



BGP Security in Partial Deployment

Is the Juice Worth the Squeeze?

Robert Lychev
Georgia Tech
Atlanta, GA, USA
rlychev@cc.gatech.edu

Sharon Goldberg
Boston University
Boston, MA, USA
goldbe@cs.bu.edu

Michael Schapira
Hebrew University
Jerusalem, Israel
schapiram@huji.ac.il

ABSTRACT

As the rollout of secure route origin authentication with the RPKI slowly gains traction among network operators, there is a push to standardize secure path validation for BGP (*i.e.*, S*BGP: S-BGP, soBGP, BGPSEC, etc.). Origin authentication already does much to improve routing security. Moreover, the transition to S*BGP is expected to be long and slow, with S*BGP coexisting in “partial deployment” alongside BGP for a long time. We therefore use theoretical and experimental approach to study the security benefits provided by partially-deployed S*BGP, vis-a-vis those already provided by origin authentication. Because routing policies have a profound impact on routing security, we use a survey of 100 network operators to find the policies that are likely to be most popular during partial S*BGP deployment. We find that S*BGP provides only meagre benefits over origin authentication when these popular policies are used. We also study the security benefits of other routing policies, provide prescriptive guidelines for partially-deployed S*BGP, and show how interactions between S*BGP and BGP can introduce new vulnerabilities into the routing system.

Categories and Subject Descriptors: C.2.2 [Computer-Communication Networks]: Network Protocols

Keywords: BGP; partial deployment; routing; security

1. INTRODUCTION

Recent high-profile routing failures [9,14,39,40] have highlighted major vulnerabilities in BGP, the Internet’s interdomain routing protocol. To remedy this, secure origin authentication [10,35,37] using the RPKI [32] is gaining traction among network operators, and there is now a push to standardize a path validation protocol (*i.e.*, S*BGP [27,31,46]). Origin authentication is relatively lightweight, requiring neither changes to the BGP message structure nor online cryptographic computations. Meanwhile, path validation with S*BGP could require both [31]. The deployment of origin authentication is already a significant challenge [2]; here we

ask, is the deployment of S*BGP path validation worth the extra effort? (That is, is the juice worth the squeeze?)

To answer this question, we must contend with the fact that any deployment of S*BGP is likely to coexist with legacy insecure BGP for a long time. (IPv6 and DNSSEC, for example, have been in deployment since at least 1999 and 2007 respectively.) In a realistic *partial deployment* scenario, an autonomous system (AS) that has deployed S*BGP will sometimes need to accept insecure routes sent via legacy BGP; otherwise, it would lose connectivity to the parts of the Internet that have not yet deployed S*BGP [31]. Most prior research has ignored this issue, either by assuming that ASes will never accept insecure routes [6,11], by studying only the *full deployment* scenario where every AS has already deployed S*BGP [10,21], or by focusing on creating incentives for ASes to adopt S*BGP in the first place [11,19].

We consider the security benefits provided by partially-deployed S*BGP vis-a-vis those already provided by origin authentication. Fully-deployed origin authentication is lightweight and already does much to improve security, even against attacks it was not designed to prevent (*e.g.*, propagation of bogus AS-level paths) [21]. We find that, given the routing policies that are likely to be most popular during partial deployment, S*BGP can provide only meagre improvements to security over what is already possible with origin authentication; we find that other, less popular policies can sometimes provide tangible security improvements. (“Popular” routing policies were found using a survey of 100 network operators [18].) However, we also show that security improvements can come at a risk; complex interactions between BGP and S*BGP can introduce new instabilities and vulnerabilities into the routing system.

1.1 Security with partially-deployed S*BGP.

With BGP, an AS learns AS-level paths to destination ASes (and their IP prefixes) via routing announcements from neighboring ASes; it then selects one path per destination by applying its local *routing policies*. Origin authentication ensures that the destination AS that announces a given IP prefix is really authorized to do so. S*BGP ensures that the AS-level paths learned actually exist in the network.

In S*BGP partial deployment, security will be profoundly affected by the routing policies used by individual ASes, the AS-level topology, and the set of ASes that are *secure* (*i.e.*, have deployed S*BGP). Suppose a secure AS has a choice between a *secure route* (learned via S*BGP) and an *insecure route* (learned via legacy BGP) to the same destination. While it seems natural that the AS should always prefer the secure route over the insecure route, a network

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

SIGCOMM’13, August 12–16, 2013, Hong Kong, China.

Copyright is held by the owner/author(s). Publication rights licensed to ACM.

Copyright 2013 ACM 978-1-4503-2056-6/13/08 ...\$15.00.

operator must balance security against economic and performance concerns. As such, a *long* secure route through a *costly* provider might be less desirable than a *short* insecure route through a *revenue-generating* customer. Indeed, the BGPSEC standard is careful to provide maximum flexibility, stating the relationship between an AS’s routing policies and the security of a route “is a matter of local policy” [31].

While this flexibility is a prerequisite for assuring operators that S*BGP will not disrupt existing traffic engineering or network management policies, it can have dire consequences on security. Attackers can exploit routing policies that prioritize economic and/or length considerations above security. In a *protocol downgrade attack*, for example, an attacker convinces a secure AS with a secure route to downgrade to a *bogus* route sent via legacy BGP, simply because the bogus route is shorter, or less costly (Section 3.2).

1.2 Methodology & paper roadmap.

This paper summarizes our results. Extended analysis, robustness tests, and proofs are in the full version [34].

Three routing models. In Section 2 we develop models for routing with partially-deployed S*BGP, based on classic models of AS business relationships and BGP [16, 17, 23–25]. Our *security 1st model* supposes that secure ASes *always* prefer secure routes over insecure ones; while this is most natural from a security perspective, a survey of 100 network operators [18] suggests that it is least popular in partial deployment. In our *security 2nd model*, a secure route is preferred only if no *less-costly* insecure route is available. The survey confirms that our *security 3rd model* is most popular in partial deployment [18]; here a secure route is preferred only if there is no *shorter* or *less-costly* insecure route. This paper works within these models; the full version assesses robustness to assumptions made in these models.

Threat model & metric. Sections 3–4.1 introduce our threat model, and a metric to quantify security within this threat model; our metric measures the *average* fraction of ASes using a legitimate route when a destination is attacked.

Deployment invariants. The vast number of choices for the set S of ASes that adopt S*BGP makes evaluating security challenging. Section 4 therefore presents our (arguably) most novel methodological contribution; a framework that bounds the *maximum* improvements in security possible for each routing model, for *any* deployment scenario S .

Deployment scenarios. How close do real S*BGP deployments S come to these bounds? While a natural objective would be to determine the “optimal” deployment S , we prove that this is NP-hard. Instead, Sections 5–6 use simulations on empirical AS-level graphs to quantify security in scenarios suggested in the literature [6, 11, 19, 41], and determine root causes for security improvements (or lack thereof).

Algorithms & experimental robustness. We designed parallel simulation algorithms to deal with the large space of parameters that we explore, *i.e.*, attackers, destinations, deployment scenarios S , and routing policies, (full version). We also controlled for empirical pitfalls, including (a) variations in routing policies (full version) (b) the fact that empirical AS-level graphs tend to miss many peering links at Internet eXchange Points (IXPs) [3, 5, 42], (Section 2.2) (c) a large fraction of the Internet’s traffic originates at a few ASes [30] (Sections 2.2, 4.5, 5.2.2, 5.3.1). While our analysis

cannot predict exactly how individual ASes would react to routing attacks, we do report on strong aggregate trends.

1.3 Results.

Our theoretical and experimental analyses indicate that:

Downgrades are a harsh reality. We find that protocol downgrade attacks (Sections 1.1, 3.2) can be extremely effective; so effective, in fact, that they render deployments of S*BGP at large Tier 1 ISPs almost useless in the face of attacks (Sections 4.6 and 5.3.1).

New vulnerabilities. We find that the interplay between topology and routing policies can cause some ASes to fall victim to attacks they would have avoided if S*BGP had *not* been deployed. Fortunately, these troubling phenomena occur less frequently than phenomena that protect ASes from attacks during partial deployment (Section 6).

New instabilities. We show that undesirable phenomena (BGP Wedgies [22]) can occur if ASes prioritize security inconsistently (Section 2.3).

Prescriptive deployment guidelines. Other than suggesting that (1) ASes should prioritize security in the same way in order to avoid routing instabilities, our results (2) confirm that deploying lightweight *simplex* S*BGP [19, 31] (instead of full-fledged S*BGP) at stub ASes at the edge of the Internet does not harm security (Section 5.3.2). Moreover, while [6, 11, 19] suggest that Tier 1s should be early adopters of S*BGP, our results do not support this; instead, we suggest that (3) Tier 2 ISPs should be among the earliest adopters of S*BGP (Section 4.6, 5.2.3, 5.3.1).

Is the juice worth the squeeze? We use our metric to compare S*BGP in a partial deployment S to the baseline scenario where no AS is secure (*i.e.*, $S = \emptyset$ and only origin authentication is in place). We find that large partial deployments of S*BGP provide excellent protection against attacks when ASes use routing policies that prioritize security 1st (Section 5.2.3); however, [18] suggests that network operators are less likely to use these routing policies. Meanwhile, the policies that operators most favor (*i.e.*, security 3rd) provide only meagre improvements over origin authentication (Section 4.4). This is not very surprising, since S*BGP is designed to prevent path-shortening attacks and when security is 3rd, ASes prefer (possibly-bogus) short insecure routes over longer secure routes.

However, it is less clear what happens in security is 2nd, where route security is prioritized over route length. Unfortunately, even when S*BGP is deployed at 50% of ASes, the benefits obtained in the security 2nd model lag significantly behind those available when security is 1st. While some destinations can obtain tangible benefits when security is 2nd, for others (especially Tier 1s) the security 2nd model behaves much like the security 3rd model (Section 5.2). We could only find clear-cut evidence of large overall improvements in security when ASes prioritize security 1st.

2. SECURITY & ROUTING POLICIES

S*BGP allows an AS to validate the correctness of the AS-level path information it learns from its neighbors [10]. (S-BGP [27] and BGPSEC [31] validate that every AS on a path sent a routing announcement for that path; soBGP [46] validates that all the edges in a path announcement physically exist in the AS-level topology. As we shall see in Section 3, our analysis applies to all these protocols.) However,

Tier 1	13 ASes with high customer degree & no providers
Tier 2	100 top ASes by customer degree & with providers
Tier 3	Next 100 ASes by customer degree & with providers
CPs	17 Content providers: AS 15169, 8075, 20940, 22822, 32934, 15133, 16265, 16509, 2906, 23286, 40428, 714, 10310, 38365, 14907 13414, and 4837.
Small CPs	Top 300 ASes by peering degree (other than Tier 1, 2, 3, and CP)
Stubs-x	ASes with peers but no customers
Stubs	ASes with no customers & no peers
SMDG	Remaining non-stub ASes

Table 1: Tiers.

for S*BGP to prevent routing attacks, validation of paths alone is not sufficient. ASes also need to use information from path validation to make their routing decisions. We consider three alternatives for incorporating path validation into routing decisions, and analyze the security of each.

2.1 Dilemma: Where to place security?

An AS that adopts S*BGP must be able to process and react to insecure routing information, so that it can still route to destination ASes that have not yet adopted S*BGP. The BGPSEC standard is such that a router only learns a path via BGPSEC if every AS on that path has adopted BGPSEC; otherwise, the path is learnt via legacy BGP. (The reasoning for this is in [45] and Appendix A of [19]):

Secure routes. We call an AS that has adopted S*BGP a *secure AS*, and a path learned via S*BGP (*i.e.*, a path where every AS is secure) a *secure path* or *secure route*; all other paths are called *insecure*.

If a secure AS can learn both secure and insecure routes, what role should security play in route selection?

2.2 S*BGP routing models.

While it is well known that BGP routing policies differ between ASes and are often kept private, we need a concrete model of ASes' routing policies so as to analyze and simulate their behaviors during attacks. The following models of routing with S*BGP are variations of the well-studied models from [7, 16, 17, 19, 23–25].

AS-level topology. The AS-level topology is represented by an undirected graph $G = (V, E)$; the set of vertices V represents ASes and the set of links (edges) E represents direct BGP links between neighboring ASes. We will sometimes also refer to the “tiers” of ASes [15] in Table 1; the list of 17 content providers (CPs) in Table 1 was culled from recent empirical work on interdomain traffic volumes [4, 28–30, 44].

ASes' business relationships. Each edge in E is annotated with a business relationship: either (1) *customer-to-provider*, where the customer purchases connectivity from its provider (our figures depict this with an arrow from customer to provider), or (2) *peer-to-peer*, where two ASes transit each other's customer traffic for free (an undirected edge).

Empirical AS topologies. All simulations and examples described in this paper were run on the UCLA AS-level topology from 24 September 2012 [12]. Because empirical AS graphs often miss many of peer-to-peer links in Internet eXchange Points (IXP) [3, 5, 42], we constructed a second graph where we augmented the UCLA graph with over 550K peer-to-peer edges between ASes listed as members of the same IXP on voluntary online sources (IXPs websites, EuroIX, Peering DB, Packet Clearing House, *etc.*). Because

not all ASes at an IXP peer with each other [3], our augmented graph is an upper bound on the number of missing links in the AS graph. When we repeated our simulations on this second graph, we found that all the aggregate trends we discuss in subsequent sections still hold, which suggests they are robust to missing IXP edges. (Results in full version.)

S*BGP routing. ASes running BGP compute routes to each *destination* AS $d \in V$ independently. For every destination AS $d \in V$, each *source* AS $s \in V \setminus \{d\}$ repeatedly uses its local *BGP decision process* to select a single “best” route to d from routes it learns from neighboring ASes. s then announces this route to a subset of its neighbors according to its *local export policy*. An AS s *learns a route* or has an *available route* R if R was announced to s by one of its neighbors; AS s *has* or *uses a route* R if it chooses R from its set of available routes. AS s has *customer* (*resp.*, *peer*, *provider*) route if its neighbor on that route is a customer (*resp.*, *peer*, *provider*); see *e.g.*, AS 29518 in Figure 1 left.

2.2.1 Insecure routing policy model.

When choosing between many routes to a destination d , each *insecure* AS executes the following (in order):

Local pref (LP): Prefer customer routes over peer routes. Prefer peer routes over provider routes.

AS paths (SP): Prefer shorter routes over longer routes.

Tiebreak (TB): Use intradomain criteria (*e.g.*, geographic location, device ID) to break ties among remaining routes.

After selecting a single route as above, an AS announces that route to a subset of its neighbors:

Export policy (Ex): In the event that the route is via a customer, the route is exported to all neighbors. Otherwise, the route is exported to customers only.

The relative ranking of the **LP**, **SP**, and **TB** are standard in most router implementations [13]. The **LP** and **Ex** steps are based on the classical economic model of BGP routing [16, 17, 24, 25]. **LP** captures ASes' incentives to send traffic along revenue-generating customer routes, as opposed to routing through peers (which does not increase revenue), or routing through providers (which comes at a monetary cost). **Ex** captures ASes's willingness to transit traffic only when paid to do so by a customer.

Robustness to LP model. While this paper reports results for the above **LP** model, we also test their robustness to other models for **LP**; results are in the full version.

2.2.2 Secure routing policy models.

Every *secure* AS also adds this step to its routing policy.

Secure paths (SecP): Prefer a secure route over an insecure route.

We consider three models for incorporating the **SecP** step:

Security 1st. The **SecP** is placed before the **LP** step; this model supposes security is an AS's highest priority.

Security 2nd. The **SecP** step comes between the **LP** and **SP** steps; this model supposes that an AS places economic considerations above security concerns.

Security 3rd. The **SecP** step comes between **SP** and **TB** steps; this model, also used in [19], supposes security is prioritized below business considerations and AS-path length.

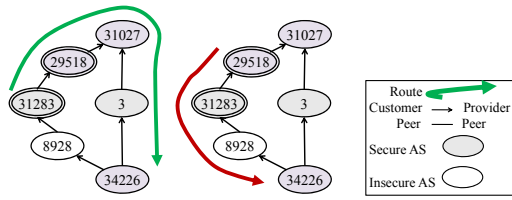


Figure 1: S*BGP Wedgie.

2.2.3 The security 1st model is unpopular.

While the security 1st model is the most “idealistic” from the security perspective, it is likely the least realistic. During incremental deployment, network operators are expected to cautiously incorporate S*BGP into routing policies, placing security 2nd or 3rd, to avoid disruptions due to (1) changes to traffic engineering, and (2) revenue lost when expensive secure routes are chosen instead of revenue-generating customer routes. The security 1st model might be used only once these disruptions are absent (*e.g.*, when most ASes have transitioned to S*BGP), or to protect specific, highly-sensitive IP prefixes. Indeed, a survey of 100 network operators [18] found that 10% would rank security 1st, 20% would rank security 2nd and 41% would rank security 3rd. (The remaining operators opted not to answer this question.)

2.3 Mixing the models?

It is important to note that in each of our S*BGP routing models, the prioritization of the **SecP** step in the route selection process is consistent across ASes. The alternative—lack of consensus amongst network operators as to where to place security in the route selection process—can lead to more than just confusion; it can result in a number of undesirable phenomena that we discuss next.

2.3.1 Disagreements can lead to BGP Wedgies.

Figure 1. Suppose that all ASes in the network, except AS 8928, have deployed S*BGP. The Swedish ISP AS 29518 places security below **LP** in its route selection process, while the Norwegian ISP AS 31283 prioritizes security above all else (including **LP**). Thus, while AS 29518 prefers the customer path through AS 31283, AS 31283 prefers the secure path through its provider AS 29518. The following undesirable scenario, called a “BGP Wedgie” [22] can occur. Initially, the network is in an intended *stable routing state*¹, in which AS 31283 uses the secure path through its provider AS 29518 (left). Now suppose the link between AS 31027 and AS 3 fails. Routing now converges to a *different* stable state, where AS 29518 prefers the customer path through AS 31283 (right). When the link comes back up, BGP does not revert to the original stable state, and the system is stuck in an unintended routing outcome.

“BGP Wedgies” [22] cause unpredictable network behavior that is difficult to debug. (Sami *et al.* [43] also showed that the existence of two stable states, as in Figure 1, implies that persistent routing oscillations are possible.)

2.3.2 Agreements imply convergence.

In the full version we prove that when all ASes prioritize secure routes the same way, convergence to a single stable state is guaranteed, *regardless* of which ASes adopt S*BGP:

¹A routing state, *i.e.*, the route chosen by each AS $s \in V \setminus \{d\}$ to destination d , is *stable* if any AS s that re-runs its route selection algorithm does not change its route [23].

THEOREM 2.1. *S*BGP convergence to a unique stable routing state is guaranteed in all three S*BGP routing models even under partial S*BGP deployment.*

This holds even in the presence of the attack of Section 3.1, *cf.*, [33]. This suggests a prescriptive guideline for S*BGP deployment: ASes should all prioritize security in the same way. (See Section 5.3 for more guidelines.) The reminder of this paper supposes that ASes follow this guideline.

3. THREAT MODEL

To quantify “security” in each of our three models, we first need to discuss what constitutes a routing attack. We focus on a future scenario where RPKI and origin authentication are deployed, and the challenge is engineering global S*BGP adoption. We therefore disregard attacks that are prevented by origin authentication, *e.g.*, prefix- and subprefix-hijacks [7, 9, 10, 14, 36] (when an attacker originates a prefix, or more specific subprefix, when not authorized to do so). Instead, we focus on attacks that are effective even in the presence of origin authentication, as these are precisely the attacks that S*BGP is designed to prevent.

Previous studies on S*BGP security [6, 11, 21] focused on the endgame scenario, where S*BGP is fully deployed, making the crucial assumption that *any secure AS that learns an insecure route from one of its neighbors can safely ignore that route*. This assumption is invalid in the context of a partial deployment of S*BGP, where S*BGP coexists alongside BGP. In this setting, some destinations may only be reachable via insecure routes. Moreover, even a secure AS may prefer to use an insecure route for economic or performance reasons (as in our security 2nd or 3rd models). Therefore, propagating a bogus AS path using legacy insecure BGP [21, 40] (an attack that is effective against fully-deployed origin authentication) can *also* work against some *secure* ASes when S*BGP is partially deployed.

3.1 The attack.

We focus on the scenario where a single attacker AS m attacks a single destination AS d ; all ASes except m use the policies in Section 2.2. The attacker m ’s objective is to maximize the number of source ASes that send traffic to m , rather than d . This commonly-used objective function [7, 20, 21] reflects m ’s incentive to attract (and therefore tamper / eavesdrop / drop) traffic from as many source ASes as possible. (We deal with the fact that ASes can source different amounts of traffic [30] in Sections 4.5, 5.2.2, 5.3.1.)

Attacker’s strategy. The attacker m wants to convince ASes to route to m , instead of the legitimate destination AS d that is authorized to originate the prefix under attack. It will do this by sending bogus AS-path information using legacy BGP. What AS path information should m propagate? A straightforward extension of the results in [21] to our models shows it is NP-hard for m to determine a bogus route to export to each neighbor that maximizes the number of source ASes it attracts. As such, we consider the arguably simplest, yet very disruptive [7, 21], attack: the attacker, which is not actually a neighbor of the destination d , pretends to be directly connected to d . Since there is no need to explicitly include IP prefixes in our models, this translates to a single attacker AS m announcing the bogus AS-level path “ m, d ” using legacy BGP to *all* its neighbor ASes. Since the path is announced via legacy BGP, recipient ASes will

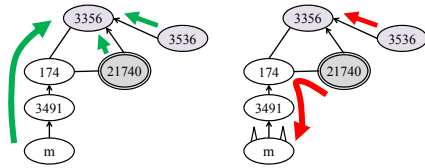


Figure 2: Protocol downgrade attack; Sec 2nd.

not validate it with S*BGP, and thus will not learn that it is bogus. (This attack is equally effective against partially-deployed soBGP, S-BGP and BGPSEC. With soBGP, the attacker claims to have an edge to d that does not exist in the graph. With S-BGP or BGPSEC the attacker claims to have learned a path “ m, d ” that d never announced.)

3.2 Are secure ASes subject to attacks?

Ideally, we would like a secure AS with a secure route to be protected from a routing attack. Unfortunately, however, this is not always the case. We now discuss a troubling aspect of S*BGP in partial deployment [26]:

Protocol downgrade attack. In a protocol downgrade attack, a source AS that uses a secure route to the legitimate destination under normal conditions, downgrades to an insecure bogus route *during* an attack.

The best way to explain this is via an example:

Figure 2. We show how AS 21740, a webhosting company, suffers a protocol downgrade attack, in the security 2nd (or 3rd) model. Under normal conditions (left), AS 21740 has a secure provider route directly to the destination Level 3 AS 3356, a Tier 1 ISP. (AS 21740 does *not* have a peer route via AS 174 due to **Ex.**) During the attack (right), m announces that it is directly connected to Level3, and so AS 21740 sees a bogus, insecure 4-hop peer route, via his peer AS 174. Importantly, AS 21740 has no idea that this route is bogus; it looks just like any other route that might be announced with legacy BGP. In the security 2nd (and 3rd) model, AS 21740 prefers an insecure *peer* route over a secure *provider* route, and will therefore downgrade to the bogus route.

Downgrades are avoided in the security 1st model. Protocol downgrade attacks can happen in the security 2nd and 3rd models, but not when security is 1st:

THEOREM 3.1. *In the security 1st model, for every attacker AS m , destination AS d , and AS s that, in normal conditions, has a secure route to d that does not go through m , s will use a secure route to d even during m ’s attack.*

While the theorem holds only if the attacker m is not on AS s ’s route, this is not a severe restriction because, otherwise, m would attract traffic from s to d even without attacking.

4. INVARIANTS TO DEPLOYMENT

Given the vast number of possible configurations for a partial deployment of S*BGP, we present a framework for exploring the security benefits of S*BGP vis-a-vis origin authentication, *without making any assumptions about which ASes are secure*. To do this, we show how to quantify security (Section 4.1), discuss how to determine an *upper bound* on security available with *any* S*BGP deployment for any routing model (Section 4.3.1), finally compare it to the security available with origin authentication (Section 4.2, 4.4).

happy	Chooses a legitimate secure/insecure route to d .
unhappy	Chooses a bogus insecure route to m .
immune	Happy <i>regardless of which ASes are secure</i> .
doomed	Unhappy <i>regardless of which ASes are secure</i> .
protectable	Neither immune nor doomed.

Table 2: Status of source s when m attacks d .

4.1 Quantifying security: A metric.

We quantify improvements in “security” by determining the fraction of ASes that avoid attacks (per Section 3.1). The attacker’s goal is to attract traffic from as many ASes as possible; our metric therefore measures the average fraction of ASes that do *not* choose a route to the attacker.

Metric. Suppose the ASes in set S are secure and consider an attacker m that attacks a destination d . Let $H(m, d, S)$ be the number of “happy” source ASes that choose a legitimate route to d instead of a bogus route to m . (See Table 2). Our metric is:

$$H_{M,D}(S) = \frac{1}{|D|(|M|-1)(|V|-2)} \sum_{m \in M} \sum_{d \in D \setminus \{m\}} H(m, d, S)$$

Since we cannot predict where an attack will come from, or which ASes it will target, the metric averages over all attackers in a set M and destinations in a set D ; we can choose M and D to be any subset of the ASes in the graph, depending on (i) where we expect attacks to come from, and (ii) which destinations we are particularly interested in protecting. When we want to capture the idea that all destinations are of equal importance, we average over all destinations; note that “China’s 18 minute mystery” of 2010 [14] fits into this framework well, since the hijacker targeted prefixes originated by a large number of (seemingly random) destination ASes. However, we can also zoom in on important destinations D (e.g., content providers [9, 30, 39]) by averaging over those destinations only. We can, analogously, zoom in on certain types of attackers M by averaging over them only.

Tiebreaking & bounds on the metric. Recall from Section 2.2 that our model fully determines an AS’s routing decision up to the tiebreak step **TB** of its routing policy. Since computing $H_{M,D}(S)$ only requires us to distinguish between “happy” and “unhappy” ASes, the tiebreak step matters only when a source AS s has to choose between (1) an *insecure* route(s) to the legitimate destination d (that makes it happy), and (2) an *insecure* bogus route(s) to m (that makes it unhappy). Importantly, s has no idea which route is bogus and which is legitimate, as both of them are insecure. Therefore, to avoid making uninformed guesses about how ASes choose between equally-good *insecure* routes, we will compute upper and lower bounds on our metric; to get a lower bound, we assume that every AS s in the aforementioned situation will always choose to be unhappy (i.e., option (2)); the upper bound is obtained by assuming s always chooses to be happy (i.e., (1)).

4.2 Origin authentication gives good security.

At this point, we could compute the metric for various S*BGP deployment scenarios, show that most source ASes are “happy”, argue that S*BGP has improved security, and conclude our analysis. This, however, would not give us the full picture, because it is possible that most of the happy ASes would have been happy *even if S*BGP had not been deployed*. Thus, to understand if the juice is worth the squeeze,

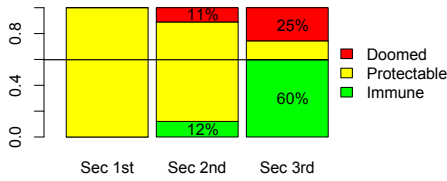


Figure 3: Partitions

we need to ask how many more attacks are prevented by a particular S*BGP deployment scenario, relative to those already prevented by RPKI with origin authentication. More concretely, we need to compare the fraction of happy ASes *before and after the ASes in S deploy S*BGP*. To do this, we compare the metric for a deployment scenario S against the “baseline scenario”, where RPKI and origin authentication are in place, but no AS has adopted S*BGP, so that the set of secure ASes is $S = \emptyset$.

In [21], the authors evaluated the efficacy of origin authentication against attacks that it was not designed to prevent — namely, the “ m, d ” attack of Section 3.1. They randomly sampled pairs of attackers and destinations and plotted the distribution of the fraction of “unhappy” source ASes (ASes that route through the attacker, see Table 2). Figure 3 of [21] shows that attacker is able to attract traffic from less than half of the source ASes in the AS graph, on average. We now perform a computation and obtain a result that is similar in spirit; rather than randomly sampling pairs of attackers and destinations as in [21], we instead compute a *lower bound* on our metric over all possible attackers and destinations. We find that $H_{V,V}(\emptyset) \geq 60\%$ on the basic UCLA graph, and $H_{V,V}(\emptyset) \geq 62\%$ on our IXP-augmented graph.

It is striking that both our and [21]’s result indicate more than half of the AS graph is *already* happy even *before* S*BGP is deployed. To understand why this is the case, recall that with origin authentication, an attacking AS m must announce a bogus path “ m, d ” that is one hop longer than the path “ d ” announced by the legitimate destination AS d . When we average over all (m, d) pairs and all the source ASes, bogus paths through m will appear longer, on average, than legitimate paths through d . Since path length plays an important role in route selection, on average, more source ASes choose the legitimate route.

4.3 Does S*BGP give better security?

How much further can we get with a partial deployment of S*BGP? We now obtain bounds on the improvements in security that are possible for a given routing policy model, but *for any* set S of secure ASes.

We can obtain these bounds thanks to the following crucial observation: ASes can be partitioned into three distinct categories with respect to each attacker-destination pair (m, d) . Some ASes are *doomed* to route through the attacker regardless of which ASes are secure. Others are *immune* to the attack regardless of which ASes are secure. Only the remaining ASes are *protectable*, in the sense that whether or not they route through the attacker depends on which ASes are secure (see Table 2).

To bound our metric $H_{M,D}(S)$ for a given routing policy model (i.e., security 1^{st} , 2^{nd} , or 3^{rd}) and *across all partial-deployment scenarios S* , we first partition source ASes into categories — doomed, immune, and protectable — for each (m, d) pair and each routing policy model. By computing the average fraction of immune ASes across all $(m, d) \in$

$M \times D$ for a given routing model, we get a lower bound on $H_{M,D}(S) \forall S$ and that routing model. We similarly get an upper bound on $H_{M,D}(S)$ by computing the average fraction of ASes that are *not* doomed.

4.3.1 Partitions: Doomed, protectable & immune.

We return to Figure 2 to explain our partitioning:

Doomed. A source AS s is *doomed* with respect to pair (m, d) if s routes through m no matter which set S of ASes is secure. AS 174 in Figure 2 is doomed when security is 2^{nd} (or 3^{rd}). If security is 2^{nd} (or 3^{rd}), AS 174 *always* prefers the bogus customer route to the attacker over a (possibly secure) peer path to the destination AS 3356, for every S .

Immune. A source AS s is *immune* with respect to pair (m, d) if s will route through d no matter which set S of ASes is secure. AS 3536 in Figure 2 is one example; this single-homed stub customer of the destination AS 3356 can *never* learn a bogus route in any of our security models. When security is 2^{nd} or 3^{rd} , another example of an immune AS is AS 10310 in Figure 7; its customer route to the legitimate destination AS 40426 is always more attractive than its provider route to the attacker in these models.

Protectable. AS s is protectable with respect to pair (m, d) if it can either choose the legitimate route to d , or the bogus one to m , depending on S . With security 1^{st} , AS 174 in Figure 2 becomes protectable. If it has a secure route to the destination AS 3356, AS 174 will choose it and be happy; if not, it will choose the bogus route to m .

4.3.2 Which ASes are protectable?

The intuition behind the following partitioning of ASes is straightforward. The subtleties involved in proving that an AS is doomed/immune are discussed in the full version.

Security 1^{st} . Here, we suppose that all ASes are protectable; the few exceptions (e.g., the single-homed stub of Figure 2) have little impact on the count of protectable ASes.

Security 2^{nd} . Here, an AS is doomed if it has a route to the attacker with better local preference **LP** than every available route to the legitimate destination; (e.g., the bogus *customer* route offered to AS 174 in Figure 2 has higher **LP** than the legitimate *peer* route). An immune AS has a route to the destination that has higher **LP** than every route to the attacker. For protectable AS, its best available routes to the attacker and destination have *exactly the same LP*.

Security 3^{rd} . Here, a doomed AS has a path to m with (1) better **LP** OR (2) equal **LP** and shorter length **SP**, than every available path to d . The opposite holds for an immune AS. A protectable AS has best available routes to m and d with equal **LP** and path length **SP**.

4.4 Bounding security for all deployments.

For each routing model, we found the fraction of doomed/protectable / immune source ASes for each attacker destination pair (m, d) , and took the average over all $(m, d) \in V \times V$. We used these values to get upper- and lower bounds on $H_{V,V}(S)$ for *all deployments S* , for each routing model.

Figure 3: The colored parts of each bar represent the average fraction of immune, protectable, and doomed source ASes, averaged over all $O(|V|^2)$ possible pairs of attackers and destinations. Since $H_{V,V}(S)$ is an average of the fraction of happy source ASes over all pairs of attackers and

destinations, the upper bound on the metric $H_{V,V}(S) \forall S$ is the average fraction of source ASes that are *not* doomed. The upper bound on the metric $H_{V,V}(S) \forall S$ is therefore: $\approx 100\%$ with security 1^{st} , 89% with security 2^{nd} , and 75% with security 3^{rd} . (The same figure computed on our IXP-edge-augmented graph looks almost exactly the same, with the proportions being $\approx 100\%$, 90% and 77% .) Meanwhile, the heavy solid line is the lower bound on the metric $H_{V,V}(\emptyset)$ in the baseline setting where $S = \emptyset$ and there is only origin authentication; in Section 4.2 we found that $H_{V,V}(\emptyset) = 60\%$ (and 62% for the IXP-edge-augmented graph). Therefore, we can bound the maximum change in our security metric $H_{V,V}(S) \forall S$ for each routing policy model by computing the distance between the solid line and the boundary between the fraction of doomed and protectable ASes. We find:

Security 3^{rd} : Little improvement. Figure 3 shows that the maximum gains over origin authentication that are provided by the security 3^{rd} model are quite slim — at most 15% — *regardless* of which ASes are secure. (This follows because the upper bound on the metric $H_{V,V}(S) \leq 75\%$ for any S while the lower bound on the baseline setting is $H_{V,V}(\emptyset) \geq 60\%$.) Moreover, these are the *maximum* gains $\forall S$; in a realistic S*BGP deployment, the gains are likely to be much smaller. This result is disappointing, since the security 3^{rd} model is likely to be the most preferred by network operators (Section 2.2.3), but it is not especially surprising. S*BGP is designed to prevent path shortening attacks; however, in the security 3^{rd} model ASes prefer short (possibly bogus) insecure routes over a long secure routes, so it is natural that this model realizes only minimal security benefits.

Security 2^{nd} : More improvement. Meanwhile, route security is prioritized above route length with the security 2^{nd} model, so we could hope for better security benefits. Indeed, Figure 3 confirms that the *maximum gains* over origin authentication are better: $89 - 60 = 29\%$. But can these gains be realized in realistic partial-deployment scenarios?

Decreasing numbers of immune ASes? The fraction of immune ASes in the security 2^{nd} (12%) and 1^{st} ($\approx 0\%$) models is (strangely) lower than the fraction of happy ASes in the baseline scenario (60%). How is this possible? In Section 6.1.1 we explain this counterintuitive observation by showing that *more* secure ASes can sometimes result in *less* happy ASes; these “collateral damages”, that occur only in the security 1^{st} and 2^{nd} models, account for the decrease in the number of immune ASes.

4.5 Robustness to destination tier.

Thus far, we have been averaging our results over all possible attacker-destination pairs in the graph. However, some destination ASes might be particularly important to secure, perhaps because they source important content (*e.g.*, the content provider ASes (CPs)) or transit large volumes of traffic (the Tier 1 ASes). As such, we broke down the metric over destinations in each *tier* in Table 1.

Figure 4. We show the partitioning into immune / protectable / doomed ASes in the security 3^{rd} model, but this time averaged individually over all destinations in each tier, and all possible attackers V . The thick horizontal line over each vertical bar again shows the corresponding lower bound on our metric $H_{V,Tier}(\emptyset)$ when no AS is secure. Apart from the Tier 1s (discussed next), we observe similar trends as in Section 4.4, with the improvement in security ranging from

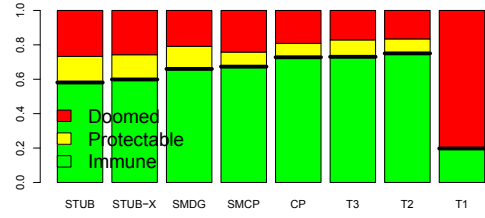


Figure 4: Partitions by destination tier. Sec 3^{rd} .

$8 - 15\%$ for all tiers; the same holds for the security 2^{nd} model. (Figure in full version).

4.6 It’s difficult to protect Tier 1 destinations.

Strangely enough, Figure 4 shows that when Tier 1 destinations are attacked in the security 3^{rd} model, the vast majority ($\approx 80\%$) of ASes are doomed, and only a tiny fraction are protectable; the same holds when security is 2^{nd} (not shown). Therefore, in these models, S*BGP can do little to blunt attacks on Tier 1 destinations.

How can it be that Tier 1s, the largest and best connected (at least in terms of customer-provider edges) ASes in our AS graph, are the most vulnerable to attacks? Ironically, it is the Tier 1s’ very connectivity that harms their security. Because the Tier 1s are so well-connected, they can charge most of their neighbors for Internet service. As a result, most ASes reach the Tier 1s via costly provider paths that are the least preferred type of path according to the **LP** step in our routing policy models. Meanwhile, it turns out that when a Tier 1 destination is attacked, most source ASes will learn a bogus path to the attacker that is *not* through a provider, and is therefore preferred over the (possibly secure) provider route to the T1 destination in the security 2^{nd} or 3^{rd} models. In fact, this is exactly what lead to the protocol downgrade attack on the Tier 1 destination AS 3356 in Figure 2. We will later (Section 5.3.1) find that this is a serious hurdle to protecting Tier 1 destinations.

5. DEPLOYMENT SCENARIOS

In Section 4.4 we presented upper bounds on the improvements in security from S*BGP deployment for choice of secure ASes S . We found that while only meagre improvements over origin authentication are possible in the security 3^{rd} model, better results are possible in the security 2^{nd} and 1^{st} models. However, achieving the bounds in Section 4.4 could require full S*BGP deployment at every AS. What happens in more realistic deployment scenarios? First, we find that the security 2^{nd} model often behaves disappointingly like the security 3^{rd} model. We also find that Tier 1 destinations remain most vulnerable to attacks when security is 2^{nd} or 3^{rd} . We conclude the section by presenting prescriptive guidelines for partial S*BGP deployment.

Robustness to missing IXP edges. In the full version, we repeat the analysis in Section 5.2-5.3 on IXP-edge-augmented AS graph, and see almost identical trends.

5.1 It’s hard to decide whom to secure.

We first need to decide which ASes to secure. Ideally, we could choose the smallest set of ASes that maximizes the value of the metric. To formalize this, consider the following computational problem, that we call “Max-k-Security”: Given an AS graph, a specific attacker-destination pair (m, d) ,

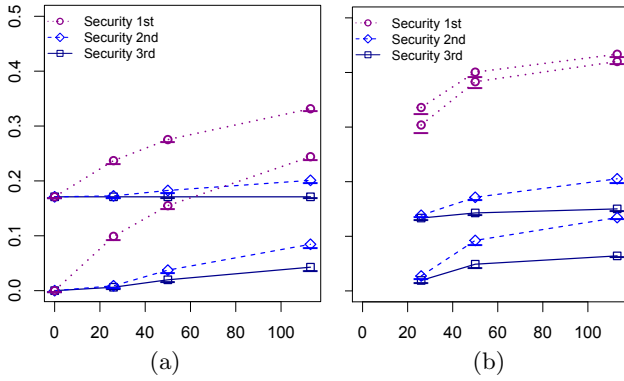


Figure 5: Tier 1+2 rollout: For each step S in rollout, upper and lower bounds on (a) $H_{M',V}(S) - H_{M',V}(\emptyset)$ and (b) $H_{M',V}(S) - H_{M',d}(\emptyset)$ averaged over all $d \in S$. The x -axis is the number of non-stub ASes in S . The “error bars” are explained in Section 5.3.2.

and a parameter $k > 0$, find a set S of secure ASes of size k that maximizes the total number of happy ASes. Then:

THEOREM 5.1. *Max- k -Security is NP-hard in all three routing policy models.*

The proof is in the full version. This result can be extended to the problem of choosing the set of secure ASes that maximize the number of happy ASes over *multiple* attacker-destination pairs (which is what our metric computes).

5.2 Large partial deployments.

Instead of focusing on choosing the optimum set S of ASes to secure (an intractable feat), we will instead consider a few partial deployment scenarios among high-degree ASes S , as suggested in practice [41] and in the literature [6, 11, 19].

Non-stub attackers. We now suppose that the set of attackers is the set of non-stub ASes in our graph M' (i.e., not “Stubs” or “Stubs- x ” per Table 1). Ruling out stub ASes is consistent with the idea that stubs cannot launch attacks if their providers perform prefix filtering [10, 21], a functionality that can be achieved via IRRs [1] or even the RPKI [38], and does not require S*BGP.

5.2.1 Security across all destinations.

Gill *et al.* [19] suggest bootstrapping S*BGP deployment by having secure ISPs deploy S*BGP in their customers that are stub ASes. We therefore consider this “rollout”:

Tier 1 & Tier 2 rollout. Other than the empty set, we consider three different secure sets. We secure X Tier 1’s and Y Tier 2’s and all of their stubs, where $(X, Y) \in \{(13, 13), (13, 37), (13, 100)\}$; this corresponds to securing about 33%, 40%, and 50% of the AS graph.

The results are shown in **Figure 5(a)**, which plots, for each routing policy model, the increase in the upper- and lower bound on $H_{M',V}(S)$ (Section 4.1) for each set S of secure ASes in the rollout (y -axis), versus the number of non-stub ASes in S (x -axis). We make a few important observations:

Tiebreaking can seal an AS’s fate. Even with a large deployment of S*BGP, the improvement in security is highly dependent on the vagaries of the intradomain tiebreaking criteria used to decide between *insecure* routes. (See also

Section 4.1’s discussion on tiebreaking.) Even when we secure 50% of ASes in the security 1st model (the last step of our rollout), there is still a gap of more than 10% between the lower and upper bounds of our metric. Thus, in a partial S*BGP deployment, there is a large fraction of ASes that are balanced on a knife’s edge between an insecure legitimate route and an insecure bogus route; only the (unknown-to-us) intradomain routing policies of these ASes can save them from attack. This is inherent to any partial deployment of S*BGP, even in the security 1st model.

Meagre improvements even when security is 2nd. As expected, the biggest improvements come in the security 1st model, where ASes make security their highest priority and deprecate all economic and operational considerations. When security is 1st and 50% of the AS graph is secure (at the last step in the rollout), the improvement over the baseline scenario is significant; about 24%. While we might hope that the security 2nd model would present improvements that are similar to those achieved when security is 1st, this is unfortunately not the case. In both the security 2nd and 3rd models we see similarly disappointing increases in our metric. We explain this observation in Section 6.2.

5.2.2 Focus on the content providers?

Since much of the Internet’s traffic originates at the content providers (CPs), we might consider the impact of S*BGP deployment on CPs only. We considered the same rollout as above, but with all 17 CPs secure, and computed the metric *over CP destinations only*, i.e., $H_{M',CP}(S)$. Results, shown in the full version, were similar to Figure 5(a).

5.2.3 Different destinations see different benefits.

Thus far, we have looked at the impact of S*BGP in aggregate across all destinations $d \in V$ (or $d \in CP$). Because secure routes can only exist to secure destinations, we now look at the impact of S*BGP on *individual secure destinations* $d \in S$, by considering $H_{M',d}(S)$.

Figure 5(b). We plot the upper and lower bounds on the *change* in the metric, i.e., $H_{M',d}(S) - H_{M',d}(\emptyset)$, averaged across *secure destinations only*, i.e., $d \in S$. As expected, we find large improvements when security is 1st, and small improvements when security is 3rd. Interestingly, however, when security is 2nd the metric does increase by 13 – 20% by the last step in the rollout; while this is still significantly smaller than what is possible when security is 1st, it does suggest that at least some secure destinations benefit more when security is 2nd, rather than 3rd.

For more insight, we zoom in on this last step in our rollout:

Figure 6. For the last step in our rollout, we plot upper and lower bounds on the *change* in the metric, i.e., $H_{M',d}(S) - H_{M',d}(\emptyset)$, for each *individual* secure destination $d \in S$. For each of our three models, the lower bound for each $d \in S$ is plotted as a non-decreasing sequence; these are the three “smooth” lines. The corresponding upper bound for each $d \in S$ was plotted as well. For security 1st, the upper and lower bounds are almost identical, and for security 2nd and 3rd, the upper bounds are the “clouds” that hover over the lower bounds. A few observations:

Security 1st provides excellent protection. We find that when security is 1st, a *secure* destination can reap the full benefits of S*BGP even in (a large) partial deployment. To see this, we computed the *true value* of $H_{M',d}(S)$ for all

secure destinations $d \in S$, and found that it was between 96.8 – 97.9% on average (across all $d \in S$).

Security 2^{nd} and 3^{rd} are similar for many destinations. Figure 6 also reveals that many destinations obtain roughly the same benefits from S*BGP when security is 2^{nd} as when security is 3^{rd} . Indeed, 93% of 7500 secure destinations that see $< 4\%$ (lower-bound) improvement in Figure 6 when security is 3^{rd} , do the same when security is 2^{nd} as well. What is the reason for this? There are certain types of protocol downgrade attacks that succeed *both* when security is 2^{nd} and when security is 3^{rd} (*i.e.*, when the bogus path has better **LP** than the legitimate path, see *e.g.*, Figure 2). In Section 6.2 we shall show that protocol downgrade attacks are the most significant reason for the metric to degrade; therefore, for destinations where these “LP-based” protocol downgrade attacks are most common, the security 2^{nd} model looks much like the security 3^{rd} model.

Tier 1s do best when security is 1^{st} , and worst when it is 2^{nd} or 3^{rd} . When security is 1^{st} , our data also shows that the secure destinations that obtain the largest ($> 40\%$) increases in their security metric $H_{M',d}(S)$ (relative to the baseline setting $H_{M',d}(\emptyset)$) include: (a) all 13 Tier 1s, and (b) $\geq 99\%$ of “Tier 1 stub” destinations (*i.e.*, stub ASes such that all their providers are Tier 1 ASes). On the other hand, these same destinations experience the *worst* improvements when security is 2^{nd} or 3^{rd} (*i.e.*, a lower bound of $< 3\%$).

To explain this, recall from Section 4.6 that when security is 2^{nd} or 3^{rd} , most source ASes that want to reach a Tier 1 destination are *doomed*, because of protocol downgrade attacks like the one shown in Figure 2. This explains the meagre benefits these destinations obtain when security is 2^{nd} or 3^{rd} . On the other hand, protocol downgrade attacks fail when security is 1^{st} . Therefore, in the security 1^{st} model, the Tier 1 destinations (and by extension, Tier 1 stub destinations) obtain excellent security when S*BGP is partially deployed; moreover, they see most significant gains simply because they were so highly vulnerable to attacks in the absence of S*BGP (Figure 4, Section 4.6).

Security 2^{nd} helps some secure destinations. Finally, when security is 2^{nd} , about half of the secure destinations $d \in S$ see benefits that are discernibly better than what is possible when security is 3^{rd} , though not quite as impressive as those when security is 1^{st} . These destinations include some Tier 2s and their stubs, but never any Tier 1s. Similar observations hold for earlier steps in our “rollout”.

5.3 Prescriptive deployment guidelines.

Section 2.3 suggested that ASes use consistent routing policies. We now suggest a few more deployment guidelines.

5.3.1 On the choice of early adopters.

Previous work [6, 11, 19] suggests that Tier 1s should be the earliest adopters of S*BGP. However, the discussion in Sections 4.6 and 5.2.3 suggests that securing Tier 1s might not lead to good security benefits at the early adoption stage, when ASes are most likely to rank security 2^{nd} or 3^{rd} . We now confirm this.

All Tier 1s and their stubs. Even in a deployment that includes *all* 13 Tier 1 ASes and their stubs (*i.e.*, 7872 ASes or $\approx 20\%$ of the AS graph), improvements in security were almost imperceptible. With security 2^{nd} or 3^{rd} , the average

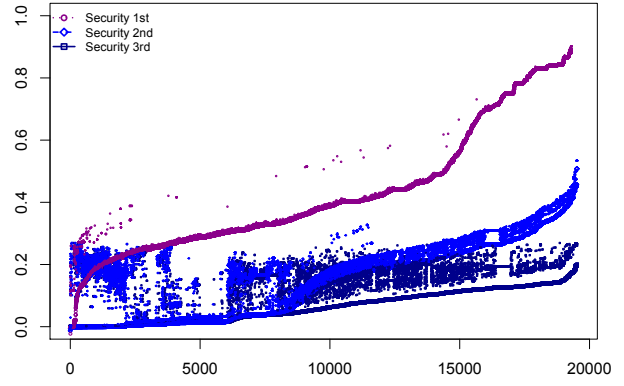


Figure 6: Non-decreasing sequence of $H_{M',d}(S) - H_{M',d}(\emptyset) \forall d \in S$. S is all T1s, T2s, and their stubs.

change in $H_{M',d}(S) - H_{M',d}(\emptyset)$ over secure destinations $d \in S$ causes the metric to increase by $< 0.2\%$.

Tier 1s, their stubs, and content providers. Following [19, 41], we consider securing the CPs, the Tier 1s and all of their stubs, and obtained similar results.

Choose Tier 2s as early adopters. We found that early deployments at the Tier 2 ISPs actually fare better than those at the larger, and better connected Tier 1s. For example, securing the 13 largest Tier 2s (in terms of customer degree) and all their stubs (a total of 6918 ASes), the average change in $H_{M',d}(S) - H_{M',d}(\emptyset)$ over secure destinations $d \in S$ is $\approx 1\%$ when security is 2^{nd} or 3^{rd} .

5.3.2 Use simplex S*BGP at stubs.

Next, we consider [19, 31]’s suggestion for reducing complexity by securing stubs with *simplex S*BGP*.

Simplex S*BGP. Stub ASes have no customers of their own, and therefore (by **Ex**) they will never send S*BGP announcements for routes through other ASes. They will, however, announce routes to their own IP prefixes. For this reason [19, 31] suggests either (1) allowing ISPs to send S*BGP messages on behalf of their stub customers or (2) allowing stubs to deploy S*BGP in a unidirectional manner, sending outgoing S*BGP messages but receiving legacy BGP messages. Since a stub propagates only outgoing BGP announcements for a very small number of IP prefixes (namely, the prefixes owned by that stub), simplex mode can decrease computational load, and make S*BGP adoption less costly. Given that 85% of ASes are stubs, does this harm security?

Figure 5(a)-5(b). The “error bars” in Figure 5(a)-5(b) show what happens when we suppose that all stubs run simplex S*BGP. There is little change in the metric. To explain this, we note that (1) a stub’s routing decision does not affect any other AS’s routing decision, since by **Ex** stubs do not propagate BGP routes from one neighbor to another, and (2) a stub’s routing decisions are limited by the decisions made by its providers, so if its providers avoid attacks, so will the stub, but (3) the stub acts like a secure destination, and therefore (nonstub) ASes establishing routes to the stub still benefit from S*BGP. These results indicate that simplex S*BGP at stubs can lower the complexity of S*BGP deployment without impacting overall security.

Security model	1 st	2 nd	3 rd
Protocol downgrade attacks	X	✓	✓
Collateral benefits	✓	✓	✓
Collateral damages	✓	✓	X

Table 3: Phenomena in different security models

6. ROOT CAUSES & NON-MONOTONICITY

We now examine the reasons for the changes in our security metric as S*BGP is deployed. We start by discussing two subtle phenomena: the collateral damages and collateral benefits incurred by insecure ASes from the deployment of S*BGP at *other* ASes. We then use these phenomena in a root-cause analysis of the results of Section 5.

6.1 Security is not monotonic!

The most obvious desiderata from S*BGP deployment is that the Internet should become only more secure as more ASes adopt S*BGP. Unfortunately, however, this is *not* always the case. Security is not *monotonic*, in the sense that securing more ASes can actually make other ASes unhappy.

To explain this, we use a running example taken from the UCLA AS graph, where the destination (victim) AS d is Pandora’s AS40426 (a content provider) and the attacker m is an anonymized Tier 2 network. We consider the network *before* and *after* a partial deployment of S*BGP S and see how the set of happy ASes changes; S consists of all 100 Tier 2s, all 17 content providers, and all of their stubs.

6.1.1 Collateral Damages

Figure 7. We show how AS 52142, a Polish ISP, suffers from collateral damage when security is 2nd. On the left, we show the network prior to S*BGP deployment. AS 52142 is offered two paths, both insecure: a 3-hop path through his provider AS 5617 to the legitimate destination AS 40426, and a 5-hop bogus route to the attacker. (The route to m is really 4 hops long, but m (falsely) claims a link to AS 40426 so AS 52142 thinks it is 5 hops long.) AS 52142 will choose the legitimate route because it is shorter. On the right, we show the network after S*BGP deployment. AS 5617 has become secure and now prefers the secure route through its neighbor Cogent AS 174. However, AS 5617’s secure route is 5 hops long (right), significantly longer than the 2 hop route AS 5617 used prior to S*BGP deployment (left). Thus, after S*BGP deployment AS 52142 learns a 6-hop legitimate route through AS 5617, and a 5-hop bogus route. Since AS 52142 is insecure, it chooses the shorter route, and becomes unhappy as collateral damage.

No collateral damages in the security 3rd model: The collateral damage above occurs because AS 5617 prefers a *longer* secure route over a shorter insecure route. This can also happen in the security 1st model (but see also Appendix A), but not when security is 3rd. See Table 3.

THEOREM 6.1. *In the security 3rd model, if an AS s has a route to a destination d that avoids an attacker m when the set of secure ASes is S , then s has a route to a destination d that avoids attacker m for every set of secure ASes in $T \supset S$.*

6.1.2 Collateral Benefits

Insecure ASes can also become happy as a *collateral benefit*, because *other* ASes obtained secure routes:

Figure 7. We show how AS 5166, with the Department of Defense Network Information Center, obtains collateral

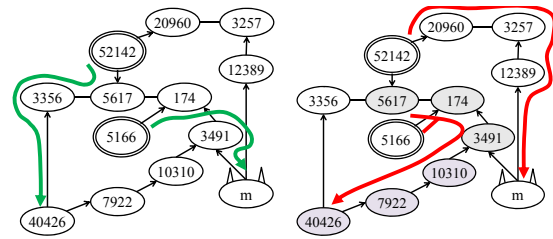


Figure 7: Collateral benefits & damages; sec 2nd.

benefits when its provider AS 174, Cogent, deploys S*BGP. On the left, we show the network prior to the deployment of S*BGP; focusing on Cogent AS 174, we see that it falls victim to the attack, choosing a bogus route through its customer AS 3491. As a result, AS 5166 routes to the attacker as well. On the right, we show the network after S*BGP deployment. Now, both AS 174 and AS 3491 are secure, and choose a longer secure customer route to the legitimate destination. As a result, AS 5166, which remains insecure, becomes happy as a collateral benefit.

Collateral benefits are possible in all three routing policy models (Table 3). Examples are in the full version.

6.2 Root-cause analysis.

Which of the phenomena in Table 3 have the biggest impact on security? We now check how these phenomena play out in the last step of the Tier 1 & Tier 2 rollout of Section 5.2.1. Recall that S is all 13 Tier 1s, all 100 Tier 2s and all of their stubs, *i.e.*, roughly 50% of the AS graph.

Figure 8 (left). We start with a root cause analysis for the security 3rd model. Recall that Theorem 6.1 showed that collateral damages do not occur in the security 3rd model, and so we do not consider them here.

Changes in secure routes. The bottom three parts of the bar show the fraction of secure routes available in normal conditions, prior to any routing attacks. (Averaging is across all V^2 sources and destinations.) During routing attacks, these routes can be broken down into three types: (1) secure routes lost to protocol downgrade attacks (lowest part of the bar), (2) secure routes that are “wasted” on ASes that would have been happy *even in the absence of S*BGP* (second lowest part), and (3) secure routes that protected ASes that were unhappy in the absence of S*BGP (third lowest part). (Averaging is, as usual, over M' and $D = V$ and all V source ASes.) Importantly, improvements in our security metric can only result from the small fraction of secure routes in class (3); the remaining secure routes either (1) disappear due to protocol downgrades, or (2) are “wasted” on ASes that would have avoided the attack even without S*BGP.

Changes in the metric. The top two parts of the bar show how (the lower bound on) the metric $H_{M',V}(S)$ grows relative to the baseline scenario $S = \emptyset$ due to: (a) secure routes in class (3), and (b) (the lower bound on) the fraction of insecure ASes that obtained collateral benefits. Figure 8(left) thus illustrates the importance of collateral benefits.

Figure 8 (right). We perform the same analysis for the security 1st model. By Theorem 3.1, protocol downgrade attacks occur only rarely in this model, so these are not visible in the figure. However, we now have to account for collateral damages (Section 6.1.1), which we depict with the smaller sliver on right of the figure. We obtain the change in the

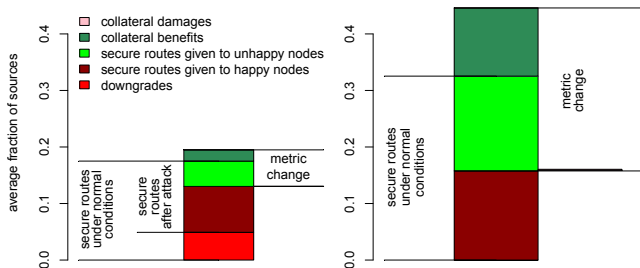


Figure 8: Changes in the metric explained. Sec 3rd (left) and Sec 1st (right).

metric by subtracting the collateral damages from the gains resulting from (a) offering secure routes to unhappy ASes and (b) collateral benefits. Fortunately, we find collateral damages to be a relatively rare phenomenon.

Results where security is 2nd look very similar to results when security is 3rd, with the addition of a small amount of collateral damage. The bottom line is, when security is 2nd or 3rd, (1) protocol downgrade attacks cause many secure routes that were available under normal conditions to disappear, and (2) those ASes that retain their secure routes during the attack would have been happy even if S*BGP had not been deployed; the result is meagre increases in the security metric. Meanwhile, when security is 1st, few downgrades occur, and the security metric is greatly improved.

7. RELATED WORK

Over the past decades several security extensions to BGP have been proposed; see [10] for a survey. However, proposals of new security extensions to BGP, and their subsequent security analyses typically assume that secure ASes will never accept insecure routes [6, 11], which is reasonable in the full deployment scenario where every AS has already deployed S*BGP [7, 10, 21]. There have also been studies on incentives for S*BGP adoption [11, 19]; these works suggest that “S*BGP and BGP will coexist in the long term” [19], which motivated our study of S*BGP in partial deployment.

Our work is most closely related to [21], which also measures “security” as the fraction of source ASes that avoid having their traffic intercepted by the attacking AS. However, [21] always assumes that the S*BGP variant is *fully deployed*. Thus, as discussed in Section 4.2, [21] also finds that fully-deployed origin authentication provides good security against attack we studied here (*i.e.*, announcing “*m, d*” using insecure BGP, see Section 3.1), but rightly assumes this attack fails against fully-deployed S*BGP. Moreover, [21] does not analyze interactions between S*BGP and BGP that arise during partial deployment (*e.g.*, Table 3).

Finally, [8] includes cryptographic analysis of S*BGP in partial deployment, and an Internet draft [26] mentions protocol downgrade attacks. However, neither explores how attacks on partially-deployed S*BGP can impact routing, or considers the number / type of ASes harmed by an attack.

8. CONCLUSION

On one hand, our results give rise to guidelines for partially-deployed S*BGP: (1) Deploying lightweight simplex S*BGP at stub ASes, instead of full-fledged S*BGP; this reduces deployment complexity at the majority of ASes without compromising overall security. (2) Incorporating S*BGP into

routing policies in a similar fashion at all ASes, to avoid introducing routing anomalies like BGP Wedgies. (3) Deploying S*BGP at Tier 2 ISPs, since deployments of S*BGP at Tier 1s can do little to improve security. On the other hand, we find that partially-deployed S*BGP provides, on average, limited security benefits over route origin authentication when ASes do not prioritize security 1st.

We hope that our work will call attention to the challenges that arise during partial deployment, and drive the development of solutions that can help surmount them. One idea is to find ways to limit protocol downgrade attacks, as these cause many of our negative results. For example, one could add “hysteresis” to S*BGP, so that an AS does not immediately drop a secure route when “better” insecure route appears. Alternatively, one could find deployment scenarios that create “islands” of secure ASes that agree to prioritize security 1st for routes between ASes in the island; the challenge is to do this without disrupting existing traffic engineering or business arrangements. Other security solutions could also be explored. For example, origin authentication with anomaly detection and prefix filtering could be easier to deploy (they can be based on the RPKI), and may be as effective as partially-deployed S*BGP.

Acknowledgments

We are grateful to BU and XSEDE for computing resources, and Kadin Tseng, Doug Sondak, Roberto Gomez and David O’Neal for helping us get our code running on various platforms. We thank Walter Willinger and Mario Sanchez for providing the list of ASes in each IXP that we used to generate our IXP-augmented AS graph, Phillipa Gill for useful discussions and sharing the results of [18] with us, and Leonid Reyzin, Gonca Gursun, Adam Udi, our shepherd Tim Griffin and the anonymous SIGCOMM reviewers for comments on drafts of this paper. This work was supported by NSF Grants S-1017907, CNS-1111723, ISF grant 420/12, Israel Ministry of Science Grant 3-9772, Marie Curie Career Integration Grant, IRG Grant 48106, the Israeli Center for Research Excellence in Algorithms, and a gift from Cisco.

9. REFERENCES

- [1] IRR power tools. <http://sourceforge.net/projects/irrpt/>, 2011.
- [2] Working group 6 Secure BGP Deployment Report. Technical report, FCC CSRIC http://transition.fcc.gov/bureaus/pshs/advisory/csric3/CSRICIII_9-12-12_WG6-Final-Report.pdf, 2012.
- [3] B. Ager, N. Chatzis, A. Feldmann, N. Sarrar, S. Uhlig, and W. Willinger. Anatomy of a large european IXP. In *SIGCOMM’12*, 2012.
- [4] Alexa. The top 500 sites on the web. <http://www.alexa.com/topsites>, October 1 2012.
- [5] B. Augustin, B. Krishnamurthy, and W. Willinger. IXPs: Mapped? In *IMC’09*, 2009.
- [6] I. Avramopoulos, M. Suchara, and J. Rexford. How small groups can secure interdomain routing. Technical report, Princeton University Comp. Sci., 2007.
- [7] H. Ballani, P. Francis, and X. Zhang. A study of prefix hijacking and interception in the Internet. In *SIGCOMM’07*, 2007.
- [8] A. Boldyreva and R. Lychev. Provable security of S-BGP and other path vector protocols: model, analysis and extensions. In *CCS’12*, pages 541–552.
- [9] M. A. Brown. Rensys Blog: Pakistan hijacks YouTube. http://www.rensys.com/blog/2008/02/pakistan_hijacks_youtube_1.shtml.

- [10] K. Butler, T. Farley, P. McDaniel, and J. Rexford. A survey of BGP security issues and solutions. *Proceedings of the IEEE*, 2010.
- [11] H. Chang, D. Dash, A. Perrig, and H. Zhang. Modeling adoptability of secure BGP protocol. In *SIGCOMM'06*, 2006.
- [12] Y.-J. Chi, R. Oliveira, and L. Zhang. Cyclops: The Internet AS-level observatory. *SIGCOMM CCR*, 2008.
- [13] Cisco. BGP best path selection algorithm: How the best path algorithm works. Document ID: 13753, May 2012. http://www.cisco.com/en/US/tech/tk365/technologies_tech_note09186a0080094431.shtml#bestpath.
- [14] J. Cowie. Rensys blog: China's 18-minute mystery. <http://www.rennesys.com/blog/2010/11/chinas-18-minute-mystery.shtml>.
- [15] A. Dhamdhere and C. Dovrolis. Twelve years in the evolution of the internet ecosystem. *Trans. Netw.*, 19(5):1420–1433, 2011.
- [16] L. Gao, T. Griffin, and J. Rexford. Inherently safe backup routing with BGP. *IEEE INFOCOM*, 2001.
- [17] L. Gao and J. Rexford. Stable Internet routing without global coordination. *Trans. Netw.*, 2001.
- [18] P. Gill, S. Goldberg, and M. Schapira. A survey of interdomain routing policies. NANOG'56, October 2012.
- [19] P. Gill, M. Schapira, and S. Goldberg. Let the market drive deployment: A strategy for transitioning to BGP security. *SIGCOMM'11*, 2011.
- [20] S. Goldberg, S. Halevi, A. D. Jaggard, V. Ramachandran, and R. N. Wright. Rationality and traffic attraction: Incentives for honest path announcements in BGP. In *SIGCOMM'08*, 2008.
- [21] S. Goldberg, M. Schapira, P. Hummon, and J. Rexford. How secure are secure interdomain routing protocols? In *SIGCOMM'10*, 2010.
- [22] T. Griffin and G. Huston. BGP wedgies. RFC 4264, 2005.
- [23] T. Griffin, F. B. Shepherd, and G. Wilfong. The stable paths problem and interdomain routing. *Trans. Netw.*, 2002.
- [24] G. Huston. Peering and settlements - Part I. *The Internet Protocol Journal (Cisco)*, 2(1), March 1999.
- [25] G. Huston. Peering and settlements - Part II. *The Internet Protocol Journal (Cisco)*, 2(2), June 1999.
- [26] S. Kent and A. Chi. Threat model for BGP path security. Internet draft: draft-ietf-sidr-bgpsec-threats-04, 2013.
- [27] S. Kent, C. Lynn, and K. Seo. Secure border gateway protocol (S-BGP). *JSAC*, 2000.
- [28] C. Labovitz. Arbor blog: Battle of the hyper giants. <http://asert.arbornetworks.com/2010/04/the-battle-of-the-hyper-giants-part-i-2/>.
- [29] C. Labovitz. Internet traffic 2007 - 2011. Global Peering Forum. Santi Monica, CA., April 2011.
- [30] C. Labovitz, S. Iekel-Johnson, D. McPherson, J. Oberheide, and F. Jahanian. Internet inter-domain traffic. In *SIGCOMM'10*, 2010.
- [31] M. Lepinski. Bgpsec protocol specification: draft-ietf-sidr-bgpsec-protocol-06. Internet-Draft, 2012.
- [32] M. Lepinski and S. Kent. *RFC 6480: An Infrastructure to Support Secure Internet Routing*.
- [33] R. Lychev, S. Goldberg, and M. Schapira. Network destabilizing attacks. In *PODC'12*, 2012.
- [34] R. Lychev, S. Goldberg, and M. Schapira. Is the juice worth the squeeze? BGP security in partial deployment. Technical report, Arxiv, 2013.
- [35] P. McDaniel, W. Aiello, K. Butler, and J. Ioannidis. Origin authentication in interdomain routing. *Computer Networks*, November 2006.
- [36] S. Misel. "Wow, AS7007!". Merit NANOG Archive, April 1997. <http://www.merit.edu/mail.archives/nanog/1997-04/msg00340.html>.
- [37] P. Mohapatra, J. Scudder, D. Ward, R. Bush, and R. Austein. *BGP Prefix Origin Validation*. Internet Engineering Task Force Network Working Group, 2012. <http://tools.ietf.org/html/draft-ietf-sidr-pfx-validate-09>.
- [38] P. Palse. Serving ROAs as RPSL route[6] Objects from the RIPE Database. RIPE Labs, June 2010. https://labs.ripe.net/Members/Paul_P_/content-serving-roas-rpsl-route-objects.
- [39] T. Paseka. Cloudflare blog: Why google went offline today., November 2012. <http://blog.cloudflare.com/why-google-went-offline-today-and-a-bit-about>.
- [40] A. Pilosov and T. Kapela. Stealing the Internet: An Internet-scale man in the middle attack, 2008. DEFCON'16.
- [41] Reuters. Internet providers pledge anti-botnet effort, March 22 2012.
- [42] M. Roughan, W. Willinger, O. Maennel, D. Perouli, and R. Bush. 10 lessons from 10 years of measuring and modeling the internet's autonomous systems. *JSAC*, 29(9):1810–1821, 2011.
- [43] R. Sami, M. Schapira, and A. Zohar. Searching for stability in interdomain routing. In *INFOCOM'09*, 2009.
- [44] Sandvine. Fall 2012 global internet phenomena, 2012.
- [45] K. Sriram. BGPSEC design choices and summary of supporting discussions. Internet-Draft: draft-sriram-bgpsec-design-choices-03, January 2013.
- [46] R. White. Deployment considerations for secure origin BGP (soBGP). draft-white-sobgp-bgp-deployment-01.txt, June 2003, expired.

APPENDIX

A. MORE COLLATERAL DAMAGE

Figure 7 revealed that collateral damages can be caused by secure ASes that choose *long* secure paths. When security is 1^{st} , collateral damages can also be caused by secure ASes that choose *expensive* secure paths:

Figure 9. We show how AS 4805, Orange Business in Oceania, suffers from collateral damage when security is 1^{st} . On the left, we show the network prior to S*BGP deployment. Orange Business AS4805 learns two routes: a legitimate route through its peer Optus Communications AS 7474, and a bogus route through its provider AS 2647. Since AS 4805 prefers peer routes over provider routes per our **LP** rule, it will choose the legitimate route and avoid the attack. On the right, we show what happens after S*BGP deployment. Now, Optus Communications AS 7474 has started using a secure route. However, this secure route is through its provider AS 7473. Observe that AS 7474 is no longer willing to announce a route to its peer AS 4805 as this would violate the export policy **Ex**. AS 4805 is now left with the bogus provider route through AS 2647, and becomes unhappy as collateral damage.

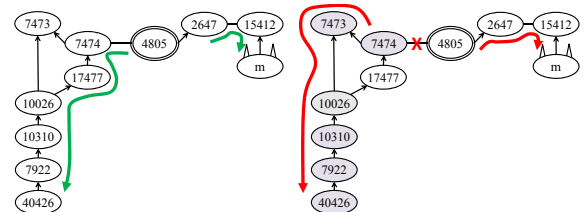


Figure 9: Collateral damages; security 1^{st} .