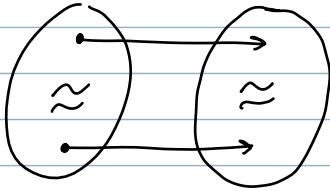


OPLSS/12

Lecture 4 - Downgrading & other future directions

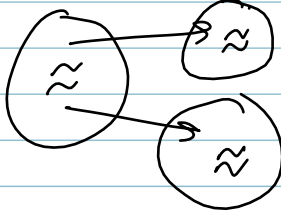
Noninterference: execution preserves equivalence



But: real systems need to release some information as part of their intended function.

- password checker: passwords
- reviewing system: reviews (eventually)
- distributed games: opponent actions (integrity)

Result: execution sometimes does this.

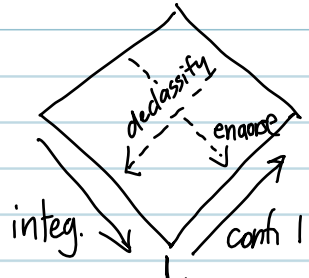


Noninterference $\Leftrightarrow$ 0 information flow

Need to be able to enforce it, but need information flow control, not just prevention.

Can add downgrading to system

declassify: lower confidentiality
integrity: ~~raise~~ integrity



Problem: Justification for downgrading is application-specific.

Ex. 1 Password checking.

```
if (declassify(guess == security, H to L)) {  
    login := true;  
}
```

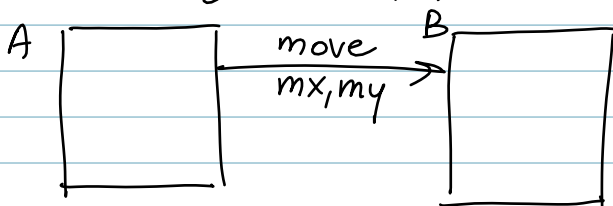
Justification: adversary learns very little.

Ex2. Auction.

```
if (auction-done)  
    released_bid = declassify(bid, H to L)
```

Justification: information no longer confidential.

Ex3 Distributed game (e.g., Battleship)



```
mx, my : A (integrity)  
if (legal_move(mx, my, board))  
endorse (mx, my, A to B)  
if (legal_move(mx, my, board)) {  
    :  
    // use mx, my at B integrity  
}
```

Justification: adversary gets to make legal moves.

Many justifications $\Rightarrow$ many different approaches to relaxing noninterference.

Sabelfeld & Sands, "Dimensions and Principles" of declassification: categorize mechanism

- who - who decides to declassify
- what - which information, or how much, is declassified
- when - ~~under what conditions~~ temporally related policies
- where - using notions of locality

All have been explored.

An early idea: restrict declassification to trusted / authorized code - selective declassification [SOSP'97] - a "who" approach.

- need to have labels that talk about principals - "decentralized labels"
- Problem: untrusted code / agents can still influence trusted code.

A "what" approach: delimited release

- specifies which expressions may be downgraded; semantic condition says $s_1 \sim_L s_2$ if s_1, s_2 equal at those expressions (& at low)

Relaxed noninterference generalizes this [Li & Zdancewicz] to specify types that say what computations may be declassified (password: $\lambda p. p = \text{guess}$)

A "when" approach: Steve Chong's
downgrading policies $l_1 \xrightarrow{c} l_2$

"flow allowed if condition c holds"

e.g., bid: $H \xrightarrow{\text{auction and}} L$

Logic c unspecified; Dimitrova et al [VMCA11]
explore instantiating c with temporal logic,
using model checking for enforcement.
Practical?

S&S also identify some useful principles
declassification mechanisms should aim for,
such as:

Semantic consistency:

- semantics-preserving program changes
shouldn't affect security judgment.
(undecidable, but a good goal)
- Want extensional security that
depends on behavior, not intensional
security based on details of code.

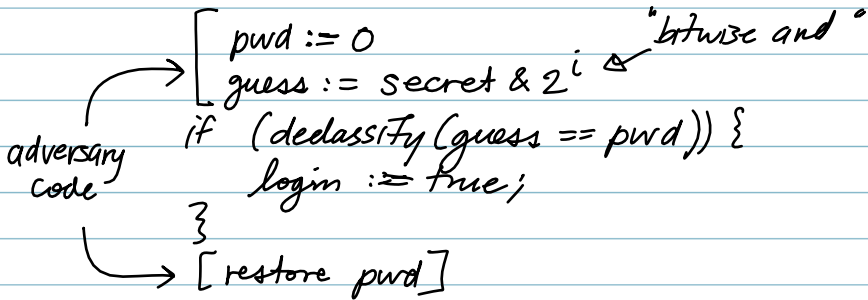
Non-occlusion



- Use of dedeclassify should not hide
other leaks. E.g. in delimited
release; $l_1 := h$; $l_2 := \text{dedeclassify}(h)$
is semantically "secure"

Laundering

- aspect of non-declassification
- Adversary can exploit declassify to create unintended information release.
- Particularly an issue in distributed systems where adversary may supply some of the code.

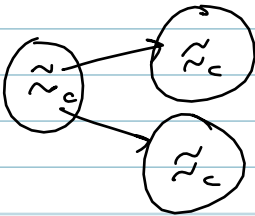


Integrity can affect confidentiality — not fully dual.

To prevent: adversary should not be able to affect

- what is declassified
- whether declassification happens
- whether endorsement happens

} rob! declass



with or without adversary.

Extensionally:

Let $s[a]$ be state s with low-integrity parts replaced by adversary "fair attack" a .

Robust declassification:

$$\forall s, s', a, a'. s[a] \approx_L s'[a] \Rightarrow s[a'] \approx s'[a']$$

↑
termination-sensitive

↑
-insensitive

- "No attack is worse than the dummy attack"
- Extensional, enforceable by a type system

To connect confidentiality & integrity
need to map integrity to the confidentiality
it enforces

Eg if $l = (p_c, p_i)$, define:

$$\text{enforces}(l) = (p_i, \perp)$$

$\Gamma, pc \vdash e : l_e \quad l_e \sqsubseteq l_1 \quad l_1 \sqsubseteq l_2 \wedge \text{enforces}(pc)$	$\Gamma, pc \vdash \text{declassify}(e, l_1 \text{ to } l_2) : l_2$
--	---

See Askarov et al (LMCS) for more precise
(progress-sensitive) security conditions.

- Most interesting directions seem to be
in the "what" and "when" dimensions.

One more direction: NI too weak!

- crypto device should forget keys
- voting machine should forget voter-ballot
linkage
- legal requirements to forget information
(medicine, finance, govt, ...)

Idea: mandatory upgrading

erasure policies require information label to increase, disappear from orig. level.

Chong: policy $l_1 \nearrow l_2$ means

- info. at level l_1 must upgrade to level l_2 when condition c holds. — including all derived info.
- vs. may downgrade policies
- can combine both to achieve interesting temporal policies
- implemented in Jif_{until} using static + dynamic enforcement.
- used to implement Civitas voting system.

Open problems / interesting directions

- Downgrading policies and their connection to extensional system security requirements.
- Integrating information flow with cryptography.
- Concurrency — how to avoid internal timing channels. Starting point: low determinism (Roscoe 95)

$$[S_1] \approx [S_2] \Leftrightarrow \forall t_1 \in S_1, \forall t_2 \in S_2, t_1 \approx t_2$$



Problem: how to enforce compositionally?

- Soundness of Jif, etc.
 - complex language: objects, exceptions, dependent labels, parametric polymorphism.
- How to get developer buy-in?

costs

- added annotation burden (Jif, Flow Caml do do inference)
- challenge of mapping system requirements to labels.
- non-local errors that are hard to diagnose.

benefits

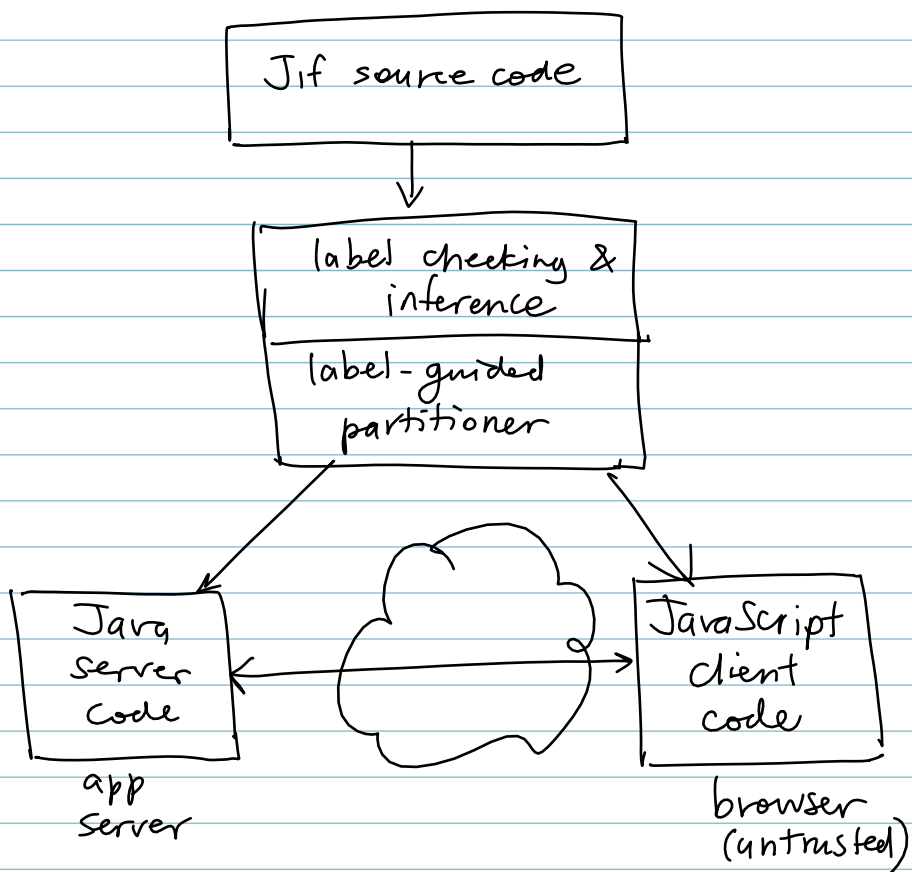
+ increased security assurance

idea: information flow enables
+ higher-level programming model.

type systems enforcement: compositional

⇒ Can transform code without breaking security.

Example: Swift: automatically securely
partitioning web applications [sosp'07]



- Partitioner keeps secrets off client, does not accept high-integrity results from client,
 - high-integrity public values may be replicated automatically.
 - e.g. input validation typically replicated to both sides
 - low latency (client side)
 - security (server side)
 - consistent

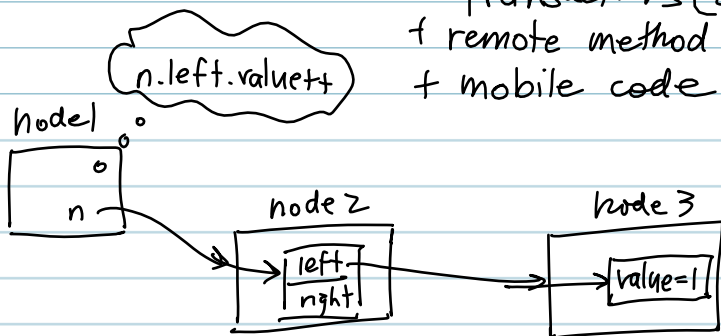
- partitioning done to minimize control transfers (= network delays)
- Fine-grained partitioning: splits individual objects, statements, according to labels. See also: work by Fournet et al. on secure partitioning

Another higher-level programming model:
Fabriz [SOSP'09, S&P'12].

Enforcing security one machine at a time makes little sense! How do we program networks as if they are computers? (+ consistency, persistence, ...)

One challenge: heterogeneous trust, generalizing beyond Swift.

Fabriz: everything is a Jif (roughly) object.
 language = Jif (Java + labels)
 + orthogonal transparent persistence (no DB!)
 + secure federated transactions (strong consistency)
 + remote method calls
 + mobile code



• Nodes are principals

atomic {		} transaction: atomic, isolated
n.left.value++;		
n.append @ node3 (n')		
}	↑	

check: $p \in \sqsubseteq \text{node3} ?$

- can write complex systems including "mashups" much more concisely.
- lots of challenges left
 - side channels from distributed protocols such as transactions (e.g. timing)
 - scalability
 - verification!

Summary: language-based security

is a great application area for languages & verification.