

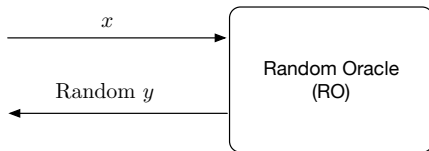
Random Oracles and Obfuscation

Pooya Farshim

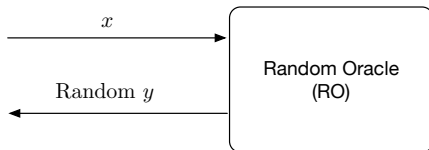
École Normale Supérieure

Summer School on Cryptology Crypto-CO

Random Oracles



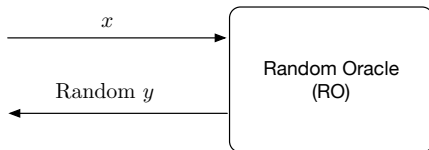
Random Oracles



- Random oracles (ROs) model ideal hash functions [BR93].
In the RO model:

All parties have oracle access to a uniformly chosen **random function**.

Random Oracles



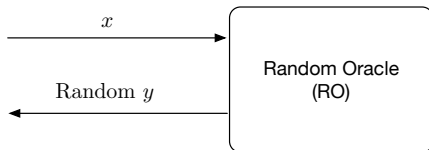
- Random oracles (ROs) model ideal hash functions [BR93].

In the RO model:

All parties have oracle access to a
uniformly chosen **random function**.

- ROs enable the security proofs of a wide range of practical and strongly secure cryptosystems: encryption & signature schemes, key exchange, disk encryption, ...

Random Oracles



- Random oracles (ROs) model ideal hash functions [BR93].

In the RO model:

All parties have oracle access to a
uniformly chosen **random function**.

- ROs enable the security proofs of a wide range of practical and strongly secure cryptosystems: encryption & signature schemes, key exchange, disk encryption, . . .
- Standard-model counterparts are often less secure and/or less efficient.

RO Uninstantiability

- Reliance on ROs, although practical, is somewhat debatable:
There are **uninstantiable** ROM schemes [CGH98].

RO Uninstantiability

- Reliance on ROs, although practical, is somewhat debatable:
There are **uninstantiable** ROM schemes [CGH98].

This means $\exists$ scheme $\text{Enc}^{\mathcal{O}}$ s.t.

- 1 $\text{Enc}^{\mathcal{RO}}$ is secure.
- 2 Enc^{Hash} is insecure for **any** concrete Hash.

RO Uninstantiability

- Reliance on ROs, although practical, is somewhat debatable:
There are **uninstantiable** ROM schemes [CGH98].

This means $\exists$ scheme Enc^O s.t.

- 1 $\text{Enc}^{\mathcal{RO}}$ is secure.
- 2 Enc^{Hash} is insecure for **any** concrete Hash.

Scheme $\text{Enc}^O(K, M)$:

- 1 Interpret M as the **description** of a hash function Hash.
- 2 If $O(x) = \text{Hash}(x)$ for $x = 1 \dots n$ append K to ciphertexts.
- 3 Else return a normal/good encryption of M .

RO Uninstantiability

- Reliance on ROs, although practical, is somewhat debatable:
There are **uninstantiable** ROM schemes [CGH98].

This means $\exists$ scheme Enc^O s.t.

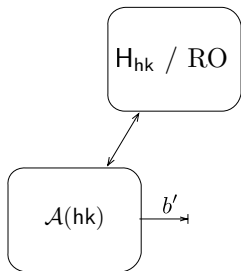
- 1 $\text{Enc}^{\mathcal{RO}}$ is secure.
- 2 Enc^{Hash} is insecure for **any** concrete Hash.

Scheme $\text{Enc}^O(K, M)$:

- 1 Interpret M as the **description** of a hash function Hash.
 - 2 If $O(x) = \text{Hash}(x)$ for $x = 1 \dots n$ append K to ciphertexts.
 - 3 Else return a normal/good encryption of M .
- Lack of a **definition** formalizing “RO-like” behavior.

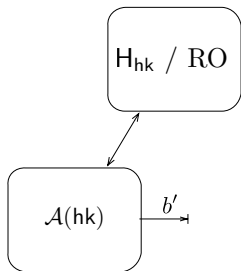
(Very) Naïve Attempt at Modeling ROs

Call a hash function “IND-RO” if:



(Very) Naïve Attempt at Modeling ROs

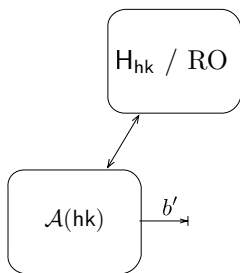
Call a hash function “IND-RO” if:



$$\mathbf{Adv}_{H, \mathcal{A}}^{\text{ind-ro}}(\lambda) := 2 \cdot \Pr[b' = b] - 1$$

(Very) Naïve Attempt at Modeling ROs

Call a hash function “IND-RO” if:



$$\text{Adv}_{H, \mathcal{A}}^{\text{ind-ro}}(\lambda) := 2 \cdot \Pr[b' = b] - 1$$

Clearly uninstantiable:

$\mathcal{A}(\text{hk})$: compute $H_{\text{hk}}(0)$ and compare to the oracle's reply.

But observe:

The adversary knows a full input, namely $(\text{hk}, 0)$.

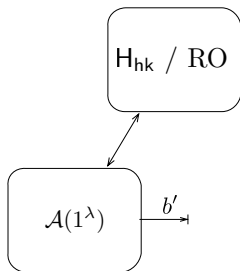
Can We Fix the Naïve Model?

Can We Fix the Naïve Model?

Let's hide hk :

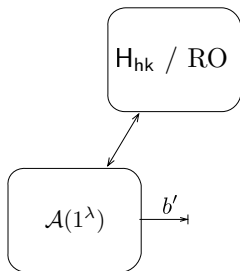
Can We Fix the Naïve Model?

Let's hide hk : PRF security:



Can We Fix the Naïve Model?

Let's hide hk : PRF security:

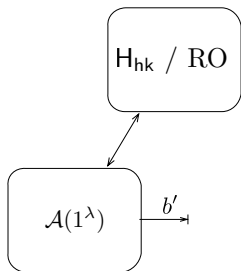


$$\mathbf{Adv}_{H, \mathcal{A}}^{\text{prf}}(\lambda) := 2 \cdot \Pr[b' = b] - 1$$

Not so useful in the context of hashing: hk is publicly available.

Can We Fix the Naïve Model?

Let's hide hk : PRF security:



$$\mathbf{Adv}_{H, \mathcal{A}}^{\text{prf}}(\lambda) := 2 \cdot \Pr[b' = b] - 1$$

Not so useful in the context of hashing: hk is publicly available.

First idea:

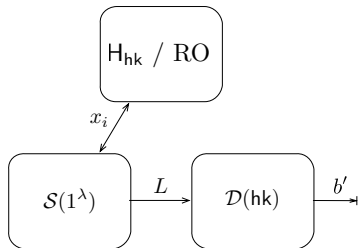
Split $\mathcal{A}$: one part gets hk and the other gets oracle access.

Modeling ROs via Split Adversaries

Call the two components of $\mathcal{A}$ the **source** $\mathcal{S}$ and the **distinguisher** $\mathcal{D}$:

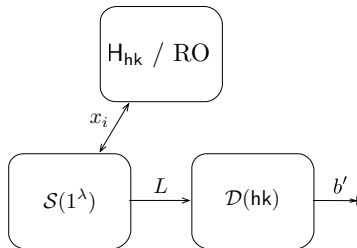
Modeling ROs via Split Adversaries

Call the two components of $\mathcal{A}$ the **source** $\mathcal{S}$ and the **distinguisher** $\mathcal{D}$:



Modeling ROs via Split Adversaries

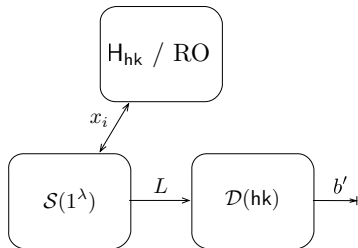
Call the two components of $\mathcal{A}$ the **source** $\mathcal{S}$ and the **distinguisher** $\mathcal{D}$:



$$\mathbf{Adv}_{H, \mathcal{S}, \mathcal{D}}^{\text{uce}}(\lambda) := 2 \cdot \Pr[b' = b] - 1$$

Modeling ROs via Split Adversaries

Call the two components of $\mathcal{A}$ the **source** $\mathcal{S}$ and the **distinguisher** $\mathcal{D}$:



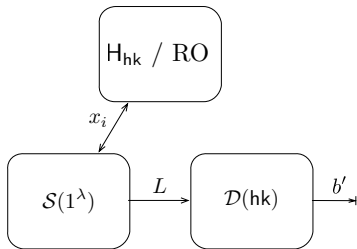
$$\mathbf{Adv}_{\mathbf{H}, \mathcal{S}, \mathcal{D}}^{\text{uce}}(\lambda) := 2 \cdot \Pr[b' = b] - 1$$

Still uninstantiable:

$\mathcal{S}$ leaks oracle's response on 0 via L , and
 $\mathcal{D}(\text{hk})$ checks where it's coming from.

Modeling ROs via Split Adversaries

Call the two components of $\mathcal{A}$ the **source** $\mathcal{S}$ and the **distinguisher** $\mathcal{D}$:



$$\text{Adv}_{H, \mathcal{S}, \mathcal{D}}^{\text{uce}}(\lambda) := 2 \cdot \Pr[b' = b] - 1$$

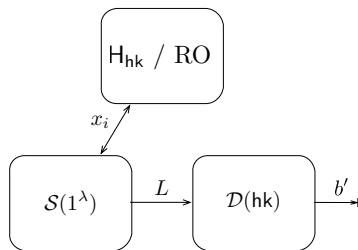
Still uninstantiable:

$\mathcal{S}$ leaks oracle's response on 0 via L , and
 $\mathcal{D}(hk)$ checks where it's coming from.

Second idea:

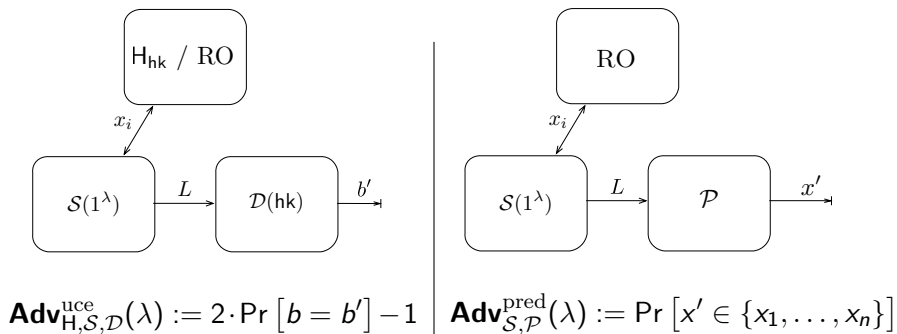
Restrict L : it must not leak any of $\mathcal{S}$'s queries.

Universal Computational Extractors (UCEs) [BHK13a]

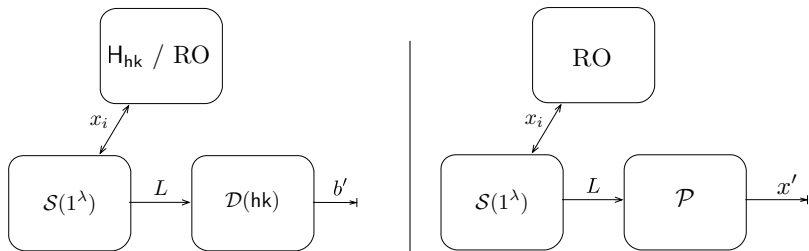


$$\mathbf{Adv}_{H,S,D}^{\text{uce}}(\lambda) := 2 \cdot \Pr[b = b'] - 1$$

Universal Computational Extractors (UCEs) [BHK13a]



Universal Computational Extractors (UCEs) [BHK13a]

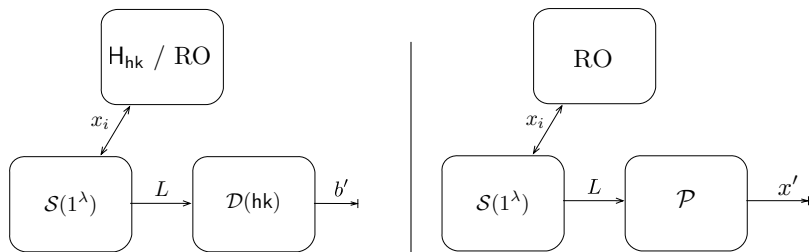


$$\mathbf{Adv}_{\mathcal{H}, \mathcal{S}, \mathcal{D}}^{\text{uce}}(\lambda) := 2 \cdot \Pr[b = b'] - 1$$

$$\mathbf{Adv}_{\mathcal{S}, \mathcal{P}}^{\text{pred}}(\lambda) := \Pr[x' \in \{x_1, \dots, x_n\}]$$

$\mathcal{S}$ is **unpredictable** iff: $\mathbf{Adv}_{\mathcal{S}, \mathcal{P}}^{\text{pred}}(\lambda) \in \text{Negl}$ for any **efficient** $\mathcal{P}$.

Universal Computational Extractors (UCEs) [BHK13a]



$$\mathbf{Adv}_{\mathcal{H}, \mathcal{S}, \mathcal{D}}^{\text{uce}}(\lambda) := 2 \cdot \Pr[b = b'] - 1 \quad \mathbf{Adv}_{\mathcal{S}, \mathcal{P}}^{\text{pred}}(\lambda) := \Pr[x' \in \{x_1, \dots, x_n\}]$$

$\mathcal{S}$ is **unpredictable** iff: $\mathbf{Adv}_{\mathcal{S}, \mathcal{P}}^{\text{pred}}(\lambda) \in \text{Negl}$ for any **efficient** $\mathcal{P}$.

$\mathcal{H}$ is **UCE1 secure** iff: $\mathbf{Adv}_{\mathcal{H}, \mathcal{S}, \mathcal{D}}^{\text{uce}}(\lambda) \in \text{Negl}$ for any **unpredictable** $\mathcal{S}$.

Applications of UCE [BHK13a]

UCE-secure hash functions can instantiate the RO in

- Deterministic PKEs
- RKA and KDM security
- Point-function obfuscation
- Message-locked encryption
- Proofs of storage
- Poly-many hard-core bits for any OWF
- OAEP, garbling schemes, ...

Applications of UCE [BHK13a]

UCE-secure hash functions can instantiate the RO in

- Deterministic PKEs
- RKA and KDM security
- Point-function obfuscation
- Message-locked encryption
- Proofs of storage
- Poly-many hard-core bits for any OWF
- OAEP, garbling schemes, ...

Bottom line:

UCEs model many RO-like properties.

Is UCE Security Instantiable?

Suppose we could “strongly” obfuscate H :

- $\mathcal{S}$: Choose a random x . Query x to get y . Leak as L the values

$$y, \quad \text{Obf}(H(\cdot, x)) .$$

- $\mathcal{D}$: Run the obfuscated code on hk . Check if the output matches y .

Is UCE Security Instantiable?

Suppose we could “strongly” obfuscate H :

- $\mathcal{S}$: Choose a random x . Query x to get y . Leak as L the values

$$y, \quad \text{Obf}(H(\cdot, x)) .$$

- $\mathcal{D}$: Run the obfuscated code on hk . Check if the output matches y .

Recall we need unpredictability for a nontrivial attack:

$\text{Obf}(H(\cdot, x))$ **must hide** x for unpredictability.

Is UCE Security Instantiable?

Suppose we could “strongly” obfuscate H :

- $\mathcal{S}$: Choose a random x . Query x to get y . Leak as L the values

$$y, \quad \text{Obf}(H(\cdot, x)) .$$

- $\mathcal{D}$: Run the obfuscated code on hk . Check if the output matches y .

Recall we need unpredictability for a nontrivial attack:

$\text{Obf}(H(\cdot, x))$ **must hide** x for unpredictability.

It's unclear if such obfuscators exist.

Is UCE Security Instantiable?

Suppose we could “strongly” obfuscate H :

- $\mathcal{S}$: Choose a random x . Query x to get y . Leak as L the values $y, \text{Obf}(H(\cdot, x))$.
- $\mathcal{D}$: Run the obfuscated code on hk . Check if the output matches y .

Recall we need unpredictability for a nontrivial attack:

$\text{Obf}(H(\cdot, x))$ **must hide** x for unpredictability.

It's unclear if such obfuscators exist.

Perhaps a different circuit and/or obfuscator might help?

Indistinguishability Obfuscation (iO)

Let C_0 and C_1 be two **functionally equivalent** circuits:

$$\forall x : C_0(x) = C_1(x) .$$

iO **security**: cannot efficiently distinguish obfuscations of such circuits:

$$\text{iO}(C_0) \stackrel{c}{\approx} \text{iO}(C_1) .$$

[GGH⁺13]: indistinguishability obfuscation for all poly-sized circuits from intractability assumptions related to multi-linear maps.

Indistinguishability Obfuscation (iO)

Let C_0 and C_1 be two **functionally equivalent** circuits:

$$\forall x : C_0(x) = C_1(x) .$$

iO **security**: cannot efficiently distinguish obfuscations of such circuits:

$$\text{iO}(C_0) \stackrel{c}{\approx} \text{iO}(C_1) .$$

[GGH⁺13]: indistinguishability obfuscation for all poly-sized circuits from intractability assumptions related to multi-linear maps.

Can we use iO to attack UCEs?

The iO Attack

$\mathcal{S}$: Leak an indistinguishably obfuscation of the Boolean circuit

$$H(\cdot, x) \stackrel{?}{=} y$$

where x is random and y is oracle's response on x .

The iO Attack

$\mathcal{S}$: Leak an indistinguishably obfuscation of the Boolean circuit

$$\boxed{H(\cdot, x) \stackrel{?}{=} y}$$

where x is random and y is oracle's response on x .

$\mathcal{D}$: Run the obfuscation on hk and return the result.

- Oracle = H : $y = H(hk, x) \implies$ check always passes.

The iO Attack

$\mathcal{S}$: Leak an indistinguishably obfuscation of the Boolean circuit

$$\boxed{H(\cdot, x) \stackrel{?}{=} y}$$

where x is random and y is oracle's response on x .

$\mathcal{D}$: Run the obfuscation on hk and return the result.

- Oracle = H : $y = H(hk, x) \implies$ check always passes.
- Oracle = $\mathcal{RO}$: y is random $\implies$ check passes with prob $2^{-|y|}$.

Advantage of $(\mathcal{S}, \mathcal{D})$ is:

$$1 - 2^{-|y|}.$$

The iO Attack

$\mathcal{S}$: Leak an indistinguishably obfuscation of the Boolean circuit

$$\boxed{H(\cdot, x) \stackrel{?}{=} y}$$

where x is random and y is oracle's response on x .

$\mathcal{D}$: Run the obfuscation on hk and return the result.

- Oracle = H : $y = H(hk, x) \implies$ check always passes.
- Oracle = $\mathcal{RO}$: y is random $\implies$ check passes with prob $2^{-|y|}$.

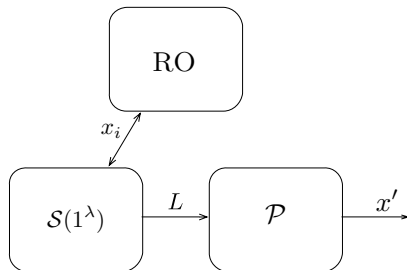
Advantage of $(\mathcal{S}, \mathcal{D})$ is:

$$1 - 2^{-|y|}.$$

Is $\mathcal{S}$ unpredictable?

Proving Unpredictability

Need to ensure obfuscation hides x when y is **truly random**:
unpredictability was defined wrt. **random oracle**.



Proving Unpredictability

Need to ensure obfuscation hides x when y is **truly random**:

unpredictability was defined wrt. **random oracle**.

For any x , by the union bound we have

$$\Pr_{y \leftarrow \mathcal{Rng}} \left[\exists hk \text{ s.t. } H(hk, x) = y \right] \leq \frac{|\mathcal{KSp}|}{|\mathcal{Rng}|} .$$

Proving Unpredictability

Need to ensure obfuscation hides x when y is **truly random**:

unpredictability was defined wrt. **random oracle**.

For any x , by the union bound we have

$$\Pr_{y \leftarrow \mathcal{Rng}} \left[\exists hk \text{ s.t. } H(hk, x) = y \right] \leq \frac{|\mathcal{KSp}|}{|\mathcal{Rng}|}.$$

If $|\mathcal{KSp}| \ll |\mathcal{Rng}|$, we get that

$$H(\cdot, x) \stackrel{?}{=} y \quad \equiv \quad \text{Zero circuit}$$

with overwhelming probability.

Proving Unpredictability

Need to ensure obfuscation hides x when y is **truly random**:

unpredictability was defined wrt. **random oracle**.

For any x , by the union bound we have

$$\Pr_{y \leftarrow \text{Rng}} \left[\exists hk \text{ s.t. } H(hk, x) = y \right] \leq \frac{|\text{KSp}|}{|\text{Rng}|}.$$

If $|\text{KSp}| \ll |\text{Rng}|$, we get that

$$H(\cdot, x) \stackrel{?}{=} y \quad \equiv \quad \text{Zero circuit}$$

with overwhelming probability. Now by the security of iO

$$\text{iO}(H(\cdot, x) \stackrel{?}{=} y) \text{ leaks no more than } \text{iO}(\text{Zero}),$$

and the latter is independent of x .

Dropping $|KSp| \ll |Rng|$

Consider instead iO-ing the circuit

$$H(\cdot, x_1) \stackrel{?}{=} y_1 \wedge \cdots \wedge H(\cdot, x_n) \stackrel{?}{=} y_n$$

where x_i are random and y_i are the corresponding responses.

Dropping $|KSp| \ll |Rng|$

Consider instead iO-ing the circuit

$$H(\cdot, x_1) \stackrel{?}{=} y_1 \wedge \cdots \wedge H(\cdot, x_n) \stackrel{?}{=} y_n$$

where x_i are random and y_i are the corresponding responses.

Theorem

Under the existence of iO, UCE security is uninstantiable.

Salvaging UCE: Statistical Unpredictability

iO is only **computationally** secure. Restrict the source class by

Allowing the **predictors** to run in **unbounded** time.

This raises the question: does **statistical iO** exist?

Salvaging UCE: Statistical Unpredictability

iO is only **computationally** secure. Restrict the source class by

Allowing the **predictors** to run in **unbounded** time.

This raises the question: does **statistical iO** exist?

Theorem (Goldwasser and Rothblum [GR07])

If statistical iO exists then the polynomial hierarchy collapses.

Salvaging UCE: Statistical Unpredictability

iO is only **computationally** secure. Restrict the source class by

Allowing the **predictors** to run in **unbounded** time.

This raises the question: does **statistical iO** exist?

Theorem (Goldwasser and Rothblum [GR07])

If statistical iO exists then the polynomial hierarchy collapses.

Remark for researchers working on negative results:

Example of an impossibility result that is used positively to define a security model.

Statistical UCE and Applications

Definition

H is **statistically UCE** secure if UCE advantage is negligible for all sources that are unpredictable even wrt **unbounded** predictors.

Statistical UCE and Applications

Definition

H is **statistically UCE** secure if UCE advantage is negligible for all sources that are unpredictable even wrt **unbounded** predictors.

Applications of statistical UCEs:

- RKA and KDM security
- Proof-of-storage protocol
- point-function obfuscation, correlated-input secure hashing, and garbling schemes

Statistical UCE and Applications

Definition

H is **statistically UCE** secure if UCE advantage is negligible for all sources that are unpredictable even wrt **unbounded** predictors.

Applications of statistical UCEs:

- RKA and KDM security
- Proof-of-storage protocol
- point-function obfuscation, correlated-input secure hashing, and garbling schemes

Does **not** hold for some other applications:

Reductions to UCE rely on L that only **computationally** hide x .

Final Thoughts

- UCEs bring added provable-security guarantees to practical protocols.
- Reliance on idealized objects, replaced by UCEs.
- This work brings refinements to UCEs:
 - ▶ Computational UCE notion is uninstantiable under iO.
 - ▶ Our weaker statistical UCE is robust and still useful.
- iO attack extends to many other ROM schemes:
 - ▶ D-PKE, Hedged PKE, Fujisaki–Okamoto, KDM, MLEs, ...

Final Thoughts

- UCEs bring added provable-security guarantees to practical protocols.
- Reliance on idealized objects, replaced by UCEs.
- This work brings refinements to UCEs:
 - ▶ Computational UCE notion is uninstantiable under iO.
 - ▶ Our weaker statistical UCE is robust and still useful.
- iO attack extends to many other ROM schemes:
 - ▶ D-PKE, Hedged PKE, Fujisaki–Okamoto, KDM, MLEs, ...

Thank you.