

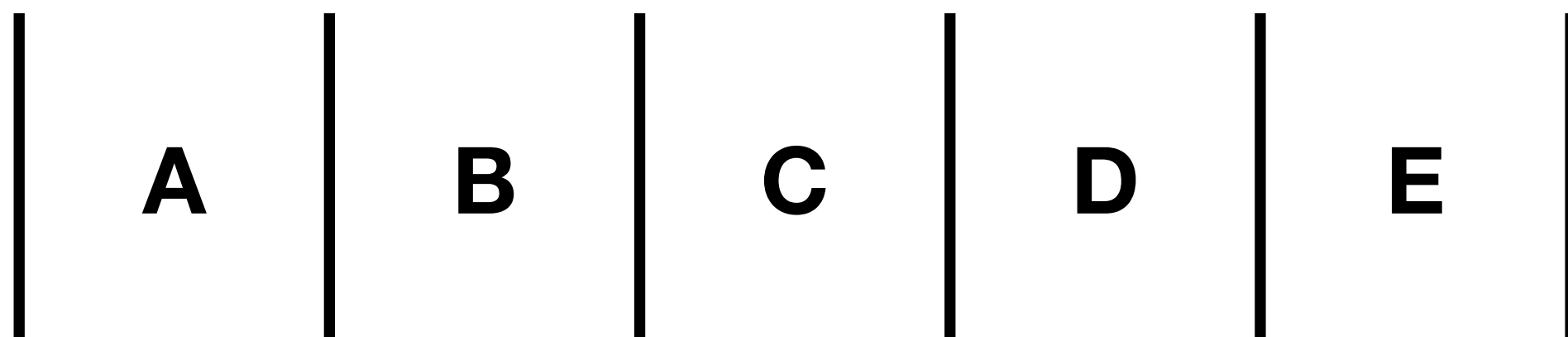
Time Efficient Swap Regret Minimization

Ashley Yu

Mentored by: Maxwell Fishelson

Online Learning

Online Learning



Online Learning

	A	B	C	D	E
1					

Online Learning

	A	B	C	D	E
1					

Online Learning

	A	B	C	D	E
1	0.1	0.3	0.9	0.5	0.7

Online Learning

	A	B	C	D	E
1	0.1	0.3	0.9	0.5	0.7
2					

Online Learning

	A	B	C	D	E
1	0.1	0.3	0.9	0.5	0.7
2					

Online Learning

	A	B	C	D	E
1	0.1	0.3	0.9	0.5	0.7
2	0.8	0.4	0	0.2	0.3

Online Learning

	A	B	C	D	E
1	0.1	0.3	0.9	0.5	0.7
2	0.8	0.4	0	0.2	0.3
3					

Online Learning

	A	B	C	D	E
1	0.1	0.3	0.9	0.5	0.7
2	0.8	0.4	0	0.2	0.3
3					

Online Learning

	A	B	C	D	E
1	0.1	0.3	0.9	0.5	0.7
2	0.8	0.4	0	0.2	0.3
3	0.1	0.2	0.3	0.4	0.5

Online Learning

	A	B	C	D	E
1	0.1	0.3	0.9	0.5	0.7
2	0.8	0.4	0	0.2	0.3
3	0.1	0.2	0.3	0.4	0.5

$$x^{(t)} \in \Delta(N)$$

Online Learning

	A	B	C	D	E
1	0.1	0.3	0.9	0.5	0.7
2	0.8	0.4	0	0.2	0.3
3	0.1	0.2	0.3	0.4	0.5

$$x^{(t)} \in \Delta(N)$$

$$u^{(t)} \in [0,1]^N$$

Online Learning

	A	B	C	D	E
1	0.1	0.3	0.9	0.5	0.7
2	0.8	0.4	0	0.2	0.3
3	0.1	0.2	0.3	0.4	0.5

$$x^{(t)} \in \Delta(N)$$

$$u^{(t)} \in [0,1]^N$$

$$\text{Total Reward} = \sum_{t=1}^T x^{(t)} \cdot u^{(t)}$$

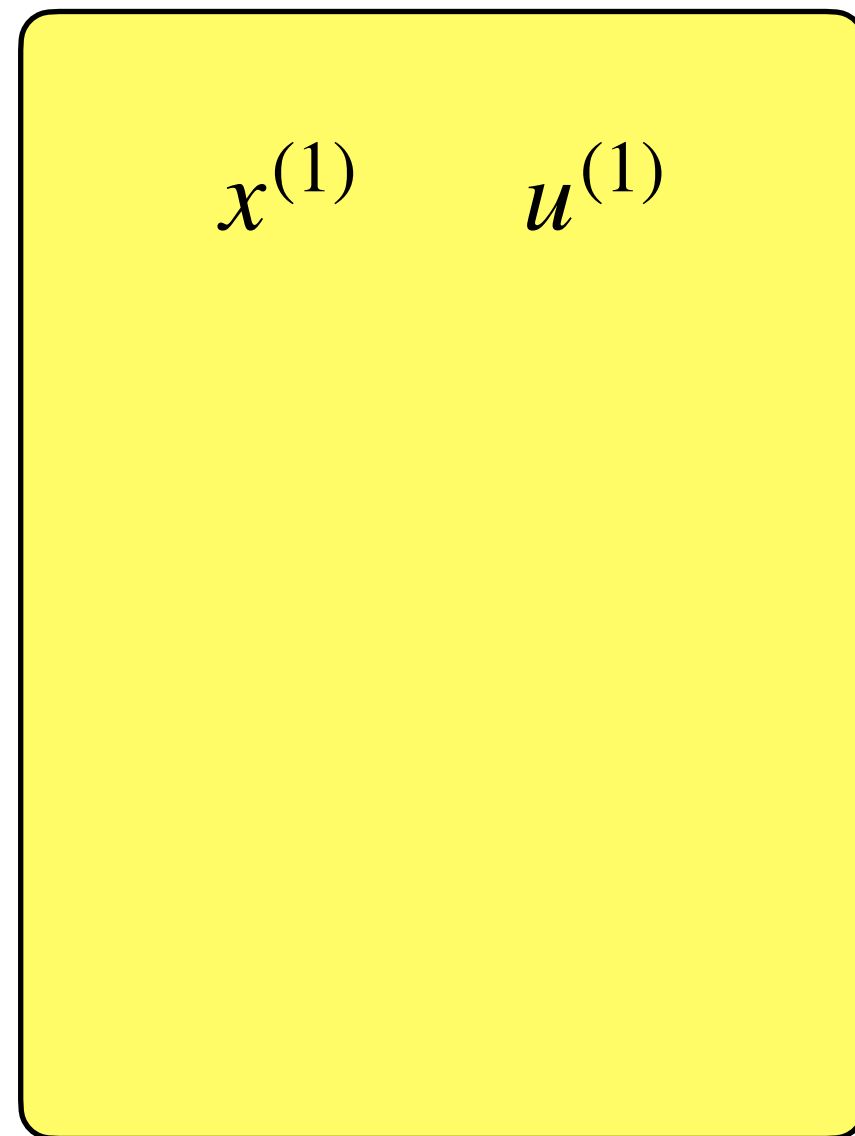
Regret

Regret

regret = reward obtained vs benchmark

Regret

regret = reward obtained vs benchmark



Regret

regret = reward obtained vs benchmark

$x^{(1)}$

$u^{(1)}$

$x^{(2)}$

$u^{(2)}$

Regret

regret = reward obtained vs benchmark

$x^{(1)}$ $u^{(1)}$

$x^{(2)}$ $u^{(2)}$

$x^{(3)}$ $u^{(3)}$

Regret

regret = reward obtained vs benchmark

$x^{(1)}$	$u^{(1)}$
$x^{(2)}$	$u^{(2)}$
$x^{(3)}$	$u^{(3)}$
$\vdots$	
$x^{(T)}$	$u^{(T)}$

Regret

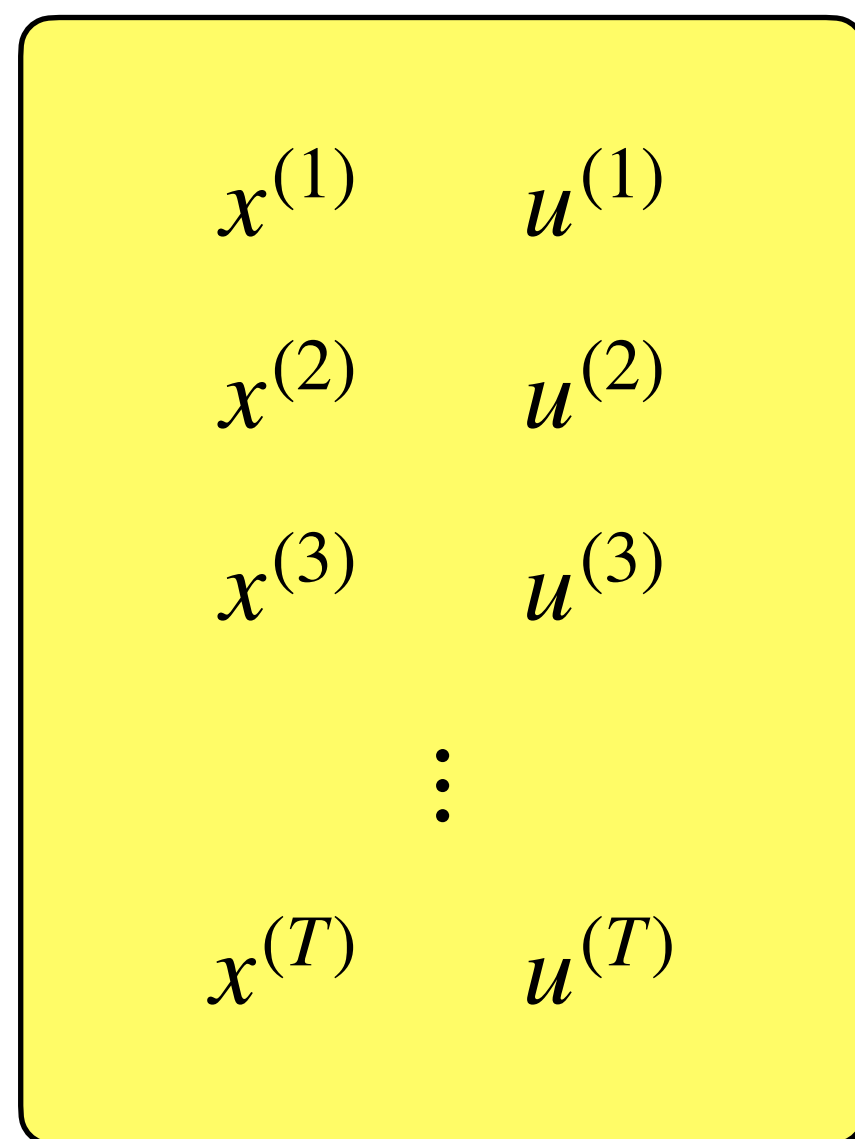
regret = reward obtained vs benchmark

$x^{(1)}$	$u^{(1)}$
$x^{(2)}$	$u^{(2)}$
$x^{(3)}$	$u^{(3)}$
$\vdots$	
$x^{(T)}$	$u^{(T)}$

$$\sum_{t=1}^T x^{(t)} \cdot u^{(t)}$$

Regret

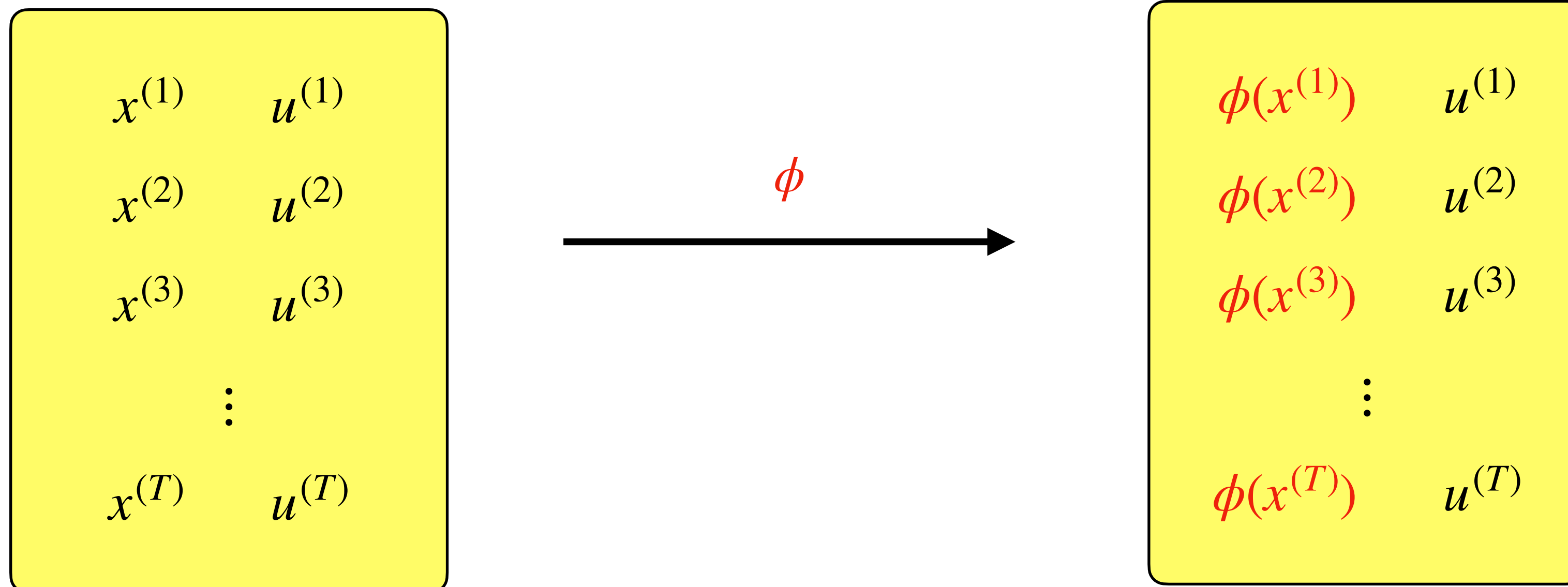
regret = reward obtained vs benchmark



$$\sum_{t=1}^T x^{(t)} \cdot u^{(t)}$$

Regret

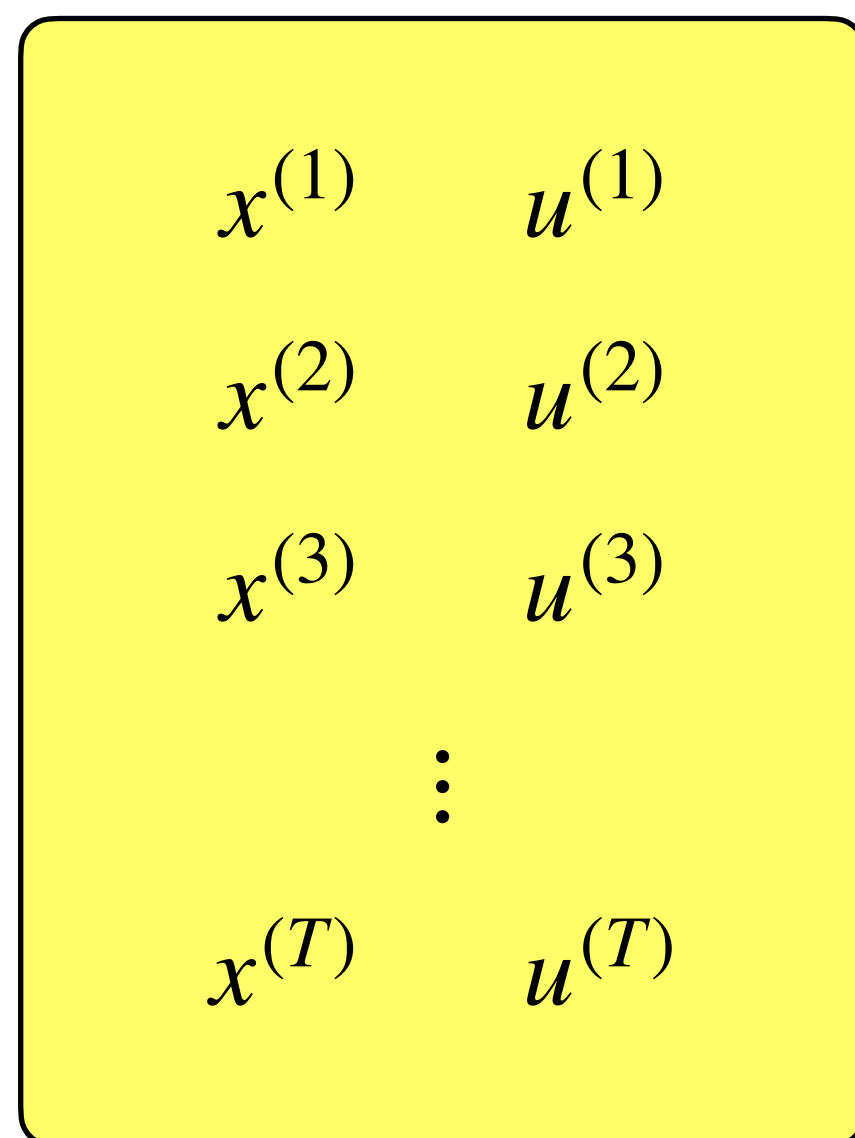
regret = reward obtained vs benchmark



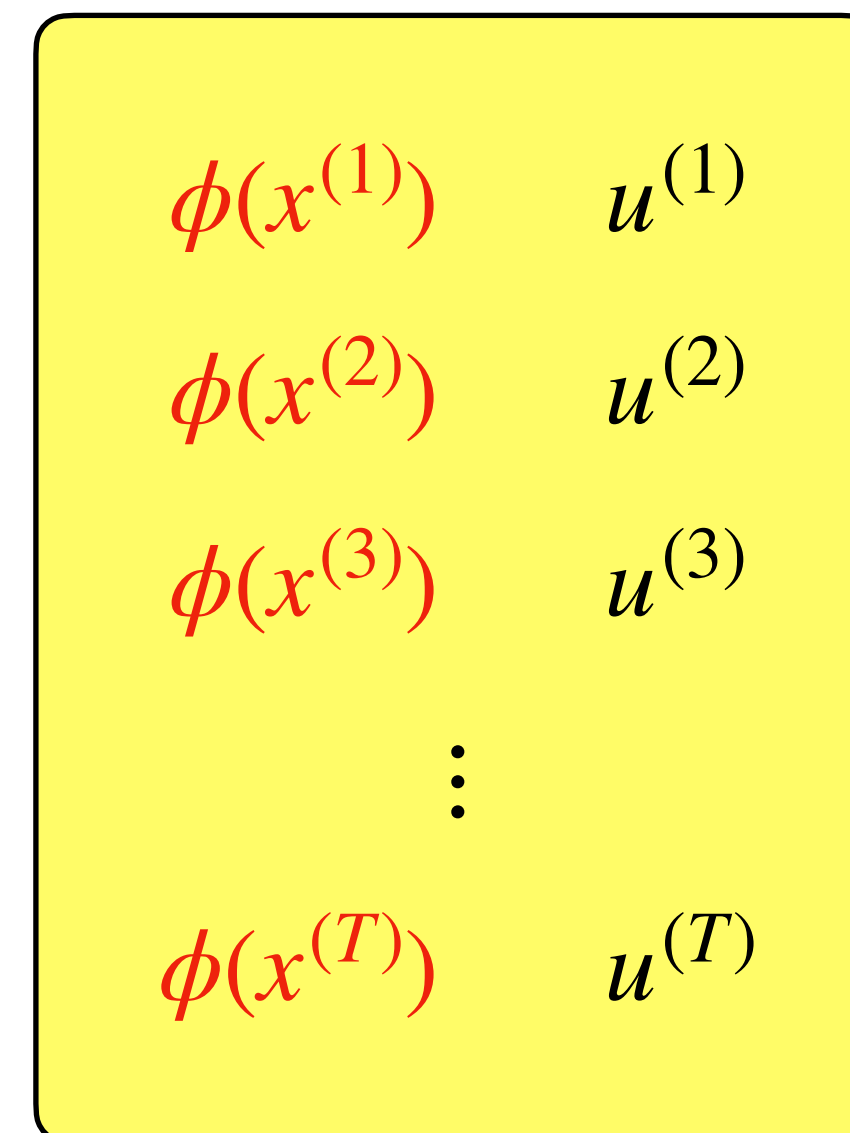
$$\sum_{t=1}^T x^{(t)} \cdot u^{(t)}$$

Regret

regret = reward obtained vs benchmark



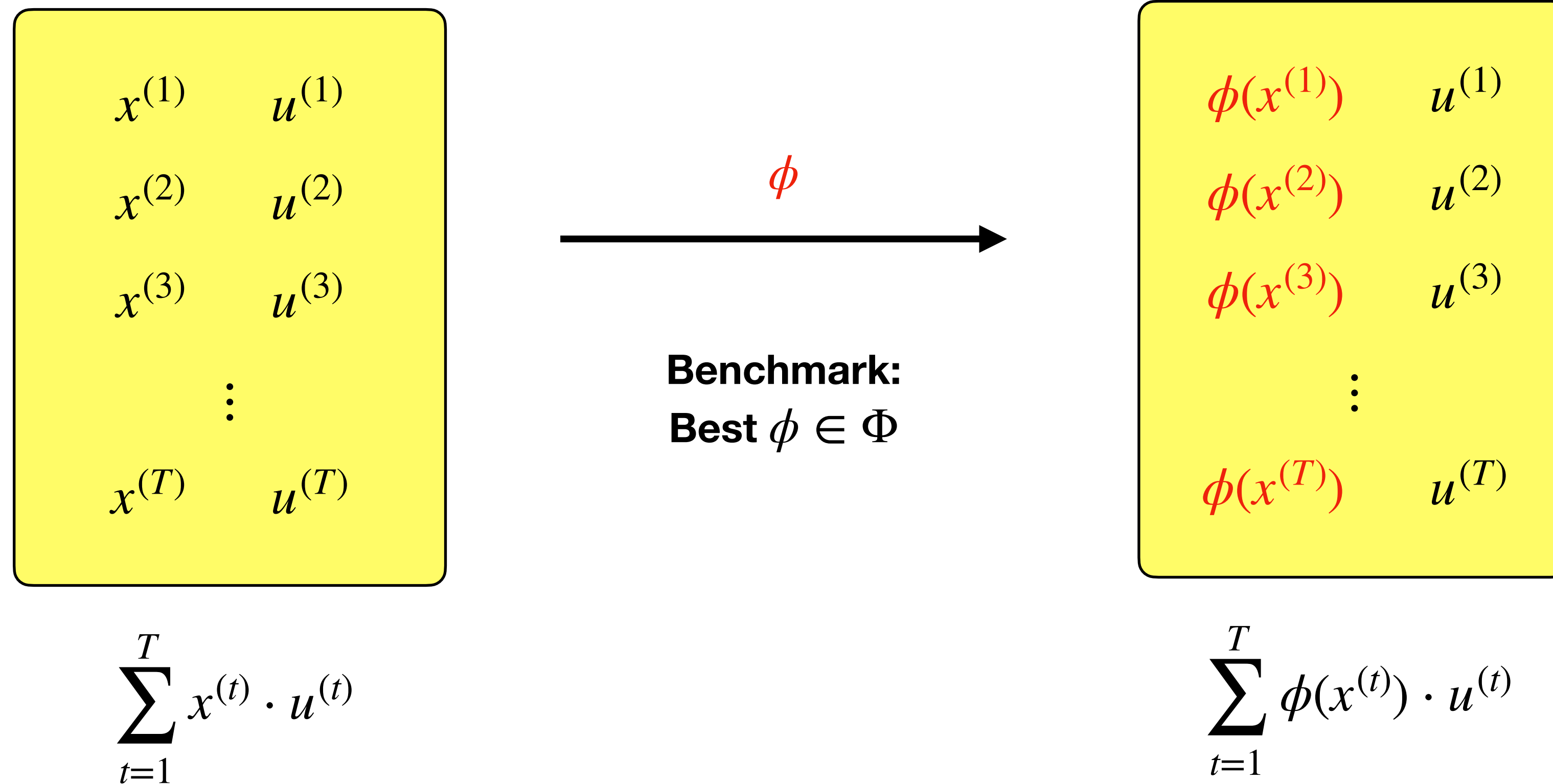
$$\sum_{t=1}^T x^{(t)} \cdot u^{(t)}$$



$$\sum_{t=1}^T \phi(x^{(t)}) \cdot u^{(t)}$$

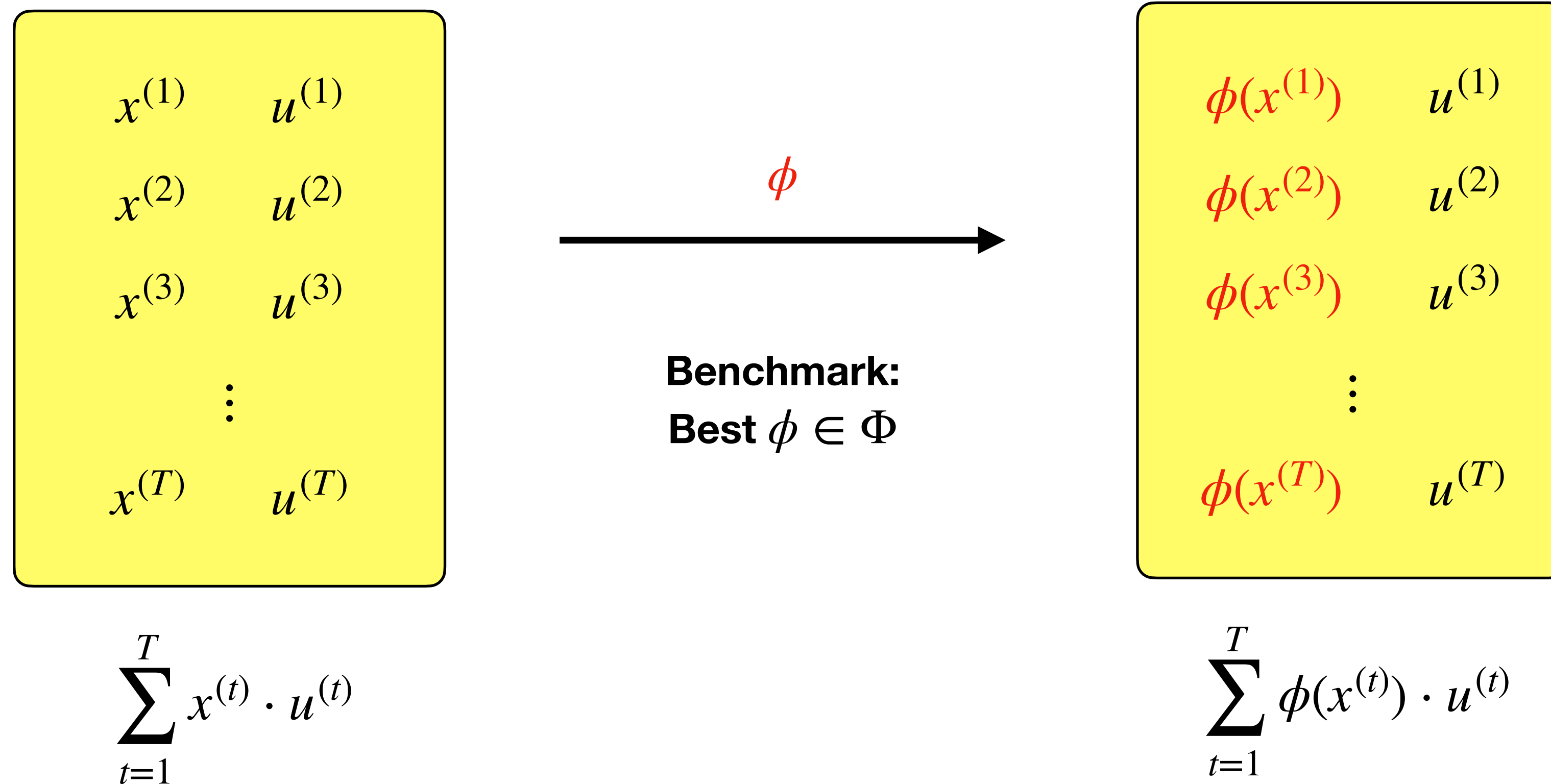
Regret

regret = reward obtained vs benchmark



Regret

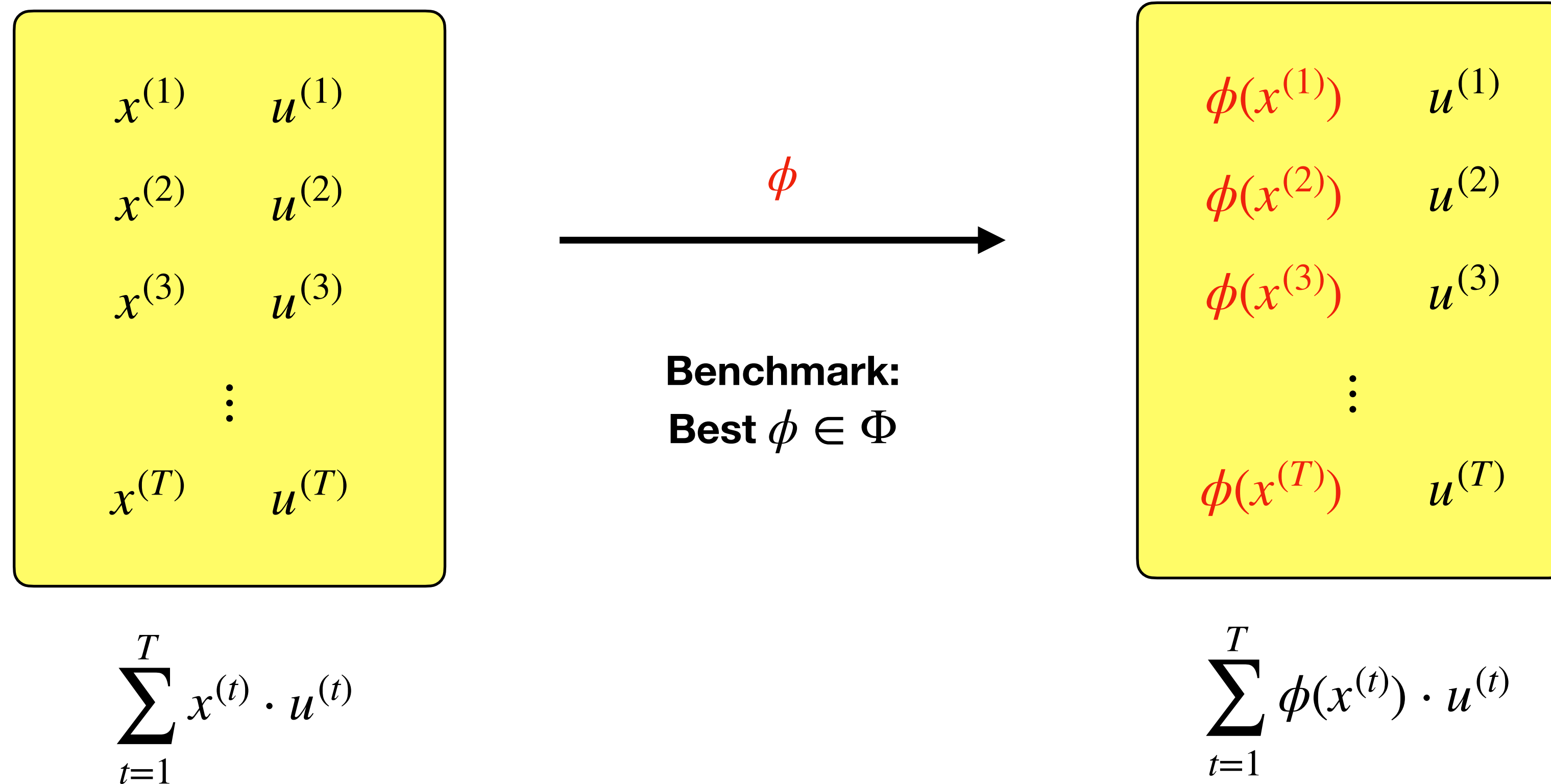
regret = reward obtained vs benchmark



$$\text{Regret} (x^{(1:T)}, u^{(1:T)}) = \max_{\phi \in \Phi} \sum_{t=1}^T \phi(x^{(t)}) \cdot u^{(t)} - \sum_{t=1}^T x^{(t)} \cdot u^{(t)}$$

Regret

regret = reward obtained vs benchmark



$$\text{Regret} (x^{(1:T)}, u^{(1:T)}) = \max_{\phi \in \Phi} \sum_{t=1}^T \phi(x^{(t)}) \cdot u^{(t)} - \sum_{t=1}^T x^{(t)} \cdot u^{(t)}$$

ϵ -Regret:
 $\leq \epsilon T$

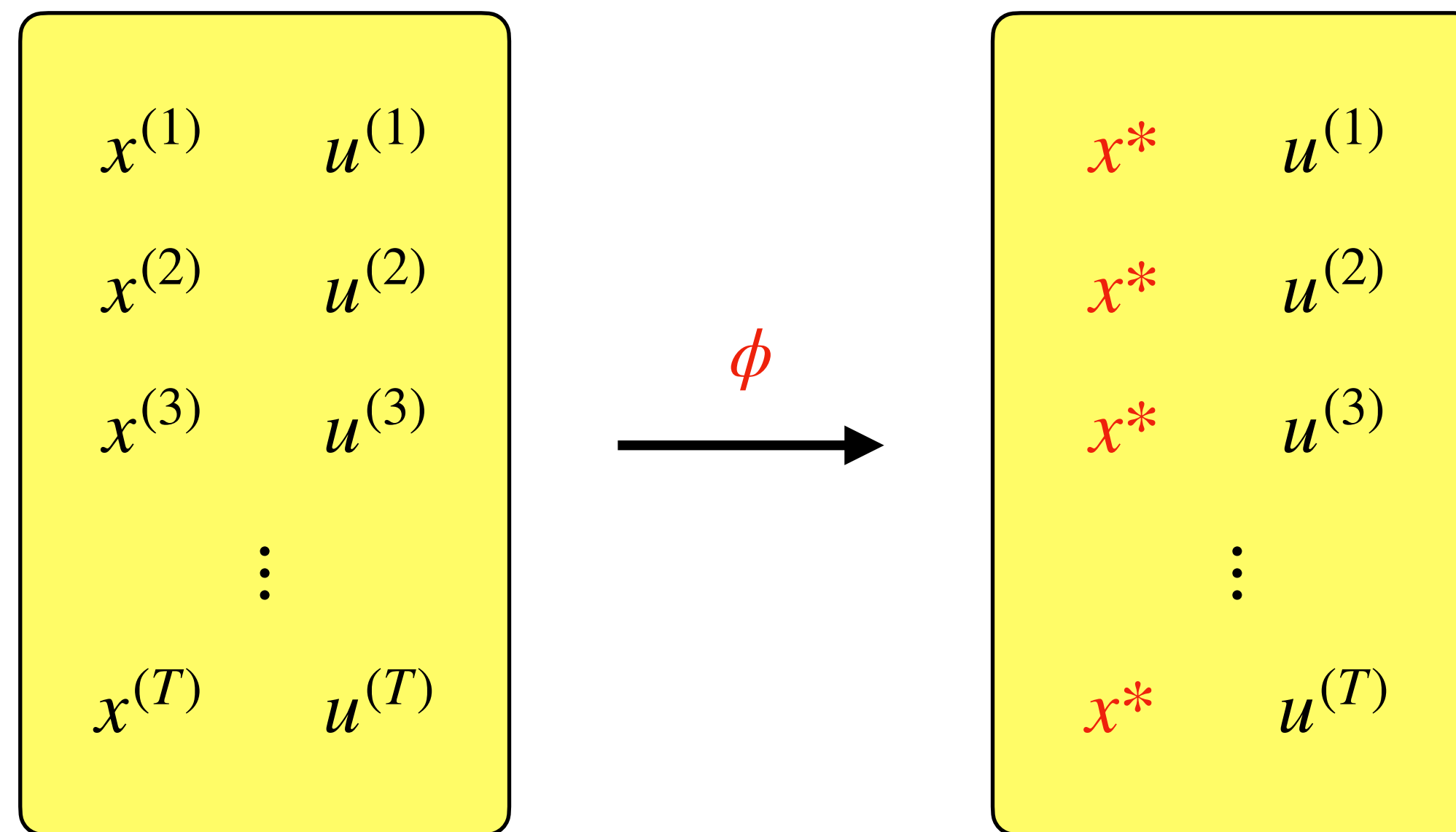
External Regret

External Regret

Φ = all constant maps

External Regret

Φ = all constant maps



Swap Regret

Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$

Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$

A $u^{(1)}$

B $u^{(2)}$

A $u^{(3)}$

$\vdots$

B $u^{(T)}$

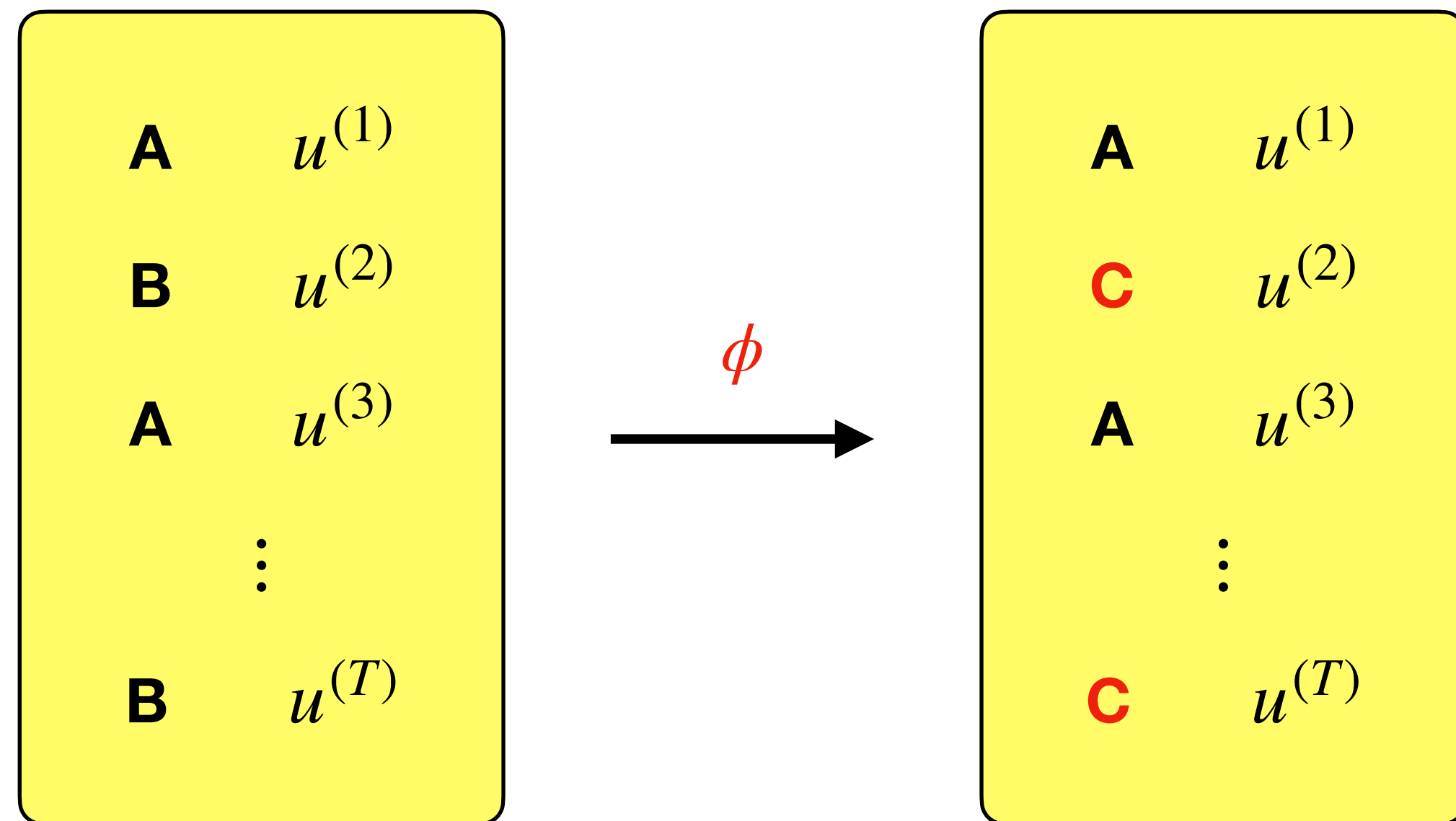
Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$

✓	A	$u^{(1)}$
✗	B	$u^{(2)}$
✓	A	$u^{(3)}$
		$\vdots$
✗	B	$u^{(T)}$

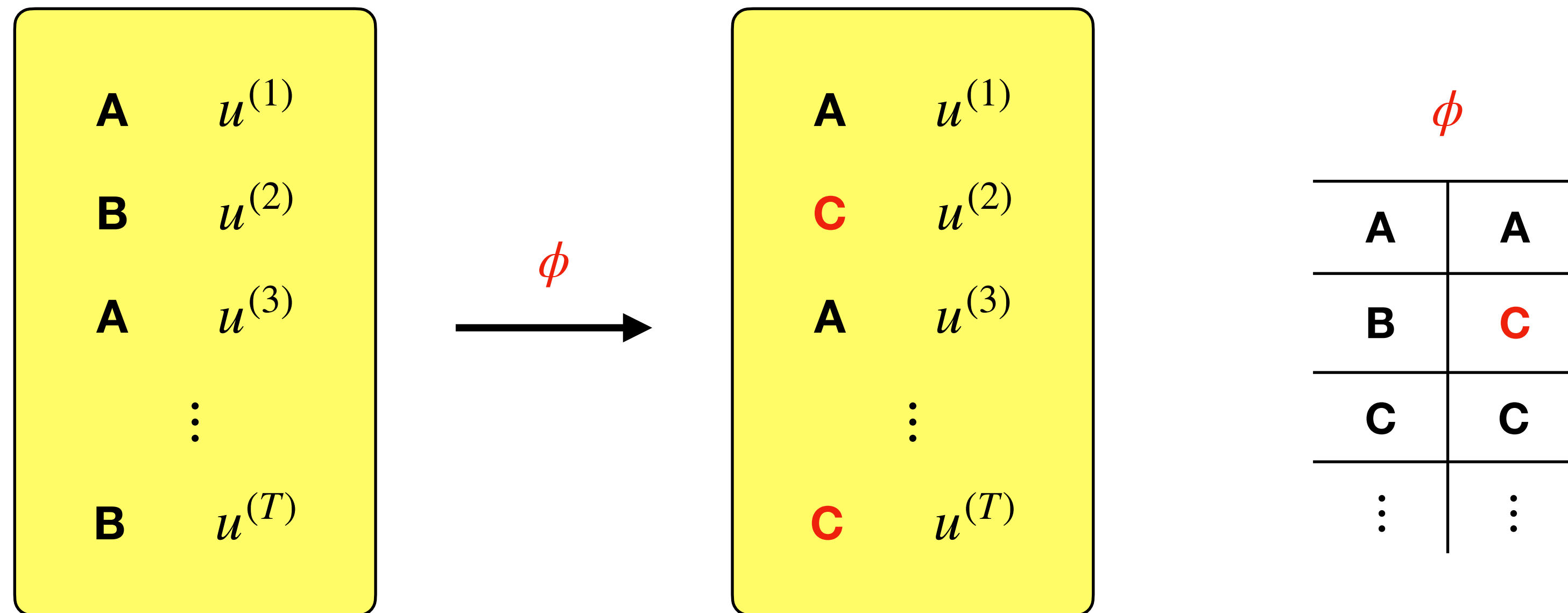
Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$



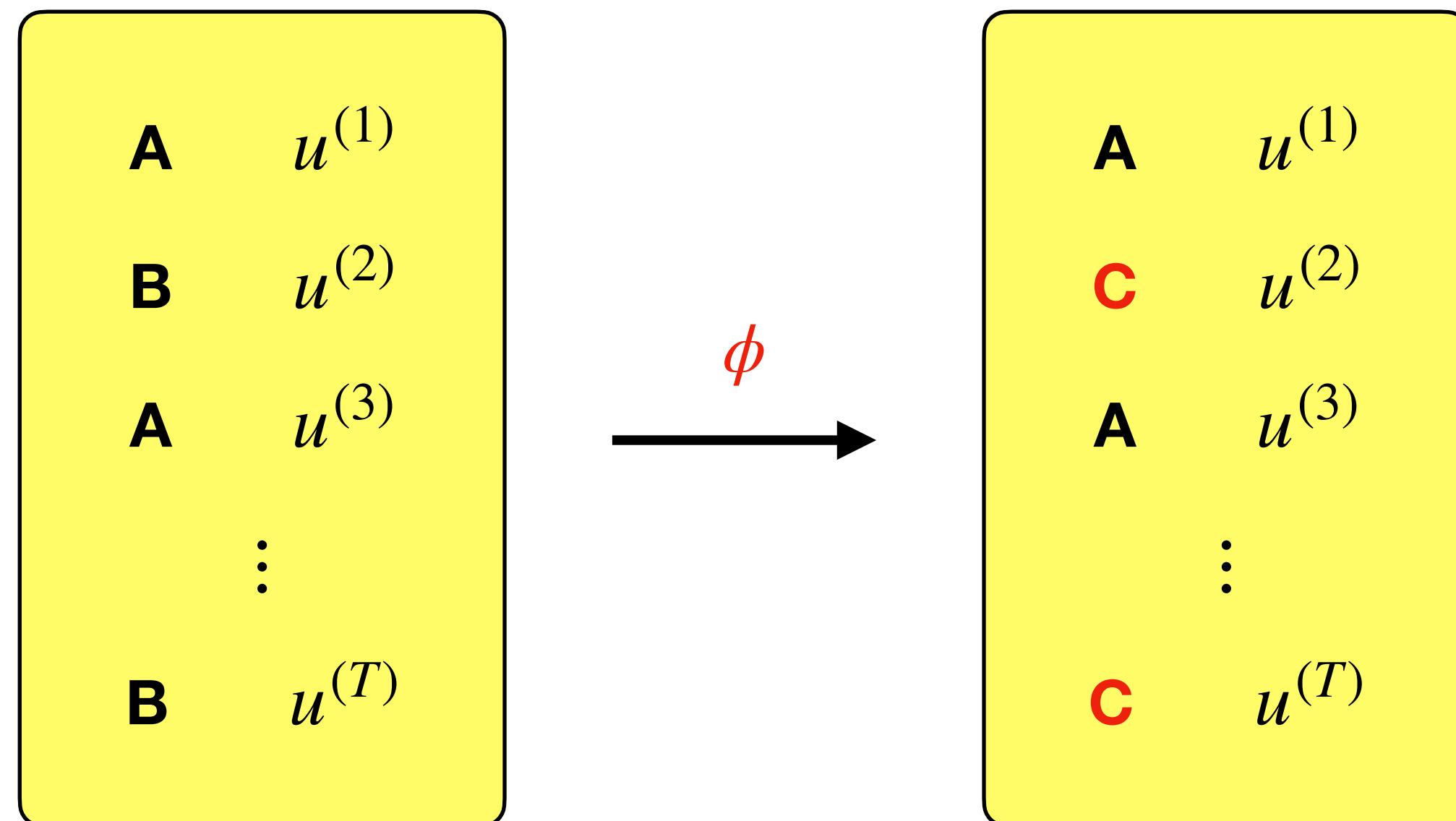
Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$



Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$



ϕ

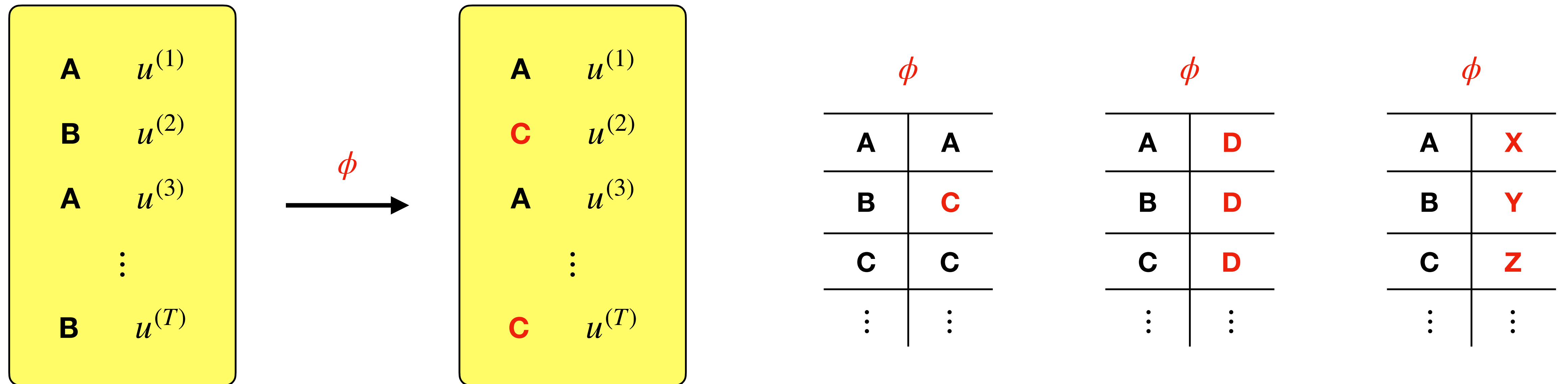
A	A
B	C
C	C
$\vdots$	$\vdots$

ϕ

A	D
B	D
C	D
$\vdots$	$\vdots$

Swap Regret

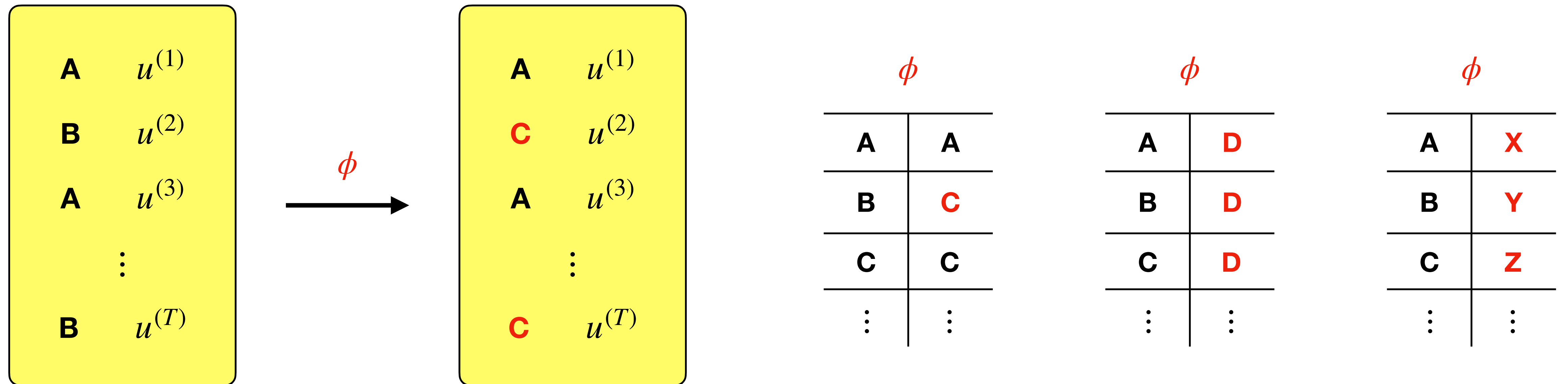
$\Phi = \text{all maps } [N] \rightarrow [N]$



Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$

$\Phi = \text{all linear maps } \Delta(N) \rightarrow \Delta(N)$

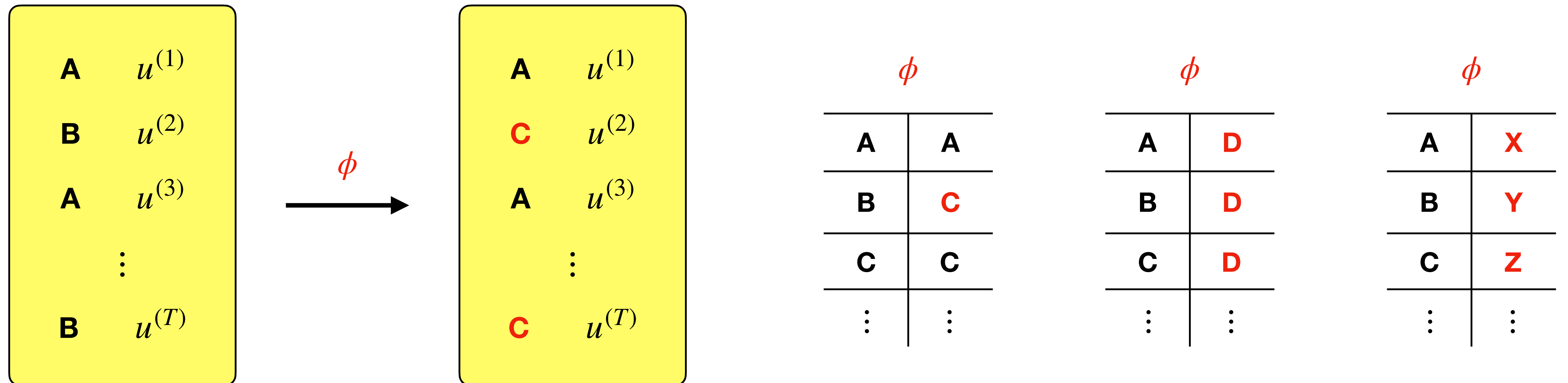


Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$

$\Phi = \text{all linear maps } \Delta(N) \rightarrow \Delta(N)$

$\Phi = \text{ALL maps } \Delta(N) \rightarrow \Delta(N)$

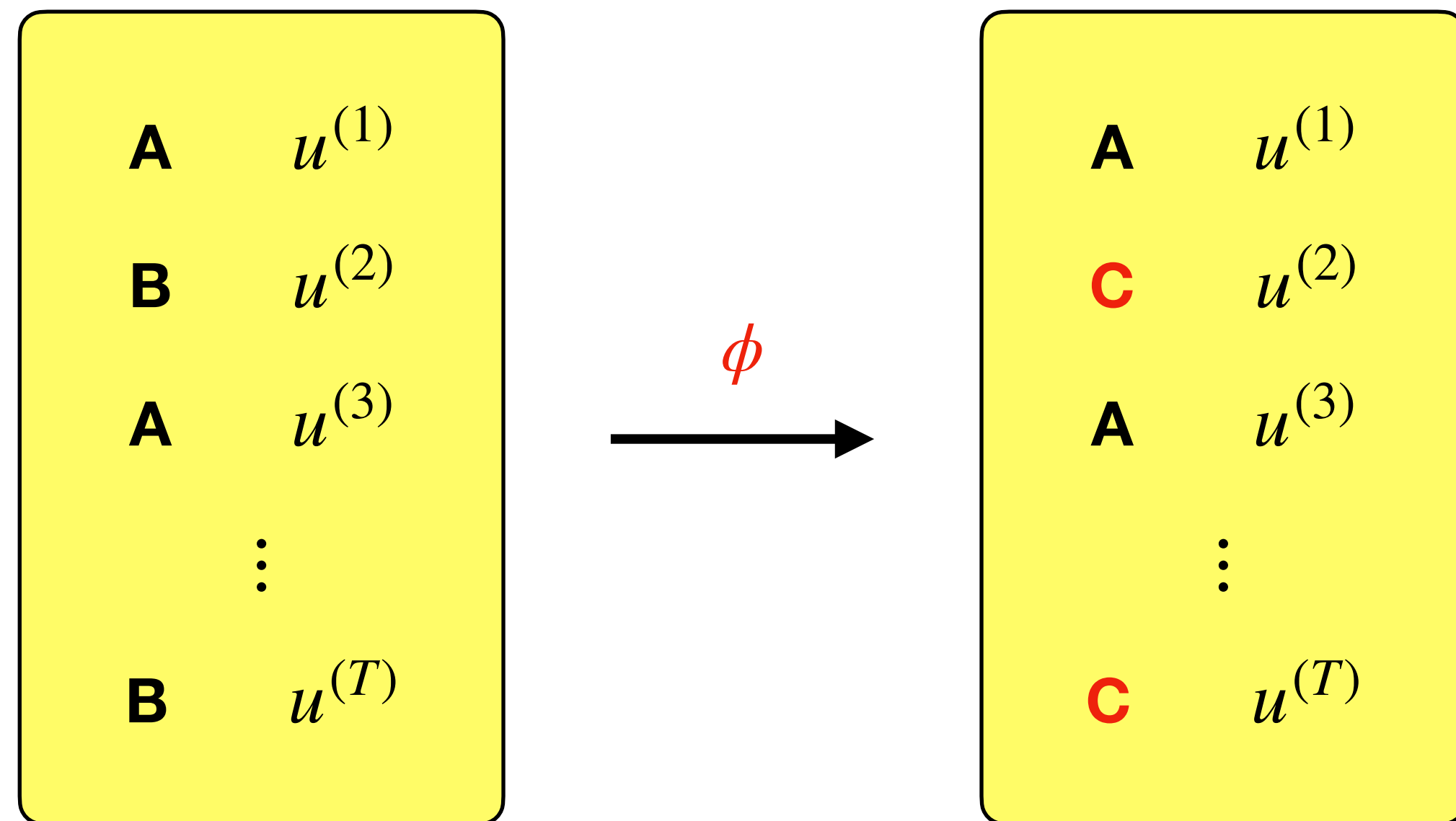


Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$

$\Phi = \text{all linear maps } \Delta(N) \rightarrow \Delta(N)$

$\Phi = \text{ALL maps } \Delta(N) \rightarrow \Delta(N)$

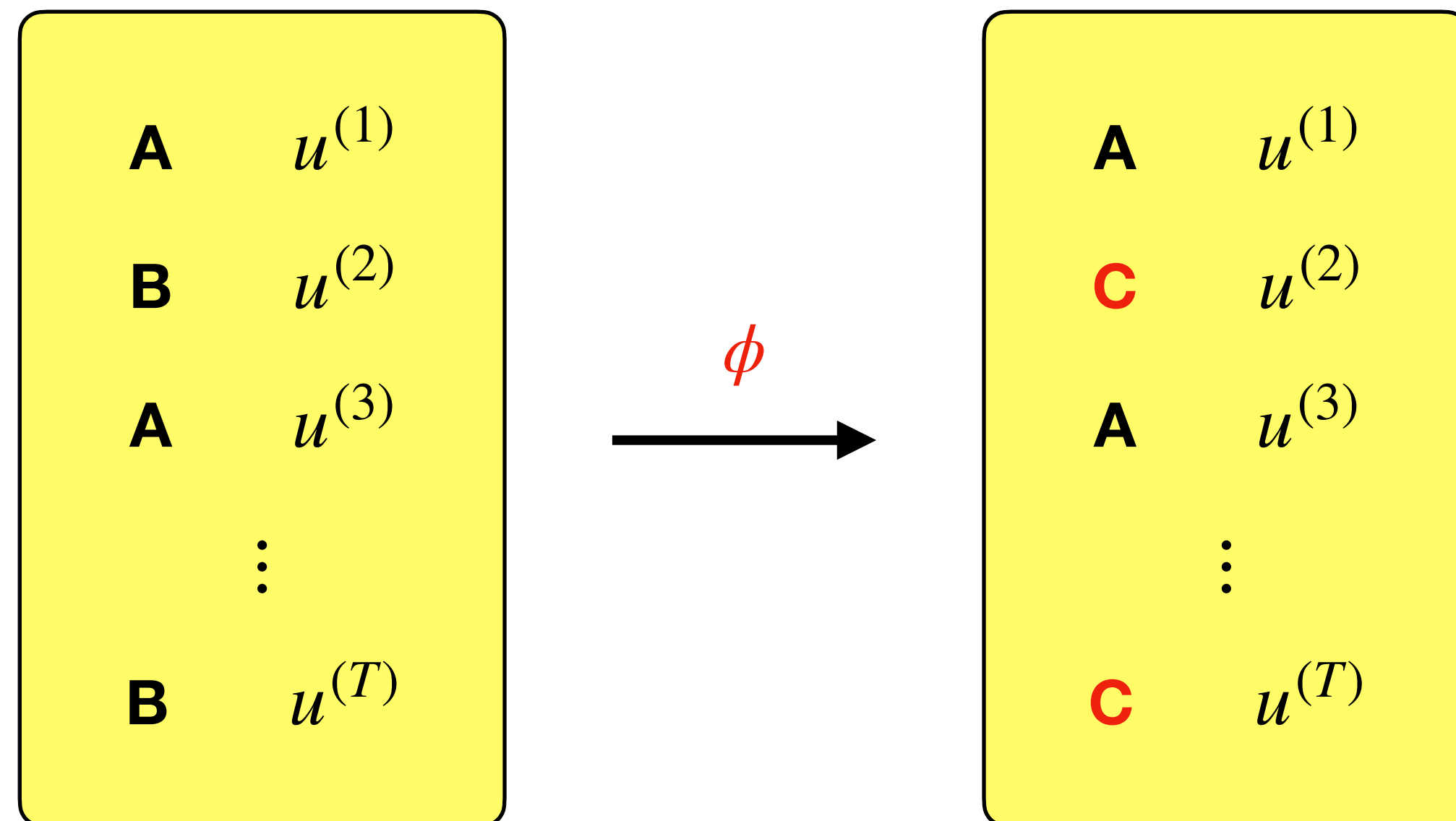


Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$

$\Phi = \text{all linear maps } \Delta(N) \rightarrow \Delta(N)$

$\Phi = \text{ALL maps } \Delta(N) \rightarrow \Delta(N)$



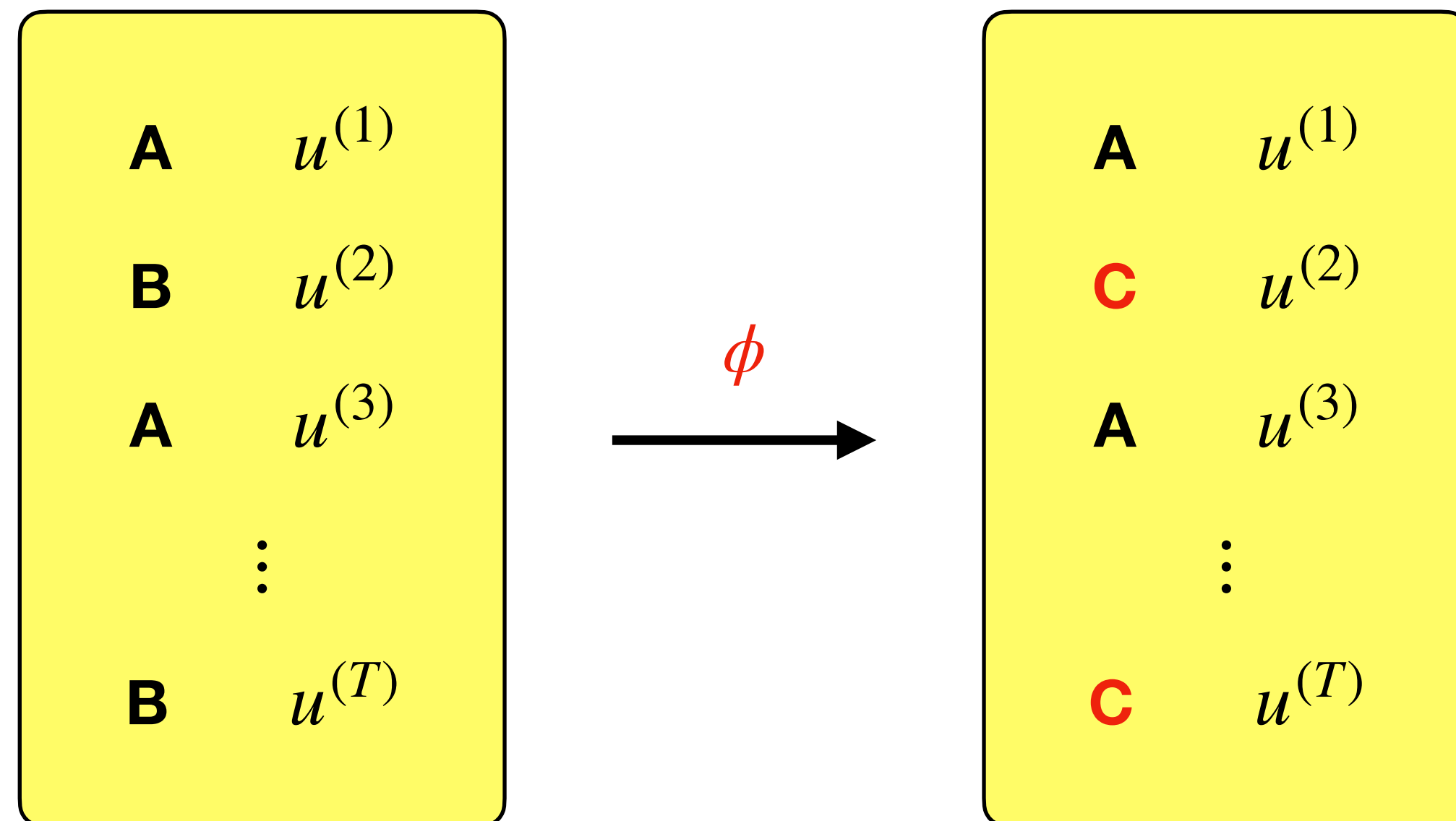
Why?

Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$

$\Phi = \text{all linear maps } \Delta(N) \rightarrow \Delta(N)$

$\Phi = \text{ALL maps } \Delta(N) \rightarrow \Delta(N)$



Why?

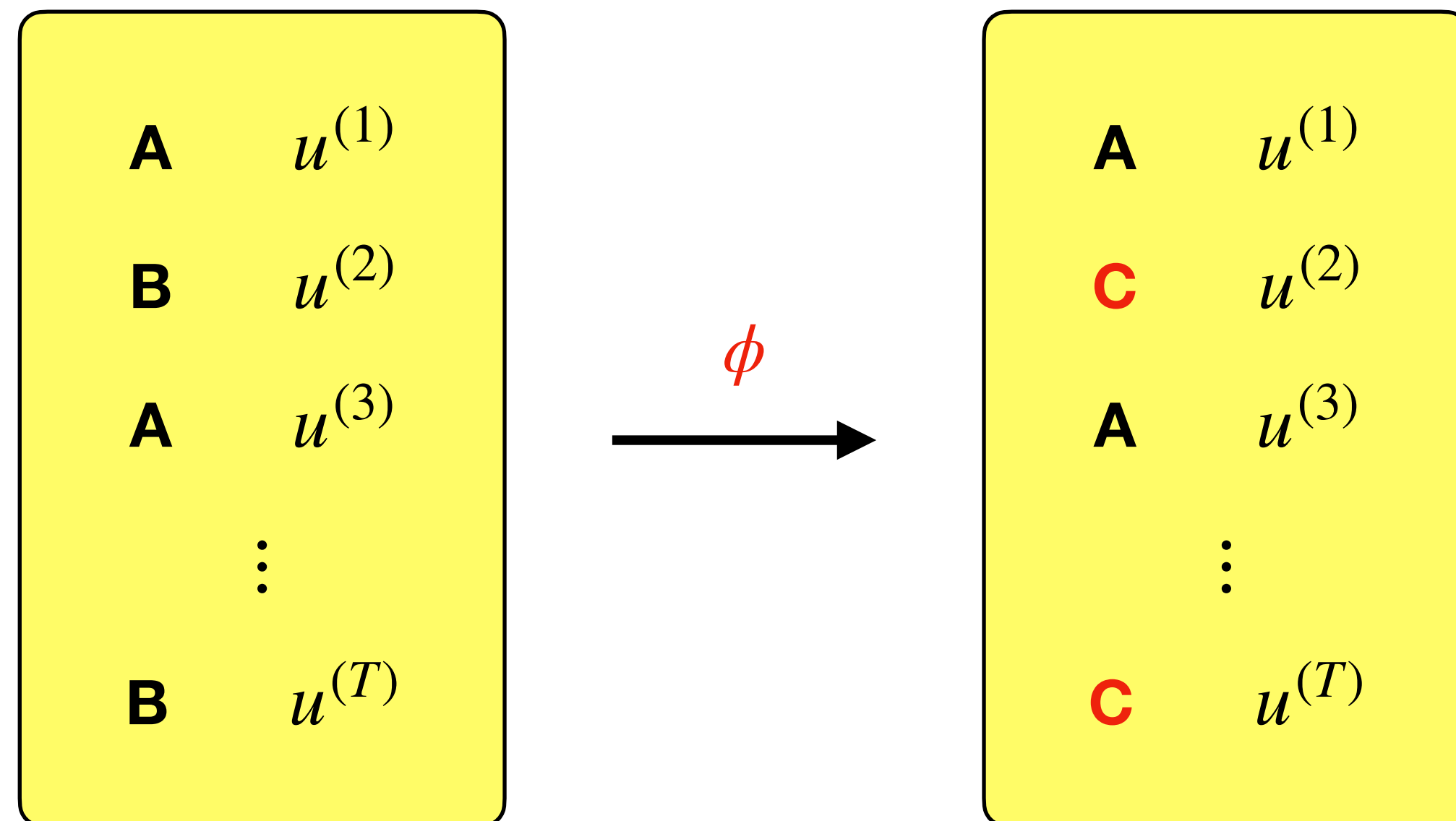
- Unexploitable

Swap Regret

$\Phi = \text{all maps } [N] \rightarrow [N]$

$\Phi = \text{all linear maps } \Delta(N) \rightarrow \Delta(N)$

$\Phi = \text{ALL maps } \Delta(N) \rightarrow \Delta(N)$



Why?

- Unexploitable
- Games!

Learning in Games

ϵ external regret over T rounds $\rightarrow$ average strategies are ϵ coarse-correlated-equilibrium

ϵ swap regret over T rounds $\rightarrow$ average strategies are ϵ correlated-equilibrium

Game Equilibrium: Strategy profile where no player wants to deviate

Coarse Deviations	
A	C
B	C

Swap Deviations	
A	A
B	C

CCE: No player wants to make a coarse deviation

So why not Swap Regret?

So why not Swap Regret?

Multiplicative Weight Updates: ϵ -External-Regret for $T = \Omega\left(\frac{\log N}{\epsilon^2}\right)$

So why not Swap Regret?

Multiplicative Weight Updates: ϵ -External-Regret for $T = \Omega\left(\frac{\log N}{\epsilon^2}\right)$

Blum-Mansour MWU [BM07]: ϵ -Swap-Regret for $T = \Omega\left(\frac{N \log N}{\epsilon^2}\right)$

So why not Swap Regret?

Multiplicative Weight Updates: ϵ -External-Regret for $T = \Omega\left(\frac{\log N}{\epsilon^2}\right)$

Blum-Mansour MWU [BM07]: ϵ -Swap-Regret for $T = \Omega\left(\frac{N \log N}{\epsilon^2}\right)$

Question: can we improve this?

So why not Swap Regret?

Blum Mansour task the learner with computing the stationary distribution of a Markov chain at each time step.

Instead, we propose estimating the stationary distribution through sampling

