

On the Adaptive Security of Free-XOR-based Garbling Schemes in the Plain Model

Anasuya Acharya¹, **Karen Azari**², Chethan Kamath³

1 - Aarhus University, Denmark

2 - University of Vienna, Faculty of Computer Science, Vienna, Austria

3 - IIT Bombay, India



AARHUS
UNIVERSITY



universität
wien



On the Adaptive Security of Free-XOR-based Garbling Schemes in the Plain Model — Contents

On the Adaptive Security of Free-XOR-based Garbling Schemes in the Plain Model — Contents

- **Garbling:**

- invented in [Yao86], modern practical schemes based on Yao's
- many theoretical and practical applications,
- central building block for threshold schemes (see NIST threshold call)

On the Adaptive Security of Free-XOR-based Garbling Schemes in the Plain Model — Contents

- **Garbling:**
 - invented in [Yao86], modern practical schemes based on Yao's
 - many theoretical and practical applications,
 - central building block for threshold schemes (see NIST threshold call)
- **Free-XOR:** Common technique for practical garbling (e.g. 'Half Gates' garbling [ZRE15])

On the Adaptive Security of Free-XOR-based Garbling Schemes in the Plain Model — Contents

- **Garbling:**
 - invented in [Yao86], modern practical schemes based on Yao's
 - many theoretical and practical applications,
 - central building block for threshold schemes (see NIST threshold call)
- **Free-XOR:** Common technique for practical garbling (e.g. 'Half Gates' garbling [ZRE15])
- **Adaptive security:** Required for offline precomputation of expensive circuit garbling

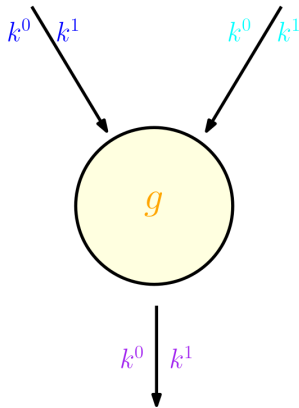
On the Adaptive Security of Free-XOR-based Garbling Schemes in the Plain Model — Contents

- **Garbling:**
 - invented in [Yao86], modern practical schemes based on Yao's
 - many theoretical and practical applications,
 - central building block for threshold schemes (see NIST threshold call)
- **Free-XOR:** Common technique for practical garbling (e.g. 'Half Gates' garbling [ZRE15])
- **Adaptive security:** Required for offline precomputation of expensive circuit garbling
- **Plain model:** ROM gives only heuristic security, focus on standard-model security

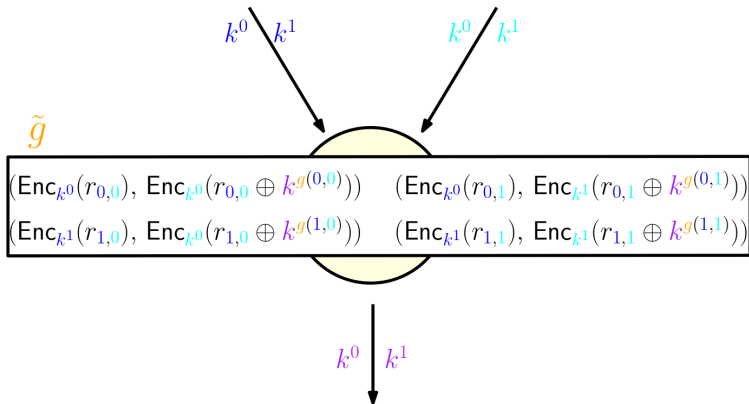
On the Adaptive Security of Free-XOR-based Garbling Schemes in the Plain Model — Contents

- **Garbling:**
 - invented in [Yao86], modern practical schemes based on Yao's
 - many theoretical and practical applications,
 - central building block for threshold schemes (see NIST threshold call)
- **Free-XOR:** Common technique for practical garbling (e.g. 'Half Gates' garbling [ZRE15])
- **Adaptive security:** Required for offline precomputation of expensive circuit garbling
- **Plain model:** ROM gives only heuristic security, focus on standard-model security
- **Our results:** **Limitations** on provable security of free-XOR based garbling in this setting (applies to [App16] and [ZRE15])

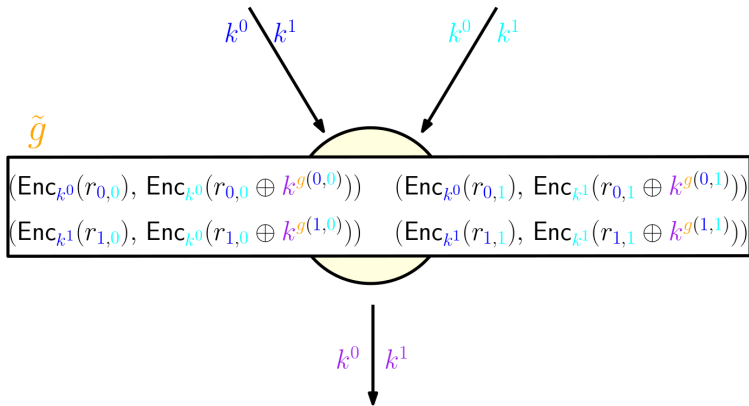
Yao's Garbling: gate-by-gate



Yao's Garbling: gate-by-gate

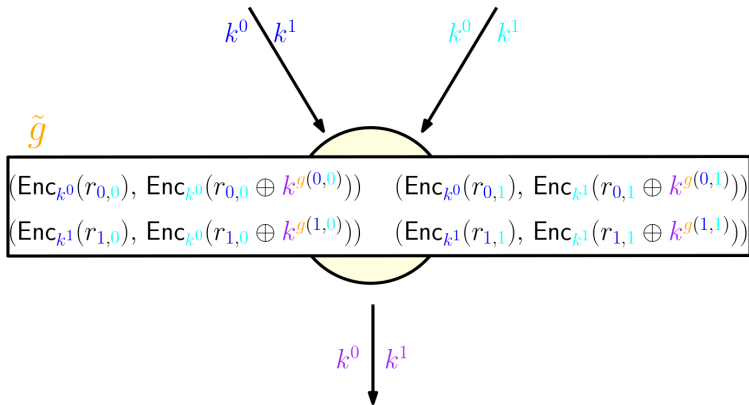


Yao's Garbling: gate-by-gate



$\tilde{C} = \{\tilde{g}\}_{g \in C}$, $K = \{k^0, k^1, k^0, k^1, \dots\}$ can be computed offline

Yao's Garbling: gate-by-gate

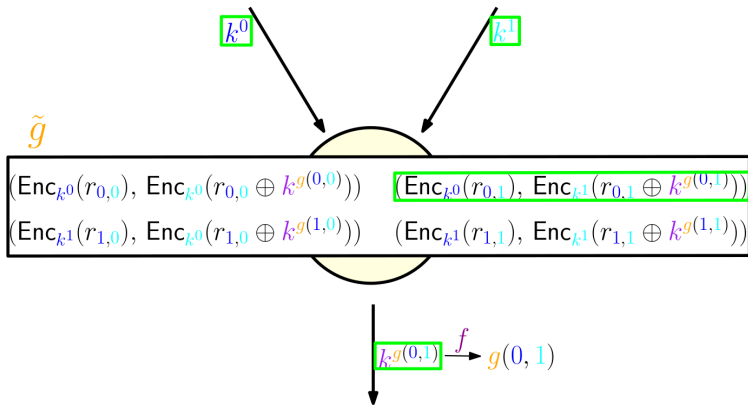


$\tilde{C} = \{\tilde{g}\}_{g \in C}$, $K = \{k^0, k^1, k^0, k^1, \dots\}$ can be computed offline

For $x = (x_1, x_2, \dots)$: $\tilde{x} = (k^{x_1}, k^{x_2}, \dots)$

Output mapping: $f = \{k^0 \rightarrow 0, k^1 \rightarrow 1, \dots\}$

Yao's Garbling: gate-by-gate

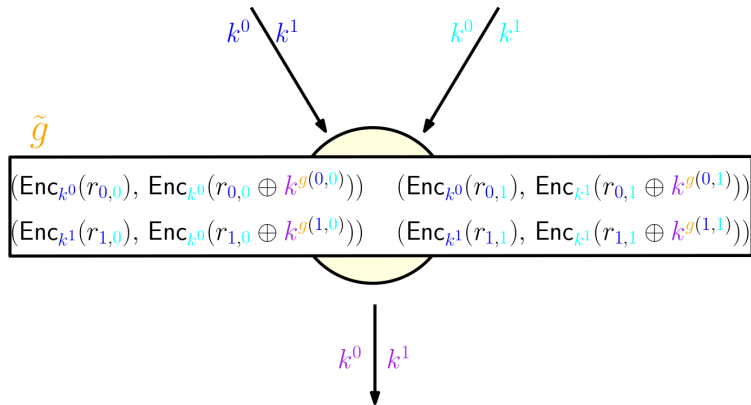


$\tilde{C} = \{\tilde{g}\}_{g \in C}$, $K = \{k^0, k^1, k^0, k^1, \dots\}$ can be computed offline

For $x = (x_1, x_2, \dots)$: $\tilde{x} = (k^{x_1}, k^{x_2}, \dots)$

Output mapping: $f = \{k^0 \rightarrow 0, k^1 \rightarrow 1, \dots\}$

Optimization for public C: free XOR

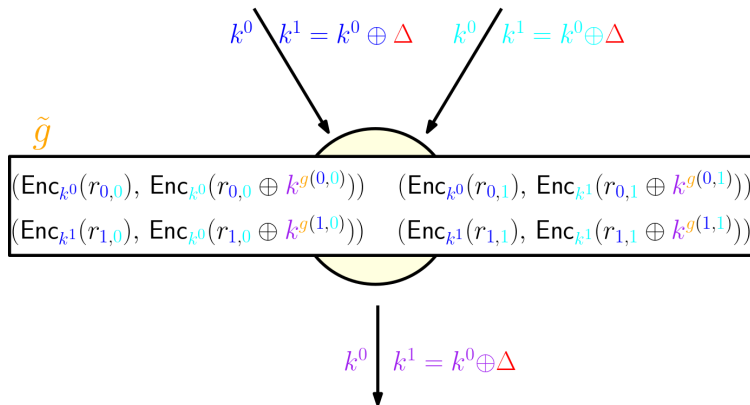


$\tilde{C} = \{\tilde{g}\}_{g \in C}$, $K = \{k^0, k^1, k^0, k^1, \dots\}$ can be computed offline

For $x = (x_1, x_2, \dots)$: $\tilde{x} = (k^{x_1}, k^{x_2}, \dots)$

Output mapping: $f = \{k^0 \rightarrow 0, k^1 \rightarrow 1, \dots\}$

Optimization for public C: free XOR

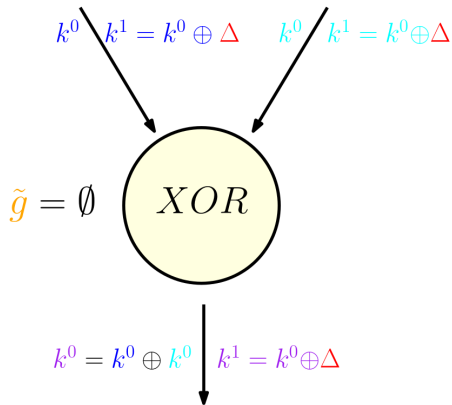


$\tilde{C} = \{\tilde{g}\}_{g \in C}$, $K = \{k^0, k^1, k^0, k^1, \dots\}$ can be computed offline

For $x = (x_1, x_2, \dots)$: $\tilde{x} = (k^{x_1}, k^{x_2}, \dots)$

Output mapping: $f = \{k^0 \rightarrow 0, k^1 \rightarrow 1, \dots\}$

Optimization for public C: free XOR

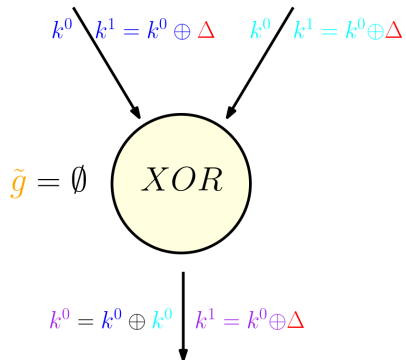


$\tilde{C} = \{\tilde{g}\}_{g \in C}$, $K = \{k^0, k^1, k^0, k^1, \dots\}$ can be computed offline

For $x = (x_1, x_2, \dots)$: $\tilde{x} = (k^{x_1}, k^{x_2}, \dots)$

Output mapping: $f = \{k^0 \rightarrow 0, k^1 \rightarrow 1, \dots\}$

Optimization for public C: free XOR



$\tilde{C} = \{\tilde{g}\}_{g \in C}$, $K = \{k^0, k^1, k^0, k^1, \dots\}$ can be computed offline

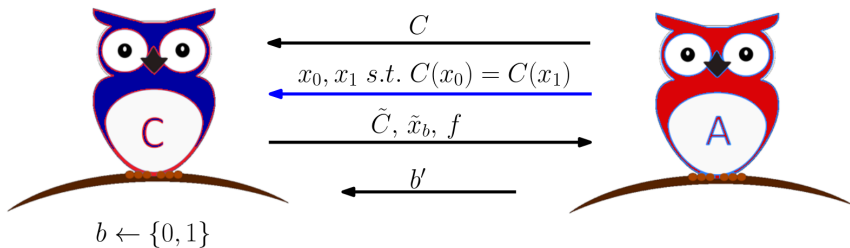
For $x = (x_1, x_2, \dots)$: $\tilde{x} = (k^{x_1}, k^{x_2}, \dots)$

Output mapping: $f = \{k^0 \rightarrow 0, k^1 \rightarrow 1, \dots\}$

Instantiations of free-XOR: [App16], “Half Gates” scheme [ZRE15]

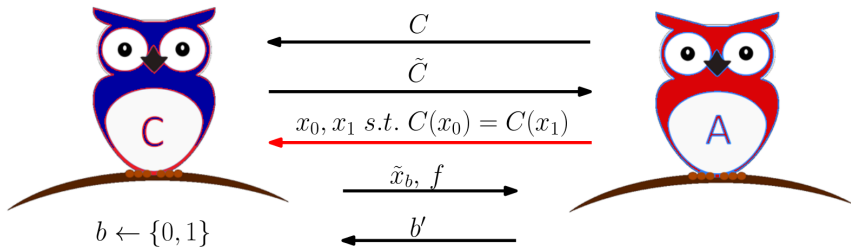
Security definition for Garbling

selective indistinguishability
(weaker than simulation-based security)



Security definition for Garbling

adaptive indistinguishability
(weaker than simulation-based security)



Required for offline precomputation

Security of Yao's Garbling

- LP09: **selective** security proof based on IND-CPA security of SKE
⇒ **adaptive** security via **guessing the input** of length n :

SKE ϵ -IND-CPA secure $\Rightarrow$ Yao's scheme $2^n \cdot \epsilon$ -secure

Security of Yao's Garbling

- LP09: **selective** security proof based on IND-CPA security of SKE
⇒ **adaptive** security via **guessing the input** of length n :

SKE ϵ -IND-CPA secure $\Rightarrow$ Yao's scheme $2^n \cdot \epsilon$ -secure

- JW16: **adaptive** security proof for circuits of **depth** D via “pebbling”:

SKE ϵ -IND-CPA secure $\Rightarrow$ Yao's scheme $2^D \cdot \epsilon$ -secure

Security of Yao's Garbling

- LP09: **selective** security proof based on IND-CPA security of SKE
⇒ **adaptive** security via **guessing the input** of length n :

SKE ϵ -IND-CPA secure $\Rightarrow$ Yao's scheme $2^n \cdot \epsilon$ -secure

- JW16: **adaptive** security proof for circuits of **depth** D via “pebbling”:

SKE ϵ -IND-CPA secure $\Rightarrow$ Yao's scheme $2^D \cdot \epsilon$ -secure

- KKPW21: Any **black-box proof** of **adaptive** security for Yao's garbling scheme for circuits with n -bit input and depth $D \leq 2n$ based on **IND-CPA secure SKE** incurs a security loss $2^{\Omega(\sqrt{D})}$.

Security of free-XOR

- App16: **selective** security proof based on **LIN-RK-KDM secure SKE**

(LIN-RK-KDM: Related-Key Key-Dependent-Message security under LINear relations)

Security of free-XOR

- App16: **selective** security proof based on **LIN-RK-KDM secure SKE**
(LIN-RK-KDM: Related-Key Key-Dependent-Message security under LINear relations)
- ZRE15: **selective** security proof for “Half Gates” based on **CCR secure hashing**
[CKKZ12] (CCR for hash functions $\approx$ LIN-RK-KDM for SKE)

Security of free-XOR

- App16: **selective** security proof based on **LIN-RK-KDM secure SKE**
(LIN-RK-KDM: Related-Key Key-Dependent-Message security under LINear relations)
- ZRE15: **selective** security proof for “Half Gates” based on **CCR secure hashing**
[CKKZ12] (CCR for hash functions $\approx$ LIN-RK-KDM for SKE)
- JO20: Pebbling techniques from JW16 likely **not** useful for **adaptive** security of free-XOR
in the standard model (due to **global** offset Δ)

Security of free-XOR

- App16: **selective** security proof based on **LIN-RK-KDM secure SKE**
(LIN-RK-KDM: Related-Key Key-Dependent-Message security under LINear relations)
- ZRE15: **selective** security proof for “Half Gates” based on **CCR secure hashing**
[CKKZ12] (CCR for hash functions $\approx$ LIN-RK-KDM for SKE)
- JO20: Pebbling techniques from JW16 likely **not** useful for **adaptive** security of free-XOR
in the standard model (due to **global** offset Δ)
- GYWYL23, BHKO23: **adaptive** security proof for “Half Gates” in the **ROM/RPM**

Security of free-XOR

- App16: **selective** security proof based on **LIN-RK-KDM** secure **SKE**
(LIN-RK-KDM: Related-Key Key-Dependent-Message security under LINear relations)
- ZRE15: **selective** security proof for “Half Gates” based on **CCR** secure hashing
[CKKZ12] (CCR for hash functions $\approx$ LIN-RK-KDM for SKE)
- JO20: Pebbling techniques from JW16 likely **not** useful for **adaptive** security of free-XOR
in the standard model (due to **global** offset Δ)
- GYWYL23, BHKO23: **adaptive** security proof for “Half Gates” in the **ROM/RPM**

Theorem (Our results (informal))

Any **black-box** proof of **adaptive** security for **free-XOR** based on **LIN-RK-KDM** secure **SKE** incurs an **exponential** security loss.

Discussion of our Results

Theorem (Our results (informal))

Any **black-box proof** of **adaptive** security for **free-XOR** based on **LIN-RK-KDM** secure **SKE** incurs an **exponential** security loss.

- applies even to **NC1 circuits**

($\neq$ Yao's scheme, i.e. proves JO20 right)

Discussion of our Results

Theorem (Our results (informal))

Any **black-box proof** of **adaptive** security for **free-XOR** based on **LIN-RK-KDM** secure **SKE** incurs an **exponential** security loss.

- applies even to **NC1 circuits** ($\neq$ Yao's scheme, i.e. proves JO20 right)
- holds for **indistinguishability** (weaker security than simulatability) and when output map f is sent *online* (AIKW13 doesn't apply here!)

Discussion of our Results

Theorem (Our results (informal))

Any **black-box proof** of **adaptive** security for **free-XOR** based on **LIN-RK-KDM** secure **SKE** incurs an **exponential** security loss.

- applies even to **NC1 circuits** ($\neq$ Yao's scheme, i.e. proves JO20 right)
- holds for **indistinguishability** (weaker security than simulatability) and when output map f is sent *online* (AIKW13 doesn't apply here!)
- holds also for **"Half Gates"** based on **CCR secure hash function**

Discussion of our Results

Theorem (Our results (informal))

Any **black-box proof** of **adaptive** security for **free-XOR** based on **LIN-RK-KDM** secure **SKE** incurs an **exponential** security loss.

- applies even to **NC1 circuits** ($\neq$ Yao's scheme, i.e. proves JO20 right)
- holds for **indistinguishability** (weaker security than simulatability) and when output map f is sent *online* (AIKW13 doesn't apply here!)
- holds also for **"Half Gates"** based on **CCR secure hash function**

$\Rightarrow$ **adaptive** security via **random guessing** essentially best we can do!

SKE ε -LIN-RK-KDM secure $\Rightarrow$ free-XOR scheme $2^n \cdot \varepsilon$ -secure

Discussion of our Results

Theorem (Our results (informal))

Any **black-box proof** of **adaptive** security for **free-XOR** based on **LIN-RK-KDM** secure **SKE** incurs an **exponential** security loss.

- applies even to **NC1 circuits** ($\neq$ Yao's scheme, i.e. proves JO20 right)
- holds for **indistinguishability** (weaker security than simulatability) and when output map f is sent *online* (AIKW13 doesn't apply here!)
- holds also for **"Half Gates"** based on **CCR secure hash function**

$\Rightarrow$ **adaptive** security via **random guessing** essentially best we can do!

$\text{SKE } \varepsilon\text{-LIN-RK-KDM secure} \Rightarrow \text{free-XOR scheme } 2^n \cdot \varepsilon\text{-secure}$

(applies only to *black-box* proofs for *specific* constructions from *specific* assumptions)

Proof Idea: oracle separation

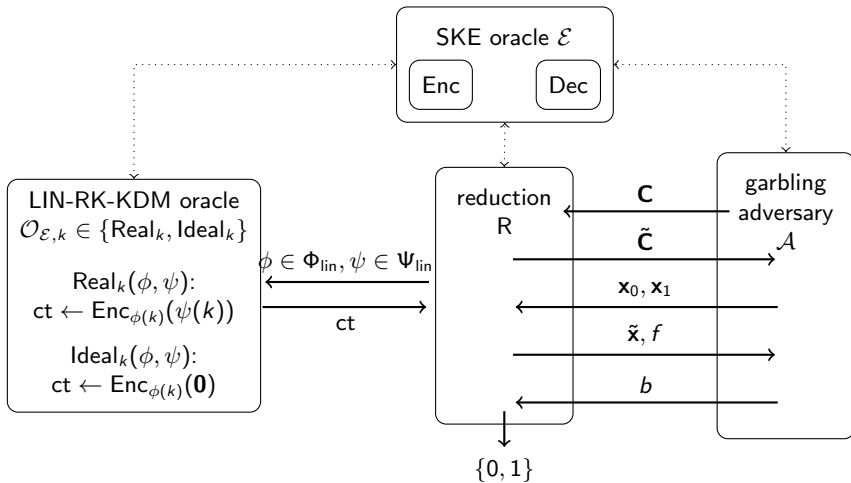
Theorem (Our results (informal))

Any **black-box proof** of **adaptive** security for **free-XOR** based on **LIN-RK-KDM** secure **SKE** incurs an **exponential** security loss.

Define oracles $\mathcal{E}$ and $\mathcal{A}$ such that

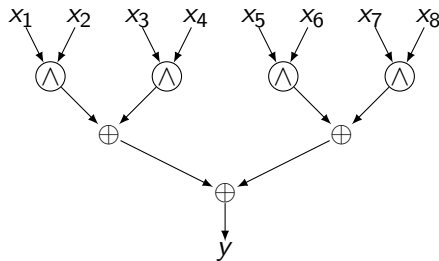
- $\mathcal{E} = (\text{Enc}, \text{Dec})$ is an **ideal SKE scheme**
- $\mathcal{A}$ is an (inefficient) **adversary** breaking adaptive security of the free-XOR scheme, but “not too helpful” in breaking $\mathcal{E}$.

Proof Idea: oracle separation



Proof Idea: the Adversary $\mathcal{A}$

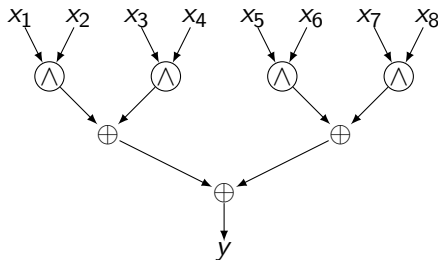
- Send log-depth circuit C :



- Receive $\tilde{C}$

Proof Idea: the Adversary $\mathcal{A}$

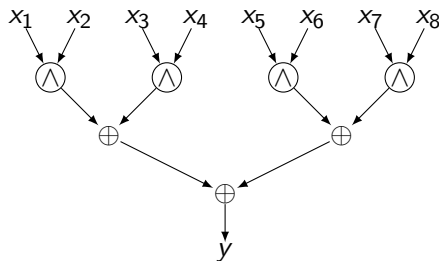
- Send log-depth circuit C :



- Receive $\tilde{C}$
- Sample $\mathbf{x}_1 \leftarrow \{0, 1\}^n$, $\mathbf{x}_0 \leftarrow \{\mathbf{x}_0 \in \{0, 1\}^n \mid C(\mathbf{x}_0) = C(\mathbf{x}_1)\}$
- Receive $\tilde{\mathbf{x}}$

Proof Idea: the Adversary $\mathcal{A}$

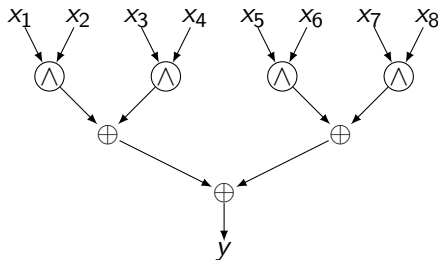
- Send log-depth circuit C :



- Receive $\tilde{C}$
- Sample $\mathbf{x}_1 \leftarrow \{0, 1\}^n$, $\mathbf{x}_0 \leftarrow \{\mathbf{x}_0 \in \{0, 1\}^n \mid C(\mathbf{x}_0) = C(\mathbf{x}_1)\}$
- Receive $\tilde{\mathbf{x}}$
- Output 1 iff
 - $\tilde{C}$ wellformed, and
 - $(\tilde{C}, \tilde{\mathbf{x}})$ consistent with $\mathbf{x}_1$.

Proof Idea: the Adversary $\mathcal{A}$

- Send log-depth circuit C :

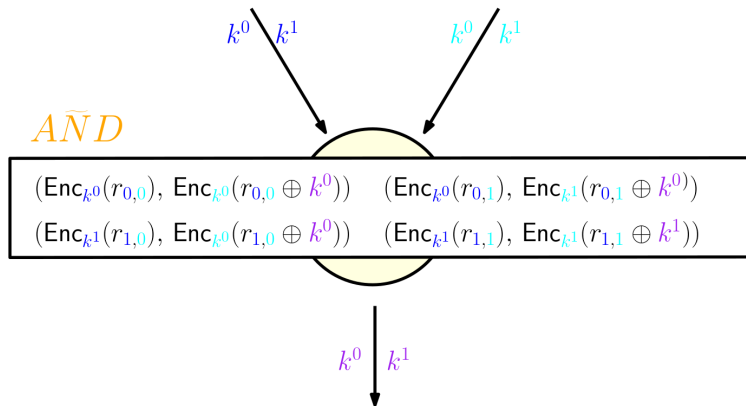


- Receive $\tilde{C}$
- Sample $\mathbf{x}_1 \leftarrow \{0, 1\}^n$, $\mathbf{x}_0 \leftarrow \{\mathbf{x}_0 \in \{0, 1\}^n \mid C(\mathbf{x}_0) = C(\mathbf{x}_1)\}$
- Receive $\tilde{\mathbf{x}}$
- Output 1 iff
 - $\tilde{C}$ wellformed, and
 - $(\tilde{C}, \tilde{\mathbf{x}})$ consistent with $\mathbf{x}_1$.

(here: consider *non-rewinding* reduction that runs $\mathcal{A}$ *once*, general case: q -wise independent hash functions)

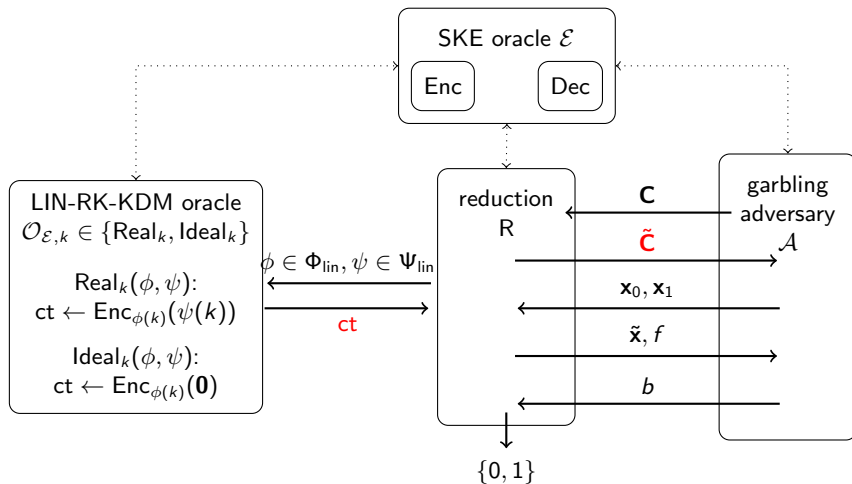
Proof Idea: the Adversary $\mathcal{A}$

Wellformedness of $\tilde{C}$ by brute-force:



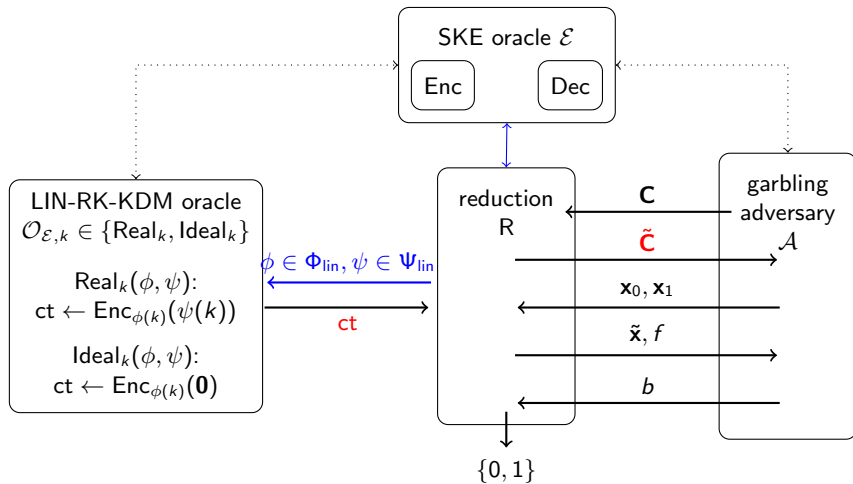
Allows to map keys to bits $\Rightarrow$ check consistency of this map with $(\tilde{x}, \mathbf{x}_1)$

Proof Idea: intuition for R



Some ct must be embedded in a garbling table of an AND gate in $\tilde{\mathbf{C}}$

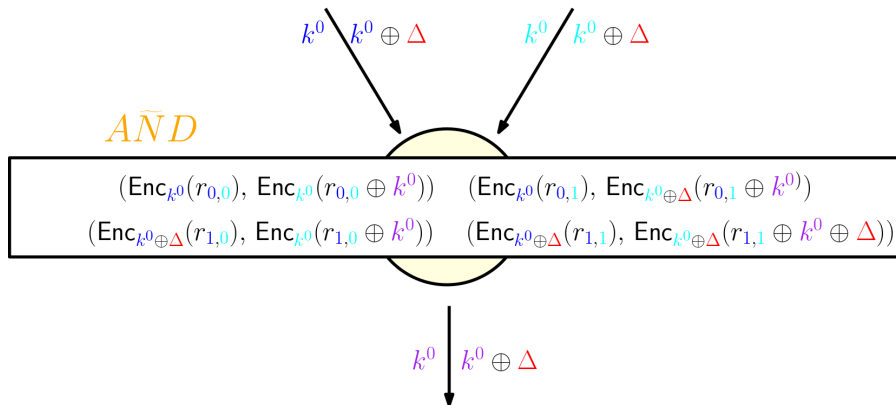
Proof Idea: uselessness of $\mathcal{A}$ for R



Some **ct** must be embedded in a garbling table of an AND gate in $\tilde{\mathbf{C}}$

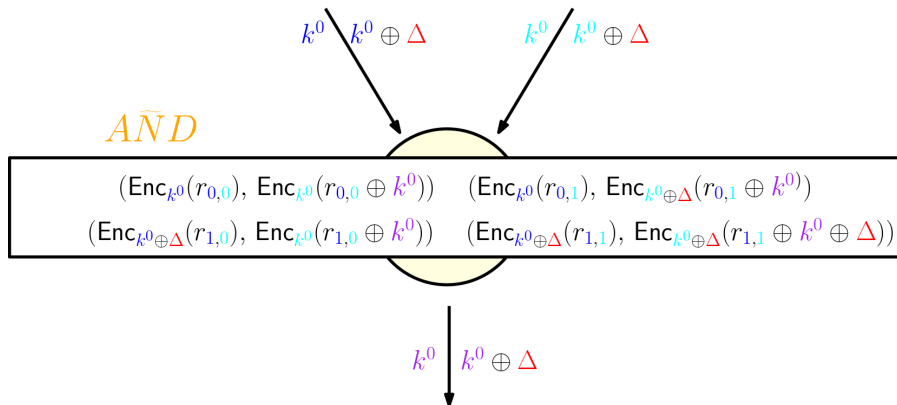
Enc random expanding function $\Rightarrow$ **all** **ct** in $\tilde{\mathbf{C}}$ through oracle queries

Proof Idea: uselessness of $\mathcal{A}$ for R



1) only secret is Δ (and enc rand. $r_{b,c}$) & 2) all ct through oracle queries

Proof Idea: uselessness of $\mathcal{A}$ for R



1) only secret is Δ (and enc rand. $r_{b,c}$) & 2) all ct through oracle queries

$\Rightarrow$ either $(\tilde{\mathbf{C}}, \tilde{\mathbf{x}})$ malformed w.r.t. $\mathbf{x}_1$, or can extract Δ from queries

(except for negl chance of embedding LIN-RK-KDM challenge key as Δ *consistent* with $\mathbf{x}_1$ (req. guessing $\mathbf{x}_1$))

Conclusion

Theorem (Our results (informal))

*Any black-box proof of **adaptive** security for free-XOR / “Half Gates” based on LIN-RK-KDM secure SKE / CCR secure hash function incurs an **exponential** security loss (even for NC1 circuits).*

⇒ free-XOR based garbling schemes **selectively** secure, but can **not** be proven **adaptively** secure using black-box reduction (i.e. standard proof approach)

Conclusion

Theorem (Our results (informal))

*Any black-box proof of **adaptive** security for free-XOR / “Half Gates” based on LIN-RK-KDM secure SKE / CCR secure hash function incurs an **exponential** security loss (even for NC1 circuits).*

⇒ free-XOR based garbling schemes **selectively** secure, but can **not** be proven **adaptively** secure using black-box reduction (i.e. standard proof approach)

- for those concrete constructions What about minor modifications?
- under given computational assumptions Stronger assumptions?
- in the standard model Easy to circumvent in ROM
- proves weakness of the schemes, but no attack/counterexample

Conclusion

Theorem (Our results (informal))

*Any black-box proof of **adaptive** security for free-XOR / “Half Gates” based on LIN-RK-KDM secure SKE / CCR secure hash function incurs an **exponential** security loss (even for NC1 circuits).*

⇒ free-XOR based garbling schemes **selectively** secure, but can **not** be proven **adaptively** secure using black-box reduction (i.e. standard proof approach)

- for those concrete constructions What about minor modifications?
- under given computational assumptions Stronger assumptions?
- in the standard model Easy to circumvent in ROM
- proves weakness of the schemes, but no attack/counterexample

THANK YOU FOR YOUR ATTENTION!

OPEN QUESTIONS?