

# The Circuit Breaker That Cried Wolf: Measuring BGP Maximum-Prefix Exceedance and Connectivity Impact

Orlando E. Martínez-Durive\*

IMDEA Networks Institute

Madrid, Spain

orlando.martinez@networks.imdea.org

Antonis Chariton

Cisco ThousandEyes

Zurich, Switzerland

acharito@cisco.com

Marco Fiore

IMDEA Networks Institute

Madrid, Spain

marco.fiore@networks.imdea.org

## Abstract

The BGP maximum-prefix limit serves as a network circuit breaker, tearing down sessions when announcements exceed a configured threshold. Despite its operational significance, little is known empirically about how accurately these limits reflect reality, how often they are exceeded, or what connectivity impact results from their exceedance. We present a first large-scale measurement study of BGP maximum-prefix limits in the wild, combining 12 months of RIPE RIS snapshots with PeeringDB metadata and CAIDA AS Rank data. Only 32.49% of RIS-visible ASes appear in PeeringDB, and 21–27% of recent limit fields carry a value of zero, indicating no declared limit; after filtering for recency and non-zero limits, 14,973 ASes form our analytical dataset. Among these, 14.55% (IPv4) and 11.24% (IPv6) exceed their declared limit intermittently, and 2.15% and 2.32% do so in every observed snapshot. These exceedances are driven by organic growth, such as new space acquisition or deaggregation, yet the mechanism responds identically to a route leak: Tier-1 and Major peers account for 18.4% of IPv4 lost sessions despite representing roughly 0.1% of all ASes. Based on our findings, we offer recommendations for operators and IXPs and call for renewed IETF standardization of dynamic limit negotiation.

## CCS Concepts

• **Networks** → **Network measurement; Routing protocols; Network reliability; Public Internet; Network manageability**; • **Security and privacy** → **Network security**.

## Keywords

BGP, network measurement, maximum-prefix limit, PeeringDB, route leaks, routing security, RIPE RIS

## ACM Reference Format:

Orlando E. Martínez-Durive, Antonis Chariton, and Marco Fiore. 2026. The Circuit Breaker That Cried Wolf: Measuring BGP Maximum-Prefix Exceedance and Connectivity Impact. In *Proceedings of the 2026 ACM Internet Measurement Conference (IMC '26)*, October 12–16, 2026, Karlsruhe, Germany. ACM, New York, NY, USA, 16 pages. <https://doi.org/10.1145/3777912.3839785>

\*Part of this work was done while the author was a Technical Intern at Cisco ThousandEyes.



This work is licensed under a Creative Commons Attribution 4.0 International License. *IMC '26, Karlsruhe, Germany*

© 2026 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-2327-8/26/10

<https://doi.org/10.1145/3777912.3839785>

## 1 Introduction

The Border Gateway Protocol (BGP) [60] is the inter-domain routing protocol of the Internet. Networks (autonomous systems, or ASes) exchange routing information through BGP sessions established between neighboring routers; each session reflects a commercial relationship, either customer-to-provider or peer-to-peer, that determines which routes are accepted and propagated [27]. Together they build a global routing table that today spans over 120,000 ASes [14], advertising over 1.1 million IPv4 and 260,000 IPv6 prefixes [29], a count that has grown steadily as IPv4 address exhaustion drives deaggregation [26] and IPv6 adoption accelerates. Although BGP is the backbone of modern Internet connectivity, it was designed in an era of small scale and mutual trust; it has no built-in mechanisms to validate the correctness or intent of the routing information it carries.

**BGP failures.** Its architectural openness makes BGP a source of both accidental and deliberate disruptions. Prior work has measured BGP dynamics broadly [7, 72], characterized operational BGP mechanisms at Internet Exchange Points (IXPs) such as blackholing [46], peering infrastructure outages [28], and route-server route diversity [21], studied Resource Public Key Infrastructure (RPKI) deployment and its operational challenges [16], and characterized PeeringDB as a coordination platform [39].

State actors have exploited BGP for censorship and traffic interception, routing domestic traffic through national chokepoints to enable surveillance and deep-packet inspection, or leaking intended blackhole routes to the global Internet [44, 61, 64]. Attackers have weaponized BGP for financial fraud by advertising more-specific prefixes that appear more attractive to routers: in April 2018, hijacked Amazon Route 53 DNS prefixes redirected MyEtherWallet users to fraudulent servers, enabling cryptocurrency theft [54].

Misconfigurations can propagate globally within minutes; software bugs can leave stale routing state (zombie routes) persisting in routing tables long after withdrawal, with outbreaks occurring more than once per day on monitored prefixes and documented cases of degradation lasting days or weeks [24, 73].

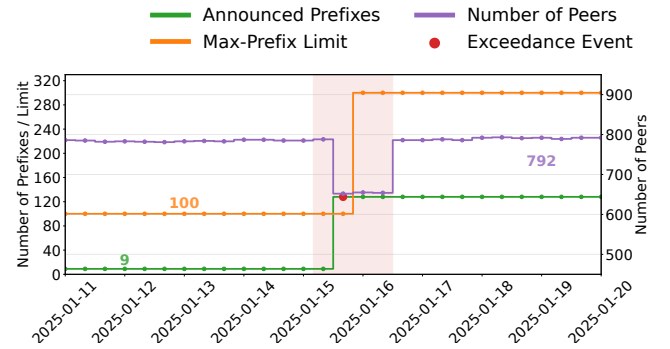
Across this spectrum of failures, the community's defensive response concentrates on routing information correctness. RPKI route-origin authorizations now cover more than half of routed prefixes [30], though the protection actually delivered is limited where route-origin validation is only partially deployed [17]. Autonomous System Provider Authorization (ASPA) [62, 74] extends validation to the AS-path, remaining effective against route leaks even under partial adoption [8, 23, 25]. Both validate the *correctness* of routing information, authorizing which AS may originate a prefix and which AS-paths are legitimate, yet neither places any bound on the *how many* of routes a neighbor may send.

**A circuit breaker against route leaks.** The operational safeguards that govern session behavior have received far less empirical scrutiny. One such mechanism stands out: the *maximum-prefix limit* [19, 60], which bounds *how many* prefixes a neighbor may send. This operator-set threshold causes a BGP session to take a protective action when a neighbor announces more prefixes than expected. The action goes from logging a warning to tearing the session down, depending on the configuration, serving as a circuit breaker against memory exhaustion and route leaks. Events where an AS re-advertises routes received from one neighbor to another in violation of its routing policy, typically a customer propagating a transit provider's routes to a second provider. The routes leaked tend to gain preference on the BGP best-path selection by through two mechanisms: by advertising *more-specific* prefixes, preferred under the longest-prefix match rule; and/or *shorter AS-paths* through the leaking AS, causing neighbors to prefer this route.

During route leaks, neighboring routers absorb tens of thousands of spurious routes, exhausting their memory; the leaking AS, suddenly carrying traffic far beyond its capacity, may itself experience congestion or outage. In some cases, resolution requires manual coordination: detecting the anomaly, identifying the responsible AS, and negotiating a withdrawal; in the 2019 Verizon incident, a BGP optimizer at a small Internet Service Provider (ISP) propagated over 20,000 illegitimate routes globally, taking Cloudflare, Amazon, and Linode offline, with Verizon still unresponsive more than eight hours later due to the absence of inbound maximum-prefix limits [70]. In 2023, Optus, Australia's second-largest telecommunications provider, experienced a 12-hour nationwide outage after core routers hit their maximum-prefix limits and self-isolated; a critical infrastructure failure leaving 10 million customers without service and blocking 2,145 emergency Triple Zero calls, resulting in a \$12 million government fine [5, 20]. As recently as January 2026, Venezuela's state-run ISP CANTV (AS8048) leaked routes received from transit provider Sparkle to peer GlobeNet across eleven separate events spanning weeks, each traced to overly permissive export policies, with affected prefixes hosting banking, telecommunications, and government infrastructure [33]. The Kirin attack [58] showed that maximum-prefix configurations are not merely fragile but exploitable: an adversary can craft announcements that make a victim's peers believe it has exceeded its limit, triggering tear-downs that isolate the target from transit and peering with no misconfiguration on its part.

**Consequences beyond the session.** The examples above show how the consequences of a *maximum-prefix limit* breach extend beyond the two networks involved in a session teardown. When a session drops, the effects on downstream customers of the affected AS range from suboptimal routing paths and increased latency or cost to, in the worst case, complete loss of reachability. Recovery is not always automatic. Most implementations require the session to be manually cleared or a timer to expire, leaving networks isolated for minutes to hours. For ISPs and content providers, these episodes mean Service Level Agreement (SLA) violations, escalations, brand damage, and lost revenue. Unlike most failures, these outages are deterministic, arising from a configuration mismatch that persists undetected until traffic exposes it.

Beyond the immediate connectivity loss, the problem is compounded by the lack of in-protocol coordination. BGP cannot signal



**Figure 1: AS44901 (BelCloud Ltd., IPv6): announced prefix count (green) crosses its suggested max-prefix limit (orange), after deaggregating its /32 block into more-specific /38 sub-prefixes, causing an immediate peer drop from ~800 to ~650.**

that a configured limit is outdated or that a neighbor's announced prefix count has exceeded it. Resolution requires out-of-band coordination, typically Network Operations Center (NOC) calls or email, with no protocol-level guidance for either party. Which side acts first depends on the peering relationship. The peer holding the limit may move quickly if the originating network is a critical provider; the originating network may reach out to request a limit increase or withdraw the new prefixes if the peer is essential for its reachability. In some cases, the session can be restored without either party addressing the root cause, a stale limit that no longer reflects the peer's actual prefix count, leaving the same mismatch.

**Stale limits and organic exceedances.** Our measurements reveal a different and more insidious failure mode: BGP sessions dropping because a network's prefix count has grown by a modest margin above a stale limit. Figure 1 illustrates this dynamic in a concrete representative example: AS44901's announced prefix count crosses its recommended limit, triggering an immediate collapse in peer count. Although the operator responds by updating their limit, connectivity takes hours to days to partially recover as peers reconfigure the new limit on their own schedules.

A systematic classification of a large number of similar critical events that we identify confirms that the trigger is routine operational growth: new prefix space acquisition for IPv4 (60.0%) and deaggregation of existing address space for IPv6 (60.3%), with median additions of just 5 (IPv4) and 2 (IPv6) prefixes at crossing. *The mechanism has no way to distinguish this routine operational practice from a route leak, and no channel exists to negotiate an updated limit; the circuit breaker fires on normal traffic engineering, and well-behaved networks bear the cost of an outdated configuration.*

To our knowledge, no public post-mortem of a max-prefix exceedance event caused by organic prefix growth has been documented: the affected operator faces no external pressure to publish one, and the session is quietly restarted once the peer reconfigures their limits or the originating operator rolls back its announcements. Nor has any study systematically measured how well this mechanism actually functions: whether declared limits are accurate, how often they are exceeded, and what connectivity impact follows.

**Contributions.** We present a first large-scale empirical study of BGP maximum-prefix limits. We combine 12 months of RIPE Routing Information Service (RIS) routing snapshots (January 1 – December 31, 2025) from 23 route collectors (25,185 RIB dumps totaling 3.4 TB of raw data) with PeeringDB metadata, covering 90,451 ASes and yielding a working dataset of 14,973 ASes with RIS visibility and a recent, non-zero PeeringDB limit. Our main findings and contributions are as follows.

- **PeeringDB coverage and data quality.** About a third of RIS-visible ASes appear in PeeringDB, and just a fifth have an entry updated since 2023; among those recent entries, roughly a quarter carry a zero (unset) limit.
- **Exceedance rates.** About 2% of ASes in both the IPv4 and IPv6 families exceed their declared limit in *every* snapshot across the year, a chronic mismatch rather than isolated incidents, while a larger 11–15% exceed only intermittently. The median prefix increase during an exceedance is just a handful of prefixes (5 IPv4, 2 IPv6), suggesting gradual growth against stale limits rather than sudden routing anomalies.
- **Peer connectivity impact.** We identify 303 IPv4 and 132 IPv6 *impactful* events, where the exceedance coincides with a peer-count drop beyond the AS's 95th-percentile churn. The typical loss is substantial, a median of 25% of peers for IPv4 and 33% for IPv6, with a tail of ASes losing over 85%. Tier-1 and Major providers, though only roughly 0.1% of all ASes, account for a disproportionate 18.4% of lost IPv4 sessions, showing how these exceedances can sever high-value transit.
- **Root-cause analysis.** We classify every impactful event by its structural signature, finding that IPv4 exceedances are driven mainly by new address-space origination and IPv6 by deaggregation of existing space (about 60% each).
- **Recommendations.** We provide concrete guidance for operators and discuss the path toward Internet Engineering Task Force (IETF) standardization of limit signaling.

The remainder of this paper is organized as follows: Section 2 describes the maximum-prefix mechanism, the history of standardization attempts, and current operator practices. Section 3 describes our data sources, measurement methodology, and dataset construction. Section 4 analyzes PeeringDB coverage and data quality. Section 5 measures how often ASes exceed their declared limits. Section 6 examines the connectivity impact of exceedance events. Section 7 offers recommendations for operators and the community.

## 2 BGP Maximum-Prefix Limit

The maximum-prefix limit is a locally configured parameter that specifies the maximum number of prefixes a BGP session will accept from a neighbor [12, 19, 37, 47, 60]; to protect against memory exhaustion on routers with limited hardware resources, and route leaks. While most implementations support two thresholds: a configurable warning level at which the router logs an alert but keeps the session alive, and a hard limit at which the session is torn down automatically. Vendors implement this mechanism differently: some, notably Juniper [37] and BIRD [12], add further limit types that differ both in *when* prefix counting occurs and in the *action* taken on exceed such as warn, drop, or restart. Table 1

summarizes vendor behavior based on Job Snijders' RIPE 77 presentation [67] and the expired IETF draft [69], updated with each vendor's latest public documentation. A *pre-policy* limit counts all received prefixes before import filters are applied; a *post-policy* limit counts only prefixes accepted after filtering; and an *outbound* limit caps what the local router advertises to the neighbor. Pre-policy limits are stricter and more predictable; post-policy limits depend on local filters and can vary across sessions with the same peer.

This fragmentation means operators must understand vendor-specific semantics, and, moreover, *no mechanism signals which type of limit is in effect to a remote peer.*

### 2.1 The IETF Standardization Saga

The importance of maximum-prefix limits is broadly recognized across the industry: the mechanism is part of the base BGP specification [60], recommended as a Best Current Practice [22], and implemented by every major router vendor (see Table 1); even outside the networking industry, government security guidance specifically recommends bounding the number of prefixes accepted from a peer [20, 45]. Despite this agreement, the IETF standardization efforts to automatically exchange or negotiate limits between peers have repeatedly stalled. The timeline below summarizes key attempts:

- **April 2004:** draft-chavali-bgp-prefixlimit [59] proposes embedding limits in BGP OPEN message capabilities. No traction.
- **December 2018:** draft-sa-grow-maxprefix [68] introduced in Global Routing Operations Working Group (GROW) by Snijders and Aelmans.
- **September 2019:** Draft migrates to Inter-Domain Routing (IDR) working group as draft-sa-idr-maxprefix [68].
- **February 2021:** Draft splits into separate inbound and outbound documents.
- **March 2022:** Working Group adoption call issued for draft-sas-idr-maxprefix-inbound [69]. Documents implementations from Huawei, Nokia, OpenBSD, and Cisco.
- **February 2023:** Inbound draft reaches version -05 [69].
- **August 2023:** Draft expires without RFC publication.
- **December 2024:** New draft by Abraitis, draft-abraitis-idr-maximum-prefix-orf [1], proposes Outbound Route Filtering (ORF)-based approach. Expires without adoption.

The lack of standardization has left operators without a mechanism to automatically exchange or negotiate limits; as late as June 2023, the twin drafts were still highlighted as significant ongoing work in operator community overviews [6]. Many factors have contributed to this impasse. Competing design approaches (OPEN message capability advertisement versus ORF-based signaling) prevented consensus on the right layer for limit exchange. The pre-policy vs. post-policy divide reflects a genuine semantic tension: transit providers prioritize raw-announcement protection while peering fabrics care about contract enforcement, making a single standard difficult to specify. Two further factors surfaced in our conversations with operators. The fact that any mechanism that signals a peer's configured limit also hands a would-be route leaker the exact headroom available before a teardown triggers, therefore in-band limit negotiation carries a security cost that a static, out-of-band value does not. On top of that, part of the community regards

**Table 1: Vendor implementations of maximum-prefix limits and their exceed actions.**

| Vendor                | Pre-policy    | Post-policy           | Outbound               | Exceed action(s)                                |
|-----------------------|---------------|-----------------------|------------------------|-------------------------------------------------|
| Juniper Junos [37]    | prefix-limit  | accepted-prefix-limit | advertise-prefix-limit | log-only / reject / teardown (idle-timeout)     |
| Cisco IOS XR/XE [18]  | —             | maximum-prefix        | —                      | warn / discard-extra / restart-timer / teardown |
| Nokia SR OS [47]      | prefix-limit  | post-import           | —                      | log-only / teardown / idle-timeout              |
| OpenBSD OpenBGPD [48] | max-prefix    | —                     | max-prefix out         | teardown / restart                              |
| NIC.CZ BIRD [12]      | receive limit | import limit          | export limit           | warn / block / restart / disable                |

limit management as operational hygiene and therefore operator’s own responsibility; holding that it does not belong in the protocol. Finally, the emergence of PeeringDB as a de facto workaround reduced urgency; imperfect but functional, it dampened pressure for a protocol solution even as it shifted the problem from the protocol layer to *manual operational hygiene*.

## 2.2 How Networks Set Their Limits

A 2021 NANOG thread [57] captures the operational confusion left by two decades of stalled standardization: “*I can hardly find any recommendations on how to arrive at a sensible limit... Do you negotiate/exchange sensible values whenever you establish a new session? Do you rely on PeeringDB?*” Quite tellingly, operators in the same thread independently proposed what the IETF had been attempting to standardize, namely embedding limits in BGP OPEN messages or encoding them as RPKI objects, suggesting broad demand that standardization never satisfied.

**2.2.1 General Practices.** Drawing on this thread, publicly available documentation, and conversations with network engineers at RIPE meetings and regional Network Operators Group (NOG)s, we identify five common approaches operators have developed in the absence of a standard:

- **PeeringDB-based:** Operators query PeeringDB `info_prefixes` fields, as they represent the network suggested limits [52, 53]<sup>1</sup> and apply a multiplier. Multipliers vary widely in practice: a common approach is 2x the PeeringDB value, falling back to 1.5x (or the raw value itself) when the result would approach full-table size [38]. More sophisticated operators combine three inputs (PeeringDB value, Internet Routing Registry (IRR) prefix expansion, and currently observed count), select the most reliable, then apply ~30% headroom [13].
- **Bucket tiering:** Some operators assign peers to one of 5–6 fixed capacity tiers based on their expected prefix count. At ~90% utilization, the peer is moved to the next tier, eliminating per-peer limit management for the vast majority of sessions [9]. Transit sessions instead use tighter, table-size-based limits.
- **Conservative thresholds:** Some operators eyeball the current observed prefix count and add a buffer, treating the limit as a last-resort failsafe for gross misconfigurations rather than a precision tool [57]; others apply a fixed blanket default (e.g., 1000 prefixes) for all but the largest networks, a practice corroborated by our measurements (Section 4, Figure 5).

<sup>1</sup>PeeringDB’s documentation is inconsistent here: its field help and API documentation describe the value as a “recommended maximum number of routes/prefixes to be configured on peering sessions for this ASN,” whereas its operator getting-started guide frames it as the number of prefixes that “peers [should] expect to see advertised by your network” [50]. We adopt the limit reading; the two coincide when operators declare a value close to what they advertise.

- **Manual negotiation:** For significant peering relationships, limits are discussed via email, phone, or at forums.

The thread also revealed data quality concerns about PeeringDB, noting that approximately 30–40% of IPv4/IPv6 entries carry a value of 0 [55], and that stale entries frequently cause newly established sessions to immediately trip limits [38]. Operators report that hygiene is handled ad-hoc: some rely on quarterly calendar reminders to audit configured limits [57], while best practice calls for daily snapshotting of peer prefix counts to catch growth before it causes session drops [57], concerns echoed across the community at the RIPE Routing Working Group [56].

**2.2.2 Automation Tools.** Network automation tools handle prefix limit configuration inconsistently, reflecting the absence of a standard data source. `bgpq4` [65] generates vendor-specific prefix-list filters from IRR data but does not produce max-prefix session limits. `ARouteServer` [71] queries PeeringDB for AS-SET information used in IRR filtering but does not read the `info_prefixes{4,6}` fields; its max-prefix support covers only enforcement actions (e.g., session shutdown or a 15-minute restart timer), not limit values. `IXP Manager` [35, 36] auto-populates limits at member onboarding using a single value applied to both IPv4 and IPv6, falling back to 250 prefixes when no per-member value is configured, independent of PeeringDB-declared values. `Peering Manager` [49] is the exception: it exposes `IPv4 Max Prefixes` and `IPv6 Max Prefixes` as configurable fields and can synchronise them from the operator’s PeeringDB record. Three of the four therefore bypass `info_prefixes` entirely, leaving the values operators maintain unused by the toolchain meant to operationalize them.

PeeringDB addresses these operational needs but relies on voluntary participation and manual data entry, with no enforcement mechanism; our work is the first empirical characterization of how well it functions in practice.

**Takeaways.** *PeeringDB [51] is the de facto community-maintained coordination mechanism for max-prefix limits, yet no mechanism enforces accuracy or currency of its entries, and vendor fragmentation means operators cannot know what limit type their peer has configured or whether the published value reflects the real status of the actual system.*

## 3 Data and Methodology

We now describe the data and methodology behind our evaluation of how accurately declared maximum-prefix limits reflect observed prefix announcements. All measurement code and processed datasets are publicly released (see Appendix B).

### 3.1 Data Sources

We construct a working dataset of ASes with observed routing behavior, recently maintained prefix limit information, and network importance rankings, combining RIPE RIS routing snapshots, PeeringDB metadata, and CAIDA AS Rank [14].

**RIPE RIS [63].** We use Routing Information Base (RIB) snapshots from all 23 RIPE RIS route collectors active during 2025, spanning January 1 to December 31. Each collector (identified as RRC00–RRC26, non-sequential) peers with BGP speakers at one or more IXP peering LANs; three of them (RRC00, RRC24, and RRC25) operate as multi-hop collectors, accepting sessions from ASes anywhere on the Internet rather than only from co-located IXP participants. Collectors produce continuous BGP update feeds (rotated every 5 minutes) and full RIB dumps every 8 hours; we use the latter. The collectors are geographically distributed across Europe, the Americas, Asia, and Africa, though coverage is uneven, with 14 of 23 collectors concentrated in Europe.

Snapshots are taken at 00:00, 08:00, and 16:00 UTC, yielding 3 dumps per collector per day and 25,185 RIB files in total across the observation period. Each dump contains prefix announcements with full AS paths, timestamps, and origin ASes, providing visibility into 90,451 unique ASes over the year. We parse all MRT-formatted RIB files using `pybgpkit` [10]. For each snapshot, we compute announced prefix counts per AS; the precise computation and visibility threshold are detailed in Section 3.2.

**PeeringDB [51].** We collected daily PeeringDB snapshots throughout 2025 from the CAIDA public PeeringDB archive [15]. For each network with an entry, we extract the suggested IPv4 and IPv6 (`info_prefixes4` and `info_prefixes6`) prefix limits, last update timestamp (updated), and network type (`info_type`). PeeringDB limits are *suggested* values: they reflect what a network *recommends* peers to configure, not the limits their peers actually deploy on routers. Because `info_prefixes` is a single value a network advertises to all of its peers rather than a per-session setting, we read it as a network-wide suggested limit.

**CAIDA AS Rank [14].** We use the CAIDA AS Rank dataset to classify ASes by their topological position in the Internet hierarchy. Each AS is assigned a rank based on its customer cone size, *i.e.*, the number of ASes it provides transit for directly or indirectly, with rank 1 being the most central.

### 3.2 Methodology

We adopt the upstream, downstream and peer terminology from `bgp.tools` [11]: an *upstream* exists when the AS reaches the Tier-1 core through a neighbor; a *downstream* is the same adjacency in the opposite direction (the neighbor reaches the core through this AS, *i.e.*, a customer); and a *peer* is any lateral relationship between two ASes. We define Tier-1 ASes following the `bgp.tools` knowledge base [11] (see Appendix D for the full list).

For each snapshot, *announced* prefixes are computed by identifying all AS paths containing a Tier-1 AS and assigning each prefix to every AS between the rightmost Tier-1 and the origin, inclusive, covering both prefixes the AS originates itself and those from downstream customers it transits toward the global backbone, following the methodology in [11]. Prefixes confined to local peering or regional transit are excluded, as peers enforcing PeeringDB-based

**Table 2: Sensitivity of exceedance and impactful-event counts to the high-visibility threshold. *Exceedance*: the announced count, visible at the given threshold, exceeds the declared limit; *Impactful*: the exceedance’s coincident peer drop exceeds the AS’s own 95th-percentile churn baseline.**

| Threshold | IPv4       |           | IPv6       |           |
|-----------|------------|-----------|------------|-----------|
|           | Exceedance | Impactful | Exceedance | Impactful |
| 80%       | 2283       | 181       | 1133       | 61        |
| 90%       | 3447       | 265       | 1195       | 75        |
| 95%       | 4314       | 303       | 1501       | 132       |
| 100%      | 5060       | 399       | 2371       | 198       |

limits would not observe them on their sessions. For the exceedance analysis in Sections 5–6, we use this count exclusively, making our exceedance counts a conservative lower bound. This count corresponds to a *pre-policy* view (*i.e.*, prefixes received before any import filters are applied), the strictest of the three limit types described in Section 2 and the most comparable across vendor implementations.

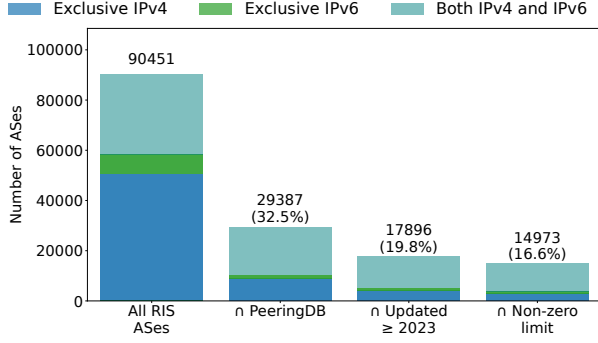
Throughout this work we use AS-level peer count as a proxy for connectivity impact: a sudden drop in visible peers signals that BGP sessions have been torn down, with direct consequences for reachability, latency, and traffic delivery; we acknowledge that, without direct traffic or latency measurements, this remains an approximation, as mentioned on Limitations on Section 3.4.

To reduce noise from transient route leaks, hijacks, and partial propagation, we apply a *high-visibility* filter adapted from [40]: a prefix is included in an AS’s effective count only if it is seen by at least 95% of the route collectors (22 out of 23) at that snapshot. The choice of 95% reflects a deliberate trade-off between precision and recall. Setting the threshold lower (*e.g.*, 70–80%) risks counting prefixes not part of an AS’s global routing footprint: prefixes selectively announced only to specific IXP route servers under looser local policies, routes advertised under a private agreement, or prefixes still propagating slowly and not yet widely visible; in each case, the prefix would not be seen by peers enforcing PeeringDB-based limits, but our analysis would incorrectly record an exceedance. Setting the threshold at 100%, on the other hand, is overly strict: a single collector outage or transient peering failure would suppress a prefix even when it is genuinely and widely propagated. At 95%, the threshold is narrow enough to exclude spurious or partially-propagated announcements yet broad enough to tolerate a single-collector outage or transient peering failure.

The two filters are complementary: the Tier-1 filter is a *quality* check (only prefixes observed in globally-propagated paths are counted) while the 95% visibility threshold is a *consistency* check (the prefix must be stably observed across the collector infrastructure, not just glimpsed from one vantage point). Together, they define a prefix that is clearly part of an AS’s global routing footprint.

Table 2 reports the number of exceedance events and how many were impactful, across four thresholds. The exceedance count increases from 80% to 100%. This is the opposite of the naive expectation that counting fewer prefixes yields fewer exceedances. However, an exceedance is a below-to-above *transition*, so the lower counts at stricter thresholds move ASes out of the chronically-above





**Figure 2: Filtering funnel from all RIPE RIS-visible ASes to the analytical sample with at least one non-zero limit. Each bar labeled  $\cap$  is a strict subset of the previous.**

regime, where they never transition and are excluded, and into the near-limit regime where they can cross. Two artifacts bound the extremes: at permissive thresholds the excluded, chronically-above ASes are largely those with stale, low PeeringDB limits, deflating the count; at 100%, single-collector outages produce transient dips below the limit that register as spurious exceedances, inflating it. Among these exceedances, we also report how many are impactful (defined in detail in Section 6.1): those whose coincident peer drop exceeds the AS’s own 95th-percentile churn baseline. The share of impactful exceedances stays within a narrow band across thresholds, 7–8% for IPv4 and 5–9% for IPv6. The 95% choice therefore rescales both columns together instead of distorting the fraction of exceedances that coincide with a meaningful peer loss.

### 3.3 Dataset Construction

We build our working dataset through a four-step filtering pipeline, illustrated in Figure 2, and an auxiliary peer graph.

**Step 1: RIPE RIS.** We observe 90,451 unique ASes in RIPE RIS during 2025: 50,507 announcing exclusively IPv4 prefixes, 7,712 exclusively IPv6, and 32,232 both.

**Step 2: PeeringDB intersection.** Of these 90,451 ASes, 29,387 (32.49%) have a corresponding PeeringDB entry. The remaining two-thirds of visible ASes publish no public prefix limit information.

While coverage may appear low, it is concentrated where limits matter: 75.5% of transit-providing ASes (those with customers) appear in PeeringDB, versus only 19.7% of single-homed stubs.

**Step 3: Recency filter.** A PeeringDB entry that has not been updated in years may reflect outdated limits rather than current practice. We restrict to entries last updated on or after January 1, 2023, a conservative threshold that excludes entries untouched for over two years. This yields 17,896 ASes (19.79% of all RIS-visible ASes); among these, 4,110 are exclusively IPv4, 1,118 exclusively IPv6, and 12,668 announce both.

**Step 4: Non-zero filter.** A PeeringDB entry with a limit of zero, while valid, carries no information for our measurement. We therefore retain only ASes with at least one non-zero limit in either address family, yielding 14,973 ASes (16.55% of all RIS-visible ASes), which form our working dataset for the exceedance and impact analyses in Sections 5–6.

**AS-level peer graph.** For each of the 1,095 RIB snapshots spanning the observation year, we construct an AS-level peer graph

using NetworkX [32], with nodes representing all visible RIS ASes and directed edges representing adjacency relationships, enabling efficient per-snapshot queries for any AS in the working dataset. From this graph we derive *peer count*, the number of unique adjacent ASes observed in a snapshot, and *relative peer loss*, the fractional decrease in this count across an exceedance. Because peer count derives from the AS-path adjacency graph, it is independent of whether the AS originates any prefixes: it remains a node with intact adjacencies as long as it appears in any observed path.

### 3.4 Limitations

Several constraints apply to our methodology and should be kept in mind when interpreting our results.

First, the `info_prefixes4` and `info_prefixes6` fields in PeeringDB reflect what a network *recommends* peers configure, not what is actually configured on routers; a router may enforce a stricter or more permissive limit than what PeeringDB advertises. Consequently, our exceedance events are those where the *announced* count crosses the *declared* PeeringDB limit; actual enforcement may occur at a different threshold. This means we may observe events that were not enforced (because the router was configured more permissively) and may miss events that were enforced (because the router was configured more strictly). Both failure modes act in the direction of measurement uncertainty rather than systematic bias: our reported enforcement-associated peer drops are conservative lower bounds on events where the PeeringDB limit matches the actual router configuration.

Second, PeeringDB does not distinguish which type of limit is intended. Our prefix counts reflect announcements visible to RIPE RIS (a pre-policy view), but operators may configure post-policy limits that differ based on their import filters. This mismatch could lead us to over- or under-estimate exceedance rates if import policies substantially reduce the announced prefix set.

Third, our analysis is restricted to RIPE RIS [63] and does not incorporate other public route collector infrastructures such as RouteViews. This is a deliberate choice: RIPE RIS and RouteViews are highly redundant for widely propagated prefixes [2, 3], and because our visibility metric requires a prefix to be present on 95% (22 of 23) of collectors, a marginal collector adds little. Emerging platforms such as `bgproutes.io` [4] impose strict per-hour data-volume limits (100 MB/h) and were still under active development during our observation period, making year-scale bulk analysis infeasible. On balance, RIPE RIS offers a suitable tradeoff between visibility, stability, and data volume at year scale; our methodology could extend to other collector infrastructures in future work. Within RIPE RIS, vantage-point coverage is uneven: sessions between networks not well-connected to RIS peers may be missed, biasing our sample toward larger, better-connected ASes. Our high-visibility filter compounds this at a finer grain: it counts route collectors rather than individual collector-peers, so collectors with many and few peers contribute equally toward the 95% threshold; a peer-weighted metric is a refinement we leave to future work.

Fourth, our methodology is based on 8-hour RIB snapshots [63], *i.e.*, point-in-time captures of routing table state at 00:00, 08:00, and 16:00 UTC, rather than continuous BGP update streams. An event that begins and fully recovers between two consecutive snapshots

leaves no trace in our data: a peer session that drops and restores within a single 8-hour window goes undetected. To understand the impact of this limitation, we examined the BGP update stream for our case studies, confirming both the exceedance and its persistence. This finer view also illustrates instances where snapshots are a sound base: the raw update stream carries session-reset *zombie* routes [24], *i.e.*, stale paths a peer stops refreshing without an explicit withdrawal, which only periodic RIB dumps reconcile away. The 8-hour snapshots under a 95% visibility threshold therefore yield a noise-robust base that update-level replay alone does not. A systematic update-stream study is left as future work.

Finally, when a peer-count drop coincides with a limit exceedance, *we treat it as a strong correlation, not definitive causal evidence*. A teardown is only one of several configurable exceed-actions, as shown in Table 1, and the only one externally visible: because warn-only enforcement leaves no trace in BGP data, our teardown-based counts are a lower bound on configured max-prefix enforcement, missing sessions where the operator chose to warn rather than drop. Conversely, we do not claim every impactful event is exceedance-caused. A peer drop coinciding with an exceedance may instead reflect deliberate traffic engineering or DDoS mitigation, and RIB-only data cannot always separate limit enforcement from an intentional withdrawal. Our per-AS churn floor (see Section 6) discards drops indistinguishable from ordinary churn, and we classify events by structural signature to flag the ambiguous ones, but we frame the results as exceedance-coincident peer loss rather than a strict causal claim. Without access to router logs, we cannot rule out concurrent unrelated events such as planned maintenance or other misconfigurations. Our update-level reconstruction of the case studies (see Section 6.3) confirms these drops at 5-minute resolution, yet the peer counts still derive from RIS control-plane best-paths rather than directly observed session state. Independent validation against live looking glasses or active reachability probes is a natural extension for future work.

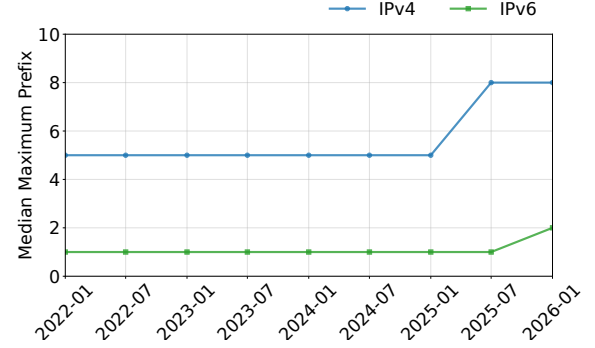
## 4 PeeringDB Coverage and Data Quality

We now characterize PeeringDB as a data source: how recent its entries are, which networks participate, and how declared limits distribute across address families and network types.

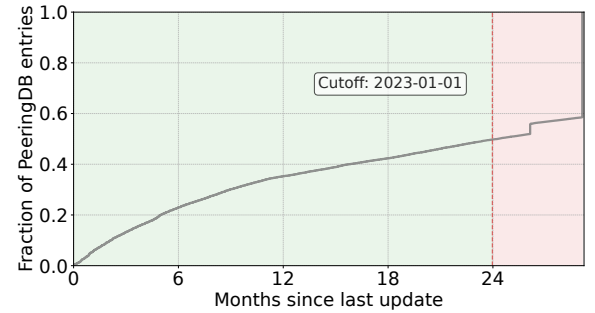
### 4.1 Staleness Analysis

The median declared limit in PeeringDB is essentially unchanged over the four-year observation window. Figure 3 shows the median limit across all entries at six-month intervals from 2022 to 2026, unfiltered. The median IPv4 limit holds at 5 prefixes across 7 of 9 snapshots, stepping to 8 from mid-2025 onward; the median IPv6 limit holds at 1 across 8 of 9 snapshots, stepping to 2 only at the start of 2026. The upward steps reflect a mix of newly registered entries carrying higher limits and incremental operator increases.

To understand why, we examine entry age. Figure 4 shows the CDF of entry ages across all 33,130 PeeringDB entries as of January 1, 2025. Two step discontinuities near months 27 and 30 correspond to PeeringDB platform bulk updates in October and July 2022 (affecting 3.3% and 34.5% of all entries, respectively), where timestamps were refreshed without any operator action.



**Figure 3: Median declared maximum-prefix limit at 6-month intervals (2022–2026), all PeeringDB entries unfiltered. IPv4 steps from 5 to 8 at mid-2025; IPv6 steps from 1 to 2 at the start of 2026.**



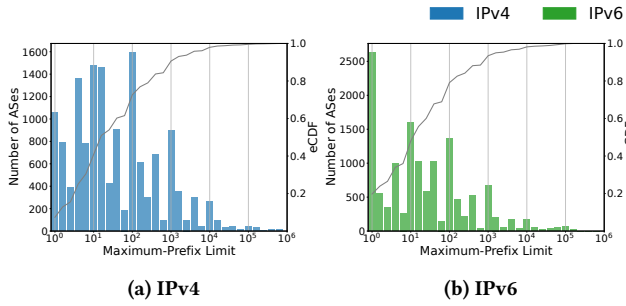
**Figure 4: CDF of PeeringDB entry ages as of January 1, 2025. Entries updated within 24 months are retained (green); entries older than our cutoff are excluded (red). Steps near months 27 and 30 reflect mid-2022 PeeringDB platform bulk resets, not operator edits.**

These bulk resets motivate our two-year exclusion threshold: cutting at January 1, 2023 excludes entries whose most recent timestamp traces to a platform event, retaining only those with genuine operator activity since 2023. A looser threshold would readmit both these bulk-refreshed timestamps and older entries whose limits pre-date current practice, confounding our exceedance measurement with values no operator has revisited in years. Of the 33,130 entries, 19,391 (58.5%) satisfy this criterion; the remaining 13,739 (41.5%) have not been updated in over two years. Intersecting with RIPE RIS yields 17,896 ASes.

A recently updated entry is necessary but not sufficient: in PeeringDB, zero is the uninitialized default for the limit field, not an explicit “no limit” declaration. Among the 19,391 recently updated entries, 4,213 IPv4 limits (21.73%) and 5,226 IPv6 limits (26.95%) are zero; 5,953 entries (30.70%) carry a zero in at least one address family, and 3,486 (17.98%) are zero in both. These rates fall below the 30–40% operator estimate from community discussions [55, 66], because our recency filter already excludes entries untouched for over two years, which disproportionately carry zero values. The zeros that remain belong to recently active networks that nonetheless never set a meaningful limit. Excluding these leaves 14,973 ASes with at least one actionable limit.

**Table 3: PeeringDB coverage as of December 2025.**

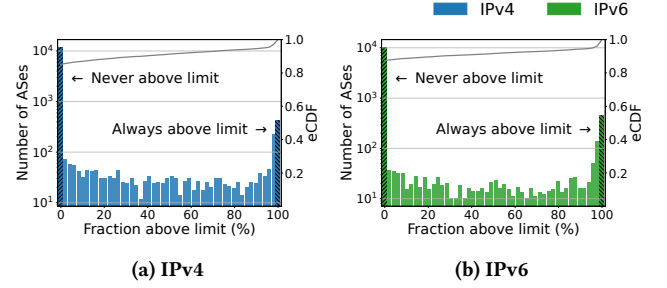
| Network Type         | Visible ASes | Percentage |
|----------------------|--------------|------------|
| Cable/DSL/ISP        | 7,104        | 47.4%      |
| NSP                  | 2,400        | 16.0%      |
| N/A                  | 1,509        | 10.1%      |
| Content              | 1,324        | 8.8%       |
| Educational/Research | 889          | 5.9%       |
| Enterprise           | 758          | 5.1%       |
| Network Services     | 553          | 3.7%       |
| Non-Profit           | 297          | 2.0%       |
| Government           | 72           | 0.5%       |
| Route Server         | 60           | 0.4%       |
| Route Collector      | 7            | 0.0%       |

**Figure 5: Distribution of declared maximum-prefix limits (log scale, December 31, 2025); spikes at exact powers of ten reveal a round-number bias in operator configuration.**

## 4.2 Coverage and Limit Distribution

Table 3 breaks down the 14,973 working-dataset ASes by PeeringDB’s `info_type` field as of December 2025. Connectivity providers (Cable, DSL, and ISP networks) dominate at 47.4%, reflecting that transit ISPs are the most active PeeringDB participants; peering coordination is a core operational concern for networks that sell transit. NSPs (16.0%) and Content providers (8.8%) follow, consistent with their reliance on coordinated peering agreements. The N/A category (10.1%) carries no declared type; its reduced share relative to the broader PeeringDB population reflects a disproportionately high zero-limit rate among these entries. It likely also masks misclassified networks that would shift type-level rates if properly labeled. Educational, Enterprise, and Government networks are underrepresented relative to their actual prevalence on the Internet. These organizations more commonly reach the global routing table through upstream transit providers rather than through IXPs, where PeeringDB participation is most active.

Beyond who participates, a natural question is what limits operators actually declare. Figure 5 shows the distribution of declared limits on a log scale across the working dataset. The distribution is heavily right-skewed, with most limits concentrated below 100 prefixes. The most striking feature is a pronounced *round-number bias*: sharp spikes at exact powers of ten (1, 10, 100, 1,000, 10,000) stand well above neighboring bins in both address families, indicating that in some instances operators select limits by Eyeballing and/or intuition rather than by measuring their actual prefix footprint.

**Figure 6: Distribution of the share of 8-hour snapshots each AS spends above its declared limit. Hatched bars mark the endpoints: ASes never exceeding (leftmost) and exceeding in every snapshot (rightmost).**

**Takeaways.** Only 32.49% of RIS-visible ASes appear in PeeringDB, 21.7–27.0% of recent limit fields carry a value of zero, and declared limits cluster at powers of ten, revealing that operators set them by Eyeballing and/or intuition rather than measurement. Together, these gaps leave the max-prefix mechanism without a reliable configuration baseline for the vast majority of Internet ASes; only 14,973 (16.55%) carry an actionable limit.

## 5 How Often Do ASes Exceed Their Limits?

Leveraging our 14,973 working-dataset ASes, we measure how their announced prefix count relates to their suggested (and recently updated) PeeringDB limit, over each of the 8-hour snapshots. Figure 6 shows the resulting distribution. Among the 13,978 IPv4 ASes: 11,643 (83.30%) never exceed their limit; 2,034 (14.55%) exceed it in 1–99% of snapshots; and 301 (2.15%) exceed it in every observed snapshot. Similarly, the 11,757 IPv6 ASes: 10,163 (86.44%) never exceed; 1,321 (11.24%) exceed intermittently; and 273 (2.32%) exceed persistently. These proportions suggest that PeeringDB is adequate for the large majority of networks, and in particular for those it is designed to serve, *i.e.*, networks that peer or are multi-homed. We deliberately study the remaining minority, the 16.70% of IPv4 and 13.56% of IPv6 ASes that exceed at least once, operators for whom prefix limits matter yet the declared value is mismatched.

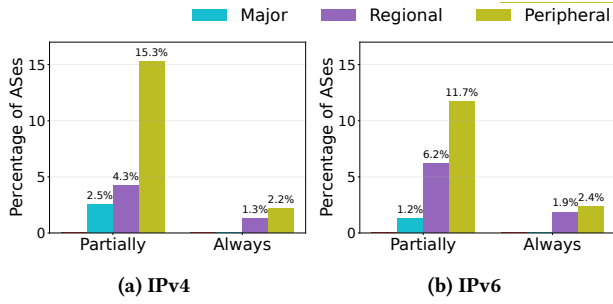
The *Partially* exceeding category is roughly 7× larger than the *Always* category for IPv4 (2,034 vs. 301 ASes) and 5× for IPv6 (1,321 vs. 273), suggesting that transient churn, not merely stale limits, drives the majority of exceedances. The 2% that are *Always* above their limit indicate a permanent mismatch. Their declared PeeringDB limit has never reflected their actual routing footprint throughout the entire observation year.

Beyond how *often* ASes exceed, we also measure by *how much*. Across all exceedances, *i.e.*, every below-to-above transition treated as one independent event, the announced count sits a median of +40% (IPv4) and +100% (IPv6) over the declared limit. In absolute terms the excess is small: *nearly half of exceedances are a single prefix over the limit* (42% IPv4, 48% IPv6), and roughly 90% are within ten. A stricter threshold of 1.5× or 2× the limit would therefore miss the majority of real exceedances, discarding 64% and 37% of events at 1.5× and 83% and 70% at 2×; we thus flag exceedance at the declared limit itself. Appendix E reports the full distribution, which



**Table 4: Exceedance rates by network type. Remaining ASes never exceeded their declared limit.**

| Type             | Partially    |       | Always      |      |
|------------------|--------------|-------|-------------|------|
|                  | IPv4         | IPv6  | IPv4        | IPv6 |
| Cable/DSL/ISP    | <b>16.9%</b> | 10.1% | 2.1%        | 1.8% |
| Network Services | 14.6%        | 8.1%  | <b>3.4%</b> | 2.0% |
| N/A              | 12.6%        | 8.4%  | 2.7%        | 1.9% |
| Enterprise       | 12.5%        | 7.0%  | 2.0%        | 2.1% |
| Content          | 11.5%        | 7.3%  | 1.8%        | 2.1% |
| NSP              | 10.0%        | 7.5%  | 1.6%        | 1.7% |
| Government       | 9.7%         | 5.6%  | 2.8%        | 1.4% |
| Non-Profit       | 7.1%         | 9.4%  | 1.0%        | 2.7% |
| Route Server     | 5.0%         | 3.3%  | 0.0%        | 0.0% |
| Edu/Research     | 4.8%         | 7.4%  | 1.0%        | 1.2% |
| Route Collector  | 0.0%         | 14.3% | 0.0%        | 0.0% |

**Figure 7: Exceedance rates by CAIDA AS Rank tier. Percentages are computed over all ASes in the working dataset with a valid rank assignment in each tier.**

empirically justifies using the declared limit itself, rather than a stricter multiple, as our exceedance threshold.

### 5.1 Exceedance by Network Type and AS Rank

Table 4 breaks down exceedance rates by PeeringDB network type. Cable/DSL/ISP networks exhibit the highest episodic exceedance rate (16.9% IPv4), consistent with their role as connectivity providers for residential and business customers whose prefix counts grow organically over time. Network Services providers show the highest persistent-misconfiguration rate (3.4% IPv4 always above limit), suggesting that operators set their *PeeringDB entries* once and rarely revisit them. In principle, their router configurations may well be maintained more actively; but the coordination data others rely on that drifts out of date. NSPs and Route Servers exhibit the best hygiene, with near-zero persistent exceedance rates (1.6% and 0.0% IPv4, respectively). Their small representation in the dataset (Table 3) and limited visibility in RIB dumps warrant caution when interpreting these rates. Route Collectors are a special case: their low counts and unusual role make comparisons unreliable.

Figure 7 breaks down exceedance rates by CAIDA AS Rank tier. All rates are computed within each tier: Peripheral ASes (rank >1,000) show the highest intermittent exceedance rate at 15.3% (IPv4) and 11.7% (IPv6), 3.6× and 1.9× higher than Regional ASes

(4.3% and 6.2%, respectively), while Major ASes sit at 2.5% (IPv4) and 1.2% (IPv6). Tier-1 ASes show no measurable exceedance in either address family. Persistent exceedance (always above) follows the same gradient but at lower absolute rates: Peripheral 2.2%/2.4%, Regional 1.3%/1.9%, Major and Tier-1 near zero. The monotonic decrease from Peripheral to Major is consistent with larger, better-resourced operators maintaining tighter coordination between their PeeringDB entries and operational routing policy.

**Takeaways.** About 2% of ASes persistently exceed their declared limits throughout the year, and about 15% (IPv4) or 11% (IPv6) do so intermittently, the signature of organic growth against stale coordination data, not route leaks. Exceedance follows a clear rank gradient: Peripheral ASes exceed intermittently while Tier-1 ASes show none, consistent with larger operators maintaining tighter coordination between their PeeringDB entries and operational policy.

## 6 Operational Impact and Causes

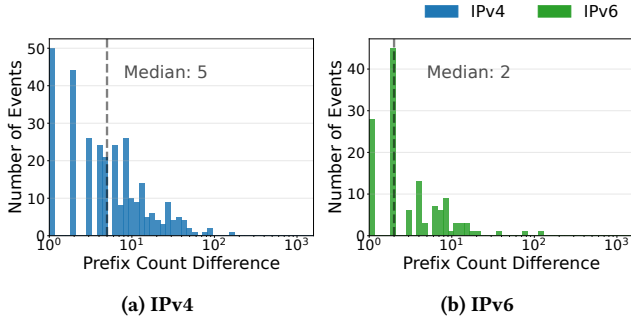
When a peer enforces the PeeringDB-suggested limit, an exceedance triggers an automatic session teardown and a direct connectivity loss. The question is not how many peers are lost, but which: severity depends on the disconnected neighbor.

### 6.1 Detecting Exceedance-Driven Peer Loss

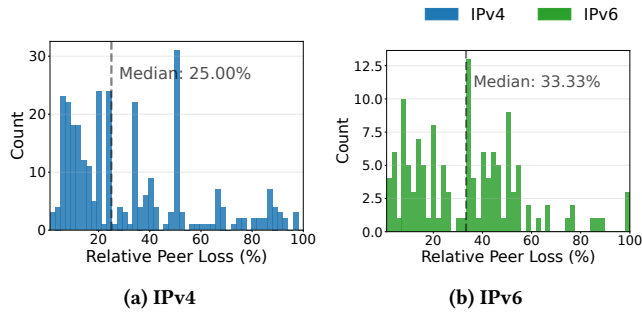
We define an *exceedance event* as a transition in which an AS’s announced prefix count crosses its PeeringDB limit from below to above within a single 8-hour observation window. We further classify an event as *impactful* when its coincident peer drop exceeds that AS’s own churn baseline. The baseline is defined as the 95th percentile of the AS’s peer drops across all non-exceedance 8-hour windows in the year. Therefore coincident peer drops that are indistinguishable from ordinary churn are discarded, and only those that clearly exceed the AS’s own historical variability are classified as impactful. Because the bar is each AS’s own churn, it is size-aware by construction, adapting to large and small networks alike rather than imposing a single fixed percentage, and it separates exceedance-driven loss from ordinary rerouting. We report event counts under the 90th, 95th, and 99th-percentile floors as a sensitivity check. Appendix C shows the impactful count is robust to the choice of percentile floor. We measure the drop over the current and the following snapshot, which accommodates enforcement that only becomes visible one snapshot after the crossing.

Across our working dataset, where a single AS may cross its limit multiple times during the year, we identify 303 IPv4 events (across 219 ASes) and 132 IPv6 events (across 81 ASes) whose coincident peer drop clears the AS’s own 95th-percentile churn floor. During these *impactful* events, the median absolute prefix increase over the prior 8-hour window is just 5 prefixes for IPv4 and 2 for IPv6, as shown in Figure 8. This pattern is consistent with gradual growth on a stale limit rather than a sudden routing anomaly.

Figure 9 shows the distribution of relative peer loss across all impactful exceedance events. The distribution is right-skewed: most events cause modest drops, but a long tail of severe incidents pulls the mean well above the median. For IPv4 the median relative peer loss is 25.0% against a mean of 32.0%; for IPv6 the median is 33.3% and the mean 33.4%. The per-AS churn baseline discards



**Figure 8: Distribution of absolute prefix count increase at impactful exceedance events (log scale).**



**Figure 9: Distribution of relative peer loss (%) immediately after impactful exceedance events.**

the many marginal fluctuations and retains only losses that clearly exceed an AS's ordinary variability, thus the surviving events are correspondingly severe. In the most severe cases, 20 IPv4 and 5 IPv6 ASes lose more than 85% of their visible peers, a near-complete loss of reachability. At the extreme, an AS losing all BGP sessions ceases to appear in the RIS observation space, a complete routing outage. Still, aggregate statistics understate operational impact: a small peer loss among 200 BGP sessions is negligible if the lost sessions are low-traffic peers, whereas the same percentage severing a Tier-1 transit link produces a network partition that effectively isolates the AS from most of the Internet.

## 6.2 Connectivity Impact by Lost Peer Rank

To understand how severe these losses are, we classify peers using CAIDA AS Rank [14] into four tiers. **Tier-1**: the 16 settlement-free backbone providers listed in Appendix D [11]. **Major**: AS Rank from 1 to 100 excluding Tier-1. **Regional**: AS Rank from 101 to 1,000. **Peripheral**: AS Rank >1,000.

Of the 303 IPv4 impactful events, 170 (56.1%) include at least one Tier-1 or Major peer in the lost-peer set, with 43 severing at least one Tier-1 provider: one in seven IPv4 events cuts a Tier-1 transit link. For IPv6, 78 of the 132 impactful events (59.1%) include at least one Tier-1 or Major peer, with 13 losing at least one Tier-1 provider. Table 5 breaks down the 1,771 IPv4 and 5,157 IPv6 peer connections lost across all impactful events by AS Rank tier, expressed as a share of each address family's total lost peers.

The majority of lost peers by count are Regional and Peripheral ASes, consistent with enforcement events affecting the full

**Table 5: AS Rank tier of peers lost during exceedances.**

| AS Rank tier | Lost peers (IPv4) | Lost peers (IPv6) |
|--------------|-------------------|-------------------|
| Tier-1       | 2.7%              | 0.3%              |
| Major        | 15.7%             | 3.1%              |
| Regional     | 46.6%             | 12.3%             |
| Peripheral   | 35.0%             | 84.0%             |
| Unknown      | 0.0%              | 0.3%              |

peering fabric rather than transit specifically. Yet Tier-1 and Major peers account for 18.4% of IPv4 and 3.4% of IPv6 lost peers, a disproportionately high fraction given that Tier-1 and Major ASes represent roughly 0.1% of all ASes. The lower top-tier fraction in IPv6 is likely attributable to IXP route-server cascades: a single enforcement event can drop hundreds of route-server-mediated sessions with peripheral participants simultaneously, inflating the Peripheral count relative to top-tier losses.

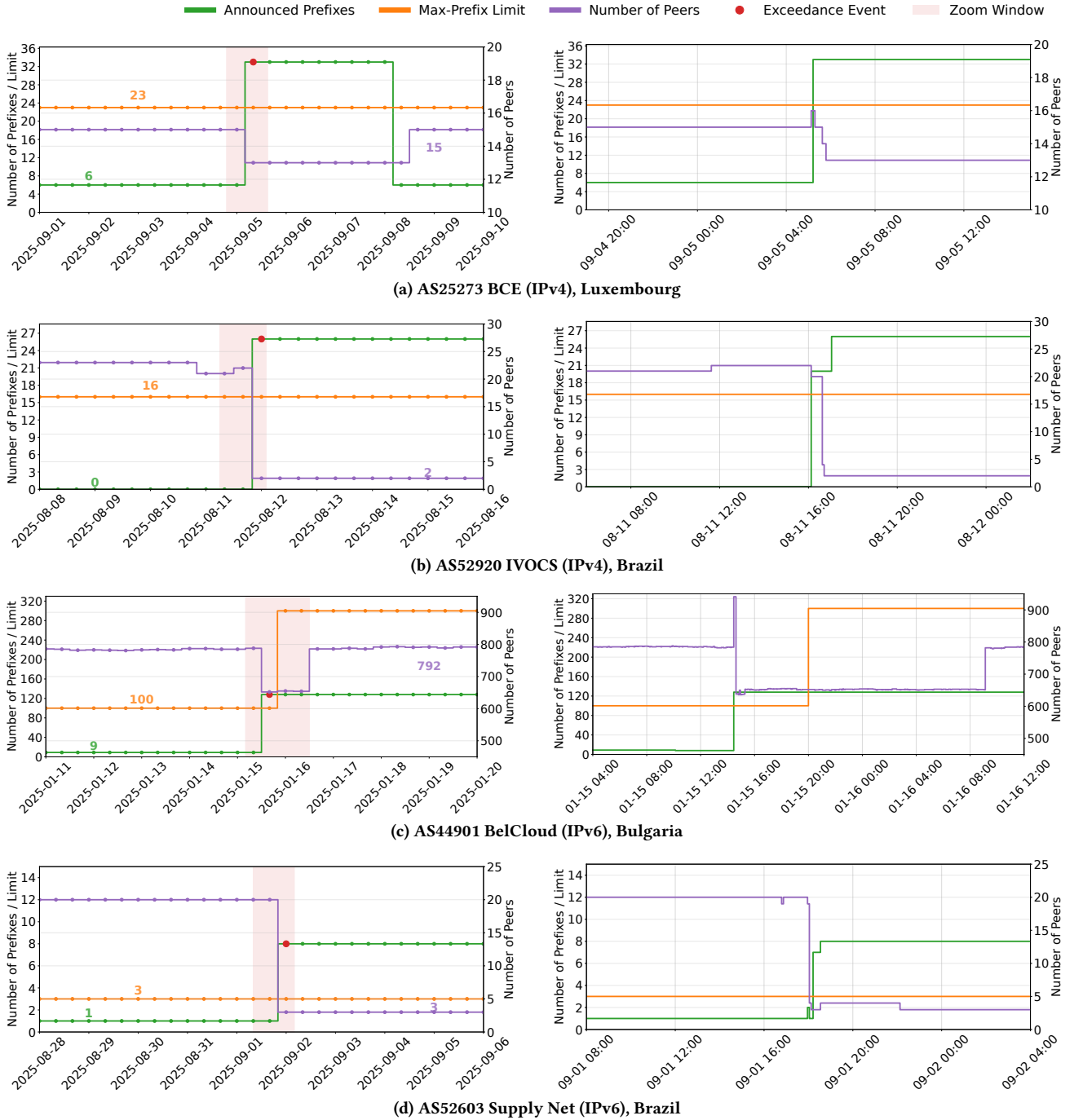
## 6.3 Case Studies

To make the operational impact more concrete, instead of reporting aggregate statistics, we examine four cases, two IPv4 and two IPv6, chosen to span what an exceedance can look like in our data: a deaggregation we cannot separate from traffic engineering, a clean enforcement, a fast operator response, and an enforcement RIBs vantage cannot fully observe. For each case we pair the 8-hour RIB view with the update stream at 5-minute resolution, re-anchoring at each snapshot to discard session-reset *zombie* routes [24] that the raw stream would otherwise carry.

### Case 1: AS25273 – Broadcasting Centre Europe (BCE/ IPv4).

Figure 10a shows that BCE, the broadcast-infrastructure subsidiary of RTL Group, had maintained a stable announced IPv4 count of roughly 6 prefixes against a PeeringDB limit of 23 for most of 2025. On September 5, the announced count jumped from 6 to 33, a deaggregation of BCE's existing address space: all 27 new prefixes are /24 sub-announcements of the five /22 and one /20 block already in their routing footprint, instantly crossing the limit. Within the same 8-hour snapshot, BCE lost three Tier-1 transit providers simultaneously: Level3/Lumen (AS3356, rank 1), Arelion (AS1299, rank 2), and Cogent (AS174, rank 4). It also gained one new session at LU-CIX (AS208374), the Luxembourg Internet Exchange route server, bringing the net peer count from 15 to 13. The update-level zoom first shows the peer count increase by one, due to the new LU-CIX session. It then resolves the loss into a staggered teardown of the three Tier-1 upstreams within 40 minutes.

This case, however, reads two ways. The deaggregation and the simultaneous LU-CIX session are equally consistent with traffic engineering or DDoS mitigation, plausible given Luxembourg's heightened threat environment in 2025: a July cyberattack on the national operator POST disrupted mobile and emergency services nationwide [31, 41]. From BGP data alone we cannot attribute the transit loss to either the three providers enforcing their limits against the 33-prefix announcement or BCE intentionally withdrawing from them, since the two are indistinguishable in best-path data.



**Figure 10: Our four case studies. In each panel, the left plot shows the 8-hour RIB: announced prefix count (green) and PeeringDB limit (orange) on the left axis, peer count (purple) on the right, and the red dot marks the exceedance snapshot. The right plot zooms into the shaded band around the exceedance, resolving the peer loss at 5-minute update resolution.**

Either way the lesson holds: whether enforced or intentional, those more-specifics crossed a suggested limit too low to carry them.

**Case 2: AS52920 — Internet e Comunicação Ltda (IVOCS / IPv4).** Figure 10b shows IVOCS over the 10-day window; its history divides into three phases. From January through mid-July,

IVOCS announced 4–7 IPv4 prefixes to roughly 22 peers: normal operation fully connected to the Default-Free Zone (DFZ). From mid-July through early August, IVOCS withdrew all originations, reporting zero prefixes; despite this, all 22 peer sessions stayed

active, indicating it kept DFZ connectivity through upstream transit, with peers' limits untriggered.

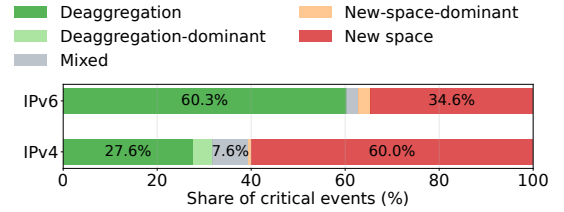
On August 12, IVOCS began originating 26 prefixes simultaneously, all within the newly assigned block 177.184.160.0/19, instantly crossing its PeeringDB limit of 16. Peer count collapsed from 22 to 2 (90.9%) within the same snapshot and remained at 2 for the rest of the year, a permanent near-total isolation. The update-level zoom resolves the jump into two steps, but the first alone clears the limit of 16 and drives essentially the entire collapse from 22 to 2; the second merely completes the announcement. Lost peers included HE (AS6939, Tier-1 rank 7), AS7195 (Major, rank 23), and AS24482 (Major, rank 89). The three-phase pattern eliminates any confound from prior instability: the collapse is a single, abrupt, and irreversible cliff edge triggered by the new address block announcement.

**Case 3: AS44901 – BelCloud (IPv6).** BelCloud's announced IPv6 prefix count was stable at 9 prefixes and roughly 788 peers until January 15, when the count surged to 128, a deaggregation of its existing 2a09:be87::/32 block into 120 more-specific /38 and /41 sub-prefixes, thereby crossing its PeeringDB limit of 100. Within the same snapshot 136 peers disconnected, leaving 652. Lost peers included AS28917 (rank 52), AS50263 (rank 58), and AS29076 (rank 93).

This case is distinctive for the operator's visible reaction: within hours of enforcement, BelCloud updated its PeeringDB entry from 100 to 300, raising the limit to accommodate its expanded announcement. This limit update is directly observable as a step in the orange (limit) line of Figure 10c, a rare end-to-end instance of the feedback loop from enforcement to corrective action.

The update-level zoom sharpens the picture and argues this is legitimate growth rather than a transient leak. The crossing localizes to ~14:30, about 1.5 hours before the 8-hour snapshot flags it, and the 120 more-specifics all originate from a single downstream customer (AS214342) announcing them within BelCloud's own 2a09:be87::/32 allocation, with the covering aggregate withdrawn as the more-specifics appear ( $9 - 1 + 120 = 128$ ). As in BCE, the peer count first ticks up at the crossing before it drops, the same visibility artifact as the 120 more-specifics reach additional collector peers. Two signatures separate this from a route leak: the added prefixes lie entirely inside BelCloud's own block rather than being third-party routes, and they persist across the rest of the window instead of being withdrawn once observed. That the operator responds by *raising* its own limit, not by withdrawing the routes, is itself consistent with intended growth: the exceedance is a false alarm of the protection mechanism, not a misconfiguration or attack. The zoom also shows the peers reconverging (from ~650 back toward ~790) over the following day once the limit is lifted.

**Case 4: AS52603 – Serviços Ltda (Supply Net / IPv6).** Supply Net advertised a single highly visible IPv6 prefix, a /32 aggregate, to roughly 20 peers for most of 2025, as shown in Figure 10d. Two less-visible /33 halves of that block had also been standing before September. Late on September 1 it deaggregated the /32 into additional /36 more-specifics, driving the visible count to 8 and crossing its PeeringDB limit of 3; peer count collapsed from 20 to 3 (85.0%), leaving Supply Net reachable only through its two transit providers. The update-level zoom exposes a subtlety the RIBs hide: the peering collapse precedes, by about ten minutes, the point at which the announced count formally exceeds the limit in our data. Supply Net reaches the RIS collectors two ways: two transit providers,



**Figure 11: Root-cause classification of 170 IPv4 and 78 IPv6 critical events: new prefixes are labeled *deaggregation* or *new space*, with mixed events split by deaggregation share.**

AS61568 and AS269646, carry its /32 aggregate, while a set of peers, Hurricane Electric (AS6939) among them, carry its two /33 halves. When Supply Net deaggregated, its peers dropped it at once, while the two providers accepted the full over-limit set and kept carrying it, so the excess reappears in our data only through those providers, about ten minutes later. We read this as the peers enforcing a tight limit as Supply Net crossed it, before the more-specifics could propagate. We cannot show it directly, because no collector receives Supply Net's announcements first-hand, not even additional vantage points at its own IXP. The ordering we record is therefore a property of where the collectors sit, not of the enforcement itself. Lost peers included HE and AS24482 (rank 89). Supply Net neither corrected its prefix count nor updated its PeeringDB entry: the high announcement persisted, and enforcement continued.

## 6.4 Root Cause Classification

The four cases cover a variety of operational patterns, from a single new block to deaggregation of existing space, and from a prompt operator response to persistent degradation. Three cases involve deaggregation of existing address space; while one involves bulk origination of a newly assigned block. To determine whether this breakdown holds at scale, we apply prefix containment checks to each of the 170 IPv4 and 78 IPv6 critical events: for each prefix present at the crossing snapshot but absent in the prior window, we test whether it is a subnet of any block already in the pre-event announcement set (*i.e.*, deaggregation) or falls entirely outside that space (*i.e.*, new space acquisition). Events where all new prefixes fall on one side are labeled *Deaggregation* or *New space*; events with both types are sub-classified by the deaggregation share into *deaggregation-dominant* ( $\geq 80\%$ ), *truly mixed* (20%–80%), and *New-space-dominant* ( $\leq 20\%$ ). We require 80% to label an event as dominant, ensuring the classification reflects a clear operational pattern rather than a marginal majority.

Figure 11 shows the full breakdown. For IPv4, pure new space dominates (60.0%), with those prefixes concentrated at /24 (63.3%), consistent with typical allocation sizes at the edge of the routing table; mixed events account for 12.4% of IPv4 cases. For IPv6, pure deaggregation dominates (60.3%), with new space acquisition accounting for a further 34.6%; consistent with address space subdivision within pre-allocated blocks, with mixed events accounting for only 5.1%. For a broader analysis that extends this classification to all impactful events and adds whether each announcement persists or quickly reverts, the reader can consult Appendix F.



**Table 6: Recommendations organized by time horizon and target audience.**

| ID    | Horizon     | Actor                 | Recommendation                                                                                                                                           |
|-------|-------------|-----------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------|
| (i)   | Near-term   | Multi-homed operators | Maintain PeeringDB entries; review annually or on significant network changes (e.g., expansions, deaggregation).                                         |
| (ii)  | Near-term   | All operators         | Set safety margins: 50–100% buffer above current announced count for peers/customers; 1.2–1.5× routing table size for transit providers.                 |
| (iii) | Near-term   | All operators         | Alert at 80–90% of configured limit; include max-prefix review in change management, especially before customer acquisitions and network expansions.     |
| (iv)  | Near-term   | IXPs                  | Build dashboards showing member prefix trajectories, current position relative to declared limits, and projected breach dates.                           |
| (v)   | Medium-term | Tools                 | Automate PeeringDB synchronization from router configuration systems; build lightweight peer visibility APIs for voluntary prefix count sharing [58].    |
| (vi)  | Long-term   | IETF / Vendors        | Revive the ORF-based draft [1]; standardize dynamic limit negotiation, staged enforcement, and pre/post-policy semantics to reduce vendor fragmentation. |

**Takeaways.** *Limit exceedances produce median peer losses of 25% (IPv4) and 33% (IPv6), but severity is driven by identity: Tier-1 and Major peers account for 18.4% of lost IPv4 sessions despite being roughly 0.1% of all ASes. Across critical events, IPv4 exceedances are predominantly new prefix space (60.0%) and IPv6 deaggregation (60.3%), routine operational growth rather than route leaks.*

## 7 Recommendations

Our findings present a largely positive view of the maximum-prefix mechanism. The large majority of ASes with a declared limit (about 83% IPv4 and 86% IPv6) never exceed it across the year. However, in the minority of cases where exceedance occurs, the consequences are severe: a moderate prefix increase can sever a meaningful fraction of an AS’s peers, including Tier-1 transit links, and isolate it from the global routing table. The recommendations below therefore target the minority for whom the declared value is stale or mismatched, and we frame them as scoped rather than absolute. Addressing these mismatches does not require waiting for protocol standardization. Meaningful improvements are available today, with longer-term gains unlocked by tooling and protocol extensions. Table 6 organizes six concrete recommendations by time horizon and target audience.

Recommendations (i)–(iii) address the root cause of the exceedances we observe: stale limits and absent monitoring. The safety margins in (ii) are grounded in our empirical violation patterns and are consistent with prior operator guidance [6, 57, 58, 67]; a 50–100% buffer would have prevented the majority of the crossing events we identify. Recommendation (iii) extends this to operational workflows. The IVOCS and Supply Net cases show that a single network event, a new block announcement or a deaggregation, can permanently cross a stale limit without the operator noticing. A max-prefix review embedded in change management would have surfaced the mismatch before enforcement. Recommendation (iv) positions IXPs as early-warning monitors. Because route servers observe prefix trajectories for all connected members simultaneously, an IXP is uniquely placed to detect approaching limit violations and notify members proactively, before enforcement forces a session teardown. Tooling (v) attacks staleness at its source: synchronizing PeeringDB directly from router configuration systems keeps the

declared limit current without manual effort, while lightweight peer-visibility APIs let networks share their live prefix count voluntarily [58]. Such sharing is especially valuable in dynamic scenarios such as on-demand Distributed Denial of Service (DDoS) scrubbing: a scrubbing provider’s prefix count spikes during mitigations and returns to a low steady state between attacks, a case where a static limit cannot capture both states. The long-term direction (vi) is in-band standardization: reviving dynamic limit negotiation so peers exchange limits automatically, alongside staged enforcement and consistent pre/post-policy semantics. The concern that signaling a negotiated limit creates an attack surface is real, but it argues for openness over obscurity. As RPKI and ASPA deployment matures, the marginal risk of limit disclosure shrinks.

Operators must also settle a fundamental design question: whether enforcement should drop the session or freeze it. Session drops risk extended outages (as in the Optus case [20]), while session freezing risks traffic loops if stale routes persist [34]. The right answer is criticality-dependent rather than universal: for most sessions a tear-down is a safe default, but for a highly critical session, keeping it up may be preferable to isolating dependent customers. We therefore favor staged enforcement, warn, then alert, then act, which several vendors (Juniper’s idle-timeout, BIRD’s warn-and-restart) already support, reserving a hard teardown for sessions where the operator judges the exceedance riskier than the disconnection.

## 8 Conclusion

We present the first large-scale empirical measurement of the BGP maximum-prefix mechanism. Community-maintained limit data in PeeringDB is missing or stale for most ASes, and even among well-maintained entries, exceedance is common. These exceedances are modest: the median count above the limit is just 5 (IPv4) and 2 (IPv6). They are driven by organic growth, such as new space or deaggregation, not routing anomalies. Yet they are treated as route leaks, with a disproportionate share of lost peers being Tier-1 and Major providers. The maximum-prefix mechanism exemplifies a broader class of infrastructure problems: critical systems sustained by informal coordination, where the gap between design intent and operational reality goes unmeasured until someone looks.

## Acknowledgments

This work was partially supported by the Pioneering AI-enhanced SEcure 6G Services framework (PAISES) project, funded by the SNS JU and the European Union's Horizon Europe research and innovation programme under Grant Agreement No 101292896. We thank our shepherd and the anonymous reviewers for their careful and constructive feedback, which substantially improved this work.

## References

- [1] Donatas Abraitis. 2024. *Maximum Prefix Outbound Route Filter for BGP*. Internet-Draft draft-abraitis-idr-maximum-prefix-orf-00. Internet Engineering Task Force. <https://datatracker.ietf.org/doc/draft-abraitis-idr-maximum-prefix-orf/>
- [2] Thomas Alfroy, Thomas Holterbach, Thomas Krenc, KC Claffy, and Cristel Pelsser. 2023. Internet Science Moonshot: Expanding BGP Data Horizons. In *Proceedings of the 22nd ACM Workshop on Hot Topics in Networks* (Cambridge, MA, USA) (*HotNets '23*). Association for Computing Machinery, New York, NY, USA, 102–108. <https://doi.org/10.1145/3626111.3628202>
- [3] Thomas Alfroy, Thomas Holterbach, Thomas Krenc, K. C. Claffy, and Cristel Pelsser. 2024. The Next Generation of BGP Data Collection Platforms. In *Proceedings of the ACM SIGCOMM 2024 Conference* (Sydney, NSW, Australia) (*ACM SIGCOMM '24*). Association for Computing Machinery, New York, NY, USA, 794–812. <https://doi.org/10.1145/3651890.3672251>
- [4] Thomas Alfroy, Thomas Holterbach, Thomas Krenc, K. C. Claffy, and Cristel Pelsser. 2024. The Next Generation of BGP Data Collection Platforms. In *Proceedings of the ACM SIGCOMM 2024 Conference* (Sydney, NSW, Australia) (*ACM SIGCOMM '24*). Association for Computing Machinery, New York, NY, USA, 794–812. <https://doi.org/10.1145/3651890.3672251>
- [5] Australian Communications and Media Authority. 2024. Optus pays \$12 million penalty for Triple Zero outage. ACMA Media Release. <https://www.acma.gov.au/articles/2024-11/optus-pays-12-million-penalty-triple-zero-outage> Accessed: April 2026.
- [6] Ignas Bagdonas. 2023. What is going on with BGP [RIPE Routing WG mailing list]. RIPE Routing Working Group Mailing List. <https://www.ripe.net/ripe/mail/archives/routing-wg/2023-June/004748.html> Accessed: 2026-04-23.
- [7] Marcelo Bagnulo, Alberto García-Martínez, Stefano Angieri, Andra Lutu, and Jinze Yang. 2022. Practicable route leak detection and protection with ASIRIA. *Comput. Netw.* 211, C (jul 2022), 12 pages. <https://doi.org/10.1016/j.comnet.2022.108966>
- [8] Renan Paredes Barreto, Leandro Márcio Bertholdo, and Pedro de Botelho Marcos. 2026. Routing Under Siege: How Traffic Engineering Decisions Facilitate Prefix Hijackings. In *International Conference on Passive and Active Network Measurement*. Springer, 44–70.
- [9] Tom Beecher. 2021. Re: Setting sensible max-prefix limits. NANOG Mailing List. <https://seclists.org/nanog/2021/Aug/182> Accessed: 2026-04-08.
- [10] BGPKIT. 2024. pybgpkit: Python bindings for the bgpkit-t-parser MRT parsing library. <https://github.com/bgpkit/pybgpkit> Accessed: April 2026.
- [11] bgp.tools. [n. d.]. How do we calculate Down/Upstreams? bgp.tools Knowledge Base. <https://bgp.tools/kb/what-is-a-upstream> Accessed: 2026-04-08.
- [12] BIRD Project. 2024. BIRD 3.0.4 User's Guide. CZ.NIC. <https://bird.nic.cz/doc/bird-3.0.4.html> Accessed: 2026-08-16.
- [13] Dale W. Carder. 2021. Re: Setting sensible max-prefix limits. NANOG Mailing List. <https://seclists.org/nanog/2021/Aug/188> Accessed: 2026-04-08.
- [14] Center for Applied Internet Data Analysis. 2025. CAIDA AS Rank. <https://asrank.caida.org/> Accessed: April 2026.
- [15] Center for Applied Internet Data Analysis. 2025. CAIDA Public PeeringDB Dataset. CAIDA Data Repository. <https://publicdata.caida.org/datasets/peeringdb/> Accessed: April 2026.
- [16] Taejoong Chung, Emile Aben, Tim Bruijnzeels, Balakrishnan Chandrasekaran, David Clark, Willem de Bruijn, Wes Hardaker, Doris Mao, Danny McPherson, and Arthur Berger. 2019. RPKI is Coming of Age: A Longitudinal Study of RPKI Deployment and Invalid Route Origins. In *Proceedings of the 2019 ACM Internet Measurement Conference*. 406–419.
- [17] Taejoong Chung, Yuze Li, and Weitong Li. 2025. ImpROV: measurement and practical mitigation of collateral damage in RPKI route origin validation. In *Proceedings of the 34th USENIX Conference on Security Symposium* (Seattle, WA, USA) (*SEC '25*). USENIX Association, USA, Article 187, 17 pages.
- [18] Cisco. 2024. BGP Maximum-Prefix on IOS XE. Cisco IOS XE 17.x IP Routing Configuration Guide. [https://www.cisco.com/c/en/us/td/docs/routers/ios/config/17-x/ip-routing/b-ip-routing/m\\_irg-max-prefix.html](https://www.cisco.com/c/en/us/td/docs/routers/ios/config/17-x/ip-routing/b-ip-routing/m_irg-max-prefix.html) Accessed: 2026-08-16.
- [19] Cisco Systems. 2023. Configure the BGP Maximum-Prefix Feature. Cisco Technical Documentation, Document ID 25160. <https://www.cisco.com/c/en/us/support/docs/ip/border-gateway-protocol-bgp/25160-bgp-maximum-prefix.html> Accessed: 2026-04-08.
- [20] Australian Communications and Media Authority. 2023. Investigation Report. <https://www.acma.gov.au/sites/default/files/2024-11/Investigation%20report%20-%20Optus%20outage%201Nov23%20%28redacted%29.pdf> Accessed: 2026-01-15.
- [21] David de Andres Hernandez, Yasin Alhamwy, Daniel Wagner, Maximilian Stephan, Matthias Wichtlhuber, and Oliver Hohlfeld. 2026. LARS: Keeping Local Traffic Local with Latency-Aware Route Servers at IXPs. In *Proceedings of the ACM SIGCOMM 2026 Conference* (Denver, CO, USA) (*SIGCOMM '26*). Association for Computing Machinery, New York, NY, USA, 1009–1023. <https://doi.org/10.1145/3789240.3829211>
- [22] Jerome Durand, Ivan Pepelnjak, and Gert Doering. 2015. *BGP Operations and Security*. RFC 7454. Internet Engineering Task Force (IETF). <https://doi.org/10.17487/RFC7454> BCP 194.
- [23] Mauro Farina, Martino Trevisan, and Alberto Bartoli. 2025. Poster: Evaluating CAIDA's Inferred AS Relationships Through ASPAs. In *Proceedings of the 21st International Conference on Emerging Networking Experiments and Technologies (CoNEXT '25)*. 7–8. <https://doi.org/10.1145/3765515.3771747>
- [24] Romain Fontugne and Emile Aben. 2019. BGP Zombies. RIPE Labs. [https://labs.ripe.net/author/romain\\_fontugne/bgp-zombies/](https://labs.ripe.net/author/romain_fontugne/bgp-zombies/) Accessed: April 2026.
- [25] Justin Furunes, Cameron Morris, Reynaldo Morillo, Arvind Kasiliya, Bing Wang, and Amir Herzberg. 2025. Securing BGP ASAP: ASPA and other Post-ROV Defenses. (2025).
- [26] Julien Gamba, Romain Fontugne, and Emile Aben. 2017. BGP Table Fragmentation: What & Who? RIPE Labs. [https://labs.ripe.net/author/julien\\_gamba/bgp-table-fragmentation-what-who/](https://labs.ripe.net/author/julien_gamba/bgp-table-fragmentation-what-who/) Accessed: April 2026.
- [27] Lixin Gao. 2001. On Inferring Autonomous System Relationships in the Internet. *IEEE/ACM Transactions on Networking* 9, 6 (2001), 733–745.
- [28] Vasileios Giotsas, Christoph Dietzel, Georgios Smaragdakis, Anja Feldmann, Arthur W. Berger, and Emile Aben. 2017. Detecting Peering Infrastructure Outages in the Wild. In *Proceedings of the Conference of the ACM Special Interest Group on Data Communication (SIGCOMM '17)*. 446–459. <https://doi.org/10.1145/3098822.3098855>
- [29] Deepak Gouda, Alberto Dainotti, and Cecilia Testart. 2025. Prefix2Org: Mapping BGP Prefixes to Organizations. In *Proceedings of the 2025 ACM Internet Measurement Conference (USA) (IMC '25)*. Association for Computing Machinery, New York, NY, USA, 397–414. <https://doi.org/10.1145/3730567.3764485>
- [30] Deepak Gouda, Romain Fontugne, and Cecilia Testart. 2025. ru-RPKI-ready: the Road Left to Full ROA Adoption. In *Proceedings of the 2025 ACM Internet Measurement Conference*. 415–429.
- [31] Government of the Grand Duchy of Luxembourg. 2025. Cyberattack against POST Luxembourg: convening of the crisis unit. High Commission for National Protection (HCPN). [https://meco.gouvernement.lu/en/actualites/gouvernement2024+en+actualites+toutes\\_actualites+communiqués+2025+07-juillet+25-celulle-crise-post.html](https://meco.gouvernement.lu/en/actualites/gouvernement2024+en+actualites+toutes_actualites+communiqués+2025+07-juillet+25-celulle-crise-post.html) Accessed: 2026-08-23.
- [32] Aric A. Hagberg, Daniel A. Schult, and Pieter J. Swart. 2008. Exploring Network Structure, Dynamics, and Function using NetworkX. *Proceedings of the 7th Python in Science Conference (SciPy)* (2008), 11–15.
- [33] Bryton Herdes. 2026. A closer look at a BGP anomaly in Venezuela. Cloudflare Blog. <https://blog.cloudflare.com/bgp-route-leak-venezuela/> Accessed: 2026-04-28.
- [34] Geoff Huston. 2025. Notes from RIPE 91. APNIC Blog. <https://blog.apnic.net/2025/10/31/notes-from-ripe-91/> Accessed: April 2026.
- [35] IXP Manager contributors. 2024. IXP Manager Documentation: Customers. <https://docs.ixpmanager.org/6.4/usage/customers/> Accessed: April 2026.
- [36] IXP Manager contributors. 2024. IXP Manager: Open-source software for running Internet Exchange Points. <https://www.ixpmanager.org/> Accessed: April 2026.
- [37] Juniper Networks. 2024. maximum-prefixes: Junos OS CLI Reference. Juniper Networks documentation. <https://www.juniper.net/documentation/us/en/software/junos/cli-reference/topics/ref/statement/maximum-prefixes-edit-routing-options.html> Accessed: 2026-08-16.
- [38] Jon Lewis. 2021. Re: Setting sensible max-prefix limits. NANOG Mailing List. <https://seclists.org/nanog/2021/Aug/197> Accessed: 2026-04-08.
- [39] Aemen Lodhi, Natalie Larson, Amogh Dhamdhere, Constantine Dovrolis, and Kc Claffy. 2014. Using peeringDB to understand the peering ecosystem. *ACM SIGCOMM Computer Communication Review* 44, 2 (April 2014), 21–27.
- [40] Andra Lutu, Marcelo Bagnulo, and Olaf Maennel. 2013. The BGP Visibility Scanner. In *2013 IEEE Conference on Computer Communications Workshops (INFOCOM WKSHPS)*. 115–120. <https://doi.org/10.1109/INFOCOMW.2013.6562877>
- [41] Alexander Martin. 2025. Luxembourg probes reported attack on Huawei tech that caused nationwide telecoms outage. The Record from Recorded Future News. <https://therecord.media/luxembourg-telecom-outage-reported-cyberattack-huawei-tech> Accessed: 2026-08-23.
- [42] Orlando Eduardo Martínez-Durive. 2026. What Happens When You Overshare? A Look into the BGP Maximum-Prefix Feature. RIPE SEE 14 Meeting. <https://pretalx.ripe.net/see-14/talk/3PHMTW/>
- [43] Orlando Eduardo Martínez-Durive and Antonios Chariton. 2025. What Happens When You Overshare? A Look into the BGP Maximum-Prefix Feature. RIPE 91 Meeting, Routing Session, Amsterdam. <https://ripe91.ripe.net/programme/>

- meeting-plan/sessions/30/JDHZfZ/ Slides: [https://pretalx.ripe.net/media/ripe91/submissions/JDHZfZ/resources/slides\\_UDSaivF.pdf](https://pretalx.ripe.net/media/ripe91/submissions/JDHZfZ/resources/slides_UDSaivF.pdf).
- [44] Declan McCullagh. 2008. How Pakistan Knocked YouTube Offline (and How to Make Sure It Never Happens Again). CNET. <https://www.cnet.com/culture/how-pakistan-knocked-youtube-offline-and-how-to-make-sure-it-never-happens-again/> Accessed: 2026-04-08.
  - [45] National Security Agency. 2018. *A Guide to Border Gateway Protocol (BGP) Best Practices*. Cybersecurity Report. NSA Cybersecurity. <https://media.defense.gov/2019/Jul/16/2002158110/-1/-1/0/CTR-GUIDE-TO-BORDER-GATEWAY-PROTOCOL-BEST-PRACTICES.PDF> Accessed: 2026-08-11.
  - [46] Marcin Nawrocki, Jeremias Blendin, Christoph Dietzel, Thomas C. Schmidt, and Matthias Wählisch. 2019. Down the Black Hole: Dismantling Operational Practices of BGP Blackholing at IXPs. In *Proceedings of the Internet Measurement Conference* (Amsterdam, Netherlands) (IMC '19). Association for Computing Machinery, New York, NY, USA, 435–448. <https://doi.org/10.1145/3355369.3355593>
  - [47] Nokia. 2023. BGP Prefix Limit per Address Family. Nokia SR OS 23.7.R2 Classic CLI documentation. <https://documentation.nokia.com/acg/23-7-2/books/classic-cli-part-i/c147-bgp-pfx-limit.html> Accessed: 2026-08-16.
  - [48] OpenBSD. 2022. `bgpd.conf(5)`: Border Gateway Protocol daemon configuration file. OpenBSD 7.2 manual pages. <https://man.openbsd.org/OpenBSD-7.2/bgpd.conf.5> Accessed: 2026-08-16.
  - [49] Peering Manager contributors. 2024. Peering Manager: An open-source BGP peering management system. <https://peering-manager.net/> Accessed: April 2026.
  - [50] PeeringDB. 2026. Getting Started as a Network Operator. <https://docs.peeringdb.com/howto/get-started-operator/>. “[T]he maximum number of IPv4 and/or IPv6 prefixes [that] peers [should] expect to see advertised by your network.” Accessed: 2026-08-22.
  - [51] PeeringDB. 2026. PeeringDB. Online database, <https://www.peeringdb.com>. Accessed: 2026-04-08.
  - [52] PeeringDB. 2026. PeeringDB API Documentation. <https://www.peeringdb.com/apidocs/>. `info_prefixes4`: “Recommended maximum number of IPv4 routes/prefixes to be configured on peering sessions for this ASN.” Accessed: 2026-08-22.
  - [53] PeeringDB. 2026. PeeringDB Network Record for AS766 (RedIRIS): IPv4/IPv6 Prefixes Field. <https://www.peeringdb.com/net/2813>. Field help text: “Recommended maximum number of routes/prefixes to be configured on peering sessions for this ASN.” Accessed: 2026-08-22.
  - [54] Louis Poinssignon. 2018. BGP Leaks and Cryptocurrencies. Cloudflare Blog. <https://blog.cloudflare.com/bgp-leaks-and-crypto-currencies/> Accessed: April 2026.
  - [55] Lars Prehn. 2021. Re: Setting sensible max-prefix limits. NANOG Mailing List. <https://seclists.org/nanog/2021/Aug/179> Accessed: 2026-04-08.
  - [56] Lars Prehn. 2022. routing-wg Detecting, mitigating, and preventing distributed large-scale prefix de-aggregation attacks. RIPE Routing Working Group Mailing List. <https://www.ripe.net/ripe/mail/archives/routing-wg/2022-October/004643.html> Accessed: 2026-04-17.
  - [57] Lars Prehn et al. 2021. Setting sensible max-prefix limits [NANOG mailing list thread]. NANOG Mailing List. <https://seclists.org/nanog/2021/Aug/171> Accessed: 2026-04-08.
  - [58] Lars Prehn, Paweł Foremski, and Oliver Gasser. 2024. Kirin: Hitting the Internet with Distributed BGP Announcements. In *Proceedings of the 19th ACM Asia Conference on Computer and Communications Security*. 19–34.
  - [59] Vasile Radoaca, Mo Miri, Luyuan Fang, Susan Hares, and Srikanth Chavali. 2004. *Peer Prefix Limits Exchange in BGP*. Internet-Draft draft-chavali-bgp-prefixlimit-02. Internet Engineering Task Force. <https://datatracker.ietf.org/doc/draft-chavali-bgp-prefixlimit/>
  - [60] Y. Rekhter, T. Li, and S. Hares. 2006. *A Border Gateway Protocol 4 (BGP-4)*. RFC 4271. RFC Editor. <https://www.rfc-editor.org/rfc/rfc4271>
  - [61] RIPE NCC. 2008. YouTube Hijacking: A RIPE NCC RIS Case Study. RIPE NCC. <https://www.ripe.net/about-us/news/youtube-hijacking-a-ripe-ncc-ris-case-study/> Accessed: 2026-04-28.
  - [62] RIPE NCC. 2024. Autonomous System Provider Authorization (ASPA). RIPE NCC. <https://www.ripe.net/manage-ips-and-asns/resource-management/rpki/aspa/> Accessed: 2026-04-28.
  - [63] RIPE NCC. 2024. Route Collection Raw Data: MRT Files. RIPE Network Coordination Centre. <https://ris.ripe.net/docs/mrt/> Last updated March 29, 2024. Accessed: 2026-04-08.
  - [64] Loqman Salamatian, Frédéric Douzet, Kavé Salamatian, and Kévin Limonier. 2021. The Geopolitics Behind the Routes Data Travel: A Case Study of Iran. *Journal of Cybersecurity* 7, 1 (2021), tyab018.
  - [65] Seirdy and contributors. 2024. `bgpq4`: BGP filter generation tool. <https://github.com/bgp/bgpq4> Accessed: April 2026.
  - [66] Job Snijders. 2014. PeeringDB Accuracy. NANOG 58, New Orleans. <https://archive.nanog.org/sites/default/files/wed.general.peeringdb.accuracy.snijders.14.pdf> Accessed: April 2026.
  - [67] Job Snijders. 2018. Robust Routing Policy Architecture. RIPE 77 Meeting, Amsterdam. [https://ripe77.ripe.net/presentations/59-RIPE77\\_Snijders\\_Routing\\_Policy\\_Architecture.pdf](https://ripe77.ripe.net/presentations/59-RIPE77_Snijders_Routing_Policy_Architecture.pdf) Accessed: 2026-04-08.
  - [68] Job Snijders and Melchior Aelmans. 2019. *BGP Maximum Prefix Limits*. Internet-Draft draft-sa-grow-maxprefix-02. Internet Engineering Task Force. <https://datatracker.ietf.org/doc/draft-sa-grow-maxprefix/> Expired Internet-Draft; migrated to IDR WG as draft-sa-idr-maxprefix.
  - [69] Job Snijders, Melchior Aelmans, and Massimiliano Stucchi. 2023. *Inbound BGP Maximum Prefix Limits*. Internet-Draft draft-sas-idr-maxprefix-inbound-05. Internet Engineering Task Force. <https://datatracker.ietf.org/doc/draft-sas-idr-maxprefix-inbound/> Expired Internet-Draft.
  - [70] Tom Strickx. 2019. How Verizon and a BGP Optimizer Knocked Large Parts of the Internet Offline Today. Cloudflare Blog. <https://blog.cloudflare.com/how-verizon-and-a-bgp-optimizer-knocked-large-parts-of-the-internet-offline-today/> Accessed: April 2026.
  - [71] Pier Carlo Stucchi. 2024. ARouteServer: A tool to automatically build route server configurations. <https://arouteserver.readthedocs.io/> Accessed: April 2026.
  - [72] Cecilia Testart, Philipp Richter, Alistair King, Alberto Dainotti, and David Clark. 2019. Profiling BGP Serial Hijackers: Capturing Persistent Misbehavior in the Global Routing Table. In *Proceedings of the Internet Measurement Conference* (Amsterdam, Netherlands) (IMC '19). Association for Computing Machinery, New York, NY, USA, 420–434. <https://doi.org/10.1145/3355369.3355581>
  - [73] Iliana Xygiokou, Antonis Chariton, Xenofontas Dimitropoulos, and Alberto Dainotti. 2025. A First Look into Long-lived BGP Zombies. In *Proceedings of the 2025 ACM Internet Measurement Conference*. 826–834.
  - [74] Mingwei Zhang and Bryton Herdes. 2026. ASPA: making Internet routing more secure. <https://blog.cloudflare.com/aspa-secure-internet/>.

## A Ethics

This study relies exclusively on publicly available data: RIPE RIS routing tables, PeeringDB API snapshots, and CAIDA AS Rank data. We do not collect or analyze traffic content, intercept communications, or access private network configurations or operator logs. The study requires no IRB review under standard passive-measurement exemptions. ASes named in case studies are identified by their public ASN and organization name from PeeringDB and Regional Internet Registry (RIR) databases; we do not anonymize them, since our data are openly accessible, our methodology reproducible, and the documented behaviors arise from voluntary participation in a public coordination system. We frame all findings as systemic measurement results rather than attribution of blame to individual operators. We shared preliminary results with the network operator community prior to submission [42, 43] to gather feedback and ensure our methodology reflects operational realities.

## B Data Availability

Our code, processed datasets, and analysis notebooks (per-AS exceedance records, crossing event logs, and the RIS∩PeeringDB datasets) are publicly available under an open-source license<sup>2</sup>. Raw RIPE RIS RIBs (3.4 TB) and case-study updates (144 GB) are not redistributed but are publicly accessible from RIPE [63]; the repository includes scripts to reproduce the pipeline.

## C Robustness of the Impactful Threshold

Our impactful definition (Section 6.1) flags an exceedance event when its coincident peer drop exceeds the AS’s own 95th-percentile non-exceedance churn. Table 7 shows how the impactful count varies with the choice of percentile: moving the floor from the 90th to the 99th percentile changes the count monotonically, as expected, but leaves the qualitative picture unchanged, with the 95th percentile (our choice) sitting between the permissive and conservative ends.

<sup>2</sup>Available at <https://github.com/nds-group/max-prefix-limit-bgp>

**Table 7: Impactful events under per-AS churn floors.**

| IP   | Crossings | @90th | @95th | @99th |
|------|-----------|-------|-------|-------|
| IPv4 | 4,314     | 391   | 303   | 131   |
| IPv6 | 1,501     | 150   | 132   | 66    |

## D Tier-1 ASes

Table 8 lists the 16 Tier-1 ASes used to identify paths reaching the global routing core, following the definition in [11], annotated with their CAIDA AS Rank [14] (from the `asns.jsonl` dataset). An AS is considered Tier-1 if it has no upstream provider, *i.e.*, it exchanges traffic exclusively through peering and does not pay for transit.

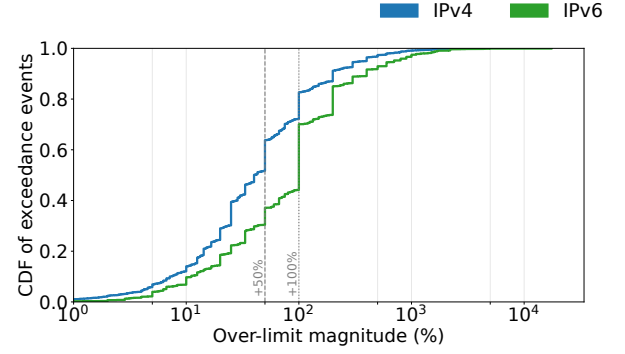
We acknowledge that this classification is generally accepted but not universally agreed upon: Hurricane Electric (AS6939), for instance, is widely regarded as Tier-1 for IPv6, where it maintains settlement-free peering with all major providers, but its IPv4 status is debated due to a long-standing peering dispute with Cogent (AS174) [11]. We follow the `bgp.tools` classification and include HE in our list, consistent with common operator practice.

**Table 8: Tier-1 ASes used in this study.**

| Network Name                     | ASN   | AS Rank |
|----------------------------------|-------|---------|
| Cogent Communications            | 174   | 4       |
| Verizon Enterprise (UUNET)       | 701   | 35      |
| Arelion (formerly Telia Carrier) | 1299  | 2       |
| NTT Communications (Verio)       | 2914  | 5       |
| GTT Communications               | 3257  | 3       |
| Deutsche Telekom Global Carrier  | 3320  | 18      |
| Lumen Technologies (Level 3)     | 3356  | 1       |
| PCCW Global                      | 3491  | 10      |
| Orange S.A.                      | 5511  | 12      |
| TATA Communications              | 6453  | 6       |
| Zayo Group                       | 6461  | 8       |
| Telecom Italia Sparkle/Seabone   | 6762  | 9       |
| Liberty Global                   | 6830  | 31      |
| Hurricane Electric               | 6939  | 7       |
| AT&T                             | 7018  | 32      |
| Telxius International            | 12956 | 13      |

## E Exceedance Magnitude

A natural objection to flagging exceedance at the declared limit is that operators may tolerate some overshoot: community discussions on NANOG [57] and RIPE [67] point to a  $1.5\times$  or  $2\times$  multiple of the declared limit as a more common enforcement threshold. We therefore examine how far over the limit ASes go at the moment they cross it, treating each below-to-above transition as one independent event pooled across all ASes, so a persistently-exceeding AS contributes one point per new exceedance rather than one per 8-hour snapshot (this also avoids treating the time-varying exceedance value as static). Figure 12 shows the resulting distribution. The medians are  $+40\%$  (IPv4) and  $+100\%$  (IPv6); a  $2\times$  threshold would discard the majority of exceedances in both families, and even a  $1.5\times$  threshold a large share (Section 5). The pronounced relative



**Figure 12: CDF of the over-limit magnitude at each exceedance, across all below-to-above transitions. A stricter  $1.5\times/2\times$  threshold would miss a large share of exceedances, motivating our choice to flag at the declared limit.**

steps at  $+100\%$ ,  $+200\%$ , and  $+300\%$  hide how small the absolute values are: 45% of IPv6 exceedances occur against a limit of a single prefix, so one extra prefix already registers as a 100% exceedance.

## F Event Classification: Cause and Persistence

In Section 6 we classify the critical events by cause. Here we broaden to *all* impactful events (303 IPv4, 132 IPv6) and add a second axis, whether the exceedance *persistent* or *reverts*, to isolate the cases RIB-only data cannot distinguish from traffic engineering or DDoS scrubbing. **Cause** labels each event by prefix containment (as in Section 6): *Deaggregation* ( $\geq 80\%$  of new prefixes are more-specifics of existing space), *New space* ( $\leq 20\%$ ), or *Mixed*. **Persistence** marks an event *persistent* if the announced count stays above its pre-crossing baseline for most of the following week, else *reverting*; we use the baseline rather than the limit, since an operator can raise its limit while keeping the prefixes (BelCloud, Figure 10c). Only the *deaggregation-that-reverts* cell (Table 9, shaded) is ambiguous under RIB-only data: a quickly-withdrawn more-specific is equally consistent with enforcement, traffic engineering, or scrubbing, and is where BCE sits. That cell holds just 12% of IPv4 and 35% of IPv6 impactful events; the rest are new-space growth or *persistent* deaggregation, consistent with genuine enforcement rather than a transient anomaly.

**Table 9: Impactful events by cause and persistence. Only the shaded deaggregation-that-reverts cell is ambiguous under RIB-only data (where BCE sits).**

| Cause         | IPv4    |          | IPv6    |          |
|---------------|---------|----------|---------|----------|
|               | Reverts | Persists | Reverts | Persists |
| Deaggregation | 37      | 47       | 46      | 21       |
| New space     | 97      | 99       | 43      | 20       |
| Mixed         | 11      | 12       | 1       | 1        |