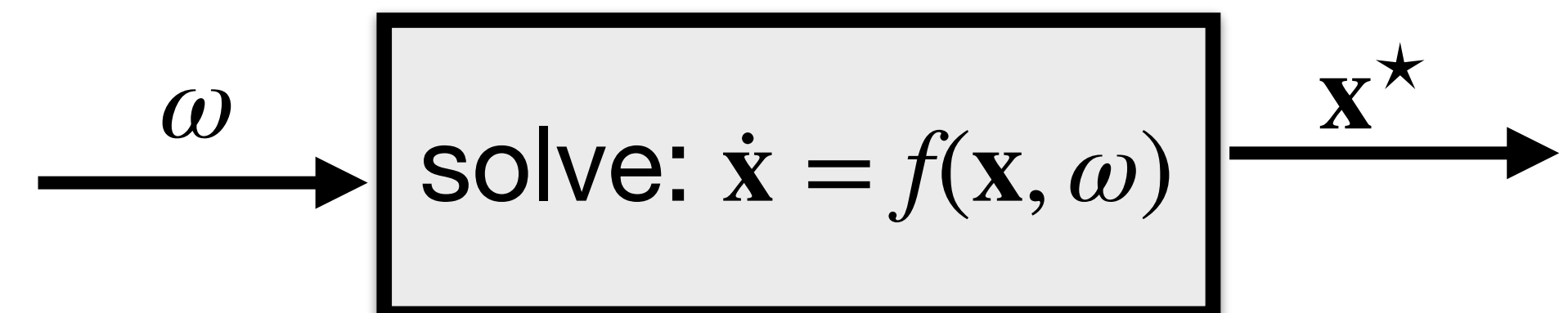
Unconstrained optimizer:



Constrained optimizer:

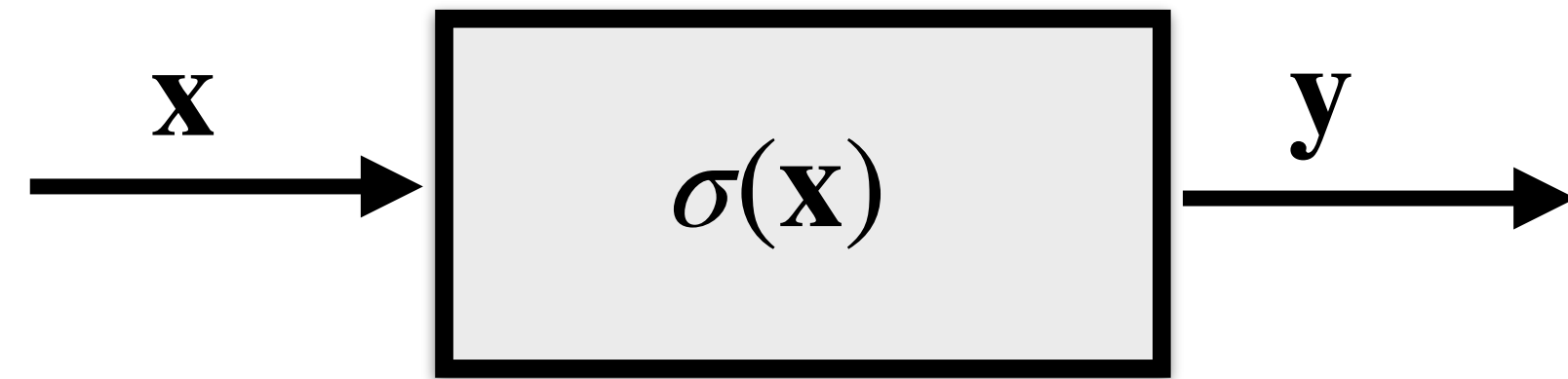


ODE solver:



Explicit layers

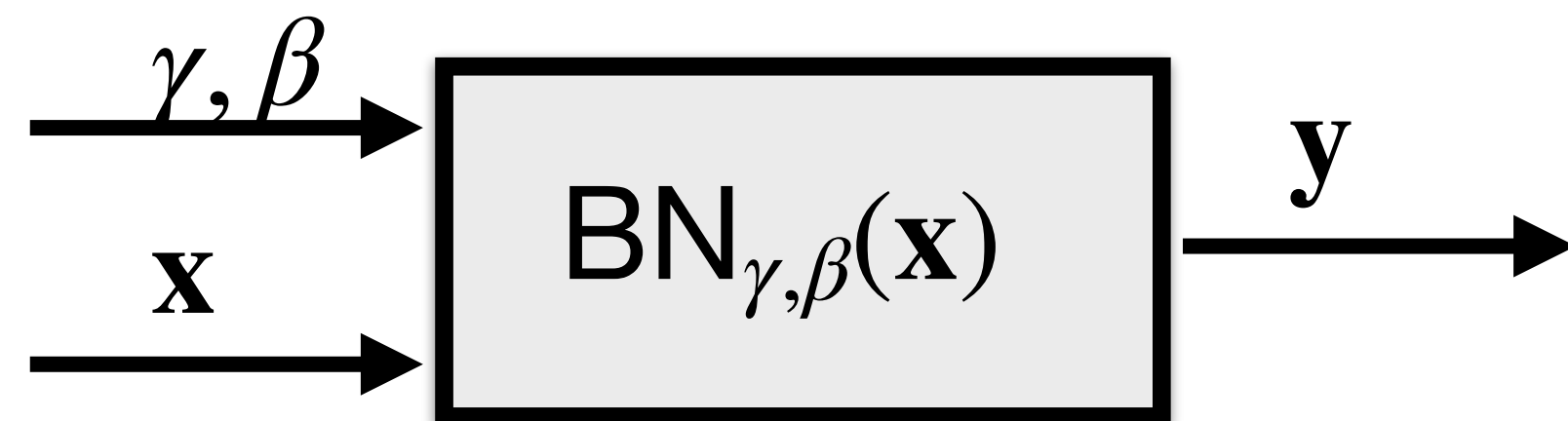
Sigmoid:



Convolution:

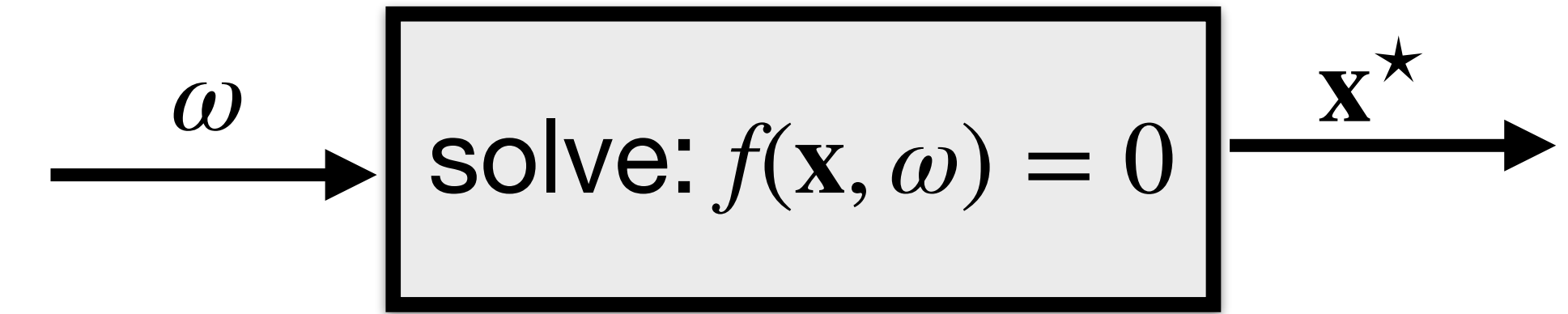


Batch-norm:

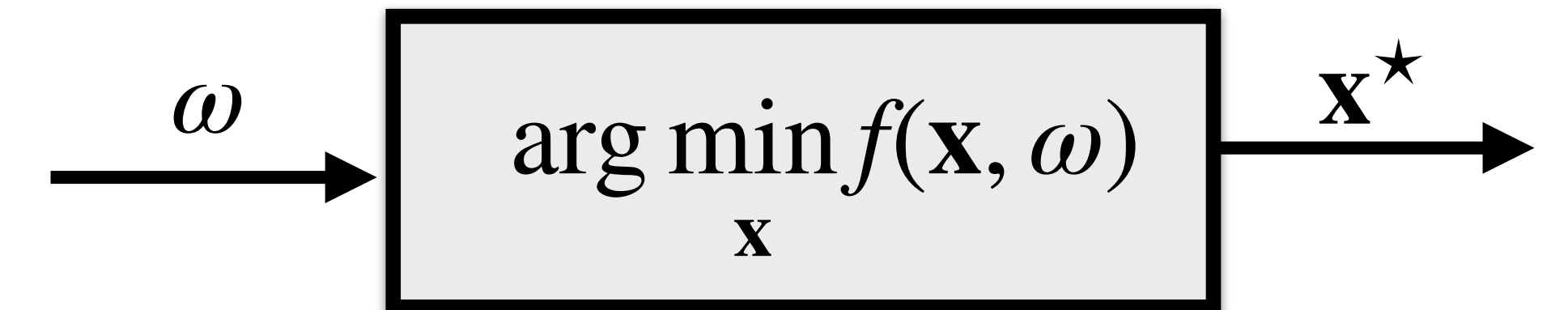


Implicit layers

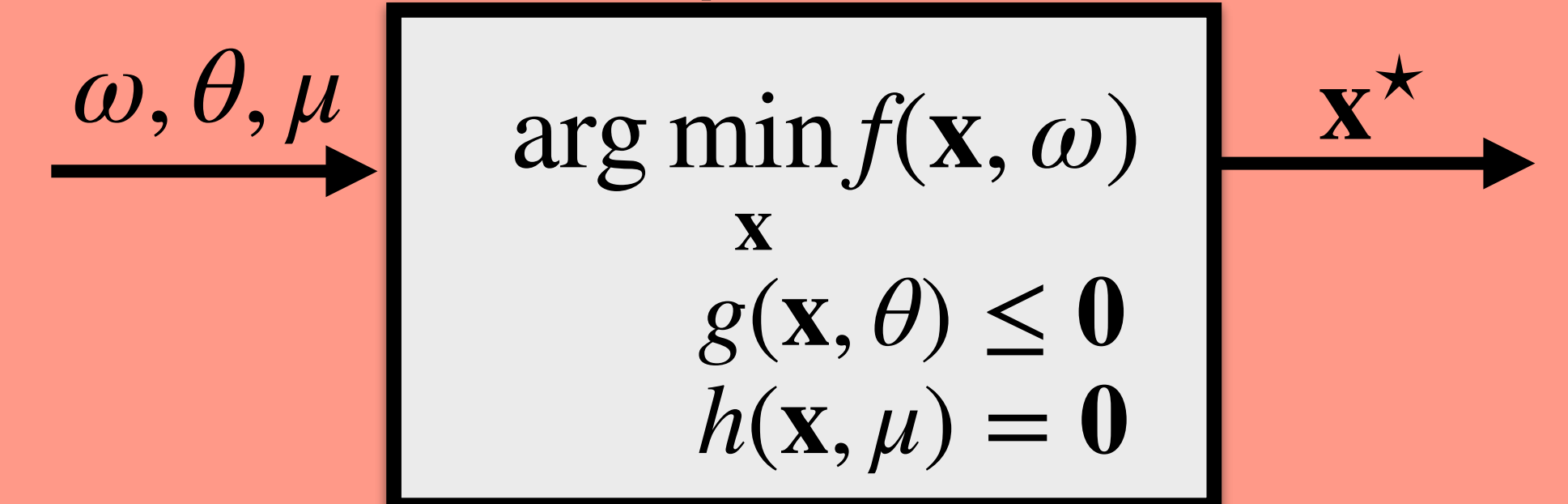
Root finder:



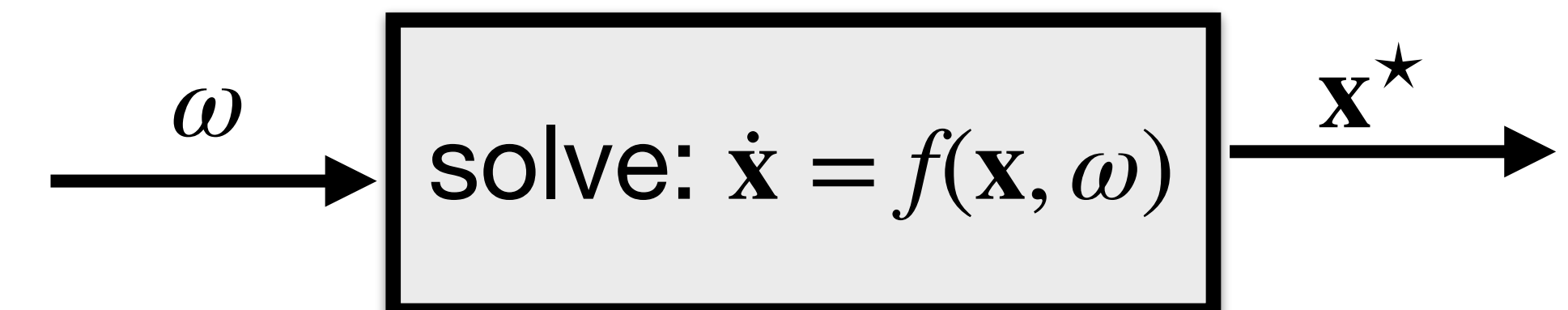
Unconstrained optimizer:



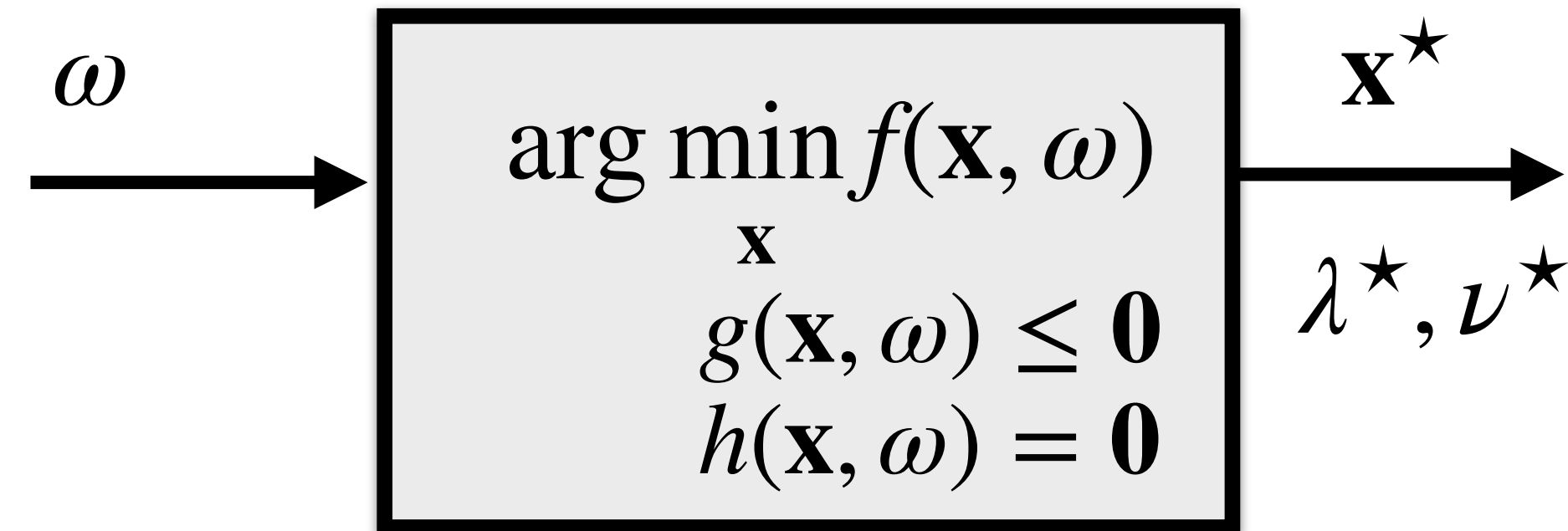
Constrained optimizer:



ODE solver:



Constrained optimizer



- Given parameters ω , typical optimizer provides primal-dual solution $\mathbf{x}^*, \lambda^*, \nu^*$ that satisfy necessary and sufficient optimality conditions (KKT):

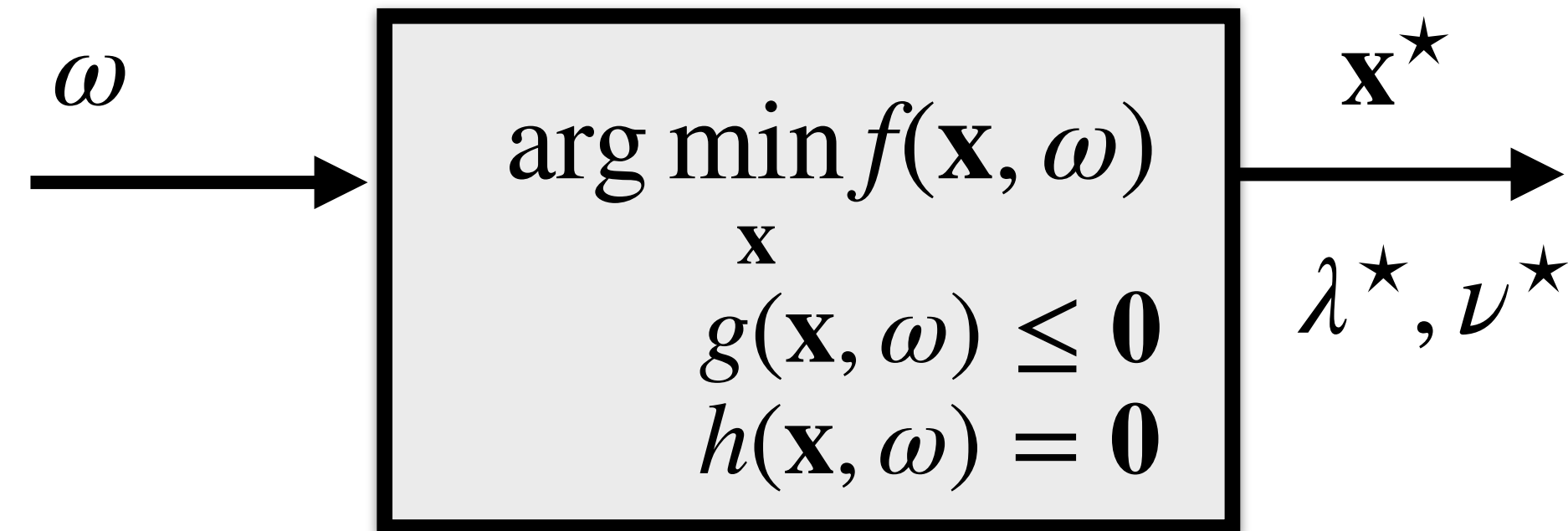
Lagrangian:
$$L(\mathbf{x}, \lambda, \nu) = f(\mathbf{x}, \omega) + \sum_i \lambda_i g_i(\mathbf{x}, \omega) + \sum_i \nu_i h_i(\mathbf{x}, \omega)$$

Stationarity:
$$\frac{\partial L(\mathbf{x}^*, \lambda^*, \nu^*)}{\partial \mathbf{x}^*} = \frac{\partial f(\mathbf{x}^*, \omega)}{\partial \mathbf{x}^*} + \sum_i \lambda_i \frac{\partial g_i(\mathbf{x}, \omega)}{\partial \mathbf{x}^*} + \sum_i \nu_i \frac{\partial h_i(\mathbf{x}, \omega)}{\partial \mathbf{x}^*} = \mathbf{0}$$

Complementarity:
$$\lambda^* \circ g(\mathbf{x}^*, \omega) = \mathbf{0}$$

Primal feasibility:
$$\begin{aligned} g(\mathbf{x}^*, \omega) &\leq \mathbf{0} \\ h(\mathbf{x}^*, \omega) &= \mathbf{0} \end{aligned}$$

Constrained optimizer



- Given parameters ω , typical optimizer provides primal-dual solution $\mathbf{x}^*, \lambda^*, \nu^*$ that satisfy necessary and sufficient optimality conditions (KKT):

Lagrangian:
$$L(\mathbf{x}, \lambda, \nu) = f(\mathbf{x}, \omega) + \sum_i \lambda_i g_i(\mathbf{x}, \omega) + \sum_i \nu_i h_i(\mathbf{x}, \omega)$$

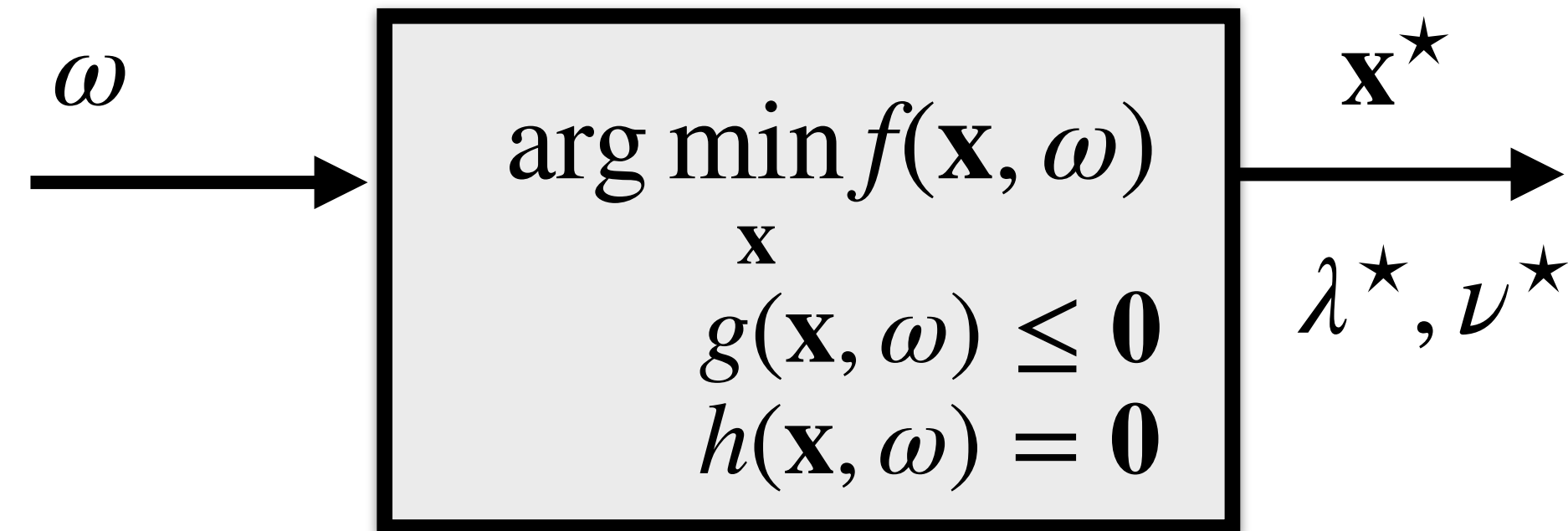
Stationarity:
$$\frac{\partial L(\mathbf{x}^*, \lambda^*, \nu^*)}{\partial \mathbf{x}^*} = \frac{\partial f(\mathbf{x}^*, \omega)}{\partial \mathbf{x}^*} + \sum_i \lambda_i \frac{\partial g_i(\mathbf{x}, \omega)}{\partial \mathbf{x}^*} + \sum_i \nu_i \frac{\partial h_i(\mathbf{x}, \omega)}{\partial \mathbf{x}^*} = \mathbf{0}$$

Complementarity:
$$\lambda^* \circ g(\mathbf{x}^*, \omega) = \mathbf{0}$$

Primal feasibility:
$$\begin{aligned} g(\mathbf{x}^*, \omega) &\leq \mathbf{0} \\ h(\mathbf{x}^*, \omega) &= \mathbf{0} \end{aligned}$$

Dual feasibility:
$$\lambda^* \geq \mathbf{0}$$

Constrained optimizer



- Given parameters ω , typical optimizer provides primal-dual solution $\mathbf{x}^*, \lambda^*, \nu^*$ that satisfy necessary and sufficient optimality conditions (KKT):

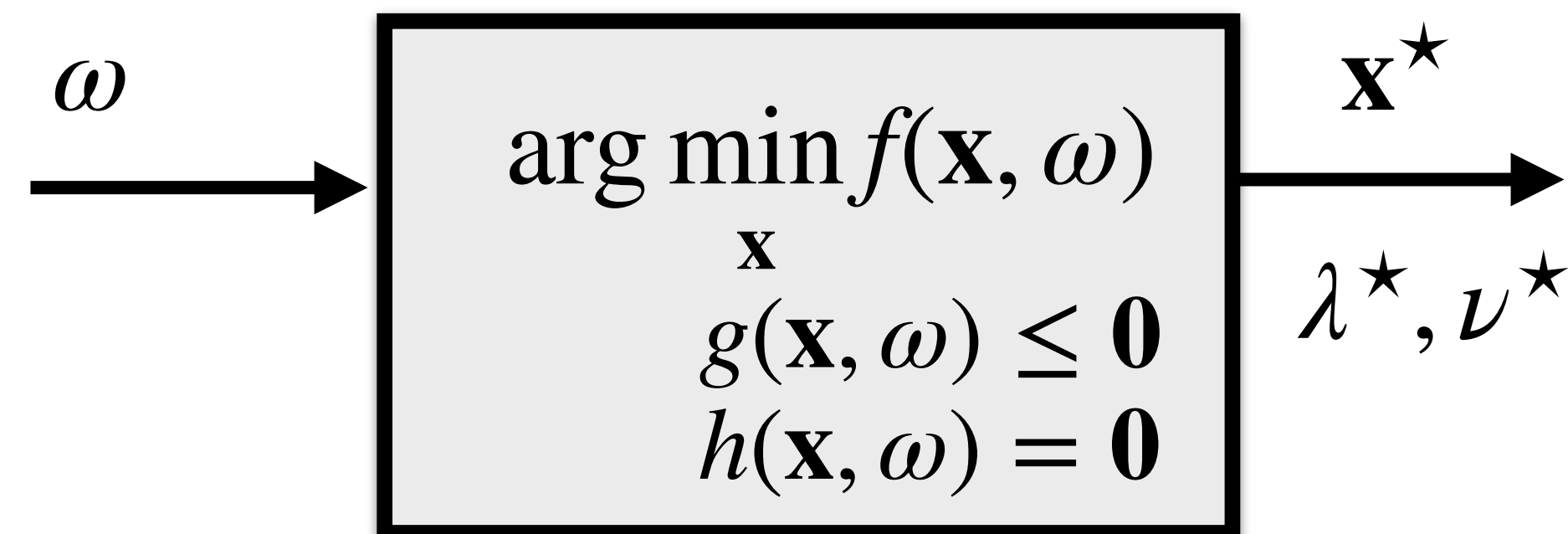
$$G(\mathbf{x}^*, \lambda^*, \nu^*, \omega) = \begin{bmatrix} \frac{\partial f(\mathbf{x}^*, \omega)}{\partial \mathbf{x}^*} + \sum_i \lambda_i^* \frac{\partial g_i(\mathbf{x}, \omega)}{\partial \mathbf{x}^*} + \sum_i \nu_i^* \frac{\partial h_i(\mathbf{x}, \omega)}{\partial \mathbf{x}^*} \\ \lambda^* \circ g(\mathbf{x}^*, \omega) \\ h(\mathbf{x}^*, \omega) \end{bmatrix} = \mathbf{0}$$

$$g(\mathbf{x}^*, \omega) \leq \mathbf{0}$$

$$\lambda^* \geq \mathbf{0}$$

Inequalities can be ignored due to complementarity condition (i.e. c.s. implies that one of them have to be zero, if $g(\mathbf{x}^*, \omega) = \mathbf{0}$ and $\lambda > \mathbf{0}$ then small change of λ will preserve positivity of λ and the other way around)

Constrained optimizer



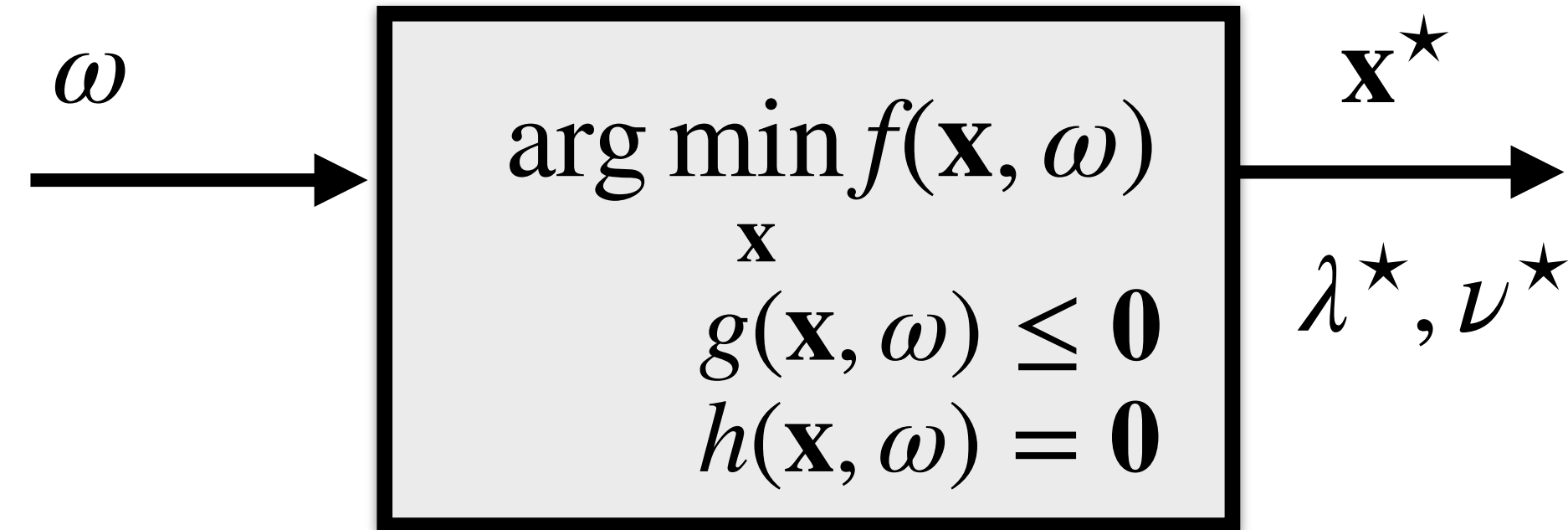
- Given parameters ω , typical optimizer provides primal-dual solution $\mathbf{x}^*, \lambda^*, \nu^*$ that satisfy necessary and sufficient optimality conditions (KKT):

$$G(\mathbf{x}^*, \lambda^*, \nu^*, \omega) = \begin{bmatrix} \frac{\partial f(\mathbf{x}^*, \omega)}{\partial \mathbf{x}^*} + \sum_i \lambda_i^* \frac{\partial g_i(\mathbf{x}, \omega)}{\partial \mathbf{x}^*} + \sum_i \nu_i^* \frac{\partial h_i(\mathbf{x}, \omega)}{\partial \mathbf{x}^*} \\ \lambda^* \circ g(\mathbf{x}^*, \omega) \\ h(\mathbf{x}^*, \omega) \end{bmatrix} = \mathbf{0}$$

- Input params ω can change output solution $(\mathbf{x}^*, \lambda^*, \nu^*)(\omega) = [\mathbf{x}^*(\omega), \lambda^*(\omega), \nu^*(\omega)]$ only in direction that preserves these conditions, therefore:

$$\frac{\partial G(\mathbf{x}^*(\omega), \lambda^*(\omega), \nu^*(\omega), \omega)}{\partial \omega} = \mathbf{0}$$

Constrained optimizer



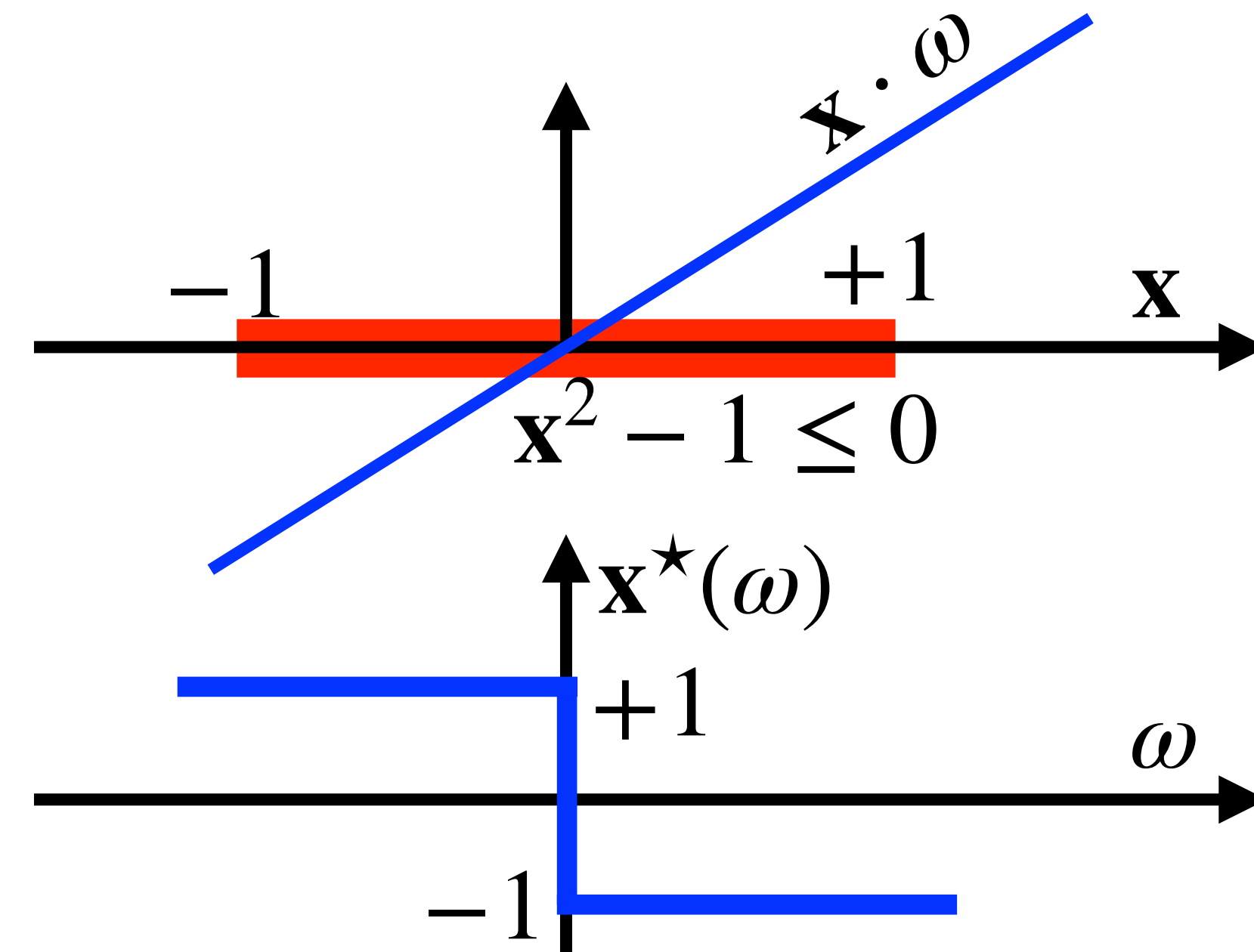
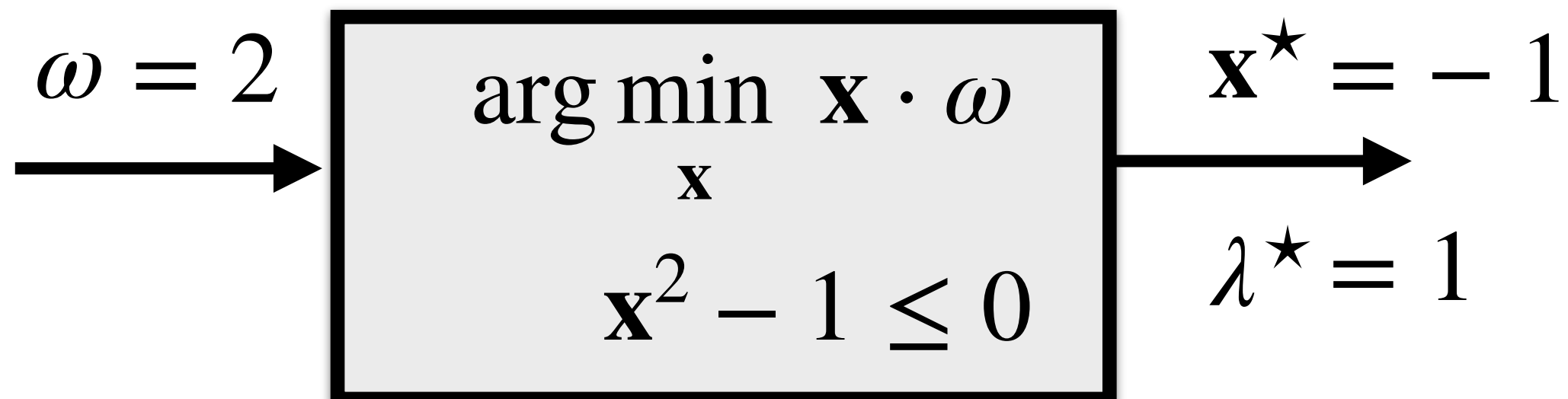
$\frac{\partial G(\mathbf{x}^*(\omega), \lambda^*(\omega), \nu^*(\omega), \omega)}{\partial \omega} = \mathbf{0}$
Computing derivative of the compound function yields desired equation for the gradient

$$\partial_{0,1,2} G(\mathbf{x}^*(\omega), \lambda^*(\omega), \nu^*(\omega), \omega) \frac{\partial(\mathbf{x}^*, \lambda^*, \nu^*)(\omega)}{\partial \omega} + \partial_3 G(\mathbf{x}^*(\omega), \lambda^*(\omega), \nu^*(\omega), \omega) = \mathbf{0}$$

$$\frac{\partial(\mathbf{x}^*, \lambda^*, \nu^*)(\omega)}{\partial \omega} = - \left[\partial_{0,1,2} G(\mathbf{x}^*(\omega), \lambda^*(\omega), \nu^*(\omega), \omega) \right]^{-1} \partial_3 G(\mathbf{x}^*(\omega), \lambda^*(\omega), \nu^*(\omega), \omega)$$

Constrained optimizer

Example:



$$G(\mathbf{x}^*(\omega), \lambda^*(\omega), \omega) = \begin{bmatrix} \frac{\partial f(\mathbf{x}^*, \omega)}{\partial \mathbf{x}^*} + \lambda \frac{\partial g(\mathbf{x}, \omega)}{\partial \mathbf{x}^*} \\ \lambda g(\mathbf{x}^*, \omega) \end{bmatrix} = \begin{bmatrix} \omega + 2\lambda \mathbf{x} \\ \lambda(\mathbf{x}^2 - 1) \end{bmatrix}$$

$$\frac{\partial(\mathbf{x}^*, \lambda^*)(\omega)}{\partial \omega} = - \left[\partial_{0,1} G(\mathbf{x}^*(\omega), \lambda^*(\omega), \omega) \right]^{-1} \partial_2 G(\mathbf{x}^*(\omega), \lambda^*(\omega), \omega)$$

$$= - \begin{bmatrix} 2\lambda^* & 2\mathbf{x}^* \\ 2\lambda^* \mathbf{x}^* & (\mathbf{x}^*)^2 - 1 \end{bmatrix}^{-1} \cdot \begin{bmatrix} 1 \\ 0 \end{bmatrix} = - \begin{bmatrix} 2 & -2 \\ -2 & 0 \end{bmatrix}^{-1} \cdot \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 1/2 \end{bmatrix}$$

$$-\text{inv}([2 \ -2 \ ; \ -2 \ 0])^* [1;0]$$

Constrained optimizer

Example2:

ω

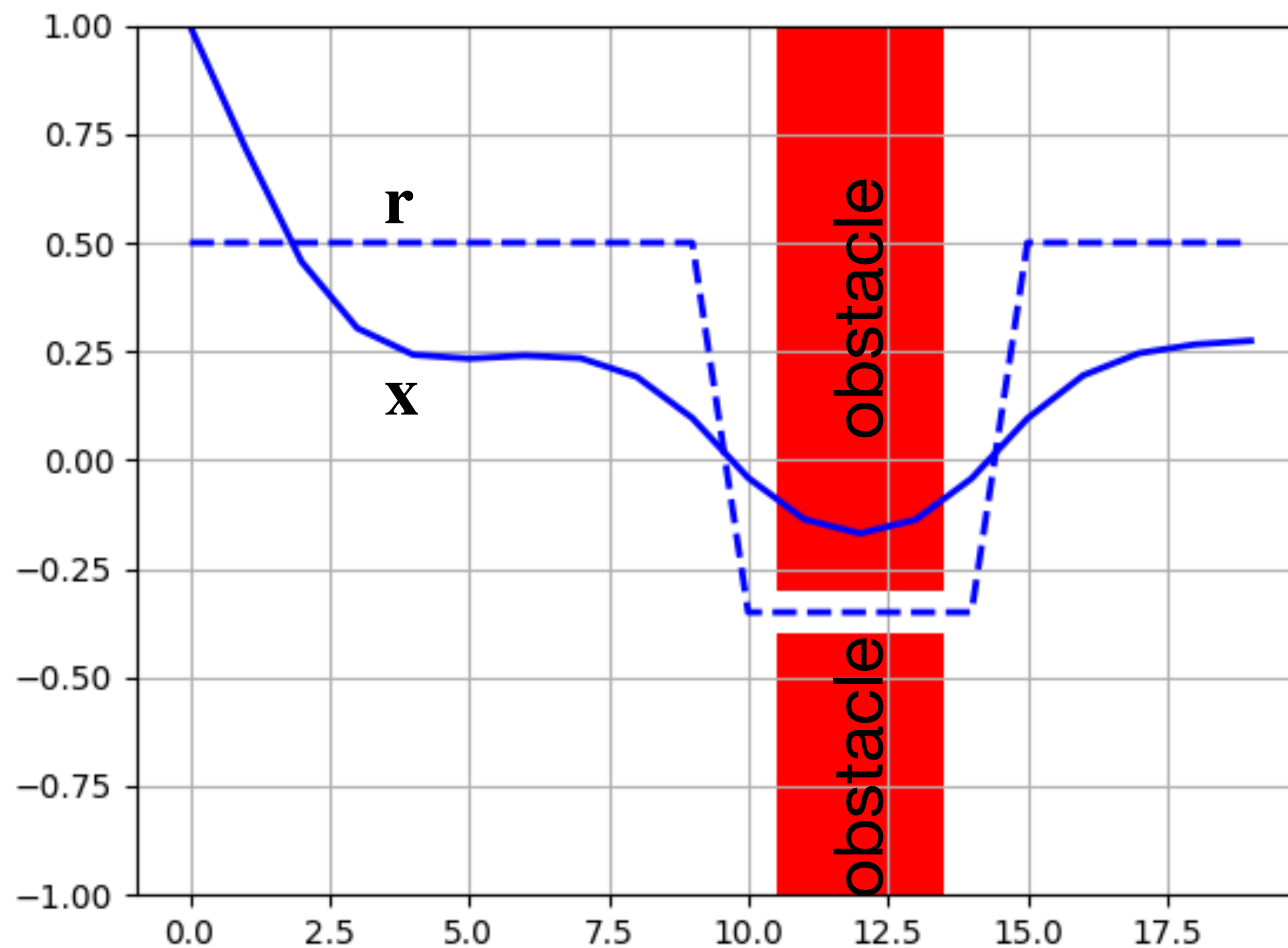
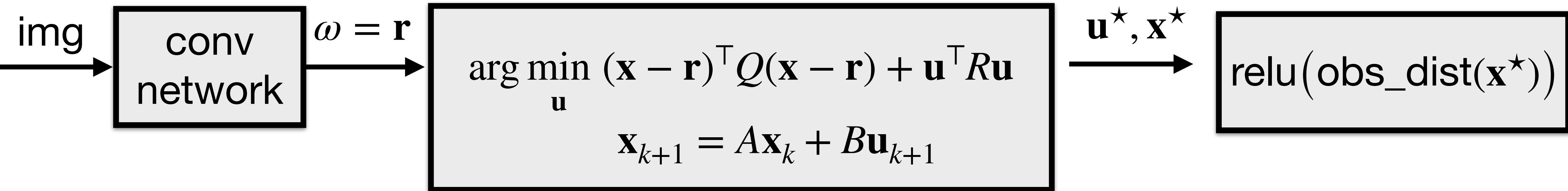
$$\begin{aligned} \arg \min_{\mathbf{x}} \quad & \frac{1}{2} \mathbf{x}^\top \mathbf{Q} \mathbf{x} + \mathbf{p}^\top \mathbf{x} \\ \text{subject to} \quad & \mathbf{G} \mathbf{x} \leq \mathbf{h} \\ & \mathbf{A} \mathbf{x} = \mathbf{b} \end{aligned}$$

$\mathbf{x}^\star$
 $\lambda^\star$

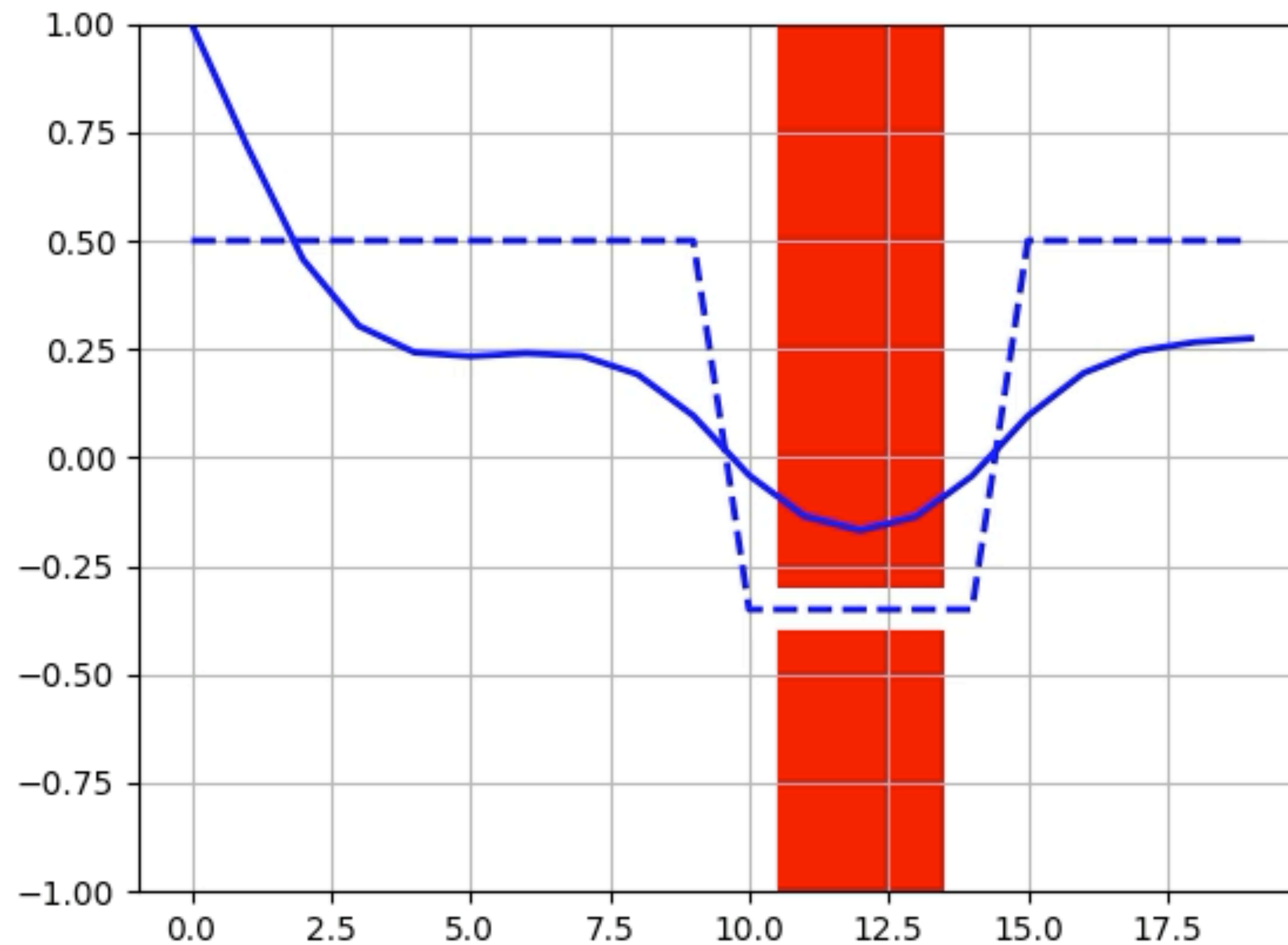
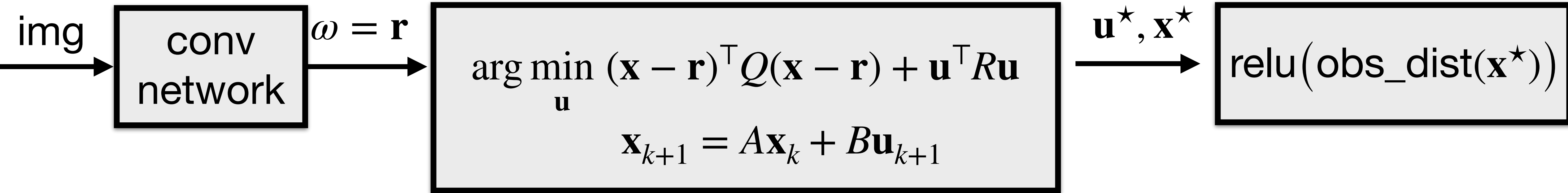
$$\partial_{0,1,2} G(\mathbf{x}^\star, \lambda^\star, \nu^\star, \omega) = \begin{bmatrix} \mathbf{Q} & \mathbf{G}^\top & \mathbf{A}^\top \\ \text{diag}(\lambda^\star) \mathbf{G} & \text{diag}(\mathbf{G} \mathbf{x}^\star - \mathbf{h}) & \mathbf{0} \\ \mathbf{A} & \mathbf{0} & \mathbf{0} \end{bmatrix}$$

$$\frac{\partial(\mathbf{x}^\star, \lambda^\star, \nu^\star)(\omega)}{\partial \omega} = - \left[\partial_{0,1,2} G(\mathbf{x}^\star, \lambda^\star, \nu^\star, \omega) \right]^{-1} \partial_3 G(\mathbf{x}^\star, \lambda^\star, \nu^\star, \omega)$$

Application: MPC with non-linear system



Application: MPC with non-linear system



PyTorch embedded modeling language for convex optimization problems

<https://www.cvxpy.org/>

[Boyd, et al, NIPS, 2019]

```
cvxpylayer = CvxpyLayer(problem, parameters=[A, b], variables=[x])  
solution, = cvxpylayer(A, b) # feed-forward pass (solve the problem)  
solution.sum().backward()   # backward pass (how A,b influence solution)
```

CVXpy

$$\begin{aligned} x^*(\theta) = \operatorname{argmin}_x f(x; \theta) \\ \text{subject to } g(x; \theta) \leq 0 \\ h(x; \theta) = 0 \end{aligned}$$

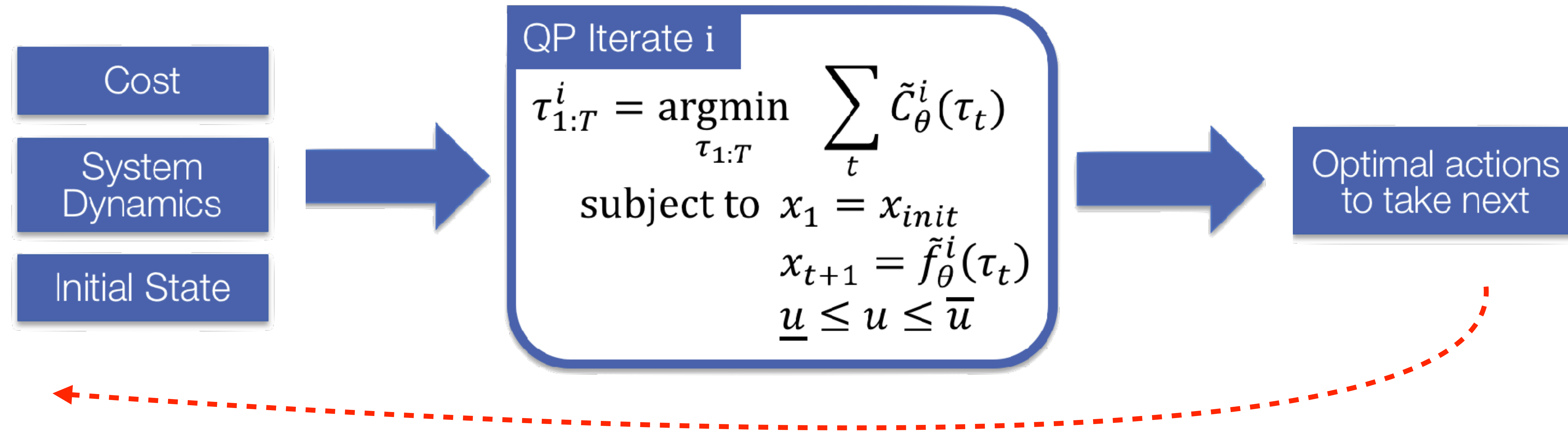


 PyTorch



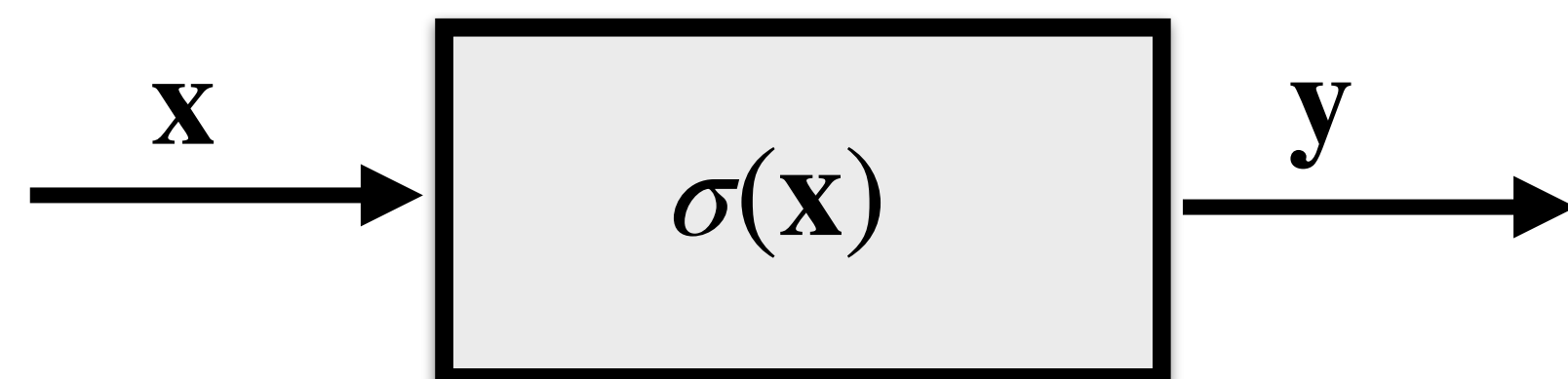
TensorFlow

Also optnet: <https://github.com/locuslab/optnet>



Explicit layers

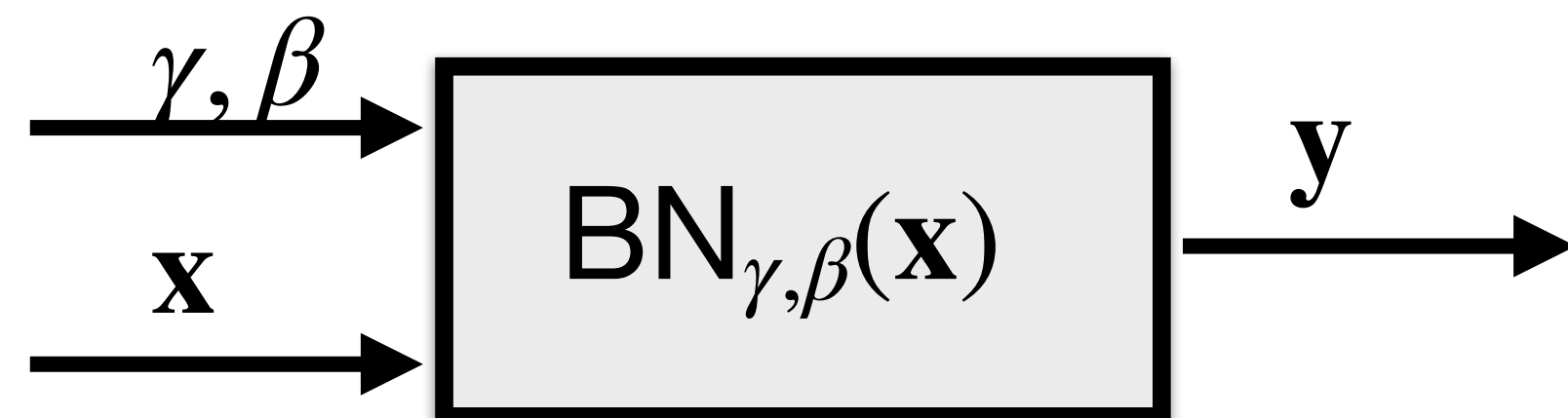
Sigmoid:



Convolution:

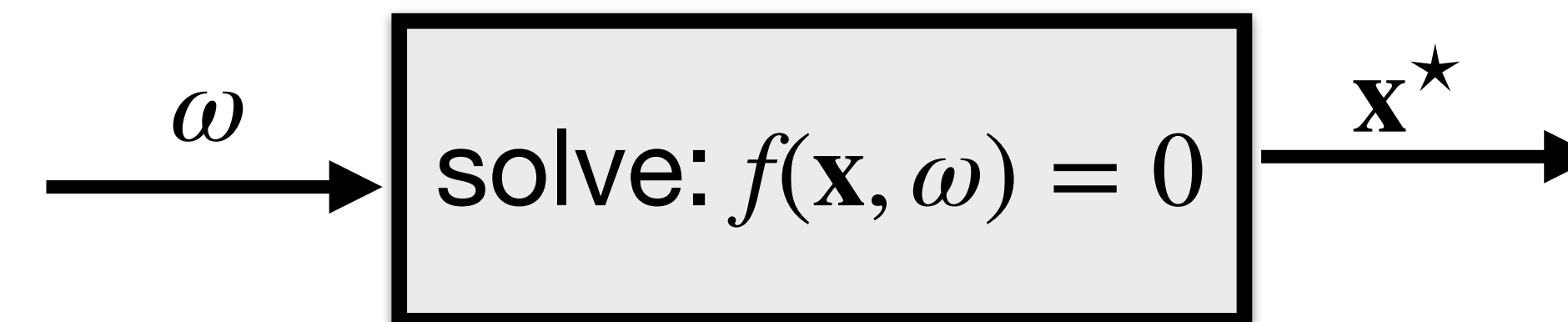


Batch-norm:

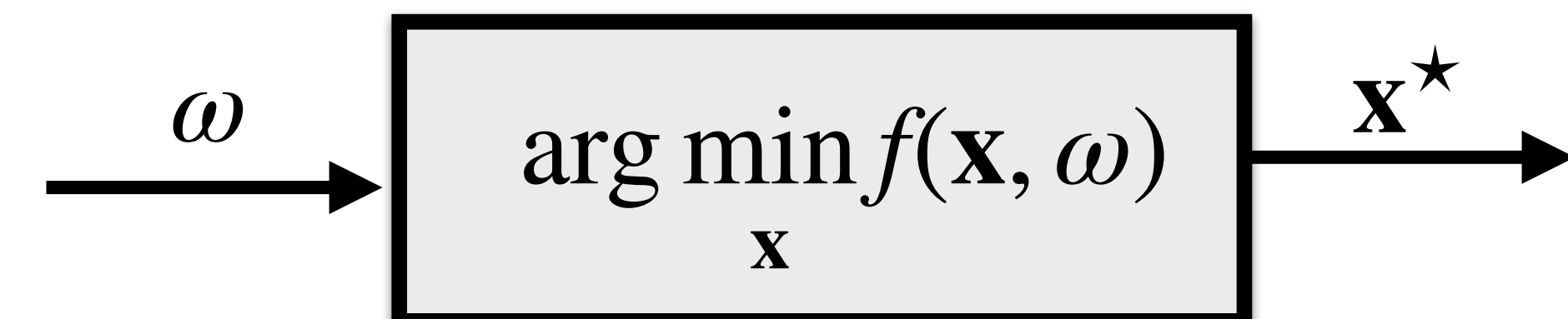


Implicit layers

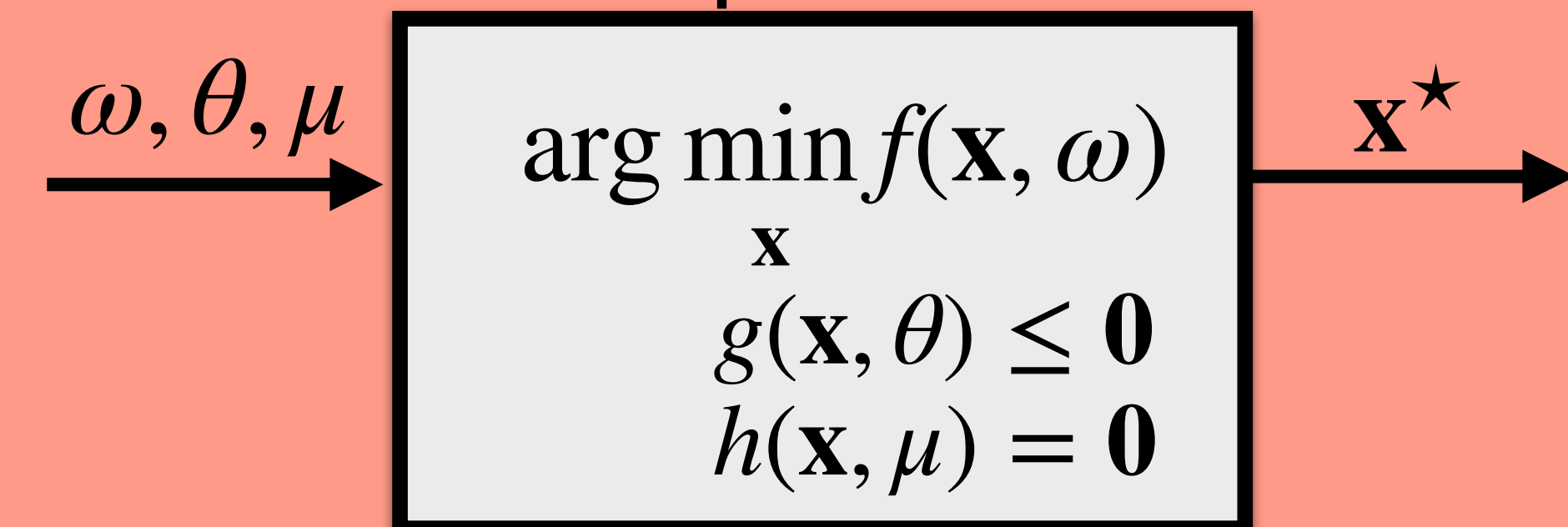
Root finder:



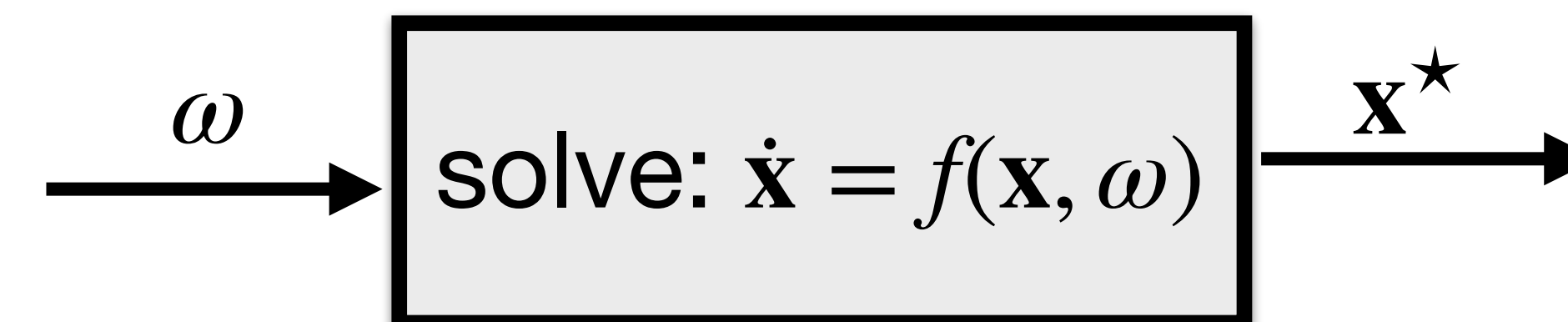
Unconstrained optimizer:



Constrained optimizer:

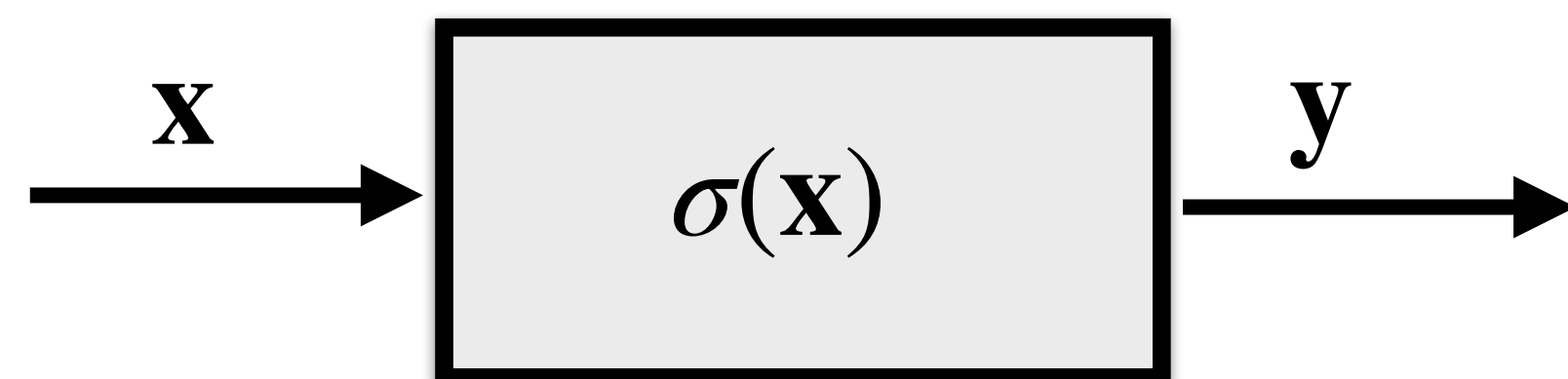


ODE solver:



Explicit layers

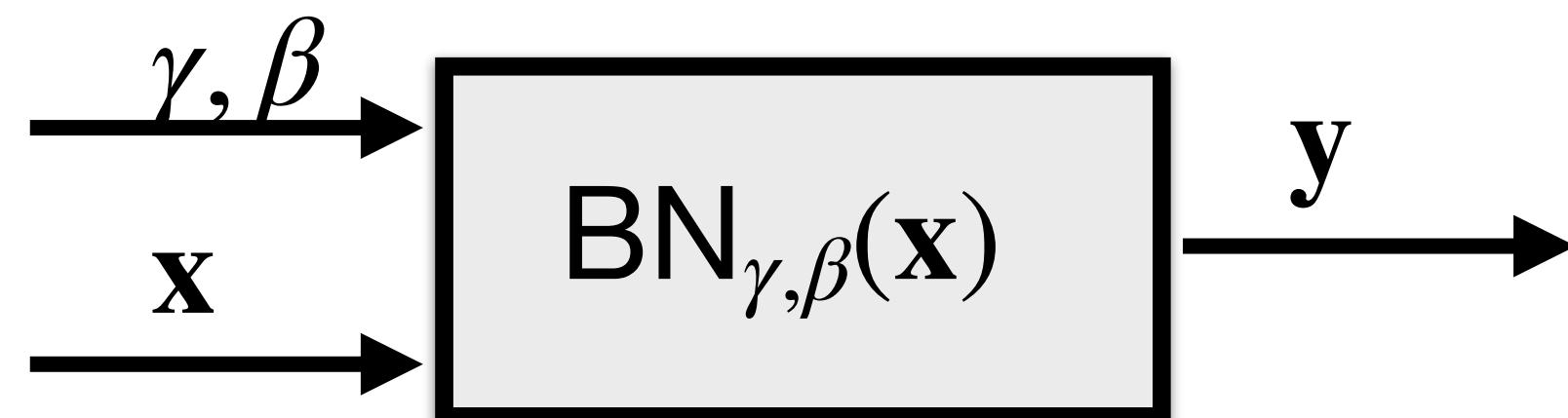
Sigmoid:



Convolution:

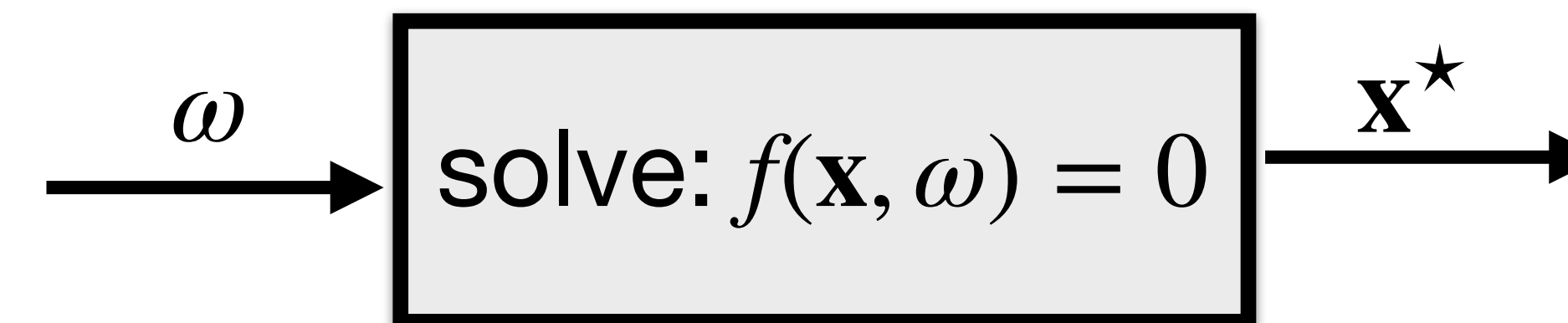


Batch-norm:

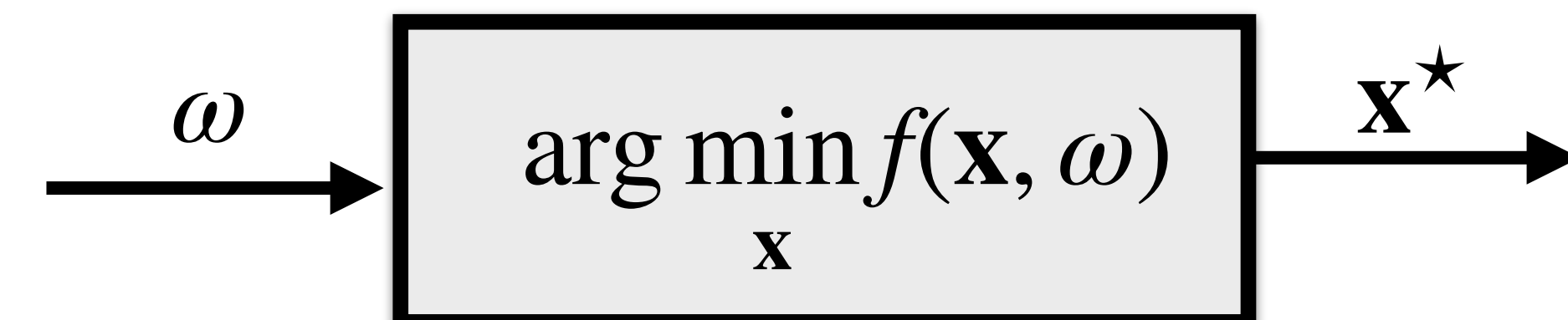


Implicit layers

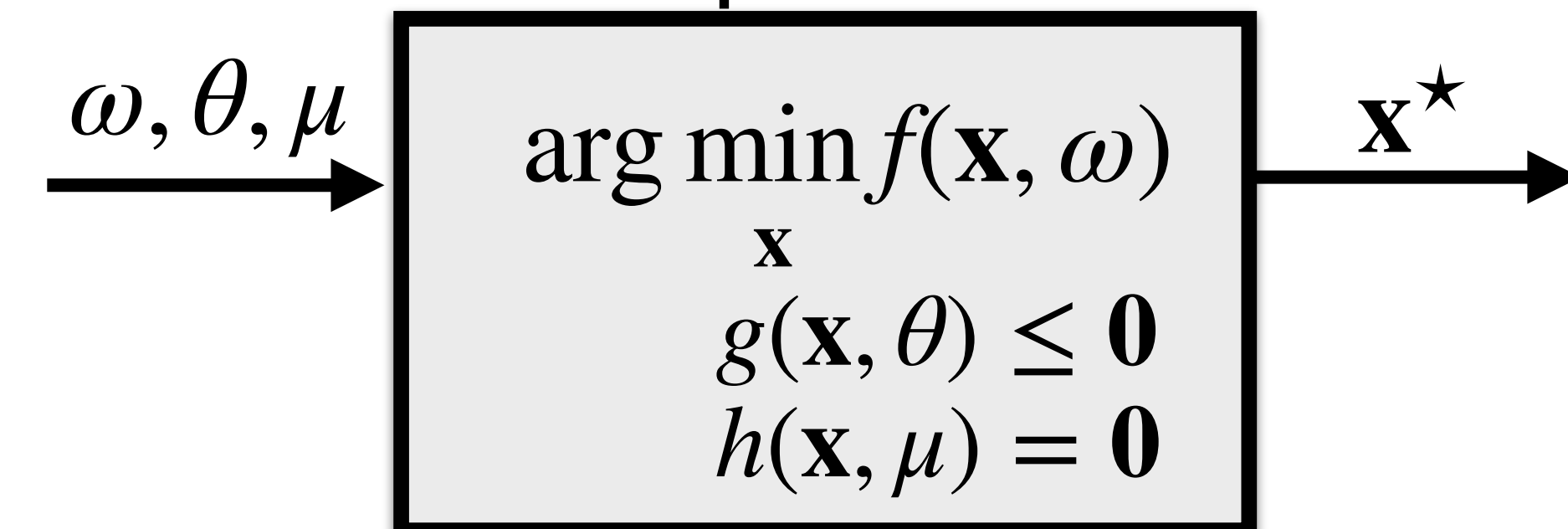
Root finder:



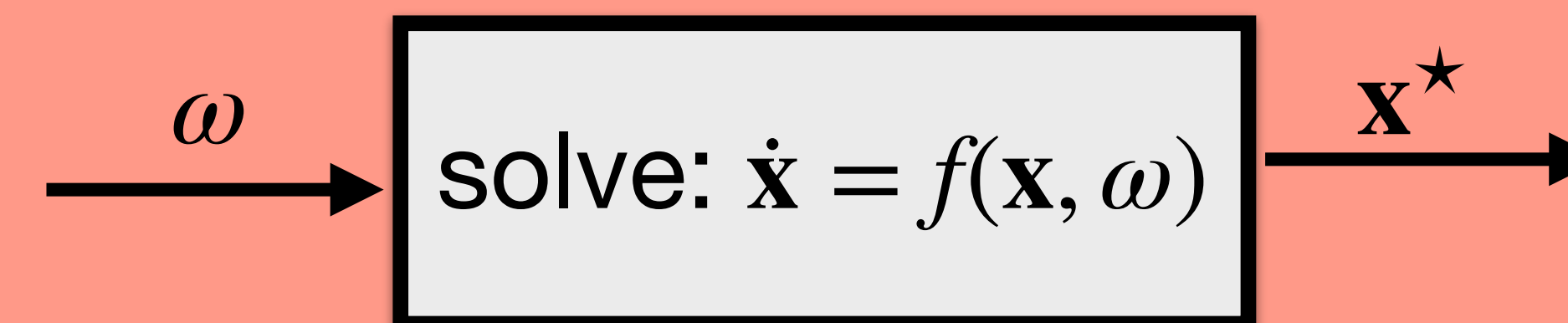
Unconstrained optimizer:



Constrained optimizer:



ODE solver:



ODE solver

$$\dot{\mathbf{x}} = f(\mathbf{x}, \omega)$$

Any solution $\mathbf{x}^\star(t, \omega)$ (obtained from an arbitrary ODE solver) have to follow:

$$\frac{\partial}{\partial t} \mathbf{x}^\star(t, \omega) - f(\mathbf{x}^\star(t, \omega), \omega) = 0 \quad \text{Given } \omega, \text{ the ODE solution } \mathbf{x}^\star(\omega) \text{ satisfy this eq.}$$

$$\frac{\partial}{\partial \omega} \frac{\partial}{\partial t} \mathbf{x}^\star(t, \omega) = \frac{\partial f(\mathbf{x}^\star(t, \omega), \omega)}{\partial \omega} \quad \text{Change of } \omega \text{ generates ODE solution } \mathbf{x}^\star(\omega) \text{ that satisfy this equation.}$$

$$\frac{\partial}{\partial t} \frac{\partial}{\partial \omega} \mathbf{x}^\star(t, \omega) = \partial_0 f(\mathbf{x}^\star(t, \omega), \omega) \frac{\partial \mathbf{x}^\star(t, \omega)}{\partial \omega} + \partial_1 f(\mathbf{x}^\star(t, \omega), \omega)$$

$\mathbf{z}(t, \omega)$

$\mathbf{z}(t, \omega)$

$g(\mathbf{z}(t, \omega), \omega)$

$$\frac{\partial}{\partial t} \mathbf{z}(t, \omega) = g(\mathbf{z}(t, \omega), \omega, t)$$

$$\dot{\mathbf{z}} = g(\mathbf{z}, \omega, t)$$

Since solvers often use adaptive step length, both ODEs solved jointly

ODE solver

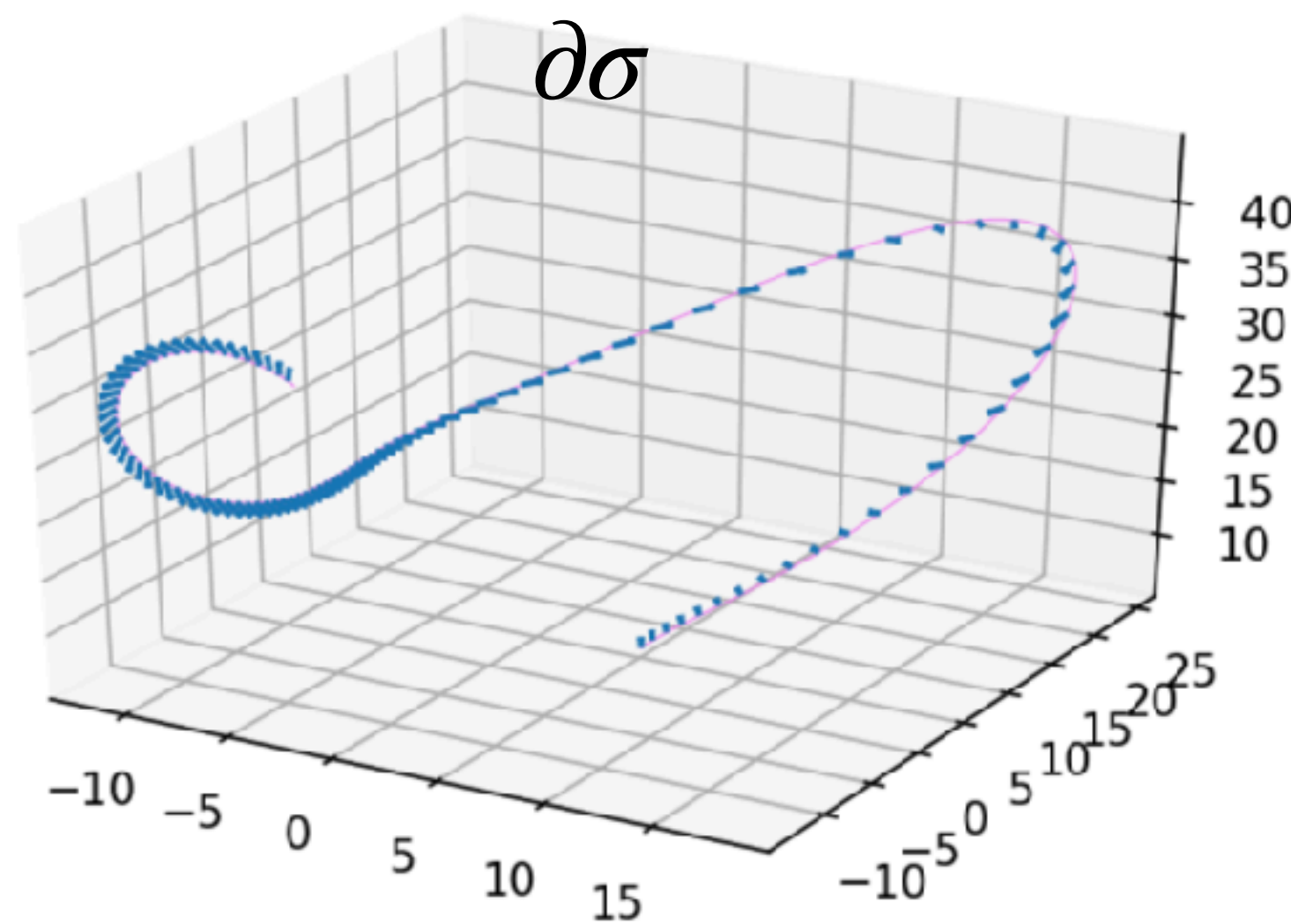
Since solvers often use adaptive step length, both ODEs solved jointly

$$\begin{aligned} \dot{\mathbf{x}} &= f(\mathbf{x}, \omega) \\ \dot{\mathbf{z}} &= g(\mathbf{z}, \omega, t) \end{aligned} \Rightarrow \begin{bmatrix} \dot{\mathbf{x}} \\ \dot{\mathbf{z}} \end{bmatrix} = \begin{bmatrix} f(\mathbf{x}, \omega) \\ \partial_0 f(\mathbf{x}, \omega) \mathbf{z} + \partial_1 f(\mathbf{x}, \omega) \end{bmatrix} \Rightarrow \mathbf{z}^\star = \frac{\partial \mathbf{x}^\star(t, \omega)}{\partial \omega}$$

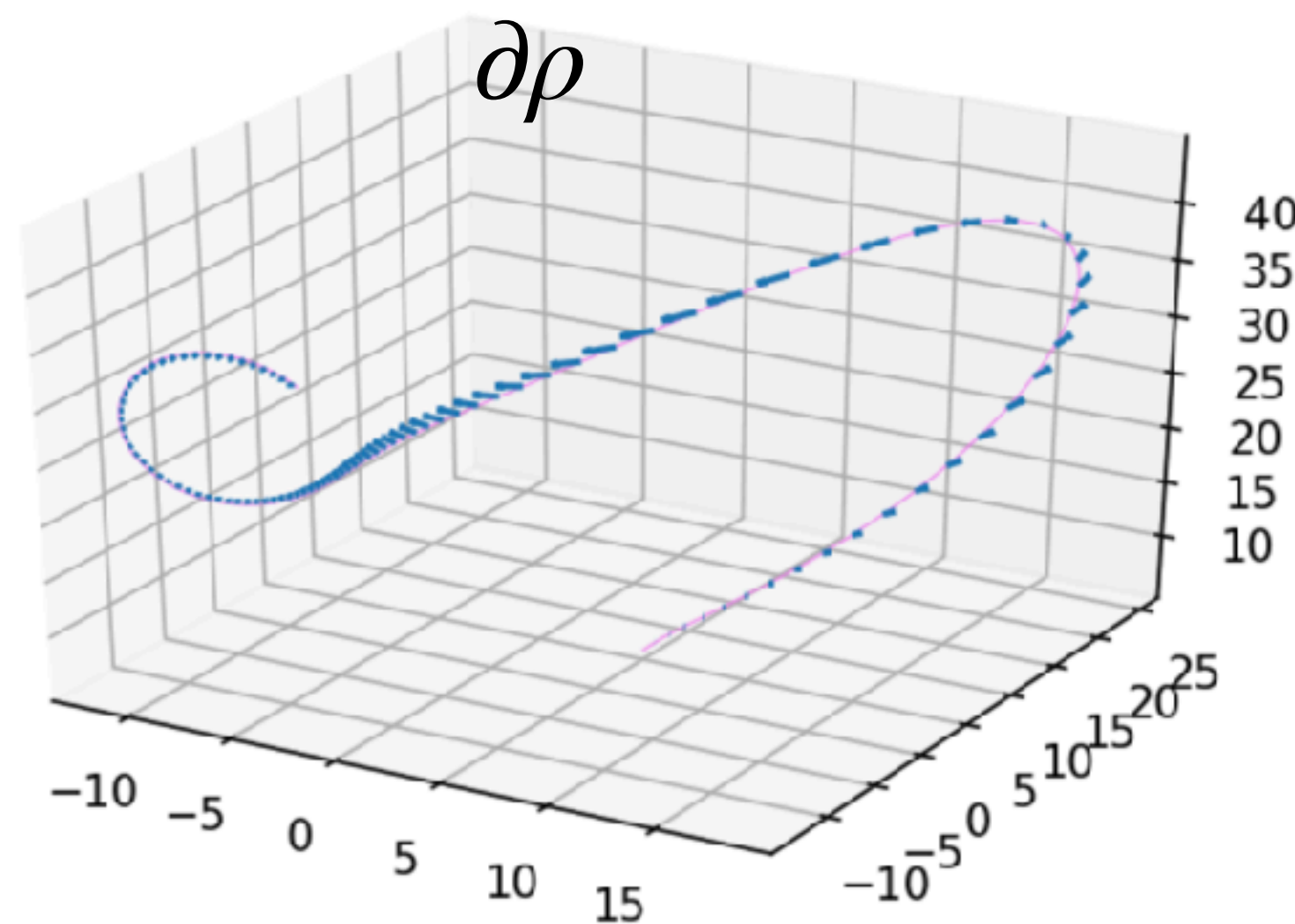
Example:
$$\begin{bmatrix} \dot{x} \\ \dot{y} \\ \dot{z} \end{bmatrix} = f\left(\begin{bmatrix} x \\ y \\ z \end{bmatrix}, \begin{bmatrix} \sigma \\ \rho \\ \beta \end{bmatrix}\right) = \begin{bmatrix} \sigma^*(y - x) \\ x(\rho - z) - y \\ xy - \beta z \end{bmatrix}$$

ODE solver: $\mathbf{x}^\star(\sigma, \rho, \beta)$

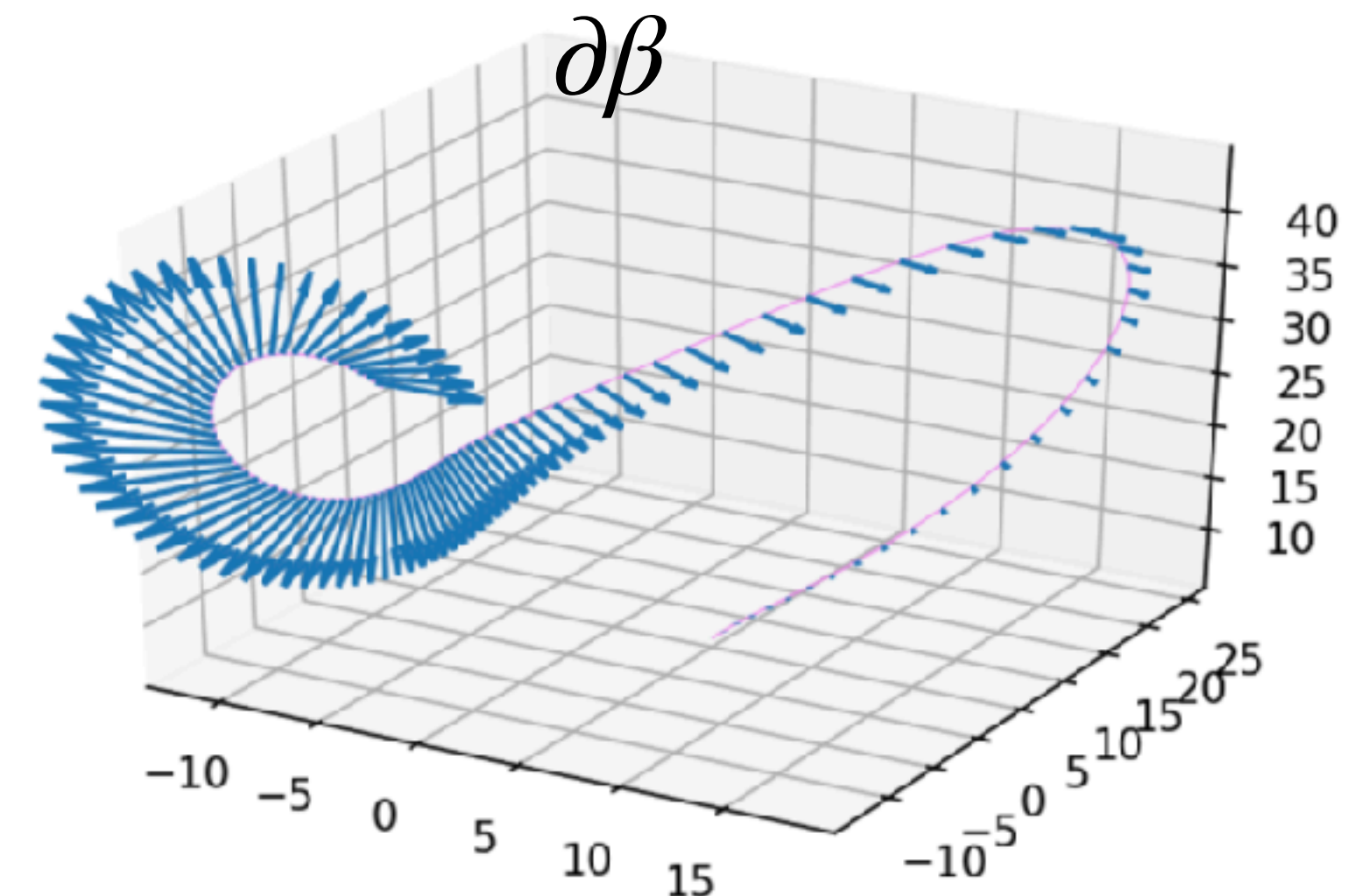
$$\frac{\partial \mathbf{x}^\star(\sigma, \rho, \beta)}{\partial \sigma}$$



$$\frac{\partial \mathbf{x}^\star(\sigma, \rho, \beta)}{\partial \rho}$$



$$\frac{\partial \mathbf{x}^\star(\sigma, \rho, \beta)}{\partial \beta}$$



Motivation

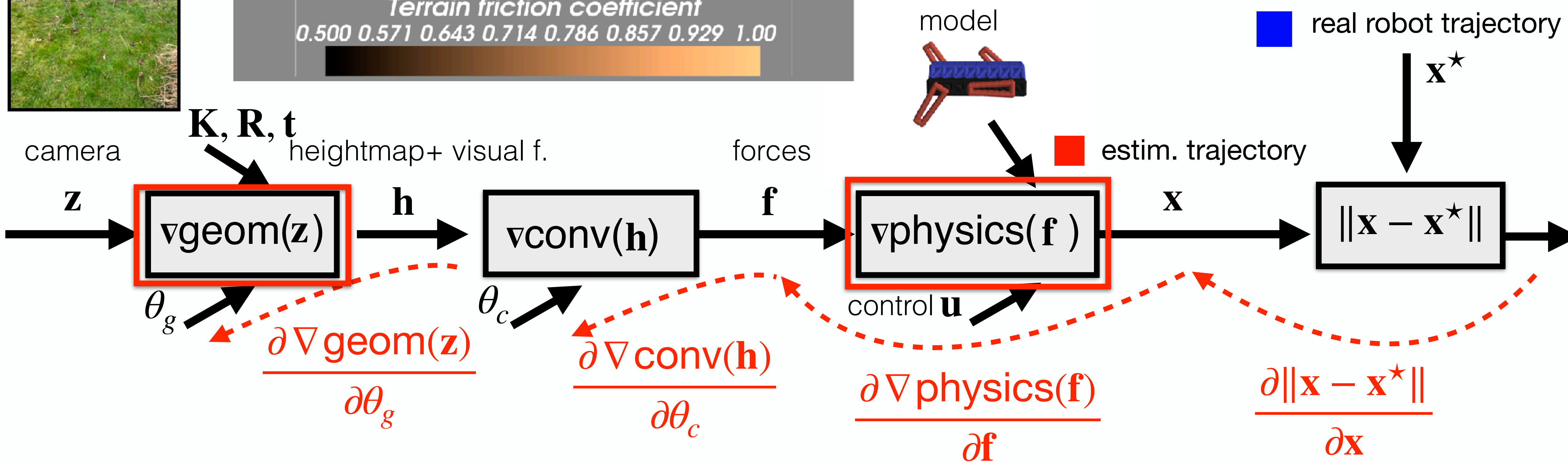
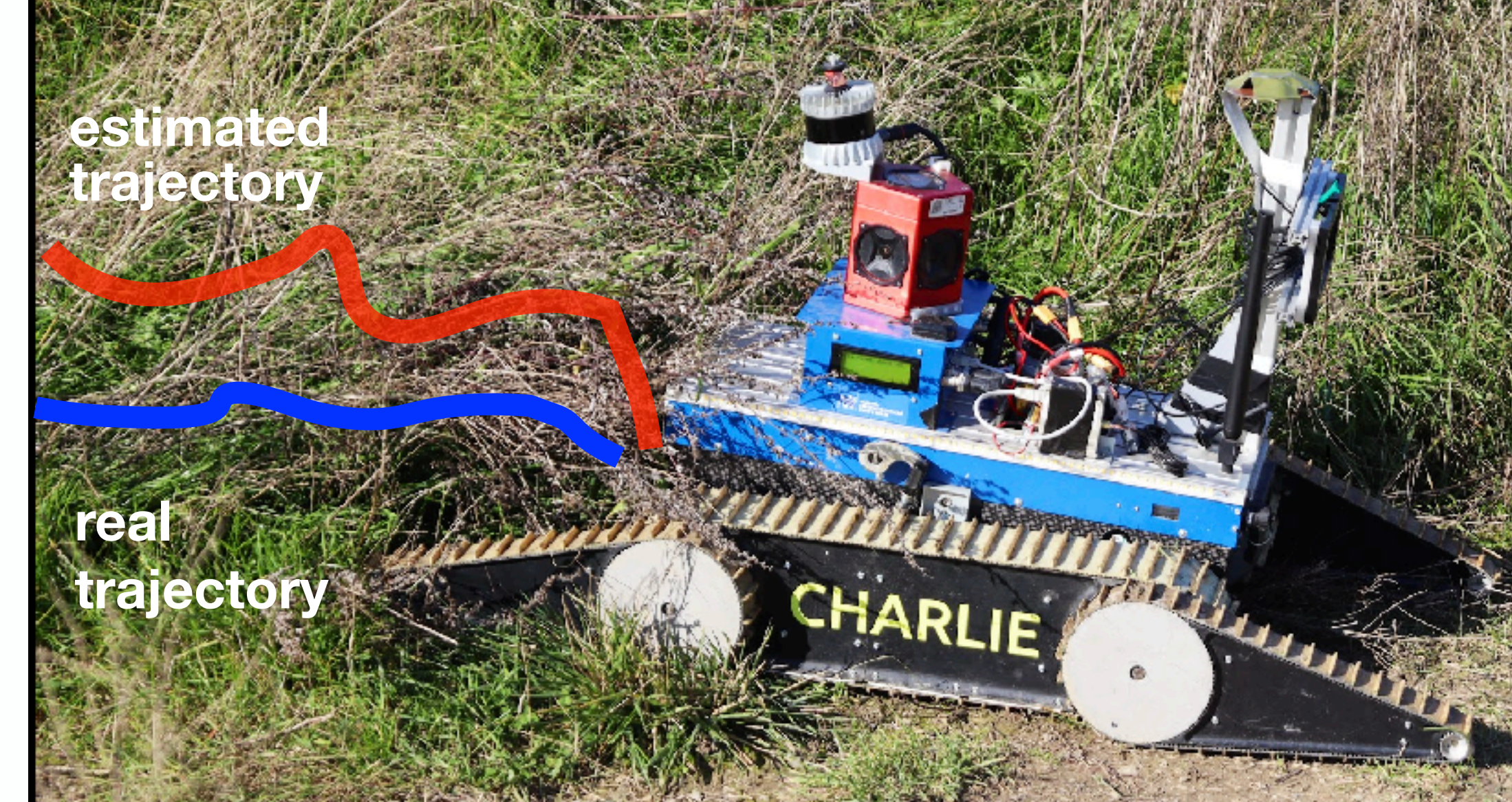
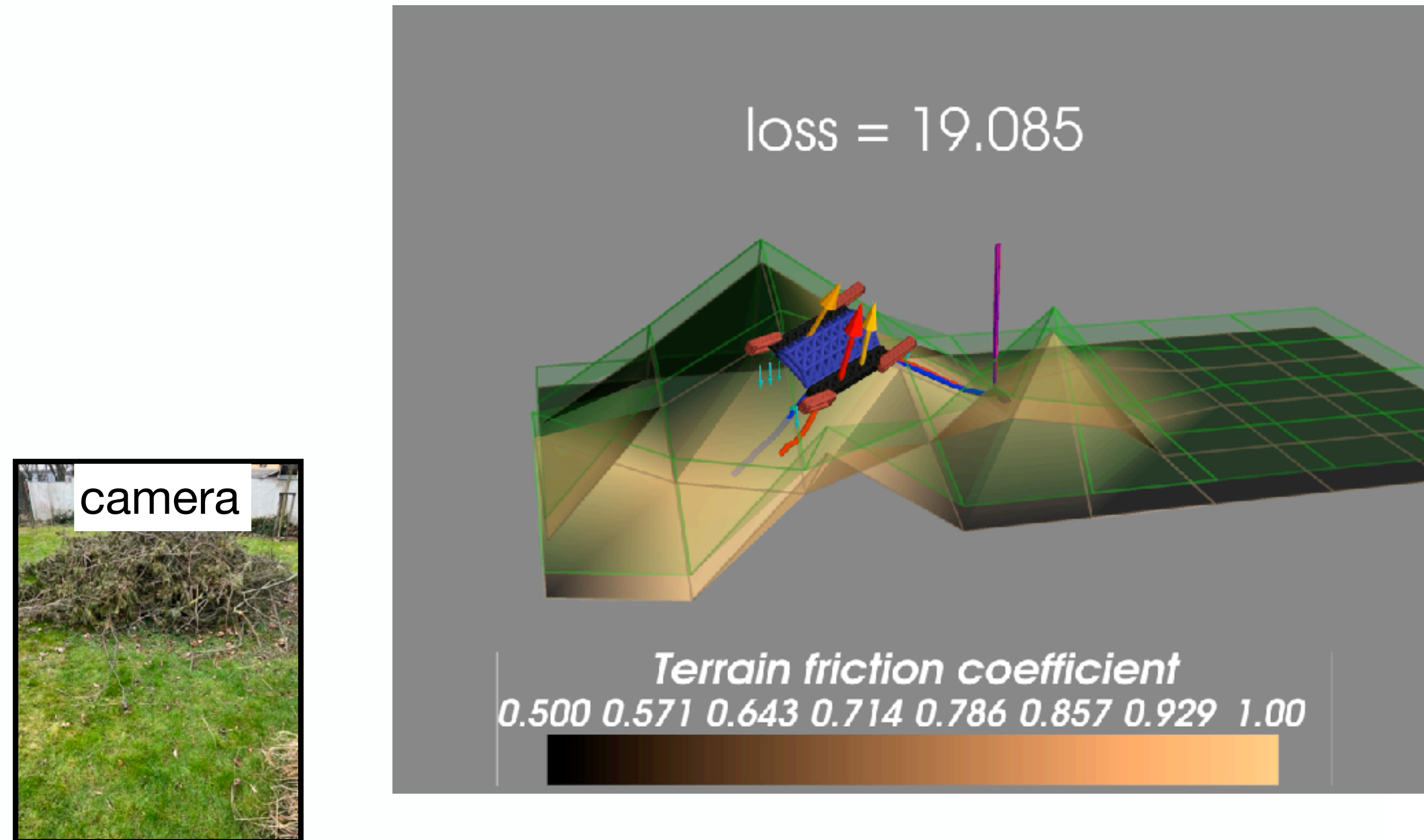
- **Missing piece preventing reliable deployment in wild outdoor environment is a robust model that predicts robot-terrain behaviour from onboard sensors.**



Unique wheel self-clean capability



Excellent obstacle crossing ability



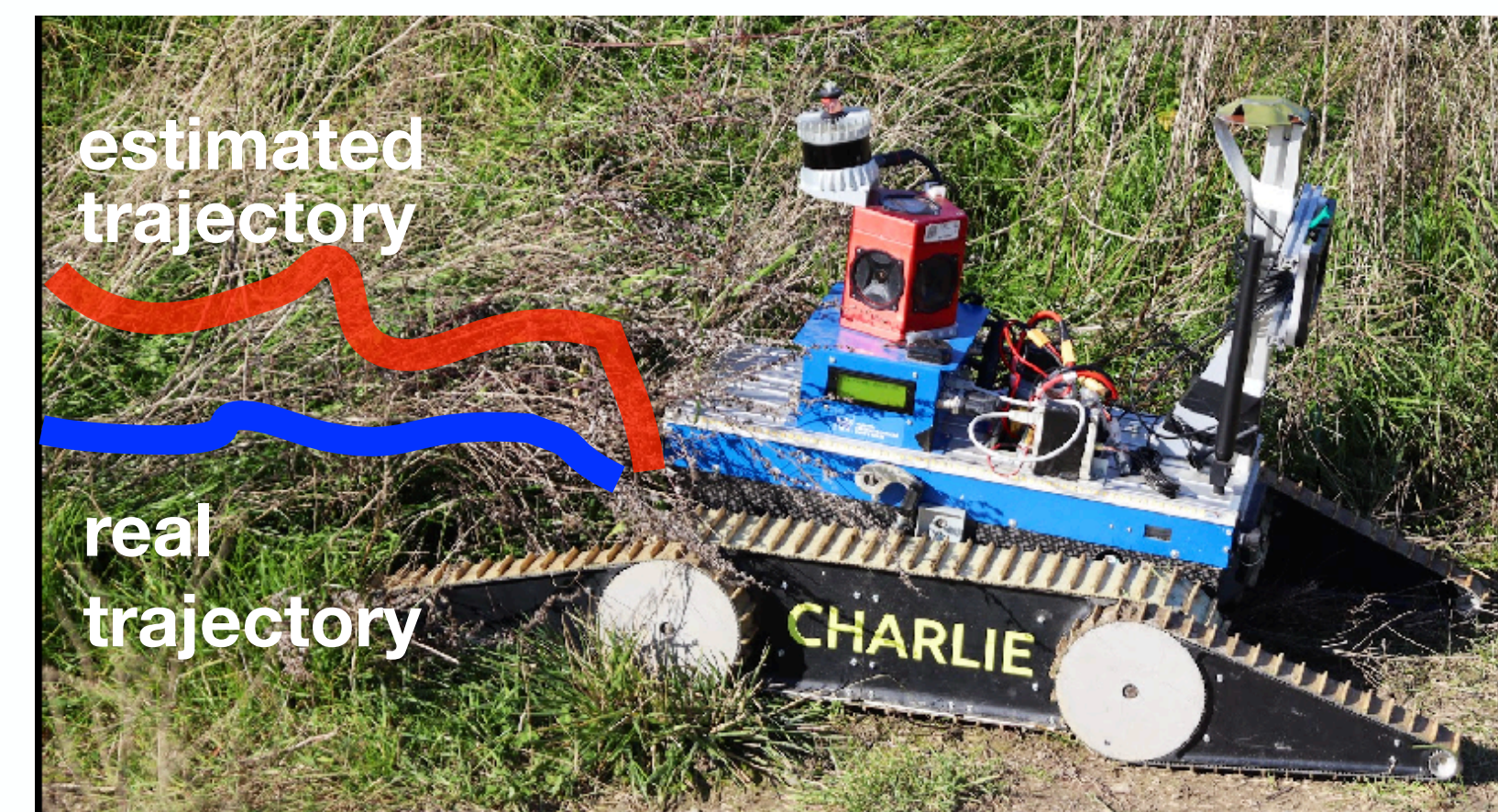
What do we offer to robotics community?

camera
recalibration

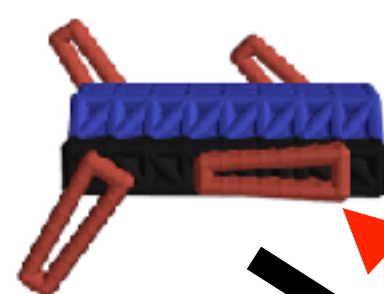
model
identification

parameter
learning

control, RL
planning



$\mathbf{K}, \mathbf{R}, \mathbf{t}$



$\theta_c \theta_g$

$\mathbf{u}$

$\mathbf{z}$
onboard
camera

end-to-end differentiable, $10^5 \times$ real time,
image conditioned simulation

estim. trajectory
 $\mathbf{x}$

real robot trajectory
 $\mathbf{x}^*$

$\|\mathbf{x} - \mathbf{x}^*\|$

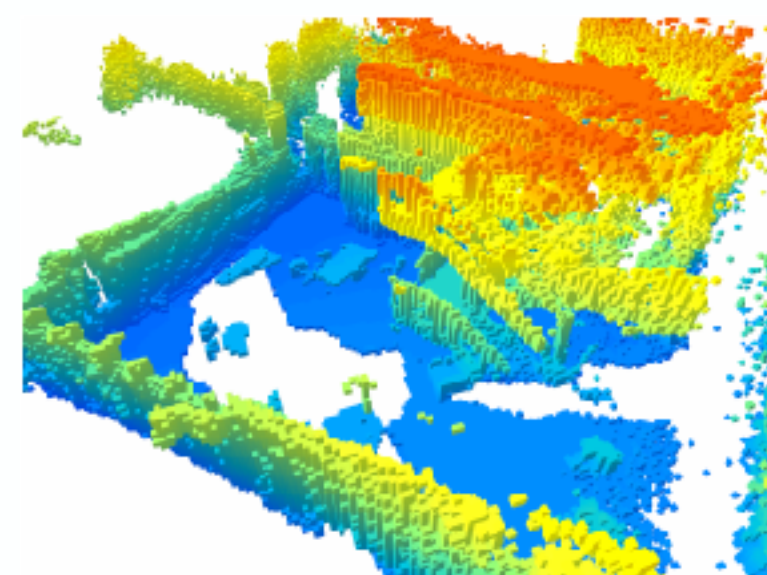
Interpretable

Intermediate: heightmaps forces contacts terrain properties

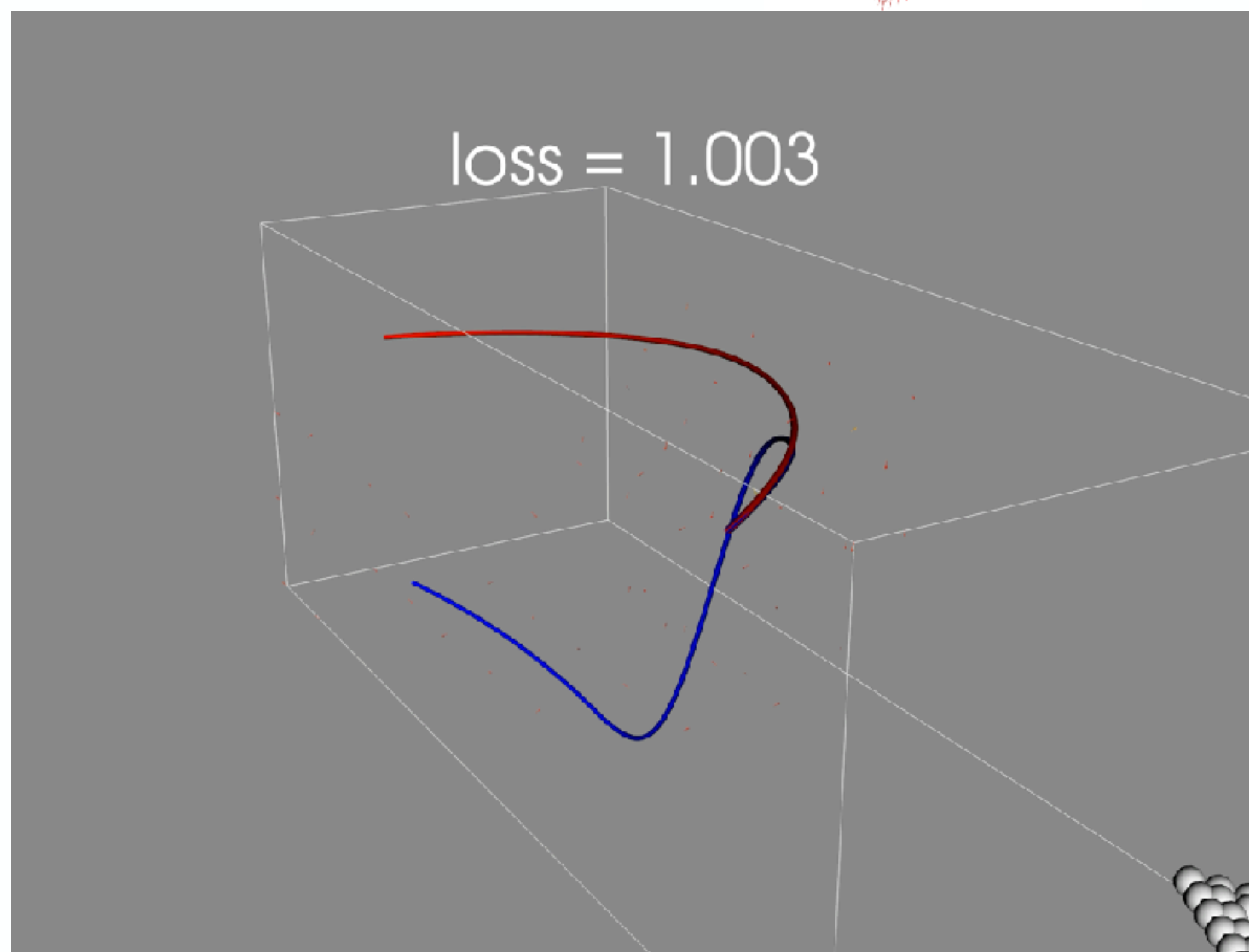
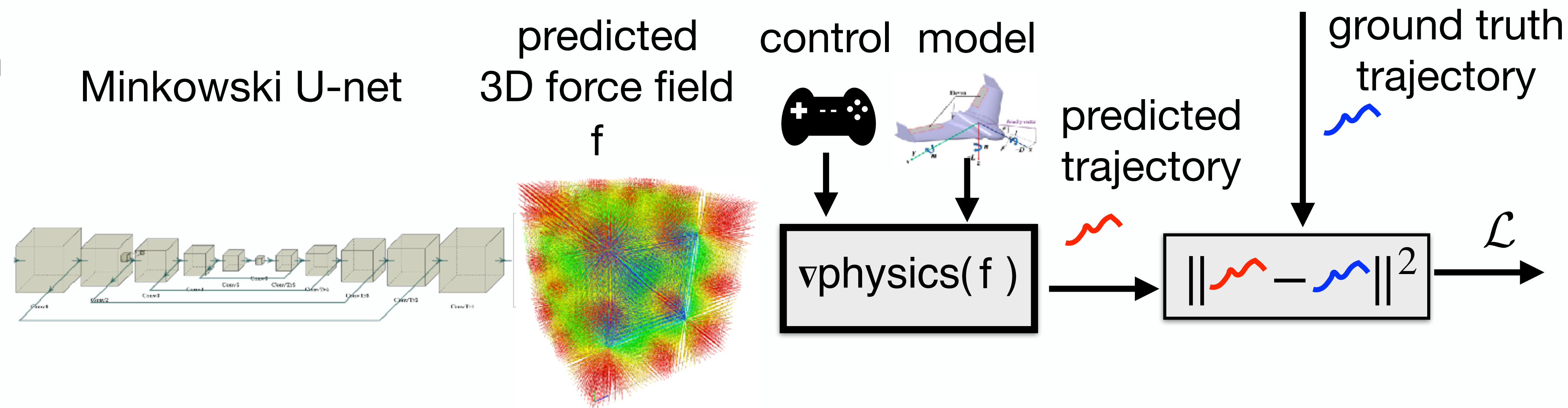
Outputs

self-supervision

3D gridmap with temperatures



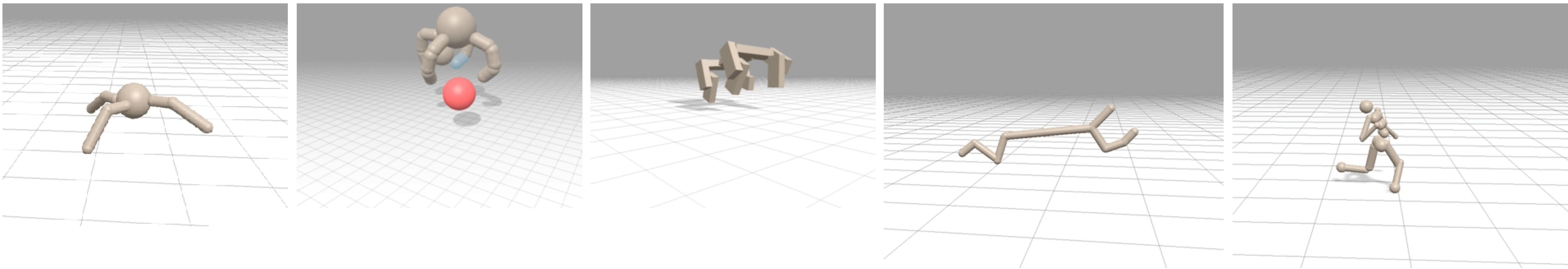
Fixedwing model



- predicted flow have to satisfy Navier-Stokes equations
- add NS-loss: $\|NS(f)\|$

Google's BRAX/ Nvidia's WARP - differentiable physics engine

<https://github.com/google/brax>



Brax simulates these environments at millions of physics steps per second on TPU



NVIDIA WARP

<https://developer.nvidia.com/warp-python>



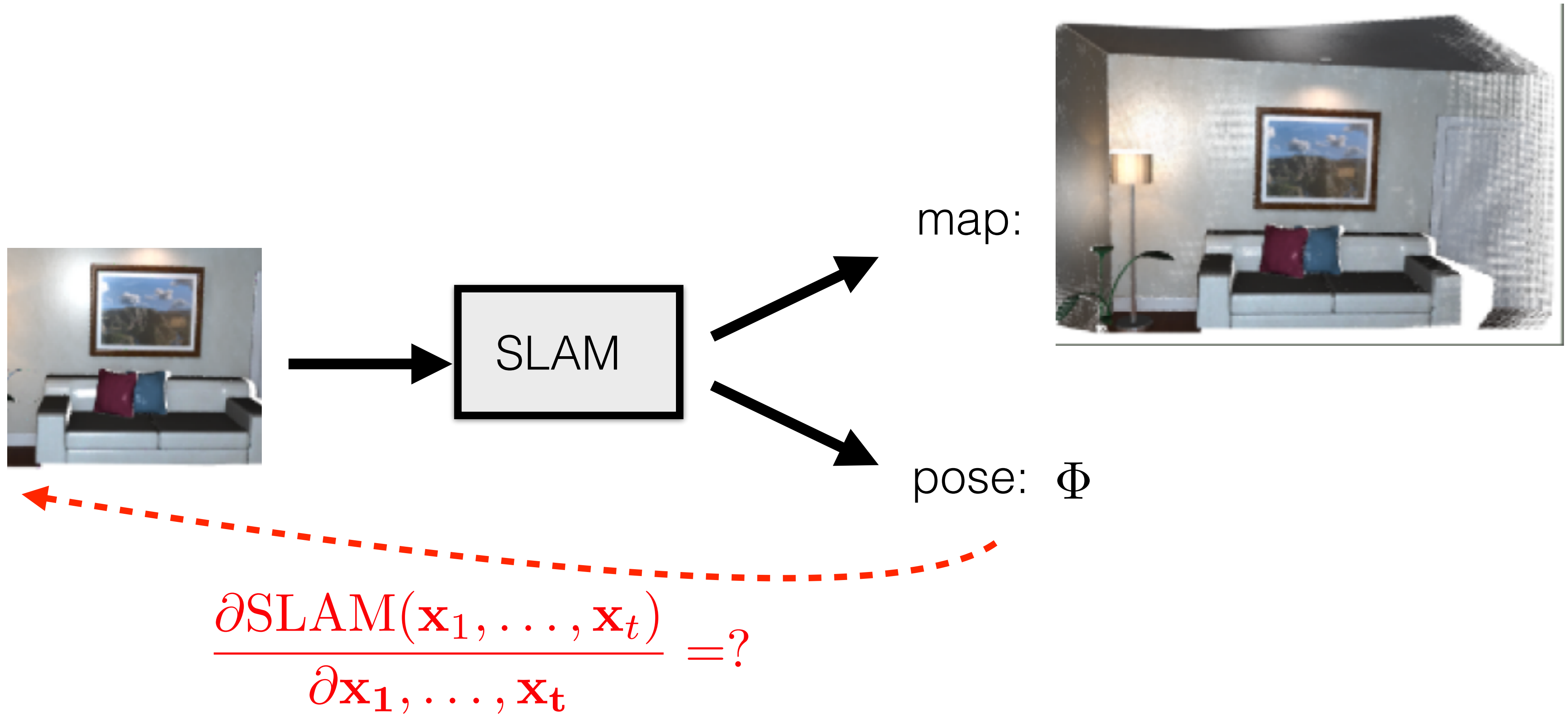
Cloth simulation



Particle-based simulation

Grad SLAM [Murthy, ICRA, 2021]

<https://gradslam.github.io/>

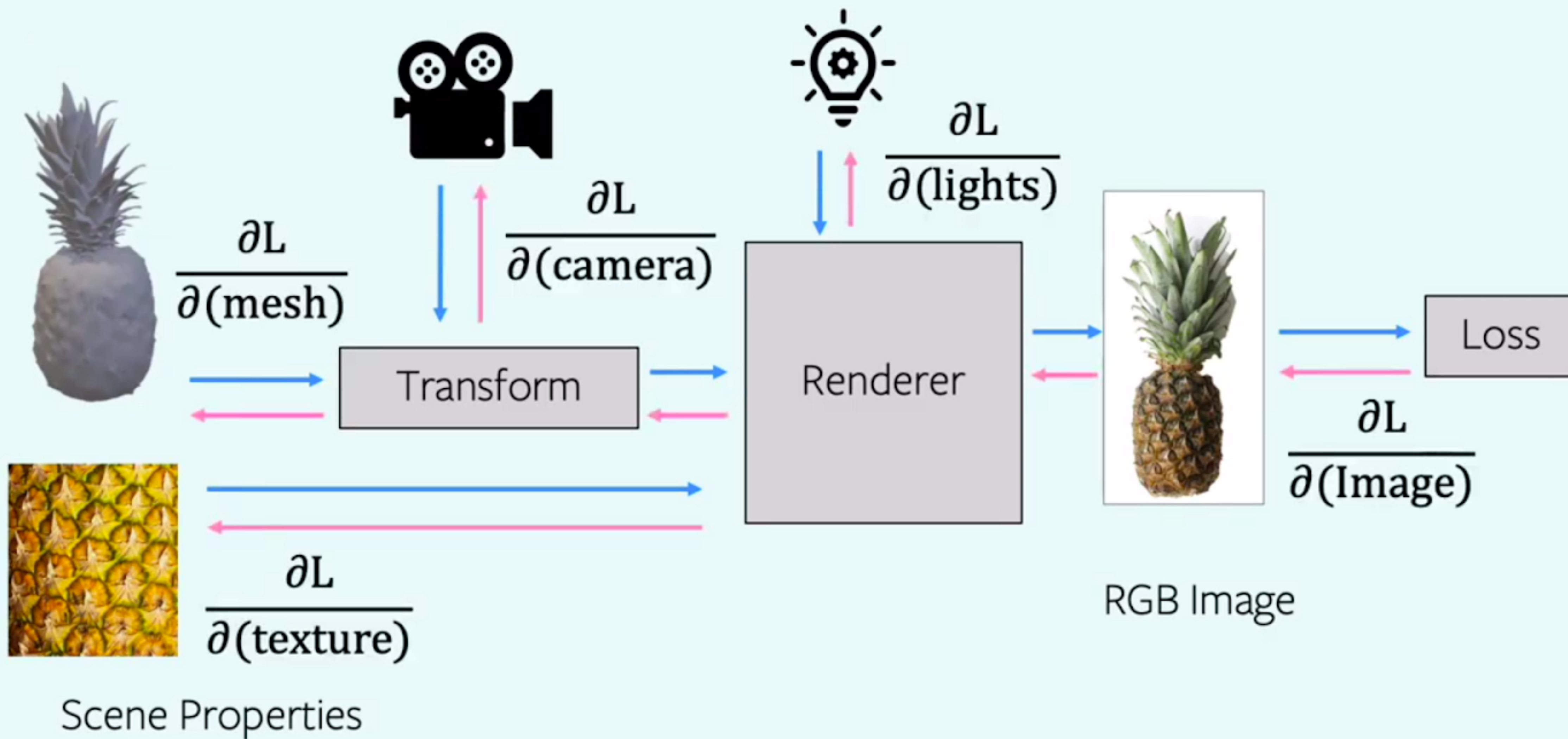




PyTorch3D

<https://pytorch3d.org/>

DIFFERENTIABLE RENDERING



Summary

- If you know that the mapping is **rendering** use ∇ **rendering** layer
- If you know that the mapping is **solution of opt. problem** use ∇ **solver** layer
- If you know that the mapping is **ODE solution** use ∇ **ODE** layer
- Always use the **right** tool ;-)



References to differentiable physics

1D/2D examples:

<https://towardsdatascience.com/physics-informed-neural-networks-pinns-an-intuitive-guide-fff138069563>

<https://medium.com/@theo.wolf/physics-informed-neural-networks-a-simple-tutorial-with-pytorch-f28a890b874a>

https://github.com/benmoseley/harmonic-oscillator-pinn-workshop/blob/main/PINN_intro_workshop.ipynb

Stable simulation:

<https://arxiv.org/pdf/2207.05060>

https://www.researchgate.net/publication/344464310_Detailed_Rigid_Body_Simulation_using_Extended_Position_Based_Dynamics