

A Calculus for Flow-Limited Authorization

Owen Arden^{a,*}, Anitha Gollamudi^b, Ethan Cecchetti^c, Stephen Chong^b, and
Andrew C. Myers^c

^a *UC Santa Cruz, Santa Cruz, CA, USA*

E-mail: owen@soe.ucsc.edu

^b *Harvard University, Cambridge, MA, USA*

E-mail: agollamudi@g.harvard.edu, chong@seas.harvard.edu

^c *Cornell University, Ithaca, NY, USA*

E-mail: {ethan, andru}@cs.cornell.edu

Abstract. Real-world applications routinely make authorization decisions based on dynamic computation. Reasoning about dynamically computed authority is challenging. Integrity of the system might be compromised if attackers can improperly influence the authorizing computation. Confidentiality can also be compromised by authorization, since authorization decisions are often based on sensitive data such as membership lists and passwords. Previous formal models for authorization do not fully address the security implications of permitting trust relationships to change, which limits their ability to reason about authority that derives from dynamic computation. Our goal is a way to construct dynamic authorization mechanisms that do not violate confidentiality or integrity.

We introduce the Flow-Limited Authorization Calculus (FLAC), which is both a simple, expressive model for reasoning about dynamic authorization and also an information flow control language for securely implementing various authorization mechanisms. FLAC combines the insights of two previous models: it extends the Dependency Core Calculus with features made possible by the Flow-Limited Authorization Model. FLAC provides strong end-to-end information security guarantees even for programs that incorporate and implement rich dynamic authorization mechanisms. These guarantees include noninterference and robust declassification, which prevent attackers from influencing information disclosures in unauthorized ways. We prove these security properties formally for all FLAC programs and explore the expressiveness of FLAC with several examples.

1. Introduction

Authorization mechanisms are critical components in all distributed systems. The policies enforced by these mechanisms constrain what computation may be safely executed, and therefore an expressive policy language is important. Expressive mechanisms for authorization have been an active research area. A variety of approaches have been developed, including authorization logics [27,2,38], often implemented with cryptographic mechanisms [19,13]; role-based access control (RBAC) [20]; and trust management [28,39,9].

However, the security guarantees of authorization mechanisms are usually analyzed using formal models that abstract away the computation and communication performed by the system. Developers must

*Corresponding author. E-mail: owen@soe.ucsc.edu

take great care to faithfully preserve the (often implicit) assumptions of the model, not only when implementing authorization mechanisms, but also when employing them. Simplifying abstractions can help extract formal security guarantees, but abstractions can also obscure the challenges of implementing and using an abstraction securely. This disconnect between abstraction and implementation can lead to vulnerabilities and covert channels that allow attackers to leak or corrupt information.

A common blind spot in many authorization models is confidentiality. Most models cannot express authorization policies that are confidential or are based on confidential data. Real systems, however, use confidential data for authorization all the time: users on social networks receive access to photos based on friend lists, frequent fliers receive tickets based on credit card purchase histories, and doctors exchange patient data while keeping doctor–patient relationships confidential. While many models can ensure, for instance, that only friends are permitted to access a photo, few can say anything about the secondary goal of preserving the confidentiality of the friend list. Such authorization schemes may fundamentally require *some* information to be disclosed, but failing to detect these disclosures can lead to unintentional leaks.

Authorization without integrity is meaningless, so formal models are typically better at enforcing integrity. However, many formal models make unreasonable or unintuitive assumptions about integrity. For instance, in many models (e.g., [27], [2], [28]) authorization policies either do not change or change only when modified by a trusted administrator. This is a reasonable assumption in centralized systems where such an administrator will always exist, but in decentralized systems, there may be no single entity that is trusted by all other entities.

Even in centralized systems, administrators must be careful when performing updates based on partially trusted information, since malicious users may try to use the administrator to carry out an attack on their behalf. Unfortunately, existing models offer little help to administrators that need to reason about how attackers may have influenced security-critical update operations.

Developers need a better programming model for implementing expressive dynamic authorization mechanisms. Errors that undermine the security of these mechanisms are common [29], so we want to be able to verify their security. We argue that information flow control is a lightweight, useful tool for building secure authorization mechanisms. Using information flow control is attractive since it offers compositional, end-to-end security guarantees. However, applying information flow control to these mechanisms in a meaningful way requires building on a theory that integrates authority and information security. In this work, we show how to embed such a theory into a programming model, so that dynamic authorization mechanisms—as well as the programs that employ them—can be statically verified.

Approaching the verification of dynamic authorization mechanisms from this perspective is attractive for two reasons. First, it gives a model for building secure authorization mechanisms by construction rather than verifying them after the fact. This model offers programmers insight into the sometimes subtle interaction between information flow and authorization, and helps programmers address problems early, during the design process. Second, it addresses a core weakness lurking at the heart of existing language-based security schemes: that the underlying policies may change in a way that breaks security. By statically verifying the information security of dynamic authorization mechanisms, we expand the real-world scenarios in which language-based information flow control is useful and strengthen its security guarantees.

We demonstrate that such an embedding is possible by presenting a core language for authorization and information flow control, called the Flow-Limited Authorization Calculus (FLAC). FLAC is a functional language for designing and verifying decentralized authorization protocols. FLAC is inspired by

the Polymorphic Dependency Core Calculus [2] (DCC).¹ Abadi develops DCC as an authorization logic, but DCC is limited to static trust relationships defined externally to DCC programs by a lattice of principals. FLAC supports dynamic authorization by building on the Flow-Limited Authorization Model (FLAM) [7], which unifies reasoning about authority, confidentiality, and integrity. Furthermore, FLAC is a language for information flow control. It uses FLAM’s principal model and FLAM’s logical reasoning rules to define an operational model and type system for authorization computations that preserve information security.

The types in a FLAC program can be considered propositions [42] in an authorization logic, and the programs can be considered proofs that the proposition holds. Well-typed FLAC programs are not only proofs of authorization, but also proofs of secure information flow, ensuring the confidentiality and integrity of authorization policies and of the data those policies depend upon.

FLAC is useful from a logical perspective, but also serves as a core programming model for real language implementations. Since FLAC programs can dynamically authorize computation and flows of information, FLAC applies to more realistic settings than previous authorization logics. Thus FLAC offers more than a type system for proving propositions—FLAC programs do useful computation.

This paper makes the following contributions.

- We define FLAC, a language, type system, and semantics for dynamic authorization mechanisms with strong information security:
 - * Programs in low-integrity contexts exhibit *noninterference*, ensuring attackers cannot leak or corrupt information, and cannot subvert authorization mechanisms.
 - * Programs in higher-integrity contexts exhibit *robust declassification*, ensuring attackers cannot influence authorized disclosures of information.
- We present two authorization mechanisms implemented in FLAC, commitment schemes and bearer credentials, and demonstrate that FLAC ensures the programs that use these mechanisms preserve the desired confidentiality and integrity properties.

We have organized our discussion of FLAC as follows. Section 2 introduces commitment schemes and bearer credentials, two examples of dynamic authorization mechanisms we use to explore the features of FLAC. Section 3 reviews the FLAM principal lattice [7], and Section 4 defines the FLAC language and type system. FLAC implementations of the dynamic authorization examples are presented in Section 5, and their properties are examined. Section 6 explores aspects of FLAC’s proof theory, and Section 7 discusses semantic security guarantees of FLAC programs, including noninterference and robust declassification. We explore related work in Section 8 and conclude in Section 9.

Many of the contributions of this paper were previously published by Arden and Myers [5]. This article expands upon and strengthens the formal results, as well as corrects several technical errors in that work. Most of the semantic security proofs in Section 7 are either new to this work or were redeveloped from scratch to account for changes made to the semantics and type system.

Furthermore, subsequent work based on FLAC by Ceccetti *et al.* [15] and Gollamudi *et al.* [?] exposed alternate design decisions that improve the connection between FLAC and cryptographic implementations of its protection abstractions. We feel some of these changes are objectively better than the original abstractions, and so incorporate them into the core FLAC formalism so that future work based on FLAC may benefit. Significant departures from the original formalization are footnoted, but minor changes and corrections are included without comment.

¹DCC was first presented in [4]. We use the abbreviation DCC to refer to the extension to polymorphic types in [2].

2. Dynamic authorization mechanisms

Dynamic authorization is challenging to implement correctly since authority, confidentiality, and integrity interact in subtle ways. FLAC helps programmers securely implement both authorization mechanisms and code that uses them. FLAC types support the definition of compositional security abstractions, and vulnerabilities in the implementations of these abstractions are caught statically. Further, the guarantees offered by FLAC simplify reasoning about the security properties of these abstractions.

We illustrate the usefulness and expressive power of FLAC using two important security mechanisms: commitment schemes and bearer credentials. We show in Section 5 that these mechanisms can be implemented using FLAC, and that their security goals are easily verified in the context of FLAC.

2.1. Commitment schemes

A commitment scheme [35] allows one party to give another party a “commitment” to a secret value without revealing the value. The committing party may later reveal the secret in a way that convinces the receiver that the revealed value is the value originally committed.

Commitment schemes provide three essential operations: `commit`, `reveal`, and `open`. Suppose p wants to commit to a value to principal q . First, p applies `commit` to the value and provides the result to q which does not reveal the value to q . When p wishes to reveal the value, it supplies q with a `reveal` operation for the commitment it sent previously. Then q uses `reveal` to open the received value, revealing it to q .

A commitment scheme must have several properties in order to be secure. First, q should not be able to open a value that hasn’t been committed to by p , since this could allow q to manipulate p to open a value it had not committed to. Second, q should not be able to learn a secret of p that has not been committed to or revealed by p . Third, p should not be able to modify the committed value after it is received by q .

One might wonder why a programmer would bother to create high-level *implementations* of operations like `commit`, `reveal`, and `open`. Why not simply treat these as primitive operations and give them type signatures so that programs using them can be type-checked with respect to those signatures? The answer is that an error in a type signature could lead to a serious vulnerability. Therefore, we want more assurance that the type signatures are correct. Implementing such operations in FLAC is often easy and ensures that the type signature is consistent with a set of assumptions about existing trust relationships and the information flow context the operations are used within. These FLAC-based implementations serve as language-based models of the security properties achieved by implementations that use cryptography or trusted third parties.

2.2. Bearer credentials with caveats

A bearer credential is a capability that grants authority to any entity that possesses it. Many authorization mechanisms used in distributed systems employ bearer credentials in some form. Browser cookies that store session tokens are one example: after a website authenticates a user’s identity, it gives the user a token to use in subsequent interactions. Since it is infeasible for attackers to guess the token, the website grants the authority of the user to any requests that include the token.

Bearer credentials create an information security conundrum for authorization mechanisms. Though they efficiently control access to restricted resources, they create vulnerabilities and introduce covert channels when used incorrectly. For example, suppose Alice shares a remotely-hosted photo with her friends by giving them a credential to access the photo. Giving a friend such a credential doesn’t disclose their friendship, but each friend that accesses the photo implicitly discloses the friendship to the hosting

service. Such covert channels are pervasive, both in classic distributed authorization mechanisms like SPKI/SDSI [19], as well as in more recent ones like Macaroons [13].

Bearer credentials can also lead to vulnerabilities if they are leaked. If an attacker obtains a credential, it can exploit the authority of the credential. Thus, to limit the authority of a credential, approaches like SPKI/SDSI and Macaroons provide *constrained delegation* in which a newly issued credential attenuates the authority of an existing one by adding *caveats*. Caveats require additional properties to hold for the bearer to be granted authority. Session tokens, for example, might have a caveat that restricts the source IP address or encodes an expiration time. As pointed out by Birgisson et al. [13], caveats themselves can introduce covert channels if the properties reveal sensitive information.

FLAC is an effective framework for reasoning about bearer credentials with caveats since it captures the flow of credentials in programs as well as the sensitivity of the information the credentials and caveats derive from. We can reason about credentials and the programs that use them in FLAC with an approach similar to that used for commitment schemes. That we can do so in a straightforward way is somewhat remarkable: prior formalizations of credential mechanisms (e.g., [13,24,11]) usually do not consider confidentiality nor provide end-to-end guarantees about credential propagation.

3. The FLAM Principal Lattice

Like many models, FLAM uses *principals* to represent the authority of all entities relevant to a system. However, FLAM's principals and their algebraic properties are richer than in most models, so we briefly review the FLAM principal model and notation. Further details are found in the earlier paper [7].

Primitive principals such as Alice, Bob, etc., are represented as elements n of a (potentially infinite) set of names $\mathcal{N}$.² In addition, FLAM uses $\top$ to represent a universally trusted principal and $\perp$ to represent a universally untrusted principal. The combined authority of two principals, p and q , is represented by the conjunction $p \wedge q$, whereas the authority of either p or q is the disjunction $p \vee q$.

Unlike principals in other models, FLAM principals also represent information flow policies. The confidentiality of principal p is represented by the principal $p^\rightarrow$, called p 's confidentiality projection. It denotes the authority necessary to *learn* anything p can learn. The integrity of principal p is represented by $p^\leftarrow$, called p 's integrity projection. It denotes the authority to *influence* anything p can influence. All authority may be represented as some combination of confidentiality and integrity. For instance, principal p is equivalent to the conjunction $p^\rightarrow \wedge p^\leftarrow$, and in fact any FLAM principal can be written $p^\rightarrow \wedge q^\leftarrow$ for some p and q . The closure of the set of names $\mathcal{N}$ plus $\top$ and $\perp$ under the operators³ $\wedge, \vee, \leftarrow, \rightarrow$ forms a lattice $\mathcal{L}$ ordered by an *acts-for relation* $\succcurlyeq$, defined by the inference rules in Figure 1. We write operators $\leftarrow, \rightarrow$ with higher precedence than $\wedge, \vee$; for instance, $p \wedge q^\leftarrow$ is equal to $p^\rightarrow \wedge (p \wedge q)^\leftarrow$. Projections distribute over $\wedge$ and $\vee$ so, for example, $(p \wedge q)^\leftarrow = (p^\leftarrow \wedge q^\leftarrow)$. The confidentiality and integrity authority of principals are disjoint, so the confidentiality projection of an integrity projection is $\perp$ and vice-versa: $(p^\leftarrow)^\rightarrow = \perp = (p^\rightarrow)^\leftarrow$.

An advantage of this model is that secure information flow can be defined in terms of authority. An information flow policy q is at least as *restrictive* as a policy p if q has at least the confidentiality au-

²Using $\mathcal{N}$ as the set of all names is convenient in our formal calculus, but a general-purpose language based on FLAC may wish to dynamically allocate names at runtime. Since knowing or using a principal's name holds no special privilege in FLAC, this presents no fundamental difficulties. To use dynamically allocated principals in type signatures, however, the language's type system should support types in which principal names may be existentially quantified.

³FLAM defines an additional set of operators called *ownership projections*, which we omit here to simplify our presentation.

$$\boxed{\mathcal{L} \models p \geq q}$$

$$\begin{array}{c}
\text{[BOT]} \mathcal{L} \models p \geq \perp \quad \text{[TOP]} \mathcal{L} \models \top \geq p \quad \text{[REFL]} \mathcal{L} \models p \geq p \quad \text{[PROJ]} \frac{\mathcal{L} \models p \geq q}{\mathcal{L} \models p^\pi \geq q^\pi} \quad \text{[PROJR]} \mathcal{L} \models p \geq p^\pi \\
\\
\text{[CONJL]} \frac{\mathcal{L} \models p_k \geq p}{\mathcal{L} \models p_1 \wedge p_2 \geq p} \quad \text{[CONJR]} \frac{\mathcal{L} \models p \geq p_1}{\mathcal{L} \models p \geq p_1 \wedge p_2} \quad \text{[DISJL]} \frac{\mathcal{L} \models p_1 \geq p}{\mathcal{L} \models p_1 \vee p_2 \geq p} \\
\\
\text{[DISJR]} \frac{\mathcal{L} \models p \geq p_k}{\mathcal{L} \models p \geq p_1 \vee p_2} \quad \text{[TRANS]} \frac{\mathcal{L} \models p \geq q \quad \mathcal{L} \models q \geq r}{\mathcal{L} \models p \geq r}
\end{array}$$

Fig. 1. Static principal lattice rules, adapted from FLAM [7]. The projection π may be either confidentiality ($\rightarrow$) or integrity ($\leftarrow$).

thority $p^\rightarrow$ and p has at least the integrity authority $q^\leftarrow$. This relationship between the confidentiality and integrity of p and q reflects the usual duality seen in information flow control [12]. As in [7], we use the following shorthand for relating principals by policy restrictiveness:

$$\begin{aligned}
p \sqsubseteq q &\triangleq (p^\leftarrow \wedge q^\rightarrow) \geq (q^\leftarrow \wedge p^\rightarrow) \\
p \sqcup q &\triangleq (p \wedge q)^\rightarrow \wedge (p \vee q)^\leftarrow \\
p \sqcap q &\triangleq (p \vee q)^\rightarrow \wedge (p \wedge q)^\leftarrow
\end{aligned}$$

Thus, $p \sqsubseteq q$ indicates the direction of secure information flow: from p to q . The information flow join $p \sqcup q$ is the least restrictive principal that both p and q flow to, and the information flow meet $p \sqcap q$ is the most restrictive principal that flows to both p and q .

In FLAM, the ability to “speak for” another principal is an integrity relationship between principals. This makes sense intuitively, because speaking for another principal influences that principal’s trust relationships and information flow policies. FLAM defines the *voice* of a principal p , written $\nabla(p)$, as the integrity necessary to speak for that principal. Given a principal expressed in normal form⁴ as $q^\rightarrow \wedge r^\leftarrow$, the voice of that principal is

$$\nabla(q^\rightarrow \wedge r^\leftarrow) \triangleq q^\leftarrow \wedge r^\leftarrow$$

For example, the voice of Alice, $\nabla(\text{Alice})$, is $\text{Alice}^\leftarrow$. The voice of Alice’s confidentiality $\nabla(\text{Alice}^\rightarrow)$ is also $\text{Alice}^\leftarrow$.

4. Flow-Limited Authorization Calculus

FLAC uses information flow to reason about the security implications of dynamically computed authority. Like previous information-flow type systems [37], FLAC incorporates types for reasoning about

⁴In normal form, a principal is the conjunction of a confidentiality principal and an integrity principal. See [7] for details.

$n \in \mathcal{N}$ (primitive principals)
 $x \in \mathcal{V}$ (variable names)

$$\begin{aligned}
 p, \ell, pc &::= n \mid \top \mid \perp \mid p^\rightarrow \mid p^\leftarrow \mid p \wedge p \mid p \vee p \\
 \tau &::= (p \geq p) \mid \text{unit} \mid \tau + \tau \mid \tau \times \tau \\
 &\quad \mid \tau \xrightarrow{pc} \tau \mid \ell \text{ says } \tau \mid X \mid \forall X[pc]. \tau \\
 v &::= () \mid \langle w, w \rangle \mid \langle p \geq p \rangle \mid \bar{\eta}_\ell w \\
 &\quad \mid \text{inj}_i w \mid \lambda(x:\tau)[pc]. e \mid \Lambda X[pc]. e \\
 w &::= v \mid w \text{ where } v \\
 e &::= x \mid w \mid e e \mid \langle e, e \rangle \mid \eta_\ell e \\
 &\quad \mid e \tau \mid \text{proj}_i e \mid \text{inj}_i e \\
 &\quad \mid \text{case } e \text{ of } \text{inj}_1(x). e \mid \text{inj}_2(x). e \\
 &\quad \mid \text{bind } x = e \text{ in } e \mid \text{assume } e \text{ in } e \\
 &\quad \mid e \text{ where } v
 \end{aligned}$$

Fig. 2. FLAC syntax. Terms using where are syntactically prohibited in the source language and are produced only during evaluation.

$e \longrightarrow e'$

[E-APP] $(\lambda(x:\tau)[pc]. e) w \longrightarrow e[x \mapsto w]$ [E-UNPAIR] $\text{proj}_i \langle w_1, w_2 \rangle \longrightarrow w_i$ [E-BINDM] $\text{bind } x = \bar{\eta}_\ell w \text{ in } e \longrightarrow e[x \mapsto w]$ [E-UNITM] $\eta_\ell w \longrightarrow \bar{\eta}_\ell w$	[E-TAPP] $(\Lambda X[pc]. e) \tau \longrightarrow e[X \mapsto \tau]$ [E-CASE] $(\text{case } (\text{inj}_i w) \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2) \longrightarrow e_i[x \mapsto w]$ [E-ASSUME] $\text{assume } \langle p \geq q \rangle \text{ in } e \longrightarrow e \text{ where } \langle p \geq q \rangle$ [E-EVAL] $\frac{e \longrightarrow e'}{E[e] \longrightarrow E[e']}$
---	---

$$\begin{aligned}
 E &::= [\cdot] \mid E e \mid w E \mid E \tau \mid \langle E, e \rangle \mid \langle w, E \rangle \mid \text{proj}_i E \mid \text{inj}_i E \mid \eta_\ell E \\
 &\quad \mid \text{bind } x = E \text{ in } e \mid \text{assume } E \text{ in } e \mid \text{case } E \text{ of } \text{inj}_1(x). e \mid \text{inj}_2(x). e \mid E \text{ where } v
 \end{aligned}$$

Fig. 3. FLAC operational semantics

information flow, but FLAC's type system goes further by using Flow-Limited Authorization [7] to ensure that principals cannot use FLAC programs to exceed their authority, or to leak or corrupt information. FLAC is based on DCC [2], but unlike DCC, FLAC supports reasoning about authority deriving from the evaluation of FLAC terms. In contrast, all authority in DCC derives from trust relationships defined by a fixed, external lattice of principals. Thus, using an approach based on DCC in systems where trust relationships change dynamically could introduce vulnerabilities like delegation loopholes, probing and poaching attacks, and authorization side channels [7].

$\Pi; \Gamma; pc \vdash e : \tau$		
[VAR]	$\Pi; \Gamma, x : \tau, \Gamma'; pc \vdash x : \tau$	[UNIT] $\Pi; \Gamma; pc \vdash () : \text{unit}$ [DEL] $\Pi; \Gamma; pc \vdash \langle p \geq q \rangle : (p \geq q)$
[LAM]	$\frac{\Pi; \Gamma, x : \tau_1; pc' \vdash e : \tau_2}{\Pi; \Gamma; pc \vdash \lambda(x : \tau_1)[pc']. e : (\tau_1 \xrightarrow{pc'} \tau_2)}$	[TLAM] $\frac{\Pi; \Gamma, X; pc' \vdash e : \tau}{\Pi; \Gamma; pc \vdash \Lambda X[pc']. e : \forall X[pc']. \tau}$
[APP]	$\frac{\Pi; \Gamma; pc \vdash e : (\tau_1 \xrightarrow{pc'} \tau_2) \quad \Pi; \Gamma; pc \vdash e' : \tau_1 \quad \Pi \Vdash pc \sqsubseteq pc'}{\Pi; \Gamma; pc \vdash (e e') : \tau_2}$	[TAPP] $\frac{\Pi; \Gamma; pc \vdash e : \forall X[pc']. \tau \quad \Pi \Vdash pc \sqsubseteq pc'}{\Pi; \Gamma; pc \vdash (e \tau') : \tau[X \mapsto \tau']} \quad \tau' \text{ well-formed in } \Gamma$
[PAIR]	$\frac{\Pi; \Gamma; pc \vdash e_1 : \tau_1 \quad \Pi; \Gamma; pc \vdash e_2 : \tau_2}{\Pi; \Gamma; pc \vdash \langle e_1, e_2 \rangle : \tau_1 \times \tau_2}$	[UNPAIR] $\frac{\Pi; \Gamma; pc \vdash e : \tau_1 \times \tau_2}{\Pi; \Gamma; pc \vdash (\text{proj}_i e) : \tau_i}$
[INJ]	$\frac{\Pi; \Gamma; pc \vdash e : \tau_i \quad i \in \{1, 2\}}{\Pi; \Gamma; pc \vdash (\text{inj}_i e) : \tau_1 + \tau_2}$	[CASE] $\frac{\Pi; \Gamma; pc \vdash e : \tau_1 + \tau_2 \quad \Pi \vdash pc \sqsubseteq \tau \quad \Pi; \Gamma, x : \tau_1; pc \vdash e_1 : \tau \quad \Pi; \Gamma, x : \tau_2; pc \vdash e_2 : \tau}{\Pi; \Gamma; pc \vdash \text{case } e \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2 : \tau}$
[UNITM]	$\frac{\Pi; \Gamma; pc \vdash e : \tau \quad \Pi \Vdash pc \sqsubseteq \ell}{\Pi; \Gamma; pc \vdash \eta_\ell e : \ell \text{ says } \tau}$	[SEALED] $\frac{\Pi; \Gamma; pc \vdash v : \tau}{\Pi; \Gamma; pc \vdash \bar{\eta}_\ell v : \ell \text{ says } \tau}$
[BINDM]	$\frac{\Pi; \Gamma; pc \vdash e : \ell \text{ says } \tau' \quad \Pi; \Gamma, x : \tau'; pc \sqcup \ell \vdash e' : \tau \quad \Pi \vdash pc \sqcup \ell \sqsubseteq \tau}{\Pi; \Gamma; pc \vdash \text{bind } x = e \text{ in } e' : \tau}$	
[ASSUME]	$\frac{\Pi; \Gamma; pc \vdash e : (p \geq q) \quad \Pi \Vdash pc \geq \nabla(q) \quad \Pi \Vdash \nabla(p^{\rightarrow}) \geq \nabla(q^{\rightarrow}) \quad \Pi, \langle p \geq q \rangle; \Gamma; pc \vdash e' : \tau}{\Pi; \Gamma; pc \vdash \text{assume } e \text{ in } e' : \tau}$	
[WHERE]	$\frac{\Pi; \Gamma; pc \vdash v : (p \geq q) \quad \Pi \Vdash \bar{pc} \geq \nabla(q) \quad \Pi \Vdash \nabla(p^{\rightarrow}) \geq \nabla(q^{\rightarrow}) \quad \Pi, \langle p \geq q \rangle; \Gamma; pc \vdash e : \tau}{\Pi; \Gamma; pc \vdash (e \text{ where } v) : \tau}$	

Fig. 4. FLAC type system.

Figure 2 defines the FLAC syntax. The core FLAC operational semantics and evaluation contexts [44] in Figure 3 are mostly standard except for E-ASSUME and E-UNITM, which we discuss below, along with additional rules that handle the propagation of the where terms introduced by E-ASSUME.

The core FLAC type system is presented in Figure 4. FLAC typing judgments have the form $\Pi; \Gamma; pc \vdash e : \tau$. The *delegation context*, Π , contains a set of labeled dynamic trust relationships $\langle p \geq q \rangle$ where $p \geq q$ (read as “ p acts for q ”) is a delegation from q to p . The *typing context*, Γ , is a map from variables to types, and pc is the *program counter label*, a FLAM principal representing the confidentiality and integrity of control flow. The type system makes frequent use of judgments adapted from FLAM’s inference rules [7]. These rules, adapted to FLAC, are presented in Figure 5.⁵

⁵FLAM’s rules [7] also include *query* and *result* labels as part of the judgment context that represent the confidentiality and integrity of a FLAM query context and result, respectively. These labels are unnecessary in FLAC because we use FLAM

[R-STATIC]	$\frac{\mathcal{L} \models p \geq q}{\Pi \Vdash p \geq q}$	[R-ASSUME]	$\frac{\langle p \geq q \rangle \in \Pi \quad \Pi \Vdash \nabla(p^{\rightarrow}) \geq \nabla(q^{\rightarrow})}{\Pi \Vdash p \geq q}$
[R-CONJL]	$\frac{\Pi \Vdash p_k \geq q \quad k \in \{1, 2\}}{\Pi \Vdash p_1 \wedge p_2 \geq q_2}$	[R-CONJR]	$\frac{\Pi \Vdash p \geq q_1 \quad \Pi \Vdash p \geq q_2}{\Pi \Vdash p \geq q_1 \wedge q_2}$
[R-DISJL]	$\frac{\Pi \Vdash p_1 \geq q \quad \Pi \Vdash p_2 \geq q}{\Pi \Vdash p_1 \vee p_2 \geq q}$	[R-DISJR]	$\frac{\Pi \Vdash p \geq q_k \quad k \in \{1, 2\}}{\Pi \Vdash p \geq q_1 \vee q_2}$
[R-TRANS]	$\frac{\Pi \Vdash p \geq q \quad \Pi \Vdash q \geq r \quad \Pi \Vdash pc \geq \nabla(r^{\rightarrow})}{\Pi \Vdash p \geq r}$		

Fig. 5. Inference rules for robust assumption, adapted from FLAM [7].

Since FLAC is a pure functional language, it might seem odd for FLAC to have a label for the program counter; such labels are usually used to control implicit flows through assignments (e.g., in [36,32]). The purpose of FLAC's pc label is to control a different kind of side effect: changes to the delegation context, Π .⁶ In order to control what information can influence whether a new trust relationship is added to the delegation context, the type system tracks the confidentiality and security of control flow. Viewed as an authorization logic, FLAC's type system has the unique feature that it expresses deduction constrained by an information flow context. For instance, if we have $\varphi \xrightarrow{p^{\leftarrow}} \psi$ and φ , then (via APP) we may derive ψ in a context with integrity $p^{\leftarrow}$, but not in contexts that don't flow to $p^{\leftarrow}$. This feature offers needed control over how principals may apply existing facts to derive new facts.

Many FLAC terms are standard, such as pairs $\langle e_1, e_2 \rangle$, projections $\text{proj}_i e$, variants $\text{inj}_i e$, polymorphic type abstraction, $\Lambda X[pc]. e$, and case expressions. Function abstraction, $\lambda(x : \tau)[pc]. e$, includes a pc label that constrains the information flow context in which the function may be applied. The rule APP ensures that function application respects these policies, requiring that the robust FLAM judgment $\Pi \Vdash pc \sqsubseteq pc'$ holds. This judgment ensures that the current program counter label, pc , flows to the function label, pc' .

Branching occurs in case expressions, which conditionally evaluate one of two expressions. The rule CASE ensures that both expressions have the same type and thus the same protection level. The premise $\Pi \vdash pc \sqsubseteq \tau$ ensures that this type protects the current pc label.⁷

Like DCC, FLAC uses monadic operators to track dependencies. The monadic unit term $\eta_\ell v$ (UNITM) says that a value v of type τ is *protected at level* ℓ . This protected value has the type ℓ says τ , meaning that it has the confidentiality and integrity of principal ℓ . Because v could implicitly reveal information about the dependencies of the computation that produced it, UNITM requires that $\Pi \Vdash pc \sqsubseteq \ell$.⁸

judgments only in the type system—these “queries” only occur at compile time and do not create information flows. We formalize the connection between FLAC and FLAM in Appendix A.

⁶DFLATE [?], an extension of FLAC, uses the same pc label to control implicit flows due to communication. Extensions of FLAC to support mutable references could control implicit flows similarly.

⁷This premise simplifies our proofs, but does not appear to be strictly necessary; BINDM ensures the same property.

⁸Neither DCC [4] nor the original FLAC formalization [5] included this premise. DCC does not maintain a pc label at all. FLAC originally used a version of the DCC rule, but Cecchetti et al. [15] and Gollamudi et al. [?] added the pc restriction in support of a non-commutative says. See Section 6.1 for additional details.

Computation on protected values must occur in a protected context (“in the monad”), expressed using a monadic bind term. The typing rule BINDM ensures that the result of the computation protects the confidentiality and integrity of protected values. For instance, the expression $\text{bind } x = \eta_\ell v \text{ in } \eta_{\ell'} x$ is only well-typed if ℓ' protects values with confidentiality and integrity ℓ . Since case expressions may use the variable x for branching, BINDM raises the pc label to $pc \sqcup \ell$ to conservatively reflect the control-flow dependency.

Protection levels are defined by the set of inference rules in Figure 7, adapted from [41]. Expressions with unit type (P-UNIT) do not propagate any information, so they protect information at any ℓ . Product types protect information at ℓ if both components do (P-PAIR). Function types protect information at ℓ if the return type and function label does (P-FUN), and polymorphic types protect information at whatever level the abstracted type and type function label does (P-TFUN). Finally, if ℓ flows to ℓ' , then ℓ' says τ protects information at ℓ (P-LBL).⁹

Most of the novelty of FLAC lies in its delegation values and assume terms. These terms enable expressive reasoning about authority and information flow control. A delegation value serves as evidence of trust. For instance, the term $\langle p \geq q \rangle$, read “ p acts for q ”, is evidence that q trusts p . Delegation values have *acts-for* types; $\langle p \geq q \rangle$ has type $(p \geq q)$.¹⁰ The assume term enables programs to use evidence securely to create new flows between protection levels. In the typing context $\emptyset; x : p^\leftarrow$ says $\tau; q^\leftarrow$ (i.e., $\Pi = \emptyset, \Gamma = x : p^\leftarrow$ says τ , and $pc = q^\leftarrow$), the following expression is not well typed:

$$\text{bind } x' = x \text{ in } \eta_{q^\leftarrow} x'$$

since $p^\leftarrow$ does not flow to $q^\leftarrow$, as required by the premise $\Pi \vdash \ell \sqsubseteq \tau$ in rule BINDM. Specifically, we cannot derive $\Pi \vdash p^\leftarrow \sqsubseteq q^\leftarrow$ says τ since P-LBL requires the FLAM judgment $\Pi \Vdash p^\leftarrow \sqsubseteq q^\leftarrow$ to hold.

However, the following expression is well typed:

$$\text{assume } \langle p^\leftarrow \geq q^\leftarrow \rangle \text{ in } \text{bind } x' = x \text{ in } \eta_{q^\leftarrow} x'$$

The difference is that the assume term adds a trust relationship, represented by an expression with an *acts-for* type, to the delegation context. In this case, the expression $\langle p^\leftarrow \geq q^\leftarrow \rangle$ adds a trust relationship that allows $p^\leftarrow$ to flow to $q^\leftarrow$. This is secure since $pc = q^\leftarrow$, meaning that only principals with integrity $q^\leftarrow$ have influenced the computation. With $\langle p^\leftarrow \geq q^\leftarrow \rangle$ in the delegation context, added via the ASSUME rule, the premises of BINDM are now satisfied, so the expression type-checks.

Creating a delegation value requires no special privilege because the type system ensures only high-integrity delegations are used as evidence that enable new flows. Using low-integrity evidence for authorization would be insecure since attackers could use delegation values to create new flows that reveal secrets or corrupt data. The premises of the ASSUME rule ensure the integrity of dynamic authorization computations that produce values like $\langle p^\leftarrow \geq q^\leftarrow \rangle$ in the example above.¹¹ The second premise, $\Pi \Vdash pc \geq \nabla(q)$, requires that the pc has enough integrity to be trusted by q , the principal whose security is affected. For instance, to make the assumption $p \geq q$, the evidence represented by the term e must have at least the integrity of the voice of q , written $\nabla(q)$. Since the pc bounds the restrictiveness of the

⁹DCC [4] and the original FLAC formalization included an additional protection rule that considered ℓ to be protected by ℓ' says τ if τ protects ℓ (even if ℓ' does not). This rule was removed by Ceccetti *et al.* and Gollamudi *et al.* [?] to make says non-commutative.

¹⁰This correspondence with delegation values makes *acts-for* types a kind of singleton type [18].

¹¹These premises are related to the robust FLAM rule LIFT.

$\boxed{e \longrightarrow e'}$	
[W-APP]	$(w \text{ where } v) e \longrightarrow (w e) \text{ where } v$
[W-TAPP]	$(w \text{ where } v) \tau \longrightarrow (w \tau) \text{ where } v$
[W-UNPAIR]	$\text{proj}_i (w \text{ where } v) \longrightarrow (\text{proj}_i w) \text{ where } v$
[W-CASE]	$(\text{case } (w \text{ where } v) \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2) \longrightarrow (\text{case } w \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2) \text{ where } v$
[W-BINDM]	$\text{bind } x = (w \text{ where } v) \text{ in } e \longrightarrow (\text{bind } x = w \text{ in } e) \text{ where } v$
[W-ASSUME]	$\text{assume } (w \text{ where } v) \text{ in } e \longrightarrow (\text{assume } w \text{ in } e) \text{ where } v$

Fig. 6. Propagation of where terms

dependencies of e , this ensures that only information with integrity $\nabla(q)$ or higher may influence the evaluation of e . The third premise, $\Pi \Vdash \nabla(p^\rightarrow) \geq \nabla(q^\rightarrow)$, ensures that principal p has sufficient integrity to be trusted to enforce q 's confidentiality, $q^\rightarrow$. This premise means that q permits data to be relabeled from $q^\rightarrow$ to $p^\rightarrow$.¹²

Assumption terms evaluate to where expressions (rule E-ASSUME). These expressions are a purely formal bookkeeping mechanism (i.e., they would be unnecessary in a FLAC-based implementation) for tracking implicit flows due to dynamic delegations. They record and maintain the authorization evidence used to justify new flows of information during evaluation. They are not part of the source language and generated only by the evaluation rules. The term $e \text{ where } \langle p \geq q \rangle$ records that e is evaluated in a context that includes the delegation $\langle p \geq q \rangle$.

The rule WHERE gives a typing rule for where terms; though similar to ASSUME, it requires only that $\nabla(q)$ delegate to the distinguished label $\overline{pc}$. The use of $\overline{pc}$ is purely technical: our proofs in Section 7 use $\overline{pc}$ to help reason about what new flows may have created by assume terms. The only requirement is that $\overline{pc}$ be as trusted as the pc label used to type-check the source program (or programs) that generated the where term. Since the pc increases monotonically when typing subexpressions, the typical choice for $\overline{pc}$ for a source program e such that $\Pi; \Gamma; pc \vdash e : \tau$ would be pc .¹³

Figure 6 presents evaluation rules for where terms. The rules are designed to treat where values like the value they enclose. For instance, applying a where term (rule W-APP) simply moves the value it is applied to inside the where term. If the where term was wrapping a lambda expression, then it may now be applied via APP. Otherwise, further reduction steps via W-APP may be necessary. We use the syntactic category w (see Figure 2) to refer to fully-evaluated where terms, or where *values*. In other words, a where value w is an expression consisting of a value v enclosed by one or more where clauses. A where value usually behaves like a value, but it is occasionally convenient to distinguish them.¹⁴

¹²More precisely, it means that the voice of q 's confidentiality, $\nabla(q^\rightarrow)$, permits data to be relabeled from $q^\rightarrow$ to $p^\rightarrow$. Recall that $\nabla(\text{Alice}^\rightarrow)$ is just Alice's integrity projection: $\text{Alice}^\leftarrow$.

¹³Choosing $\overline{pc}$ to be $\top^\leftarrow$ is always valid, but selecting a more restrictive label permits finer-grained reasoning about what downgrades may occur in e .

¹⁴The original FLAC formalization did not distinguish where values and values, and did not include the rules in Figure 6. Unfortunately, this resulted in stuck terms when where terms were not propagated appropriately.

$$\boxed{\Pi \vdash \ell \sqsubseteq \tau}$$

$$\begin{array}{lll}
[\text{P-UNIT}] \quad \Pi \vdash \ell \sqsubseteq \text{unit} & [\text{P-PAIR}] \quad \frac{\Pi \vdash \ell \sqsubseteq \tau_1 \quad \Pi \vdash \ell \sqsubseteq \tau_2}{\Pi \vdash \ell \sqsubseteq \tau_1 \times \tau_2} & [\text{P-FUN}] \quad \frac{\Pi \vdash \ell \sqsubseteq \tau_2 \quad \Pi \vdash \ell \sqsubseteq pc'}{\Pi \vdash \ell \sqsubseteq \tau_1 \xrightarrow{pc'} \tau_2} \\
[\text{P-TFUN}] \quad \frac{\Pi \vdash \ell \sqsubseteq \tau \quad \Pi \vdash \ell \sqsubseteq pc'}{\Pi \vdash \ell \sqsubseteq \forall X[pc']. \tau} & & [\text{P-LBL}] \quad \frac{\Pi \vdash \ell \sqsubseteq \ell'}{\Pi \vdash \ell \sqsubseteq \ell' \text{ says } \tau}
\end{array}$$

Fig. 7. Type protection levels

$$\begin{aligned}
&\text{commit} : \forall N[p^{\leftarrow}]. \forall X[p^{\leftarrow}]. N \xrightarrow{p^{\leftarrow}} p^{\rightarrow} \text{ says } X \xrightarrow{p^{\leftarrow}} p \text{ says } (N, X) \\
&\text{commit} = \Lambda N[p^{\leftarrow}]. \Lambda X[p^{\leftarrow}]. \lambda(n:N)[p^{\leftarrow}]. \lambda(x:p^{\rightarrow} \text{ says } X)[p^{\leftarrow}]. \\
&\quad \text{assume } \langle \perp^{\leftarrow} \geq p^{\leftarrow} \rangle \text{ in bind } x' = x \text{ in } \eta_p(n, x') \\
&\text{reveal} : \forall N[\nabla(p^{\rightarrow})]. \forall X[q^{\leftarrow}]. p \text{ says } (N, X) \xrightarrow{q^{\leftarrow}} q^{\rightarrow} \wedge p^{\leftarrow} \text{ says } (N, X) \\
&\text{reveal} = \Lambda N[\nabla(p^{\rightarrow}) \wedge p^{\leftarrow}]. \text{assume } \langle \nabla(q^{\rightarrow}) \geq \nabla(p^{\rightarrow}) \rangle \text{ in assume } \langle q^{\leftarrow} \geq p^{\leftarrow} \rangle \text{ in} \\
&\quad \text{assume } \langle q^{\rightarrow} \geq p^{\rightarrow} \rangle \text{ in } \Lambda X[q^{\leftarrow}]. \lambda(x:p \text{ says } (N, X))[q^{\leftarrow}]. \text{bind } x' = x \text{ in } \eta_{q^{\rightarrow} \wedge p^{\leftarrow}} x' \\
&\text{open} : \forall N[q^{\leftarrow}]. \forall X[q^{\leftarrow}]. (\forall Y[q^{\leftarrow}]. p \text{ says } (N, Y) \xrightarrow{q^{\leftarrow}} q^{\rightarrow} \wedge p^{\leftarrow} \text{ says } (N, Y)) \xrightarrow{q^{\leftarrow}} \\
&\quad p \text{ says } (N, X) \xrightarrow{q^{\leftarrow}} q^{\leftarrow} \text{ says } (q^{\rightarrow} \wedge p^{\leftarrow} \text{ says } (N, X))) \\
&\text{open} = \Lambda N[q^{\leftarrow}]. \Lambda X[q^{\leftarrow}]. \lambda(f : (\forall Y[q^{\leftarrow}]. p \text{ says } (N, Y) \xrightarrow{q^{\leftarrow}} q^{\rightarrow} \wedge p^{\leftarrow} \text{ says } (N, Y)))[q^{\leftarrow}]. \\
&\quad \lambda(x:p \text{ says } (N, X))[q^{\leftarrow}]. \eta_{q^{\leftarrow}}(f \ X \ x)
\end{aligned}$$

Fig. 8. FLAC implementations of commitment scheme operations.

5. Examples revisited

We can now implement our examples from Section 2 in FLAC. Using FLAC ensures that authority and information flow assumptions are explicit, and that programs using these abstractions are secure with respect to those assumptions. In this section, we discuss at a high level how FLAC types help enforce specific end-to-end security properties for commitment schemes and bearer credentials. Section 7 formalizes the semantic security properties of all well-typed FLAC programs.

5.1. Commitment Schemes

Figure 8 contains the essential operations of a commitment scheme—`commit`, `reveal`, and `open`—implemented in FLAC. Principal p commits to pairs of the form (n, v) where n is a term that encodes a

type-level integer N .¹⁵ To prevent q from influencing which committed values are revealed, p commits to a single value for each integer type.

The `commit` operation takes a value of any type (hence $\forall X$) with confidentiality $p^\rightarrow$ and produces a value with confidentiality and integrity p . In other words, p *endorses* [47] the value to have integrity $p^\leftarrow$.

Attackers should not be able to influence whether principal p commits to a particular value. The pc constraint on `commit` ensures that only principal p and principals trusted with at least p 's integrity, $p^\leftarrow$, may apply `commit` to a value.¹⁶ Furthermore, if the programmer omitted this constraint or instead chose $\perp^\leftarrow$, say, then `commit` would be rejected by the type system. Specifically, the `assume` term would not type-check via rule `ASSUME` since the pc does not act for $\nabla(p^\leftarrow)$, which is equal to $p^\leftarrow$.

When p is ready to reveal a previously committed value, it instantiates the function `reveal` with the integer type N of the committed value and sends the result to q . The body of `reveal` creates delegations $\langle \nabla(q^\rightarrow) \geq \nabla(p^\rightarrow) \rangle$ and $\langle q^\rightarrow \geq p^\rightarrow \rangle$, permitting the inner function to relabel its argument from p to $q^\rightarrow \wedge p^\leftarrow$. The outer `assume` term establishes that principals speaking for $q^\rightarrow$ also speak for $p^\rightarrow$ by creating an integrity relationship between their voices. With this relationship in place, the inner `assume` term may delegate p 's confidentiality to q .¹⁷

Only principals trusted by $\nabla(p^\rightarrow)$ and $p^\leftarrow$ may instantiate `reveal` with the integer N of the commitment, therefore q cannot invoke `reveal` until it is authorized by p . If, for example, the type function had pc annotation q instead of $\nabla(p^\rightarrow)$, it would be rejected by the type system since the `assume` term would not check under the `ASSUME` rule. Note however, that once the delegation is established, the inner function may be invoked by q since it has the pc annotation $q^\leftarrow$. Without the delegation, however, the `bind` term would fail to type-check under `BINDM` since p would not flow to $q^\rightarrow \wedge p^\leftarrow$.

Given a function of the type instantiated by `reveal` and a committed value of type p says X , the open function may be invoked by q to relabel the committed value to a type observable by q . Because `open` may only be invoked by principals trusted by $q^\leftarrow$, p cannot influence which value of type p says (N, X) `open` is applied to. If q only accepts a one value for each integer type N , then the argument must be the one originally committed to by p . Furthermore, since the `reveal` function provided to p is parametric with respect to X , it cannot return a value other than (a relabeled version of) its argument. The return type $q^\leftarrow$ says $q^\rightarrow \wedge p^\leftarrow$ says X protects this relationship between the committed value and the opened one: it indicates that q trusts that $q^\rightarrow \wedge p^\leftarrow$ says X is the one originally committed to by p .

In DCC, functions are not annotated with pc labels and may be applied in any context. So a DCC function analogous to `reveal` might have type

$$\text{dcc_reveal} : \forall N. \forall X. p \text{ says } (N, X) \rightarrow q^\rightarrow \wedge p^\leftarrow \text{ says } (N, X)$$

However, `dcc_reveal` would not be appropriate for a commitment scheme since any principal could use it to relabel information from p -confidential ($p^\rightarrow$) to q -confidential ($q^\rightarrow$).

The real power of FLAC is that the security guarantees of well-typed FLAC functions like those above are compositional. The FLAC type system ensures the security of both the functions themselves and the programs that use them. For instance, the following code should be rejected because it would permit q to

¹⁵Any reasonable integer encoding is permissible provided that each integer N is a singleton type.

¹⁶We make the reasonable assumption that an untrusted programmer cannot modify high-integrity code, thus the influence of attackers is captured by the pc and the protection levels of values. Enforcing this assumption is beyond the scope of FLAC, but has been explored in [6].

¹⁷Specifically, it satisfies the `ASSUME` premise $\Pi \Vdash \nabla(p^\rightarrow) \geq \nabla(q^\rightarrow)$.

reveal p 's commitments.

$$\Lambda N[q^{\leftarrow}]. \Lambda X[q^{\leftarrow}]. \lambda(x:p \wedge q^{\leftarrow} \text{ says } X)[q^{\leftarrow}]. \text{assume } \langle q \geq p \rangle \text{ in reveal } N \ X \ x$$

The pc constraints on this function all have the integrity of q , but the body of the function uses `assume` to create a new delegation from p to q . If this `assume` was permitted by the type system, then q could use the function to reveal p 's committed values. Since `ASSUME` requires the pc to speak for p , this function is rejected.

FLAC's guarantees make it possible to state general security properties of all programs that use the above commitment scheme, even if those programs are malicious. For example, suppose we have $pc_p = \nabla(p)$, $pc_q = \nabla(q)$, and

$$\Gamma_{cro} = \text{commit, reveal, open, (reveal } N'), x:p^{\rightarrow} \text{ says } \tau, y:p \text{ says } (N, \tau), z:p \text{ says } (N', \tau)$$

Intuitively, pc_p and pc_q are execution contexts under the control of p or q , respectively. Γ_{cro} is a typing context for programs using the commitment scheme. For presentation purposes, we have omitted the types of `commit`, `reveal`, and `open` in Γ_{cro} . Their types are as defined previously. The variable x represents an uncommitted value with p 's confidentiality, whereas y and z are a committed value. The type $(\text{reveal } N')$ is shorthand for `reveal` instantiated with integer type N' and represents that z has been revealed to q by p .

Since we are interested in properties that hold for all principals p and q , we want the properties to hold in an empty delegation context: $\Pi = \emptyset$. Below, we omit the delegation context altogether for brevity.

FLAC's noninterference guarantee helps rule out information that an attacker can influence or learn. Using results presented in Section 7, we can prove that:

- **q cannot open a value that hasn't been committed.** For any e and τ' such that $\Gamma_{cro}; pc_q \vdash e : q^{\rightarrow} \wedge p^{\leftarrow} \text{ says } \tau'$, the result of evaluating e is independent of x ; specifically, for any v_1 and v_2 , if $e[x \mapsto v_1] \longrightarrow^* v'_1$ and $e[x \mapsto v_2] \longrightarrow^* v'_2$, then $v'_1 = v'_2$.
- **q cannot learn a value that hasn't been revealed.** For any e , ℓ , and τ' such that $\Gamma_{cro}; pc_q \vdash e : \ell \sqcap q^{\rightarrow} \text{ says } \tau'$, then the result of e is independent of x and y .

For these properties we consider programs using our commitment scheme that q might invoke, hence we consider FLAC programs that type-check in the $\Gamma_{cro}; pc_q$ context. In the first property, we are concerned with programs that produce values protected by policy p . Since such programs produce values with the integrity of p but are invoked by q , we want to ensure that no program exists that enables q to obtain a value with p 's integrity that depends on x , which is a value without p 's integrity. The second property concerns programs that produces values at $\ell \sqcap q^{\rightarrow}$ for any ℓ ; these are values readable by q . Therefore, we want to ensure that no program exists that enables q to produce such a value that depends on x or y , which are not readable by q .

Each of these properties hold by a straightforward application of our noninterference theorems (Theorem 2). This result is strengthened by our robust declassification theorem (Theorem 1), which ensures that attacks by q on p 's programs cannot subvert the intended declassifications.

5.2. Bearer Credentials

We can also use FLAC to implement bearer credentials, our second example of a dynamic authorization mechanism. We represent a bearer credential with authority k in FLAC as a term with the type

$$\forall X[pc]. k^{\rightarrow} \text{ says } X \xrightarrow{pc} k^{\leftarrow} \text{ says } X$$

which we abbreviate as $k^{\rightarrow} \xRightarrow{pc} k^{\leftarrow}$. These terms act as bearer credentials for a principal k since they may be used as a proxy for k 's confidentiality and integrity authority. Recall that $k^{\leftarrow} = k^{\leftarrow} \wedge \perp^{\rightarrow}$ and $k^{\rightarrow} = k^{\rightarrow} \wedge \perp^{\leftarrow}$. Then secrets protected by $k^{\rightarrow}$ can be declassified to $\perp^{\rightarrow}$, and untrusted data protected by $\perp^{\leftarrow}$ can be endorsed to $k^{\leftarrow}$. Thus this term wields the full authority of k , and if $pc = \perp^{\leftarrow}$, the credential may be used in any context—any “bearer” may use it. From such credentials, more restricted credentials can be derived. For example, the credential $k^{\rightarrow} \xRightarrow{pc} \perp^{\rightarrow}$ grants the bearer authority to declassify k -confidential values, but no authority to endorse values.

We postpone an in-depth discussion of terms with types of the form $k^{\rightarrow} \xRightarrow{pc} k^{\leftarrow}$ until Section 6.2, but it is interesting to note that an analogous term in DCC is only well-typed if k is equivalent to $\perp$. This is because the function takes an argument with $k^{\rightarrow}$ confidentiality and no integrity, and produces a value with $k^{\leftarrow}$ integrity and no confidentiality. Suppose $\mathcal{L}$ is a security lattice used to type-check DCC programs with suitable encodings for k 's confidentiality and integrity. If a DCC term has a type analogous to $k^{\rightarrow} \Rightarrow k^{\leftarrow}$, then $\mathcal{L}$ must have the property $k^{\rightarrow} \sqsubseteq \perp$ and $\perp \sqsubseteq k^{\leftarrow}$. This means that k has no confidentiality and no integrity. That FLAC terms may have this type for any principal k makes it straightforward to implement bearer credentials and demonstrates a useful application of FLAC's extra expressiveness.

The pc of a credential $k^{\rightarrow} \xRightarrow{pc} k^{\leftarrow}$ acts as a sort of caveat: it restricts the information flow context in which the credential may be used. We can add more general caveats to credentials by wrapping them in lambda terms. To add a caveat ϕ to a credential with type $k^{\rightarrow} \xRightarrow{pc} k^{\leftarrow}$, we use a wrapper:

$$\lambda(x: k^{\rightarrow} \xRightarrow{pc} k^{\leftarrow})[pc]. \Lambda X[pc]. \lambda(y: \phi)[pc]. xX$$

which gives us a term with type

$$\forall X[pc]. \phi \xrightarrow{pc} k^{\rightarrow} \text{ says } X \xrightarrow{pc} k^{\leftarrow} \text{ says } X$$

This requires a term with type ϕ (in which X may occur) to be applied before the authority of k can be used. Similar wrappers allow us to chain multiple caveats; i.e., for caveats $\phi_1 \dots \phi_n$, we obtain the type

$$\forall X[pc]. \phi_1 \xrightarrow{pc} \dots \xrightarrow{pc} \phi_n \xrightarrow{pc} k^{\rightarrow} \text{ says } X \xrightarrow{pc} k^{\leftarrow} \text{ says } X$$

which abbreviates to

$$k^{\rightarrow} \xRightarrow{\phi_1 \times \dots \times \phi_n; pc} k^{\leftarrow}$$

We will also use the syntax $(k \xrightarrow{\phi_1 \times \dots \times \phi_n; pc} k^\leftarrow) \tau$ (suggesting a type-level application) as an abbreviation for

$$\phi_1[X \mapsto \tau] \xrightarrow{pc} \dots \xrightarrow{pc} \phi_n[X \mapsto \tau] \xrightarrow{pc} k^\rightarrow \text{ says } \tau \xrightarrow{pc} k^\leftarrow \text{ says } \tau$$

Like any other FLAC terms, credentials may be protected by information flow policies. So a credential that should only be accessible to Alice might be protected by the type $\text{Alice}^\rightarrow \text{ says } (k^\rightarrow \xrightarrow{\phi; pc} k^\leftarrow)$. This confidentiality policy ensures the credential cannot accidentally be leaked to an attacker. A further step might be to constrain uses of this credential so that only Alice may invoke it to relabel information. If we require $pc = \text{Alice}^\leftarrow$, this credential may only be used in contexts trusted by Alice: $\text{Alice}^\rightarrow \text{ says } (k^\rightarrow \xrightarrow{\phi; \text{Alice}^\leftarrow} k^\leftarrow)$.

A subtle point about the way in which we construct caveats is that the caveats are polymorphic with respect to X , the same type variable the credential ranges over. This means that each caveat may constrain what types X may be instantiated with. For instance, suppose isEduc is a predicate for educational films; it holds (has a proof term with type $\text{isEduc } X$) for types like Bio and Doc , but not RomCom . Adding $\text{isEduc } X$ as a caveat to a credential would mean that the bearer of the credential could use it to access biographies and documentaries, but could not use it to access romantic comedies. Since no term of type $\text{isEduc } \text{RomCom}$ could be applied, the bearer could only satisfy isEduc by instantiating X with Bio or Doc . Once X is instantiated with Bio or Doc , the credential cannot be used on a RomCom value. Thus we have two mechanisms for constraining the use of credentials: information flow policies to constrain propagation, and caveats to establish prerequisites and constrain the types of data covered by the credential.

As a more in-depth example of using such credentials, suppose Alice hosts a file sharing service. For a simpler presentation, we use free variables to refer to these files; for instance, $x_1 : (k_1 \text{ says ph})$ is a variable that stores a photo (type ph) protected by k_1 . For each such variable x_1 , Alice has a credential $k_1^\rightarrow \xrightarrow{\perp^\leftarrow} k_1^\leftarrow$, and can give access to users by providing this credential or deriving a more restricted one. To access x_1 , Bob does not need the full authority of Alice or k_1 —a more restricted credential suffices. Alice can provide Bob with a credential c of type $(k_1 \xrightarrow{\text{Bob}^\leftarrow} k_1^\leftarrow) \text{ ph}$. By applying this credential to x_1 , Bob is able to access the result of type $k_1^\leftarrow \text{ says ph}$ since its confidentiality is now public.

This example demonstrates an advantage of bearer credentials: access to x_1 can be provided to principals other than k_1 in a decentralized way, without changing the policy on x_1 . Suppose Alice has a credential with type $k_1^\rightarrow \xrightarrow{\perp^\leftarrow} k_1^\leftarrow$ and wants to issue the above credential to Bob. Alice can create such a credential using a wrapper that derives the more constrained credential from her original one.

$$\begin{aligned} & \lambda(c : k_1^\rightarrow \xrightarrow{\perp^\leftarrow} k_1^\leftarrow) [\text{Alice}^\leftarrow]. \\ & \lambda(y : k_1 \text{ says ph}) [\text{Bob}^\leftarrow]. \\ & \text{bind } y' = y \text{ in } (c \text{ ph}) (\eta_{k^\rightarrow} y') \end{aligned}$$

This function has type $(k_1^{\rightarrow} \xRightarrow{\perp^{\leftarrow}} k_1^{\leftarrow}) \xrightarrow{\text{Alice}^{\leftarrow}} (k_1 \xRightarrow{\text{Bob}^{\leftarrow}} k_1^{\leftarrow})$ ph: given her root credential $k_1^{\rightarrow} \xRightarrow{\perp^{\leftarrow}} k_1^{\leftarrow}$, Alice (or someone she trusts) can create a restricted credential that allows Bob (or someone he trusts) to access values of type ph protected under k_1 .

Bob can also use this credential to share photos with friends. For instance, the function

$$\begin{aligned} & \lambda(c : (k_1 \xRightarrow{\text{Bob}^{\leftarrow}} k_1^{\leftarrow}) \text{ ph}) [\text{Bob}^{\leftarrow}]. \\ & \text{assume } \langle \text{Carol}^{\leftarrow} \geq \text{Bob}^{\leftarrow} \rangle \text{ in} \\ & \lambda(y : k_1 \text{ says ph}) [\text{Carol}^{\leftarrow}]. \\ & \text{bind } y' = y \text{ in } (c \text{ ph}) (\eta_{k \rightarrow} y') \end{aligned}$$

creates a wrapper around a Bob's credential that is invocable by Carol. Using the assume term, Bob delegates authority to Carol so that his credential may be invoked at pc Carol[←].

The properties of FLAC let us prove many general properties about such bearer-credential programs; here, we examine three properties. For $i \in \{1..n\}$, let

$$\Gamma_{bc} = x_i : k_i \text{ says } \tau_i, c_i : \text{Alice says } (k_i^{\leftarrow} \xRightarrow{\perp^{\leftarrow}} k_i^{\leftarrow})$$

where k_i is a primitive principal protecting the i^{th} resource of type τ_i , and c_i is a credential for the i^{th} resource and protected by Alice. Assume $k_i \notin \{\text{Alice}, \text{Friends}, p\}$ for all i where p represents a (potentially malicious) user of Alice's service, and Friends is a principal for Alice's friends, (e.g., Friends = (Bob $\vee$ Carol)). Also, define $pc_p = p^{\leftarrow}$ and $pc_A = \text{Alice}^{\leftarrow}$.

- **p cannot access resources without a credential.** For any e , ℓ , and τ' such that $\Gamma_{bc}; pc_p \vdash e : \ell \sqcap p^{\rightarrow} \text{ says } \tau'$, the value of e is independent of x_i for all i .
- **p cannot use unrelated credentials to access resources.** For any e , ℓ , and τ' such that

$$\Gamma_{bc}, c_p : (k_1^{\leftarrow} \xRightarrow{\perp^{\leftarrow}} k_1^{\leftarrow}); pc_p \vdash e : \ell \sqcap p^{\rightarrow} \text{ says } \tau'$$

the value e computes is independent of x_i for $i \neq 1$.

- **Alice cannot disclose secrets by issuing credentials.** For all i and $j \neq 1$, define

$$\Gamma'_{bc} = x_i : k_i \text{ says } \tau_i, c_i : \text{Alice says } (k_j^{\leftarrow} \xRightarrow{\perp^{\leftarrow}} k_j^{\leftarrow}),$$

$$c_F : \text{Friends says } (k_1^{\leftarrow} \xRightarrow{\perp^{\leftarrow}} k_1^{\leftarrow})$$

Then if $\Gamma'_{bc}; pc_A \vdash e : \ell \sqcap p^{\rightarrow} \text{ says } (k_j^{\leftarrow} \xRightarrow{\perp^{\leftarrow}} k_j^{\leftarrow})$ for some e , ℓ , and τ' , the value of e is independent of x_1 .

These properties demonstrate the power of FLAC's type system. The first two ensure that credentials really are necessary for p to access protected resources, even indirectly. In the first, p has no credentials, and the type system ensures that p cannot invoke a program that produces a value p can read (represented

by $\ell \sqcap p^{\rightarrow}$) that depends on any variable x_i . In the second, a credential c_p with type $k_1^{\leftarrow} \xrightarrow{\perp^{\leftarrow}} k_1^{\leftarrow}$ is accessible to p , but p cannot use it to access other variables. The third property eliminates covert channels like the one discussed in Section 2.2. It implies that credentials issued by Alice do not leak information, in this case about Alice's friends. By implementing bearer credentials in FLAC, we can demonstrate these three properties with relatively little effort.

6. FLAC Proof theory

6.1. Properties of says

FLAC's type system constrains how principals apply existing facts to derive new facts. For instance, a property of says in other authorization logics (e.g., Lampson et al. [27] and Abadi [2]) is that implications that hold for top-level propositions also hold for propositions of any principal ℓ :

$$\vdash (\tau_1 \rightarrow \tau_2) \rightarrow (\ell \text{ says } \tau_1 \rightarrow \ell \text{ says } \tau_2)$$

The pc annotations on FLAC function types refine this property. Each implication (in other words, each function) in FLAC is annotated with an upper bound on the information flow context it may be invoked within. To lift such an implication to operate on propositions protected at label ℓ , the label ℓ must flow to the pc of the implication. Thus, for all ℓ and τ_i ,

$$\vdash (\tau_1 \xrightarrow{\ell} \tau_2) \xrightarrow{\ell} (\ell \text{ says } \tau_1 \xrightarrow{\ell} \ell \text{ says } \tau_2)$$

This judgment is a FLAC typing judgment in *logical form*, where terms have been omitted. We write such judgments with an empty typing context (as above) when the judgment is valid for any Π , Γ , and pc . A judgment in logical form is valid if a *proof term* exists for the specified type, proving the type is inhabited. The above type has proof term

$$\lambda(f : (\tau_1 \xrightarrow{\ell} \tau_2))[\ell]. \\ \lambda(x : \ell \text{ says } \tau_1)[\ell]. \text{bind } x' = x \text{ in } \eta_{\ell} f x'$$

In order to apply f , we must first bind x , so according to rules BINDM and APP, the function f must have a label at least as restrictive as ℓ , and UNITM requires the label of the returned value must also be as restrictive as ℓ . We can actually prove a slightly more general version of the above theorem:

$$(\tau_1 \xrightarrow{pc \sqcup \ell} \tau_2) \xrightarrow{\ell} (\ell \text{ says } \tau_1 \xrightarrow{pc} pc \sqcup \ell \text{ says } \tau_2)$$

This version permits using the implications in more restrictive contexts, but doesn't map as well to a DCC theorem since the principal of the return type differs from the argument type.

These refinements of DCC's theorems are crucial for supporting applications like commitment schemes and bearer credentials. Recall from Sections 5.1 and 5.2 that the security of these mechanisms relied in part on restricting the pc to a specific principal's integrity. Without such refinements, principal q could open principal p 's commitments using open, or create credentials with p authority: $p^{\rightarrow} \xrightarrow{pc} p^{\leftarrow}$.

Other properties of says common to DCC and other logics (cf. [1] for examples) are similarly refined by pc bounds. Two examples are: $\vdash \tau \xrightarrow{\ell} \ell \text{ says } \tau$ which has proof term: $\lambda(x:\tau)[\ell]. \eta_{\ell} \tau$ and

$$\vdash \ell \text{ says } (\tau_1 \xrightarrow{\ell} \tau_2) \xrightarrow{\ell} (\ell \text{ says } \tau_1 \xrightarrow{\ell} \ell \text{ says } \tau_2)$$

with proof term:

$$\begin{aligned} & \lambda(f:\ell \text{ says } (\tau_1 \xrightarrow{\ell} \tau_2))[\ell]. \text{bind } f' = f \text{ in} \\ & \lambda(y:\ell \text{ says } \tau_1)[\ell]. \text{bind } y' = y \text{ in } \eta_{\ell} f' y' \end{aligned}$$

Some theorems of DCC cannot be obtained in FLAC, due to the pc restriction on UNITM as well as the more restrictive protection relation. For example, chains of says are not commutative in FLAC. Given ℓ_1, ℓ_2 , and pc ,

$$\not\vdash \ell_1 \text{ says } \ell_2 \text{ says } \tau \xrightarrow{pc} \ell_2 \text{ says } \ell_1 \text{ says } \tau$$

unless $\Pi \Vdash \ell_1 \sqcup pc \sqsubseteq \ell_2$ and $\Pi \Vdash \ell_2 \sqcup pc \sqsubseteq \ell_1$, which implies ℓ_1, ℓ_2 , and pc must be equivalent in Π . CCD [3], a logic related to DCC, is also non-commutative with respect to says, but does not have an associated term language.

Preserving the nesting order of says types is attractive for authorization settings since it encodes the provenance of statements. It also enables modeling of cryptographic mechanisms in FLAC (cf. [?]) where $\ell \text{ says } \tau$ is interpreted as a value of type τ protected by encryption key $\ell^{\rightarrow}$ and signing key $\ell^{\leftarrow}$. Preserving the order of nested types in this context is critical for modeling decryption and verification of protected values.

6.2. Dynamic Hand-off

Many authorization logics support delegation using a “hand-off” axiom. In DCC, this axiom is actually a provable theorem:

$$\vdash (q \text{ says } (p \Rightarrow q)) \rightarrow (p \Rightarrow q)$$

where $p \Rightarrow q$ is shorthand for

$$\forall X. (p \text{ says } X \rightarrow q \text{ says } X)$$

However, $p \Rightarrow q$ is only inhabited if $p \sqsubseteq q$ in the security lattice. Thus, DCC can reason about the consequences of $p \sqsubseteq q$ (whether it is true for the lattice or not), but a DCC program cannot produce a term of type $p \Rightarrow q$ unless $p \sqsubseteq q$.

FLAC programs, on the other hand, can create new trust relationships from delegation expressions using assume terms. The type analogous to $p \Rightarrow q$ in FLAC is

$$\forall X[pc]. (p \text{ says } X \xrightarrow{pc} q \text{ says } X)$$

which we wrote as $p \xRightarrow{pc} q$ in Section 5.2. FLAC programs construct terms of this type from proofs of authority, represented by terms with acts-for types. This feature enables a more general form of hand-off, which we state formally below.

Proposition 1 (Dynamic hand-off). *For all ℓ and pc' , let $pc = \ell^{\rightarrow} \wedge \nabla(p^{\rightarrow}) \wedge q^{\leftarrow} \wedge \nabla(pc')$*

$$\begin{aligned} &(\nabla(q^{\rightarrow}) \geq \nabla(p^{\rightarrow})) \xrightarrow{pc} (p \sqsubseteq q) \xrightarrow{pc} (pc' \sqsubseteq q) \xrightarrow{pc} \\ &\forall X[pc']. (p \text{ says } X \xrightarrow{pc'} q \text{ says } X) \end{aligned}$$

Proof term.

$$\begin{aligned} &\lambda(pf_1 : (\nabla(q^{\rightarrow}) \geq \nabla(p^{\rightarrow}))[pc]. \text{assume } pf_1 \text{ in} \\ &\lambda(pf_2 : (p \sqsubseteq q))[pc]. \text{assume } pf_2 \text{ in} \\ &\lambda(pf_3 : (pc' \sqsubseteq q))[pc]. \text{assume } pf_3 \text{ in} \\ &\Lambda X[pc']. \lambda(x : p \text{ says } X)[pc']. \text{bind } x' = x \text{ in } \eta_q x' \end{aligned}$$

The principal $pc = \ell^{\rightarrow} \wedge \nabla(p^{\rightarrow}) \wedge q^{\leftarrow} \wedge \nabla(pc')$ restricts delegation (hand-off) to contexts with sufficient integrity to authorize the delegations made by the assume terms. In other words, the context that creates these delegations must be authorized by the combined authority of $\nabla(p^{\rightarrow})$, $q^{\leftarrow}$, and $\nabla(pc')$.

The three arguments are proofs of authority with acts-for types: a proof of $\nabla(q^{\rightarrow}) \geq \nabla(p^{\rightarrow})$, a proof of $p \sqsubseteq q$, and a proof of $pc' \sqsubseteq q$. The pc ensures that the proofs have sufficient integrity to be used in assume terms since it has the integrity of both $\nabla(p^{\rightarrow})$ and $q^{\leftarrow}$. Note that low-integrity or confidential delegation values must first be bound via bind before the above term may be applied. Thus the pc would reflect the protection level of both arguments. Principals $\ell^{\rightarrow}$ and pc' are unconstrained, but the third proof argument ensures that flows from pc' to q are authorized, since a principal with access to q 's secrets could infer something about the context (protected at pc') in which the hand-off function is applied. Other dynamic hand-off formulations are possible, Proposition 1 simply has the fewest assumptions. Other formulations can eliminate the need for proofs pf_1 , pf_2 , and/or pf_3 if these relationships already exist in the context defining the hand-off.

Dynamic hand-off terms give FLAC programs a level of expressiveness and security not offered by other authorization logics. Observe that pc' may be chosen independently of the other principals. This means that although the pc prevents low-integrity principals from creating hand-off terms, a high-integrity principal may create a hand-off term and provide it to an arbitrary principal. Hand-off terms in FLAC, then, are similar to capabilities since even untrusted principals may use them to change the protection level of values. Unlike in most capability systems, however, the propagation of hand-off terms can be constrained using information flow policies.

Terms that have types of the form in Proposition 1 illustrate a subtlety of enforcing information flow in an authorization mechanism. Because these terms relabel information from one protection level to another protection level, the transformed information implicitly depends on the proofs of authorization. FLAC ensures that the information security of these proofs is protected—like that of all other values—even as the policies of other information are being modified. Hence, authorization proofs cannot be used as a side channel to leak information.

Syntax

$$e ::= \dots \mid \langle e \rangle_\ell$$
Evaluation contexts

$$E ::= \dots \mid \langle E \rangle_\ell$$
Evaluation rules

$$[\text{E-APP}^*] \quad (\lambda(x:\tau)[pc]. e) w \longrightarrow \langle e[x \mapsto w] \rangle_{pc} \quad [\text{E-TAPP}^*] \quad (\Lambda X[pc]. e) \tau \longrightarrow \langle e[X \mapsto \tau] \rangle_{pc}$$

$$[\text{E-BINDM}^*] \quad \text{bind } x = \bar{\eta}_\ell w \text{ in } e \longrightarrow \langle e[x \mapsto w] \rangle_\ell \quad [\text{O-CTX}] \quad \langle w \rangle_\ell \longrightarrow w$$
Typing rules

$$[\text{CTX}] \quad \frac{\Pi; \Gamma; pc' \vdash e : \tau \quad \Pi \Vdash pc \sqsubseteq pc'}{\Pi; \Gamma; pc \vdash \langle e \rangle_{pc'} : \tau}$$

Fig. 9. Extensions to support protection contexts.

7. Semantic security properties of FLAC*7.1. Trace indistinguishability*

We express our semantic results in terms of the *traces* of a program observable to an attacker. FLAC traces are simply the sequence of terms under the $\longrightarrow$ relation. That is, each evaluation step of the form $e \longrightarrow e'$ generates a new trace element e' . We write the trace generated by taking n steps from e to e' as $e \xrightarrow{t}^* e'$, where $t[0] = e$ and $t[n] = e'$.

FLAC traces are not fully observable to an attacker. Trace elements generated by protected information such as sealed values or protected contexts such as within a bind or lambda term are hidden from the attacker. This approach models the scenario where sealed values are implemented by encrypting and signing data, and computation may only be done on these values by hosts with access to the necessary cryptographic keys.¹⁸ An insecure program in this model allows an attacker to learn information by allowing protected information to flow to sealed values or contexts that are observable by the attacker.

Before formally defining what portions of program trace are available to an attack, we must first add some additional bookkeeping to our semantic rules. Specifically, we need additional notation for evaluating some intermediate expressions. The notation $\langle e \rangle_\ell$ denotes an intermediate expression e evaluated in a context protected at ℓ . We extend our grammar in Figure 2 with this notation as below:

Figure 9 presents rules for introducing and eliminating protected contexts. Rules E-APP*, E-TAPP*, and E-BINDM* replace their counterparts in Figure 3 with rules that introduce a protected context based

¹⁸DFLATE [?] makes this connection with cryptographic enforcement more explicit for a distributed extension of FLAC.

$$\begin{aligned}
\mathcal{O}(x, \Pi, p, \pi) &= x \\
\mathcal{O}(\langle a \geq b \rangle, \Pi, p, \pi) &= \circ \\
\mathcal{O}(\langle \rangle, \Pi, p, \pi) &= \circ \\
\mathcal{O}(\eta_\ell e, \Pi, p, \pi) &= \eta_\ell \mathcal{O}(e, \Pi, p, \pi) \\
\mathcal{O}(\bar{\eta}_\ell w, \Pi, p, \pi) &= \begin{cases} \bar{\eta}_\ell \mathcal{O}(w, \Pi, p, \pi) & \Pi \Vdash p \geq \ell^\pi \\ \circ & \text{otherwise} \end{cases} \\
\mathcal{O}(\lambda(x:\tau)[pc].e, \Pi, p, \pi) &= \begin{cases} \lambda(x:\tau)[pc].\mathcal{O}(e, \Pi, p, \pi) & \Pi \Vdash p \geq pc^\pi \\ \circ & \text{otherwise} \end{cases} \\
\mathcal{O}(\Lambda X[pc].e, \Pi, p, \pi) &= \begin{cases} \Lambda X[pc].\mathcal{O}(e, \Pi, p, \pi) & \Pi \Vdash p \geq pc^\pi \\ \circ & \text{otherwise} \end{cases} \\
\mathcal{O}(\llbracket w \rrbracket_\ell, \Pi, p, \pi) &= \begin{cases} \llbracket \mathcal{O}(w, \Pi, p, \pi) \rrbracket_\ell & \Pi \Vdash p \geq \ell^\pi \\ \circ & \text{otherwise} \end{cases} \\
\mathcal{O}(e e', \Pi, p, \pi) &= \mathcal{O}(e, \Pi, p, \pi) \mathcal{O}(e', \Pi, p, \pi) \\
\mathcal{O}(\langle e_1, e_2 \rangle, \Pi, p, \pi) &= \begin{cases} \circ & \mathcal{O}(e_i, \Pi, p, \pi) = \circ \\ \langle \mathcal{O}(e_1, \Pi, p, \pi), \mathcal{O}(e_2, \Pi, p, \pi) \rangle & \text{otherwise} \end{cases} \\
\mathcal{O}(\text{proj}_i e, \Pi, p, \pi) &= \text{proj}_i \mathcal{O}(e, \Pi, p, \pi) \\
\mathcal{O}(\text{inj}_i e, \Pi, p, \pi) &= \begin{cases} \circ & \mathcal{O}(e, \Pi, p, \pi) = \circ \\ \text{inj}_i \mathcal{O}(e, \Pi, p, \pi) & \text{otherwise} \end{cases} \\
\mathcal{O}(\text{case } e \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2, \Pi, p, \pi) &= \text{case } \mathcal{O}(e, \Pi, p, \pi) \text{ of} \\
&\quad \text{inj}_1(x). \mathcal{O}(e_1, \Pi, p, \pi) \\
&\quad \mid \text{inj}_2(x). \mathcal{O}(e_2, \Pi, p, \pi) \\
\mathcal{O}(\text{bind } x = e \text{ in } e', \Pi, p, \pi) &= \text{bind } x = \mathcal{O}(e, \Pi, p, \pi) \text{ in } \mathcal{O}(e', \Pi, p, \pi) \\
\mathcal{O}(\text{assume } e \text{ in } e', \Pi, p, \pi) &= \begin{cases} \circ & \mathcal{O}(e, \Pi, p, \pi) = \circ \text{ and } \mathcal{O}(e', \Pi, p, \pi) = \circ \\ \text{assume } \mathcal{O}(e, \Pi, p, \pi) \text{ in } \mathcal{O}(e', \Pi, p, \pi) & \text{otherwise} \end{cases} \\
\mathcal{O}(e \text{ where } v, \Pi, p, \pi) &= \begin{cases} \circ & \mathcal{O}(e, \Pi \cdot v, p, \pi) = \circ \\ \mathcal{O}(e, \Pi \cdot v, p, \pi) \text{ where } v & \text{otherwise} \end{cases}
\end{aligned}$$

(a) Observation function for intermediate FLAC terms.

$$\begin{aligned}
\mathcal{O}([e], \Pi, p, \pi) &= [\mathcal{O}(e, \Pi, p, \pi)] \\
\mathcal{O}([e; e'], \Pi, p, \pi) &= \begin{cases} [\mathcal{O}(e, \Pi, p, \pi)] & \mathcal{O}(e, \Pi, p, \pi) = \mathcal{O}(e', \Pi, p, \pi) \\ [\mathcal{O}(e, \Pi, p, \pi); \mathcal{O}(e', \Pi, p, \pi)] & \text{otherwise} \end{cases} \\
\mathcal{O}([e; e'] \cdot t, \Pi, p, \pi) &= \mathcal{O}(\mathcal{O}([e; e'], \Pi, p, \pi) \cdot t, \Pi, p, \pi)
\end{aligned}$$

(b) Observation function for traces.

Fig. 10. Observation function definitions

on the relevant label. Rule O-CTX eliminates the protected context when the intermediate expression is fully evaluated.

The FLAC programs observable to an attacker are defined by an observation function $\mathcal{O}(e, \Pi, p, \pi)\pi$, defined in Figure 10a. An expression e is observable by p in delegation context Π depending on the authority of p relative to the protected terms in e . The projection π specifies whether to consider the confidentiality of protected terms ($\pi = \rightarrow$) or the integrity ($\pi = \leftarrow$). For example, sealed values such as $\bar{\eta}_\ell w$ are observable by principal p if p acts for ℓ , otherwise they are erased. Protected contexts like $\langle e \rangle_\ell$ are treated similarly. To simplify our proofs, we collapse terms whose subterms have been erased. For example $\langle \circ, \circ \rangle$ is collapsed to $\circ$. For values with singleton types like $()$ and $\langle a \geq b \rangle$ we erase the value since no information is learned from the value itself.

Erasing delegations in the context of where terms essentially hides the delegations that justify a value's flow from the attacker. This is consistent with the idea that attackers may learn *implicit* information from flows that introduce new delegations. In other words, by observing differences in outputs the attacker may infer the existence of different delegations, but they are not *explicit* in the output. In DFLATE [?], these delegations are explicit in the output, modeling values that carry certified justifications of why their flow was authorized by the program.

The observability of trace elements is defined in Figure 10b in terms of the observability function for terms. Note that duplicate entries (which may occur due to evaluation in protected contexts) are removed.

We now can define trace indistinguishability. Two traces t and t' are indistinguishable to a principal p^π in delegation context Π if the observable elements of t and t' are equal.

Definition 1 (Trace indistinguishability).

$$t \approx_{p^\pi}^\Pi t' \iff \mathcal{O}(t, \Pi, p, \pi) = \mathcal{O}(t', \Pi, p, \pi)$$

7.2. Robust declassification

Robust declassification [45] requires disclosures of secret information to be independent of low-integrity information. Robust declassification permits some confidential information to be disclosed to an attacker, but attackers can influence neither the decision to disclose information nor the choice of what information is disclosed. Therefore, robust declassification is a more appropriate security condition than noninterference when programs are intended to disclose information.

Following Myers *et al.* [33], we extend our set of terms with a “hole” term $[\bullet]$ representing portions of a program that are under the control of an attacker. We extend the type system with the following rule for holes:

$$[\text{HOLE}] \quad \frac{\Pi \Vdash H^\leftarrow \geq pc^\leftarrow \quad \Pi \Vdash pc^\rightarrow \sqsubseteq \Delta(H^\leftarrow)}{\Pi; \Gamma; pc \vdash \bullet : \tau}$$

Where $\Delta(H^\leftarrow)$ is the *view* of principal $H^\leftarrow$. The view of a principal is a dual notion to voice, introduced by Cecchetti *et al.* [15] to represent the confidentiality observable to a principal. Given a principal in normal form $q^\rightarrow \wedge r^\leftarrow$, the view of that principal is

$$\Delta(q^\rightarrow \wedge r^\leftarrow) \triangleq q^\rightarrow \wedge r^\rightarrow$$

In other words $\Delta(H^\leftarrow)$ represents the confidentiality of information observable to the attacker $H^\leftarrow$. Therefore, the HOLE premises $\Pi \Vdash H^\leftarrow \geq pc^\leftarrow$ and $\Pi \Vdash pc^\rightarrow \sqsubseteq \Delta(H^\leftarrow)$ require that holes be inserted only in contexts that are controlled by and observable to the attacker.

We write $e[\bullet]$ to denote a program e with holes. Let an *attack* be a vector $\vec{a}$ of terms and $e[\vec{a}]$ be the program where a_i is substituted for $\bullet_i$.

Definition 2 (Π -fair attacks). An attack $\vec{a}$ is a Π -fair attack on a well-typed program with holes $\Pi; \Gamma; pc \vdash e[\bullet] : \tau$ if the program $e[\vec{a}]$ is also well typed (thus $\Pi; \Gamma; pc \vdash e[\vec{a}] : \tau$), and furthermore, for each $a_i \in \vec{a}$, we have $\Pi; \Gamma'; pc \vdash [a_i] : \tau'$ where each entry in Γ' has the form $y : \ell$ says τ'' and $\forall y : \ell$ says $\tau'' \in \Gamma' . \Pi \Vdash \ell^\rightarrow \sqsubseteq \Delta(H^\leftarrow)$

By specifying the relationships between confidential information labeled $H^\rightarrow$ and the integrity of the attacker $H^\leftarrow$, we can precisely express the authority an attacker is able to wield in its Π -fair attacks against protected information. Proposition 2 states the conditions under which a Π -fair attack is noninterfering. Condition 1 specifies that τ' contains confidential information (labeled $H^\rightarrow$), Condition 2 specifies that this information is not visible to the attacker (a principal with integrity $H^\leftarrow$) since it does not flow to the attacker's view. Finally, Condition 3 specifies that the attacker $H^\leftarrow$ has insufficient integrity to declassify information at label $H^\rightarrow$.

Proposition 2 (Noninterfering fair attacks). *For H such that*

1. $\Pi \vdash H^\rightarrow \sqsubseteq \tau'$
2. $\Pi \not\Vdash H^\rightarrow \sqsubseteq \Delta(H^\leftarrow)$
3. $\Pi \not\Vdash H^\leftarrow \geq \nabla(H^\rightarrow)$,

let $\Pi; \Gamma, x : \tau', \Gamma'; pc \vdash \bullet : \tau$. Then for all Π -fair attacks a where $\Pi'; \Gamma, x : \tau', \Gamma'; pc \vdash a : \tau$ for $\Pi' \supseteq \Pi$, and $a[x \mapsto v_i] \longrightarrow a'_i, a'_1 = a'_2$.

Proof. By Definition 2, there exists some Γ'' such that $\Pi; \Gamma''; pc \vdash a_i : \tau'$ where $\forall y : \tau'' \in \Gamma' . \Pi \vdash \tau'' \sqsubseteq \Delta(H^\leftarrow)$. Therefore since $\Pi \vdash H^\rightarrow \sqsubseteq \tau'$, we have $x \notin \Gamma''$ and thus $a[x \mapsto v_i] = a$. Consequently, $a'_1 = a'_2$. $\square$

In other words, Π -fair attacks must be well-typed on variables observable to the attacker in delegation context Π . Under the three conditions in Proposition 2, these attacks are noninterfering. Unfair or interfering attacks give the attacker enough power to break security directly by creating new declassifications without exploiting existing ones. Restricting attacks to noninterfering Π -fair attacks also rule out nonsensical scenarios such as when the “attacker” has the authority to read the confidential information. Noninterfering fair attacks represent an upper-bound on the power of the attacker over low-integrity portions of the program.

Our robust declassification theorem is stated in terms of an attacker's delegation context Π_H which is used to specify the noninterfering Π_H -fair attacks on a program. The program itself is typed under a distinct delegation context Π that may be extended by assume terms to permit new flows of information, including flows observable to the attacker.

Theorem 1 (Robust declassification). *Suppose $\Pi; \Gamma, x : \tau', \Gamma'; pc \vdash e[\bullet] : \tau$. For Π_H and H such that*

1. $\Pi_H \vdash H^\rightarrow \sqsubseteq \tau'$
2. $\Pi_H \not\Vdash H^\rightarrow \sqsubseteq \Delta(H^\leftarrow)$
3. $\Pi_H \not\Vdash H^\leftarrow \geq \nabla(H^\rightarrow)$,

Then for all Π_H -fair attacks $\vec{a}_1$ and $\vec{a}_2$ such that $\Pi; \Gamma, x:\tau', \Gamma'; pc \vdash e[\vec{a}_i] : \tau$, if $e[a_j][x \mapsto v_i] \longrightarrow e'_{ij}$ for $i, j \in \{1, 2\}$, then the for the traces $t_{ij} = (e[a_j][x \mapsto v_i]) \cdot e'_{ij}$, we have

$$t_{11} \approx_{\Delta(H^{\leftarrow})}^{\Pi} t_{21} \iff t_{12} \approx_{\Delta(H^{\leftarrow})}^{\Pi} t_{22}$$

Proof. By induction on the evaluation of e .

Case E-APP* ($e = (\lambda(y:\tau'')[pc'].e'')[a_i]$): By Proposition 2 x is not free in a_i , therefore

$$e[x \mapsto v_i] = \lambda(y:\tau'')[pc'].e''[x \mapsto v_i][a_i]$$

Suppose $e''[x \mapsto v_1] \approx_{\Delta(H^{\leftarrow})}^{\Pi} e''[x \mapsto v_2]$ then

$$(\lambda(y:\tau'')[pc'].e''[x \mapsto v_1])[a_j] \approx_{\Delta(H^{\leftarrow})}^{\Pi} (\lambda(y:\tau'')[pc'].e''[x \mapsto v_2])[a_j]$$

Since v_1 and v_2 are closed, we also have

$$e''[x \mapsto v_1][y \mapsto a_j] \approx_{\Delta(H^{\leftarrow})}^{\Pi} e''[x \mapsto v_2][y \mapsto a_j]$$

By a symmetric argument, we also have that $e''[x \mapsto v_1] \not\approx_{\Delta(H^{\leftarrow})}^{\Pi} e''[x \mapsto v_2]$ implies

$$e''[x \mapsto v_1][y \mapsto a_j] \not\approx_{\Delta(H^{\leftarrow})}^{\Pi} e''[x \mapsto v_2][y \mapsto a_j]$$

Then it remains to show that if $\mathcal{O}(\lambda(y:\tau'')[pc'].e'', \Pi, \Delta(H^{\leftarrow}), \rightarrow) = \circ$ then

$$\mathcal{O}(\llbracket e''[x \mapsto v_1][y \mapsto a_1] \rrbracket_{pc'}, \Pi, \Delta(H^{\leftarrow}), \rightarrow) = \mathcal{O}(\llbracket e''[x \mapsto v_2][y \mapsto a_1] \rrbracket_{pc'}, \Pi, \Delta(H^{\leftarrow}), \rightarrow)$$

$$\iff$$

$$\mathcal{O}(\llbracket e''[x \mapsto v_1][y \mapsto a_2] \rrbracket_{pc'}, \Pi, \Delta(H^{\leftarrow}), \rightarrow) = \mathcal{O}(\llbracket e''[x \mapsto v_2][y \mapsto a_2] \rrbracket_{pc'}, \Pi, \Delta(H^{\leftarrow}), \rightarrow)$$

This holds by the definition of $\mathcal{O}$ since $\mathcal{O}(\lambda(y:\tau'')[pc'].e'', \Pi, \Delta(H^{\leftarrow}), \rightarrow) = \circ$ implies that $\Pi \not\models \Delta(H^{\leftarrow}) \geq pc'$.

Case E-APP* ($e = a_j x$): If $\mathcal{O}(v_1, \Pi, \Delta(H^{\leftarrow}), \rightarrow) \neq \mathcal{O}(v_2, \Pi, \Delta(H^{\leftarrow}), \rightarrow)$ then by definition $t[0]_{1j} \not\approx_{\Delta(H^{\leftarrow})}^{\Pi} t[0]_{2j}$. Thus $t_{1j} \not\approx_{\Delta(H^{\leftarrow})}^{\Pi} t_{2j}$. Therefore assume $\mathcal{O}(v_1, \Pi, \Delta(H^{\leftarrow}), \rightarrow) = \mathcal{O}(v_2, \Pi, \Delta(H^{\leftarrow}), \rightarrow)$. Since a_j has the form $\lambda(y:\tau)[pc'].e''_j$, we need to show that

$$\mathcal{O}(\llbracket e''_j[y \mapsto v_1] \rrbracket_{pc'}, \Pi, \Delta(H^{\leftarrow}), \rightarrow) = \mathcal{O}(\llbracket e''_j[y \mapsto v_2] \rrbracket_{pc'}, \Pi, \Delta(H^{\leftarrow}), \rightarrow)$$

$$\iff$$

$$\mathcal{O}(\llbracket e''_2[y \mapsto v_1] \rrbracket_{pc'}, \Pi, \Delta(H^{\leftarrow}), \rightarrow) = \mathcal{O}(\llbracket e''_2[y \mapsto v_2] \rrbracket_{pc'}, \Pi, \Delta(H^{\leftarrow}), \rightarrow)$$

Without loss of generality, assume $\Pi \Vdash \Delta(H^{\leftarrow}) \geq pc^{\rightarrow}$. Then since $\mathcal{O}(v_1, \Pi, \Delta(H^{\leftarrow}), \rightarrow) = \mathcal{O}(v_2, \Pi, \Delta(H^{\leftarrow}), \rightarrow)$, by the definition of $\mathcal{O}$, we also have

$$\mathcal{O}(\llbracket e''_j[y \mapsto v_1] \rrbracket_{pc'}, \Pi, \Delta(H^{\leftarrow}), \rightarrow) = \mathcal{O}(\llbracket e''_j[y \mapsto v_2] \rrbracket_{pc'}, \Pi, \Delta(H^{\leftarrow}), \rightarrow)$$

Case E-TAPP*: Similar to E-APP*.

Case O-CTX: ($e = \langle a_j \rangle_\ell$) By Proposition 2.

Case O-CTX: ($e = \langle x \rangle_\ell$) Trivial.

Case E-CASE ($e = \text{case } x \text{ of } \text{inj}_1(y). a_j^1 \mid \text{inj}_2(y). a_j^2$): Not well typed since x cannot be a sum type.

Case E-CASE ($e = \text{case } a_j \text{ of } \text{inj}_1(y). e_1 \mid \text{inj}_2(y). e_2$): Observe

$$e[x \mapsto v_i] = \text{case } a_j \text{ of } \text{inj}_1(y). e_1[x \mapsto v_i] \mid \text{inj}_2(y). e_2[x \mapsto v_i]$$

Suppose

$$\text{case } a_j \text{ of } \text{inj}_1(y). e_1[x \mapsto v_1] \mid \text{inj}_2(y). e_2[x \mapsto v_1]$$

$$\approx_{\Delta(H^{\leftarrow})}^{\Pi}$$

$$\text{case } a_j \text{ of } \text{inj}_1(y). e_1[x \mapsto v_2] \mid \text{inj}_2(y). e_2[x \mapsto v_2]$$

then for $k \in \{1, 2\}$ we have

$$e_k[x \mapsto v_1] \approx_{\Delta(H^{\leftarrow})}^{\Pi} e_k[x \mapsto v_2]$$

therefore since v_1 and v_2 are closed, we also have

$$e_k[x \mapsto v_1][y \mapsto a_j] \approx_{\Delta(H^{\leftarrow})}^{\Pi} e_k[x \mapsto v_2][y \mapsto a_j]$$

Now suppose

$$\text{case } a_j \text{ of } \text{inj}_1(y). e_1[x \mapsto v_1] \mid \text{inj}_2(y). e_2[x \mapsto v_1]$$

$$\not\approx_{\Delta(H^{\leftarrow})}^{\Pi}$$

$$\text{case } a_i \text{ of } \text{inj}_1(y). e_1[x \mapsto v_2] \mid \text{inj}_2(y). e_2[x \mapsto v_2]$$

then by definition $t[0]_{1j} \not\approx_{\Delta(H^{\leftarrow})}^{\Pi} t[0]_{2j}$. Thus $t_{1j} \not\approx_{\Delta(H^{\leftarrow})}^{\Pi} t_{2j}$.

Case E-BINDM* ($e = \text{bind } y = [a_j] \text{ in } e'$) Follows by an argument similar to that for E-App.

Case E-BINDM* ($e = \text{bind } y = x \text{ in } e'[a_j]$) By BINDM (wlog), $\Pi; \Gamma, x : H^{\rightarrow} \text{ says } \tau', \Gamma', y : \tau''; pc \wedge H^{\rightarrow} \vdash [a_j] : \tau$ and by HOLE, $\Pi \Vdash pc \wedge H^{\rightarrow} \geq \Delta(H^{\leftarrow})$. Therefore the value of x is distinguishable by $\Delta(H^{\leftarrow})$ and $t[0]_{1j} \not\approx_{\Delta(H^{\leftarrow})}^{\Pi} t[0]_{2j}$.

Case E-ASSUME ($e = \text{assume } \langle p \geq q \rangle \text{ in } e[a_j]$): By Proposition 2, x is not free in a_j , therefore $e[a_j][x \mapsto v_i] = e[x \mapsto v_i][a_j]$. If

$$\mathcal{O}(e[x \mapsto v_1][a_j], \Pi, \Delta(H^{\leftarrow}), \rightarrow) \neq \mathcal{O}(e[x \mapsto v_2][a_j], \Pi, \Delta(H^{\leftarrow}), \rightarrow)$$

then by definition $t_{1j} \not\approx_{\Delta(H^{\leftarrow})}^{\Pi} t_{2j}$. If instead we have

$$\mathcal{O}(e[x \mapsto v_1][a_j], \Pi, \Delta(H^{\leftarrow}), \rightarrow) = \mathcal{O}(e[x \mapsto v_2][a_j], \Pi, \Delta(H^{\leftarrow}), \rightarrow)$$

then we need to prove that

$$\mathcal{O}(e'_{11}, \Pi, \Delta(H^\leftarrow), \rightarrow) = \mathcal{O}(e'_{21}, \Pi, \Delta(H^\leftarrow), \rightarrow) \Leftrightarrow \mathcal{O}(e'_{12}, \Pi, \Delta(H^\leftarrow), \rightarrow) = \mathcal{O}(e'_{22}, \Pi, \Delta(H^\leftarrow), \rightarrow)$$

where e' is of the form $e'_{ij} = e[x \mapsto v_i][a_j]$ where $\langle p \geq q \rangle$. Since x is not free in a_j , if

$$\mathcal{O}(e[x \mapsto v_1][a_j], \Pi \cdot \langle p \geq q \rangle, \Delta(H^\leftarrow), \rightarrow) = \mathcal{O}(e[x \mapsto v_2][a_j], \Pi \cdot \langle p \geq q \rangle, \Delta(H^\leftarrow), \rightarrow)$$

for $j = 1$, then it holds for $j = 2$ as well. Similarly, if

$$\mathcal{O}(e[x \mapsto v_1][a_j], \Pi \cdot \langle p \geq q \rangle, \Delta(H^\leftarrow), \rightarrow) \neq \mathcal{O}(e[x \mapsto v_2][a_j], \Pi \cdot \langle p \geq q \rangle, \Delta(H^\leftarrow), \rightarrow)$$

for $j = 1$, then it holds for $j = 2$.

Case W-APP: Similar to the assume case.

Case W-TAPP: Similar to the assume case.

Case W-UNPAIR: Trivial.

Case W-CASE: Similar to the assume case.

Case W-BINDM: Similar to the assume case.

Case W-ASSUME: Similar to the assume case.

□

Our formulation of robust declassification is in some sense more general than previous definitions since it permits some endorsements. Previous definitions of robust declassification [33,45] forbid endorsement altogether; *qualified robustness* [33] permits endorsement but offers only possibilistic security.

Theorem 1 permits some declassifications that prior definitions reject. For example, our definition admits declassifications of x even if $\Pi \Vdash H^\leftarrow \sqsubseteq \tau'$. In other words, even though low-integrity attacks cannot influence declassification, it is possible to declassify a secret input that has low-integrity. Therefore, an attacker that is permitted to influence the secret input could affect how much information is revealed. This is an example of a *malleable information flow* [15], which neither FLAC nor prior definitions (in the presence of endorsement) prevent in general. Cecchetti *et al.* [15] present a language based on FLAC that replaces assume with restricted declassification and endorsement terms to enforce nonmalleable information flow.

7.3. Bracketed semantics

Theorem 1 demonstrates that the indistinguishability of FLAC program traces is invariant under fair attacks. However, this result is only meaningful if FLAC programs that process protected information actually generate indistinguishable traces when we expect them to. These expectations are formalized in our noninterference theorem, but to understand and prove that theorem, we need additional formal tools.

Our noninterference proof is based on the bracketed semantics approach used by Pottier and Simonet [36]. This approach extends FLAC with *bracketed expressions* which represent two executions of a program, and allows us to reason about noninterference in FLAC, a *2-safety hyperproperty* [17], as *type safety* in the extended language.

Figure 11 presents the bracket extensions for FLAC. Where-values w and expressions e may be bracketed. New B-* evaluation rules propagate brackets out of subexpressions and ensure bracketed terms do not get stuck unless the unbracketed terms do. B-STEP evaluates expressions inside of brackets.

Syntax

$$w ::= \dots \mid (w \mid w)$$

$$e ::= \dots \mid (e \mid e)$$

Evaluation rules

$$[\text{B-STEP}] \quad \frac{e_i \longrightarrow e'_i \quad e'_j = e_j \quad \{i, j\} = \{1, 2\}}{(e_1 \mid e_2) \longrightarrow (e'_1 \mid e'_2)} \quad [\text{B-APP}] \quad (w_1 \mid w_2) w \longrightarrow (w_1 [w]_1 \mid w_2 [w]_2)$$

$$[\text{B-TAPP}] \quad (w \mid w') \tau \longrightarrow (w \tau \mid w' \tau) \quad [\text{B-UNPAIR}] \quad \text{proj}_i (w \mid w') \longrightarrow (\text{proj}_i w \mid \text{proj}_i w')$$

$$[\text{B-BINDM}] \quad \text{bind } x = (w \mid w') \text{ in } e \longrightarrow (\text{bind } x = w \text{ in } [e]_1 \mid \text{bind } x = w' \text{ in } [e]_2)$$

$$[\text{B-CASE}] \quad \frac{\{i, j\} = \{1, 2\}}{\text{case } (w \mid w') \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2 \longrightarrow (\text{case } w \text{ of } \text{inj}_1(x). [e_1]_1 \mid \text{inj}_2(x). [e_2]_1 \mid \text{case } w' \text{ of } \text{inj}_1(x). [e_1]_2 \mid \text{inj}_2(x). [e_2]_2)}$$

$$[\text{B-ASSUME}] \quad \text{assume } (w \mid w') \text{ in } e \longrightarrow (\text{assume } w \text{ in } [e]_1 \mid \text{assume } w' \text{ in } [e]_2)$$

Typing rules

$$[\text{BRACKET}] \quad \frac{e_1 = v_1 \iff e_2 \neq v_2 \quad \Pi \Vdash (H^\pi \sqcup pc) \sqsubseteq pc' \quad \Pi; \Gamma; pc' \vdash e_1 : \tau \quad \Pi; \Gamma; pc' \vdash e_2 : \tau \quad \Pi \vdash H^\pi \sqsubseteq \tau}{\Pi; \Gamma; pc \vdash (e_1 \mid e_2) : \tau}$$

$$[\text{BRACKET-VALUES}] \quad \frac{\Pi \vdash H^\pi \sqsubseteq \tau \quad \Pi; \Gamma; pc \vdash w_1 : \tau \quad \Pi; \Gamma; pc \vdash w_2 : \tau}{\Pi; \Gamma; pc \vdash (w_1 \mid w_2) : \tau}$$

Observation function

$$\mathcal{O}((e_1 \mid e_2), \Pi, p, \pi) = \begin{cases} \circ & \mathcal{O}(e_i, \Pi, p, \pi) = \circ \\ (\mathcal{O}(e_1, \Pi, p, \pi) \mid \mathcal{O}(e_2, \Pi, p, \pi)) & \text{otherwise} \end{cases}$$

Fig. 11. Extensions for bracketed semantics

The typing rules BRACKET and BRACKET-VALUES illustrate the primary purpose of the bracketed semantics: to link distinguishable evaluations of an expression to the expression's type. BRACKET requires that the bracketed expressions e_1 and e_2 are typable at a pc' that protects $H^\pi \sqcup pc$ and pc' and have a type τ protects H^π . BRACKET-VALUES relaxes the restriction on pc for where-values.

Our first result on the bracketed semantics is that they are sound and complete with respect to the unbracketed semantics. By soundness, we mean that given a step in the bracketed execution, then at least one of the left or right projections take a step such that they are in relation with the bracketed execu-

tion. By completeness, we mean that given a left and right execution, we can construct a corresponding bracketed execution. We now prove the soundness of the bracketed semantics.

Lemma 1 (Soundness). *If $e \longrightarrow e'$ then $[e]_k \longrightarrow^* [e']_k$ for $k \in \{1, 2\}$.*

Proof. By induction on the evaluation of e . Observe that all bracketed rules in Figure ?? except B-STEP only expand brackets, so $[e]_k = [e']_k$ for $k \in \{1, 2\}$. For B-STEP, $[e]_i \longrightarrow [e']_i$ and $[e]_j = [e']_j$. $\square$

Before proving completeness, we prove if the bracketed term is stuck then either left or right execution are stuck.

Lemma 2 (Stuck expressions). *If e gets stuck then $[e]_i$ is stuck for some $i \in \{1, 2\}$.*

Proof. By induction on the structure of e . See Appendix A for complete proof. $\square$

Lemma 3 (Completeness). *If $[e]_1 \longrightarrow^* v_1$ and $[e]_2 \longrightarrow^* v_2$, then there exists some v such that $e \longrightarrow^* v$.*

Proof. Assume $[e]_1 \longrightarrow^* v_1$ and $[e]_2 \longrightarrow^* v_2$. The extended set of rules in Figure ?? always move brackets out of subterms, and therefore can only be applied a finite number of times. Therefore, by Lemma 1, if e diverges, either $[e]_1$ or $[e]_2$ diverge; this contradicts our assumption.

Furthermore, by Lemma 2, if the evaluation of e gets stuck, either $[e]_1$ or $[e]_2$ gets stuck. Therefore, since we assumed $[e]_i \longrightarrow^* v_i$, then e must terminate. Thus, there exists some v such that $e \longrightarrow^* v$. $\square$

Using the soundness and completeness results, we can relate properties of programs executed under the bracketed semantics to executions under the non-bracketed semantics. In particular, since bracketed expressions represent distinguishable executions of a program, and may only have types that protect $H^\rightarrow$ (or $H^\leftarrow$), it is important to establish that the type of a FLAC expression is preserved by evaluation. Lemma 4 states this formally. This lemma helps us reason about whether an expression is bracketed based on its type.

Lemma 4 (Subject Reduction). *Let $\Pi; \Gamma; pc \vdash e : \tau$. If $e \longrightarrow e'$ then $\Pi; \Gamma; pc \vdash e' : \tau$.*

Proof. By induction on the evaluation of e . See Appendix A for complete proof and supporting lemmas. $\square$

7.4. Delegation invariance

FLAC programs dynamically extend trust relationships, enabling new flows of information. Nevertheless, well-typed programs have end-to-end semantic properties that enforce strong information security. These properties derive primarily from FLAC's control of the delegation context. The ASSUME rule ensures that only high-integrity proofs of authorization can extend the delegation context, and furthermore that such extensions occur only in high-integrity contexts.

That low-integrity contexts cannot extend the delegation context turns out to be a crucial property. This property allows us to state a useful invariant about the evaluation of FLAC programs. Recall that assume terms evaluate to where terms in the FLAC semantics. Thus, FLAC programs typically compute values containing a hierarchy of nested where terms. The terms record the values whose types were used to extend the delegation context during type checking.

For a well-typed FLAC program, we can prove that certain trust relationships could not have been added by the program. Characterizing these relationships requires a concept of the minimal authority required to cause one principal to act for another. Although similar, this idea is distinct from the voice of a principal. Consider the relationship between a and $a \wedge b$. The voice of $a \wedge b$, $\nabla(a \wedge b)$, is sufficient integrity to add a delegation $a \wedge b$ to a so that $a \geq a \wedge b$. Alternatively, having only the integrity of $\nabla(b)$ is sufficient to add a delegation $a \geq b$, which also results in $a \geq a \wedge b$. To precisely characterize which trust relationships cannot be added by a program, we need to identify this minimal integrity $\nabla(b)$ given the pair of principals a and $a \wedge b$. The following definitions are in service of this goal.

The first definition formalizes the idea that two principals are considered equivalent in a given context if they act for each other.

Definition 3 (Principal Equivalence). We say that two principals p and q are *equivalent* in Π , denoted $\Pi \Vdash p \equiv q$, if

$$\Pi \Vdash p \geq q \text{ and } \Pi \Vdash q \geq p.$$

Next, we define the *factorization* of two principals in a given context. For two principals, p and q , their factorization involves representing q as the conjunction of two principals $q_s \wedge q_d$ such that $p \geq q_s$ in the desired context. Note that p need not act for q_d .

Definition 4 (Factorization). A Π -factorization of an ordered pair of principals (p, q) is a tuple (p, q_s, q_d) such that $\Pi \Vdash q \equiv q_s \wedge q_d$ and $\Pi \Vdash p \geq q_s$.

A factorization is *static* if $\Pi = \emptyset$ (and thus $\mathcal{L} \models p \geq q_s$).

Finally, the minimal factorization of p and q is a q_s and q_d such that q_s has greater authority and q_d has less authority than any other factorization of p and q in the same context.

Definition 5 (Minimal Factorization). A Π -factorization (p, q_s, q_d) of (p, q) is *minimal* if for any Π -factorization (p, q'_s, q'_d) of (p, q) ,

$$\Pi \Vdash q_s \geq q'_s \text{ and } \Pi \Vdash q'_d \geq q_d$$

The minimal factorization (p, q_s, q_d) of p and q for a given Π and pc identifies the authority necessary to cause p to act for q . Because q_s is the principal with the greatest authority such that $p \geq q_s$ and $q \equiv q_s \wedge q_d$, then speaking for q_d is sufficient authority to cause p to act for q since adding the delegation $p \geq q_d$ would imply that $p \geq q$. This intuition also matches with the fact that $\Pi \Vdash p \geq q_d$ if and only if $q_d = \perp$, which is the case if and only if $\Pi \Vdash p \geq q$. Observe also that minimal Π -factorizations are also trivially unique up to equivalence.

Since the q_d component of minimal factorization can be thought of as the “gap” in authority between two principals, we use q_d to define the notion of principal *subtraction*.

Definition 6 (Principal Subtraction). Let (p, q_s, q_d) be the minimal Π -factorization of (p, q) . We define $q - p$ in Π to be q_d . That is, $\Pi \Vdash q - p \equiv q_d$. Note that $q - p$ is not defined outside of a judgement context.

Lemma 5 proves that minimal factorizations exist for all contexts and principals, so principal subtract is well defined.

Lemma 5 (Minimal Factorizations Exist). *For any context Π and principals p, q , there exists a minimal Π -factorization of (p, q) .*

Proof. Given (p, q) , we first let $q_s = p \vee q$. By definition, $\Pi \Vdash p \geq p \vee q$, and for all factorizations (p, q'_s, q'_d) , $\Pi \Vdash p \geq q'_s$ and $\Pi \Vdash q \geq q'_s$, so $\Pi \Vdash q_s \geq q'_s$.

Now let $D = \{r \in \mathcal{L} \mid \Pi \Vdash q \equiv q_s \wedge r\}$. All principals in $\mathcal{L}$ can be represented as a finite set of meets and joins of names in $\mathcal{N}$ so q and q_s are finite. Π is also finite, adding only finitely-many dynamic equivalences, so D is finite up to equivalence. Moreover, since $\Pi \Vdash q \equiv (p \vee q) \wedge q$ (by absorption) we have $q \in D$. Therefore D is always non-empty and we can define $q_d = \bigvee D$.

Now let (p, q'_s, q'_d) be any Π -factorization of (p, q) . We must show that $\Pi \Vdash q'_d \geq q_d$.

First, see that $\Pi \Vdash q \equiv q_s \wedge q'_d$. for one direction, observe that $\Pi \Vdash q \geq q_s$ and $\Pi \Vdash q \geq q_d$ (by Definition 3). For the other direction, since $\Pi \Vdash q_s \geq q'_s$, we have $\Pi \Vdash q_s \wedge q'_d \geq q'_s \wedge q'_d$, so $\Pi \Vdash q_s \wedge q'_d \geq q$.

Therefore, by the definition of D , we know $q'_d \in D$, so by the definition of $\vee$ and q_d , $\Pi \Vdash q'_d \geq q_d$. Thus (p, q_s, q_d) is a minimal Π -factorization of (p, q) . $\square$

We can now state precisely which trust relationships may change in a given information flow context.¹⁹

Lemma 6 (Delegation Invariance). *Let $\Pi; \Gamma; pc \vdash e : \tau$ such that $e \longrightarrow^* e'$ where v . Then*

1. *There exist $r, t \in \mathcal{L}$ and $\Pi' = \Pi, \langle r \geq t \rangle$ such that $\Pi; \Gamma; pc \vdash v : (r \geq t)$, $\Pi \Vdash pc \geq \nabla(t)$ and $\Pi'; \Gamma; pc \vdash e' : \tau$.*
2. *For all principals p and q , if $\Pi \nVdash p \geq q$ and $\Pi \nVdash pc \geq \nabla(q - p)$, then $\Pi' \nVdash p \geq q$.*

Proof. We prove the first part using subject reduction (Lemma 4) to get $\Pi; \Gamma; pc \vdash e' : \tau$. Then, from WHERE we have $\Pi'; \Gamma; pc \vdash e' : \tau$.

We prove the second part by contradiction. Assume that

$$\Pi' \Vdash p \geq q \tag{1}$$

Let (p, q_s, q_d) be the minimal Π -factorization of (p, q) . So, $\Pi \Vdash p \geq q_s$. Since $\Pi \nVdash p \geq q$, this implies that $\Pi \nVdash p \geq q_d$ but from (1), we have that

$$\Pi' \Vdash p \geq q_d \tag{2}$$

So any derivation of (2) must involve the derivation of $\Pi' \Vdash r \geq t$ via R-ASSUME, since R-ASSUME is the only rule that uses the contents of Π . Therefore, without loss of generality, we can assume that either (a) $\Pi \Vdash t \geq q_d$, or (b) for some q_1 and q_2 such that $\Pi \Vdash q_d \equiv q_1 \wedge q_2$, $\Pi \Vdash t \geq q_1$ and $\Pi \nVdash t \geq q_2$; otherwise $\Pi' \Vdash r \geq t$ is unnecessary to prove (2).

In fact, it must be the case that $\Pi \Vdash t \geq q_d$ (or equivalently $\Pi \Vdash t \geq q_1 \wedge q_2$). To see why, assume $\Pi \nVdash t \geq q_1 \wedge q_2$: specifically, $\Pi \Vdash t \geq q_1$ and $\Pi \nVdash t \geq q_2$. Since $\Pi' \Vdash p \geq q_s \wedge (q_1 \wedge q_2)$ but $\Pi \nVdash t \geq q_2$, it must be that $\Pi \nVdash p \geq q_s \wedge q_2$, otherwise $\Pi' \Vdash p \geq q_s \wedge (q_1 \wedge q_2)$ wouldn't hold.

Therefore $(p, q_s \wedge q_2, q_1)$ is a Π -factorization of (p, q) . Since (p, q_s, q_d) is a *minimal* Π -factorization, we have that $\Pi \Vdash q_2 \geq q_d$. From R-TRANS, we now have that $\Pi \Vdash t \geq q_d$, but since we assumed $\Pi \nVdash t \geq q_d$, we have a contradiction. Hence $\Pi \Vdash t \geq q_d$.

¹⁹The original delegation invariance lemma [5] was flawed due to a case where a minor delegation could have a cascading effect that enabled new delegations, breaking the desired invariant. This new (slightly more restrictive) formulation, stated in terms of principal subtraction, addresses more precisely the connection between the pc and the invariant trust relationships.

By the monotonicity of $\nabla(\cdot)$ with respect to $\geq$ (Proven in Coq for FLAM [7]), we have $\Pi \Vdash \nabla(t) \geq \nabla(q_d)$. As shown above, $\Pi \Vdash pc \geq \nabla(t)$, so by R-TRANS, $\Pi \Vdash pc \geq \nabla(q_d)$. However, recall that $\Pi \Vdash q - p \equiv q_d$ (Definition 6) contradicting our assumption that $\Pi \nVdash pc \geq \nabla(q - p)$.

Therefore, we obtain that $\Pi \nVdash p \geq q$. $\square$

7.5. Noninterference

Lemma 6 is critical for our proof of *noninterference*, a result that states that public and trusted output of a program cannot depend on restricted (secret or untrustworthy) information. Our proof of noninterference for FLAC programs relies on a proof of subject reduction under a bracketed semantics, based on the proof technique of Pottier and Simonet [36]. This technique is mostly standard, so we omit it here. The complete proof of subject reduction and other results are found in Appendix A.

In other noninterference results based on bracketed semantics, including [36], noninterference follows almost directly from the proof of subject reduction. This is because the subject reduction proof shows that evaluating a term cannot change its type. In FLAC, however, subject reduction alone is insufficient; evaluation may enable flows from secret or untrusted inputs to public and trusted types.

To see how, suppose e is a well-typed program according to $\Pi; \Gamma, x : s; pc \vdash e : \tau'$. Furthermore, let H be a principal such that $\Pi \vdash H \sqsubseteq \tau$ and $\Pi \nVdash H \sqsubseteq \tau'$. In other words, x is a “high” variable (more restrictive; secret and untrusted), and e evaluates to a “low” result (less restrictive; public and trusted). In [36], executions that differ only in secret or untrusted inputs must evaluate to the same value, since otherwise the value would not be well typed. In FLAC, however, if the pc has sufficient integrity, then an assume term could cause $\Pi'; pc \vdash H \sqsubseteq \tau'$ to hold in a delegation context Π' of a subterm of e .

The key to proving our result relies on using Lemma 6 to constrain the assumptions that can be added to Π' . Thus noninterference in FLAC is dependent on H and its relationship to pc and the type τ' . For confidentiality, most of this reasoning occurs in the proof of Lemma 7, which does most of the heavy lifting for the noninterference proof. Specifically, Lemma 7 states that for a well-typed program e , if the pc is insufficiently trusted to create new flows from $H^\rightarrow$ to $\ell^\rightarrow$, then if the portion of $[e]_1$ observable to $\ell^\rightarrow$ under Π is equal to the observable portion of $[e]_2$, then if $e \longrightarrow^* e'$, the observable portions of e' are still equal. A corresponding lemma for integrity holds by a similar argument.

Lemma 7 (Erasure conservation). *Suppose $\Pi; \Gamma; pc \vdash e : \tau$, then for some H and ℓ such that $\Pi \nVdash pc \geq \nabla(H^\rightarrow - \ell^\rightarrow)$, if $\mathcal{O}([e]_1, \Pi, \ell^\rightarrow, \rightarrow) = \mathcal{O}([e]_2, \Pi, \ell^\rightarrow, \rightarrow)$, then $e \longrightarrow^* e'$ implies $\mathcal{O}([e']_1, \Pi, \ell^\rightarrow, \rightarrow) = \mathcal{O}([e']_2, \Pi, \ell^\rightarrow, \rightarrow)$.*

Proof. By induction on the evaluation of e , using the definition of $\mathcal{O}$ and Lemma 6.

Case E-APP*: We have $\mathcal{O}([\lambda(x : \tau)[pc']. e' w]_1, \Pi, \ell^\rightarrow, \rightarrow) = \mathcal{O}([\lambda(x : \tau)[pc']. e' w]_2, \Pi, \ell^\rightarrow, \rightarrow)$.

We have to show that $\mathcal{O}([\llbracket e'[x \mapsto w] \rrbracket_{pc'}]_1, \Pi, \ell^\rightarrow, \rightarrow) = \mathcal{O}([\llbracket e'[x \mapsto w] \rrbracket_{pc'}]_2, \Pi, \ell^\rightarrow, \rightarrow)$.

Depending on the observability of pc' , we have two subcases.

Case $\Pi \nVdash \ell^\rightarrow \geq pc'^\rightarrow$: Here $\mathcal{O}([w]_1, \Pi, \ell^\rightarrow, \rightarrow) = \mathcal{O}([w]_2, \Pi, \ell^\rightarrow, \rightarrow)$ and $\mathcal{O}([e']_1, \Pi, \ell^\rightarrow, \rightarrow) = \mathcal{O}([e']_2, \Pi, \ell^\rightarrow, \rightarrow)$. And so, $\mathcal{O}([\llbracket e'[x \mapsto w] \rrbracket_{pc'}]_1, \Pi, \ell^\rightarrow, \rightarrow) = \mathcal{O}([\llbracket e'[x \mapsto w] \rrbracket_{pc'}]_2, \Pi, \ell^\rightarrow, \rightarrow)$.

Case $\Pi \Vdash \ell^\rightarrow \geq pc'^\rightarrow$: It remains to show that if $\mathcal{O}([\lambda(x : \tau)[pc']. e']_i, \Pi, \ell^\rightarrow, \rightarrow) = \circ$ then $\mathcal{O}([\llbracket e'[x \mapsto w] \rrbracket_{pc'}]_i, \Pi, \ell^\rightarrow, \rightarrow) = \circ$. This holds by the definition of $\mathcal{O}$ since $\mathcal{O}([\lambda(x : \tau)[pc']. e']_i, \Pi, \ell^\rightarrow, \rightarrow) = \circ$ implies that $\Pi \nVdash \ell^\rightarrow \geq pc'^\rightarrow$.

Case E-TAPP*: Similar to the E-APP* case.

Case E-BINDM*: We have

$$\mathcal{O}([\text{bind } x = \bar{\eta}_{\ell'} w \text{ in } e']_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([\text{bind } x = \bar{\eta}_{\ell'} w \text{ in } e']_2, \Pi, \ell^{\rightarrow}, \rightarrow)$$

Case $\Pi \not\models \ell^{\rightarrow} \geq \ell'^{\rightarrow}$: Here $\mathcal{O}([e']_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([e']_2, \Pi, \ell^{\rightarrow}, \rightarrow)$ and $\mathcal{O}([w]_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([w]_2, \Pi, \ell^{\rightarrow}, \rightarrow)$.

Case $\Pi \models \ell^{\rightarrow} \geq \ell'^{\rightarrow}$: It remains to show that if $\mathcal{O}([\bar{\eta}_{\ell'} w]_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$ then $\mathcal{O}([\ell e'[x \mapsto w]]_{\ell'}]_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$. This holds by the definition of $\mathcal{O}$ since $\mathcal{O}([\bar{\eta}_{\ell'} w]_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$ implies that $\Pi \not\models \ell^{\rightarrow} \geq \ell'^{\rightarrow}$.

Case O-CTX: Given $\mathcal{O}([\ell w]_{\ell'}]_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([\ell w]_{\ell'}]_2, \Pi, \ell^{\rightarrow}, \rightarrow)$. We have to prove that

$$\mathcal{O}([w]_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([w]_2, \Pi, \ell^{\rightarrow}, \rightarrow)$$

We have two cases: either $\mathcal{O}([\ell w]_{\ell'}]_i, \Pi, \ell^{\rightarrow}, \rightarrow) \neq \circ$ or $\mathcal{O}([\ell w]_{\ell'}]_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$. If $\mathcal{O}([\ell w]_{\ell'}]_i, \Pi, \ell^{\rightarrow}, \rightarrow) \neq \circ$, then the proof is straightforward (since the $\mathcal{O}$ function is invoked on same term w). If $\mathcal{O}([\ell w]_{\ell'}]_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$, then by inducting on the structure of w , we will show that $\mathcal{O}([w]_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$. By CTX, we know $\Pi; \Gamma; pc \vdash w : \tau$. If $[\ell w]_{\ell'}]_1 \neq [\ell w]_{\ell'}]_2$, then w contains a bracket subexpression $(w_1 \mid w_2)$. The “root” values v_i of w_i have one of the following forms:

Subcase $v_i = \lambda(x:\tau')[pc'].e'$: By Lemma 4, WHERE, and LAM, we have $\Pi'; \Gamma; pc; \vdash \lambda(x:\tau')[pc'].e' : \tau' \xrightarrow{pc'} \tau''$ (thus $\tau = \tau' \xrightarrow{pc'} \tau''$), where $\Pi' \supseteq \Pi$. By BRACKET-VALUE and P-FUN, $\Pi' \vdash H^{\rightarrow} \sqsubseteq pc'$ and $\Pi' \vdash H^{\rightarrow} \sqsubseteq \tau''$. By assumption, we have $\Pi \not\models pc \geq \nabla(H^{\rightarrow} - \ell^{\rightarrow})$. Therefore we know that $\Pi \not\models \ell^{\rightarrow} \geq H^{\rightarrow}$ since otherwise (by Definition 6) $\Pi \models \nabla(H^{\rightarrow} - \ell^{\rightarrow}) \equiv \perp$. Invoking DI (Lemma 6), we have that $\Pi' \not\models \ell^{\rightarrow} \geq pc'^{\rightarrow}$. By the definition of $\mathcal{O}$, we have $\mathcal{O}(\lambda(x:\tau')[pc'].e', \Pi', \ell^{\rightarrow}, \rightarrow) = \circ$ and so $\mathcal{O}(w_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$ (See the case where for the $\mathcal{O}$ function). Hence proved.

Subcase $v_i = \Lambda X[pc'].e'$: Similar to above case.

Subcase $v_i = \bar{\eta}_{\ell'} w'$: By Lemma 4, WHERE, and SEALED, we have $\Pi'; \Gamma; pc; \vdash \bar{\eta}_{\ell'} w' : \tau$, where $\Pi' \supseteq \Pi$. By BRACKET-VALUE and P-LBL, $\Pi' \vdash H^{\rightarrow} \sqsubseteq \ell'$. By assumption, we have $\Pi \not\models pc \geq \nabla(H^{\rightarrow} - \ell^{\rightarrow})$. Therefore we know that $\Pi \not\models \ell^{\rightarrow} \geq H^{\rightarrow}$ since otherwise (by Definition 6) $\Pi \models \nabla(H^{\rightarrow} - \ell^{\rightarrow}) \equiv \perp$. Invoking DI (Lemma 6), we have that $\Pi' \not\models \ell^{\rightarrow} \geq \ell'^{\rightarrow}$. Therefore, $\mathcal{O}(\bar{\eta}_{\ell'} w', \Pi', \ell^{\rightarrow}, \rightarrow) = \circ$ and $\mathcal{O}(w_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$.

Subcase $v_i = \langle w_1, w_2 \rangle$: Straightforward application of induction hypothesis.

Case ASSUME: Given that $\mathcal{O}([\text{assume } \langle p \geq q \rangle \text{ in } e]_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([\text{assume } \langle p \geq q \rangle \text{ in } e]_2, \Pi, \ell^{\rightarrow}, \rightarrow)$. We have to prove that

$$\mathcal{O}([e \text{ where } \langle p \geq q \rangle]_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([e \text{ where } \langle p \geq q \rangle]_2, \Pi, \ell^{\rightarrow}, \rightarrow)$$

There are two cases: either $\mathcal{O}([\text{assume } \langle p \geq q \rangle \text{ in } e]_i, \Pi, \ell^{\rightarrow}, \rightarrow) \neq \circ$ or $\mathcal{O}([\text{assume } \langle p \geq q \rangle \text{ in } e]_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$. Consider the former case. From the given conditions, we have $\mathcal{O}([e]_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([e]_2, \Pi, \ell^{\rightarrow}, \rightarrow)$. Therefore, $\mathcal{O}([e]_1, \Pi, \langle p \geq q \rangle, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([e]_2, \Pi, \langle p \geq q \rangle, \ell^{\rightarrow}, \rightarrow)$. Hence proved.

Consider the latter case. That is, $\mathcal{O}([\text{assume } \langle p \geq q \rangle \text{ in } e]_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$. By the definition of $\mathcal{O}$ and $[\cdot]_i$, we know that $\mathcal{O}([e]_i, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$. We will show that $\mathcal{O}([e \text{ where } \langle p \geq q \rangle]_i, \Pi, \ell^{\rightarrow}, \rightarrow)$

) = $\circ$ by inducting on the structure of e . From subject reduction (Lemma 4) and WHERE, we know $\Pi, \langle p \geq q \rangle; \Gamma, pc \vdash e : \tau$. The interesting cases are the terms whose observability depend on an acts-for judgment. Since $\Pi \not\vdash pc \geq \nabla(H^\rightarrow - \ell^\rightarrow)$, we have that $\Pi \not\vdash \ell^\rightarrow \geq H^\rightarrow$. We now show for each case that DI (Lemma 6) implies $\Pi, \langle p \geq q \rangle \not\vdash \ell^\rightarrow \geq H^\rightarrow$ that the delegation $\langle p \geq q \rangle$ cannot reveal new terms visible to $\ell^\rightarrow$.

Subcase $[e]_i = \lambda(x:\tau')[pc']$. e' : Suppose $\Pi \Vdash \ell^\rightarrow \geq pc'^\rightarrow$, then $\Pi, \langle p \geq q \rangle \Vdash \ell^\rightarrow \geq pc'^\rightarrow$.

Now suppose $\Pi \not\vdash \ell^\rightarrow \geq pc'^\rightarrow$. Since $\Pi \not\vdash pc \geq \nabla(H^\rightarrow - \ell^\rightarrow)$, then by Lemma 6, $\Pi, \langle p \geq q \rangle \not\vdash \ell^\rightarrow \geq H^\rightarrow$. But by BRACKET-VALUES, LAM, and P-FUN, we know $\Pi, \langle p \geq q \rangle \Vdash pc'^\rightarrow \geq H^\rightarrow$. Therefore, we have that $\Pi, \langle p \geq q \rangle \not\vdash \ell^\rightarrow \geq pc'^\rightarrow$, since otherwise we could prove $\Pi, \langle p \geq q \rangle \not\vdash \ell^\rightarrow \geq H^\rightarrow$ by transitivity.

Subcase $[e]_i = \Lambda X[pc']$. e' : Similar to above case.

Subcase $e = (\bar{\eta}_{\ell'} w'_1 \mid \bar{\eta}_{\ell'} w'_2)$: Suppose $\Pi \Vdash \ell^\rightarrow \geq \ell'^\rightarrow$, then $\Pi, \langle p \geq q \rangle \Vdash \ell^\rightarrow \geq \ell'^\rightarrow$ also.

Now suppose $\Pi \not\vdash \ell^\rightarrow \geq \ell'^\rightarrow$. Since $\Pi \not\vdash pc \geq \nabla(H^\rightarrow - \ell^\rightarrow)$, then by Lemma 6, $\Pi, \langle p \geq q \rangle \not\vdash \ell^\rightarrow \geq H^\rightarrow$. But by BRACKET-VALUES, SEALED, and P-LBL, we know $\Pi, \langle p \geq q \rangle \Vdash \ell'^\rightarrow \geq H^\rightarrow$. Therefore, we have that $\Pi, \langle p \geq q \rangle \not\vdash \ell^\rightarrow \geq \ell'^\rightarrow$, since otherwise we could prove $\Pi, \langle p \geq q \rangle \not\vdash \ell^\rightarrow \geq H^\rightarrow$ by transitivity.

Subcase $[e]_i = (\ell' \triangleright_{e'})$: Similar to above case.

Case W-*: Similar to ASSUME case.

Case B-STEP: Straightforward application of induction hypothesis.

Case B-*: Observe that, with the exception of B-STEP, all B-* rules step from $e \longrightarrow^* e'$ such that $[e]_i = [e']_i$. Therefore the lemma trivially holds.

E-Unpair: Trivial.

E-Case: Trivial.

E-UnitM: Trivial.

E-Eval: Straightforward application of induction hypothesis.

□

For more precision, we specialize our noninterference theorem on confidentiality, but a noninterference theorem for integrity holds by a similar argument. Theorem 2 states that for some principal $H^\rightarrow$ that flows to τ' (Condition 1) but not ℓ says τ (Condition 2), if pc is insufficiently trusted to create new flows from $H^\rightarrow$ to $\ell^\rightarrow$ (Condition 3), then an execution of e that differs only in the value of τ' -typed inputs, the computed values must be equal.

Theorem 2 (Confidentiality Noninterference). *Let $\Pi; \Gamma; x : \tau'; pc \vdash e : \ell$ says τ for some H such that*

1. $\Pi \Vdash H^\rightarrow \wedge \top^\leftarrow \sqsubseteq \tau'$
2. $\Pi \not\vdash H^\rightarrow \sqsubseteq \ell^\rightarrow$
3. $\Pi \not\vdash pc \geq \nabla(H^\rightarrow - \ell^\rightarrow)$

Then for all $v_z, z \in \{1, 2\}$ such that $\Pi; \Gamma; pc \vdash v_z : \tau'$, if $e[x \mapsto v_z] \xrightarrow{t_z^} w'_z$, then $t_1 \approx_{\ell^\rightarrow}^\Pi t_2$.*

Proof. From the soundness and completeness properties of the bracketed language (Lemmas 1 and 3), we can construct a bracketed execution $e[x \mapsto (v_1 \mid v_2)] \xrightarrow{t^*} v'$ such that $[v']_z = v'_z$ and $[t]_z = t_z$ for $z = \{1, 2\}$. We will occasionally write v or v' as shorthand for $(v_1 \mid v_2)$ or $(v'_1 \mid v'_2)$.

Since $\Pi; \Gamma; pc \vdash v_z : \tau'$ for $z \in \{1, 2\}$, then we have $\Pi; \Gamma; pc \vdash (v_1 \mid v_2) : \tau'$ via BRACKET-VALUES. Therefore, by Lemma 15 (Variable Substitution) of $\Pi; \Gamma; x : \tau'; pc \vdash e : \ell$ says τ , we have $\Pi; \Gamma; pc \vdash$

$e[x \mapsto (v_1 \mid v_2)] : \ell$ says τ . Then by repeatedly applying Lemma 4 for each step $e[x \mapsto v] \xrightarrow{t}^* e'$, we have $\Pi; \Gamma; pc \vdash e' : \ell$ says τ .

We now want to prove that $t_1 \approx_{\ell}^{\Pi} t_2$. First, consider $e[x \mapsto (v_1 \mid v_2)]$. To prove that $\mathcal{O}(e[x \mapsto v_1], \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}(e[x \mapsto v_2], \Pi, \ell^{\rightarrow}, \rightarrow)$, it suffices to show that $\mathcal{O}(v_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}(v_2, \Pi, \ell^{\rightarrow}, \rightarrow)$. We do this by induction on the structure of v_z .

Case $v_z = ()$: Trivial.

Case $v_z = \mathbf{inj}_i e'$: Not well typed since $\Pi \Vdash H^{\rightarrow} \sqsubseteq \tau'$ cannot hold for types of the form $\tau_1 + \tau_2$.

Case $v_z = \bar{\eta}_{\ell'} w$: Since $\Pi \Vdash H^{\rightarrow} \sqsubseteq \tau'$, by UNITM and P-LBL we know $\Pi \Vdash H^{\rightarrow} \sqsubseteq \ell'$. Therefore we have $\Pi \not\vdash \ell^{\rightarrow} \geq \ell'^{\rightarrow}$, since otherwise we would have $\Pi \Vdash H^{\rightarrow} \sqsubseteq \ell^{\rightarrow}$, contradicting our assumption that $\Pi \not\vdash H^{\rightarrow} \wedge \top^{\leftarrow} \sqsubseteq \ell$. By the definition of $\mathcal{O}$, $\Pi \not\vdash \ell^{\rightarrow} \geq \ell'^{\rightarrow}$ implies that $\mathcal{O}(\bar{\eta}_{\ell'} w, \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$ for $z \in \{1, 2\}$.

Case $v_z = \lambda(x:\tau'')[pc'] . e'$: Since $\Pi \Vdash H^{\rightarrow} \sqsubseteq \tau'$, by LAM and P-FUN we know $\Pi \Vdash H^{\rightarrow} \sqsubseteq pc'$. Therefore we have $\Pi \not\vdash \ell^{\rightarrow} \geq pc'^{\rightarrow}$, since otherwise we would have $\Pi \Vdash H^{\rightarrow} \sqsubseteq \ell^{\rightarrow}$, contradicting our assumption that $\Pi \not\vdash H^{\rightarrow} \wedge \top^{\leftarrow} \sqsubseteq \ell$. By the definition of $\mathcal{O}$, $\Pi \not\vdash \ell^{\rightarrow} \geq pc'^{\rightarrow}$ implies that $\mathcal{O}(\lambda(x:\tau'')[pc'] . e', \Pi, \ell^{\rightarrow}, \rightarrow) = \circ$ for $z \in \{1, 2\}$.

Case $v_z = \Lambda X[pc'] . e'$: Similar to above case.

Case $v_z = \langle v_1'', v_2'' \rangle$: Straightforward application of the induction hypothesis.

Thus $\mathcal{O}([e[x \mapsto v]]_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([e[x \mapsto v]]_2, \Pi, \ell^{\rightarrow}, \rightarrow)$. Now, by Lemma 7, for any e' such that $e \xrightarrow{*} e'$, we have $\mathcal{O}([e']_1, \Pi, \ell^{\rightarrow}, \rightarrow) = \mathcal{O}([e']_2, \Pi, \ell^{\rightarrow}, \rightarrow)$. Therefore $t_1 \approx_{\ell}^{\Pi} t_2$. $\square$

Noninterference is a key tool for obtaining many of the security properties we seek. For instance, noninterference is essential for verifying the properties of commitment schemes discussed in Section 5.1.

Unlike some definitions of noninterference, our definition does not prohibit all downgrading. Instead, Conditions 1 and 2 define relationships between principals H and ℓ , and Condition 3 ensures (via Lemma 6) those relationships remain unchanged by delegations made within the program, and thus by Lemma 7, the trace of a program observable to an attack remains unchanged during evaluation.

8. Related Work

Many languages and systems for authorization or access control have combined aspects of information security and authorization (e.g., [43,23,30,39,31,9]) in dynamic settings. However, almost all are susceptible to security vulnerabilities that arise from the interaction of information flow and authorization [7]: probing attacks, delegation loopholes, poaching attacks, and authorization side channels.

DCC [4,2] has been used to model both authorization and information flow, but not simultaneously. DCC programs are type-checked with respect to a static security lattice, whereas FLAC programs can introduce new trust relationships during evaluation, enabling more general applications.

Boudol [14] defines terms that enable or disable flows for a lexical scope, similar to assume terms, but does not restrict their usage. Rx [39] and RTI [9] use labeled roles to represent information flow policies. The integrity of a role restricts who may change policies. However, information flow in these languages is not robust [33]: attackers may indirectly affect how flows change when authorized principals modify policies.

Some prior approaches have sought to reason about the information security of authorization mechanisms. Becker [10] discusses *probing attacks* that leak confidential information to an attacker. Garg and

Pfenning [21] present a logic that ensures assertions made by untrusted principals cannot influence the truth of statements made by other principals.

Previous work has studied information flow control with higher-order functions and side effects. In the SLam calculus [22], implicit flows due to side effects are controlled via *indirect reader* annotations on types. Zdancewic and Myers [46] and Flow Caml [36] control implicit flows via *pc* annotations on function types. FLAC also controls side effects via a *pc* annotation, but here the side effects are changes in trust relationships that define which flows are permitted. Tse and Zdancewic [41] also extend DCC with a program-counter label but for a different purpose: their *pc* tracks information about the protection context, permitting more terms to be typed.

DKAL* [25] is an executable specification language for authorization protocols, simplifying analysis of protocol implementations. FLAC may be used as a specification language, but FLAC offers stronger guarantees regarding the information security of specified protocols. Errors in DKAL* specifications could lead to vulnerabilities. For instance, DKAL* provides no intrinsic guarantees about confidentiality, which could lead to authorization side channels or probing attacks.

The Jif programming language [32,34] supports dynamically computed labels through a simple dependent type system. Jif also supports dynamically changing trust relationships through operations on principal objects [16]. Because the signatures of principal operations (e.g., to add a new trust relationship) are missing the constraints imposed by FLAC, authorization can be used as a covert channel. FLAC shows how to close these channels in languages like Jif.

Dependently-typed languages are often expressive enough to encode authorization policies, information flow policies, or both. The F* [40] type system is capable of enforcing information flow and authorization policies. Typing rules like those in FLAC could probably be encoded within its type system, but so could incorrect, insecure rules. Thus, FLAC contributes a model for encodings that enforce strong information security. Aura [26] embeds a DCC-based proof language and type system in a dependently-typed general-purpose functional language. As in DCC, Aura programs may derive new authorization proofs using existing proof terms and a monadic bind operator. However, since Aura only tracks dependencies between proofs, it is ill-suited for reasoning about the end-to-end information-flow properties of authorization mechanisms.

9. Conclusion

Existing security models do not account fully for the interactions between authorization and information flow. The result is that both the implementations and the uses of authorization mechanisms can lead to insecure information flows that violate confidentiality or integrity. The security of information flow mechanisms can also be compromised by dynamic changes in trust. This paper has proposed FLAC, a core programming language that coherently integrates these two security paradigms, controlling the interactions between dynamic authorization and secure information flow. FLAC offers strong guarantees and can serve as the foundation for building software that implements and uses authorization securely. Further, FLAC can be used to reason compositionally about secure authorization and secure information flow, guiding the design and implementation of future security mechanisms.

Acknowledgments

We thank Mike George, Elaine Shi, and Fred Schneider for helpful discussions, our anonymous reviewers for their comments and suggestions, and Jed Liu, Matt Stillerman, and Priyanka Mondal for

feedback on early drafts. This work was supported by grant N00014-13-1-0089 from the Office of Naval Research, by MURI grant FA9550-12-1-0400, by grant FA9550-16-1-0351 from the Air Force Office of Scientific Research, by an NDSEG fellowship, by the Intelligence Advanced Research Projects Activity (2019-19-020700006), and by grants from the National Science Foundation (CCF-0964409, CNS-1565387, CNS-1704845, CCF-1750060). This paper does not necessarily reflect the views of any of these sponsors.

References

- [1] M. Abadi. Logic in access control. In *Proceedings of the 18th Annual IEEE Symposium on Logic in Computer Science, LICS '03*, pages 228–233, Washington, DC, USA, 2003. IEEE Computer Society. URL <http://dx.doi.org/10.1109/LICS.2003.1210062>.
- [2] M. Abadi. Access control in a core calculus of dependency. In *11th ACM SIGPLAN Int'l Conf. on Functional Programming*, pages 263–273, New York, NY, USA, 2006. ACM. doi: 10.1145/1159803.1159839. URL <http://doi.acm.org/10.1145/1159803.1159839>.
- [3] M. Abadi. Variations in access control logic. In R. van der Meyden and L. van der Torre, editors, *Deontic Logic in Computer Science*, volume 5076 of *Lecture Notes in Computer Science*, pages 96–109. Springer Berlin Heidelberg, 2008. ISBN 978-3-540-70524-6. doi: 10.1007/978-3-540-70525-3_9. URL http://dx.doi.org/10.1007/978-3-540-70525-3_9.
- [4] M. Abadi, A. Banerjee, N. Heintze, and J. Riecke. A core calculus of dependency. In *26th ACM Symp. on Principles of Programming Languages (POPL)*, pages 147–160, Jan. 1999. URL <http://dl.acm.org/citation.cfm?id=292555>.
- [5] O. Arden and A. C. Myers. A calculus for flow-limited authorization. In *29th IEEE Symp. on Computer Security Foundations (CSF)*, pages 135–147, June 2016. URL <http://www.cs.cornell.edu/andru/papers/flac>.
- [6] O. Arden, M. D. George, J. Liu, K. Vikram, A. Askarov, and A. C. Myers. Sharing mobile code securely with information flow control. In *IEEE Symp. on Security and Privacy*, pages 191–205, May 2012. URL <http://www.cs.cornell.edu/andru/papers/mobile.html>.
- [7] O. Arden, J. Liu, and A. C. Myers. Flow-limited authorization. In *28th IEEE Symp. on Computer Security Foundations (CSF)*, pages 569–583, July 2015. URL <http://www.cs.cornell.edu/andru/papers/flam>.
- [8] O. Arden, J. Liu, and A. C. Myers. Flow-limited authorization: Technical report. Technical Report 1813–40138, Cornell University Computing and Information Science, May 2015. URL <http://hdl.handle.net/1813/40138>.
- [9] S. Bandhakavi, W. Winsborough, and M. Winslett. A trust management approach for flexible policy management in security-typed languages. In *Computer Security Foundations Symposium, 2008*, pages 33–47, 2008. URL <http://ieeexplore.ieee.org/xpl/articleDetails.jsp?arnumber=4556677>.
- [10] M. Y. Becker. Information flow in trust management systems. *Journal of Computer Security*, 20(6):677–708, 2012.
- [11] M. Y. Becker, C. Fournet, and A. D. Gordon. SecPAL: Design and semantics of a decentralized authorization language. *Journal of Computer Security*, 18(4):619–665, 2010. URL <http://research.microsoft.com/apps/pubs/default.aspx?id=70334>.
- [12] K. J. Biba. Integrity considerations for secure computer systems. Technical Report ESD-TR-76-372, USAF Electronic Systems Division, Bedford, MA, Apr. 1977. (Also available through National Technical Information Service, Springfield Va., NTIS AD-A039324.).
- [13] A. Birgisson, J. G. Politz, Ú. Erlingsson, A. Taly, M. Vrabie, and M. Lenczner. Macaroons: Cookies with contextual caveats for decentralized authorization in the cloud. In *Network and Distributed System Security Symp.*, 2014. URL <http://theory.stanford.edu/~ataly/Papers/macaroons.pdf>.
- [14] G. Boudol. Secure information flow as a safety property. In *Formal Aspects in Security and Trust (FAST)*, pages 20–34. Springer, 2008.
- [15] E. Cecchetti, A. C. Myers, and O. Arden. Nonmalleable information flow control. In *24th ACM Conf. on Computer and Communications Security (CCS)*, pages 1875–1891, Oct. 2017.
- [16] S. Chong, K. Vikram, and A. C. Myers. SIF: Enforcing confidentiality and integrity in web applications. In *16th USENIX Security Symp.*, Aug. 2007. URL <http://www.cs.cornell.edu/andru/papers/sif.pdf>.
- [17] M. R. Clarkson and F. B. Schneider. Hyperproperties. In *IEEE Symp. on Computer Security Foundations (CSF)*, pages 51–65, June 2008.
- [18] R. A. Eisenberg and S. Weirich. Dependently typed programming with singletons. In *2012 Haskell Symposium*, pages 117–130. ACM, Sept. 2012. doi: 10.1145/2364506.2364522. URL <http://doi.acm.org/10.1145/2364506.2364522>.
- [19] C. Ellison, B. Frantz, B. Lampson, R. Rivest, B. Thomas, and T. Ylonen. SPKI certificate theory. Internet RFC-2693, Sept. 1999.
- [20] D. Ferraiolo and R. Kuhn. Role-based access controls. In *15th National Computer Security Conference*, 1992.

- [21] D. Garg and F. Pfenning. Non-interference in constructive authorization logic. In *19th IEEE Computer Security Foundations Workshop (CSFW)*, New Jersey, USA, 2006. IEEE. URL <http://dl.acm.org/citation.cfm?id=1155684>.
- [22] N. Heintze and J. G. Riecke. The SLam calculus: Programming with secrecy and integrity. In *25th ACM Symp. on Principles of Programming Languages (POPL)*, pages 365–377, San Diego, California, Jan. 1998.
- [23] M. Hicks, S. Tse, B. Hicks, and S. Zdancewic. Dynamic updating of information-flow policies. In *Foundations of Computer Security Workshop*, 2005. URL <http://www.cs.umd.edu/~mwh/papers/secupdate.pdf>.
- [24] J. Howell and D. Kotz. A formal semantics for SPKI. In *ESORICS 2000*, volume 1895 of *Lecture Notes in Computer Science*, pages 140–158. Springer Berlin Heidelberg, 2000. URL http://dx.doi.org/10.1007/10722599_9.
- [25] J.-B. Jeannin, G. de Caso, J. Chen, Y. Gurevich, P. Naldurg, and N. Swamy. DKAL*: Constructing executable specifications of authorization protocols. In *Engineering Secure Software and Systems*, pages 139–154. Springer, 2013. URL http://link.springer.com/chapter/10.1007/978-3-642-36563-8_10.
- [26] L. Jia, J. A. Vaughan, K. Mazurak, J. Zhao, L. Zarko, J. Schorr, and S. Zdancewic. Aura: A programming language for authorization and audit. In *13th ACM SIGPLAN Int'l Conf. on Functional Programming*, Sept. 2008.
- [27] B. Lampson, M. Abadi, M. Burrows, and E. Wobber. Authentication in distributed systems: Theory and practice. In *13th ACM Symp. on Operating System Principles (SOSP)*, pages 165–182, Oct. 1991. *Operating System Review*, 253(5).
- [28] N. Li, J. C. Mitchell, and W. H. Winsborough. Design of a role-based trust-management framework. In *IEEE Symp. on Security and Privacy*, pages 114–130, 2002. URL <http://dl.acm.org/citation.cfm?id=829514.830539>.
- [29] B. Martin, M. Brown, A. Paller, D. Kirby, and S. Christey. 2011 cwe/sans top 25 most dangerous software errors. *Common Weakness Enumeration*, 7515, 2011. URL <http://cwe.mitre.org/top25/>.
- [30] K. Minami and D. Kotz. Secure context-sensitive authorization. *Journal of Pervasive and Mobile Computing*, 1(1):123–156, March 2005. doi: 10.1016/j.pmcj.2005.01.004. URL <http://www.cs.dartmouth.edu/~dfk/papers/minami-jcsa.pdf>.
- [31] K. Minami and D. Kotz. Scalability in a secure distributed proof system. In *4th International Conference on Pervasive Computing*, volume 3968 of *Lecture Notes in Computer Science*, pages 220–237, Dublin, Ireland, May 2006. Springer-Verlag. doi: 10.1007/11748625_14. URL <http://www.cs.dartmouth.edu/~dfk/papers/minami-scalability.pdf>.
- [32] A. C. Myers. JFlow: Practical mostly-static information flow control. In *26th ACM Symp. on Principles of Programming Languages (POPL)*, pages 228–241, Jan. 1999. URL <http://www.cs.cornell.edu/andru/papers/pop199/pop199.pdf>.
- [33] A. C. Myers, A. Sabelfeld, and S. Zdancewic. Enforcing robust declassification and qualified robustness. *Journal of Computer Security*, 14(2):157–196, 2006. URL <http://www.cs.cornell.edu/andru/papers/robdecl-jcs>.
- [34] A. C. Myers, L. Zheng, S. Zdancewic, S. Chong, and N. Nystrom. Jif 3.0: Java information flow. Software release, <http://www.cs.cornell.edu/jif>, July 2006. URL <http://www.cs.cornell.edu/jif>.
- [35] M. Naor. Bit commitment using pseudorandomness. *Journal of cryptology*, 4(2):151–158, 1991.
- [36] F. Pottier and V. Simonet. Information flow inference for ML. *ACM Trans. on Programming Languages and Systems*, 25(1), Jan. 2003.
- [37] A. Sabelfeld and A. C. Myers. Language-based information-flow security. *IEEE Journal on Selected Areas in Communications*, 21(1):5–19, Jan. 2003. URL <http://www.cs.cornell.edu/andru/papers/jsac/sm-jsac03.pdf>.
- [38] F. B. Schneider, K. Walsh, and E. G. Sirer. Nexus Authorization Logic (NAL): Design rationale and applications. *ACM Trans. Inf. Syst. Secur.*, 14(1):8:1–8:28, June 2011. doi: 10.1145/1952982.1952990. URL <http://doi.acm.org/10.1145/1952982.1952990>.
- [39] N. Swamy, M. Hicks, S. Tse, and S. Zdancewic. Managing policy updates in security-typed languages. In *19th IEEE Computer Security Foundations Workshop (CSFW)*, pages 202–216, July 2006. URL <http://www.cs.umd.edu/projects/PL/rx/rx.pdf>.
- [40] N. Swamy, J. Chen, C. Fournet, P.-Y. Strub, K. Bhargavan, and J. Yang. Secure distributed programming with value-dependent types. In *16th ACM SIGPLAN Int'l Conf. on Functional Programming*, ICFP '11, pages 266–278, New York, NY, USA, 2011. ACM. URL <http://doi.acm.org/10.1145/2034773.2034811>.
- [41] S. Tse and S. Zdancewic. Translating dependency into parametricity. In *9th ACM SIGPLAN Int'l Conf. on Functional Programming*, pages 115–125, 2004. doi: 10.1145/1016850.1016868. URL <http://doi.acm.org/10.1145/1016850.1016868>.
- [42] P. Wadler. Propositions as types. *Communications of the ACM*, 2015.
- [43] W. H. Winsborough and N. Li. Safety in automated trust negotiation. In *IEEE Symp. on Security and Privacy*, pages 147–160, May 2004. doi: 10.1109/SECPRI.2004.1301321. URL <http://dl.acm.org/citation.cfm?id=1178623>.
- [44] A. K. Wright and M. Felleisen. A syntactic approach to type soundness. *Information and Computation*, 115(1):38–94, 1994. ISSN 0890-5401. doi: 10.1006/inco.1994.1093. URL <http://dx.doi.org/10.1006/inco.1994.1093>.
- [45] S. Zdancewic and A. C. Myers. Robust declassification. In *14th IEEE Computer Security Foundations Workshop (CSFW)*, pages 15–23, June 2001. doi: 10.1109/CSFW.2001.930133. URL <http://www.cs.cornell.edu/andru/papers/csaw01.pdf>.

- [46] S. Zdancewic and A. C. Myers. Secure information flow via linear continuations. *Higher-Order and Symbolic Computation*, 15(2–3):209–234, Sept. 2002. ISSN 1388-3690. doi: 10.1023/A:1020843229247. URL <http://dx.doi.org/10.1023/A:1020843229247>.
- [47] S. Zdancewic, L. Zheng, N. Nystrom, and A. C. Myers. Secure program partitioning. *ACM Trans. on Computer Systems*, 20(3):283–328, Aug. 2002. URL <http://www.cs.cornell.edu/andru/papers/sosp01/spp-tr.pdf>.

A. New proof

A.1. Proofs for Noninterference

Before proving noninterference, we prove few supporting lemmas required for proving type preservation.

Lemma 8 (FLAC implies FLAM). *Let $\mathcal{H}(c) = \{\langle p \triangleright q \mid \perp \rightarrow \wedge \top \leftarrow \rangle \mid \langle p \triangleright q \rangle \in \Pi\}$. If $\Pi \Vdash p \triangleright q$, then $\mathcal{H}; c; \perp \rightarrow \wedge \top \leftarrow; \perp \rightarrow \wedge \top \leftarrow \Vdash p \triangleright q$.*

Proof. Proof is by induction on the derivation of the robust assumption $\Pi \Vdash p \triangleright q$. Interesting case is R-ASSUME.

Case R-ASSUME: From the premises, we have that

$$\langle p \triangleright q \rangle \in \Pi \tag{3}$$

$$\Pi; pc; \ell \Vdash \nabla(p \rightarrow) \triangleright \nabla(q \rightarrow) \tag{4}$$

From (3) and DEL, we have that

$$\mathcal{H}; c; \perp \rightarrow \wedge \top \leftarrow; \perp \rightarrow \wedge \top \leftarrow \vdash p \triangleright q$$

From [Weaken], we get that $\mathcal{H}; c; \wedge \nabla(q) \vdash p \triangleright q$. From the (4) and R-LIFT we thus have $\mathcal{H}; c; \perp \rightarrow \wedge \top \leftarrow; \perp \rightarrow \wedge \top \leftarrow \Vdash p \triangleright q$.

Case R-Static: Since $\mathcal{L} \models p \triangleright q$, we have from FLAM R-STATIC that $\mathcal{H}; c; \perp \rightarrow \wedge \top \leftarrow; \perp \rightarrow \wedge \top \leftarrow \Vdash p \triangleright q$.

Case R-ConjR: Trivial.

Case R-DisjL: Trivial.

Case R-Trans: Trivial.

Case R-Weaken: Trivial.

□

Lemma 9 (FLAM implies FLAC). *For a trust configuration such that, for all $n \neq c$, $\mathcal{H}(n) = \emptyset$, and for all $\langle p, q, \ell \rangle \in \mathcal{H}(c)$, $\ell = \perp \rightarrow \wedge \top \leftarrow$ and $\mathcal{H}; c; \perp \rightarrow \wedge \top \leftarrow; \perp \rightarrow \wedge \top \leftarrow \Vdash \nabla(p \rightarrow) \triangleright \nabla(q \rightarrow)$, if $\mathcal{H}; c; \perp \rightarrow \wedge \top \leftarrow; \perp \rightarrow \wedge \top \leftarrow \Vdash p \triangleright q$, then $\Pi \Vdash p \triangleright q$.*

Proof. (Sketch) By induction on the FLAM derivation. Without loss of generality, we assume that the derivation of $\mathcal{H}; c; \perp \rightarrow \wedge \top \leftarrow; \perp \rightarrow \wedge \top \leftarrow \Vdash p \triangleright q$ contains no applications of R-WEAKEN or R-FWD. Since all delegations are local to c , public, and trusted, they are unnecessary.

Next, observe that because delegations are public and trusted, any non-robust FLAM derivation may be lifted to a robust one by adding an application of FLAM's R-LIFT rule wherever a DEL rule occurs, and

applying relevant robust rules in place of the non-robust rules. Rules corresponding to each non-robust rule are either part of the core robust rules or have been proven admissible [8]. The rule for R-TRANS is an exception, since it requires an additional premise to be satisfied, but since the query label is always $\perp \rightarrow \wedge \top \leftarrow$ this is trivially satisfied. $\square$

Lemma 10 (Robust Assumption). *If $\Pi \Vdash pc \geq \nabla(b)$, then $\Pi \cdot \langle a \geq b \rangle \Vdash pc \geq \nabla(b)$ for any $a, b \in \mathcal{L}$.*

Proof. By inspection of the rules in Figure 5. $\square$

Lemma 11 (Robust Protection). *If $\Pi \vdash pc \sqsubseteq \tau$, then $\Pi \cdot \langle a \geq b \rangle \vdash pc \sqsubseteq \tau$ for any $a, b \in \mathcal{L}$.*

Proof. By Lemma 10 and inspection of the rules in Figure 7. $\square$

Lemma 12 (Π Extension). *If $\Pi; \Gamma; pc \vdash e : \tau$ then $\Pi \cdot \langle p \geq q \rangle; \Gamma; pc \vdash e : \tau$ for any $p, q \in \mathcal{L}$.*

Proof. By Lemma 10, Lemma 11 and inspection of the typing rules. $\square$

The following lemma is required to prove the type preservation. It says that a expression e well-typed at pc is still well-typed at a reduced pc' that satisfies the clearance.

Lemma 13 (PC Reduction). *Let $\Pi; \Gamma; pc \vdash e : \tau$. For all pc, pc' , such that $\Pi \Vdash pc' \sqsubseteq pc$ then $\Pi; \Gamma; pc' \vdash e : \tau$ holds.*

Proof. Proof is by induction on the derivation of the typing judgment.

Case VAR: Straightforward from the corresponding typing judgment.

Case UNIT: Straightforward from the corresponding typing judgment.

Case DEL: Straightforward from the corresponding typing judgment.

Case LAM: Straightforward from the corresponding typing judgment.

Case APP: Given, $\Pi; \Gamma; pc \vdash e e' : \tau$. From APP, we have

$$\Pi; \Gamma; pc \vdash e : \tau_1 \xrightarrow{pc''} \tau_2 \quad (5)$$

$$\Pi; \Gamma; pc \vdash e' : \tau_1 \quad (6)$$

$$\Pi \Vdash pc \sqsubseteq pc'' \quad (7)$$

Applying induction to the premises we have $\Pi; \Gamma; pc' \vdash e : \tau_1 \xrightarrow{pc''} \tau_2$ and $\Pi; \Gamma; pc' \vdash e' : \tau_1$.

From R-TRANS, we have $\Pi \Vdash pc' \sqsubseteq pc''$. Hence we have all the premises.

Case TLAM: Straightforward from the corresponding typing judgment.

Case TAPP: Similar to App case.

Case PAIR: Straightforward from the corresponding typing judgment.

Case UNPAIR : Straightforward from the corresponding typing judgment.

Case INJ: Straightforward from the corresponding typing judgment.

Case CASE: Straightforward from the corresponding typing judgment.

Case UNITM: Given $\Pi; \Gamma; pc \vdash \eta_\ell e : \tau$, by UNITM we have $\Pi \Vdash pc \sqsubseteq \ell$ and $\Pi; \Gamma; pc \vdash e : \tau$. By the induction hypothesis, we have $\Pi; \Gamma; pc' \vdash e : \tau$, and since $\Pi \Vdash pc' \sqsubseteq pc$, then by R-TRANS, we have $\Pi \Vdash pc' \sqsubseteq \ell$. Therefore by UNITM we have $\Pi; \Gamma; pc' \vdash \eta_\ell e : \tau$.

Case SEALED: Straightforward from the corresponding typing judgment.

Case BINDM: Given $\Pi; \Gamma; pc \vdash \text{bind } x = e \text{ in } e' : \tau$, by BINDM we have

$$\Pi; \Gamma; pc \vdash e : \ell \text{ says } \tau' \quad (8)$$

$$\Pi; \Gamma, x : \tau'; pc \sqcup \ell \vdash e' : \tau \quad (9)$$

$$\Pi \vdash pc \sqcup \ell \sqsubseteq \tau \quad (10)$$

Since $\Pi \Vdash pc' \sqsubseteq pc$. By the monotonicity of join with respect to $\sqsubseteq$, we also have $\Pi \Vdash pc' \sqcup \ell \sqsubseteq pc \sqcup \ell$. Therefore, by the induction hypothesis applied to 8 and 9, we have

$$\Pi; \Gamma; pc' \vdash e : \ell \text{ says } \tau' \quad (11)$$

$$\Pi; \Gamma, x : \tau'; pc' \sqcup \ell \vdash e' : \tau \quad (12)$$

and by R-TRANS we get $\Pi \vdash pc \sqcup \ell \sqsubseteq \tau$. Then via BINDM we get

$$\Pi; \Gamma; pc' \vdash \text{bind } x = e \text{ in } e' : \tau$$

Case ASSUME: Given, $\Pi; \Gamma; pc \vdash \text{assume } e \text{ in } e' : \tau$, by ASSUME we have

$$\Pi; \Gamma; pc \vdash e : (p \geq q) \quad (13)$$

$$\Pi \Vdash pc \geq \nabla(q) \quad (14)$$

$$\Pi \Vdash \nabla(p^\rightarrow) \geq \nabla(q^\rightarrow) \quad (15)$$

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash e' : \tau \quad (16)$$

Applying induction hypothesis to (13) and (16) we have $\Pi; \Gamma; pc' \vdash e : (p \geq q)$ and $\Pi, \langle p \geq q \rangle; \Gamma; pc' \vdash e' : \tau$. Since $\Pi \Vdash pc' \sqsubseteq pc$, we have $\Pi \Vdash pc' \geq \nabla(q)$. Combining, we have all the premises for ASSUME and thus $\Pi; \Gamma; pc' \vdash \text{assume } e \text{ in } e' : \tau$

Case WHERE: Given, $\Pi; \Gamma; pc \vdash v \text{ where } e : \tau$, by WHERE we have

$$\Pi; \Gamma; pc \vdash v : (p \geq q) \quad (17)$$

$$\Pi \Vdash \overline{pc} \geq \nabla(q) \quad (18)$$

$$\Pi \Vdash \nabla(p^\rightarrow) \geq \nabla(q^\rightarrow) \quad (19)$$

$$\Pi; \Gamma; pc \vdash e : \tau \quad (20)$$

Applying induction hypothesis to (17) and (20), we have $\Pi; \Gamma; pc' \vdash v : (p \geq q)$ and $\Pi, \langle p \geq q \rangle; \Gamma; pc' \vdash e : \tau$. Then by WHERE, we have $\Pi; \Gamma; pc' \vdash v \text{ where } e : \tau$.

Case BRACKET: The premise $\Pi \Vdash H^\pi \sqcup pc \sqsubseteq pc''$ implies $\Pi \Vdash H^\pi \sqcup pc' \sqsubseteq pc''$.

Case BRACKET-VALUES: Applying induction to the premises gives the required conclusion.

□

Lemma 14 (Values PC). *If $\Pi; \Gamma; pc \vdash w : \tau$, then $\Pi; \Gamma; pc' \vdash w : \tau$ for any pc' .*

Proof. By induction on the typing derivation of w . Observe that only APP, CASE, UNITM, BINDM, and ASSUME contain premises that constrain typing based on the judgment pc , and these rules do not apply to w terms. $\square$

Lemma 15 (Variable Substitution). *If $\Pi; \Gamma, x : \tau'; pc \vdash e : \tau$ and $\Pi; \Gamma; pc \vdash w : \tau'$, then $\Pi; \Gamma; pc \vdash e[x \mapsto w] : \tau$.*

Proof. Proof is by induction on the typing derivation of e . Observe that by Lemma 14, we have $\Pi; \Gamma; pc' \vdash w : \tau'$ for any pc' . Therefore, each inductive case follows by straightforward application of the induction hypothesis. $\square$

Lemma 16 (Type Substitution). *Let τ' be well-formed in Γ, X, Γ' . If $\Pi; \Gamma, X, \Gamma'; pc \vdash e : \tau$ then $\Pi; \Gamma, \Gamma'[X \mapsto \tau']; pc \vdash e[X \mapsto \tau'] : \tau[X \mapsto \tau']$.*

Proof. Proof is by the induction on the typing derivation of $\Pi; \Gamma, X, \Gamma'; pc \vdash e : \tau$. $\square$

Lemma 17 (Projection Preserves Types). *If $\Pi; \Gamma; pc \vdash e : \tau$, then $\Pi; \Gamma; pc \vdash [e]_i : \tau$ for $i = \{1, 2\}$.*

Proof. Proof is by induction on the typing derivation of $\Pi; \Gamma; pc \vdash e : \tau$. The interesting case is $e = (e_1 \mid e_2)$. By BRACKET, we have $\Pi; \Gamma; pc' \vdash e_i : \tau$ for some pc' such that $\Pi \Vdash (H^\pi \sqcup pc) \sqsubseteq pc'$. Therefore, by Lemma 13, we have $\Pi; \Gamma; pc \vdash e_i : \tau$. $\square$

Lemma 2 (Stuck expressions). *If e gets stuck then $[e]_i$ is stuck for some $i \in \{1, 2\}$.*

Proof. We prove by induction on the structure of e .

Case w : No reduction rules apply to terms in the syntactic category w (including $(w \mid w')$). Hence $[x]_i$ is stuck as well.

Case x : No reduction rules apply to a variable. Hence $[x]_i$ is stuck as well.

Case $(e_1 \mid e_2)$: By B-STEP, e is only gets stuck is if both e_1 and e_2 get stuck.

Case $e e'$: Since $e e'$ is stuck, then B-APP, W-APP, E-APP are not applicable. It follows that either (1) e is not of the form $(w \mid w')$, w where v , or $\lambda(x : \tau)[pc'] \cdot e$ or (2) e has the form $\lambda(x : \tau)[pc'] \cdot e$, but e' is stuck. For the first case, $[e]_i$ is also not of the form $(w \mid w')$, w where v , or $\lambda(x : \tau)[pc'] \cdot e$, so $[e e']_i$ is also stuck. For the second case, applying the induction hypothesis gives us that $[e']_i$ is stuck for some $i \in \{1, 2\}$, so $[\lambda(x : \tau)[pc'] \cdot e e']_i$ is stuck for the same i .

Case $e \tau$: Since $e e'$ is stuck, then B-TAPP, W-TAPP, E-TAPP are not applicable. It follows that e is not of the form $(w \mid w')$, w where v , or $\Lambda X[pc'] \cdot e$. Therefore, $[e]_i$ is also not of the form $(w \mid w')$, w where v , or $\lambda(x : \tau)[pc'] \cdot e$, so $[e e']_i$ is also stuck.

Case $\eta_\ell e$: Since $\eta_\ell e$ is stuck, then E-UNITM is not applicable, so e does not have the form w . Therefore, E-UNITM is also not applicable to $[\eta_\ell e]_i$. Therefore, e must be stuck. Applying induction hypothesis, it follows that $[e]_i$ is also stuck and so $[\eta_\ell e]_i$ is also stuck.

Case $\text{proj}_j e$: Similar to the above case.

Case $\text{inj}_j e$: Similar to the above case.

Case $\langle e, e \rangle$: Similar to the above case.

Case $\text{case } e \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2$: Since B-CASE, W-CASE, and E-CASE are not applicable, it follows that e is not of the form $(w \mid w')$, w where v , or $\text{inj}_j v$. It follows that $[\text{case } e \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2]_i$ is also stuck.

Case $\text{bind } x = v \text{ in } e'$: Similar to the above case.

Case assume e in e' : Similar to the above case.

Case e where v : Similar to the above case.

□

Using the above lemmas, we are now ready to prove type preservation.

Lemma 4 (Subject Reduction). *Let $\Pi; \Gamma; pc \vdash e : \tau$. If $e \longrightarrow e'$ then $\Pi; \Gamma; pc \vdash e' : \tau$.*

Proof. **Case E-APP:** Given $e = (\lambda(x:\tau)[pc'].e) w$ and $\Pi; \Gamma; pc \vdash (\lambda(x:\tau)[pc'].e) v : \tau_2$. From APP we have,

$$\Pi; \Gamma, x : \tau'; pc' \vdash e : \tau \quad (21)$$

$$\Pi; \Gamma; pc \vdash w : \tau' \quad (22)$$

$$\Pi \Vdash pc' \sqsubseteq pc \quad (23)$$

Therefore, via PC reduction (Lemma 13) and Variable Substitution (Lemma 15), we have that $\Pi; \Gamma; pc' \vdash e[x \mapsto w] : \tau$.

Case E-TAPP: Similar to above case, but using Lemma 16.

Case E-CASE1 : Given $e = \text{case } (\text{inj}_1 w) \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2$ and $e' = e_1[x \mapsto v]$. Also, $\Pi; \Gamma; pc \vdash \text{case } (\text{inj}_1 w) \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2 : \tau$. From the premises we have, $\Pi; \Gamma; pc \vdash \text{inj}_1 v : \tau' + \tau''$ and $\Pi; \Gamma, x : \tau'; pc \vdash e_1 : \tau$. From INJ, we have $\Pi; \Gamma; pc \vdash v : \tau'$. Invoking variable substitution lemma (Lemma 15), we have $\Pi; \Gamma; pc \vdash e_1[x \mapsto w] : \tau$.

Case E-CASE2 : Similar to above.

Case E-UNITM: Given $e = \eta_\ell w$ and $e' = \bar{\eta}_\ell w$. Also, $\Pi; \Gamma; pc \vdash \eta_\ell w : \ell \text{ says } \tau$. From the premises it follows that $\Pi; \Gamma; pc \vdash \bar{\eta}_\ell w : \ell \text{ says } \tau$.

Case E-BINDM: Given $e = \text{bind } x = w \text{ in } e'$ and $e' = e'[x \mapsto w]$. Also, $\Pi; \Gamma; pc \vdash \text{bind } x = w \text{ in } e' : \tau$. From the premises, we have the following:

$$\Pi; \Gamma; pc \vdash w : \tau' \quad (24)$$

$$\Pi; \Gamma, x : \tau'; pc \sqcup \ell \vdash e' : \tau \quad (25)$$

$$\Pi \vdash pc \sqcup \ell \sqsubseteq \tau \quad (26)$$

$$\Pi \Vdash p \geq pc \quad (27)$$

We have to prove that $\Pi; \Gamma; pc \vdash e'[x \mapsto w] : \tau$. Since we have that $\Pi \Vdash p \geq pc$, applying PC reduction (Lemma 13) to the premise (25), we have $\Pi; \Gamma, x : \tau'; pc \vdash e' : \tau$.

Invoking variable substitution lemma (Lemma 15), we thus have $\Pi; \Gamma; pc \vdash e'[x \mapsto w] : \tau$.

Case E-ASSUME: Given $e = \text{assume } \langle p \geq q \rangle \text{ in } e'$ and $e' = e'$ where $\langle p \geq q \rangle$. Also, $\Pi; \Gamma; pc \vdash \text{assume } \langle p \geq q \rangle \text{ in } e' : \tau$. From ASSUME, we have

$$\Pi; \Gamma; pc \vdash \langle p \geq q \rangle : (p \geq q) \quad (28)$$

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash e' : \tau \quad (29)$$

$$\Pi \Vdash pc \geq \nabla(q) \quad (30)$$

$$\Pi \Vdash \nabla(pt^*) \geq \nabla(q^*) \quad (31)$$

We need to prove:

$$\Pi; \Gamma; pc \vdash e' \text{ where } \langle p \geq q \rangle : \tau$$

Comparing with the given premises, we already have the required premises.

$$\Pi; \Gamma; pc \vdash \langle p \geq q \rangle : (p \geq q) \quad (32)$$

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash e : \tau \quad (33)$$

$$\Pi \Vdash \overline{pc} \geq \nabla(q) \quad (34)$$

$$\Pi \Vdash \nabla(p^\rightarrow) \geq \nabla(q^\rightarrow) \quad (35)$$

Hence proved.

Case E-EVAL: For $e = E[e]$ and $e' = E[e']$ where $\Pi; \Gamma; pc \vdash E[e] : \tau$, we have $\Pi'; \Gamma'; pc' \vdash e : \tau'$ for some pc', τ', Π' , and Γ' such that $\Pi' \supseteq \Pi$, $\Gamma' \supseteq \Gamma$. By the induction hypothesis we have $\Pi'; \Gamma'; pc' \vdash e' : \tau'$. Observe that with the exception of SEALED and WHERE, the premises of all typing rules use terms from the syntactic category e . Therefore if a derivation for $\Pi; \Gamma; pc \vdash E[e] : \tau$ exists, it must be the case that derivation for $\Pi; \Gamma; pc \vdash E[e'] : \tau$ exists where the derivation of $\Pi'; \Gamma'; pc' \vdash e : \tau'$ is replaced with $\Pi'; \Gamma'; pc' \vdash e' : \tau'$. Rules SEALED and WHERE have premises that use terms from the syntactic category v , but since these are fully evaluated, e cannot be equal to a v term since no e' exists such that $v \longrightarrow^* e'$.

Case W-APP: Given $e = (w \text{ where } \langle p \geq q \rangle) e$ and $e' = (w e) \text{ where } \langle p \geq q \rangle$. We have to prove that

$$\Pi; \Gamma; pc \vdash (w e) \text{ where } \langle p \geq q \rangle : \tau$$

From APP we have:

$$\Pi; \Gamma; pc \vdash w \text{ where } \langle p \geq q \rangle : \tau_1 \xrightarrow{pc'} \tau \quad (36)$$

$$\Pi; \Gamma; pc \vdash e : \tau_1 \quad (37)$$

$$\Pi \Vdash pc \sqsubseteq pc' \quad (38)$$

$$(39)$$

Rule WHERE gives us the following:

$$\Pi; \Gamma; pc \vdash \langle p \geq q \rangle : (p \geq q) \quad (40)$$

$$\Pi \Vdash \overline{pc} \geq \nabla(q) \quad (41)$$

$$\Pi \Vdash \nabla(p^\rightarrow) \geq \nabla(q^\rightarrow) \quad (42)$$

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash w : \tau_1 \xrightarrow{pc'} \tau \quad (43)$$

We now want to show that e' is well typed via WHERE. The key premise is to show that the subexpression $(w e)$ is well-typed via APP. That is,

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash w e : \tau \quad (44)$$

Applying Lemma 12 (extending delegation contexts for well-typed terms) to (37) and Lemma 10 (extending delegation contexts for assumptions) to (38), we have:

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash e : \tau_1 \quad (45)$$

$$\Pi, \langle p \geq q \rangle \Vdash pc \sqsubseteq pc' \quad (46)$$

Combining with (43), we have (44) which when combined with remaining premises ((40), (41) and (42)) give us $\Pi; \Gamma; pc \vdash (w \ e)$ where $\langle p \geq q \rangle : \tau$.

Case W-TAPP: Given $e = (w \text{ where } \langle p \geq q \rangle) \tau$ and $e' = (w \ \tau')$ where $\langle p \geq q \rangle$. We have to prove that

$$\Pi; \Gamma; pc \vdash (v \ \tau') \text{ where } \langle p \geq q \rangle : \tau[X \mapsto \tau']$$

From TAPP we have:

$$\Pi; \Gamma; pc \vdash w \text{ where } \langle p \geq q \rangle : \forall X[pc']. \tau \quad (47)$$

$$\Pi \Vdash pc \sqsubseteq pc' \quad (48)$$

Rule WHERE gives us the following:

$$\Pi; \Gamma; pc \vdash \langle p \geq q \rangle : (p \geq q) \quad (49)$$

$$\Pi \Vdash \overline{pc} \geq \nabla(q) \quad (50)$$

$$\Pi \Vdash \nabla(p^-) \geq \nabla(q^-) \quad (51)$$

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash v : \forall X[pc']. \tau \quad (52)$$

We now want to show that e' is well typed via WHERE. The key premise is to show that the subexpression $(v \ \tau')$ is well-typed via TAPP. That is,

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash w \ \tau' : \tau[X \mapsto \tau'] \quad (53)$$

Applying Lemma 10 (extending delegation context for well-typed terms) to (48), we get:

$$\Pi, \langle p \geq q \rangle \Vdash pc \sqsubseteq pc' \quad (54)$$

Combining with (52), we have (53) which when combined with remaining premises ((49), (50) and (51)) give $\Pi; \Gamma; pc \vdash (w \ \tau') \text{ where } \langle p \geq q \rangle : \tau[X \mapsto \tau']$.

Case W-UNPAIR: Given $e = \text{proj}_i (\langle w_1, w_2 \rangle \text{ where } \langle p \geq q \rangle)$ and $e' = (\text{proj}_i \langle w_1, w_2 \rangle) \text{ where } \langle p \geq q \rangle$. We have to prove that

$$\Pi; \Gamma; pc \vdash (\text{proj}_i \langle w_1, w_2 \rangle) \text{ where } \langle p \geq q \rangle : \tau_i$$

From UNPAIR, we have:

$$\Pi; \Gamma; pc \vdash \langle w_1, w_2 \rangle \text{ where } \langle p \geq q \rangle : \tau_1 \times \tau_2 \quad (55)$$

From (55) and WHERE, we have:

$$\Pi; \Gamma; pc \vdash \langle p \geq q \rangle : (p \geq q) \quad (56)$$

$$\Pi \Vdash \overline{pc} \geq \nabla(q) \quad (57)$$

$$\Pi \Vdash \nabla(p^\rightarrow) \geq \nabla(q^\rightarrow) \quad (58)$$

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash \langle w_1, w_2 \rangle : \tau_1 \times \tau_2 \quad (59)$$

From (59) and UNPAIR, we have: $\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash \text{proj}_i \langle w_1, w_2 \rangle : \tau_i$. Combining with remaining premises ((56) to (59)) we have $\Pi; \Gamma; pc \vdash \text{proj}_i \langle w_1, w_2 \rangle$ where $\langle p \geq q \rangle : \tau_i$.

Case W-CASE: Given $e = \text{case } (w \text{ where } \langle p \geq q \rangle) \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2$ and $e' = (\text{case } w \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2)$ where $\langle p \geq q \rangle : \tau$. We have to prove that

$$\Pi; \Gamma; pc \vdash (\text{case } w \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2) \text{ where } \langle p \geq q \rangle : \tau$$

From CASE we have:

$$\Pi; \Gamma; pc \vdash w \text{ where } \langle p \geq q \rangle : \tau_1 + \tau_2 \quad (60)$$

$$\Pi \vdash pc \sqcup \ell \sqsubseteq \tau \quad (61)$$

$$\Pi \Vdash pc \sqsubseteq \ell \quad (62)$$

$$\Pi; \Gamma, x : \tau_1; pc \sqcup \ell \vdash e_1 : \tau \quad (63)$$

$$\Pi; \Gamma, x : \tau_2; pc \sqcup \ell \vdash e_2 : \tau \quad (64)$$

From (60) and rule WHERE, we get the following:

$$\Pi; \Gamma; pc \vdash \langle p \geq q \rangle : (p \geq q) \quad (65)$$

$$\Pi \Vdash \overline{pc} \geq \nabla(q) \quad (66)$$

$$\Pi \Vdash \nabla(p^\rightarrow) \geq \nabla(q^\rightarrow) \quad (67)$$

$$\Pi; \Gamma; pc \vdash w : \tau_1 + \tau_2 \quad (68)$$

The key premise to prove is that:

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash (\text{case } w \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2) : \tau$$

which follows from (68) and extending delegations in equations (61) to (64) (Lemma 12). Combining with remaining premises, we have $\Pi; \Gamma; pc \vdash (\text{case } w \text{ of } \text{inj}_1(x). e_1 \mid \text{inj}_2(x). e_2) \text{ where } \langle p \geq q \rangle : \tau$.

Case W-BINDM: Given $e = \text{bind } x = (w \text{ where } \langle p \geq q \rangle) \text{ in } e$ and $e' = (\text{bind } x = w \text{ in } e) \text{ where } \langle p \geq q \rangle$. We have to prove that

$$\Pi; \Gamma; pc \vdash (\text{bind } x = w \text{ in } e) \text{ where } \langle p \geq q \rangle : \tau$$

From BINDM we have:

$$\Pi; \Gamma; pc \vdash (w \text{ where } \langle p \geq q \rangle) : \ell \text{ says } \tau' \quad (69)$$

$$\Pi; \Gamma, x : \tau'; pc \sqcup \ell \vdash e : \tau \quad (70)$$

$$\Pi \vdash pc \sqcup \ell \sqsubseteq \tau \quad (71)$$

From (69) and rule WHERE, we get the following:

$$\Pi; \Gamma; pc \vdash \langle p \geq q \rangle : (p \geq q) \quad (72)$$

$$\Pi \Vdash \overline{pc} \geq \nabla(q) \quad (73)$$

$$\Pi \Vdash \nabla(p^\rightarrow) \geq \nabla(q^\rightarrow) \quad (74)$$

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash w : \ell \text{ says } \tau' \quad (75)$$

We now want to show that e' is well typed via WHERE. That is, we need the following premises,

$$\Pi, \langle p \geq q \rangle; \Gamma; pc \vdash \text{bind } x = w \text{ in } e : \tau \quad (76)$$

$$\Pi; \Gamma; pc \vdash \langle p \geq q \rangle : (p \geq q) \quad (77)$$

$$\Pi \Vdash \overline{pc} \geq \nabla(q) \quad (78)$$

$$\Pi \Vdash \nabla(p^\rightarrow) \geq \nabla(q^\rightarrow) \quad (79)$$

Extending the delegation context (Lemma 12) in the premises (70), (71) and from (75) we have (76).

We already have (77) from (72); (78) from (73); (79) from (74). Combining, we have $\Pi; \Gamma; pc \vdash (\text{bind } x = w \text{ in } e) \text{ where } \langle p \geq q \rangle : \tau$.

Case W-ASSUME: Given $e = \text{assume } w \text{ where } \langle a \geq b \rangle \text{ in } e$ and $e' = \text{assume } w \text{ in } e \text{ where } \langle a \geq b \rangle$.

From ASSUME, we have

$$\Pi; \Gamma; pc \vdash v \text{ where } \langle a \geq b \rangle : (r \geq s) \quad (80)$$

$$\Pi, \langle r \geq s \rangle; \Gamma; pc \vdash e : \tau \quad (81)$$

$$\Pi \Vdash pc \geq \nabla(s) \quad (82)$$

$$\Pi \Vdash \nabla(r^\rightarrow) \geq \nabla(s^\rightarrow) \quad (83)$$

Expanding the first premise using WHERE, we have

$$\Pi; \Gamma; pc \vdash \langle a \geq b \rangle : (a \geq b) \quad (84)$$

$$\Pi, \langle a \geq b \rangle; \Gamma; pc \vdash v : (r \geq s) \quad (85)$$

$$\Pi \Vdash \overline{pc} \geq \nabla(b) \quad (86)$$

$$\Pi \Vdash \nabla(a^\rightarrow) \geq \nabla(b^\rightarrow) \quad (87)$$

We want to show

$$\Pi; \Gamma; pc \vdash \langle a \geq b \rangle : (a \geq b) \quad (88)$$

$$\Pi, \langle a \geq b \rangle; \Gamma; pc \vdash \text{assume } v \text{ in } e : \tau \quad (89)$$

$$\Pi \Vdash \overline{pc} \geq \nabla(b) \quad (90)$$

$$\Pi \Vdash \nabla(a^\rightarrow) \geq \nabla(b^\rightarrow) \quad (91)$$

Extending delegation context (Lemma 12) in the premises (81), (82), (87) and combining with topmost premise we have the (89). Remaining premises follow from (84), (86) and (87).

Case B-STEP: Given $e = (e_1 \mid e_2)$ and $e' = (e'_1 \mid e'_2)$. Also $\Pi; \Gamma; \Theta; p; pc \vdash (e_1 \mid e_2) : \tau$. We have to prove

$$\Pi; \Gamma; \Theta; p; pc \vdash (e'_1 \mid e'_2) : \tau$$

Without loss of generality, let $i = 1$. Thus from the premises of B-STEP, we have $e_1 \longrightarrow e'_1$ and $e'_2 = e_2$. Since sealed values can not take a step, inverting the well-typedness of bracket is only possible through BRACKET and not through BRACKET-VALUES. From BRACKET, we have

$$\Pi; \Gamma; \Theta; p; pc' \vdash e_1 : \tau \quad (92)$$

$$\Pi; \Gamma; \Theta; p; pc' \vdash e_2 : \tau \quad (93)$$

$$\Pi \Vdash (H^\pi \sqcup pc) \sqsubseteq pc' \quad (94)$$

$$\Pi \vdash H^\pi \sqsubseteq \tau \quad (95)$$

Since (92) holds, applying induction to the premise $e_1 \longrightarrow e'_1$, we have that $\Pi; \Gamma; \Theta; p; pc' \vdash e'_1 : \tau$. Combining with remaining premises ((93) to (95)) we thus have that $\Pi; \Gamma; \Theta; p; pc \vdash (e'_1 \mid e'_2) : \tau$.

Case B-APP: Given $e = (w_1 \mid w_2) w'$ and $e' = (w_1 [w']_1 \mid w_2 [w']_2)$. Also given that $\Pi; \Gamma; pc \vdash (w_1 \mid w_2) w' : \tau_2$ is well-typed, from APP, we have the following:

$$\Pi; \Gamma; pc \vdash (w_1 \mid w_2) : \tau_1 \xrightarrow{pc''} \tau_2 \quad (96)$$

$$\Pi; \Gamma; pc \vdash w' : \tau_1 \quad (97)$$

$$\Pi \Vdash pc \sqsubseteq pc'' \quad (98)$$

Thus from BRACKET-VALUES, we have $\Pi \Vdash H^\pi \sqsubseteq \tau_1 \xrightarrow{pc''} \tau_2$. From P-FUN, we thus have

$$\Pi \Vdash H^\pi \sqsubseteq \tau_2 \quad (99)$$

$$\Pi \Vdash H^\pi \sqsubseteq pc'' \quad (100)$$

We need to prove

$$\Pi; \Gamma; pc \vdash (w_1 [w']_1 \mid w_2 [w']_2) : \tau_2$$

That is we need the following premises of BRACKET.

$$\Pi; \Gamma; pc' \vdash w_1 [w']_1 : \tau_2 \quad (101)$$

$$\Pi; \Gamma; pc' \vdash w_2 [w']_2 : \tau_2 \quad (102)$$

$$\Pi \Vdash H^\pi \sqcup pc \sqsubseteq pc' \quad (103)$$

$$\Pi \Vdash H^\pi \sqsubseteq \tau_2 \quad (104)$$

Let $pc' = pc''$. We have (103) from (98) and (100). We already have (104) from (99). To prove (101), we need the following premises:

$$\Pi; \Gamma; pc' \vdash w_1 : \tau_1 \xrightarrow{pc''} \tau_2 \quad (105)$$

$$\Pi; \Gamma; pc' \vdash [w']_1 : \tau_2 \quad (106)$$

$$\Pi \Vdash pc' \sqsubseteq pc'' \quad (107)$$

The last premise (107) holds trivially (from reflexivity). Applying Lemma 14 (values can be typed under any pc) to (96) we have (105). Applying Lemma 14 (values can be typed under any pc) and Lemma (17) (projection preserves typing) to (97) we have (106). Thus from APP, we have (101). Similarly, (102) holds. Hence proved.

Case B-TAPP: Similar to above (B-APP) case.

Case B-UNPAIR: Given $e = \text{proj}_i (\langle w_{11}, w_{12} \rangle \mid \langle w_{21}, w_{22} \rangle)$ and $e' = (w_{1i} \mid w_{2i})$. Also $\Pi; \Gamma; pc \vdash \text{proj}_i (\langle w_{11}, w_{12} \rangle \mid \langle w_{21}, w_{22} \rangle) : \tau_i$. We have to prove

$$\Pi; \Gamma; pc \vdash (w_{1i} \mid w_{2i}) : \tau_i$$

From UNPAIR, we have:

$$\Pi; \Gamma; pc \vdash (\langle w_{11}, w_{12} \rangle \mid \langle w_{21}, w_{22} \rangle) : \tau_1 \times \tau_2 \quad (108)$$

Since they are already values, they can be inverted using BRACKET-VALUES. This approach is more conservative.

$$\Pi; \Gamma; pc \vdash \langle w_{11}, w_{12} \rangle : \tau_1 \times \tau_2 \quad (109)$$

$$\Pi; \Gamma; pc \vdash \langle w_{21}, w_{22} \rangle : \tau_1 \times \tau_2 \quad (110)$$

$$\Pi \vdash H^\pi \sqsubseteq \tau_1 \times \tau_2 \quad (111)$$

From (109), (110) and UNPAIR, we have $\Pi; \Gamma; pc \vdash w_{1i} : \tau_i$ and $\Pi; \Gamma; pc \vdash w_{2i} : \tau_i$ for $i = \{1, 2\}$. From (111) and P-PAIR, we have $\Pi \vdash H^\pi \sqsubseteq \tau_i$. Combining with other premises, from BRACKET-VALUES, we have $\Pi; \Gamma; pc \vdash (w_{1i} \mid w_{2i}) : \tau_i$.

Case B-BINDM: Given $e = \text{bind } x = (\bar{\eta}_\ell w_1 \mid \bar{\eta}_\ell w_2) \text{ in } e$. Here $e' = (\text{bind } x = \bar{\eta}_\ell w_1 \text{ in } [e]_1 \mid \text{bind } x = \bar{\eta}_\ell w_2 \text{ in } [e]_2)$. Also $\Pi; \Gamma; pc \vdash \text{bind } x = (\bar{\eta}_\ell w_1 \mid \bar{\eta}_\ell w_2) \text{ in } e : \tau$. From BINDM, we have

$$\Pi; \Gamma; pc \vdash (\bar{\eta}_\ell w_1 \mid \bar{\eta}_\ell w_2) : \ell \text{ says } \tau' \quad (112)$$

$$\Pi; \Gamma, x : \tau'; pc \sqcup \ell \vdash e : \tau \quad (113)$$

$$\Pi \vdash pc \sqcup \ell \sqsubseteq \tau \quad (114)$$

From (112) and BRACKET-VALUES, we have

$$\Pi; \Gamma; pc \vdash \bar{\eta}_\ell w_1 : \ell \text{ says } \tau' \quad (115)$$

$$\Pi; \Gamma; pc \vdash \bar{\eta}_\ell w_2 : \ell \text{ says } \tau' \quad (116)$$

$$\Pi \Vdash H^\pi \sqsubseteq \ell \quad (117)$$

We have to prove that

$$\Pi; \Gamma; pc \vdash (\text{bind } x = \bar{\eta}_\ell w_1 \text{ in } [e]_1 \mid \text{bind } x = \bar{\eta}_\ell w_2 \text{ in } [e]_2) : \tau$$

For some $\hat{pc}$ we need the following premises to satisfy BRACKET:

$$\Pi; \Gamma; \hat{pc} \vdash \text{bind } x = \bar{\eta}_\ell w_1 \text{ in } [e]_1 : \tau \quad (118)$$

$$\Pi; \Gamma; \hat{pc} \vdash \text{bind } x = \bar{\eta}_\ell w_2 \text{ in } [e]_2 : \tau \quad (119)$$

$$\Pi \Vdash (H^\pi \sqcup pc) \sqsubseteq \hat{pc} \quad (120)$$

$$\Pi \vdash H^\pi \sqsubseteq \tau \quad (121)$$

A natural choice for $\hat{pc}$ is $pc \sqcup \ell$. From Lemma 14 (values can be typed under any pc), we have

$$\Pi; \Gamma; \hat{pc} \vdash \bar{\eta}_\ell w_i : \tau'$$

Applying Lemma 17 (bracket projection preserves typing) to (113), we have

$$\Pi; \Gamma, x : \tau'; \hat{pc} \vdash [e]_i : \tau$$

From BINDM, we therefore have (118) and (119). Applying R-TRANS to (117) and (114), we have (121). Thus we have all required premises.

Case B-CASE: Does not occur. Not well-typed.

Case B-ASSUME: Does not occur. Not well-typed.

□