

CODoH: Privacy-Preserving Caching for Oblivious DNS over HTTPS

Pankaj Niroula
William & Mary
Virginia, USA
pniroula@wm.edu

Lily Gloudemans
William & Mary
Virginia, USA
algloudemans@wm.edu

Aashutosh Poudel
William & Mary
Virginia, USA
apoudel01@wm.edu

Collin MacDonald
William & Mary
Virginia, USA
cmacdonald01@wm.edu

Stephen Herwig
William & Mary
Virginia, USA
smherwig@wm.edu

Abstract

Oblivious DNS over HTTPS (ODOH) enhances DNS privacy by routing queries through a proxy so that no single party can link a client’s IP address with its DNS queries. However, because ODOH encrypts each query individually, the proxy cannot cache responses. We present CODoH (Cacheable Oblivious DNS over HTTPS), a cacheable extension to ODOH that reduces DNS resolution latency while preserving ODOH’s separation of knowledge. CODoH places a proxy-side cache inside a trusted execution environment—namely, an Intel SGX enclave—and uses end-to-end encryption to ensure that the proxy never learns plaintext queries or cached responses. To prevent caching from becoming a new inference surface, CODoH defends against cache-state probing and set-difference (bracketing) attacks with (i) cover responses supplied by the resolver, (ii) batched cache updates that mix many clients’ inserts, and (iii) an oblivious RAM (ORAM) backend that hides cache access patterns. CODoH also enforces correctness and freshness of cached records via resolver authorization and replay protection. We implement CODoH as CoreDNS extensions and evaluate it on an Azure SGX testbed. CODoH reduces median latency by 2.7× over ODOH on cache hits (18.0 vs. 48.0 ms) and by 2.2× under Zipf traffic (22.3 vs. 48.6 ms), with under 1 ms of SGX overhead.

Keywords

DNS, ODOH, Intel SGX, ORAM

1 Introduction

The *Domain Name System (DNS)* [38, 39] is a critical piece of Internet infrastructure: nearly every application session begins by translating human-readable domain names into IP addresses. DNS’s original design emphasized performance and deployability, relying heavily on caching and hierarchical delegation, but it provided neither security nor privacy. Although mechanisms such as DNSSEC [24] improve integrity and authentication, they do not conceal who is querying or what is being queried. Encrypted transports such as

DNS-over-TLS (DoT) [25] and DNS-over-HTTPS (DoH) [23] protect DNS traffic from passive observers on the network path, but they still concentrate sensitive information at the resolver, which learns both the client’s network identifier and the plaintext domain name. This concentration matters in practice because a small number of large public resolvers serve a significant fraction of global DNS traffic [17, 45], creating powerful vantage points for profiling, censorship, and monetization.

Oblivious DNS over HTTPS (ODOH) [29, 46] eliminates the residual privacy gap in encrypted DNS by splitting trust between two roles: an oblivious proxy that learns the client’s identity but not the plaintext query, and an oblivious target resolver that learns the query but not the client. Apple’s iCloud Private Relay [3] uses ODOH for DNS resolution, where Apple operates the proxy, and a partner such as Cloudflare operates the target [16]. This separation, however, comes at a cost: ODOH forfeits one of DNS’s primary performance advantages: shared caching near clients. Each query is encrypted end-to-end to the target, which prevents the proxy from caching plaintext responses. The resulting lack of proxy-side caching increases latency, raises upstream load, and limits deployability, particularly for popular domains and latency-sensitive applications [27, 49]. Restoring proxy-side caching without reintroducing linkability would preserve ODOH’s privacy guarantees while recovering much of DNS’s practical performance benefits.

Caching introduces new side channels. Caching, however, fundamentally changes the security story. A cache introduces persistent, query-dependent state at the proxy, and that state becomes a side channel even if the proxy cannot decrypt queries. In the ODOH setting, an adversarial proxy can manipulate scheduling and concurrency, inject probe traffic as a client, restart cache components to create “fresh windows,” and selectively omit cache inserts—all with the goal of learning something about a victim’s query from how cache state changes over time. Worse, naïve caching enables classic bracketing (set-difference) attacks in which the proxy isolates a victim request, compares cache state before and after, and infers which name was inserted. These attacks do not arise in baseline ODOH because ODOH traffic remains intentionally non-cacheable; they arise precisely because caching creates state. Any solution therefore must (i) keep the proxy from learning plaintext queries and responses, (ii) ensure cached answers remain authorized and fresh, and (iii) prevent cache behavior itself—hits, misses, inserts,

This work is licensed under the Creative Commons Attribution 4.0 International License. To view a copy of this license visit <https://creativecommons.org/licenses/by/4.0/> or send a letter to Creative Commons, PO Box 1866, Mountain View, CA 94042, USA.



Proceedings on Privacy Enhancing Technologies YYYY(X), 1–20
© YYYY Copyright held by the owner/author(s).
<https://doi.org/XXXXXXX.XXXXXXX>

and access patterns—from becoming a reliable distinguisher under an active proxy adversary.

This paper. Our key insight is that ODoH can support caching by isolating only the cache inside a small, remotely attestable trusted component co-located with the proxy, while masking cache-induced side channels so that even an active proxy cannot link a client’s query to specific cache state changes. We present *CODOH* (*Cacheable Oblivious DNS over HTTPS*), an extension that reduces DNS latency while preserving ODoH’s separation of knowledge. CODOH places the cache inside an Intel SGX enclave [26, 36] alongside the proxy and introduces a protocol that allows the target to authorize cache inserts without revealing cache contents.

To address cache-specific inference risks, CODOH combines three techniques: (1) it structures responses to avoid exposing explicit cache hit/miss signals, (2) it expands and mixes cache updates using cover responses and batched commits to prevent attribution to any single request, and (3) it hides cache access patterns with an *oblivious RAM* (ORAM) [20, 47] backend to block architectural leakage. Together, these mechanisms recover caching’s performance benefits while maintaining ODoH’s trust split.

Contributions. We make the following contributions:

- **Design:** We design CODOH, a cacheable extension to ODoH that places proxy-side caching inside an Intel SGX enclave while preserving ODoH’s separation between client identity and query contents.
- **Security mechanisms:** We develop cache-specific defenses against active proxy attacks, including cover responses, batched cache updates, restart/omission handling, and ORAM-backed oblivious caching to mitigate set-difference and cache-probing inference.
- **Correctness and freshness:** We enforce that cached answers remain authorized and fresh via target-authenticated cache inserts and replay protection compatible with SGX’s time limitations.
- **Prototype and evaluation:** We implement CODOH as extensions to CoreDNS [5]—a graduated Cloud Native Computing Foundation (CNCF) [14] project—and evaluate it on an Azure SGX testbed using macrobenchmarks, microbenchmarks, and page-load experiments. CODOH reduces median latency by 2.7× over ODoH on cache hits and by 2.2× under Zipf traffic, with under 1 ms of SGX overhead.

2 Background and Motivation

2.1 ODoH

ODoH mitigates resolver-side linkability by decoupling client identity from query contents. It introduces two roles: a *proxy* that relays encrypted messages and a *target* resolver that decrypts and answers them. Using *Hybrid Public Key Encryption* (HPKE) [6], the client encrypts each query under the target’s public key and sends it through the proxy, which sees only ciphertext and the client’s IP address, while the target sees the plaintext query but only the proxy’s IP. As long as the proxy and target do not collude (or an attacker does not compromise both), neither can link a user to a query. This design, however, comes at a cost: ODoH encrypts each exchange with fresh randomness, preventing the proxy from

caching responses and thereby sacrificing one of DNS’s primary performance advantages: serving popular records from a nearby shared cache.

2.2 Intel SGX

Trusted execution environments (TEEs) protect code and data from a potentially hostile software stack on the same machine. Intel SGX [26, 36] is a TEE that extends the x86-64 architecture to support isolated execution in protected regions called *enclaves*, allowing user-space applications to run securely even in the presence of a malicious operating system or hypervisor, as well as some physical attackers. SGX encrypts and integrity-protects enclave memory when stored in DRAM and decrypts it only within the trusted CPU package, preventing privileged software from observing or modifying enclave state. Although SGX minimizes the trusted computing base by isolating only selected portions of an application, it can introduce compatibility challenges with legacy software, which middleware frameworks [4, 7, 22, 50] help mitigate. Like other TEEs, SGX’s security guarantees can be weakened by side-channel attacks [8, 9, 32–34, 51, 52].

A core feature of SGX is remote attestation, which enables a remote party to verify that specific software is running inside a genuine enclave on authentic hardware. During attestation, the processor signs a measurement of the enclave’s initial state (MRENCLAVE) and the identity of the signing entity (MRSIGNER), producing a verifiable *quote*. The enclave can also bind runtime data (such as a public key) to the quote via a signed *user-data* field, allowing clients to securely bootstrap an authenticated channel to the enclave.

2.3 Prior Work on DNS Privacy

CODOH builds on a long line of DNS-privacy research that tries to break the link between *who* a client is and *what* name it resolves, without making DNS too slow or too expensive to deploy. Prior systems typically trade among three levers: (1) hiding a query among many (with range queries / cover queries), (2) hiding the sender among many (mix networks), and (3) hiding the query from the infrastructure (using trusted hardware or cryptography). CODOH builds directly on these ideas: it preserves ODoH’s low-latency, proxy-based separation of identity and query contents while restoring proxy-side caching using techniques that dilute and obscure per-query observables, without incurring the deployability constraints or high latency of earlier approaches.

Range queries and cover traffic. Early privacy-preserving DNS systems used range (dummy) queries [10, 56, 57] to hide a user’s true lookup among additional queries, thereby creating an anonymity set in which an observer cannot distinguish the real query from decoys. In practice, these approaches incur significant bandwidth and latency overhead (since the client must wait for all resolutions), and require carefully chosen dummy hostnames to avoid revealing inter-query dependencies. CODOH avoids such overheads and client-side concerns by forgoing dummy queries and instead injecting only dummy responses.

Mix networks. Mix networks [11] address a complementary problem by concealing a query’s origin through delayed and shuffled

forwarding across multiple relays. This design offers strong resistance to traffic analysis but incurs substantial latency, making it impractical for DNS, where even small delays degrade user experience. Onion-routing approaches such as DNS-over-Tor [1, 40] similarly impose comparatively high latency for DNS workloads. Hybrid designs [19] therefore often reserve mixing for long-tail queries while serving popular names locally using alternative mechanisms, sometimes including broadcast-style techniques. Rather than multi-hop mixing or broadcast, CODoH pursues the same intuition—make any single query hard to isolate—through cache-centric mechanisms that aggregate and mask cache state changes.

Trusted execution environments for DNS. PDoT [42] takes a different approach by running the resolver inside an Intel SGX enclave to limit what the operator can observe, while using remote attestation to assure clients they are communicating with the protected code. (PDoT also demonstrates an important residual-leakage issue for recursive resolution: even if the recursive resolver is protected, traffic to authoritative name servers can reveal information.) In contrast, CODoH seeks to minimize the trusted computing base by confining the TEE to narrowly scoped caching functionality and combining it with ODoH’s proxy–target separation.

Private Information Retrieval. *Private Information Retrieval (PIR)* [12] offers an alternative “cryptography-only” approach: a client retrieves a record from a server’s database without revealing which record it requested. Applied to DNS, PIR can prevent a resolver from learning query contents even without a proxy. However, PIR-based DNS designs [37, 54] face practical friction: DNS data are dynamic (TTL expiration, frequent updates), caches are not fixed databases, and handling misses and updates tends to require additional protocol machinery. Moreover, PIR can impose significant computation and bandwidth costs at resolver scale. These challenges make PIR an informative point of comparison, but not the design basis for CODoH’s approach.

Summary. Taken together, this prior work suggests a recurring design pattern: systems that provide strong privacy rely on either (i) large anonymity sets created by aggregation (range queries / mixing), or (ii) strong server-side opacity (TEEs or heavy cryptography). CODoH composes these methods while preserving the deployability and latency properties of modern encrypted DNS.

3 Goals, Assumptions, and Threat Model

CODoH *inherits* ODoH’s goals, assumptions, and baseline adversary model, and then *adds* new requirements and adversarial capabilities that arise specifically from caching.

3.1 ODoH Goals and Assumptions (Inherited)

ODoH privacy goal. ODoH aims to prevent any single non-colluding party from learning both (i) a client identifier (i.e., IP address) and (ii) the plaintext DNS query/response. It achieves this separation by splitting roles between a proxy and a target resolver.

ODoH adversary model. ODoH assumes a powerful active network attacker that can observe, replay, delay, drop, and inject messages, but that cannot break standard cryptographic primitives.

ODoH further assumes compromise of *either* the proxy *or* the target may occur, but not both simultaneously (equivalently, the proxy and target do not collude).

ODoH operational assumptions. ODoH makes the following operational assumptions:

- Clients obtain and authenticate the target’s public key material (the bootstrapping mechanism is out of scope of the specification).
- The proxy and target operate under independent administrative control; if they collude or are jointly compromised, unlinkability is lost by design.
- Proxies do not add client-identifying information when forwarding requests (e.g., Forwarded headers).
- The target follows the protocol and does not intentionally introduce client-specific identifiers (e.g., via client-unique configuration) that would enable linkability.

3.2 CODoH: Added Requirements and Stronger Adversary (New)

Caching introduces persistent state at the proxy and makes cache behavior a potential side channel. CODoH therefore adds both a *performance requirement* and *cache-specific security goals*.

Performance requirement. CODoH should reduce client-perceived latency for cacheable DNS records relative to baseline ODoH; client latency for uncached DNS records should be similar to ODoH.

Primary cache-specific attack vector. A cache enables *state-based inference* attacks that do not exist in standard ODoH because ODoH traffic is intentionally non-cacheable. The canonical example is a *set-difference (bracketing) attack* in which a malicious proxy:

- (1) Learns (or approximates) the current cache state.
- (2) Allows one victim query to proceed while queuing other traffic.
- (3) Learns the cache state again.
- (4) Computes the difference to isolate which entry was added, thereby narrowing (and potentially identifying) the victim’s query.

CODoH’s security goals below target this class of attacks and related cache-manipulation strategies.

Security goals. CODoH targets the following goals in addition to ODoH’s baseline unlinkability:

- G1 *Authorized and fresh cached answers.* A malicious proxy must not be able to cause clients to accept cached DNS answers that the target did not authorize, or that are no longer valid. This goal rules out cache poisoning, tampering, and stale-answer replay that would create split views or incorrect resolutions.
- G2 *Cache-state indistinguishability under active probing and scheduling.* The proxy must not be able to learn meaningful information about a particular client’s query by manipulating or comparing cache state across time—e.g., via set-difference/bracketing, cache priming, selective withholding, restart-based “reset windows,” or similar strategies that rely on proxy-controlled scheduling and concurrency.
- G3 *Query equality and profiling resistance.* The proxy must not be able to determine whether two client requests correspond to the same DNS query (or otherwise build per-client/per-group

query profiles) from cache behavior, beyond what is already available from observing client network identifiers and coarse traffic features (as allowed by the out-of-scope items in §3.3).

Threat model. CODoH retains ODoH’s non-collusion assumption. Within that assumption, CODoH considers the following *in-scope adversary capabilities*:

- *Network control at the proxy*: Observe, delay, drop, replay, re-order, inject protocol messages on proxy-facing links.
- *Active proxy deviation*: Arbitrary protocol deviation to bias cache behavior (e.g., priming, withholding inserts, selective forwarding).
- *Scheduling control*: Manipulate concurrency and queuing at the proxy to amplify inference (e.g., isolate a victim in an “epoch”).
- *Cache lifecycle control*: Restart the cache component to reset state.
- *Proxy-as-client probing*: Send probe queries to try to enumerate or test cache contents.

CODoH also considers *architectural* leakage that arises from the cache component’s externally observable behavior, including:

- Page-level access patterns visible via architectural mechanisms (e.g., page faults / paging activity).
- Timing differences between different logical cache operations.
- Message size differences.

CODoH aims to prevent these channels from enabling G2/G3-style inferences.

3.3 Out-of-Scope Threats and Limitations

CODoH inherits several limitations from ODoH and encrypted DNS more broadly, and it makes additional scoping choices for clarity.

Traffic correlation and website fingerprinting. As in ODoH, CODoH does not attempt to defeat global traffic-analysis adversaries that can correlate flows across multiple network vantage points or over long time horizons (e.g., by combining DNS observations with application-layer behavior). Furthermore, CODoH assumes client-side indistinguishability at the network layer. Distinctive client traffic patterns—such as bursts of dependent queries and characteristic timing—may still enable website fingerprinting or related correlation attacks by an observer at or near the proxy.

Salari et al. [44] demonstrate this risk for ODoH: using deep learning over encrypted traffic features (packet sizes, counts, and timing), they show that a passive adversary near the proxy can identify visited websites with high accuracy, and propose TLS-level defenses based on padding, fragmentation, and dummy traffic. CODoH is complementary: whereas Salari et al. analyze and defend against fingerprinting in ODoH, we address the privacy and deployability challenges introduced by adding caching. CODoH can (and does) incorporate traffic-perturbation techniques such as padding to reduce protocol-level distinguishers, but it does not claim to eliminate inference driven by higher-layer client behavior.

Microarchitectural attacks on trusted hardware. CODoH does not defend against microarchitectural side channels or fault attacks [41] on the cache’s trusted-execution substrate, including transient-execution attacks (e.g., Spectre/Meltdown [30, 35, 51])

and other CPU-level leakage [53]. This is distinct from the architectural channels in §3.2—such as page-level access patterns and timing variation—which CODoH mitigates with ORAM and constant-operation cache accesses (see §4.7). Accordingly, microarchitectural side-channel attacks are explicitly out of scope. CODoH’s security therefore depends on the underlying TEE and its mitigations: if such attacks compromise enclave confidentiality or integrity, CODoH’s cache protections may no longer hold.

Malicious target and proxy–target collusion. As in ODoH, CODoH does not ensure DNS correctness against a malicious target, and it does not address active deanonymization in which a malicious target crafts answers to elicit identifying client behavior. Finally, if the proxy and target collude, ODoH’s privacy goal fails by construction; CODoH cannot prevent this and instead relies on operational separation between the two roles.

4 Design

We now incrementally develop the CODoH protocol, introducing mechanisms as needed. Figure 1 shows the system architecture, and Algorithms 1 and 2 present pseudocode for the main steps.

4.1 Syntax and Notation

Table 1 summarizes the notation for the AEAD, HPKE, hash, and signature schemes used in CODoH. In particular, HPKE provides algorithms to establish a context and context methods to seal, open, and derive (export) keys. We use $\$ \rightarrow$ to denote randomized algorithms and $\rightarrow$ for deterministic algorithms.

In ODoH, the client invokes HPKE.SetupS and ctx.Seal to encrypt a DNS query end-to-end to the target’s public key, while the target invokes HPKE.SetupR and ctx.Open to decrypt the query. The domain-separation string `info` ensures that keys derived for different purposes are cryptographically independent, preventing cross-protocol or cross-use key reuse; for example, ODoH uses the fixed value `"odoh query"` when deriving keys for query encryption. In both ODoH and CODoH, the `aad` field binds unencrypted protocol metadata—such as nonces and cryptographic suite identifiers—to the encrypted payload. We use the input-generic function `MakeAAD` in Algorithm 1 to encapsulate `aad` generation. Because the nonce (an input to the `aad`) may either be available to the higher-level code or internal to the HPKE context, `MakeAAD` abstracts over the specific inputs and avoids complicating the pseudocode.

4.2 Base Protocol

Bootstrap. As in baseline ODoH, we assume that clients securely obtain the target’s public HPKE key pk_T before issuing queries. For CODoH, we also assume the target generates a signing key pair $(\text{pk}_{T_\sigma}, \text{sk}_{T_\sigma})$, and that clients securely obtain the corresponding public key. In addition, upon startup the enclave generates a key pair $(\text{pk}_E, \text{sk}_E)$ and produces an attestation quote that cryptographically binds the public key pk_E to the enclave’s measurement as *user-data*. Clients retrieve the enclave’s public key and validate its attestation quote before use. The mechanism for discovering and distributing this keying material is outside the scope of this paper, but likely involves fetching it from a well-known service.

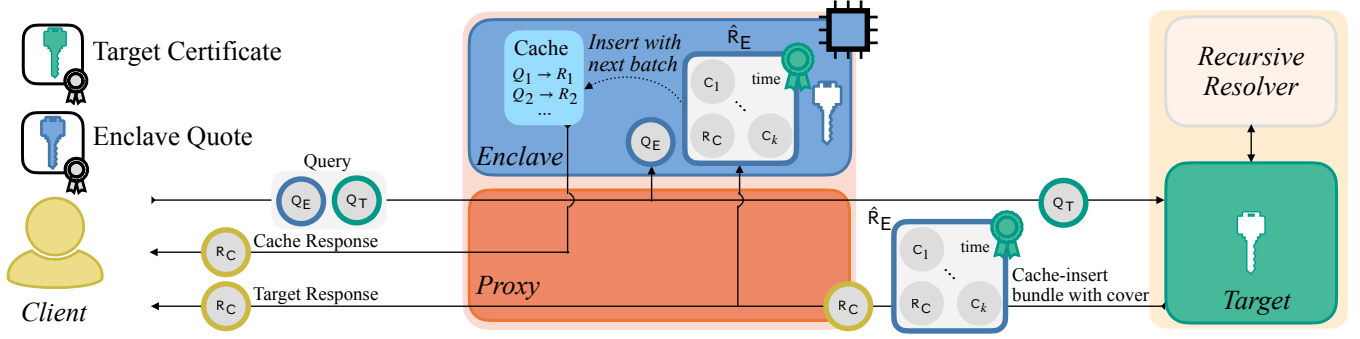


Figure 1: CODoH architecture. Prior to issuing queries, the client retrieves and verifies the target’s certificate and the enclave’s attestation quote (binding their respective public keys). In this diagram, the target is co-located with a recursive resolver.

Table 1: Syntax for AEAD, HPKE, hashes, and signatures.

AEAD	
AEAD.Enc($k, \text{nonce}, \text{aad}, m$) $\rightarrow c$	On input a secret key k , a nonce, associated data aad , and a plaintext m , outputs a ciphertext c that authenticates both aad and m .
AEAD.Dec($k, \text{nonce}, \text{aad}, c$) $\rightarrow m \mid \perp$	On input k , nonce, aad , and ciphertext c , outputs the plaintext m if authentication succeeds, or $\perp$ otherwise.
HPKE	
HPKE.SetupS(pk_R, info) $\rightarrow (\text{enc}, \text{ctx}_S)$	Creates a sender context ctx_S and an encapsulation enc for receiver public key pk_R under domain separation string info .
HPKE.SetupR($\text{sk}_R, \text{info}, \text{enc}$) $\rightarrow \text{ctx}_R$	Creates the matching receiver context ctx_R given private key sk_R , info , and enc .
ctx.Seal(aad, pt) $\rightarrow \text{ct}$	AEAD-encrypts plaintext pt with associated data aad using the context’s internal key/nonce schedule.
ctx.Open(aad, ct) $\rightarrow \text{pt} \mid \perp$	AEAD-decrypts and authenticates. On authentication failure, returns the error symbol $\perp$.
ctx.Export(label, ℓ) $\rightarrow k$	Deterministically derives ℓ bytes of key material from the context using an exporter label (domain separation within context).
Hashing & Signatures	
H : $\{0, 1\}^* \rightarrow \{0, 1\}^\ell$	Denotes a cryptographic hash function that maps arbitrary bit strings to strings of ℓ bits.
Sign(sk, hash) $\rightarrow \sigma$	Signs a hash with private key sk , producing signature σ .
We assume nonces are unique per key, and that ciphertexts include the AEAD authentication tag.	

Query. CODoH uses two independent HPKE exchanges per query: one to the target and one to the enclave, resulting in ciphertexts Q_T and Q_E , respectively.

Response. Upon receiving a request, the proxy forwards Q_E to the enclave, which decrypts it to recover the query and checks its in-memory cache. The enclave derives a response key from the request’s HPKE context via Export, encrypts either the cached

response or a cache-miss indicator under that key, and returns the resulting ciphertext to the proxy, which forwards it to the client.

In tandem with forwarding the request to the enclave, the proxy forwards Q_T to the target. The target responds as in baseline ODoH by encrypting the DNS response to the client, forming R_C . In addition, it constructs a *cache-insert bundle* $\hat{R}_E$ encrypted to the enclave’s public key pk_E . This bundle contains the DNS response plus its TTL and timestamp (see §4.4), and is authenticated with the target’s signature. Upon receiving the reply, the proxy forwards the baseline ODoH response R_C to the client and delivers the enclave-encrypted $\hat{R}_E$ to the enclave.

As in baseline DNS, the response includes the original question, which allows the enclave to unambiguously associate each response with its query. The enclave decrypts $\hat{R}_E$, verifies the target’s signature, extracts the query from the DNS response, and inserts the resulting (query, response) pair into the cache.

Message padding. Differences in size (and processing time) between cache-hit and cache-miss responses could reveal cache status to the proxy or a network observer. To prevent this leakage, CODoH pads responses to fixed-size buckets (e.g., the next 1 KB, 2 KB, or 4 KB boundary), and performs the same HPKE operations on the resultant message regardless of whether the enclave’s response is a cache hit or miss message. Since this padding also helps mitigate website fingerprinting attacks (such as those demonstrated by Salari et al. [44]), CODoH likewise pads the query messages.

4.3 Preventing Set-Difference Attacks

Returning both a cache response and a resolver response prevents the proxy from trivially determining whether a query caused a cache hit or miss. However, this mechanism does not stop a *set-difference attack*: the proxy can still (1) enumerate the cache, (2) process a single victim query while queuing all other queries, (3) re-enumerate the cache, and (4) compute the difference between the two states, uncovering the victim query.

Cover responses. A first step in mitigating set-difference attacks is to inject *cover responses* into the cache so that cache state changes are larger and harder to attribute to a single query. Since the enclave cannot generate responses on its own, CODoH uses the resolver as a source of cover traffic. For each query, the cache-insert bundle $\hat{R}_E$

Algorithm 1 Base Protocol

```

function Client.MakeQuery(pkT, pkE, query)
  ▷ Encrypt to target
  (encT, ctxT) ← $ HPKE.SetupS(pkT, "odoh query")
  aadT ← MakeAAD(ctxT)
  ctT ← $ ctxT.Seal(aadT, Pad(query))
  QT ← (encT, ctT)
  kt ← ctxT.Export("odoh response", ℓ)
  ▷ Encrypt to enclave
  (encE, ctxE) ← $ HPKE.SetupS(pkE, "codoh cache query")
  aadE ← MakeAAD(ctxE)
  ctE ← $ ctxE.Seal(aadE, Pad(query))
  QE ← (encE, ctE)
  ke ← ctxE.Export("codoh cache response", ℓ)
  ▷ Return the two HPKE ciphertexts and response keys
  return (QT, kt, QE, ke)

```

```

function Enclave.Lookup(QE, aadquery)
  ctxE ← HPKE.SetupR(skE, "codoh cache query", QE.enc)
  query ← ctxE.Open(aadquery, QE.ct)
  resp ← Cache.Get(Unpad(query))
  ke ← ctxE.Export("codoh cache response", ℓ)
  nonce ← $ {0, 1}ℓ
  aadresp ← MakeAAD(nonce)
  RC ← $ AEAD.Enc(ke, nonce, aadresp, Pad(resp))
  return RC

```

```

function Target.Resolve(QT, aadquery)
  ctxT ← HPKE.SetupR(skT, "odoh query", QT.enc)
  query ← ctxT.Open(aadquery, QT.ct)
  resp ← Resolve(Unpad(query))
  kt ← ctxT.Export("odoh response", ℓ)
  nonce ← $ {0, 1}ℓ
  aadCresp ← MakeAAD(nonce)
  RC ← $ AEAD.Enc(kt, nonce, aadCresp, Pad(resp))
  yield RC
  covers ← ResolveCovers(k)
  bundle ← (resp, covers, timestamp)
  σ ← Sign(skTσ, H(bundle))
  (encE, ctxE) ← $ HPKE.SetupS(pkE, "codoh cache insert")
  aadCresp ← MakeAAD(ctxE)
  ctE ← $ ctxE.Seal(aadCresp, bundle || σ)
  R̂E ← (encE, ctE)
  return R̂E

```

that the resolver returns contains not only the real response, but also k cover responses for randomly selected domains.

Critically, cover domains must be indistinguishable from real queries. If the resolver samples covers uniformly from a "Top 1M" list while the real query targets a long-tail domain, a proxy could identify the long-tail entry as genuine during a set-difference attack. To prevent this, the resolver should draw covers from a mixture distribution that spans both popular and long-tail names. Formally, the cover distribution D must approximate the empirical query

Algorithm 2 Proxy Protocol

```

function Proxy.OnQuery(request)
  (QT, QE) ← request
  SendToTarget(QT)
  SendToEnclave(QE)

function Proxy.AwaitResponses
  ▷ The select statement with an empty condition comes from Go, where it is used to wait on multiple concurrent events.
  switch
    case RCE ← await RecvFromEnclave
    SendToClient(RCE)
    case RCT ← await RecvFromTarget
    SendToClient(RCT)
    case R̂E ← await RecvFromTarget
    SendToEnclave(R̂E)

```

distribution Q , such as by bounding the likelihood ratios $Q(x)/D(x)$ across domain names x (empirically load-bearing; see App. G).

A natural alternative is to cache only an allowlist of popular domains. This reduces cache-state sensitivity but limits utility and adaptability: it misses long-tail reuse, wastes capacity when popularity shifts, and may not match the popularity distribution of specific client populations.

Batching cache updates. Even with cover responses, a set-difference attack remains possible because a single victim query produces only a small cache update. CODoH mitigates this by delaying insertions and committing them in *batches* for a configurable number of queries B , so that each victim's entries are mixed with many others. Specifically, the enclave accumulates candidate insertions from recent queries— B real responses and their Bk associated covers—and applies them atomically in batches of size $B(1+k)$. This per-query $1+k$ upstream fan-out is the dominant bandwidth cost of the defense; §7 quantifies it on the wire.

Batch boundaries define the "epoch" over which cache state changes. If the proxy can force or deterministically predict these epochs, set-difference attacks become significantly easier. For example, if commits occur exactly every B arrivals, the proxy could queue other clients, allow a single victim query to enter a fresh batch, and fill the remaining $B-1$ slots with its own probes, reducing uncertainty about the inserted entries. It could also align cache enumeration with commit times to minimize noise and improve attribution. To prevent this, the enclave selects commit points independently of proxy-controlled events (e.g., probabilistically choosing a commit point around B).¹

To prevent low traffic from delaying inserts and thus committing records with expired TTLs (see Appendix A for an analysis of real-world TTL values), CODoH uses both size and time thresholds for batching: the enclave commits when either (1) the batch accumulates roughly B queries or (2) a maximum delay $T_{\max}$ has

¹In principle, a commit could create a burst of (ORAM, see §4.7) accesses observable to the host. Our design does not require commits to appear as a contiguous "write storm": an implementation can amortize inserts using ORAM batch-write interfaces or spread writes over subsequent fixed-rate maintenance accesses, so that epoch boundaries are not trivially exposed by instantaneous access volume.

elapsed since the first pending insert.² Let B_{eff} denote the number of real queries included in a time-triggered commit; CODoH requires $B_{\text{eff}} \geq B_{\text{min}}$ whenever cache hits are enabled. This yields large, well-mixed batches under high load while bounding delay under low load. When traffic is insufficient to form a safe batch (fewer than B_{min} distinct client queries within T_{max}), the enclave disables cache hits and temporarily falls back to baseline ODoH, thereby avoiding small, easily isolatable cache updates.

4.4 TTLs and Replay Protection

As in baseline DNS, each cached record carries a *time-to-live (TTL)* that determines how long it remains valid. However, Intel SGX does not provide a trusted clock, since all time sources originate outside the enclave. This limitation complicates safe cache expiration and replay protection.

CODoH addresses this challenge by having the resolver include a signed timestamp with every cache insert. Inside the enclave, CODoH maintains a monotonic notion of time by tracking the largest timestamp accepted so far, denoted t_{latest} . The enclave evaluates expiration and freshness relative to this value. To defend against replay attacks—such as a malicious proxy attempting to reinsert stale records to maintain incorrect answers or create split views—the enclave rejects any cache insert with timestamp $t < t_{\text{latest}} - \delta$, where δ is a configurable tolerance (e.g., a few seconds) accounting for benign network delay or message reordering.

4.5 Restart Handling

A malicious proxy can restart the enclave to clear its state and create fresh cache windows that simplify set-difference attacks. To limit this risk, the enclave generates a new key pair (pk'_E, sk'_E) on restart and produces an attestation report that cryptographically binds the new public key pk'_E to the enclave’s measurement. Publishing a new key enables accountability, since clients and the target can detect frequent restarts.

After a restart, the enclave rejects messages encrypted under the old key pk_E and returns a key error. The enclave then enters a warm-up period during which it performs enough batch insertions to refill the cache to a configurable threshold before serving cache hits. During this period, the enclave always returns a cache miss, effectively reverting to baseline ODoH.

4.6 Cache Omission Attacks

A malicious proxy can withhold resolver responses from being inserted into the cache. While this mainly harms performance and accountability, it also creates a privacy risk: by selectively omitting inserts, the proxy can keep the cache in a controlled, predictable state. For instance, during an enumeration attack, the proxy could suppress the responses for its probing queries to avoid polluting the cache with those entries and their associated cover responses, making cache-state changes easier to isolate.

To mitigate this threat, the enclave tracks queries for which it has not received matching resolver responses. If the number of outstanding queries exceeds a threshold, the enclave enters the

same defensive mode used during a restart warm-up, temporarily disabling cache hits and reverting to baseline ODoH behavior.

4.7 Oblivious Caching

Although SGX protects data confidentiality and integrity, it does not conceal memory access patterns. An OS-level adversary can mount controlled-channel attacks [55], observing which pages the enclave accesses and linking repeated queries to the same cache entry—a form of query-equality leakage absent in standard ODoH. Because the proxy can also issue its own queries, it can compare these access patterns with its probes to identify the client’s query, effectively turning the cache into a query-equality oracle. CODoH mitigates this by implementing the cache atop ORAM—which makes access patterns computationally indistinguishable across logical queries—and by performing a fixed number of ORAM operations per lookup or insert (or using constant-time mode) to reduce timing variation.

5 Security Analysis

This section argues that CODoH meets its cache-specific security goals (G1–G3) against the stronger proxy-side adversary introduced by caching. Our arguments assume standard cryptographic security of HPKE and authenticated encryption, unforgeability of the target’s signature on cache inserts, and the access-pattern indistinguishability of the ORAM construction.

5.1 Authorized and Fresh Cached Answers (G1)

Authorized cached answers. CODoH ensures that the proxy cannot cause clients to accept a cached DNS response that the target did not authorize. Concretely, the target produces a cache-insert bundle $\hat{R}_E$ encrypted to the enclave’s public key and authenticated with a target signature. The enclave verifies this signature before inserting entries into the cache. Under EUF-CMA unforgeability of the target’s signature scheme, a malicious proxy cannot generate a new cache-insert bundle that the enclave will accept without the target’s signing key. Under IND-CCA security of HPKE and the authenticated encryption protecting this ciphertext, the proxy also cannot tamper with a valid cache insert—such as modifying the embedded response—without detection during verification. Consequently, any answer served from the enclave cache must originate from a target-authorized DNS response.

Freshness and replay resistance. A malicious proxy could replay stale cache inserts to retain outdated records or create inconsistent views. CODoH prevents such replay by requiring the target to attach a signed timestamp to each cache insert and by maintaining a monotonic “latest accepted time” within the enclave. Let t_{latest} denote the largest timestamp accepted so far. The enclave rejects any insert with timestamp $t < t_{\text{latest}} - \delta$, where δ allows for minor reordering and network delay. Since the proxy cannot forge the target’s signature, it cannot create fresh-looking stale records; it can only replay old inserts, which fall outside the acceptance window as t_{latest} advances. This limits replay to a bounded window of size δ and prevents indefinite stale reinsertion. If the proxy restarts the enclave and clears its state, the enclave generates a new key pair (pk'_E, sk'_E) on restart, thus rejecting any stale inserts encrypted under the old key.

²To avoid making time-triggered commits fully predictable under low traffic, the enclave may randomize the timeout (e.g., $T_{\text{max}} + \Delta$, for small Δ drawn from a bounded distribution).

5.2 Cache-state Indistinguishability under Active Probing (G2)

We analyze CODoH’s primary cache-specific threat: a set-difference attack in which a malicious proxy attempts to isolate a victim query by comparing cache state across time. CODoH reduces the value of any cache-state snapshot by ensuring that each commit event changes cache state by many entries whose identities remain unknown to the proxy.

Model. CODoH’s policy is to commit a batch to the cache only when it can form a safe batch (at least $B_{\min}$ unique client queries within $T_{\max}$). When the enclave cannot form a safe batch within $T_{\max}$, it disables cache hits and reverts to baseline ODoH; in this mode, the proxy does not obtain a cache-hit oracle. Let $B_{\text{eff}} \geq B_{\min}$ denote the number of effective queries in the batch commit.

Worst-case adversary scheduling. CODoH explicitly considers a proxy that can queue traffic and inject its own queries. In the worst case, the proxy causes a batch to contain exactly one victim query and $B_{\text{eff}} - 1$ adversary-chosen queries. The proxy knows the domains it queried for those $B_{\text{eff}} - 1$ adversary queries, but it does not know any of the k cover domains selected by the target (for either victim or adversary queries), because covers appear only inside the target-to-enclave encrypted bundle.

Let S denote the inserted-domain multiset produced by the batch commit, and let the proxy remove from S the $B_{\text{eff}} - 1$ known adversary real domains. The remaining candidate multiset S' contains: (1) the victim’s real domain (one element), and (2) the $B_{\text{eff}}k$ cover domains. Ignoring accidental duplicate domains (we discuss collisions below), $|S'| = 1 + B_{\text{eff}}k$.

Identification bound under indistinguishable covers. If cover domains are indistinguishable from real queried domains from the proxy’s perspective (i.e., they follow the same effective distribution once padding removes obvious size-based cues), then the victim’s real domain is exchangeable with the $B_{\text{eff}}k$ cover domains in S' . In this case, no adversary can identify the victim domain from S' with probability better than random guessing:

$$\Pr[\text{proxy identifies victim domain}] \leq \frac{1}{1 + B_{\text{eff}}k} \leq \frac{1}{1 + B_{\min}k}.$$

This bound captures the core privacy–performance knob: increasing $B_{\min}$ (mixing more queries per commit) and/or increasing k (adding more covers per query) decreases the proxy’s best-possible success probability even under highly adversarial scheduling.

Membership testing and false positives. A proxy may instead try to test whether the victim queried a particular target domain d^* . In the worst-case cache-delta model, the proxy can check whether $d^* \in S'$. However, cover inserts intentionally induce false positives: even if the victim did not query d^* , covers may include d^* . If covers are sampled independently from distribution D , then the probability that d^* appears at least once among the $B_{\text{eff}}k$ covers is:

$$\Pr[d^* \text{ appears in covers}] = 1 - (1 - D(d^*))^{B_{\text{eff}}k}.$$

Thus, increasing $B_{\text{eff}}k$ increases the rate at which any specific domain appears as a cover, directly reducing the proxy’s confidence in membership testing.

Collisions and effective anonymity. Domain collisions among covers reduce $|S'|$ and therefore weaken the $1/(1 + B_{\text{eff}}k)$ bound. In practice, collisions remain rare when the cover-domain set is large and $B_{\text{eff}}k$ stays moderate, but the system can also sample covers without replacement within a batch or enforce deduplication inside the enclave when constructing S' . Either approach increases the number of distinct candidates and strengthens the anonymity set.

Target misbehavior. Our set-difference analysis assumes that the target samples cover domains from a distribution D that approximates the system’s empirical query distribution Q (e.g., with bounded likelihood ratios $Q(x)/D(x)$). This ensures that covers are indistinguishable from real inserts from the proxy’s perspective. If the target instead biases covers toward a narrow subset—such as only highly popular domains—then out-of-distribution inserts become more likely to be real, weakening the privacy bound even without explicit proxy–target collusion. Although CODoH does not protect against a malicious target, deployments can mitigate accidental or covert bias by statistically auditing cover selection, for example by estimating divergence from a published distribution D ; the enclave itself could perform such checks against an agreed-upon reference distribution.

If the target provides incorrect freshness metadata—such as inaccurate TTLs or timestamps—this primarily impacts correctness and cache effectiveness. Stale timestamps cause inserts to be rejected, while extreme forward jumps can prematurely age out otherwise valid entries, reducing hit rates and potentially disabling caching. These behaviors do not enable the proxy to learn plaintext queries, but they can reduce the mixing benefit of batching by suppressing inserts and thus shrink the effective anonymity set. We therefore treat correct cover selection and freshness metadata as part of the target’s trusted behavior.

Empirical residual leakage. The bound above treats the proxy as observing one batch in isolation; real workloads issue page-load query bursts and, over days, accumulate correlated observations whose intersection narrows the candidate set. We quantify both effects with a trace-driven simulator and report the resulting operating-region recommendations in §7.8.

5.3 Query Equality and Profiling Resistance (G3)

A malicious proxy may try to infer whether two client requests correspond to the same query by observing cache behavior (e.g., hit/miss patterns, timing differences, or access-pattern leakage). CODoH addresses these channels in two ways. First, CODoH avoids exposing an explicit hit/miss signal by ensuring that requests still trigger resolver activity and by structuring client-visible behavior so that the proxy cannot trivially distinguish cache hits from misses. Second, CODoH uses ORAM to hide cache access patterns from a malicious host: even if the proxy can monitor page faults, memory traffic, or other architectural signals outside the enclave, ORAM makes the access pattern computationally indistinguishable across different logical queries. Together, these mechanisms prevent the proxy from turning caching into a query-equality oracle.

6 Implementation

We implement CODoH as a set of plugins for CoreDNS [5], an extensible DNS server written in Go. CoreDNS’s plugin architecture lets each protocol variant—plain DoH, ODoH, and CODoH—run as a distinct plugin configuration, enabling controlled comparisons under identical server infrastructure. The enclave component runs as a standalone binary built with EGo [48], a framework for developing Intel SGX enclaves in Go. On the client side, we extend Cloudflare’s `odoh-client-go` library [15] with CODoH support.

6.1 System Architecture

CODoH deploys as three processes: a *proxy*, an *enclave*, and a *target*. The proxy and target are CoreDNS plugins (`codohproxy` and `codohtarget`); the enclave is a standalone EGo binary co-located with the proxy.

The proxy communicates with the enclave over a Unix domain socket using a compact binary IPC protocol with a narrow interface: forwarding encrypted client queries, delivering signed cache-insert bundles, and retrieving the enclave’s public key.

The target operates as a standard ODoH resolver. When it receives a query relayed through a CODoH proxy (indicated by the presence of the enclave’s public key in the request), it additionally resolves k cover domains, assembles the real and cover responses into a signed bundle, encrypts it to the enclave’s public key, and POSTs the resulting blob back to the proxy’s cache-insert endpoint. This asynchronous path runs after the ODoH response has already been sent to the client, so it does not add latency to the critical path.

6.2 Cryptographic Instantiation

We instantiate the HPKE suite as DHKEM(X25519, HKDF-SHA256) with HKDF-SHA256 and AES-128-GCM, matching the mandatory-to-implement ciphersuite in RFC 9180 [6]. Cache-insert bundles are signed with Ed25519. The target signs the SHA-256 digest of the plaintext bundle before HPKE encryption, and the enclave verifies this signature after decryption, ensuring that the proxy cannot forge or tamper with cache contents.

The per-query response keys (k_e and k_t) are each 16-byte values derived via the HPKE Export interface. The enclave encrypts cached responses under k_e using AES-128-GCM with a random 12-byte nonce. On a cache miss, the enclave instead returns a dummy payload of identical size (284 bytes: 12-byte nonce, 256-byte random fill, and 16-byte authentication tag), which is indistinguishable from a real encrypted response after bucket padding.

6.3 ORAM Cache

We implement the oblivious cache using Path ORAM [47] with a bucket size of 5 and constant-time mode enabled to prevent timing side channels on individual accesses. Each cache entry is serialized into a fixed-size ORAM block containing the canonical query, the DNS response, an insertion timestamp, and a TTL.

TTL expiry is lazy: on each lookup, the enclave checks whether the entry’s insertion time plus its TTL exceeds the current logical time t_{latest} , and treats expired entries as misses. Batch insertions use the ORAM library’s batch-write interface, which amortizes path reads and evictions across multiple writes within a single commit.

6.4 Cover Responses

For each query, the target samples k cover domains and resolves them alongside the real query. Cover domains are drawn from the Cisco Umbrella Top-1M list [13] using a two-tier mixture: with probability p , a domain is sampled from the popular top- N tier, and with probability $1 - p$, from the remaining long-tail tier. Cover domains are resolved in parallel over UDP. If a cover response returns NXDOMAIN, the enclave applies RFC 2308 [2] negative caching rules and derives the TTL from the zone’s SOA record, capped by an enclave-defined upper bound to prevent excessively long lifetimes.

The real response and k cover responses are serialized into a single bundle, encrypted to pk_E , and signed with the target’s Ed25519 key. The enclave receives and processes the bundle atomically, so individual entries are never exposed to the proxy.

6.5 Padding and Wire Format

We pad all queries and responses to fixed-size buckets to prevent length-based distinguishers. Each padded message consists of a 2-byte little-endian length prefix, the original data, and random bytes up to the bucket boundary. Based on our analysis of Cisco Umbrella response sizes (Appendix A), we pad enclave responses—both cache hits and dummy payloads—to a 2 KB bucket, which accommodates the largest HPKE-encrypted response observed among the top-10k domains (1616 bytes). For queries, both Q_T and Q_E are padded to a single 256-byte bucket.

6.6 SGX Attestation and Key Provisioning

On startup, the enclave generates a fresh HPKE key pair and obtains a remote attestation quote from the SGX hardware that binds pk_E to the enclave’s measurement (MRENCLAVE) via the quote’s user-data field. The proxy caches this public key and serves it to clients at a well-known endpoint.

Before the target sends cache-insert bundles, it must verify the enclave’s identity. The target fetches the attestation quote, verifies it against Intel’s attestation infrastructure, checks that the MRSIGNER matches the expected enclave signer, and extracts pk_E from the user-data field. Once verified, the target provisions its Ed25519 signing public key to the enclave, completing the mutual key exchange needed for authenticated cache inserts.

6.7 Client

The client extends the Cloudflare `odoh-client-go` library [15] with CODoH support. For each query, the client constructs two independent ciphertexts: Q_T via the standard ODoH encryption path, and Q_E by HPKE-encrypting the canonical query string (e.g., "example.com") to the enclave’s public key. The client sends Q_T as the HTTP POST body and Q_E as a Base64-encoded X-CODoH-Query header, both padded to 256 bytes.

The proxy returns the enclave and target responses as tagged chunks that identify their source, and the client accepts the first valid response. If the target’s response arrives first, the client still attempts to decrypt the enclave’s chunk to detect key rotation (e.g., after an enclave restart).³ Specifically, the proxy signals rotation

³In this case, if both responses validate but differ, the client treats it as an error.

Table 2: Benchmark configurations: protocol baselines (1–2), a strawman caching proxy (3), and CODoH with and without SGX (4–5).

#	Name	Procs	Cache	Covers	Batching	SGX
C1	DoH	1	—	—	—	—
C2	ODoH	2	—	—	—	—
C3	Cached ODoH	2	LRU	✗	✗	✓
C4	CODoH-noTEE	3	ORAM	✓	✓	✗
C5	CODoH	3	ORAM	✓	✓	✓

via the `X-CODoH-Key-Rotated` header, which alerts the client to refresh `pkE` before issuing its next query.

7 Evaluation

We evaluate CODoH’s end-to-end latency overhead relative to standard DoH and ODoH. Our experiments answer three questions: (1) How does even naïve caching affect latency, and how much further overhead do CODoH’s defense mechanisms (covers, ORAM, batching) impose? (2) What is the performance of CODoH relative to standard ODoH? (3) How much of CODoH’s overhead is attributable to the TEE runtime versus the protocol-level defenses?

7.1 Experimental Setup

Testbed. We deploy the system across three Azure VMs in separate regions to capture realistic wide-area network latencies: the *client* in US North Central, the *proxy* in US East, and the *target* in US Central. The proxy runs on a Standard DC4s v2 (4 vCPUs, 16 GiB, 112 MiB EPC) with Intel SGX support, the target on a Standard D2s v3 (2 vCPUs, 8 GiB), and the client on a Standard D2alds v6 (2 vCPUs, 4 GiB). All VMs run Ubuntu 24.04 LTS with Go 1.25. The enclave is built and signed with EGo 1.8.1.

Upstream resolver. The target VM runs a local Unbound [31] recursive resolver on port 5353, configured with default caching. Between each workload, we flush Unbound’s cache (using the command `unbound-control flush_zone .`) to ensure independent measurements across workloads.

Domain lists. We draw query domains from the Cisco Umbrella Top-1M popularity list. We pre-filter the list to retain only domains that resolve successfully (return at least one A record), producing a 10,000-domain list for cold and cover workloads and a 1,000-domain subset for Zipf workloads.

7.2 Configurations

We compare five configurations, summarized in Table 2. Each successive CODoH configuration adds protocol complexity, enabling us to attribute overhead to specific defense mechanisms.

Config 1: DoH. A single CoreDNS instance serves DNS-over-HTTPS, forwarding queries directly to the upstream resolver. This configuration provides the best-case latency without any oblivious encryption and serves as a performance ceiling.

Config 2: ODoH. A standard Oblivious DoH deployment with two CoreDNS instances: an ODoH proxy and an ODoH target. This is the privacy baseline against which we measure CODoH’s overhead.

Config 3: Cached ODoH. A strawman caching proxy that adds an SGX-hosted LRU cache to the ODoH proxy without any of CODoH’s defense mechanisms: no cover responses, no ORAM, and no batched insertions. The enclave runs as a combined proxy-cache process: it terminates the client’s HTTPS connection, performs HPKE decryption to check the cache, and forwards the ODoH-encrypted query to the target. This two-process architecture (proxy and target) reflects the simplest viable caching design—before the defenses described in §4 necessitate the three-process split used in Configs 4 and 5. Comparing Configs 2 and 3 isolates the benefit of adding a cache to ODoH, while comparing Configs 3 and 5 shows the cost of CODoH’s full defense stack over naïve caching.

Config 4: CODoH-noTEE. The full CODoH protocol with all defense mechanisms—ORAM cache ($N=1024$), cover responses ($k=3$), batched insertions ($B=10$ —chosen to limit queuing delay given the sequential single-client workload—commit probability $p=0.1$), and bucketed response padding—but with the cache process running as a plain userspace process rather than inside an SGX enclave. The cache process and proxy communicate over a Unix domain socket; the proxy and target communicate over a remote TLS connection. Comparing Configs 4 and 5 cleanly isolates the performance cost of SGX hardware, since both share identical defense mechanisms and code paths.

Config 5: CODoH. The complete CODoH system: identical to Config 4 but with the cache process running inside an Intel SGX enclave via EGo. On startup, the enclave produces a remote attestation quote binding its HPKE public key; the target verifies this quote before provisioning its Ed25519 signing key. This is the production configuration as described in §4 and §6. Comparing Configs 2 and 5 gives the total overhead of CODoH over standard ODoH.

7.3 Workloads

For each configuration, we measure latency under three workloads that span the cache-performance spectrum:

Cold. 10,000 unique domains queried in rank order from the pre-filtered Umbrella top-10k list. Every query is a cache miss, requiring full upstream resolution. This workload represents the worst case for CODoH, as the cache provides no benefit and all defense overhead (cover resolution, ORAM writes, encryption) is fully exposed.

Zipf. 10,000 queries drawn from a Zipf distribution ($s=1.0$) over 1,000 domains. This models realistic DNS traffic, where a small number of popular domains account for most queries [28]. Under this workload, the CODoH cache absorbs repeated queries, and the hit rate determines how much latency is saved relative to ODoH.

Warm. 10,000 queries for a single domain (`google.com`). After the first query populates the cache, all subsequent queries are cache hits, measuring the pure protocol overhead of CODoH with no upstream resolution on the critical path.

7.4 Measurement Methodology

Client tool. We use a Go benchmarking client, forked from Cloudflare’s `odoh-client-go` [15] with added CODoH support, that issues sequential HTTPS requests and records per-query wall-clock latency, including client-side encryption and response decryption where applicable. TLS sessions are reused across queries within each workload. The proxy maintains a persistent HTTPS connection pool to the target, so TLS sessions between proxy and target are reused across queries. For each workload, the client writes a CSV of individual query latencies and a JSON summary with percentile statistics (p50, p95, p99), mean latency, and throughput.

Isolation protocol. Between configurations, the orchestrator terminates all server processes and performs a clean restart. Between workloads within a configuration, the orchestrator flushes the upstream Unbound cache to prevent cross-workload contamination. Workloads run in order: cold, Zipf, warm. The enclave’s internal cache is not flushed between workloads, so the warm workload benefits from entries inserted during prior workloads.

Sensitivity analysis. To evaluate the impact of key defense parameters, we re-run Config 5 with varied ORAM capacity ($N \in \{256, 512, 1024, 2048\}$) and cover count ($k \in \{1, 3, 5\}$), holding all other parameters at their default values. Full results across all three workloads appear in Appendix B.

7.5 Results

Table 3 summarizes the end-to-end latency for all configurations under each workload. Speedup factors (in parentheses) are relative to ODoH (Config 2). Figure 5 (Appendix C) shows the resulting per-query CDFs.

Warm workload. On the warm workload (all cache hits), CODoH achieves $0.37\times$ ODoH latency at the median (17.96 vs. 47.97 ms)—a $2.67\times$ speedup—and maintains $0.43\times$ ODoH at p99 (21.77 vs. 51.13 ms). Cache hits are served locally by the enclave, eliminating the proxy-to-target round trip that dominates ODoH. Relative to DoH (8.79 ms), CODoH adds 9.17 ms—mostly the different network path (client to proxy in US East, rather than directly to the target in US Central) plus protocol overhead (HPKE, ORAM, IPC ≈ 0.29 ms)—while providing oblivious routing that DoH lacks entirely.

Zipf workload. Under Zipf traffic, CODoH maintains $0.46\times$ ODoH at the median (22.34 vs. 48.59 ms). The cache absorbs repeated queries for popular domains, and the benefit grows at the tails: CODoH’s p99 is 89.90 ms ($0.47\times$ ODoH) compared to 193.18 ms for ODoH, where every query—including repeats—incurs full upstream resolution. Compared to Cached ODoH, CODoH has a slightly higher median (22.34 vs. 17.31 ms): batched insertion lowers the steady-state hit rate from 86.5% under immediate-insert (C3) to 62.8% under CODoH, so a larger fraction of Zipf queries traverse the upstream path. This trade-off recovers in the tail—CODoH’s p99 (89.90 ms) is substantially below C3’s (173.41 ms)—because at p99 both configurations are sampling cache misses, and CODoH’s parallel fan-out overlaps upstream resolution with proxy work whereas C3 pays it sequentially. C3’s higher hit rate improves the median but does not mask the heavy-tailed upstream variance of its remaining misses.

Cold workload. Under the cold workload, upstream recursive resolution dominates and all configurations converge: medians range from 67.60–110.01 ms, p95 from 355.67–400.68 ms, and p99 from 785.57–868.35 ms. CODoH achieves $0.90\times$ ODoH at the median and $0.90\text{--}0.92\times$ at the tails, confirming that its defense mechanisms impose no meaningful cost on cache-miss queries.

C3 vs. CODoH: streaming fan-out. At first glance the warm and cold medians are counterintuitive: CODoH (C5) and CODoH-noTEE (C4)—which add ORAM, cover insertions, batching, and padding on top of Cached ODoH (C3)—are *faster* at p50 (warm: 17.96/17.38 vs. 21.54 ms; cold: 97.20/99.85 vs. 110.01 ms). The advantage is architectural rather than cryptographic. CODoH’s proxy dispatches the enclave query (local IPC, <1 ms) and the target query (cross-region HTTPS, ~ 25 ms one-way) in parallel and streams tagged response chunks back to the client over HTTP chunked transfer. On a cache hit the enclave chunk is flushed to the client immediately, before the target reply lands; on a miss the enclave emits an indistinguishable padded dummy with the same wire format, so the proxy begins transmitting the response body *during* the upstream resolution rather than after it. C3, by contrast, consults its cache sequentially and—on a miss—then issues a standard request/response exchange to the target, so its hit latency is bounded below by a full proxy round-trip and its miss latency by a full upstream round-trip without overlap. This explains the two anomalies in Table 3: C4/C5 beat C3 on warm because streaming fan-out removes a sequential round-trip on hits, and C4/C5 beat ODoH (C2) on cold because parallel dispatch overlaps proxy compute with upstream resolution.

SGX overhead. Comparing CODoH-noTEE (C4) and CODoH (C5) isolates the cost of TEE execution. On hit-path latency (the obliviousness-relevant regime), SGX overhead is small: hit-only p50 is 0.58 ms higher on warm (17.96 vs. 17.38 ms) and 4.5 ms higher on Zipf (21.94 vs. 17.47 ms), and hit-only p99 differs by 1.6 ms on Zipf (24.65 vs. 23.04 ms) and is statistically indistinguishable on warm (21.72 vs. 22.79 ms). The wider overall p99 gap on Zipf in Table 3 (89.90 vs. 71.50 ms) is dominated not by SGX cost on hits but by miss-path variance: C4 and C5 take identical paths on misses, but C5 exhibited a 2-percentage-point higher miss rate in our run (37.2% vs. 35.0%), so its overall p99 sits deeper in the long miss tail. On cold, the two configurations are within noise (97.20 vs. 99.85 ms p50, 785.57 vs. 789.40 ms p99), as upstream variance dwarfs any TEE cost. Overall, SGX adds modest overhead at the median (<1 ms on cache-hit workloads) that is dominated by network and resolution variance in practice.

7.6 Microbenchmarks

To explain the macrobenchmark results, we measure the latency of individual operations on the critical path of a CODoH query. All microbenchmarks run on the proxy VM and report the median of 10,000 iterations.

Cryptographic operations. Table 4 reports per-operation latency for the cryptographic primitives used in CODoH. HPKE dominates the per-query crypto cost: encrypting a query (KEM setup + AEAD seal) takes $68\ \mu\text{s}$ under SGX, while the symmetric AES-GCM operations are sub-microsecond outside the enclave. Under SGX, AES-GCM Seal rises to $3.8\ \mu\text{s}$ ($4.2\times$); this is on a sub-microsecond

Table 3: End-to-end latency (ms). C3 = Cached ODoH, C4 = CODoH-noTEE, C5 = CODoH; speedups relative to ODoH (10,000 queries, first 1,000 dropped as warmup).

Config	Cold			Zipf			Warm		
	p50	p95	p99	p50	p95	p99	p50	p95	p99
DoH	67.60	355.67	804.98	14.30	42.53	129.29	8.79	9.07	10.17
ODoH	108.49	399.09	868.35	48.59	95.06	193.18	47.97	48.86	51.13
C3	110.01 (1.01×)	400.68 (1.00×)	842.49 (0.97×)	17.31 (0.36×)	79.45 (0.84×)	173.41 (0.90×)	21.54 (0.45×)	22.05 (0.45×)	22.51 (0.44×)
C4	99.85 (0.92×)	374.42 (0.94×)	789.40 (0.91×)	21.68 (0.45×)	48.12 (0.51×)	71.50 (0.37×)	17.38 (0.36×)	22.16 (0.45×)	22.96 (0.45×)
C5	97.20 (0.90×)	368.47 (0.92×)	785.57 (0.90×)	22.34 (0.46×)	58.97 (0.62×)	89.90 (0.47×)	17.96 (0.37×)	20.80 (0.43×)	21.77 (0.43×)

Table 4: Per-operation cryptographic latency (μ s). Indented rows decompose AES-GCM Seal’s SGX cost.

Operation	Plain (μ s)	SGX (μ s)	Overhead
HPKE Encrypt (query)	57.0	67.7	1.19×
HPKE Decrypt (query)	36.8	41.8	1.14×
AES-GCM Seal (resp.)	0.90	3.82	4.24×
<i>AES-GCM Seal (primitive)</i>	0.14	0.15	1.07×
<i>Nonce gen (crypto/rand)</i>	0.04	2.81	70×
AES-GCM Open (resp.)	0.59	0.77	1.31×
Ed25519 Sign	21.5	21.0	0.98×
Ed25519 Verify	49.3	50.6	1.03×

Table 5: Cache read (Get) latency vs. capacity N (μ s). LRU is $O(1)$; ORAM is $O(\log N)$ per access.

N	LRU	ORAM (Plain)	ORAM (SGX)	Overhead
256	0.11	108	223	2.1×
512	0.11	103	222	2.2×
1024	0.12	99	228	2.3×
2048	0.13	107	377	3.5×

base and the increase is dominated by per-call `crypto/rand` nonce generation rather than the AEAD primitive itself.⁴ Ed25519 signing and verification show near-zero SGX overhead ($\leq 1.03\times$), as they are CPU-bound, deterministic, and do not pressure enclave memory.

Cache access latency. Table 5 compares the per-access read (Get) latency of the ORAM cache against a plain LRU cache across capacities $N \in \{256, 512, 1024, 2048\}$. LRU access is sub-microsecond and constant regardless of N , while ORAM access ranges from 99–108 μ s outside SGX. Inside the enclave, ORAM latency roughly doubles for $N \leq 1024$ but jumps to 3.5× at $N=2048$, suggesting EPC paging pressure at that capacity. For ORAM, Put latencies are within 10% of Get; LRU Put is also sub-microsecond, though up to 33% higher than Get at larger N .

Oblivious primitive: PathORAM vs. Linear scan. The cache must hide its access patterns from the host (§4.7); both PathORAM and a full linear scan—which touches every entry on each access—achieve this. Table 6 compares them under SGX at block size 4096 B. Linear scan completes in under 4 μ s, while PathORAM ranges from

⁴The indented rows in Table 4 show this: the AEAD primitive is SGX-neutral, while `crypto/rand` is the SGX-sensitive op. We report the full Seal cost in the parent row since every cache write pays it.

Table 6: Oblivious-primitive SGX latency: PathORAM ($O(\log N)$) vs. linear scan ($O(N)$). Block size 4096 B; medians of 100 iterations across 3 trials.

N	LinearScan (μ s)	PathORAM (μ s)	PathORAM/LinearScan
256	0.107	219.795	2,054×
1024	0.270	309.170	1,145×
4096	0.926	508.627	549×
16384	3.770	1014.634	269×

Table 7: IPC round-trip latency (μ s).

IPC path	Plain (μ s)	SGX (μ s)	Overhead
Health check (baseline)	9.0	121	13.4×
Cache-hit lookup	86	292	3.4×

220 μ s to over 1 ms—269× to 2,054× slower—because its tree exceeds our 112 MiB EPC and incurs paging. On hardware with a multi-gigabyte EPC, the comparison flips; Appendix D confirms this in plain mode, where the crossover lies between $N=524,288$ and $N=786,432$.

IPC round-trip. Table 7 reports the round-trip latency of the Unix domain socket IPC channel between the proxy and enclave processes. We measure two paths: a health check (minimal processing, isolating serialization and socket overhead) and a cache-hit lookup (HPKE decryption, ORAM read, response encryption). The warm-workload macrobenchmark median for CODoH is 17.96 ms, so the SGX cache-hit lookup of 292 μ s (≈ 0.29 ms) contributes only $\sim 1.6\%$ of end-to-end latency; the remaining ~ 17.67 ms is the client-to-proxy round trip (network RTT, TLS, and HTTP processing). This confirms that CODoH’s defense mechanisms—ORAM, HPKE, padding—are not the bottleneck; network latency dominates. The 121 μ s health-check baseline isolates pure IPC overhead, showing that the cryptographic and ORAM work within a cache-hit adds ~ 171 μ s on top.

Padding overhead. To quantify the bandwidth cost of bucketed padding (§4), we resolve the Umbrella top-10k domains and measure the wire size of each query and response before and after padding to a single 256-byte (query) or 2048-byte (response) bucket (Table 8). Query overhead is moderate: the mean Q_E ciphertext is 108 B, yielding 57.5% waste in a 256 B bucket. Response overhead is higher: the mean encrypted response is 274 B, yielding 86.6% waste

Table 8: Mean padding overhead per message (Umbrella top-10k, single-bucket scheme). Waste is the mean unused bytes per padded message.

Message	Bucket (B)	Size (B)	Waste (B)	Overhead
Q_E (to enclave)	256	108	147	57.5%
Q_T (to target)	256	95	160	62.6%
Response (to client)	2048	274	1773	86.6%

in a 2048 B bucket. This is the cost of maximum privacy: fixed-size buckets reveal nothing about the query content.

Resolver amplification. Cover traffic imposes a fixed $1+k$ upstream fan-out: every client query causes the target to resolve one real and k cover domains against the recursive resolver, with all $1+k$ responses bundled into the cache-insert object $\hat{R}_E$ delivered to the enclave (§4). At the default $k=3$, cache misses at the target trigger four authoritative lookups upstream, a $4\times$ multiplier relative to unprotected DoH or baseline ODoH. The dominant on-wire cost is the bundle returned to the proxy: with a 2048 B response bucket (Table 8), each cover adds 2048 B, expanding the bundle from 2048 B under ODoH to $2048(1+k)=8192$ B at $k=3$ on the target-to-proxy return path. The target-to-upstream leg adds a further k unpadded DNS queries (~ 80 B each) and their responses, small relative to the bundle. This overhead is inherent to the cover-based defense: k trades resolver load and bandwidth against the per-batch identification bound of §5.2. §7.8 shows that reducing the cover count to $k=1$ allows a multi-day intersection attack to succeed against $\sim 37\%$ of users (11 of 30) within seven days, so the default $k=3$ is the conservative point on the privacy/bandwidth curve.

7.7 Page Load Benchmarks

We evaluate CODoH’s effect on real-world page load times using Playwright [43] to drive browsers through dnscrypt-proxy [18] on the same Azure testbed as the macrobenchmarks. We sample 100 sites from the Umbrella top-2,000 successfully resolving domains, load each five times in randomized order, and report the trimmed mean (discarding min and max). Figure 2 shows CODoH and ODoH performing identically on simple sites; for sites with many distinct DNS lookups, CODoH’s faster DNS resolution yields shorter total page-load times (e.g., *conviva.com*: 3.0 s $\rightarrow$ 2.5 s; *opera.com*: 1.3 s $\rightarrow$ 1.18 s).

7.8 Leakage under Realistic Workloads

The bound in §5.2 considers a single batch in isolation and assumes the victim’s query is exchangeable with the covers paired with it inside one commit. Real client behavior breaks that idealization in two ways: a single page load issues many related DNS queries that may *share or span batches*, and a user revisits the same sites day after day, so attacker observations *intersect across batches* and narrow the candidate set even when no single batch is identifying. An additional concern is whether the resolver’s cover sampler tracks the empirical query distribution at all—if it does not, the indistinguishability premise itself fails.

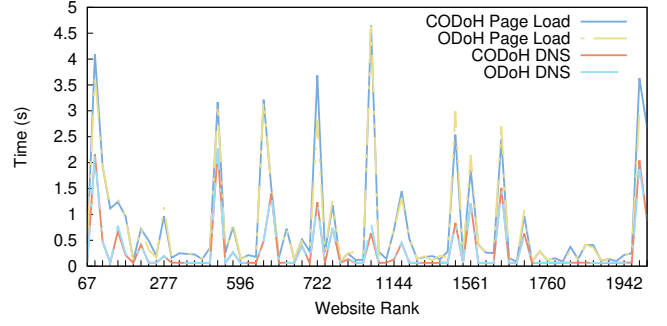


Figure 2: Comparison of average page-load and DNS-fetch times between CODoH and ODoH.

We address these issues with a trace-driven simulator that replays Playwright DNS captures of CrUX [21] top-10k websites⁵ through a faithful model of CODoH’s batching, cover sampling, cache, and commit logic, and scores the proxy’s residual observation against the strong attacker of §5.2. We use it to answer four questions: (i) How well does a single batch protect popular versus tail queries? (ii) Does the union of batches occupied by one page load leak the page? (iii) How quickly does a stable daily repertoire become identifiable from intersected observations? (iv) How much do (i)–(iii) depend on the cover distribution matching the empirical query distribution? The simulator design, parameter ranges, and full quantitative sweeps are in Appendix E and F; here we report qualitative findings and safe values for $B_{\min}$ and $T_{\max}$.

Single batch: popular vs. tail queries. We vary the popularity of the victim’s query under our Zipf model and ask whether one batch alone identifies the page. At small batch sizes ($B \leq 20$) the attacker narrows the victim to a short list containing the true page in 8–14% of batches; at $B \geq 50$ this collapses to below 1%, and at $B \geq 100$ to zero. Popular queries are slightly more exposed than tail queries (by a few percentage points), but only in the small- B regime; the gap vanishes at $B \geq 50$.

Within a page load. A single page load typically issues several DNS queries, which can span multiple batches. Combining everything the attacker observes across those batches amplifies leakage at small batch sizes, but the $B \geq 50$ threshold from above still protects a full page load. The maximum batch hold time $T_{\max}$ is effectively a fallback rather than a tuning knob—under realistic background traffic the batch fills by size well before $T_{\max}$ elapses, so $T_{\max}$ only matters at very low traffic.

Returning user: cross-day intersection. We simulate a user who visits the same 10 sites every day for a week and let the attacker intersect its observations across days to narrow down the user’s habitual repertoire. Across every swept batch size, the candidate set never collapses to just those 10 sites: it narrows substantially but always contains real sites mixed with decoys. Cover count k

⁵We use CrUX rather than Cisco Umbrella here because CrUX is a list of websites a user would actually visit in a browser, where each visit triggers follow-up DNS queries for embedded resources—the precise workload an empirical leakage study needs. Cisco Umbrella ranks popular domains overall (including many CDN, ad, and back-office origins never directly visited) and does not capture this page-and-followups structure.

is the decisive parameter here: $k = 1$ is the only setting where a meaningful fraction of users are fingerprinted, justifying $k = 3$ as our default.

Cover-distribution drift. The three scenarios above assume the resolver’s cover distribution matches the empirical query distribution ($D \approx Q$). Replacing the matched sampler with an unweighted (uniform over the full list) or stale (uniform over a frozen, outdated subset) draw breaks page-load privacy at every B we tested, so the §5.2 indistinguishability claim and the operator-tier table below presuppose $D \approx Q$ as an empirical deployment requirement; see Appendix G for the full cross-scenario sweep.

Operator-tier recommendation. We map findings onto three privacy tiers, each setting a minimum candidate-set size for the cross-day attacker and the matching B : *Strict* (≥ 200 candidates, $B \geq 256$), *Moderate* (≥ 100 , $B \geq 50$), *Permissive* (≥ 50 , $B \geq 20$). $T_{\max}$ should be set as a fallback safety bound (we use 300 s); short-session privacy (within a single batch or one page load) alone—without cross-day intersection—is already satisfied at $B \geq 50$, so operators running predominantly short-lived sessions can adopt the looser $B \geq 50$ setting. Appendix F gives the full per-parameter decomposition and sensitivity sweeps.

8 Conclusion

In this paper, we presented CODoH, a cacheable extension to Oblivious DNS-over-HTTPS that restores proxy-side caching to improve performance while preserving ODoH’s separation between client identity and query contents. CODoH builds on prior DNS-privacy work by combining cover traffic and batching with a TEE-protected, ORAM-backed cache to mitigate cache-induced inference without incurring the latency of mix-based designs. To support further research, we will be making our code publicly available.

Acknowledgments

We thank the anonymous reviewers and the revision editor for their helpful feedback. The authors of this work were supported, in part, by NSF grant CNS-2348130. The authors used ChatGPT-5.2 and Opus-4.6 to revise the text in this paper to correct typos, grammatical errors, and awkward phrasing.

References

- [1] 2025. DNS over Tor. <https://developers.cloudflare.com/1.1.1.1/other-ways-to-use-1.1.1.1/dns-over-tor/> Accessed: 2026-01-26.
- [2] Mark P. Andrews. 1998. Negative Caching of DNS Queries (DNS NCACHE). RFC 2308. <https://www.rfc-editor.org/info/rfc2308>
- [3] Apple iCloud Private Relay 2023. About iCloud Private Relay. <https://support.apple.com/en-us/102602>. Accessed: 2026-02-24.
- [4] Sergei Arnaudov, Bohdan Trach, Franz Gregor, Thomas Knauth, Andre Martin, Christian Priebe, Joshua Lind, Divya Muthukumaran, Dan O’Keeffe, Mark L Stillwell, David Goltzsche, Dave Eysers, Rüdiger Kapitza, Peter Pietzuch, and Christof Fetzter. 2016. SCONe: Secure Linux containers with Intel SGX. In *Symposium on Operating Systems Design and Implementation (OSDI)*.
- [5] The CoreDNS Authors. 2025. CoreDNS: DNS and Service Discovery. <https://coredns.io>. Accessed: 2025-07-25.
- [6] Richard Barnes, Karthikeyan Bhargavan, Benjamin Lipp, and Christopher A. Wood. 2022. Hybrid Public Key Encryption. RFC 9180. <https://www.rfc-editor.org/info/rfc9180>
- [7] Andrew Baumann, Marcus Peinado, and Galen Hunt. 2014. Shielding Applications from an Untrusted Cloud with Haven. In *Symposium on Operating Systems Design and Implementation (OSDI)*.

- [8] Pietro Borrello, Andreas Kogler, Martin Schwarzl, Moritz Lipp, Daniel Gruss, and Michael Schwarz. 2022. *ÆPIC Leak: Architecturally Leaking Uninitialized Data from the Microarchitecture*. In *USENIX Security Symposium*.
- [9] Robert Bühren, Hans-Niklas Jacob, Thilo Krachenfels, and Jean-Pierre Seifert. 2021. One Glitch to Rule Them All: Fault Injection Attacks Against AMD’s Secure Encrypted Virtualization. In *ACM Conference on Computer and Communications Security (CCS)*.
- [10] Sergio Castillo-Perez and Joaquin Garcia-Alfaro. 2008. Anonymous Resolution of DNS Queries. In *On the Move to Meaningful Internet Systems (OTM)*.
- [11] David L. Chaum. 1981. Untraceable electronic mail, return addresses, and digital pseudonyms. *Commun. ACM* 24, 2 (feb 1981).
- [12] Benny Chor, Oded Goldreich, Eyal Kushilevitz, and Madhu Sudan. 1995. Private Information Retrieval. In *Symposium on Foundations of Computer Science (FOCS)*.
- [13] Cisco. 2026. Umbrella Popularity List. <http://s3-us-west-1.amazonaws.com/umbrella-static/top-1m.csv.zip>. Accessed: 2026-02-28.
- [14] Cloud Native Computing Foundation 2026. Cloud Native Computing Foundation. <https://cncf.io>. Accessed: 2026-02-24.
- [15] Cloudflare. 2021. odoh-client-go. <https://github.com/cloudflare/odoh-client-go>.
- [16] Cloudflare Oblivious DNS over HTTPS 2025. Oblivious DNS over HTTPS. <https://developers.cloudflare.com/1.1.1.1/encryption/oblivious-dns-over-https/>. Accessed: 2026-02-24.
- [17] Wouter B. de Vries, Roland van Rijswijk-Deij, Pieter-Tjerk de Boer, and Aiko Pras. 2020. Passive Observations of a Large DNS Service: 2.5 Years in the Life of Google. *IEEE Transactions on Network and Service Management* 17, 1 (2020).
- [18] DNSCrypt. 2026. DNSCrypt. <https://dnscrypt.info/>. Accessed: 2026-02-28.
- [19] Hannes Federrath, Karl-Peter Fuchs, Dominik Herrmann, and Christopher Pioceny. 2011. Privacy-Preserving DNS: Analysis of Broadcast, Range Queries and Mix-based Protection Methods. In *European Symposium on Research in Computer Security (ESORICS)*.
- [20] Oded Goldreich. 1987. Towards a Theory of Software Protection and Simulation by Oblivious RAMs. In *ACM Symposium on Theory of Computing (STOC)*.
- [21] Google. 2026. Chrome User Experience Report. <https://developer.chrome.com/docs/crux>. Accessed: 2026-05-13.
- [22] Stephen Herwig, Christina Garman, and Dave Levin. 2020. Achieving Keyless CDNs with Conclaves. In *USENIX Security Symposium*.
- [23] P. Hoffman, ICANN, P. McManus, and Mozilla. 2018. DNS Queries over HTTPS (DoH). RFC 8484. <https://www.ietf.org/rfc/rfc8484.txt>
- [24] Paul E. Hoffman. 2023. DNS Security Extensions (DNSSEC). RFC 9364. <https://www.rfc-editor.org/info/rfc9364>
- [25] Zi Hu, Liang Zhu, John Heidemann, Allison Mankin, Duane Wessels, and Paul E. Hoffman. 2016. Specification for DNS over Transport Layer Security (TLS). RFC 7858. <https://www.rfc-editor.org/info/rfc7858>
- [26] Intel 2014. *Intel Software Guard Extensions Programming Reference*. Intel.
- [27] Jaeyeon Jung, Emil Sit, Hari Balakrishnan, and Robert Morris. 2002. DNS Performance and the Effectiveness of Caching. *IEEE/ACM Transactions on Networking* 10, 5 (2002).
- [28] Jaeyeon Jung, Emil Sit, Hari Balakrishnan, and Robert Morris. 2002. DNS Performance and the Effectiveness of Caching. In *IEEE/ACM Transactions on Networking*, Vol. 10. IEEE, 589–603.
- [29] Eric Kinnear, Patrick McManus, Tommy Pauly, Tanya Verma, and Christopher A. Wood. 2022. Oblivious DNS over HTTPS. RFC 9230. <https://www.rfc-editor.org/info/rfc9230>
- [30] Paul Kocher, Jann Horn, Anders Fogh, Daniel Genkin, Daniel Gruss, Werner Haas, Mike Hamburg, Moritz Lipp, Stefan Mangard, Thomas Prescher, Michael Schwarz, and Yuval Yarom. 2019. Spectre Attacks: Exploiting Speculative Execution. In *IEEE Symposium on Security and Privacy*.
- [31] NLnet Labs. 2026. Unbound. <https://www.nlnetlabs.nl/projects/unbound/about/>. Accessed: 2026-02-28.
- [32] Mengyuan Li. 2022. *Understanding and Exploiting Design Flaws of AMD Secure Encrypted Virtualization*. Ph.D. Dissertation. <https://etd.ohiolink.edu/>.
- [33] Mengyuan Li, Yinqian Zhang, Zhiqiang Lin, and Yan Solihin. 2019. Exploiting Unprotected I/O Operations in AMD’s Secure Encrypted Virtualization. In *USENIX Security Symposium*.
- [34] Mengyuan Li, Yinqian Zhang, Huibo Wang, Kang Li, and Yueqiang Cheng. 2021. CIPHERLEAKS: Breaking Constant-time Cryptography on AMD SEV via the Ciphertext Side Channel. In *USENIX Security Symposium*.
- [35] Moritz Lipp, Michael Schwarz, Daniel Gruss, Thomas Prescher, Werner Haas, Anders Fogh, Jann Horn, Stefan Mangard, Paul Kocher, Daniel Genkin, Yuval Yarom, and Mike Hamburg. 2018. Meltdown: Reading Kernel Memory from User Space. In *USENIX Security Symposium*.
- [36] Frank McKeen, Ilya Alexandrovich, Alex Berenzon, Carlos V Rozas, Hisham Shafi, Vedvyas Shanbhogue, and Uday R Savagaonkar. 2013. Innovative Instructions and Software Model for Isolated Execution. In *Workshop on Hardware and Architectural Support for Security and Privacy (HASP)*.
- [37] Samir Jordan Menon and David J. Wu. 2022. Spiral: Fast, High-Rate Single-Server PIR via FHE Composition. In *IEEE Symposium on Security and Privacy*.
- [38] Paul Mockapetris. 1987. Domain Names - Concepts and Facilities. RFC 1034. <https://www.rfc-editor.org/info/rfc1034>

- [39] Paul Mockapetris. 1987. Domain Names - Implementation and Specification. RFC 1035. <https://www.rfc-editor.org/info/rfc1035>
- [40] Alec Muffett. 2020. DoHoT: making practical use of DNS over HTTPS over Tor. <https://medium.com/@alecmuffett/dohot-making-practical-use-of-dns-over-https-over-tor-ef58d04ca06a> Accessed: 2026-01-26.
- [41] Kit Murdock, David Oswald, Flavio D. Garcia, Jo Van Bulck, Daniel Gruss, and Frank Piessens. 2020. Plundervolt: Software-based Fault Injection Attacks against Intel SGX. In *IEEE Symposium on Security and Privacy*.
- [42] Yoshimichi Nakatsua, Andrew Paverd, and Gene Tsudik. 2019. PDoT: Private DNS-over-TLS with TEE Support. In *Annual Computer Security Applications Conference (ACSAC)*.
- [43] Playwright. 2026. Playwright. <https://playwright.dev/>. Accessed: 2026-02-28.
- [44] Mohammad Amir Salari, Abhinav Kumar, Federico Rinaudi, Reza Tourani, Alessio Sacco, and Flavio Esposito. 2025. Privacy Analysis of Oblivious DNS over HTTPS: A Website Fingerprinting Study. In *IEEE/IFIP International Conference on Dependable Systems and Networks (DSN)*.
- [45] Kyle Schomp, Onkar Bhardwaj, Eymen Kurdoglu, Mashooq Muhaimen, and Ramesh K. Sitaraman. 2020. Akamai DNS: Providing Authoritative Answers to the World's Queries. In *ACM SIGCOMM*.
- [46] Sudheesh Singanamalla, Suphanat Chunhanya, Jonathan Hoyland, Marek Vavruša, Tanya Verma, Peter Wu, Marwan Fayed, Kurtis Heimerl, Nick Sullivan, and Christopher Wood. 2021. Oblivious DNS over HTTPS (ODoH): A Practical Privacy Enhancement to DNS. In *Privacy Enhancing Technologies Symposium (PETS)*.
- [47] Emil Stefanov, Marten van Dijk, Elaine Shi, Christopher Fletcher, Ling Ren, Xiangyao Yu, and Srinivas Devadas. 2013. Path ORAM: An Extremely Simple Oblivious RAM Protocol. In *ACM Conference on Computer and Communications Security (CCS)*.
- [48] Edgeless Systems. 2025. Welcome to EGo. <https://docs.edgeless.systems/ego>.
- [49] Martino Trevisan, Idilio Drago, Paul Schmitt, and Francesco Bronzino. 2023. Measuring the Performance of iCloud Private Relay. In *Passive and Active Measurement Conference (PAM)*.
- [50] Chia-Che Tsai, Donald E. Porter, and Mona Vij. 2017. Graphene-SGX: A Practical Library OS for Unmodified Applications on SGX. In *USENIX Annual Technical Conference (ATC)*.
- [51] Jo Van Bulck, Marina Minkin, Ofir Weisse, Daniel Genkin, Baris Kasikci, Frank Piessens, Mark Silberstein, Thomas F Wenisch, Yuval Yarom, and Raoul Strackx. 2018. Foreshadow: Extracting the Keys to the Intel SGX Kingdom with Transient Out-of-Order Execution. In *USENIX Security Symposium*.
- [52] Jo Van Bulck, Daniel Moghimi, Michael Schwarz, Moritz Lipp, Marina Minkin, Daniel Genkin, Yarom Yuval, Berk Sunar, Daniel Gruss, and Frank Piessens. 2020. LVI: Hijacking Transient Execution through Microarchitectural Load Value Injection. In *IEEE Symposium on Security and Privacy*.
- [53] Stephan Van Schaik, Alex Seto, Thomas Yurek, Adam Batori, Bader AlBassam, Daniel Genkin, Andrew Miller, Eyal Ronen, Yuval Yarom, and Christina Garman. 2024. SoK: SGX.Fail: How Stuff Gets eXposed. In *IEEE Symposium on Security and Privacy*.
- [54] Yunming Xiao, Chenkai Weng, Ruijue Yu, Peizhi Liu, Mattero Varvello, and Aleksandar Kuzmanovic. 2023. Demo: PDNS: A Fully Privacy-Preserving DNS. In *ACM SIGCOMM*.
- [55] Yuanzhong Xu, Weidong Cui, and Marcus Peinado. 2015. Controlled-Channel Attacks: Deterministic Side Channels for Untrusted Operating Systems. In *IEEE Symposium on Security and Privacy*.
- [56] Fangming Zhao, Yoshiaki Hori, and Kouichi Sakurai. 2007. Analysis of Privacy Disclosure in DNS Query. In *International Conference on Multimedia and Ubiquitous Engineering (MUE)*.
- [57] Fangming Zhao, Yoshiaki Hori, and Kouichi Sakurai. 2007. Two-Servers PIR Based DNS Query Scheme with Privacy-Preserving. In *International Conferences on Intelligent Pervasive Computing (IPC)*.

A Parameter Selection

Response padding. To motivate our choice of padding bucket sizes, we query the top 700k domains of the Cisco Umbrella Top-1M for A records and record the size of the response. Typically, this response includes the question, one or more A records, and sometimes a chain of CNAME records. Figure 3 shows the sizes of these responses as compared to possible bucket sizes of 1 KB, 2 KB, and 4 KB. Notably, only 53 of these responses exceed 1 KB, and only 7 exceed 2 KB.

TTLs and batch interval. To estimate how often batching could delay inserts past their TTL, we measured TTLs for A responses

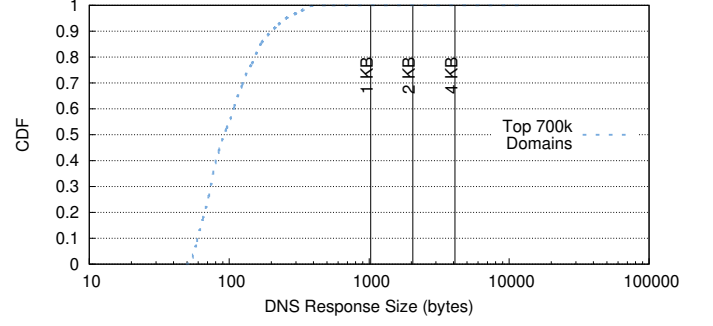


Figure 3: CDF of sizes of responses to A requests for Cisco Umbrella's top 700k domains (log scale).

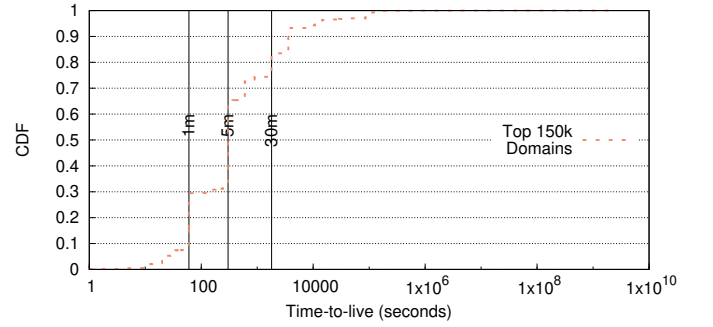


Figure 4: CDF of TTLs for A or CNAME records for Cisco Umbrella's top 150k domains (log scale).

for the top 150k domains in the Cisco Umbrella Top-1M. Because recursive resolvers return decremented TTLs, we collect authoritative TTLs by first querying NS records via a recursive resolver and then querying the corresponding authoritative nameservers for the desired record. Figure 4 plots the resulting TTL CDF, which is highly quantized with common values at 1, 5, and 30 minutes, and at 1–2 hours. These measurements suggest that the commit delay should be well below the smallest common TTL (1 minute); in practice, this means choosing B and $T_{\max}$ so that batches typically commit within tens of seconds at the expected query rate.

B Sensitivity Analysis Details

Methodology. We sweep ORAM capacity $N \in \{256, 512, 1024, 2048\}$ at fixed $k=3$, and cover count $k \in \{1, 3, 5\}$ at fixed $N=1024$; both sweeps share the default cell ($N=1024, k=3$), marked bold in Table 9. Each cell runs cold, Zipf, and warm workloads—sequential top-10 k, Zipf- $s=1.0$ over the resolvable top-1 k bucket (763 domains), and a single fixed domain—against the Cisco Umbrella Top-1M filtered to domains returning a valid A record, with Unbound flushed between workloads.

Findings. The request path itself is parameter-independent: p50 latency for each cache outcome stays within a tight band across all six cells. The metric that moves is the Zipf cache-hit rate, and only along the capacity axis: $N=256$ and $N=512$ have hit rates of 41% and

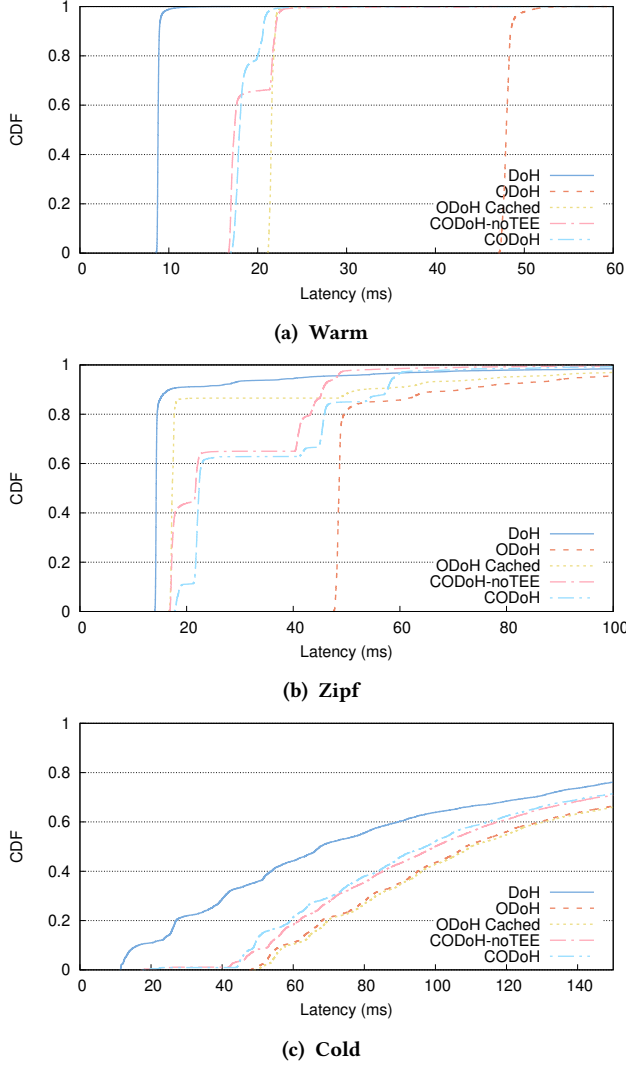


Figure 5: Per-query latency CDFs by configuration. CODoH’s streaming fan-out shifts the body of the warm and cold distributions left of Cached ODoH (C3); on Zipf, C3’s heavy tail (p99 $\approx$ 173 ms) extends past the right edge of the panel.

50%, while $N=1024$ and $N=2048$ reach 62% and 69%. Cover count has no monotonic effect on hit rate (59–62% from $k=1$ to $k=5$), and its only material latency cost concentrates in the step from $k=1$ to $k=3$ (~ 4 ms). We therefore choose $N=1024$ —the smallest capacity past which doubling N yields diminishing returns—and $k=3$, the smallest cover count that Appendix F certifies as adequate against fingerprinting.

C Per-Query Latency Distributions

Figure 5 expands upon Table 3 to show the CDF of request latencies across different cache states. On Zipf, the C3 and CODoH curves cross in the upper tail: C3’s higher hit rate gives it a lower median,

but its miss tail stays heavy, whereas CODoH’s parallel fan-out compresses miss-side variance.

D Oblivious Primitive Comparison without an EPC Ceiling

Table 10 extends the SGX comparison of §7 (Table 6) to plain (non-SGX) execution, where neither primitive incurs EPC paging and the working set is bounded only by host RAM. We sweep N up to 2^{20} at the deployed block size of 4096 B, on a Xeon Gold 5512U workstation with 503 GiB of RAM. Without the EPC ceiling, PathORAM’s per-access cost grows only as $O(\log N)$ and stays under $800 \mu\text{s}$ across the range, while linear scan grows linearly with N . The two primitives cross between $N=524,288$ and $N=786,432$, beyond which PathORAM is the faster oblivious primitive. This crossover lies above any cache size we considered for CODoH and, more importantly, lies in a regime that the EPC-constrained SGX hardware available to us could not reach in the protected setting.

E Leakage Simulator: Methodology

This appendix describes the trace-driven simulator used in §7.8.

Architecture. The simulator replays each victim page’s per-entry DNS trace through a faithful re-implementation of the enclave batching logic. A `BatchBuffer` accepts queries from the victim and from λ_{bg} concurrent background users replaying random trace pages, committing when either the size trigger B or the time trigger $T_{\max}$ fires. Each commit draws k cover queries per real query from a CrUX-top-1M Zipf distribution and then suppresses any insert (real or cover) already present in an LRU cache; the resulting hostname set with known background-user queries removed is the strong attacker’s observation S' . Against S' , the attacker scores each reference page w by the fraction of w ’s hostnames that appear in S' : $s_w = |Q_w \cap S'| / |Q_w|$, where Q_w is the set of hostnames our crawl recorded for w . With threshold $\alpha = 0.5$, the candidate set $C_\alpha = \{w : s_w \geq \alpha\}$ keeps every page whose score meets the bar.

Sweep. Table 11 lists the swept parameter grids. Single-batch and page-load cells fix $k = 3$, $\lambda_{bg} = 100$ and sweep $(B, T_{\max})$; returning-user (cross-day intersection) cells additionally sweep k and λ_{bg} at $(B = 20, T_{\max} = 300 \text{ s})$. All cells use cold-cache initialization.

Trace. We run repeated Playwright crawls of the CrUX top-10k pages over multiple days to collect each page’s DNS lookup set Q_w . After removing pages that failed to load in any run, we retain DNS sets for $\sim 8.9\text{k}$ pages, which form the attacker’s reference set.

F Detailed Leakage Results

Single batch: popular vs. tail queries. Table 12 reports two views of the attacker’s per-batch identification accuracy—the fraction of batches in which the victim is the unique highest-scoring page (top-1) and the fraction in which the victim appears in the top five candidates (top-5)—by batch size B and CrUX popularity band (ranks 1–1000; 1k–5k; 5k–10k). The top-5 fraction drops sharply and uniformly across all three bands: from 8–14% at $B \leq 20$ to under 1% at $B = 50$ and to zero at $B \geq 100$. The stricter top-1 fraction is roughly an order of magnitude lower at every B and reaches zero by $B = 50$ (1.3–2.1% at $B = 5$, 0.1–0.3% at $B = 20$). At small B , popular

Table 9: Sensitivity analysis: latency (ms) under varied ORAM capacity N and cover count k . Default configuration is $N=1024$, $k=3$ (bold). Hit/miss latencies are reported separately; “Rate” is the cache-hit percentage.

Sweep	Workload	Rate (%)	Cache Hits				Cache Misses			
			p50	p95	p99	mean	p50	p95	p99	mean
ORAM capacity sweep ($k=3$ fixed)										
$N=256$	Cold	1.4	21.8	22.8	24.6	20.9	93.1	349.9	756.6	134.0
	Warm	99.3	21.9	25.8	26.7	21.4	47.8	52.3	53.2	47.6
	Zipf	40.9	21.7	25.2	26.0	21.0	51.0	65.6	120.6	53.5
$N=512$	Cold	1.7	22.0	26.6	27.5	21.0	97.3	342.2	763.4	137.3
	Warm	99.7	17.9	22.4	23.1	18.6	47.8	48.6	49.3	47.7
	Zipf	49.6	26.4	27.4	29.2	25.6	50.1	66.5	143.2	54.4
$N=1024$	Cold	3.0	18.2	26.3	26.7	20.0	94.8	337.6	730.3	133.3
	Warm	99.9	21.8	22.8	24.5	21.9	51.8	52.6	52.7	50.4
	Zipf	61.7	22.2	26.3	26.9	22.2	52.2	71.6	138.9	54.5
$N=2048$	Cold	3.4	20.6	22.5	22.9	20.4	99.0	350.2	749.5	139.7
	Warm	99.8	21.6	23.0	24.3	20.3	46.2	47.3	49.1	46.2
	Zipf	68.5	22.4	26.7	27.4	23.0	58.7	83.1	167.2	60.8
Cover count sweep ($N=1024$ fixed)										
$k=1$	Cold	1.7	22.0	23.3	25.0	21.6	98.6	365.7	770.4	140.3
	Warm	100.0	17.9	19.0	21.0	18.0	—	—	—	—
	Zipf	59.0	22.1	23.1	24.9	22.0	52.5	93.7	176.3	58.8
$k=3$	Cold	3.0	18.2	26.3	26.7	20.0	94.8	337.6	730.3	133.3
	Warm	99.9	21.8	22.8	24.5	21.9	51.8	52.6	52.7	50.4
	Zipf	61.7	22.2	26.3	26.9	22.2	52.2	71.6	138.9	54.5
$k=5$	Cold	3.8	18.0	22.4	26.6	19.2	96.3	353.9	741.2	136.7
	Warm	98.5	22.1	23.0	24.6	21.4	52.6	53.3	54.2	51.6
	Zipf	59.1	18.8	22.8	24.1	20.1	49.2	55.3	108.9	52.1

Table 10: Oblivious-primitive plain-mode latency at block size 4096 B, on a Xeon Gold 5512U workstation. Medians of 200 iterations across 5 trials. Linear scan beats PathORAM for $N \leq 524,288$; PathORAM overtakes between $N=524,288$ and $N=786,432$.

N	LinearScan (μ s)	PathORAM (μ s)	Faster
256	0.226	593.978	LinearScan 2,629×
1,024	0.936	583.795	LinearScan 624×
4,096	4.840	596.418	LinearScan 123×
16,384	7.917	646.628	LinearScan 81.7×
65,536	82.887	483.868	LinearScan 5.84×
262,144	339.369	619.989	LinearScan 1.83×
524,288	574.803	671.115	LinearScan 1.17×
786,432	841.748	712.124	PathORAM 1.18×
1,048,576	1,026.705	734.038	PathORAM 1.40×

Table 11: Simulator sweep grids.

Parameter	Swept values
B	{5, 10, 20, 50, 100, 256}
$T_{\max}$ (s)	{30, 100, 300, 600, 1800}
k	{1, 3, 5, 10}
λ_{bg} (users)	{10, 100, 1000}

Table 12: Per-batch attacker accuracy by batch size B and CrUX popularity band: fraction of batches in which the victim is the unique highest-scoring page (top-1) or appears in the top five candidates (top-5). Averaged over $T_{\max} \in \{30, 100, 300, 600, 1800\}$ s.

B	top-1			top-5		
	top-1k	1k–5k	5k–10k	top-1k	1k–5k	5k–10k
5	0.021	0.015	0.013	0.140	0.111	0.094
20	0.003	0.002	0.001	0.132	0.092	0.077
50	0.000	0.000	0.000	0.008	0.005	0.004
100	0.000	0.000	0.000	0.000	0.000	0.000
256	0.000	0.000	0.000	0.000	0.000	0.000

pages are noticeably more identifiable than tail pages—14% vs. 9% at $B = 5$ and 13% vs. 8% at $B = 20$ (top-5)—but the gap closes by $B \geq 50$. The maximum batch hold time $T_{\max}$ is essentially absent from this table: moving $T_{\max}$ across the swept range $\{30, 100, 300, 600, 1800\}$ s leaves these numbers essentially unchanged, so the table averages over it.

Within a page load. A page load typically issues several DNS queries that span multiple batches; an attacker who unions all batches occupied by one page load can be more confident than any single batch alone. Table 13 reports the per-page-load counterpart of Table 12: top-1 and top-5 page-identification accuracy, mean /

Table 13: Per-page-load attacker accuracy and candidate-set size by batch size B . Top-1 and top-5 are the fractions of page loads in which the victim’s page is the unique highest-scoring candidate (top-1) or appears in the top five (top-5) under the unioned attacker observation; $|C|$ is the candidate-set size; $\text{frac} \leq 5$ is the fraction of page loads with $|C| \leq 5$. Averaged over $T_{\max} \in \{30, 100, 300, 600, 1800\}$ s; $k = 3$, $\lambda_{bg} = 100$, $\alpha = 0.5$, matched cover distribution, cold cache; 4500 page loads per cell.

B	top-1	top-5	mean $ C $	p_{10}	p_{90}	$\text{frac} \leq 5$
5	0.085	0.439	195	3	621	0.176
20	0.004	0.217	220	10	671	0.010
50	0.000	0.009	248	22	760	0.000
100	0.000	0.000	280	41	822	0.000
256	0.000	0.000	346	96	883	0.000

p_{10} / p_{90} of the candidate-set size $|C|$, and the fraction of page loads whose candidate set collapses to at most five sites, by batch size B . In the vulnerable regime the union amplifies leakage relative to a single batch—page-load top-5 is 44% at $B = 5$ versus the 9–14% per-batch range of Table 12, and 22% versus 8–13% at $B = 20$ —but once B crosses the $B = 50$ threshold the union adds nothing measurable: page-load top-5 falls to under 1% at $B = 50$ and to zero at $B \geq 100$, matching the per-batch numbers to within sample noise. As in Table 12, $T_{\max}$ is essentially absent from this table—moving it across the swept range $\{30, 100, 300, 600, 1800\}$ s shifts top-5 by at most 1.4 percentage points at every B —so the table averages over it.

Returning user: cross-day intersection. We simulate 30 users, each visiting the same ten habitual pages R every day for a week, drawn from the CrUX top-1k bucket; the attacker intersects its per-day candidate sets and we report the day-7 intersection size $|C_7|$ at the default $T_{\max} = 300$ s, $\alpha = 0.5$, matched cover, cold cache. No median user is fingerprinted within seven days ($|C_7| > |R| = 10$ in every cell of the grid); the dynamic range is in the size of the surviving candidate set. Figure 6 decomposes that size along the three privacy knobs. Batch size B monotonically improves privacy: median $|C_7|$ rises from 35 at $B = 5$ to 190 at $B = 100$ and 486 at $B = 256$. Cover count k has the sharpest threshold: at $k = 1$ the candidate set collapses (median 19, p_{10} 5) and 11 of 30 users cross the $|R| = 10$ threshold—the only cell in the grid where any meaningful fraction of users are fingerprinted; $k = 3$ removes this leakage (no users below threshold), and $k = 10$ pushes the median to 393 at the cost of cover bandwidth. Background traffic helps the defender: sparser λ_{bg} admits fewer non-victim hostnames per day and tightens the intersection (median 46 / 77 / 110 at $\lambda_{bg} \in \{10, 100, 1000\}$).

Background traffic and $T_{\max}$. Figure 7 confirms at the cross-day level the same observation we made in Tables 12 and 13: under realistic background traffic ($\lambda_{bg} = 100$) the batch fills by size before $T_{\max}$ elapses, so within-row variance across $T_{\max} \in \{30, 300, 1800\}$ s is below sample noise at every B in the grid. This holds even when background traffic is sparse: at $\lambda_{bg} = 10$, $B = 20$, median $|C_7|$ varies non-monotonically across $T_{\max} \in \{30, 300, 1800\}$ s (44, 46, 26) with broadly overlapping p_{10}/p_{90} bands, consistent with sampling noise

at our 30-user panel size. Because these runs set $B_{\min} = B$, a time-triggered batch that fails to fill reverts to baseline ODoH rather than committing a small, easily isolated batch, so $T_{\max}$ has no first-order effect on the cross-day intersection at any traffic level we tested—batch size B remains the dominant knob.

What the surviving candidate set looks like. The day-7 intersection size measures how much the attack narrows the search; the *repertoire accuracy* $|R \cap C_7|/|R|$ measures how well the surviving set covers the user’s true ten pages. At default $k = 3$, $\lambda_{bg} = 100$, the attack retains 6–9 of those ten in C_7 , with median repertoire accuracy rising from 0.6 at $B \leq 20$ to 0.9 at $B = 256$ —well above the random-guess baseline of $10/8862 \approx 0.001$ but below the perfect-recovery value of 1.0. The attack identifies a *region* containing the user’s repertoire, not the repertoire itself.

Score-threshold sensitivity. Re-running the per-day S'_d sets at $\alpha \in \{0.3, 0.5, 0.7\}$ shows $\alpha = 0.5$ to be the operative bar. At $\alpha = 0.7$, $|C_7|$ shrinks below threshold for many users but the repertoire accuracy collapses to 0.1–0.2: the surviving candidate set is small but mostly contains the *wrong* pages, giving the attacker false confidence rather than real identification. At $\alpha = 0.3$ the candidate set is too lax to encode meaningful information across B .

Limitations. Our reported numbers are pessimistic upper bounds on attacker advantage. All cells use cold-cache initialization; warm initialization would pre-suppress victim queries already in the cache and reduce leakage further. The cross-day construction holds the user’s ten-page repertoire fixed across all seven days and reshuffles only page order and per-query timings ($\pm 10\%$ lognormal jitter); any real-world variation in the repertoire (added/removed pages, time-of-day mix shifts) would further dilute the intersection signal. Our coverage is genuinely narrow only along the popularity axis: the trace stops at the CrUX top-10k, so a long-tail bucket is left for future work, and the cross-day sweep further restricts to the top-1k bucket.

G Cover-Distribution Sensitivity

Section 7.8 reports privacy under the assumption that the resolver’s cover-sampling distribution D matches the empirical query distribution Q that clients actually request. Here we quantify what happens when D drifts from Q . We compare three covers at the operator-tier batch sizes $B \in \{20, 50, 100\}$, holding $T_{\max} = 300$ s, $k = 3$, $\lambda_{bg} = 100$, $N = 1024$, cold cache, $\alpha = 0.5$. All three covers draw from the CrUX-1M support: *matched* weights origins by the same magnitude-band Zipf profile as Q ; *uniform* draws unweighted from the full 1M; *stale* draws unweighted from a frozen random 50% subset. The per-batch and page-load measurements use 4500 trials per cell across the three popularity buckets; the cross-day measurements use 30 simulated users over 7 days per cell.

Shape vs. support. Across all three scenarios and all three B s, *uniform* and *stale* are within sample noise of each other. Stale shrinks the cover support to a frozen 500 k-domain subset but keeps the same uniform shape; the attacker’s edge under wrong D is therefore driven by the *shape* of D relative to Q , not by the size of the cover universe. Cover-universe expansion alone (e.g., broadening from CrUX-1M to a larger zone) would not buy any additional protection if the sampler did not also re-weight to track Q .

Day-7 candidate-set intersection size — cross-day attack

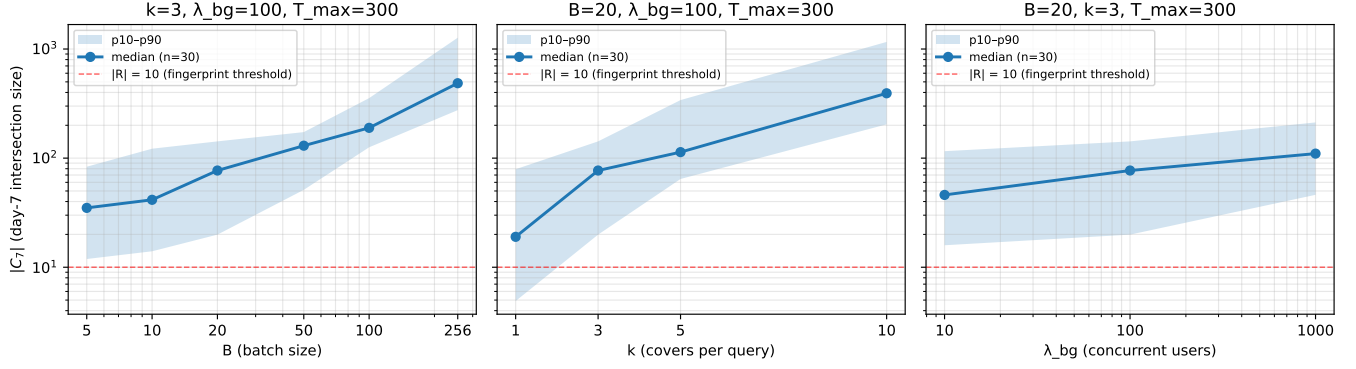


Figure 6: Cross-day intersection: median day-7 candidate-set size $|C_7|$ as a function of each privacy knob, with p10/p90 bands across 30 users. Larger is better; reference set has ~ 8.9 k pages. Left: batch size B at $k = 3, \lambda_{bg} = 100$. Center: cover count k at $B = 20, \lambda_{bg} = 100$. Right: background traffic λ_{bg} at $B = 20, k = 3$. All panels have $T_{\max} = 300$ s, $\alpha = 0.5$, and a cold cache.

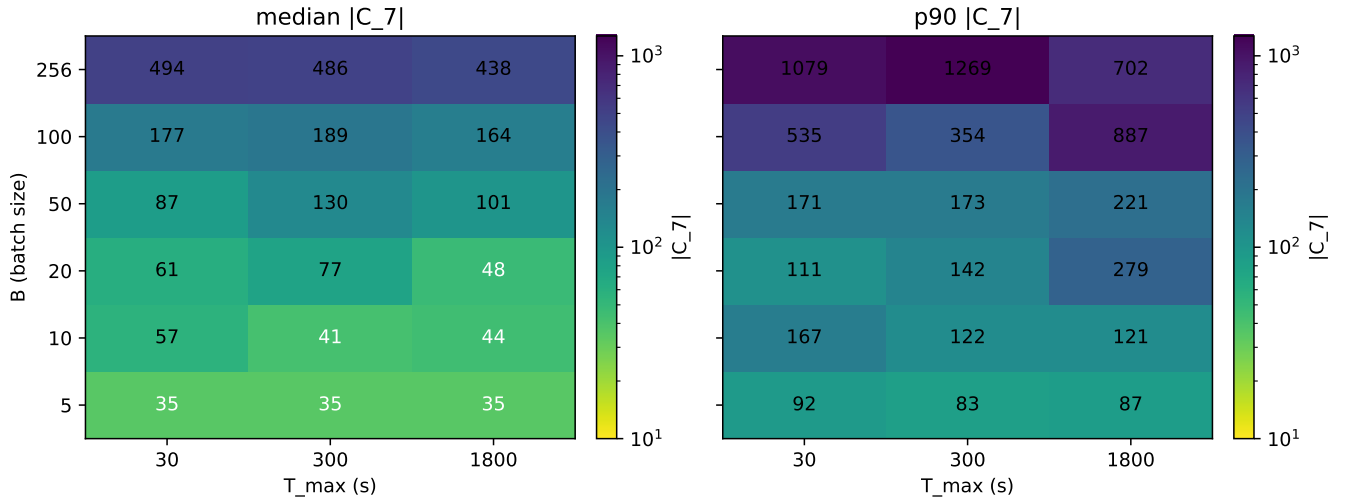
Day-7 intersection-set size — $k=3, \lambda_{bg}=100$ 

Figure 7: Cross-day intersection ($B, T_{\max}$) heatmap: day-7 candidate-set size $|C_7|$ at $k = 3, \lambda_{bg} = 100, \alpha = 0.5$, cold cache, 30 users per cell—median (left) and p90 (right) across users. Batch size dominates; varying $T_{\max}$ within $\{30, 300, 1800\}$ s leaves the candidate-set size unchanged at this background-traffic level. $D \approx Q$ assumed throughout (App. G).

Table 14: Single batch: top-5 page-identification accuracy in the top-1k bucket, by B and cover distribution. $T_{\max} = 300$ s, $k = 3, \lambda_{bg} = 100$, cold cache.

B	matched	uniform	stale
20	0.133	0.455	0.450
50	0.007	0.678	0.673
100	0.000	0.829	0.831

Table 15: Within a page load: top-5 page-identification accuracy by B and cover distribution; same axes as Table 14. Wrong- D accuracy saturates at 0.88 ± 0.01 at every B .

B	matched	uniform	stale
20	0.219	0.889	0.888
50	0.008	0.883	0.879
100	0.000	0.876	0.876

Mechanism. The strong-attacker score function is a likelihood ratio between “this batch’s contents are explained by victim queries”

and “this batch’s contents are explained by background traffic”. Under matched D , that ratio is bounded by the indistinguishability

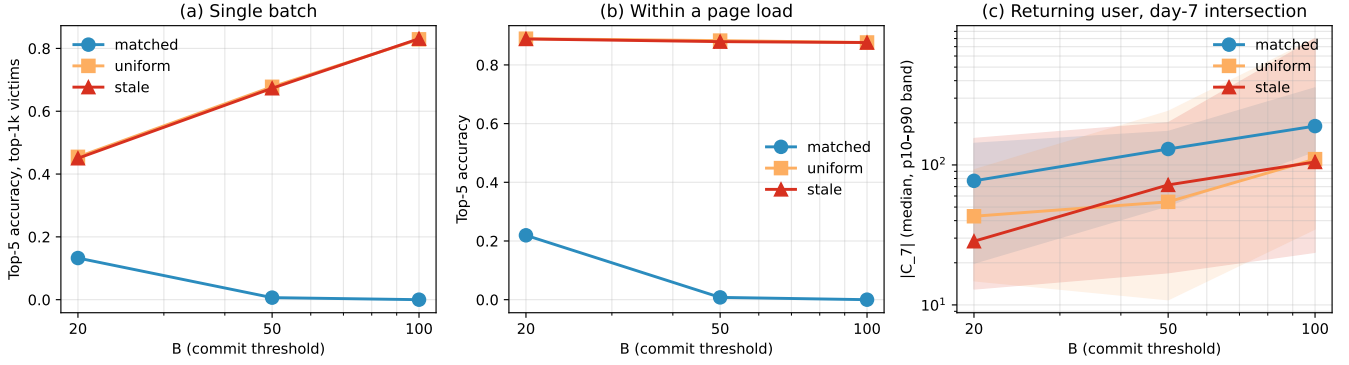
Cover-distribution sensitivity vs commit threshold B ($T_{\max}=300s$, $k=3$, $\lambda_{bg}=100$, $N=1024$)


Figure 8: Cross-scenario cover-distribution sensitivity. Each panel plots one scenario’s headline metric against B on a log axis, with one line per cover distribution. Under wrong D , per-batch top-5 accuracy rises with B (more real votes per batch from the same popular page), page-load top-5 accuracy saturates near 0.88, and the cross-day candidate set $|C_7|$ shrinks to roughly half its matched- D value. Shaded bands on the cross-day panel are p10–p90 across users.

Table 16: Returning user: day-7 candidate-set size $|C_7|$ (median, p10–p90) by B and cover distribution. Larger is better; 30 simulated users per cell.

B	matched	uniform	stale
20	77 (19–142)	43 (14–92)	28 (13–154)
50	130 (51–173)	54 (10–240)	72 (17–199)
100	189 (125–354)	110 (34–797)	105 (23–797)

condition of §5.2; under wrong D , the out-of-distribution covers contribute negligible likelihood under both hypotheses and the ratio collapses to the contribution of the real victim queries alone. This explains the per-batch B -monotone rise in Table 14: top-1k uniform top-5 climbs from 0.45 to 0.83 as B grows from 20 to 100, because more real votes per batch from the same popular page tighten the attacker’s posterior on that page. Larger batches therefore *actively hurt* privacy for popular victims when D is mis-shaped, inverting the matched- D intuition that bigger B is always safer.

Deployment recommendation. Operators deploying CODoH must continuously fit D to a recent empirical estimate of Q (e.g., a streaming top- N heavy-hitters sketch over recent client traffic, with a per-band Zipf interpolation for unseen origins). A static cover universe sampled with the wrong shape yields the wrong- D row of Tables 14–16, regardless of B . We recommend operators (i) ship a default sampler bound to a recent CrUX or Cloudflare-Radar top-1M snapshot, (ii) refit weights periodically (daily is ample given CrUX’s monthly update cycle), and (iii) treat the matched column of these tables as the achievable upper bound on privacy at the chosen B .