

Identifying Networks Vulnerable to IP Spoofing

Osvaldo Fonseca[†] Ítalo Cunha[†] Elverton Fazzion^{†‡} Wagner Meira Jr.[†]
Brivaldo Junior^{*} Ronaldo A. Ferreira^{*} Ethan Katz-Bassett[‡]

[†]Universidade Federal de Minas Gerais, Brazil

[‡]Universidade Federal de São João del-Rei, Brazil

^{*}Universidade Federal de Mato Grosso do Sul, Brazil

[‡]Columbia University

Abstract—The lack of authentication in the Internet’s data plane allows hosts to falsify (*spoof*) the source IP address in packet headers. IP source spoofing is the basis for amplification denial-of-service (DoS) attacks. Current approaches to locate sources of spoofed traffic lack coverage or are not deployable today. We propose a mechanism that a network with multiple peering links can use to coarsely locate the sources of spoofed traffic in the Internet. The idea behind our approach is that a network can monitor and map spoofed traffic arriving on a peering link to the set of sources routed toward that link. We propose mechanisms the network can use to systematically vary BGP announcement configurations to induce changes to Internet routes and to the set of sources routed to each peering link. A network using our technique can correlate observations over multiple configurations to more precisely delineate regions sending spoofed traffic. Evaluation of our techniques on the Internet shows that they can partition the Internet into small regions, allowing targeted intervention.

Index Terms—IP spoofing, security, amplification, denial-of-service, routing policies, topology discovery

I. INTRODUCTION

The lack of authentication in the Internet’s data plane allows hosts to falsify (*spoof*) the source IP addresses of their traffic and send unsolicited traffic to arbitrary destinations. These vulnerabilities form the basis for amplification denial-of-service attacks [1], which have been effectively employed against large-scale distributed service providers (e.g., [2]–[7]). The spoofed source addresses make the origins of such attacks seemingly untraceable, complicating attribution, mitigation efforts to squelch the attack, or targeted efforts to convince networks to disallow spoofed traffic.

Over the last two decades, researchers have proposed dozens of *IP traceback* techniques for identifying the routes taken by spoofed packets [8]–[14]. Approaches include temporarily congesting links to perturb (attack) traffic [8], modifying routers to encode information (usually in the IP ID field) about routers traversed by a small fraction of packets [9]–[11], modifying routers to send information about a fraction of forwarded packets towards destinations [12], or modifying routers to store packet digests and provide an interface for querying for a packet’s signature [13], [14]. Despite all the research, none of these approaches has been deployed and increased our ability to locate the origins of spoofed traffic, because they require changes to routers, cooperation from other networks, and wide deployment to provide accurate identification. Since these techniques face nearly insurmountable barriers to adoption, today’s networks get a single data point on the spoofed traffic’s route: which peering link receives the traffic.

In this paper, we explore how a network can manipulate this information source—the peering link where traffic ingresses a network—to more precisely locate sources of spoofed traffic. Our key observation is that the routes are partially under an origin network’s control, and so the network receiving the spoofed traffic has some ability to impact on which link it receives traffic, instead of relying on routers that are not under its control. We propose techniques that are fundamentally different from existing traceback approaches and can be used today, requiring no changes to deployed equipment nor cooperation from other networks. Our techniques work best when the spoofed traffic originates from few sources, as is common in amplification DoS attacks [15].

With our approach, a network announces an IP prefix through multiple peering links, a practice known as *anycast*. Each link attracts traffic from non-overlapping regions of the Internet called the link’s *catchment*. The network can infer the sources in each catchment by inspecting non-attack traffic at each ingress link and mapping the source IP addresses to their respective prefixes and controlling autonomous systems (ASes), or by sending out pings and measuring which link replies arrive at [16]. To measure the amount of spoofed traffic on each link, the network can run an amplification honeypot that does not receive legitimate traffic (e.g., AmpPot [15]) or infer the set of valid source addresses from each peering link and label the traffic from other addresses as spoofed [17], [18]. The amount of spoofed traffic arriving at each peering link can then be attributed to the sources routed toward that link. Many sources are routed toward the same peering link, however, so simply attributing an attack traffic volume to a peering link is not precise enough to isolate attack sources.

To track down sources of spoofed traffic, we present systematic approaches to vary IP prefix announcement configurations that allow networks to induce changes to routes toward their prefixes and, more importantly, in the set of ASes routed toward each peering link (the catchment). Networks using our techniques can correlate spoofed traffic observed from sets of sources across multiple announcement configurations to infer regions of the Internet sourcing spoofed traffic. Figure 1 provides intuition for how such measurements can be combined to identify networks that allow spoofed packets:

- In Configuration 1, the operator announces a prefix through three peering links with networks m , n , and p ; measures the catchment (colored polygons) and traffic arriving on each peering link; and identifies that the spoofed traffic is concentrated on the link with n , i.e.,

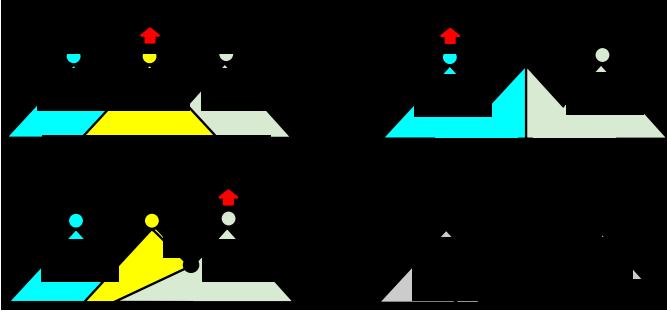


Figure 1: Example with catchments and resulting clusters for three announcement configurations performed by an origin network peering with ASes m , n , and p . On the bottom right we show the resulting clusters, denoted by Greek letters.

sent by networks in n 's catchment (red arrow).

- The operator later withdraws the announcement to n (Configuration 2), measures catchments and traffic volumes again, and identifies that the spoofed traffic is now concentrated on the peering link with m .
- Configuration 3 announces the prefix from n again, but poisoning AS u (which causes AS u to ignore the route from n and choose the route from p instead). The operator can measure catchments and traffic to identify that the spoofed traffic is concentrated on the peering link with p .
- Finally, the operator can intersect the measured catchments to partition networks into *clusters* (bottom right), and correlate clusters with observed spoofed traffic (red arrows) to identify that the spoofed traffic is concentrated on networks comprising λ .

We evaluate our techniques running experiments on the PEERING platform [19]. We deploy 1572 different announcement configurations from eight peering links, identifying multiple, different routes from each source covered in our measurements. We show that correlating information across multiple announcement configurations on PEERING allows us to partition the Internet into small regions with as few as one AS and with 1.29 ASes on average. The small size makes these regions candidate targets for countermeasures or notifications. Our results indicate that networks with peering footprints larger than PEERING's, as is the case for most regional transit networks, can more effectively manipulate routes to achieve even higher accuracy and quicker localization.

Our techniques allow identification of networks that do not employ BCP38 (ingress filtering) [20] and allow spoofed traffic, helping Internet bodies focus efforts and drive adoption of best practices. They can also be used to drive automatic DoS mitigation systems that use, e.g., BGP communities to trigger remote traffic blackholing [21] or BGP flowspec to configure traffic filters [22].

A previous version of this paper appeared on IFIP Networking [23]. This version includes several improvements over the original paper, including the collection of a larger dataset (§IV), more detailed analysis (§V), operational considerations (§VI), and deeper discussion of related work (§VIII).

II. BACKGROUND ON BGP

The BGP best-path selection algorithm defines the preferred route to an IP prefix as the route with the highest *local preference* (LocalPref), a value set by the AS according to private routing policies. If multiple routes have the same LocalPref, BGP chooses the route with the shortest AS-path length. If multiple routes remain tied for best, BGP applies other tiebreakers that include intra-domain (IGP) routing costs, hints received from neighboring ASes (MED), and route age (to reduce oscillations) [24].

An AS that controls an IP prefix can configure its BGP announcements to influence routes, e.g., to achieve traffic engineering goals [25], [26]. First, an AS can announce (anycast) an IP prefix from all or a subset of its peering links. This strategy is used by content distribution networks so remote ASes, and users therein, route to a topologically close location, improving performance and increasing reliability [25], [27]. Second, an AS can influence BGP's tie breaking at remote ASes by *prepending* its AS number to the announcement's AS-path, making the AS-path artificially longer. This strategy is used by multihomed ASes to signal on which link it prefers to receive traffic [26], [28]. Third, an AS can influence the use and propagation of its announcements through a remote AS using *BGP poisoning* [24], [29]–[31]. A poisoned announcement targets one or more ASes, and includes the target ASes' numbers in the AS-path; this triggers loop prevention and causes poisoned ASes to ignore the announcement.

III. LOCATING SOURCES OF SPOOFED TRAFFIC

We define an *announcement configuration* for an IP prefix as a triple $c = \langle \mathcal{A}_c; \mathcal{P}_c; \mathcal{Q}_c \rangle$. We denote the set of peering links of an origin AS by $\mathcal{L}$. $\mathcal{A}_c \subseteq \mathcal{L}$ is the set of locations from which the prefix is announced. Each location in $\mathcal{A}_c$ announcing the prefix will attract traffic from non-overlapping regions of the Internet that we call a *catchment*. $\mathcal{P}_c \subseteq \mathcal{A}_c$ is the set of locations where the prefix is announced with prepending, and $\mathcal{Q}_c$ is a mapping from announcement locations in $\mathcal{A}_c$ to sets of poisoned ASes. We drop the subscripts when the configuration is clear from context. For example, consider an AS with four peering links labeled from l_1 to l_4 . A configuration $c = \langle \{l_1, l_2\}; \{l_1\}; \{l_1: \emptyset, l_2: \{a, b\}\} \rangle$ means the prefix is announced through peering link l_1 with AS-path prepending, announced through peering link l_2 poisoning ASes a and b , and not announced through links l_3 and l_4 .

A. Systematic Route Changes

We propose a method that an *origin* AS with multiple peering links can use to generate announcement configurations that systematically induce route and catchment changes.

a) *Varying announcement locations*: Announcing a prefix from more peering links increases route diversity and leads to smaller catchments, on average. Smaller catchments provide better localization of spoofed traffic.

We propose that the origin AS deploy a sequence of configurations starting by announcing from all available peering links, i.e., $\mathcal{A} = \mathcal{L}$; then make announcements from all proper subsets

of available locations $\mathcal{L}$ in decreasing size order. Deploying all configurations removing up to r links from $\mathcal{L}$ is guaranteed to discover *at least* $r + 1$ routes for *all* sources in the Internet. Whenever we withdraw the prefix from the peering link a source is routed to, that source will need to be routed to an alternate link. This is a deterministic way to uncover route diversity that scales with a network’s peering footprint (i.e., the size of $\mathcal{L}$).

In Figure 1, Configuration 1 shows catchments when the origin AS announces a prefix through three peers: m , n , and p ; Configuration 2 shows catchments when the origin AS announces through m and p only.

b) Varying the AS-path length with BGP prepending:

For any given announcement configuration, a router may have multiple routes with the same LocalPref to choose from. In these cases, the router chooses the preferred route based on the AS-path length or subsequent BGP tiebreakers.

Given a configuration with a set of announcement locations $\mathcal{A} \subseteq \mathcal{L}$, we propose that the origin AS generate and deploy additional configurations prepending announcements from subsets of locations $\mathcal{P} \subseteq \mathcal{A}$, in increasing size order. To make prepended routes longer than most other routes, the origin can prepend its AS number four times, which is longer than most AS-paths in the Internet [32]. Deploying configurations that prepend announcements from all combinations of up to s locations induces BGP’s tie-breaking mechanism to choose up to s alternate routes. More precisely, prepending will cause a router to change away from its (previously shorter and preferred) route whenever an alternate route with the same LocalPref and no prepending is available.

Manipulating BGP tiebreakers like the AS-path length is a general idea. Unfortunately, BGP tiebreakers after the AS-path length cannot be controlled (e.g., IGP costs) or do not propagate to distant ASes (e.g., MED), and thus cannot be employed by the origin for route manipulation.

c) Controlling route propagation with BGP poisoning:

BGP AS-path prepending is ineffective when routers choose routes based on LocalPref, i.e., before applying BGP tiebreakers. In these cases, the origin AS can still try to induce route changes by making a remote router’s preferred route (with highest LocalPref) unavailable using BGP poisoning. The origin AS can try to induce routers in a remote AS r to change routes by poisoning r (or other intermediate ASes between itself and r) in some announcements. Target ASes for BGP poisoning can be chosen using different strategies depending on the goal [24], [29]–[31].

Figure 2a shows the routes used by each AS to reach AS o when the origin o announces (anycasts) the prefix to all neighbors without poisoning. Figure 2b shows the routes when the origin o poisons AS u on announcements through link $o-n$, announcing an AS-path starting with $O-U-O$ instead of just O . Poisoning an upstream AS u that is a neighbor of AS n will prevent routes (and traffic) from traversing the link $n-u$ (red X in Figure 2b) because it triggers BGP loop prevention when AS u observes itself in the received route, causing routing changes at all sources previously routed through link $n-u$.

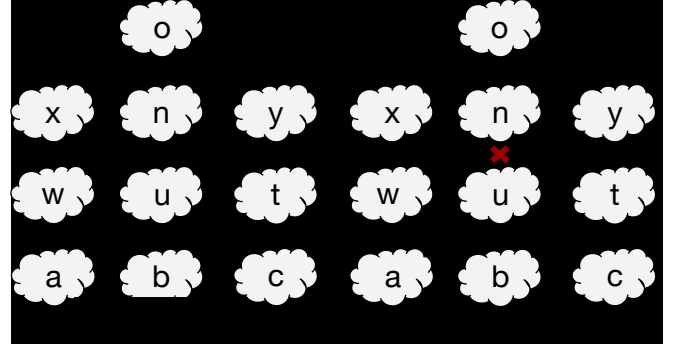


Figure 2: Example of AS o poisoning AS u through link $o-n$ to force all ASes previously routing through link $n-u$ to choose a different route.

Therefore, ASes a , b , c , and u need to find an alternate path to reach AS o (dashed lines). Configuration 3 in Figure 1 illustrates the change in catchments when we poison AS u on the announcement through n .

Poisoning is similar and complementary to our proposal to control announcement locations (choosing $\mathcal{A}$), as it attempts to control route propagation through a network’s links. In particular, poisoning directly-connected neighbors is equivalent to withdrawing the announcement. For example, in Figure 2, poisoning n on the announcement to n is equivalent to withdrawing the announcement.

Unfortunately, BGP poisoning may be ineffective: an AS may disable BGP loop prevention for traffic engineering, e.g., when interconnecting multiple sites over the Internet by announcing different prefixes from each site; and ASes may filter poisoned announcements, e.g., tier-1 ASes often filter announcements from clients whose AS-path contains other tier-1 ASes, as such announcements normally indicate a route leak [33]. We use BGP poisoning as a best-effort approach to complement the previous two techniques, which are more reliable.

B. Correlating Observations

We define a *cluster* as a set of sources that are in the same catchment across all announcement configurations. We start by placing all sources into a single cluster. We iterate over all catchments in all configurations; for each catchment α we iterate over all clusters κ identified so far and split any cluster κ that overlaps α into up to two clusters: $\kappa \cap \alpha$ and $\kappa \setminus \alpha$ (we do not split κ if $\kappa \cap \alpha = \kappa$). The bottom right corner of Figure 1 shows the clusters obtained after performing the three announcement configurations.

Our techniques generate different announcements to induce route changes with the goal of reducing the size of clusters. Small clusters allow the identification of networks responsible for sending spoofed packets and enable targeted intervention.

C. Estimating Volume of Spoofed Traffic

An origin AS can estimate the presence or volume of spoofed traffic received on each catchment by hosting a

Table I: PoPs and providers of the PEERING platform used in the experiments.

PoP	Transit Provider
AMS-IX	Bit BV (AS12859)
Georgia Tech	Georgia Institute of Technology (2637)
NEU	Northeastern University (AS156)
Seattle-IX	RGnet (AS3130)
UFMG	RNP (AS1916)
UTAH	Utah Education Network (AS210)
UW	Pacific Northwest GigaPoP (AS101)
Wisconsin	University of Wisconsin System (3128)

honeypot that emulates a service vulnerable to (but that does not contribute to) amplification attacks to attract spoofed traffic [15]. Another approach is to infer legitimate sources for each peering link and label all traffic received from other sources as spoofed [17], [18].

IV. EXPERIMENTAL SETUP

We evaluate our techniques in the Internet by making announcements from the PEERING platform [19]. PEERING is a research platform that operates an AS with multiple points-of-presence (PoPs) in various locations spread across three continents. We make announcements from eight PEERING PoPs, using one provider at each PoP. At PEERING PoPs at IXPs, which have multiple providers and peers, we choose one provider and use it throughout the experiment. Table I summarizes information about the PoPs and providers we used. We next describe how we generate configurations using our techniques and how we measure catchments.

A. Announcement Configurations

We perform exhaustive announcements to obtain a complete dataset that allows the evaluation of different aspects of the proposed heuristics. We start with a configuration that announces (anycasts) a prefix to one transit provider in each of eight active PEERING PoPs. As an attempt to induce as many path changes as possible and, consequently, reduce cluster sizes, we consider 255 possible configurations withdrawing from all possible subsets of locations; i.e., we set $r = 7$ and deploy all configurations withdrawing from up to seven locations. For each such configuration c , we generate $|\mathcal{A}_c|$ additional configurations, prepending from each active location in turn. This requires an additional $\sum_{x=0}^7 [8-x] \binom{8}{8-x} = 1024$ configurations.

We deploy 293 additional configurations poisoning targeted ASes, one at a time, while announcing from all available locations without prepending; i.e., we deploy configurations $\langle \mathcal{L}; \emptyset; \{l : \{a\}\} \rangle$, where a is the target AS and l is the location used in a 's preferred route. As a tradeoff between completeness and number of announcements, we select as targets ASes that provide transit to three or more downstream ASes in the initial configuration anycasting from all locations, i.e., $\langle \mathcal{L}; \emptyset; \emptyset \rangle$. We do not poison tier-1 and tier-2 ASes, as these large networks frequently filter announcements received from customers whose AS-paths contain other tier-1 or tier-2 ASes [33]. We deploy a total of 1572 configurations.

B. Measuring Catchments

PEERING prefixes carry no production traffic, so we cannot passively observe traffic to infer catchments. Also, concerns about executing Internet-wide scans from the PEERING platform limits our ability to issue measurements from the platform to the wide-area Internet. Instead, we measure catchments using a combination of AS-paths observed on BGP update messages toward PEERING prefixes collected from public feeds and traceroutes issued from RIPE Atlas toward PEERING prefixes [34]. We use all public BGP feeds from RouteViews [35] and RIPE RIS [36]. We partnered with RIPE and received permission to issue traceroute measurements every 20 minutes from 4800 RIPE Atlas probes, a measurement rate that is $7\times$ larger than normally supported, which helps increase the coverage of our measurements. We note the measurement rate is still low, only targets PEERING prefixes, and has not raised complaints.

We keep each announcement configuration active for 50 minutes to wait for route convergence and ensure, with high probability, that we collect at least two rounds of traceroutes *after* routes to our prefixes have converged, as convergence takes less than 2.5 minutes 99% of the time [31].

We map traceroute hops into ASes using IP-to-AS data from Team Cymru [37] and using IXP-specific data from PeeringDB [38]. In a traceroute measurement, if consecutive unresponsive hops are surrounded by responsive ones, we check whether the surrounding hops have a single sequence of responsive hops between them in other traceroutes; if that is the case, we substitute the unresponsive hops with the responsive ones. After this step, we map unresponsive hops whose surrounding responsive hops map to a single AS a to the same AS a . If surrounding hops map to different ASes, we check whether public BGP feeds have a single sequence of ASes between them in AS-paths; if that is the case, we substitute the unresponsive hops to match the public AS-paths. If we still have unmapped or unresponsive hops, we ignore those hops on the AS-level path.

C. Source Granularity

Our techniques are orthogonal to the granularity at which sources are defined. The only requirement is that each source appears in at most one catchment for each announcement configuration. For the evaluation, we define sources at the AS granularity. Different routers within an AS may choose different routes to a destination [39], e.g., routers in the US and Europe may choose different routes toward the announced prefix. In our dataset, this may also happen due to incorrect IP-to-AS mapping. Whenever we observe multiple routing decisions by an AS from multiple vantage points, we give higher priority to BGP measurements (over traceroute) to minimize errors due to IP-to-AS mapping. If multiple measurements of the same type remain, we assign the AS to the catchment most common across the available measurements. On average, we observe 4.35% of ASes in multiple catchments in an announcement configuration.

D. Source Visibility

Our dataset covers a total of 3636 ASes across both BGP and traceroute measurements. However, changing announcement configurations forces route changes and may cause ASes to route through less preferred alternate or backup paths, leading to some ASes appearing in just a few configurations (e.g., where a specific backup link is used). Whenever an AS is not observed during an announcement configuration, we cannot infer the catchment where the missing AS belongs.

We address this limitation in two steps. First, we only consider 2398 ASes that are observed in at least 80% of all configurations. This set includes all tier-1 ASes as well as 73.8% of ASes with customer cone larger than 300 ASes [40]. We use CAIDA’s inferred clique [40] as the set of tier-1 ASes.

Second, we compute the frequency that an AS a and each other AS are in the same catchment across all configurations where a was observed. We define $S[a]$ as the AS which a appears most frequently with (i.e., a and $S[a]$ route similarly). For each deployed configuration where source a was not observed, we assign a to the same catchment as $S[a]$. We apply this search recursively: if $S[a]$ is also missing, we assign a to the same catchment as $S[S[a]]$, and so on. This approach is conservative as it keeps catchments similar across configurations, and prevents separation of a and $S[a]$ into different clusters due to missing coverage.

E. Ethical Concerns

PEERING prefixes do not carry production or user traffic, so no user traffic was impacted during our experiments. In practice, we expect networks to use prefixes dedicated to the location of sources of spoofed traffic to avoid impacting real users or applications.

Our route manipulation does not affect other prefixes or networks in the Internet. We note that anycast and AS-path prepending are common traffic engineering practices and can be observed on thousands of prefixes in routing tables today. The rate at which we inject BGP updates is negligible relative to BGP churn [41].

BGP poisoning has been used for more than a decade in research as a mechanism to route around failures and traffic blackholes [31], identify static default routes [29], discover network links [30], and characterize interdomain routing policies [24]. BGP poisoning does *not* impact the poisoned AS, routes to its prefixes, or its traffic. The PEERING platform conservatively limits each announcement to two poisoned ASes. To clearly signal BGP poisoning, PEERING requires experiments to surround each poisoned AS with PEERING’s own AS47065. This avoids incorrect inference of peering links from BGP AS-paths (any false links inferred from poisoning would be with AS47065, which is easy to filter), and makes attribution to the PEERING platform trivial. PEERING maintains a blacklist of ASes that opt-out of BGP poisoning from the platform, but this list is currently empty as no ASes have complained about poisoning.

V. EVALUATION

In this section, we provide results on cluster sizes and show that we can manipulate routes to locate sources of spoofed traffic with good precision. Our results also indicate that large networks can apply the technique proposed to even greater effect, and that our techniques may potentially be used to locate sources of spoofed traffic during DDoS attacks.

A. Cluster Sizes

Figure 3 shows the complementary cumulative distribution of cluster sizes, with logarithmic scales on both axes. We show one distribution at the end of each phase (i.e., after 255 configurations varying locations, 1279 varying locations and prepending, and 1572 including all techniques). We find all our techniques are effective in reducing cluster sizes. After deploying all 1572 configurations built by the three techniques, 84.7% of clusters have a single AS (95% of clusters have one or two ASes). This indicates that, depending on the number and locations of the sources of spoofed traffic, our techniques may precisely locate them. Although most clusters are small and most ASes are in small clusters, large clusters account for a non-negligible number of ASes. After deploying all 1572 configurations, 15 clusters are larger than 5 ASes and contain 6.9% of the ASes considered in our study. Reducing cluster sizes at the tail is important to identify sources of spoofed traffic into small clusters and allow targeted intervention.

The lines in Figure 5 show the mean and maximum of cluster sizes as a function of the number of configurations. Axes use a logarithmic scale, and we sort configurations by the order in which they were deployed. We indicate when each phase finishes with vertical lines. We observe diminishing returns in our ability to reduce the mean cluster sizes (solid blue line) by deploying additional announcement configurations. However, an origin AS *can* effectively manipulate routes toward each prefix, systematically causing catchment changes even after hundreds of configurations. In particular, the results indicate we could have obtained even smaller clusters by performing more announcements. The evolution of the size of the largest cluster (dashed gray line) has a different behavior. We observe that the line may remain stable for several configurations while the mean cluster size decreases. The plateaus in the evolution of the size of the largest cluster indicate that there are some clusters harder to partition and might be more quickly split by deploying targeted announcement configurations. The small steps following the vertical bars indicate that changing techniques used to generate configurations induces different route changes (new routes) that reduce the largest cluster size.

B. Longitudinal Analysis

The dataset we analyze in this work includes more announcement configurations and has better coverage than the dataset we collected in 2019 and analyzed in an earlier version of this work [23]. Table II compares the two datasets.

Figure 4 is equivalent to Figure 3, but it shows results from the 2019 dataset. The graphs are qualitatively similar, providing evidence that running the technique multiple times

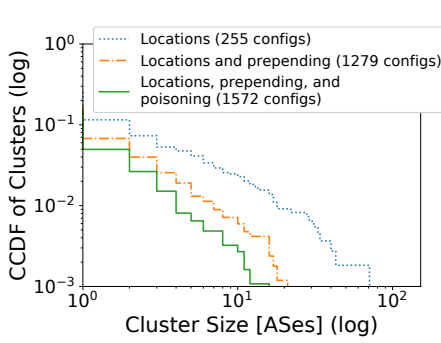


Figure 3: Distribution of cluster sizes after each phase.

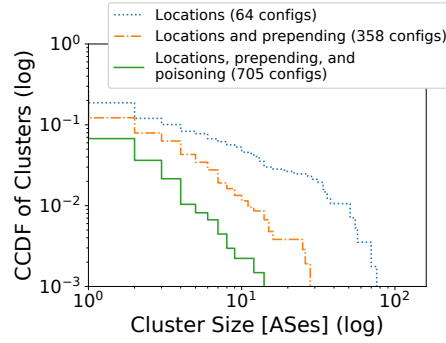


Figure 4: Distribution of cluster sizes on smaller 2019 dataset [23].

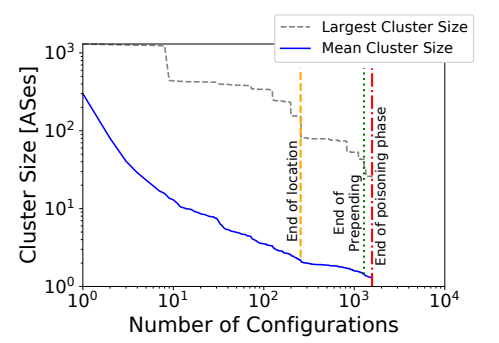


Figure 5: Cluster sizes as function of number of configurations.

Table II: Datasets Compared in Longitudinal Analysis

Dataset	PEERING PoPs	Number of Configurations			ASes Covered
		Locations	Prepending	Poisoning	
2019 [23]	7	64	294	347	1885
2020	8	255	1024	293	2398

on one network (PEERING) would yield similar results. We also note that although our new dataset has larger coverage, it still manages to separate ASes into smaller clusters than the 2019 dataset. This follows from a larger number of announcement configurations being able to force route changes that effectively contribute to splitting clusters.

C. Localization Speed

The number of possible announcement configurations grows exponentially with the number of peering links $|\mathcal{L}|$. A straightforward approach to speed up localization is to use multiple prefixes and deploy multiple configurations concurrently. This approach, however, requires spare IP space, which may be limited in IPv4. In the following we discuss heuristics that do not depend on additional resources.

When locating the sources of spoofed traffic, a network can reuse previous catchment measurements or remeasure catchments during identification at run time. For example, an origin AS employing our techniques can deploy time-consuming announcement configurations and measure catchments prior to the occurrence of an amplification DDoS attack. While an attack is ongoing, the origin AS can then assume that catchments remain unchanged since their last measurement and deploy configurations in optimal order to quickly reduce cluster sizes. This involves a trade-off between identification accuracy (reusing previous catchment measurements may incur errors due to route changes since the last measurement) and identification delay (measuring catchments during identification takes time), which depends on route stability and could be improved by resource-efficient solutions for inferring path changes.

The solid line in Figure 6 shows the mean cluster size as function of the number of announcement configurations when the origin AS chooses the sequence of configurations at ran-

dom, without repetition. The shaded area shows the variance across 1,000 random sequences. The dashed line shows the mean cluster size as a function of the number of announcement configurations deployed when the origin AS greedily chooses the configuration that results in the smallest mean cluster size before deploying each configuration. Compared to the solid line, we observe that localization can be made significantly faster if catchments are measured prior to an attack, and configurations deployed in optimal order. For example, after running ten configurations, while the random sequence yields a mean cluster size of 5.48 ASes, the optimal sequence yields a cluster size of 3.1 ASes. While our techniques can be used offline to identify networks that allow spoofed packets and drive adoption of filtering, this result indicates the techniques might be useful at run time as a source of information for active attack mitigation mechanisms.

Another approach to increase localization speed is to predict the catchments of announcement configurations and only deploy the most promising configurations, discarding or postponing configurations predicted to provide little additional information (i.e., not induce new, different route changes).

An AS in the Internet can receive multiple routes from its neighbors via BGP and choose the best according to its policy. We evaluate routing choices of ASes in our dataset according to BGP's first two decision criteria: (i) *best relationship*, which states an AS prefers routes through a client network first, through a peer second, and through a provider last; and (ii) *shortest path*, which states that, when multiple equally-preferred routes are available (tied according to the relationship criterion), the AS chooses the shortest one. We do not consider additional decision criteria as we cannot observe them from AS-paths we collect from BGP updates or traceroute measurements.

We compare the route chosen by each AS a with alternate routes AS a had available during deployment of configuration c . Because it is often impossible to know what (alternate) routes AS a has available while configuration c is deployed, we consider AS a has available all routes observed across all configurations that *could* exist in configuration c . More precisely, any route observed from AS a that goes to a location

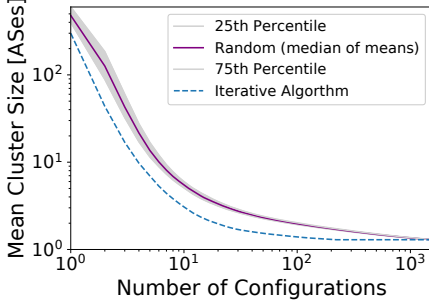


Figure 6: Mean cluster size as function of announcement schedule.

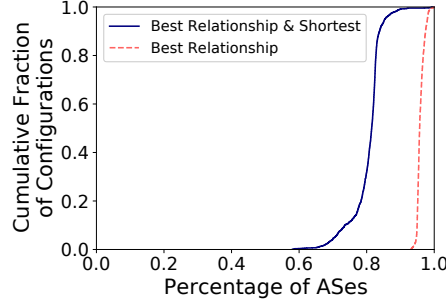


Figure 7: Percentage of ASes following well-known routing policies across configurations.

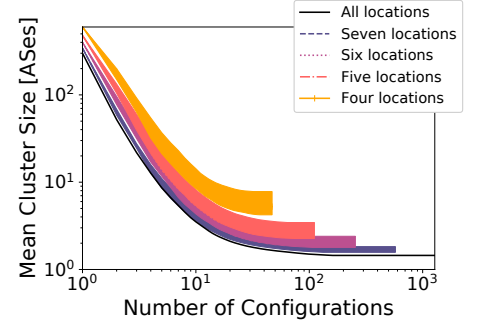


Figure 8: Mean cluster size when removing peering locations.

in $\mathcal{A}_c$, has prepending if and only if it goes to a location in $\mathcal{P}_c$, and does not traverse a poisoned AS in $\mathcal{Q}_c$ is considered as a possible alternate route. We compare the next hop neighbor on the route chosen by AS a in configuration c with the next hop neighbor on each of the possible alternate routes. We say AS a 's choice follows the best relationship criterion if the chosen next hop neighbor is more preferred or tied for most preferred when compared to all possible alternate routes. In addition, we say the choice follows the shortest path criterion if the route is shorter or tied for shortest when compared to all possible alternate routes that are tied for best relationship.

The dashed line in Figure 7 shows the distribution of the fraction of ASes observed to follow the best relationship criterion across all configurations. We observe that most ASes in the Internet follow the best relationship criterion. The solid line shows the distribution of ASes observed to follow *both* best relationship and shortest path criterion (also referred to as the Gao-Rexford model [42]). This result indicates that most ASes in the internet follow a well-defined, known behavior. Note that our methodology considers the existence of alternate routes that do not exist in practice, making these results a worst-case scenario. In practice, an even larger fraction of ASes than what we report may follow the Gao-Rexford model. Although predicting routes in the Internet is challenging [43], our techniques could benefit from future advances in catchment prediction strategies to boost attack localization speed.

D. Usefulness of Techniques

To evaluate the importance of each announcement generation technique (i.e., varying announcement locations, prepending, and poisoning), we investigate which ones were used by the first 100 announcement configurations chosen by the greedy algorithm (dashed line in Figure 6). Among the first 100 announcement configurations, we find 18 vary announcement locations, 8 add prepending to a set of locations $\mathcal{A}$ seen in a previously-deployed configuration, 50 include prepending and announce to a set of locations $\mathcal{A}$ not seen in previous configurations, and 24 poison a specific target AS. The spread across all techniques indicates that, with prior knowledge of

catchments, all techniques would be actively employed to quickly track down sources of spoofed traffic.

E. Impact of Peering Footprint

Figure 8 is similar to Figure 6, but it shows different lines for cases where we consider only a subset of our configurations, emulating networks with fewer PoPs by discarding one, two, three, or four of the eight PoPs we used. The “all locations” line includes all $255 + 1024 = 1279$ configurations using all 8 locations and prepending. The “seven locations” line includes a subset of $\sum_{x=0}^6 \left\{ \binom{7}{7-x} + [7-x] \binom{7}{7-x} \right\} = 575$ configurations using up to 7 locations. The shaded area shows the minimum and maximum mean cluster sizes across all $\binom{8}{1} = 8$ possible subsets (each subset discarding one of the 8 PEERING PoPs we used). The “six locations” line includes a subset of $\sum_{x=0}^5 \left\{ \binom{6}{6-x} + [6-x] \binom{6}{6-x} \right\} = 255$ configurations using up to 6 locations, and the shaded area shows the minimum and maximum mean cluster sizes across all $\binom{8}{2} = 28$ combinations of six locations. Lines for five and four locations are similar. We sort the announcement configurations on the x axis using the greedy algorithm to focus on the number of locations used in announcements, not their sequencing.

The graph shows that having more locations allows the generation of more configurations, leading ultimately to smaller cluster sizes, and also yields smaller cluster sizes for the same number of announcements. This result indicates that a network with a footprint larger than PEERING's could achieve even higher localization precision.

Figure 9 is similar to Figure 3 and shows the complementary distribution of cluster sizes but considering fewer announcement locations. The five lines and shaded areas correspond to the same five scenarios described in Figure 8. The “all locations”, “seven locations”, “six locations”, “five locations”, and “four locations” lines show the distributions of cluster sizes after 1279, 575, 255, 111, and 47 announcements, respectively; in other words, we show the distribution of cluster sizes at the end of the curves in Figure 8. We observe that using less announcement locations leads to larger cluster

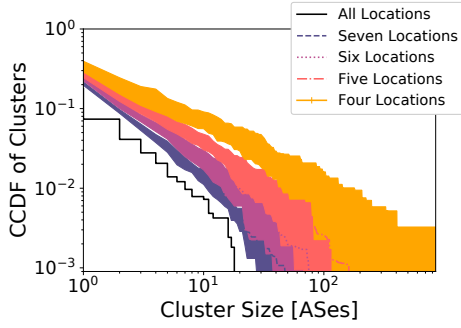


Figure 9: Distribution of cluster size after removing locations.



Figure 10: Cluster size as function of AS-hop distance from origin AS.

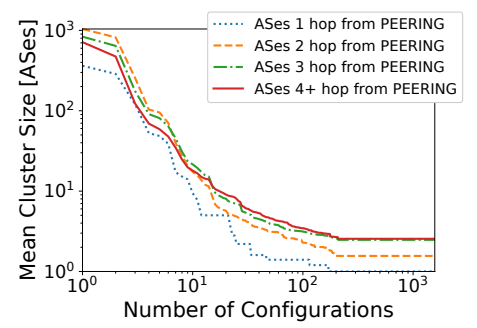


Figure 11: Cluster sizes as function of number of configuration by AS-hop distance.

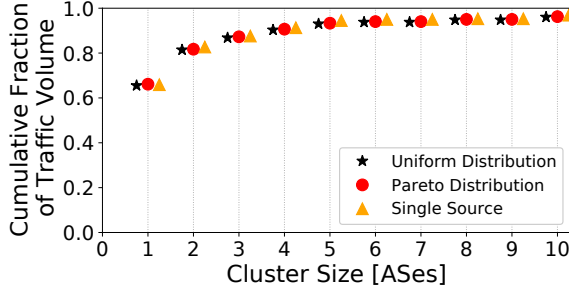


Figure 12: Distribution of cluster size as function of traffic volume for different distributions of spoofed sources.

sizes at the tail. While the “all locations” line shows 0.7% of clusters with more than 10 ASes, the “seven locations”, “six locations”, “five locations”, and “four locations” lines show 1.48%, 2.40%, 3.75%, and 6.75% of clusters with more than 10 ASes, respectively.

We evaluated the distribution of cluster sizes as a function of the distance, in number of AS-hops, between PEERING PoPs and ASes. Figure 10 shows the distribution of cluster sizes across all ASes in our dataset. We break ASes into groups based on their AS-hop distance to the closest PEERING location observed in the first anycast configuration. We find that ASes that are 1 or 2 AS-hops away from PEERING PoPs are in clusters with 1.55 ASes on average, while ASes 3 or more AS-hops away are in clusters with 2.51 ASes on average. As we expected, ASes closer to announcement locations are easier to isolate (in smaller clusters), but most ASes farther away are also in small clusters, indicating that we may still be able to identify sources of attacks that are farther away with actionable precision. Figure 10 shows that large clusters (e.g., with 10 ASes or more) are usually farther away from announcement locations.

Figure 11 shows the average cluster size as a function of the number of deployed configurations (similar to Figure 8). We group ASes by their AS-hop distance from PEERING and sort the announcement configurations on the x axis using the

greedy algorithm. We observe that the mean cluster size of ASes closer to PEERING (blue and orange dashed lines) is smaller than the mean cluster sizes of ASes farther away (green and red lines) after the deployment of a few configurations. More importantly, however, we note that cluster sizes decrease for all ASes throughout the experiment regardless of distance from PoPs. As future work, we plan to investigate targeted poisoning of distant ASes to induce specific route changes to split these distant clusters.

F. Study of Spoofed Traffic

Given the impossibility of attracting real spoofed traffic using PEERING testbed resources,¹ we study identification accuracy using simulation. We perform simulations where we choose the number of sources of spoofed traffic across ASes according to the uniform and Pareto distributions. For the Pareto distribution, we set the shape parameter such that 80% of sources of spoofed traffic are concentrated in 20% of ASes. We also run simulations with a single source of spoofed traffic placed in an AS chosen at random. We assume the volume of spoofed traffic originated in an AS is proportional to the number of sources in it. The first two scenarios are challenging scenarios, although previous work indicate that amplification attacks usually originate from a single source [15]. For each distribution, we generate and run simulations for 1000 placements.

For each distribution, Figure 12 shows the cumulative fraction of spoofed traffic (y -axis), averaged over the 1000 placements, in clusters up to a given size (x -axis). We observe that for all three distributions, most spoofed traffic originates from ASes in small clusters, which follows from Figure 3, where we showed that most clusters are small.

¹PEERING operators expressed concerns about hosting a honeypot on PEERING, subsequent blacklisting of PEERING resources, and deterioration of the platform’s future usability by the community. Some PEERING locations have bandwidth limitations which complicates hosting honeypots: although AmpPot [15] can enforce a limit on the sending rate, one cannot control the rate at which malicious packets are received from attackers.

VI. OPERATIONAL CONSIDERATIONS

In this section we discuss operational considerations for the deployment of our techniques by networks in the Internet.

A. Deployment Requirements

Our techniques generate anycast announcements. Any network with multiple peering links can deploy anycast and our techniques. We use the PEERING testbed to emulate a network with multiple peering links. Although small networks with few peering links may not be able to apply our techniques effectively, multiple small networks can cooperate to announce the same prefix and operate as a larger network that controls all of their peering links. Also, any network with a large peering footprint that deploys our techniques to identify the sources of spoofed traffic can share this information with other networks. To make an announcement, a network will also need network prefixes that preferably does not carry production traffic, as frequent announcement changes disrupts connectivity.

Our techniques are orthogonal to and compatible with the RPKI. The resource (prefix) owner needs only create a Route Origin Authorization (ROA) for the network(s) that will make the announcements; configuration generation and deployment work unmodified. BGP poisoning is also compatible with the RPKI and Route Origin Validation (ROV) as the origin of the announcement is unchanged: the poisoned AS's number is added to the middle of the AS-path.

Finally, an operator must setup infrastructure to automatically compute and deploy the sequence of announcement configurations, measure catchments (see §IV-B and §VI-C), and calculate clusters to identify the networks possibly sending the spoofed traffic. After identifying the network or set of networks responsible by the spoofed packets, the network operator can take several measures, e.g., contact network administration to notify the issue, or use either one or a combination of BGP flowspec [44] and blackholing [21] to limit propagation and impact of the spoofed traffic.

B. Requirements on Spoofed Traffic

Another requirement for deploying our techniques in practice is identifying spoofed traffic arriving on peering links. One approach is to classify traffic as either spoofed or legitimate, which is a challenging topic and active area of research [17], [45]. A practical approach is to run an amplification honeypot [15] to attract spoofed traffic and identify from which peering link the spoofed traffic is coming from. Using a honeypot has the benefit of guaranteeing accurate identification of spoofed traffic, as no benign traffic should reach the honeypots.

Our techniques can be applied even when the volume of spoofed traffic is small, as it only requires information about which peering link is receiving (most of the) spoofed traffic. However, identification of where spoofed traffic is coming from requires that sources send spoofed traffic toward the monitoring prefix. This implies identification of the sources of amplification attacks is limited to the period while the attack is ongoing. The number of configurations that can be deployed and the accuracy of the resulting clusters are proportional to

how long an attack lasts: Longer attacks allow for more precise identification of sources.

C. Measuring Catchments

In this paper we measure catchments using traceroutes from RIPE Atlas and BGP updates from public collectors (§IV-B). We chose this approach as PEERING prefixes receive very little traffic and restricts active probing using its resources. Next we discuss alternate mechanisms to measure catchments available to networks that can lift these constraints.

Networks can identify non-attack traffic received on each peering link by identifying whether packets received from a neighboring network are spoofed or not [17], [45]. This approach is limited by the accuracy of the identification, but has the benefit of working passively and not generating any measurement traffic. An operational constraint of this approach is how to attract enough traffic to identify catchments without disrupting production traffic during announcement configuration changes. In case the volume of traffic is too high, packet sampling can be used to control the processing overhead.

Another approach is to send ping measurements to IP addresses from /24 prefixes covering the globally-routed address space (e.g., using Verfploeter [16]) and calculate catchments from the responses. This approach has the benefit of not needing external vantage points (e.g., RIPE Atlas probes) to measure catchments.

D. Real-time Attack Detection

The time required to run a measurement round for an announcement configuration is decisive for real-time detection of amplification DDoS attacks. *We do not tackle real-time detection of amplification attacks in this work*, but discuss the duration of measurement rounds.

A measurement round must include time for route convergence, to measure the catchments, and to estimate where spoofed traffic is arriving from. Route convergence takes less than 2.5 minutes 99% of the time (§IV-B) and is outside the control of the network deploying our techniques.

The biggest contributor to the length of a measurement round is the time required for measuring catchments. In our study, we wait 50 minutes to collect at least two rounds of traceroutes, which is a function of the frequency at which traceroutes are triggered from RIPE Atlas. Other catchment measurement approaches, however, can operate faster. To speed up catchment measurement, an operator using pings to measure catchments can increase the probing rate to complete a scan of the IPv4 address space in less than five minutes [16], [46], [47]. Operators identifying non-spoofed traffic to measure catchments would need to adjust the time required to measure catchments as a function of the traffic volume received by their prefixes.

In most cases, we expect the time to estimate where spoofed traffic is arriving from to be shorter than the time to measure catchments, as amplification attacks send a continuous stream of spoofed packets while underway.

VII. OTHER APPLICATIONS

Although we designed our techniques to generate configurations with the goal of tracking down the sources of spoofed traffic, route manipulation has several other applications. In this section we discuss how our techniques can be adapted or extended to benefit other work. We also discuss how our dataset,² which we make publicly available, can be used as a starting point of analysis. Although PEERING and RIPE Atlas are publicly accessible, deploying hundreds of announcement configurations takes weeks, and our measurements have higher coverage than usually possible as we partnered with RIPE to collect a larger body of traceroute measurements.

Our techniques generate configurations that systematically explore routes and are applicable to previous work that manipulate BGP announcements to identify alternate paths [31], [48], trigger route changes with specific properties [49], discover network links [30], and characterize interdomain routing policies [24]. In general, our techniques and dataset can be of use to research in these areas: our dataset contains at least eight alternate routes towards PEERING for each observed AS, has thousands of route changes (with different properties), and may discover new links (particularly as a result of our poisoning experiments). More specifically, while Anwar et al. [24] generate announcement configurations to infer routing policies of a single target AS, our techniques deterministically force routing changes and explore routing decisions across all ASes in the Internet; such an approach could significantly speed up (and scale) inference of routing policies.

Research that involves prefix hijacks and defenses against it frequently deploys BGP announcements in the Internet to perform controlled hijacks and evaluate the effectiveness of their approaches in realistic scenarios (e.g., [50]). A scenario commonly studied in the literature is that of subprefix hijacks, where the hijacker announces a more specific route. This scenario, however, has a predictable outcome: the hijack is guaranteed to attract all traffic as Internet routing follows longest-prefix matching. A partial mitigation to subprefix hijacks is to announce more specific routes. In this context, the impact of a hijack depends on how competing announcements of /24 IPv4 and /48 IPv6 prefixes from a given set of locations propagate, which our announcements can be used to study. Our technique to generate configurations varying announcement locations generates *all* possible scenarios of prefix hijacking from a predefined set of announcement locations. Consider a configuration announcing from n locations: each location can be considered a legitimate announcement or an attempted hijack. Under this view, a configuration announcing from n locations covers 2^n possible hijack scenarios.

VIII. RELATED WORK

a) DDoS attacks: Denial-of-Service (DoS) attacks intend to exhaust the resources of a service, making it unavailable for legitimate users. The attack can exhaust server resources (e.g., CPU, memory, network connections) [51]–[53] or send a

torrent of requests to exceed the bandwidth capacity of the network link where the service is hosted (volumetric attacks) [1], [15], [53]. A Distributed Denial-of-Service (DDoS) attack uses a large number of devices across the network to launch the attack and overwhelm a target.

Attacks can use reflection or amplification to empower their impact. A service vulnerable to reflection replies to requests without verifying the source (e.g., classic DNS over UDP), and a service vulnerable to amplification sends large responses to short requests. A server running a service vulnerable to both reflection and amplification is usually referred to as an amplifier. Attackers exploit amplifiers by generating small requests falsifying the source IP address of the requests with the IP address of the victim, causing the amplifier to send large responses to the victim instead. Amplification reflection DDoS attacks [1], [7] are particularly challenging because they generate large volumes of traffic and are hard to track down since the source IP address of the malicious traffic is spoofed. This type of attack is perpetrated in practice, and have increased in power over time, leading to recent record-breaking attacks that reached 2,3 Tbps [7]. Our work complements efforts to mitigate amplification reflection attacks, tracking down the sources of spoofed packets.

b) DDoS mitigation: Distinguishing legitimate from malicious traffic is an essential task for many DoS mitigation techniques. Solutions to detect attacks propose to differentiate between classes of anomalous traffic (e.g., flash crowd, equipment failures, and DoS attacks) by (i) applying black-box models obtained from machine learning algorithms on traffic [54]–[56] or (ii) identifying signatures of normal and anomalous traffic, and creating white-box traffic filtering rules [57]. In practice, traffic scrubbing services [58], [59] detour the victim’s traffic through a scrubbing center, classify traffic into malicious and legitimate, discard the malicious traffic, and send only the legitimate traffic to the victim network. Although popular, this approach is very expensive since it demands high network bandwidth capacity and also advanced hardware to process and filter large volumes of packets associated with DDoS attack [60], leading many works to address the costs associated with scrubbing centers [61]–[63].

These techniques reduce the impact on the victims by detecting and filtering the malicious packages or absorbing the attacks. However, they do not address the root cause of amplification attacks, which requires identifying the source of the attacks (i.e., the sources of the spoofed packets) for targeted intervention and mitigation at the origin.

c) IP Traceback: A common approach to perform IP traceback is to *iteratively* contact operators, either directly or through mailing lists such as Outages.org or NANOG, of ASes along the path toward the source of spoofed traffic and have them identify which links are carrying the malicious traffic. This approach however, requires coordination and cooperation between several humans, incurring arbitrary identification delay and not scaling to the wide area.

Controlled flooding [8] was the first automated approach for IP traceback, and relied on temporarily congesting links

²https://homepages.dcc.ufmg.br/~osvaldo.morais/dataset_tnsm2020/

Table III: Summary of proposals for IP traceback.

Approach	Manipulates	Cooperation from networks	Router updates	Router overhead	Identification precision	Identification delay
Manual	Logs/monitoring	Required	No	No	Path prefix	Long
Flooding [8]	Packet loss	Required	No	High	Path prefix	Moderate
Marking [9]–[11]	IP ID field	Deployment	Yes	Low	Closest router	$\approx$ sampling
Out-of-band [12]	—	Deployment	Yes	High	Closest router	$\approx$ sampling
Digest-Based [13], [14]	Local state at router	Deployment	Yes	High	Closest router	Short
Routing (this paper)	Routes	No	No	No	AS	Long

to disrupt traffic on a link, allowing the victim to iteratively identify links on the path toward the attacker. Although this approach does not require upgrading routers, it is not viable today as the ability to trigger congestion at will (e.g., using UDP chargen) is considered a serious vulnerability.

Several packet *marking* approaches propose encoding information about routers traversed by a packet on a small fraction of packets (usually in the IP ID field) [9]–[11]. Similarly, routers can inform destinations using *out-of-band* data about a small fraction of packets they have forwarded toward each destination [12]. Under the assumption that attackers generate many packets toward the amplifier, the amplifier can correlate information across multiple packets and identify routers on the path to the attacker. Another approach is to compute a *digest* (e.g., a bloom filter) of packets traversing a router, and provide an interface for querying routers for a packet’s signature [13], [14]. These techniques allow for fast identification, but require upgrading routers, incur significant overhead, and require widespread deployment across the Internet to provide accurate identification.

Table III provides an overview of IP traceback proposals and compares with the proposal in this paper. Our routing-based approach manipulates routes to find the AS that originates the malicious traffic. Our proposal does not require any cooperation from remote networks nor router updates, addressing most of the drawbacks observed in previous work.

d) Locating sources of spoofed traffic: Previous approaches to locate the sources of spoofed packets either lack coverage or are not deployable. One way to prevent spoofed IP packets is filtering them at the source network. The problem with this approach is the lack of incentive that the network has to carry out the filtering of spoofed packets since the network itself does not benefit from the filtering. Although it prevents the network from participating in amplification attacks, the network is still vulnerable to these attacks. Some reports analyze the positive impact that the adoption of spoofed packet filtering can have and how to deploy the filtering [64] while others try to quantify the percentage of networks vulnerable to the use of IP spoofing [65]–[67]. A study that relied on active tests included data from volunteers at 12,500 IP addresses [66], a small fraction of the billion client IP addresses seen by large Internet services that are the victims of attacks.

A recent work proposes an approach that identifies scanning infrastructures used for amplification attacks and correlates scan events with amplification attacks to identify amplification sources [68]. The scanning infrastructure may not be used for running the attacks, which would be a limitation of the

technique. The results of this work are in line with previous works whose findings show that most amplification attacks come from a single source.

Finally, there are also studies that proposed techniques to detect spoofed IP packets in IXPs [17], [45], [69]. These techniques check if the packets entering the IXP come from a valid customer of the AS that injects the packet in the IXP fabric, i.e., checks whether a packet’s source IP address belongs to a prefix owned by an AS in the IXP member’s customer cone.

IX. CONCLUSION AND FUTURE WORK

Our control-plane traceback technique can be deployed by any network with rich connectivity today, without changes to routers, and does not require cooperation from other networks. Our results using the PEERING platform indicate that our proposed techniques to generate announcement configurations can effectively manipulate routes and induce catchment changes, allowing tracking down the sources of spoofed traffic. If sources of amplification DDoS attacks are few, as reported by analyzing logs from AmpPot honeypots [15], our techniques can map sources of spoofed traffic into sets that average 1.29 ASes. Our results indicate that precision will be higher if networks with a footprint larger than PEERING’s were to deploy our techniques.

We envision two research fronts for future work. One is to expand our techniques to reduce cluster sizes even more, e.g., designing new algorithms for choosing targets for poisoning, and using BGP communities for controlling export policies (and influence routing decisions) on remote networks. Another is to expand the system to allow identification of sources of spoofed traffic during DDoS attacks, e.g., by (i) jointly optimizing for cluster size and traffic volume, giving higher utility to reducing the size of clusters inferred to send more spoofed traffic; and (ii) improving existing catchment prediction techniques [18] to allow generation of announcement configurations without prior knowledge and reducing the need for measuring catchments in advance.

ACKNOWLEDGEMENTS

We thank the anonymous reviewers for their feedback and RIPE for the increased measurement budget used to perform traceroutes from RIPE Atlas toward PEERING prefixes. This work was supported by CAPES, CNPq, FAPEMIG, MCTIC/RNP grants 2955 and 2956; projects MASWeb, EUBra-BIGSEA, INCT-Cyber, ATMOSPHERE, TLHOP (CERT.BR/NIC.BR); and NSF grants CNS-1740883 and CNS-1835252.

REFERENCES

- [1] J. Czyz, M. Kallitsis, M. Gharaibeh, C. Papadopoulos, M. Bailey, and M. Karir, "Taming the 800 Pound Gorilla: The Rise and Decline of NTP DDoS Attacks," in *Proc. ACM IMC*, 2014.
- [2] M. Prince, "Technical Details Behind a 400Gbps NTP Amplification DDoS Attack," Feb 2014. [Online]. Available: <https://blog.cloudflare.com/technical-details-behind-a-400gbps-ntp-amplification-ddos-attack>
- [3] K. York, "Dyn Statement on 10/21/2016 DDoS Attack," 2016, <http://dyn.com/blog/dyn-statement-on-10212016-ddos-attack/>.
- [4] L. H. Newman, "Github Survived the Biggest DDoS Attack ever Recorded," *Wired*, March 2018.
- [5] V. Paxson, "An Analysis of Using Reflectors for Distributed Denial-of-service Attacks," *SIGCOMM Comput. Commun. Rev.*, vol. 31, no. 3, pp. 38–47, 2001.
- [6] "Akamai Q4 2016 State of the Internet – Security Report," <https://www.akamai.com/uk/en/our-thinking/state-of-the-internet-report/>.
- [7] E. Targett, "AWS Hit With a Record 2.3 Tbps DDoS Attack," Jun 2020. [Online]. Available: <https://www.cbronline.com/news/record-ddos-attack-aws>
- [8] H. Burch and B. Cheswick, "Tracing Anonymous Packets to Their Approximate Source," in *Proc. USENIX Conf. on System Admin.*, 2000.
- [9] S. Savage, D. Wetherall, A. Karlin, and T. Anderson, "Network Support for IP Traceback," *IEEE/ACM Trans. Netw.*, vol. 9, no. 3, pp. 226–237, 2001.
- [10] A. Belenky and N. Ansari, "On Deterministic Packet Marking," *Comput. Netw.*, vol. 51, no. 10, pp. 2677–2700, 2007.
- [11] Z. Gao and N. Ansari, "A Practical and Robust Inter-domain Marking Scheme for IP Traceback," *Comput. Netw.*, vol. 51, no. 3, pp. 732–750, 2007.
- [12] S. Bellovin, M. Leech, and T. Taylor, "ICMP Traceback Messages," 2003, <https://tools.ietf.org/html/draft-ietf-itrace-04>.
- [13] M. Sung, J. Xu, J. Li, and L. Li, "Large-scale IP Traceback in High-speed Internet: Practical Techniques and Information-theoretic Foundation," *IEEE/ACM Trans. Netw.*, vol. 16, no. 6, pp. 1253–1266, 2008.
- [14] A. C. Snoeren, C. Partridge, L. A. Sanchez, C. E. Jones, F. Tchakountio, B. Schwartz, S. T. Kent, and W. T. Strayer, "Single-packet IP Traceback," *IEEE/ACM Trans. Netw.*, vol. 10, no. 6, pp. 721–734, 2002.
- [15] L. Krämer, J. Krupp, D. Makita, T. Nishizoe, T. Koide, K. Yoshioka, and C. Rossow, "AmpPot: Monitoring and Defending Against Amplification DDoS Attacks," in *Proc. Intl. Symp. on Research in Attacks, Intrusions and Defenses (RAID)*, 2015.
- [16] W. B. de Vries, R. de O. Schmidt, W. Hardaker, J. Heidemann, P.-T. de Boer, and A. Pras, "Verfloeter: Broad and Load-Aware Anycast Mapping," in *Proc. ACM IMC*, 2017.
- [17] F. Lichtblau, F. Streibelt, T. Krüger, P. Richter, and A. Feldmann, "Detection, Classification, and Analysis of Inter-domain Traffic with Spoofed Source IP Addresses," in *Proc. ACM IMC*, 2017.
- [18] P. Sermpezis and V. Kotronis, "Inferring Catchment in Internet Routing," in *SIGMETRICS Poster Session*, 2019.
- [19] B. Schlinder, T. Arnold, I. Cunha, and E. Katz-Bassett, "PEERING: Virtualizing BGP at the Edge for Research," in *Proc. ACM CoNEXT*, 2019.
- [20] P. Ferguson and D. Senie, "Network Ingress Filtering: Defeating Denial of Service Attacks which employ IP Source Address Spoofing," 2000, <http://www.ietf.org/rfc/rfc2827.txt>.
- [21] V. Gotsas, G. Smaragdakis, C. Dietzel, P. Richter, A. Feldmann, and A. Berger, "Inferring BGP Blackholing Activity in the Internet," in *Proc. ACM IMC*, 2017.
- [22] P. Marques, N. Sheth, R. Raszk, B. Greene, J. Mauch, and D. McPherson, "RFC 5575: Dissemination of Flow Specification Rules," Aug. 2009, <http://www.ietf.org/rfc/rfc5575.txt>.
- [23] O. Fonseca, I. Cunha, E. Fazzion, W. Meira, B. Junior, R. A. Ferreira, and E. Katz-Bassett, "Tracking Down Sources of Spoofed IP Packets," in *IFIP Networking*, 2020.
- [24] R. Anwar, H. Niaz, D. R. Choffnes, I. Cunha, P. Gill, and E. Katz-Bassett, "Investigating Interdomain Routing Policies in the Wild," in *Proc. ACM IMC*, 2015.
- [25] M. Calder, A. Flavel, E. Katz-Bassett, R. Mahajan, and J. Padhye, "Analyzing the Performance of an Anycast CDN," in *Proc. ACM IMC*, 2015.
- [26] P. Sun, L. Vanbever, and J. Rexford, "Scalable Programmable Inbound Traffic Engineering," in *Proc. ACM SOSR*, 2015.
- [27] D. Cicalese and D. Rossi, "A Longitudinal Study of IP Anycast," *SIGCOMM Comput. Commun. Rev.*, vol. 48, no. 1, pp. 10–18, 2018.
- [28] B. Schlinder, H. Kim, T. Cui, E. Katz-Bassett, H. V. Madhyastha, I. Cunha, J. Quinn, S. Hasan, P. Lapukhov, and H. Zeng, "Engineering Egress with Edge Fabric: Steering Oceans of Content to the World," in *Proc. ACM SIGCOMM*, 2017.
- [29] R. Bush, O. Maennel, M. Roughan, and S. Uhlig, "Internet Optometry: Assessing the Broken Glasses in Internet Reachability," in *Proc. ACM IMC*, 2009.
- [30] L. Colitti, "Internet Topology Discovery Using Active Probing," Ph.D. dissertation, University di Roma Tre, 2006.
- [31] E. Katz-Bassett, C. Scott, D. R. Choffnes, I. Cunha, V. Valancius, N. Feamster, H. V. Madhyastha, T. Anderson, and A. Krishnamurthy, "LIFEGUARD: Practical Repair of Persistent Route Failures," in *Proc. ACM SIGCOMM*, 2012.
- [32] Y.-C. Chiu, B. Schlinder, A. B. Radhakrishnan, E. Katz-Bassett, and R. Govindan, "Are We One Hop Away from a Better Internet?" in *Proc. ACM IMC*, 2015.
- [33] J. M. Smith, K. Birkeland, T. McDaniel, and M. Schuchard, "Withdrawing the BGP Re-Routing Curtain: Understanding the Security Impact of BGP Poisoning via Real-World Measurements," in *NDSS*, 2020.
- [34] RIPE, "RIPE Atlas," 2013, <https://atlas.ripe.net>.
- [35] "The University of Oregon Routeviews Project," <http://www.routeviews.org>.
- [36] "RIPE Routing Information Service," <http://www.ripe.net/data-tools/stats/ris>.
- [37] Team Cymru, "IP to ASN Mapping," <http://www.team-cymru.org/Services/ip-to-asn.html>.
- [38] G. Nomikos and X. Dimitropoulos, "traXroute: Detecting IXPs in Traceroute Paths," in *Proc. of PAM*, 2016.
- [39] V. Gotsas, M. Luckie, B. Huffaker, and K. Claffy, "Inferring Complex AS Relationships," in *Proc. ACM IMC*, 2014.
- [40] M. Luckie, B. Huffaker, K. Claffy, A. Dhamdhere, and V. Gotsas, "AS Relationships, Customer Cones, and Validation," in *Proc. ACM IMC*, 2013.
- [41] "BGP Churn," 2018, <https://blog.apnic.net/2019/01/22/bgp-in-2018-bgp-churn/>.
- [42] L. Gao, "On Inferring Autonomous System Relationships in the Internet," *IEEE/ACM Trans. Netw.*, vol. 9, no. 6, pp. 733–745, 2001.
- [43] I. Cunha, P. Marchetta, M. Calder, Y.-C. Chiu, B. Schlinder, B. V. A. Machado, A. Pescape, V. Gotsas, H. V. Madhyastha, and E. Katz-Bassett, "Sibyl: A Practical Internet Route Oracle," in *Proc. USENIX NSDI*, 2016.
- [44] "RFC5575: Dissemination of Flow Specification Rules," <https://tools.ietf.org/html/rfc5575>.
- [45] J. Eumann, R. Hiesgen, T. C. Schmidt, and M. Wählisch, "A Reproducibility Study of "IP Spoofing Detection in Inter-Domain Traffic"," 2019.
- [46] G. F. Lyon, *Nmap network scanning: The official Nmap project guide to network discovery and security scanning*. Insecure. Com LLC (US), 2008.
- [47] Z. Durumeric, E. Wustrow, and J. A. Halderman, "ZMap: Fast Internet-wide Scanning and Its Security Applications," in *22nd USENIX Security Symposium (USENIX Security 13)*, 2013.
- [48] V. Valancius, B. Ravi, N. Feamster, and A. C. Snoeren, "Quantifying the Benefits of Joint Content and Network Routing," in *Proc. ACM SIGMETRICS*, 2013.
- [49] U. Javed, I. Cunha, D. R. Choffnes, E. Katz-Bassett, T. Anderson, and A. Krishnamurthy, "PoiRoot: Investigating the Root Cause of Interdomain Path Changes," in *Proc. ACM SIGCOMM*, 2013.
- [50] P. Sermpezis, V. Kotronis, P. Gigis, X. Dimitropoulos, D. Cicalese, A. King, and A. Dainotti, "ARTEMIS: Neutralizing BGP Hijacking Within a Minute," *IEEE/ACM Trans. Netw.*, vol. 26, no. 6, p. 2471–2486, 2018.
- [51] H. Beitollahi and G. Deconinck, "Tackling Application-layer DDoS Attacks," *Procedia Computer Science*, vol. 10, pp. 432 – 441, 2012.
- [52] K. Hong, Y. Kim, H. Choi, and J. Park, "SDN-Assisted Slow HTTP DDoS Attack Defense Method," *IEEE Communications Letters*, vol. 22, no. 4, pp. 688–691, 2018.
- [53] S. T. Zargar, J. Joshi, and D. Tipper, "A Survey of Defense Mechanisms Against Distributed Denial of Service (DDoS) Flooding Attacks," *IEEE Communications Surveys Tutorials*, vol. 15, no. 4, pp. 2046–2069, 2013.
- [54] A. Soule, K. Salamatian, and N. Taft, "Combining Filtering and Statistical Methods for Anomaly Detection," in *Proc. ACM IMC*, 2005.

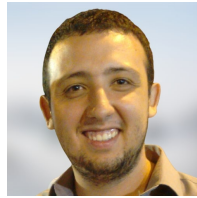
- [55] F. Silveira, C. Diot, N. Taft, and R. Govindan, "ASTUTE: Detecting a Different Class of Traffic Anomalies," in *Proc. ACM SIGCOMM*, 2010.
- [56] F. Silveira and C. Diot, "URCA: Pulling out Anomalies by their Root Causes," in *Proc. IEEE INFOCOM*, 2010.
- [57] T. Thapngam, S. Yu, W. Zhou, and G. Beliakov, "Discriminating DDoS Attack Traffic from Flash Crowd through Packet Arrival Patterns," in *IEEE Conference on Computer Communications Workshops (INFOCOM WKSHPS)*, 2011.
- [58] "CloudFlare – Advanced DDoS Protection and Mitigation," <https://www.cloudflare.com/ddos/>.
- [59] Akamai, "Why Akamai Cloud Security for DDoS Protection?" 2020, <https://www.akamai.com/us/en/products/security/ddos-protection-service.jsp>.
- [60] John Graham-Cumming, "No Scrubs: The Architecture That Made Unmetered Mitigation Possible," 2017, <https://blog.cloudflare.com/no-scrubs-architecture-unmetered-mitigation/>.
- [61] W. You, L. Jiao, J. Li, and R. Zhou, "Scheduling DDoS Cloud Scrubbing in ISP Networks via Randomized Online Auctions," in *Proc. IEEE INFOCOM*, 2020.
- [62] P. Zilberman, R. Puzis, and Y. Elovici, "On Network Footprint of Traffic Inspection and Filtering at Global Scrubbing Centers," *IEEE Transactions on Dependable and Secure Computing*, vol. 14, no. 5, pp. 521–534, 2017.
- [63] Z. Liu, Y. Cao, M. Zhu, and W. Ge, "Umbrella: Enabling ISPs to Offer Readily Deployable and Privacy-Preserving DDoS Prevention Services," *IEEE Transactions on Information Forensics and Security*, vol. 14, no. 4, pp. 1098–1108, 2019.
- [64] "RFC2827: Network Ingress Filtering," <http://www.bcp38.info/>.
- [65] R. Beverly and S. Bauer, "The spoofer project: Inferring the extent of source address filtering on the internet," in *Usenix Sruti*, 2005.
- [66] R. Beverly, A. Berger, Y. Hyun, and k. claffy, "Understanding the Efficacy of Deployed Internet Source Address Validation Filtering," in *Proc. ACM IMC*, 2009.
- [67] M. Kührer, T. Hupperich, C. Rossow, and T. Holz, "Exit from Hell? Reducing the Impact of Amplification DDoS Attacks," in *Proceedings of the 23rd USENIX Conference on Security Symposium*, 2014.
- [68] J. Krupp, M. Backes, and C. Rossow, "Identifying the Scan and Attack Infrastructures Behind Amplification DDoS Attacks," in *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*, 2016.
- [69] L. Müller, M. Luckie, B. Huffaker, K. Claffy, and M. Barcellos, "Challenges in Inferring Spoofed Traffic at IXPs," in *Proc. ACM CoNEXT*, 2019.



Osvaldo Fonseca is a Ph.D candidate of the Computer Science Department at the Universidade Federal de Minas Gerais (DCC/UFGM), Brazil. He received his Master degree at Universidade Federal de Minas Gerais in 2016 and his Bachelor degree at the same university in 2013. His research interests are in computer networks, with an emphasis on security and Internet measurement.



Ítalo Cunha is an associate professor at the Computer Science Department at UFGM, Brazil. He developed his Ph.D. research at Technicolor Research and Innovation, and graduated from Université Pierre et Marie Curie in 2011. Ítalo worked as an Associate Visiting Research Scholar at Columbia University during 2019. His research focuses on improving network performance, reliability and security. His contributions provide better visibility on Internet topology and routing dynamics; help network operators troubleshoot failures and performance problems; support identification and mitigation of cybersecurity threats; and empower novel networking and security research. He holds a research fellowship from Brazil's National Science and Technology Foundation (CNPq) and serves as a member of the National Brazilian Research and Education Network (RNP) Monitoring Working Group.



Elverton Fazzion is a lecturer at Computer Science Department in Universidade Federal de São João del Rei (DCOMP/UFSJ) and he is also a PhD student at Universidade Federal de Minas Gerais (DCC/UFGM). He received his Master degree at UFGM in 2016 and his Bachelor degree (with honors) at the same university in 2014. His research interests include computer networks, with an emphasis on Internet measurements, troubleshooting and security.



Wagner Meira Jr. obtained his PhD from the University of Rochester in 1997 and is Full Professor at the Computer Science Department at Universidade Federal de Minas Gerais, Brazil. He has published more than 300 papers in top venues and is co-author of the books *Data Mining and Analysis - Fundamental Concepts and Algorithms* and *Data Mining and Machine Learning - Fundamental Concepts and Algorithms*, both published by Cambridge University Press in 2014 and 2020, respectively. His research focuses on scalability and efficiency of large scale

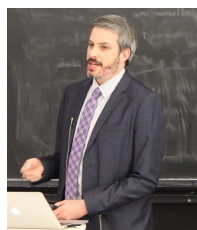
parallel and distributed systems, from massively parallel to Internet-based platforms, and on data mining algorithms, their parallelization, and application to areas such as data science, information retrieval, cybersecurity, and e-governance.



Brivaldo Alves da Silva Junior is a Ph.D. student at the Federal University of Mato Grosso do Sul (UFMS). He received his M.Sc. and B.Sc. degrees in Computer Science from UFMS, Brazil, in 2013 and 2007, respectively. He works as a Lecturer in the College of Computing at UFMS. His research interests include network, resource and service management in wired networks, future Internet architecture, and security in systems and protocols.



Ronaldo A. Ferreira is an Associate Professor of Computer Science in the College of Computing at UFMS. He received his B.Sc. from UFMS in 1992, his M.S. from the University of Campinas in 1998, and his Ph.D. from Purdue University in 2006, all in Computer Science. Ronaldo was a Visiting Research Scholar and a Visiting Associate Professor at Princeton University from 2014 to 2016. He served as the chair of the Special Interest Group on Computer Networks and Distributed Systems of the Brazilian Computing Society (SBC) from 2011 to 2013. Ronaldo was a member of the board of administration of the Brazilian National Research and Educational Network (RNP) from 2011 to 2013. He has been a member of the board of directors of the Brazilian National Laboratory of Computer Networks (LARC) since 2014. His research interests are in computer networks and distributed systems.



Ethan Katz-Bassett is an Associate Professor in the Electrical Engineering Department and an Affiliate in the Computer Science Department at Columbia University, where he is also a Member of the Computer Engineering Program and the Sense, Collect and Move Data Center at the Data Science Institute. He was previously the Andrew and Erna Viterbi Early Career Chair in the Computer Science Department at the University of Southern California. He completed his Ph.D. in the Department of Computer Science at the University of Washington and worked for half a year at Google, as part of a great team tasked with making the mobile web fast.