

From Lattices to Tensor Cores: Scaling up Private Information Retrieval

David Wu

August 2026

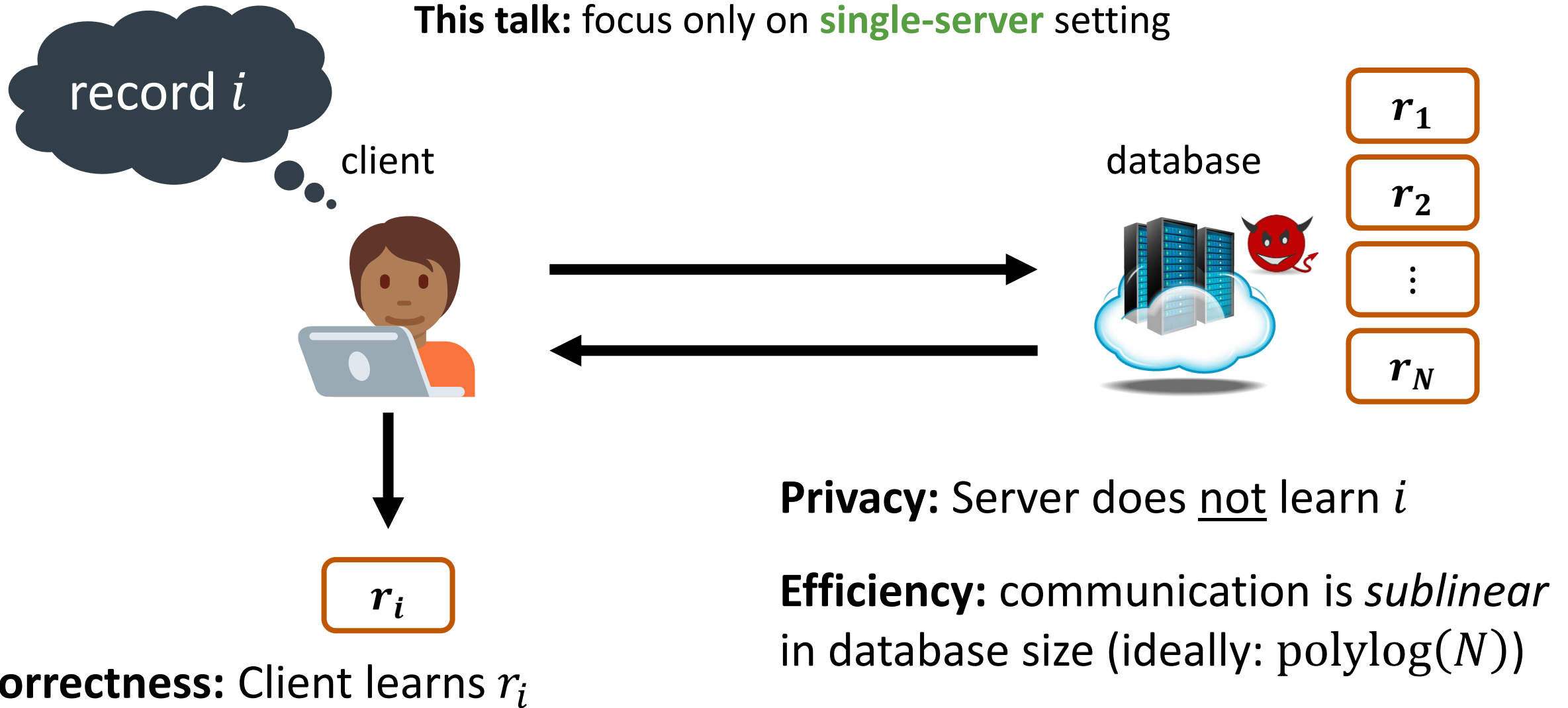
joint work with Sidaarth Sabhnani

Work funded in part by an Amazon Research Award

Private Information Retrieval (PIR)

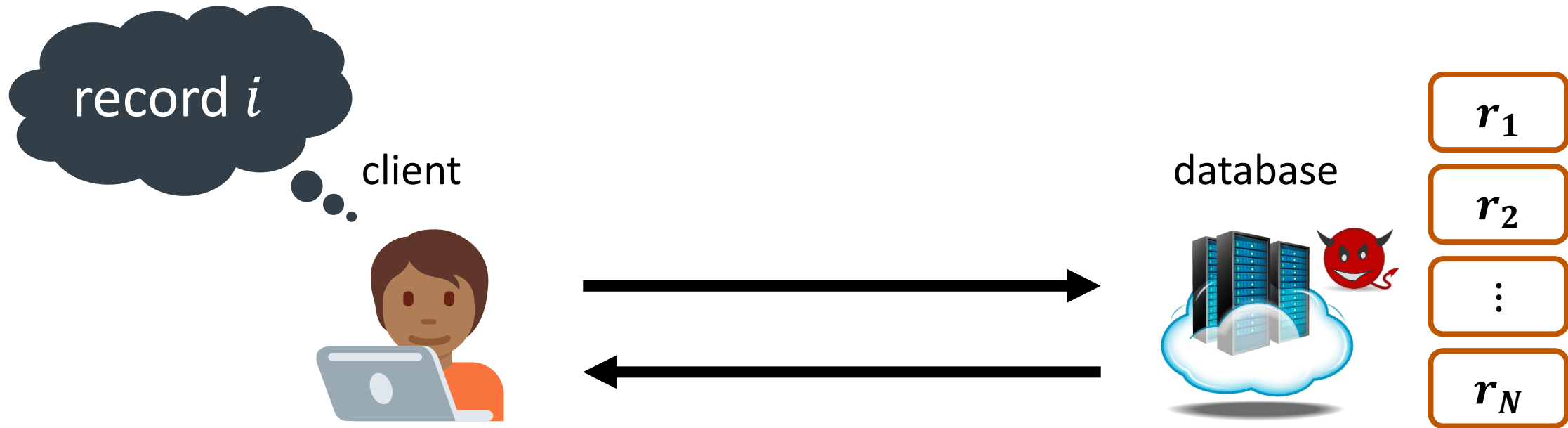
[CGKS95]

This talk: focus only on **single-server** setting



Private Information Retrieval (PIR)

[CGKS95]



Basic building block in many privacy-preserving protocols



Metadata-private messaging



Contact discovery



Private contact tracing



Certificate transparency auditing



Private web search



Private DNS



Private content delivery



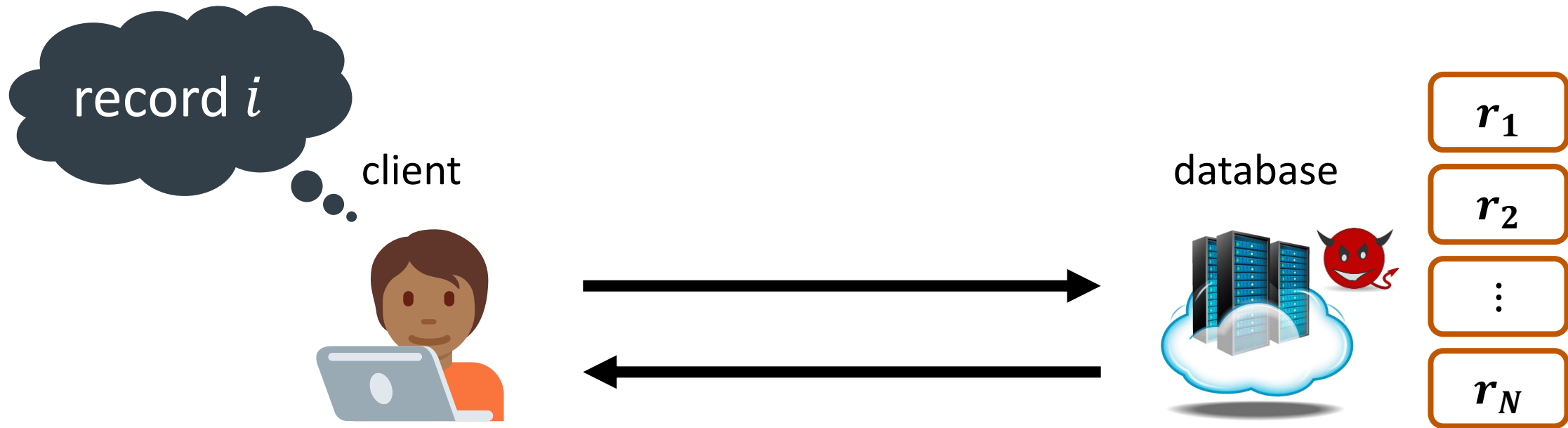
Private navigation



Private blocklist lookup

Private Information Retrieval (PIR)

[CGKS95]

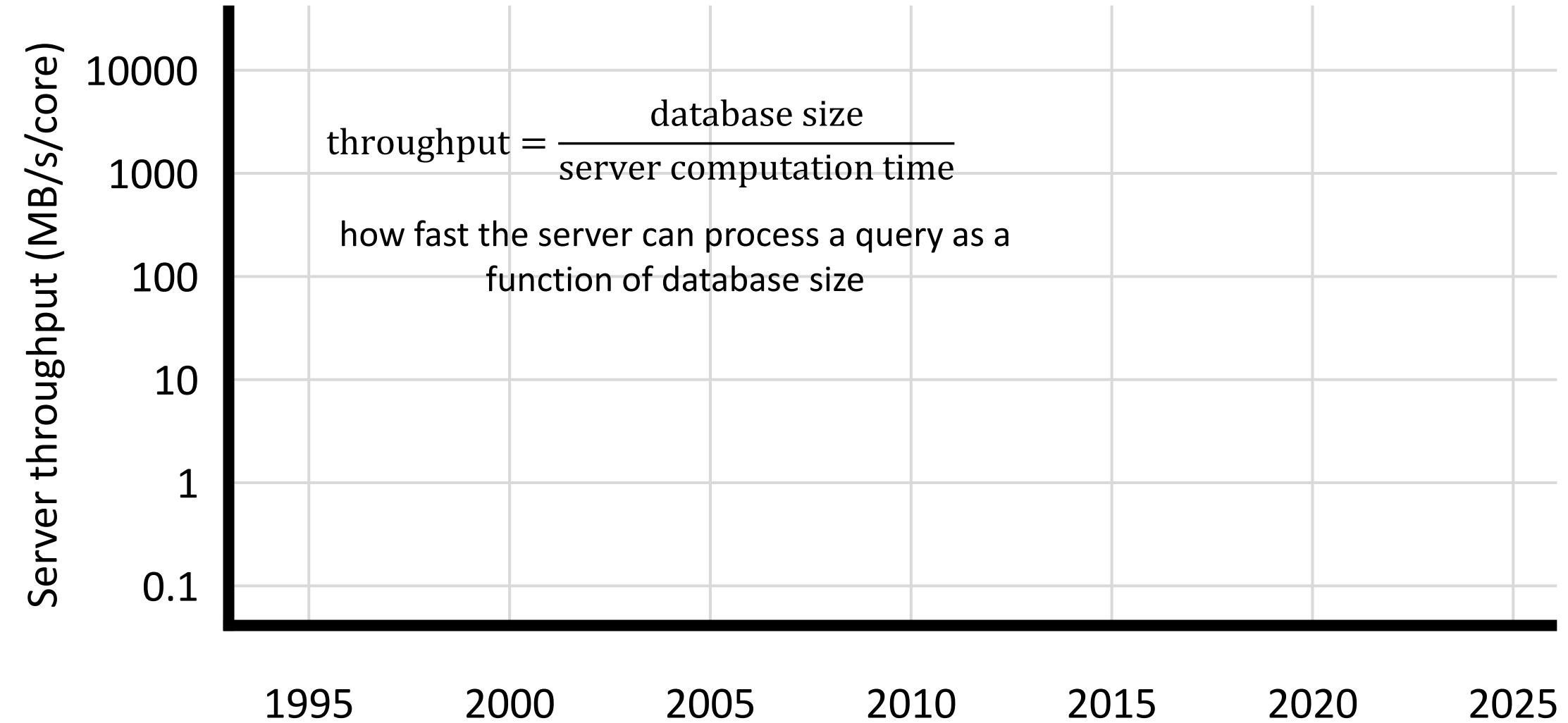


Basic building block in many privacy-preserving protocols

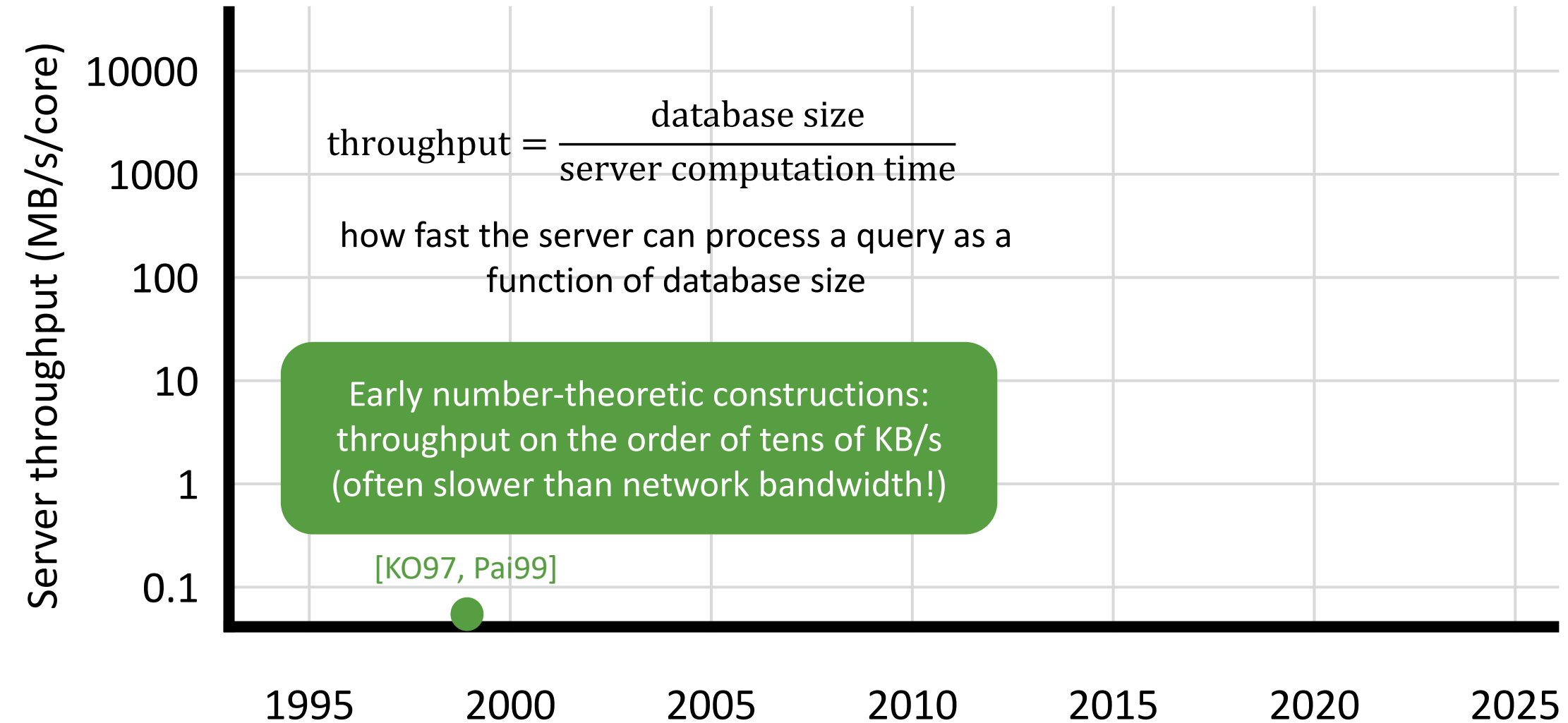
Recent deployments:

- Apple: private caller ID lookup, private image search
- Meta: advanced browsing protection

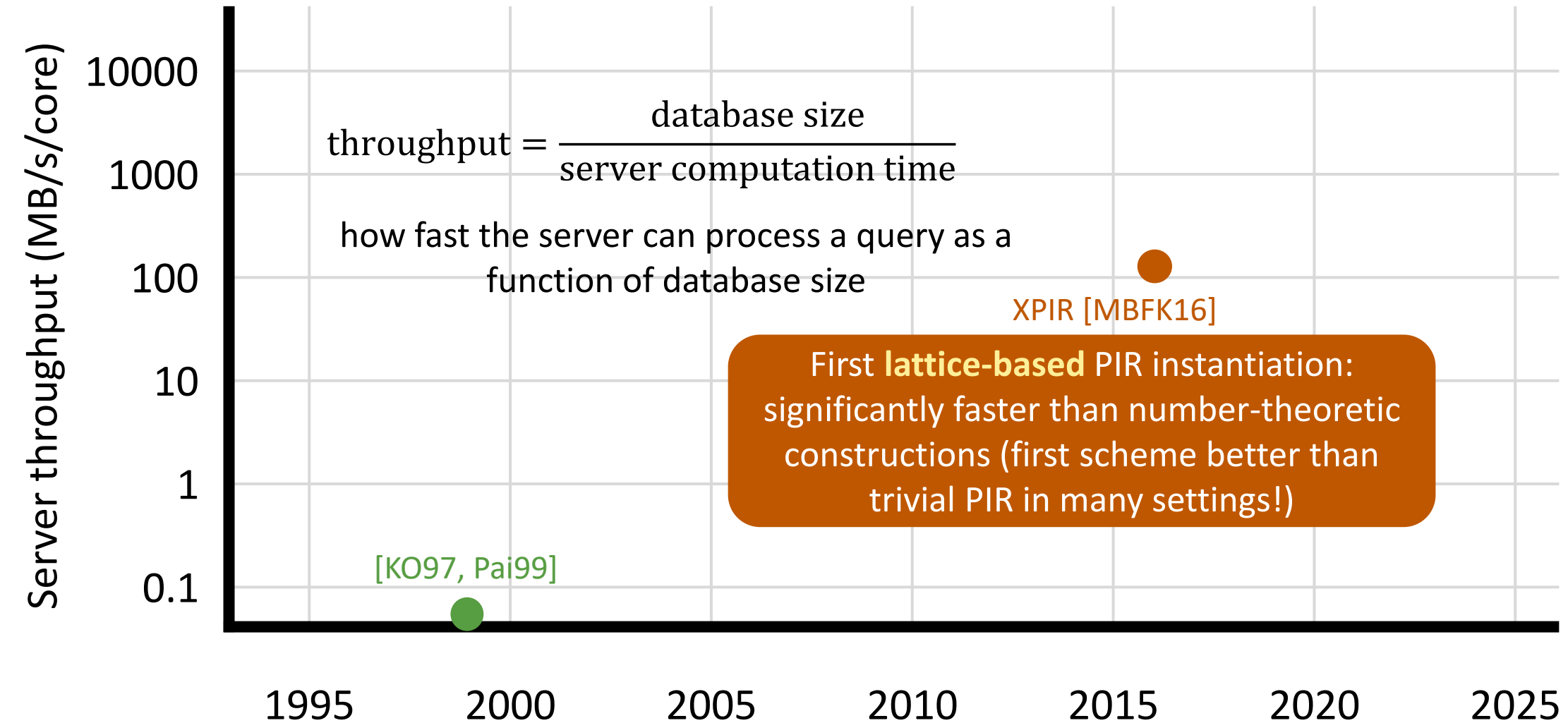
30 Years of PIR Research



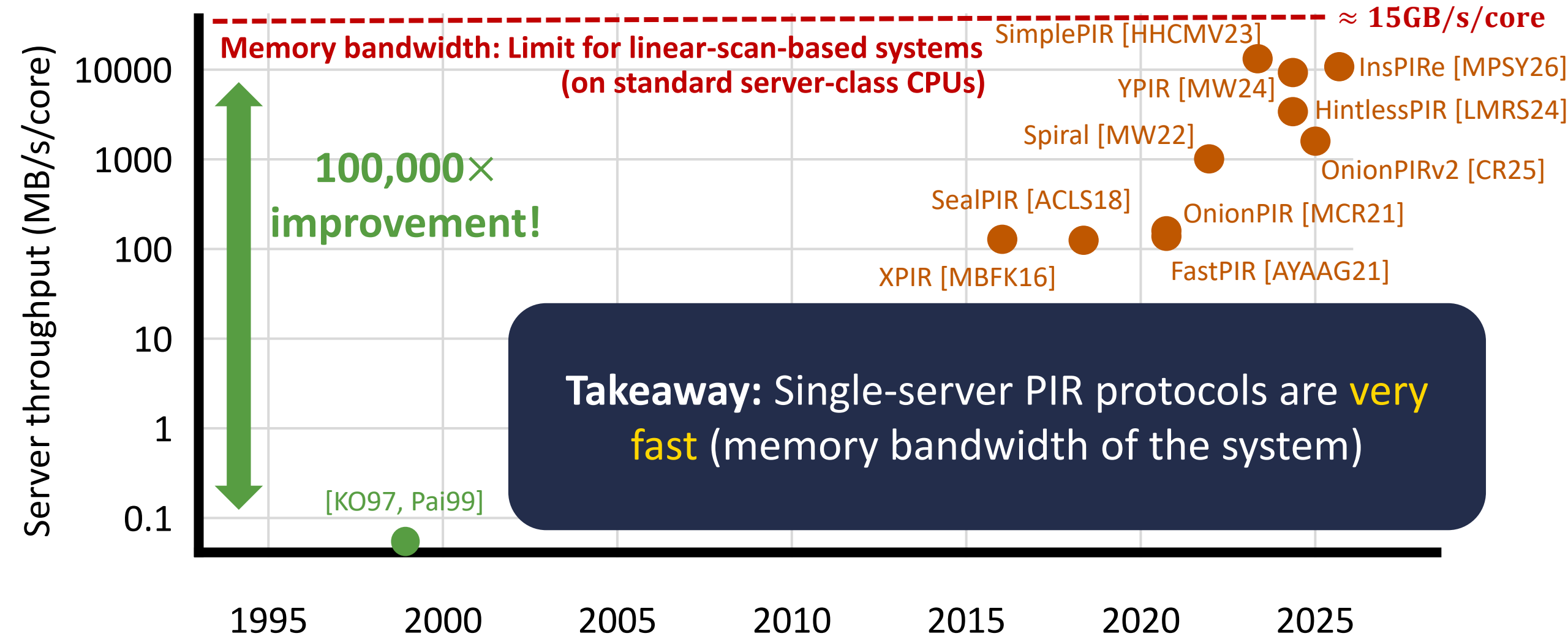
30 Years of PIR Research



30 Years of PIR Research



30 Years of PIR Research



Beyond Memory Bandwidth

[LG15, MW24, LHC25]

Can we go beyond memory bandwidth with a linear-scan PIR?

Yes, but in an **amortized** sense

Idea: if reading memory is the bottleneck, then schedule more CPU operations per byte of memory read

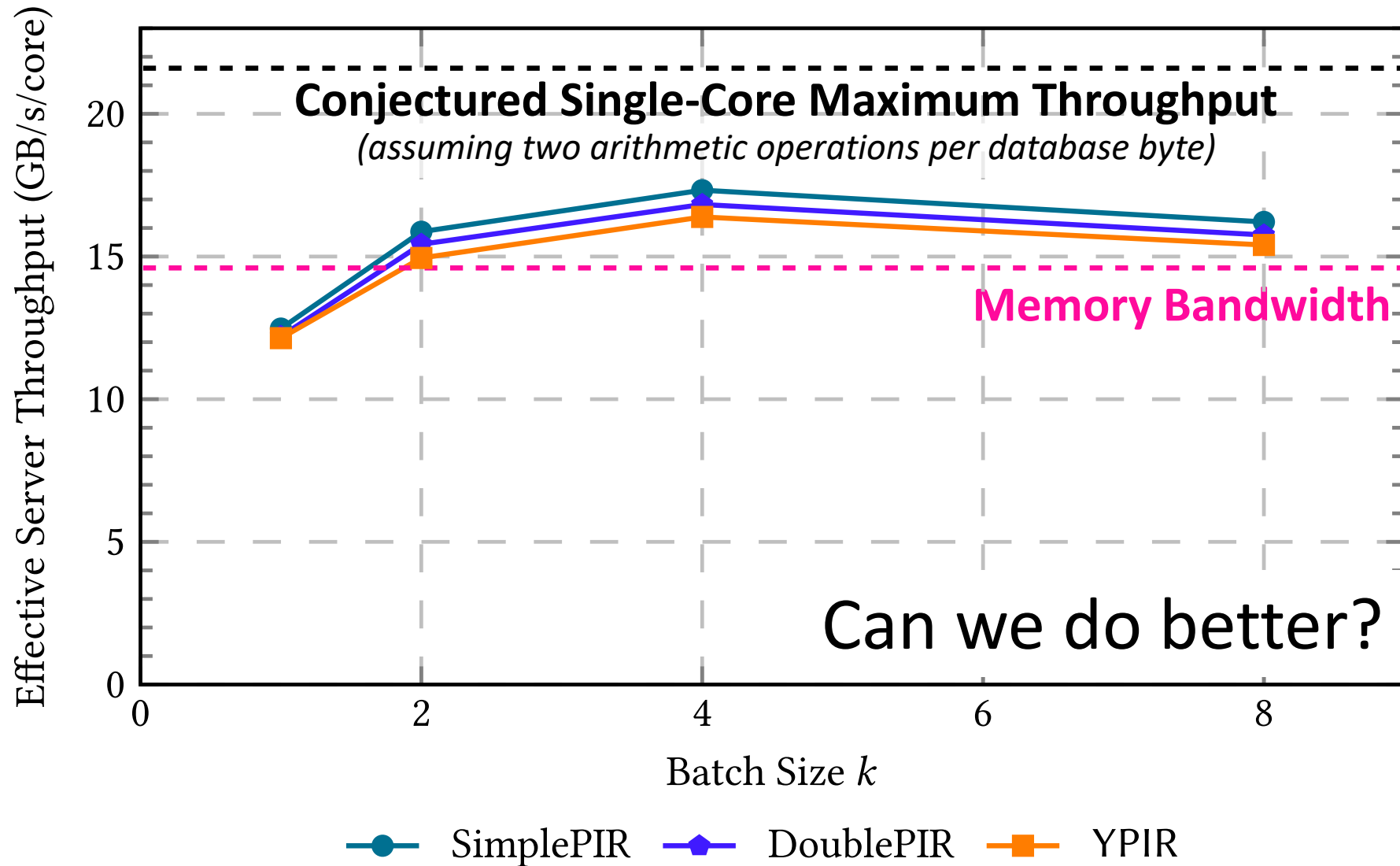
Cross-client batching: process a batch of k queries in parallel (from different, **non-coordinating** clients)

In many recent lattice-based PIR protocols, query processing is a product between a query vector and a database matrix

Cross-client batching: matrix-vector product $\Rightarrow$ matrix-matrix product

Beyond Memory Bandwidth

[MW24]



$$\frac{k \cdot |\text{DB}|}{\text{Server time}}$$

Hardware Acceleration for PIR

Main computation in PIR: matrix multiplication

NVIDIA L40S

Memory bandwidth: 864 GB/s

64-core CPU: 1.4 trillion FLOPS

Computation: 91.6 trillion 32-bit (floating point) operations/s

Tensor cores: special architecture tailored for **low-precision** matrix multiplication operations for machine learning

INT8-INT32 tensor cores: takes two input matrices A, B over $\mathbb{Z}_{2^8}$ and outputs product AB over $\mathbb{Z}_{2^{32}}$ ($\mathbb{Z}_{2^8} \times \mathbb{Z}_{2^8} \rightarrow \mathbb{Z}_{2^{32}}$)

Throughput: 733 trillion INT-8 operations/s

“tensor core model of computation”

recent Intel CPUs also support it (Intel AMX)



Hardware Acceleration for PIR

Main computation in PIR: matrix multiplication

NVIDIA L40S

Memory bandwidth: 864 GB/s

64-core CPU: 1.4 trillion FLOPS

Computation: 91.6 trillion 32-bit (floating point) operations/s

Tensor cores: special architecture tailored for **low-precision** matrix multiplication operations for machine learning

INT8-INT32 tensor cores: takes two input matrices A, B over $\mathbb{Z}_{2^8}$ and outputs product AB over $\mathbb{Z}_{2^{32}}$ ($\mathbb{Z}_{2^8} \times \mathbb{Z}_{2^8} \rightarrow \mathbb{Z}_{2^{32}}$)

Throughput: 733 trillion INT-8 operations/s



Can we leverage tensor cores to accelerate PIR?

PIR from Homomorphic Encryption

[KO97]

Starting point: a $\sqrt{N}$ construction (N = number of records)

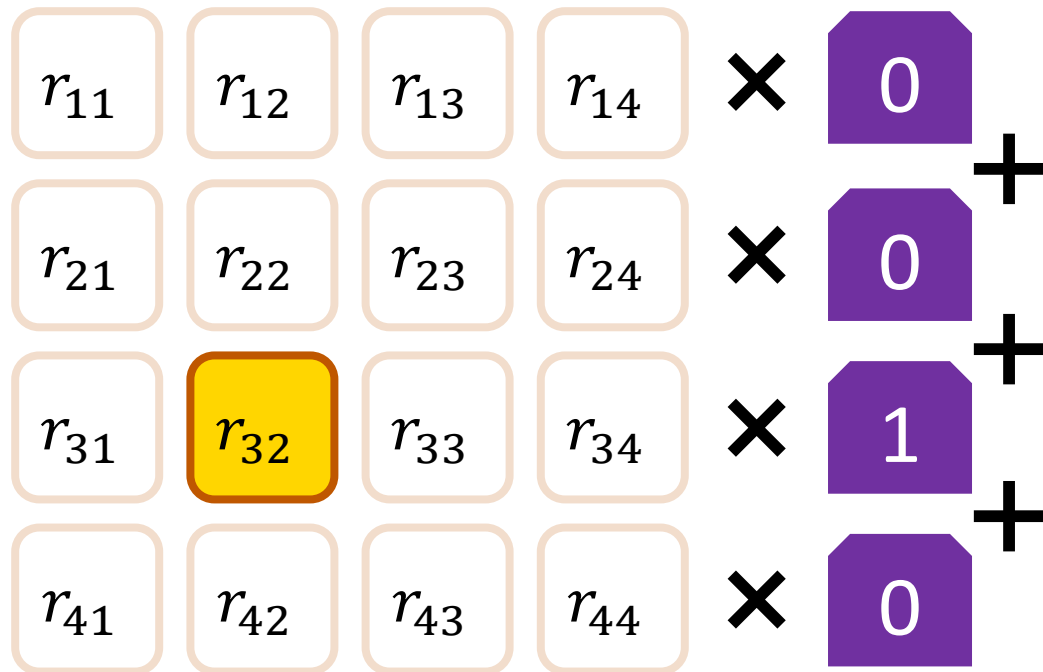
r_{11}	r_{12}	r_{13}	r_{14}
r_{21}	r_{22}	r_{23}	r_{24}
r_{31}	r_{32}	r_{33}	r_{34}
r_{41}	r_{42}	r_{43}	r_{44}

Arrange the database as a
 $\sqrt{N}$ -by- $\sqrt{N}$ matrix

PIR from Homomorphic Encryption

[KO97]

Starting point: a $\sqrt{N}$ construction (N = number of records)



Encrypt a 0/1 vector indicating the row containing the desired record

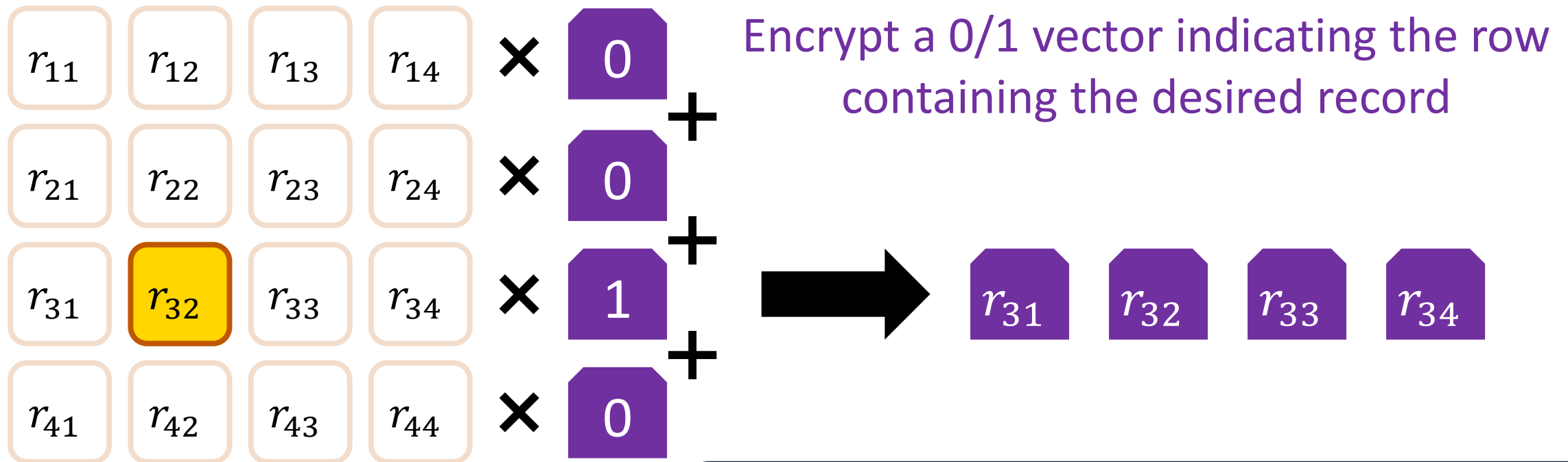
Arrange the database as a $\sqrt{N}$ -by- $\sqrt{N}$ matrix

Homomorphically compute product between query vector and database matrix

PIR from Homomorphic Encryption

[KO97]

Starting point: a $\sqrt{N}$ construction (N = number of records)



Arrange the database as a
 $\sqrt{N}$ -by- $\sqrt{N}$ matrix

Database is in the clear, so *additive*
homomorphism suffices

PIR from Homomorphic Encryption

[KO97]

Starting point: a $\sqrt{N}$ construction (N = number of records)

Client decrypts to
learn records



Encrypt a 0/1 vector indicating the row
containing the desired record



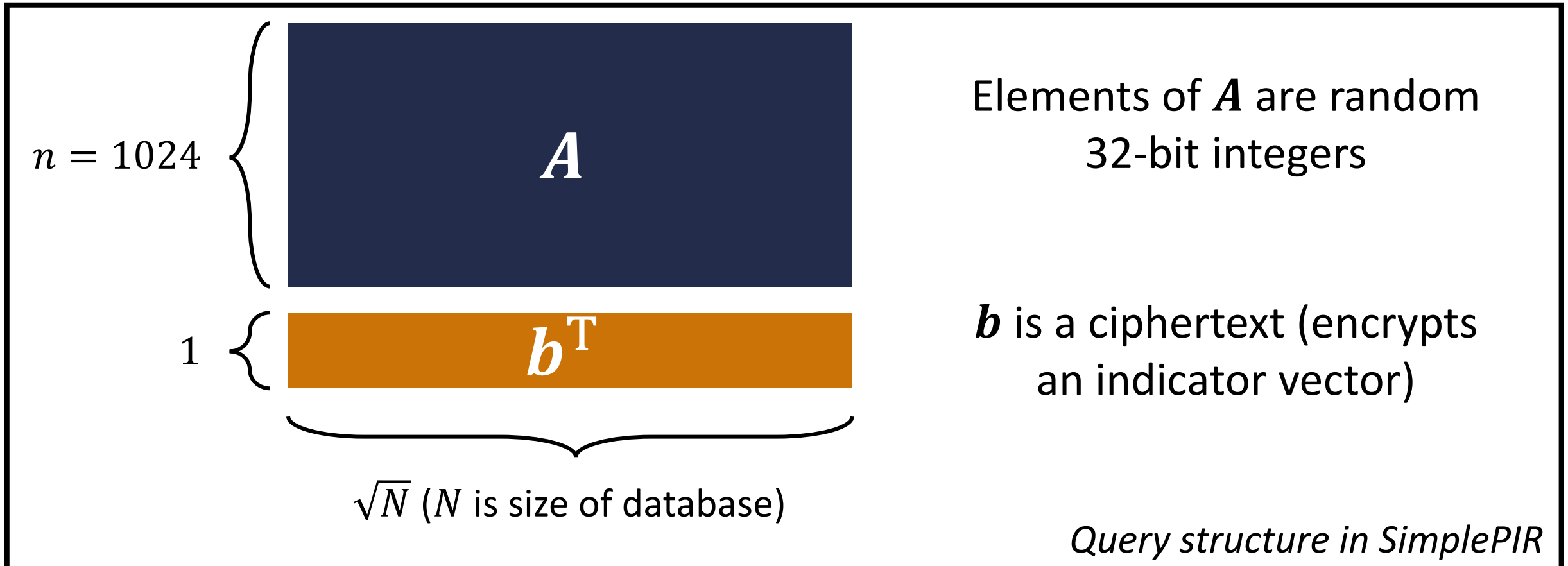
Response size: $O_\lambda(\sqrt{N})$

Homomorphically compute product
between query vector and database matrix

SimplePIR: Lightweight PIR from LWE

[HHCMV23]

Building block: Regev's linearly homomorphic encryption from LWE [Reg05]

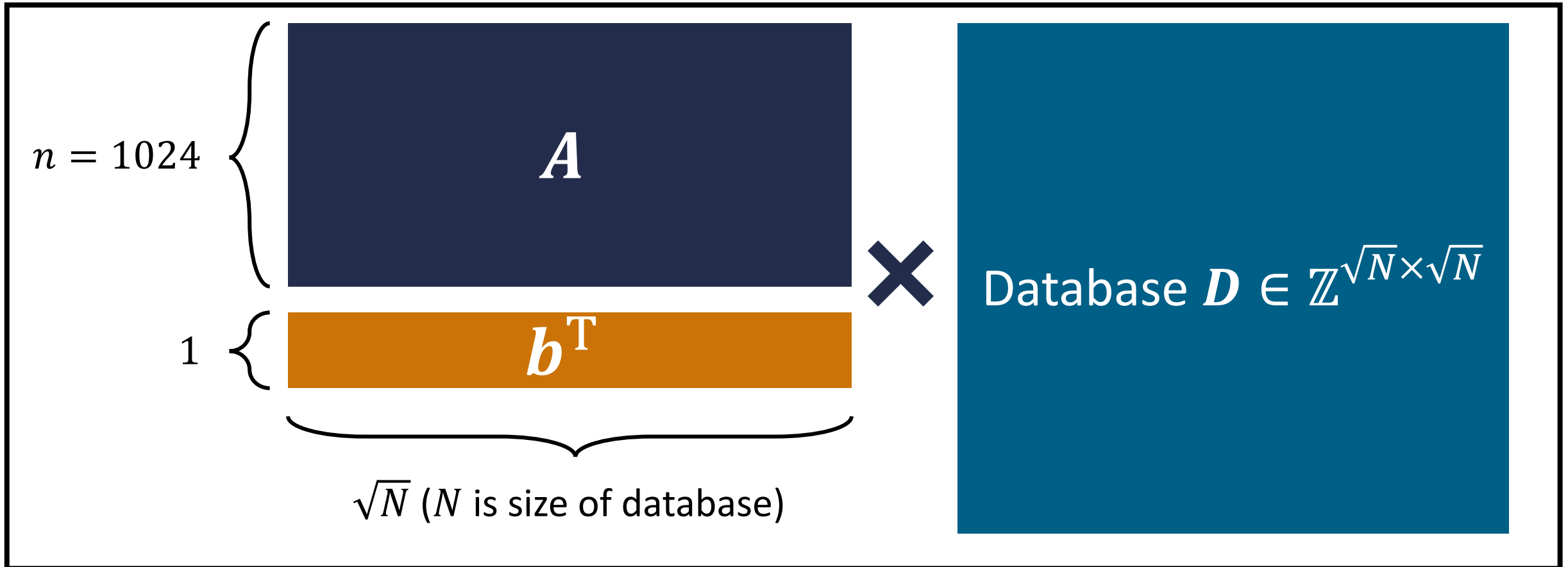


All components are elements of $\mathbb{Z}_q = \mathbb{Z}_{2^{32}}$ (i.e., 32-bit integers)

SimplePIR: Lightweight PIR from LWE

[HHCMV23]

Building block: Regev's linearly homomorphic encryption from LWE [Reg05]



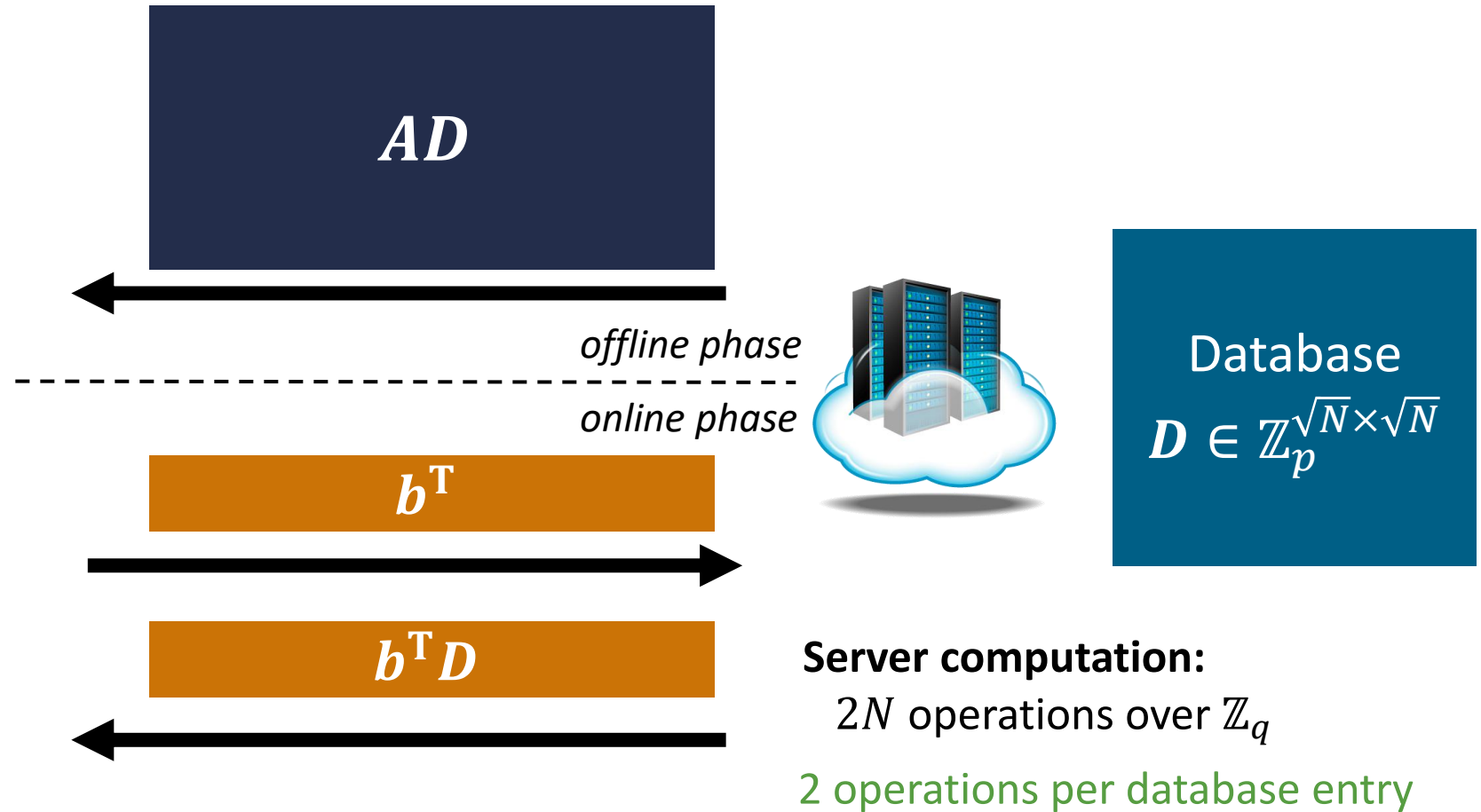
All components are elements of $\mathbb{Z}_q = \mathbb{Z}_{2^{32}}$ (i.e., 32-bit integers)

SimplePIR: Lightweight PIR from LWE

[HHCMV23]

Key insight in SimplePIR: A is query-independent so move computation of AD **offline**

Given $b^T D$ and AD , client can decrypt and learn the record



Query and response size:
 $\sqrt{N}$ elements over $\mathbb{Z}_q$

SimplePIR: Lightweight PIR from LWE

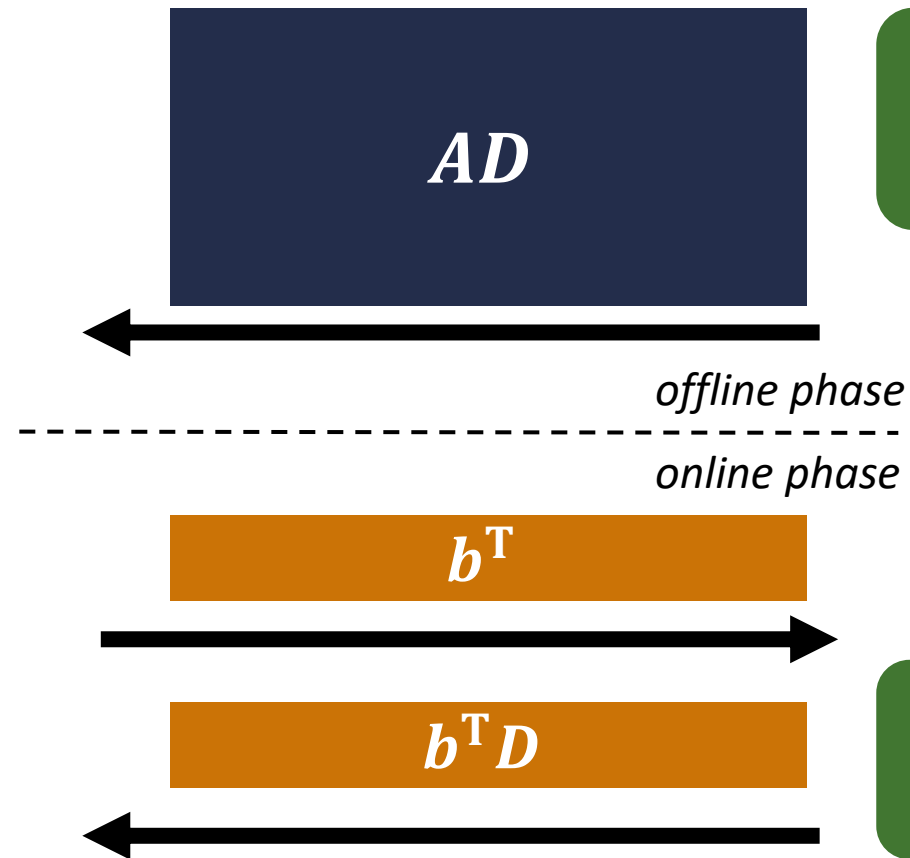
[HHCMV23]

Key insight in SimplePIR: A is query-independent so move computation of AD offline

Limitation: hint is large ($n \times$ larger than response) and needs to be updated whenever database changes



Query and response size:
 $\sqrt{N}$ elements over $\mathbb{Z}_q$



Very **GPU-friendly**: offline phase and online phase consist of matrix multiplications



Database
 $D \in \mathbb{Z}_p^{\sqrt{N} \times \sqrt{N}}$

On modern processors, can process query at *memory bandwidth*

2 operations per database entry

PIR with *Silent* Preprocessing

Silent preprocessing: allow **offline** preprocessing, but **no communication**

No need to manage hints; better suited for dynamic databases

Approach: *compress* the hint and include as part of the response

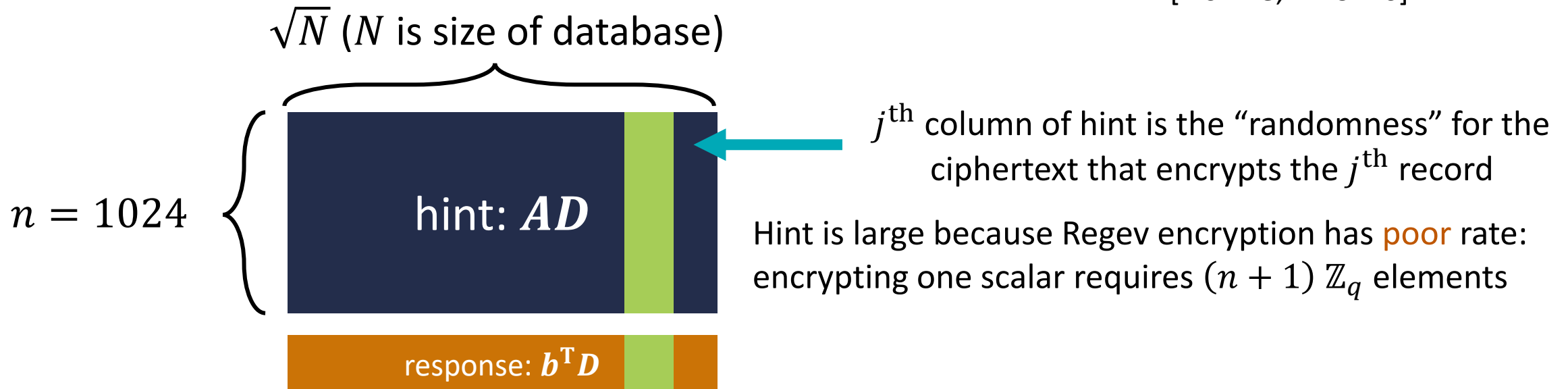
[Tiptoe; HDCZ23]

[HintlessPIR; LMRS24]

[YPIR; MW24]

[InsPIRe; MPSY26]

Why is the SimplePIR hint so large?



Higher Rate using Polynomial Rings

Common technique to get better rate: use polynomial rings

Ciphertext space:

$$\mathbb{Z}_q^{n+1} \longrightarrow R_q^2 = \left(\mathbb{Z}_q[x] / (x^n + 1) \right)^2$$

degree- n polynomials with coefficients in $\mathbb{Z}_q$

Each ciphertext now encrypts a **degree- n polynomial** (with coefficients in $\mathbb{Z}_p$)

$$\text{rate} = \frac{\text{size of plaintext}}{\text{size of ciphertext}} = \frac{\log p}{(n+1) \log q} \longrightarrow \frac{n \log p}{2n \log q} = \frac{\log p}{2 \log q}$$

*rings allow **packing** more data into each ciphertext*

Recall $n = 2^{10} = 1024$, so using rings gives over 500× reduction in size

YPIR/InsPIRe: Packing LWE Ciphertexts

[MW24, MPSY26]

a_1	b_1
a_2	b_2
$\vdots$	
a_n	b_n

LWE ciphertexts under the
secret key $\mathbf{s} \in \mathbb{Z}_q^n$

Regev invariant: $b_i - \mathbf{s}^T \mathbf{a}_i \approx \lfloor q/2 \rfloor \cdot \mu_i$

Can round to recover μ_i

YPIR/InsPIRe: Packing LWE Ciphertexts

[MW24, MPSY26]

$\mathbf{a}_1$	b_1
$\mathbf{a}_2$	b_2
$\vdots$	
$\mathbf{a}_n$	b_n

LWE ciphertexts under the
secret key $\mathbf{s} \in \mathbb{Z}_q^n$

Reinterpret $\mathbf{a}_i \in \mathbb{Z}_q^n$ as coefficients of a polynomial $\tilde{a}_i \in R_q$

$$\tilde{a}_i = a_{i,0} + a_{i,1}x + \cdots + a_{i,n-1}x^{n-1} \in R_q$$

Reinterpret $b_i \in \mathbb{Z}_q$ as a constant polynomial $\tilde{b}_i \in R_q$

Reinterpret $\mathbf{s} \in \mathbb{Z}_q^n$ as coefficients of polynomial $\tilde{s} \in R_q$ (in nega-cyclic order)

$$\tilde{s} = s_0 + s_1x^{-1} + \cdots + s_{n-1}x^{-n+1} \in R_q$$

$$R_q = (\mathbb{Z}_q[x])/(x^n - 1)$$

Observation: Constant term of $\tilde{b}_i - \tilde{s}\tilde{a}_i = b_i - \mathbf{s}^T \mathbf{a}_i \approx \lfloor q/2 \rfloor \cdot \mu_i$

$(\tilde{a}_i, \tilde{b}_i)$ is an RLWE ciphertext that decrypts to a value with the right constant term

Suppose we can homomorphically zero-out the high-order terms

$(\tilde{a}'_i, \tilde{b}'_i)$ is an RLWE ciphertext that decrypts to $b_i - \mathbf{s}^T \mathbf{a}_i$

Packed ciphertext is then $(\sum_{i \in [n]} \tilde{a}'_i x^{i-1}, \sum_{i \in [n]} \tilde{b}'_i x^{i-1})$, which decrypts to $\approx \sum_{i \in [n]} \lfloor q/2 \rfloor \cdot \mu_i x^{i-1}$

Homomorphic Trace Evaluation

Input: $(\tilde{a}_i, \tilde{b}_i)$ is an RLWE ciphertext that encrypts a polynomial $\mu_i \in R_p$

Output: $(\tilde{a}'_i, \tilde{b}'_i)$ is an RLWE ciphertext that encrypts the constant $\mu_i(0) \in \mathbb{Z}_p$

Let $R = (\mathbb{Z}[x])/(x^n + 1)$, where $n > 4$ is a power-of-two

Fact: Trace operator $\text{Tr}: R \rightarrow R$ over R maps $r \in R$ to $n \cdot r(0)$

InsPIRe [MPSY26]: Trace operation over R can be written as

$$\text{Tr}(r) = \sum_{j=0}^{n/2-1} \left(\kappa_5^j(r) + \kappa_{2n-1} \left(\kappa_5^j(r) \right) \right)$$

κ_i are automorphisms over R

$$\begin{aligned} \kappa_5: \quad r(x) &\mapsto r(x^5) \\ \kappa_{2n-1}: r(x) &\mapsto r(x^{2n-1}) \end{aligned}$$

Technically, the trace is the sum over *all* automorphisms $\kappa_0, \dots, \kappa_{n-1}$.

Over R , the automorphism group is generated by κ_5 and κ_{2n-1}

Straightforward to **homomorphically** evaluate automorphisms on RLWE ciphertexts

InsPIRe: query includes keys to perform automorphisms; server sends back packed response

Homomorphic Trace Evaluation

Input: $(\tilde{a}_i, \tilde{b}_i)$ is an RLWE ciphertext that encrypts a polynomial $\mu_i \in R_p$

Output: $(\tilde{a}'_i, \tilde{b}'_i)$ is an RLWE ciphertext that encrypts the constant $\mu_i(0) \in \mathbb{Z}_p$

Let $R = (\mathbb{Z}[x])/(x^n + 1)$, where $n > 4$ is a power-of-two

Fact: Trace operator $\text{Tr}: R \rightarrow R$ over R maps $r \in R$ to $n \cdot r(0)$

InsPIRe [MPSY26]: Trace operation over R can be written as

$$\text{Tr}(r) = \sum_{j=0}^{n/2-1} \left(\kappa_5^j(r) + \kappa_{2n-1} \left(\kappa_5^j(r) \right) \right)$$

κ_i are automorphisms over R

$$\begin{aligned} \kappa_5: \quad r(x) &\mapsto r(x^5) \\ \kappa_{2n-1}: r(x) &\mapsto r(x^{2n-1}) \end{aligned}$$



Is this a matrix multiplication?

Distributional PIR [LHC25]: perform packing on the **CPU** (considered HintlessPIR approach)

VIPIR [KJWGSA26]: design alternative GPU-friendly packing (**requires offline communication**)

Spoiler: It is!

A Deep Dive into the Trace Computation

Suppose $(\tilde{a}, \tilde{b})$ is an RLWE encryption of μ with secret key $\tilde{s}$: $\tilde{b} - \tilde{s}\tilde{a} \approx \mu$

dropping scaling factor here

For any automorphism κ , this means $\kappa(\tilde{b}) - \kappa(\tilde{s}) \cdot \kappa(\tilde{a}) \approx \kappa(\mu)$

$(\kappa(\tilde{a}), \kappa(\tilde{b}))$ is an encryption of $\kappa(\mu)$ under the transformed key $\kappa(s)$

To support homomorphic evaluation of κ , client provides a key-switching key [GHS12, BGV12]

$(\tilde{\mathbf{a}}_{\text{ks}}, \tilde{\mathbf{b}}_{\text{ks}}) = (\tilde{\mathbf{a}}_{\text{ks}}, \tilde{s}\tilde{\mathbf{a}}_{\text{ks}} + \mathbf{e} - \kappa(s) \cdot \mathbf{g})$ $\mathbf{g} = [2^0, 2^1, \dots, 2^{\lceil \log q \rceil - 1}]$ is the gadget vector [MP12]

$\mathbf{g}^{-1}(\cdot)$ is bit-decomposition operator

Key switching is a **linear** function

New ciphertext under $\tilde{s}$: $(\tilde{\mathbf{a}}_{\text{ks}}^T \cdot \mathbf{g}^{-1}(\kappa(\tilde{a})), \kappa(\tilde{b}) + \tilde{\mathbf{b}}_{\text{ks}}^T \cdot \mathbf{g}^{-1}(\kappa(\tilde{a})))$

Key observation: some components can be fixed, others depend on the query

[LMRS24]

A Deep Dive into the Trace Computation

Suppose $(\tilde{a}, \tilde{b})$ is an RLWE encryption of μ with secret key $\tilde{s}$: $\tilde{b} - \tilde{s}\tilde{a} \approx \mu$

For any automorphism κ , this means $\kappa(\tilde{b}) - \kappa(\tilde{s}) \cdot \kappa(\tilde{a}) \approx \kappa(\mu)$

$(\kappa(\tilde{a}), \kappa(\tilde{b}))$ is an encryption of $\kappa(\mu)$ under the transformed key $\kappa(s)$

To support homomorphic evaluation of κ , client provides a key-switching key [GHS12, BGV12]

$$(\tilde{\mathbf{a}}_{\text{ks}}, \tilde{\mathbf{b}}_{\text{ks}}) = (\tilde{\mathbf{a}}_{\text{ks}}, \tilde{s}\tilde{\mathbf{a}}_{\text{ks}} + \mathbf{e} - \kappa(s) \cdot \mathbf{g}) \quad \mathbf{g} = [2^0, 2^1, \dots, 2^{\lceil \log q \rceil - 1}] \text{ is the gadget vector} \quad [\text{MP12}]$$

$\mathbf{g}^{-1}(\cdot)$ is bit-decomposition operator

Key switching is a **linear** function

$$\text{New ciphertext under } \tilde{s}: \left(\tilde{\mathbf{a}}_{\text{ks}}^T \cdot \mathbf{g}^{-1}(\kappa(\tilde{a})), \kappa(\tilde{b}) + \tilde{\mathbf{b}}_{\text{ks}}^T \cdot \mathbf{g}^{-1}(\kappa(\tilde{a})) \right)$$

Key observation: some components **can be fixed**, others **depend on the query**

[LMRS24]

A Deep Dive into the Trace Computation

$$\text{Tr}(r) = \sum_{j=0}^{n/2-1} \left(\kappa_5^j(r) + \kappa_{2n-1} \left(\kappa_5^j(r) \right) \right)$$

What does this look like homomorphically? Say $(\tilde{a}, \tilde{b})$ encrypts r and (a^*, b^*) encrypts $\text{Tr}(r)$

Observation 1: a^* can be computed **offline**; only b^* needs to be computed **online**

Observation 2: Each key-switch is an inner product with something that only depends on $\tilde{a}$

$$\text{Key switching: } \left(\tilde{a}_{\text{ks}}^T \cdot g^{-1}(\kappa(\tilde{a})), \kappa(\tilde{b}) + \tilde{b}_{\text{ks}}^T \cdot g^{-1}(\kappa(\tilde{a})) \right)$$

A Deep Dive into the Trace Computation

$$\text{Tr}(r) = \sum_{j=0}^{n/2-1} \left(\kappa_5^j(r) + \kappa_{2n-1} \left(\kappa_5^j(r) \right) \right)$$

What does this look like homomorphically? Say $(\tilde{a}, \tilde{b})$ encrypts r and (a^*, b^*) encrypts $\text{Tr}(r)$

Observation 1: a^* can be computed **offline**; only b^* needs to be computed **online**

Observation 2: Each key-switch is an inner product with something that only depends on $\tilde{a}$

$$\tilde{b}^* = \tilde{b} + \sum_{j \in [n/2-1]} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_j) + \sum_{j \in [n/2-1]} \left(\kappa_{2n-1} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2+j}) + \tilde{\mathbf{b}}_{\text{ks},2n-1}^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2})$$

All of the α_j are fixed functions of $\tilde{a}, \tilde{a}_{\text{ks},5}, \tilde{a}_{\text{ks},2n-1}$



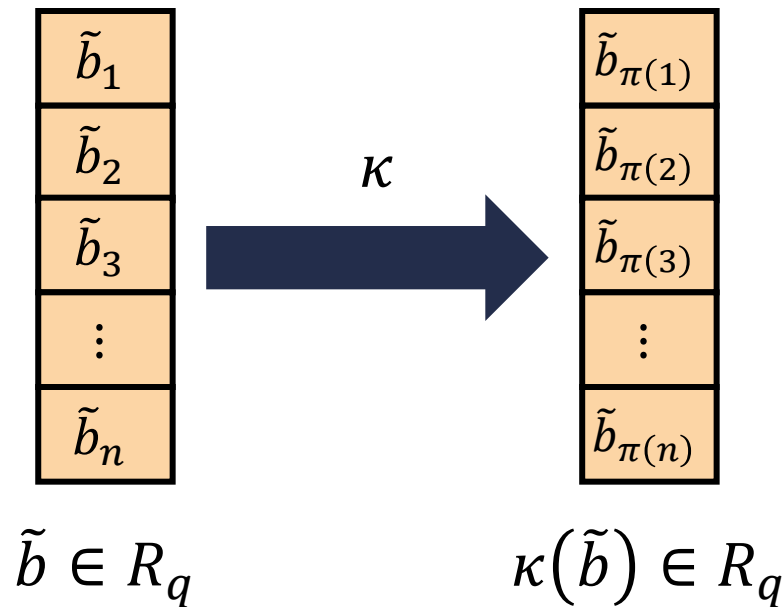
Is this a matrix multiplication?

A Deep Dive into the Trace Computation

$$\tilde{b}^* = \tilde{b} + \sum_{j \in [n/2-1]} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_j) + \sum_{j \in [n/2-1]} \left(\kappa_{2n-1} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2+j}) + \tilde{\mathbf{b}}_{\text{ks},2n-1}^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2})$$

all linear functions in $\tilde{\mathbf{b}}_{\text{ks},5}$

View polynomials in their NTT representation (i.e., in evaluation representation)



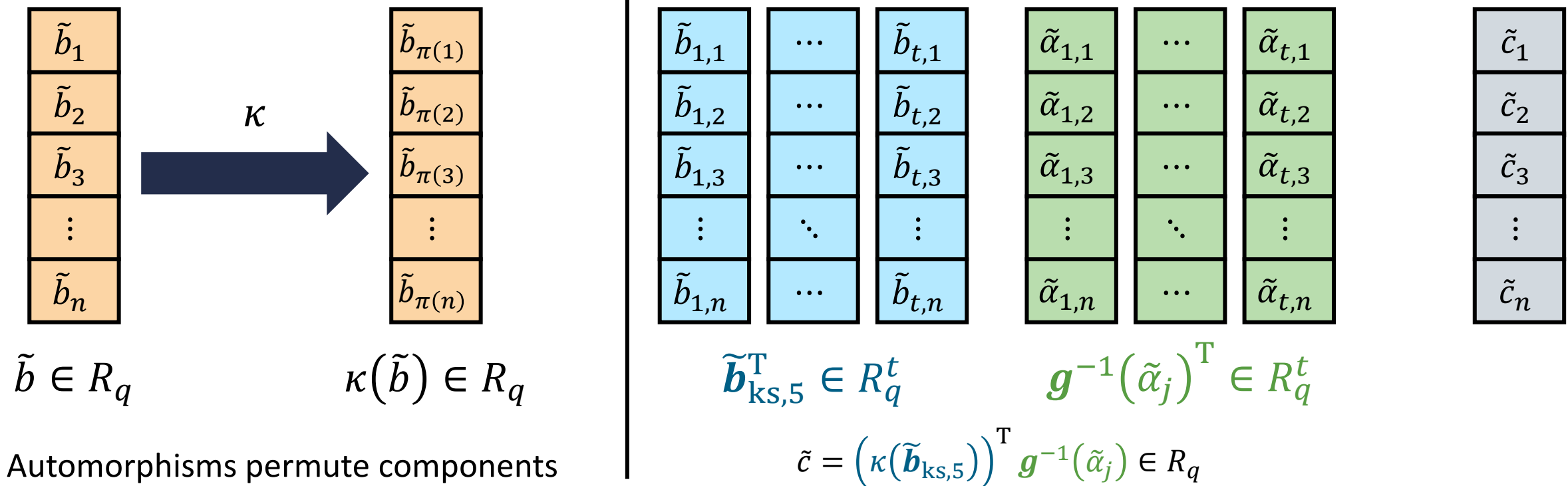
Multiplication is component-wise
in NTT representation

Automorphisms permute components

A Deep Dive into the Trace Computation

$$\tilde{b}^* = \tilde{b} + \sum_{j \in [n/2-1]} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_j) + \sum_{j \in [n/2-1]} \left(\kappa_{2n-1} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2+j}) + \tilde{\mathbf{b}}_{\text{ks},2n-1}^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2})$$

View polynomials in their NTT representation (i.e., in evaluation representation)



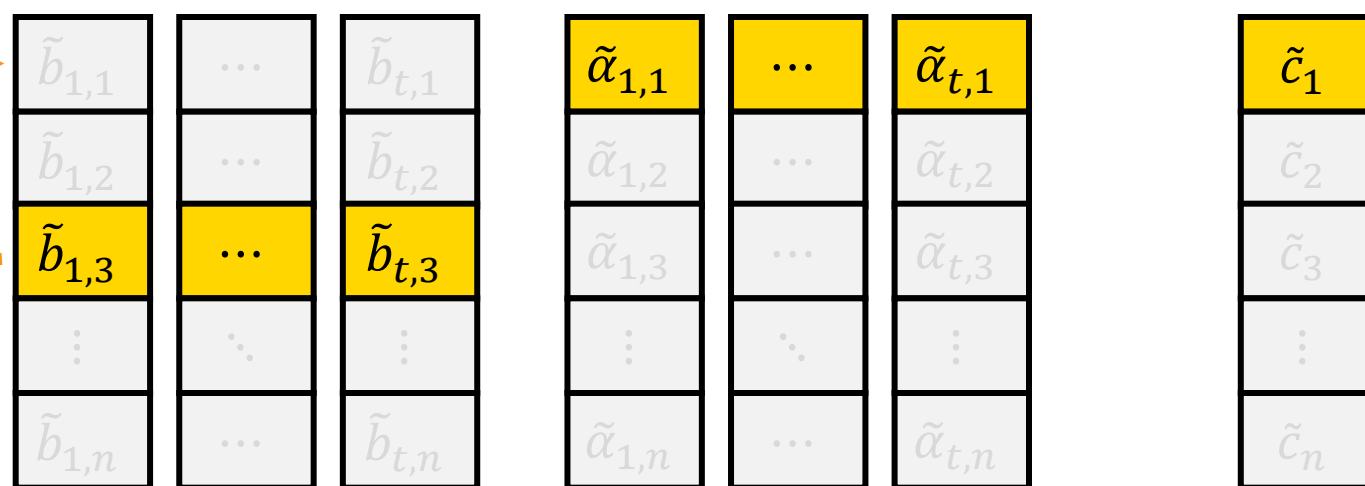
A Deep Dive into the Trace Computation

$$\tilde{b}^* = \tilde{b} + \sum_{j \in [n/2-1]} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_j) + \sum_{j \in [n/2-1]} \left(\kappa_{2n-1} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2+j}) + \tilde{\mathbf{b}}_{\text{ks},2n-1}^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2})$$

View polynomials in their NTT representation (i.e., in evaluation representation)

$$\pi(3) = 1$$

Observation: each component $\tilde{c}_i$ is an inner product between one “row” of the NTT representation of $\tilde{\mathbf{b}}_{\text{ks},5}$ and a fixed vector



$$\tilde{\mathbf{b}}_{\text{ks},5}^T \in R_q^t$$

$$\mathbf{g}^{-1}(\tilde{\alpha}_j)^T \in R_q^t$$

$$\tilde{c} = \left(\kappa(\tilde{\mathbf{b}}_{\text{ks},5}) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_j) \in R_q$$

A Deep Dive into the Trace Computation

$$\tilde{b}^* = \tilde{b} + \sum_{j \in [n/2-1]} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_j) + \sum_{j \in [n/2-1]} \left(\kappa_{2n-1} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2+j}) + \tilde{\mathbf{b}}_{\text{ks},2n-1}^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2})$$

View polynomials in their NTT representation (i.e., in evaluation representation)

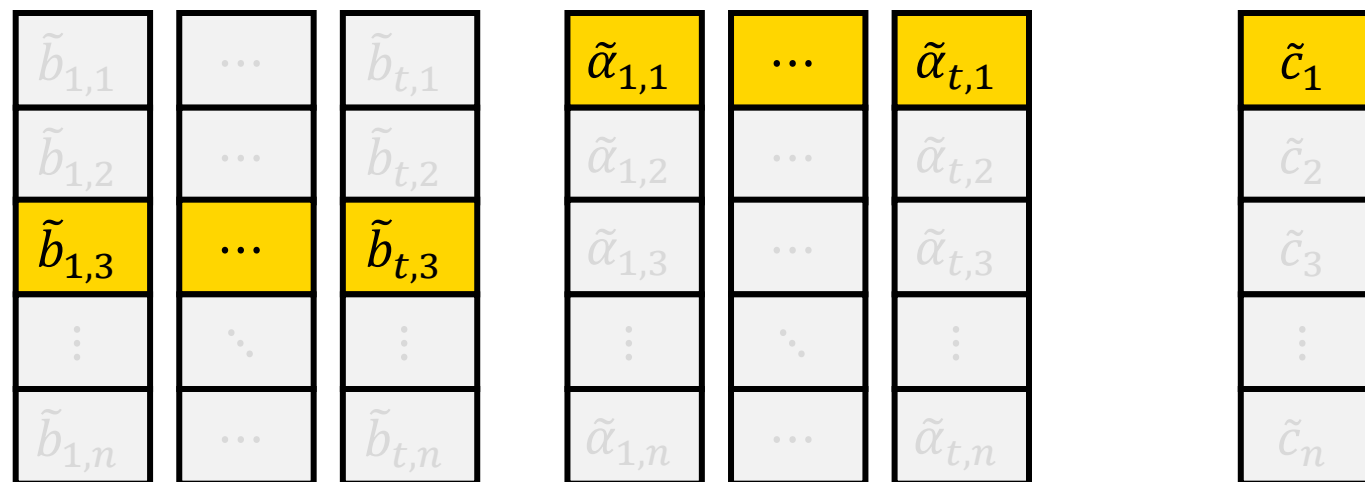
We can write the shaded term as

$$\tilde{b}_i^* = \sum_{k \in [n]} \tilde{\mathbf{b}}_k^T \mathbf{u}_{i,k}$$

$\tilde{\mathbf{b}}_k^T = k^{\text{th}}$ row of NTT representation of $\tilde{\mathbf{b}}_{\text{ks},5}^T$

$\mathbf{u}_{i,k}$: fixed vector (determined in offline phase)

This is now a matrix multiplication!



$$\tilde{\mathbf{b}}_{\text{ks},5}^T \in R_q^t$$

$$\mathbf{g}^{-1}(\tilde{\alpha}_j)^T \in R_q^t$$

$$\tilde{c} = \left(\kappa(\tilde{\mathbf{b}}_{\text{ks},5}) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_j) \in R_q$$

A Deep Dive into the Trace Computation

$$\tilde{b}^* = \tilde{b} + \sum_{j \in [n/2-1]} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_j) + \sum_{j \in [n/2-1]} \left(\kappa_{2n-1} \left(\kappa_5^{j-1}(\tilde{\mathbf{b}}_{\text{ks},5}) \right) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2+j}) + \tilde{\mathbf{b}}_{\text{ks},2n-1}^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2})$$

View polynomials in their NTT representation (i.e., in evaluation representation)

We can write the shaded term as

$$\tilde{b}_i^* = \sum_{k \in [n]} \tilde{\mathbf{b}}_k^T \mathbf{u}_{i,k}$$

$\tilde{\mathbf{b}}_k^T = k^{\text{th}}$ row of NTT representation of $\tilde{\mathbf{b}}_{\text{ks},5}$
 $\mathbf{u}_{i,k}$: fixed vector (determined in offline phase)

This is now a matrix multiplication!

To compute the NTT representation of the shaded term:

$$\hat{\mathbf{b}}^T = [\tilde{\mathbf{b}}_1^T \mid \cdots \mid \tilde{\mathbf{b}}_n^T] \in \mathbb{Z}_q^{nt}$$

$$\hat{\mathbf{U}} = \begin{bmatrix} \mathbf{u}_{1,1} & \cdots & \mathbf{u}_{n,1} \\ \vdots & \ddots & \vdots \\ \mathbf{u}_{1,n} & \cdots & \mathbf{u}_{n,n} \end{bmatrix} \in \mathbb{Z}_q^{nt \times n}$$

NTT representation is then $\hat{\mathbf{b}}^T \hat{\mathbf{U}}$

To handle multiple independent queries, stack $\hat{\mathbf{b}}_1^T, \dots, \hat{\mathbf{b}}_K^T$ into a matrix $\hat{\mathbf{B}}$ and compute $\hat{\mathbf{B}} \hat{\mathbf{U}}$

Back to PIR

$$\tilde{b}^* = \tilde{b} + \sum_{j \in [n/2-1]} \left(\kappa_5^{j-1}(\tilde{b}_{\text{ks},5}) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_j) + \sum_{j \in [n/2-1]} \left(\kappa_{2n-1} \left(\kappa_5^{j-1}(\tilde{b}_{\text{ks},5}) \right) \right)^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2+j}) + \tilde{b}_{\text{ks},2n-1}^T \mathbf{g}^{-1}(\tilde{\alpha}_{n/2})$$

SimplePIR
response

Packing component

Packing
component

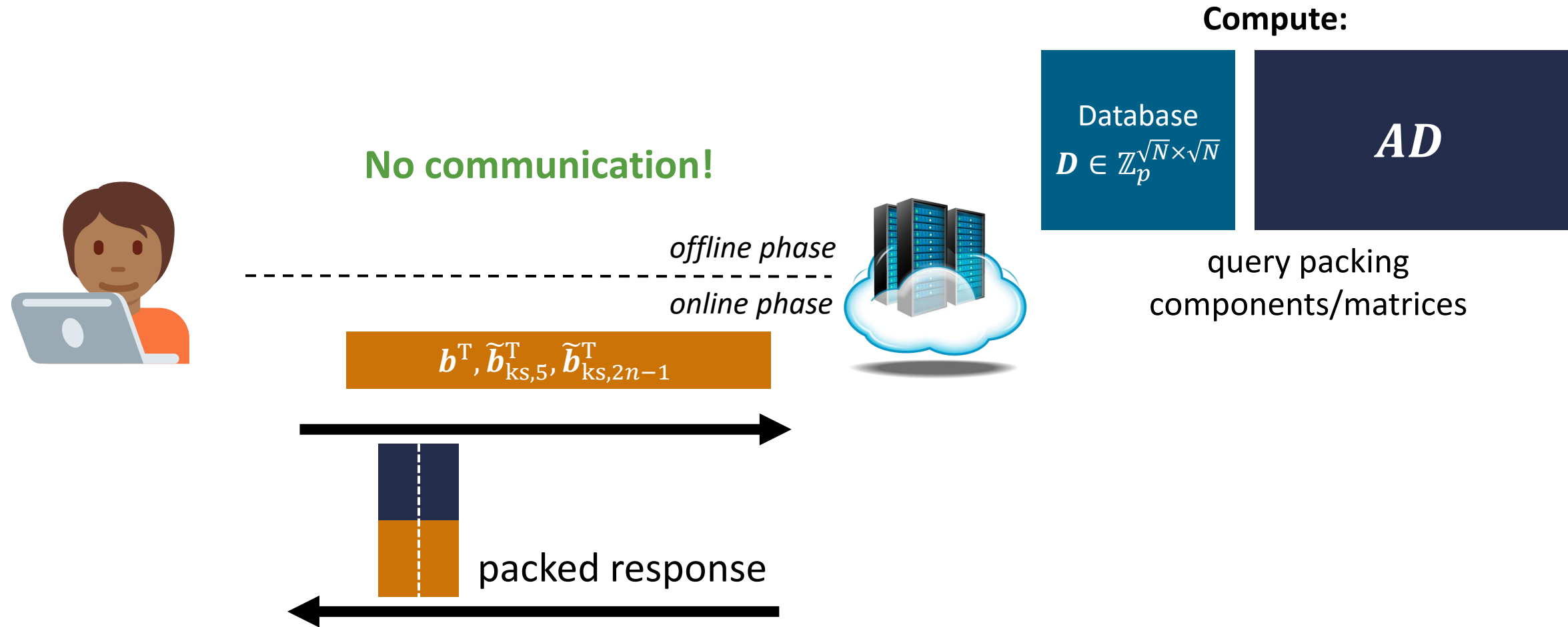
Both the SimplePIR response and the packing components are matrix-vector product between the **client query** and either the **database** or a **matrix in the public parameters**

Both operations are GPU-friendly (in **parallel** even)

Many independent queries: matrix-vector multiplication $\rightarrow$ matrix-matrix multiplication

For fast implementation, need to use modulus switching (power-of-two modulus for SimplePIR, NTT-friendly modulus for packing)

Overall Protocol Flow



Performance Comparison

Setup: 4 GB database, 32 KB records

		CPU-based		Implementation from [LHC25]	GPU-based	
		HintlessPIR	InsPIRe		SimplePIR	This Work
Offline	Download	—	—		144 MB	—
	Upload	722 KB	1.04 MB		512 KB	576 KB
Online	Download	2.97 MB	96 KB		40 KB	112 KB
	Response Time	1.4 s	718 ms		33 ms	8.2 ms
	Throughput	2.8 GB/s	5.6 GB/s		123 GB/s	488 GB/s
Cost (per million queries)		\$489.89	\$114.39		\$20.49	\$14.56

Excludes hint download costs

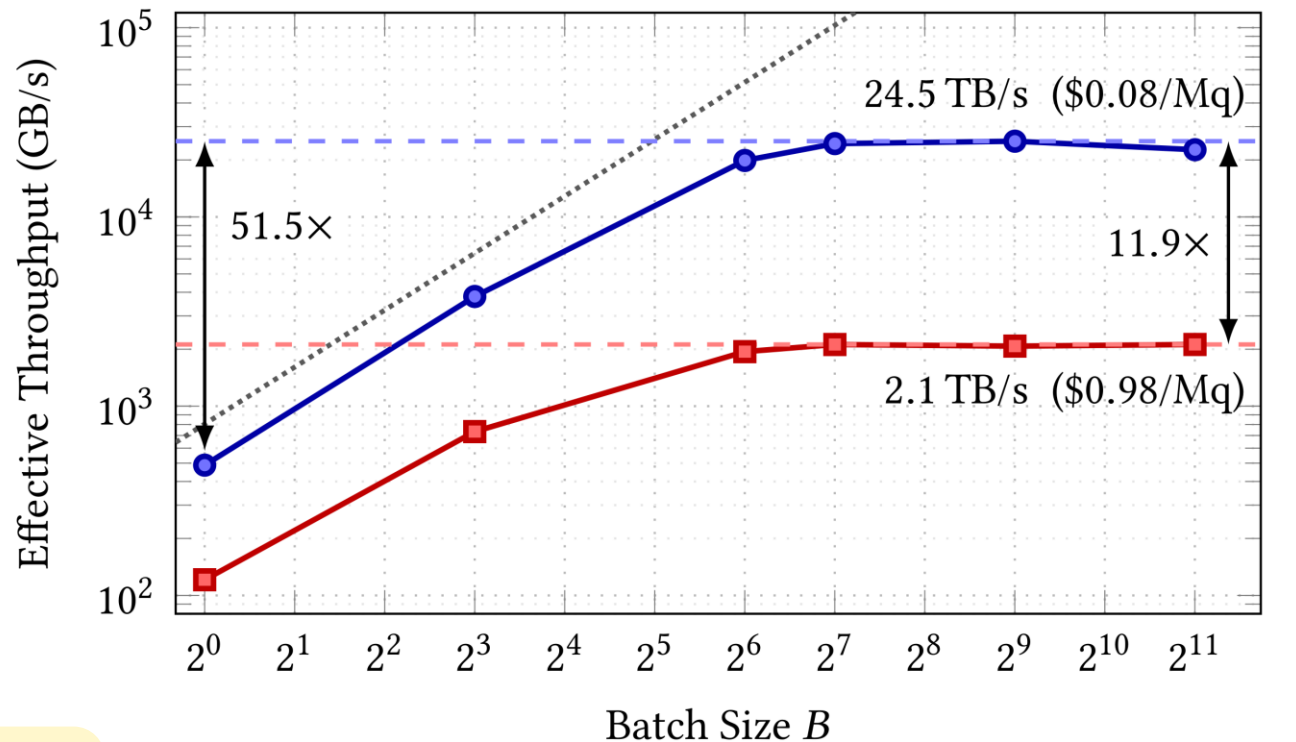
Performance Comparison

Setup: 4 GB database, 32 KB records

Setup: 4 GB database, 32 KB records						
		CPU-based		Implementation from [LHC25]	GPU-based	
		HintlessPIR	InsPIRe	SimplePIR		
Offline	Download	—	—	144 MB	72 KB is actually fixed (can be reused across queries); reduces per-query cost to \$7.93	
	Upload	Over 88 × higher throughput than CPU-based implementations for single-query setting (but we are not compute-bound yet)		512 KB		
Online	Download			40 KB		112 KB
	Response Time			33 ms		8.2 ms
	Throughput	2.8 GB/s	5.6 GB/s	123 GB/s		488 GB/s
Cost (per million queries)		\$489.89	\$114.39	\$20.49	\$14.56	
Excludes hint download costs						

Cross-Client Batching

Cross-client batching: process multiple independent queries at once (more compute saturation)



Effective throughput
peaks around 25 TB/s

On typical CPUs, we expect
 $\approx 25 \text{ GB/s/core}$

*(could be faster with
special instruction sets)*

$k \cdot |\text{DB}|$
Server time

—●— SandwichPIR
—■— SimplePIR (int32 Kernel)
(int32: not using tensor cores)
..... Memory bandwidth

Setup: 4 GB database, 32 KB records

Application to Private Wikipedia

Goal: read (text-only) Wikipedia over PIR

As of November 2023, English Wikipedia contains $6.4 \cdot 10^6$ articles
(8 GB after compression)

Record size is 128 KB (compressed articles are all under 128 KB)

Communication: 320 KB upload, 160 KB download

Server processing: ≈ 60 ms



<https://www.cs.utexas.edu/~dwu4/pir-demo>

Hardware Acceleration for **Cryptography**

Modern hardware is shaped by machine learning applications

For machine learning applications, elementary operation is matrix multiplication



*In addition to Boolean circuits, arithmetic circuits, RAM programs, we should also think about “**matrix multiplication programs**”*

Modern hardware designed for “**matrix multiplication programs**”

For PIR: using GPU tensor cores, we can now obtain (effective) throughputs of 25-275 TB/s (275 TB/s uses fleet of 8 GPUs over a 256 GB database)

Thank You!

(paper coming soon)