

metAScritic: Reframing AS-Level Topology Discovery as a Recommendation System

Loqman Salamatian
Columbia University
New York, United States

Kevin Vermeulen
LAAS-CNRS
Toulouse, France

Italo Cunha
Universidade Federal de Minas Gerais
Belo Horizonte, Brazil

Vasilis Giotsas
Cloudflare
London, United Kingdom

Ethan Katz-Bassett
Columbia University
New York, United States

Abstract

Despite prior efforts, the vast majority of the AS-level topology of the Internet remains hidden from BGP and traceroute vantage points. In this work, we introduce metAScritic, a novel system inspired by recommender system literature, designed to infer interconnections within a given metro. metAScritic uses the intuition that the connectivity matrix at a given metro is a low-rank system, since ASes employ similar peering strategies according to their infrastructures, traffic profiles, and business models. This approach allows metAScritic to accurately reconstruct the complete peering connectivity by measuring a strategic subset of interconnections that capture ASes' underlying peering strategies. We evaluate metAScritic's performance across six large metropolitan areas, achieving an average F-score of 0.88 on various validation datasets, including ground truth. metAScritic measures more than 86K edges and infers more than 368K edges, compared to the 13K edges observed for this subset of ASes in public BGP feeds – an increase of 24× what is currently seen. We study the impact of our inferred links on Internet properties, illustrating the extent of the Internet's flattening and demonstrating our ability to better predict the impact of route leaks and prefix hijacks, compared to relying only on the existing public view.

CCS Concepts

• **Networks** → **Public Internet**; *Network economics*; **Network measurement**.

Keywords

Active Internet Measurements; AS Topology; Peering Links; Matrix Completion

ACM Reference Format:

Loqman Salamatian, Kevin Vermeulen, Italo Cunha, Vasilis Giotsas, and Ethan Katz-Bassett. 2024. metAScritic: Reframing AS-Level Topology Discovery as a Recommendation System. In *Proceedings of the 2024 ACM Internet Measurement Conference (IMC '24)*, November 4–6, 2024, Madrid, Spain. ACM, New York, NY, USA, 28 pages. <https://doi.org/10.1145/3646547.3688429>

Publication rights licensed to ACM. ACM acknowledges that this contribution was authored or co-authored by an employee, contractor or affiliate of a national government. As such, the Government retains a nonexclusive, royalty-free right to publish or reproduce this article, or to allow others to do so, for Government purposes only. Request permissions from owner/author(s).

IMC '24, November 4–6, 2024, Madrid, Spain

© 2024 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 979-8-4007-0592-2/24/11

<https://doi.org/10.1145/3646547.3688429>

1 Introduction

The Internet comprises tens of thousands of autonomous systems (ASes), with their interconnections determined by a complex interplay of technical, geographic, commercial, political, and even historical factors. An accurate and comprehensive view of the resulting topology would greatly benefit operators, researchers and policy-makers by (i) enabling the characterization of Internet properties such as resilience at a local and global scale, (ii) building more faithful models of the Internet to improve the evaluation of novel solutions, (iii) informing troubleshooting and guiding operations, and (iv) providing the basis for understanding and predicting the effects of different stakeholders' investments (e.g., state and regional actors) on the network's infrastructure.

Thus far, the research community has primarily depended on Border Gateway Protocol (BGP) feeds and traceroute measurements to create maps of the Internet's topology for various purposes, such as Internet modeling [46], detecting security threats [143], mapping geopolitical tussles [134], and quantifying transit influence [57].

Most of the Internet is invisible from available vantage points. Despite its criticality, the Internet AS topology is mostly hidden from existing measurements. ASes typically only make their peering links available to themselves and their customers [58]. Hence, the peering links are only visible to vantage points hosted inside one of the peer ASes or their customer ASes [118], and no such strategically placed vantage points are available in most ASes. The exact number of missed interconnections is unknown, but previous studies indicate that the majority of peering interconnections are invisible. Ager *et al.* found that the number of existing peering links at a single Internet eXchange Point (IXP) exceeded the total number of visible peering links on the entire Internet according to publicly available BGP monitors [7]. A 2016 survey by Packet Clearing House (PCH) found nearly 2 million peering interconnections, almost 20× what was observed by the publicly available monitors [157]. These invisible links account for a substantial portion of the total Internet traffic, and so their detection is vital for developing a comprehensive understanding of the Internet's topology [15, 87].

The quantity of invisible links has likely increased because of Internet flattening, a well-documented phenomenon of increased peering between ASes that shortcuts the upper (more visible) tiers of the Internet's traditional hierarchy [2, 5, 15, 22, 37, 64]. Since ASes lower in the Internet's hierarchy and closer to its edge have (by definition) fewer customers or no customers at all, their peering links are invisible to most of the Internet [118], and capturing

them in aggregate would require vantage points spread across an unrealistically large number of edge ASes. Assuming an optimistic model where the presence of a single vantage point¹ in any AS could measure the peering links of all its upstreams, the current set of publicly available vantage points suffice to capture the full connectivity of only 6% of the ASes announced on the Internet in September 2023. Consequently, a significant portion of the topology remains elusive, despite decades of operators and researchers making a case for increased visibility [36, 62, 72, 123, 151, 155, 156].

Inferential approaches are necessary. Because decades of effort have only brought the community dwindling visibility, we argue that we must stop relying only on *measurements* to yield our best understanding of the Internet’s topology—instead, we must employ *inferential* approaches to shed light on invisible portions of the topology. To correctly translate insights from the visible segment of the topology to the hidden one, we must consider the driving forces behind AS interconnectivity. The peering strategy of an AS relies on a complex set of factors. While we can figure out certain factors, including geographic ones (e.g., presence in shared locations), economic ones (e.g., compatibility between the business objectives and market approaches), or political ones (e.g., diplomatic tension²), others remain beyond our grasp, hidden in the evolving strategies that are known only to insiders (e.g., two network operators disagreeing on an interconnection³).

Although we cannot directly know these factors, our intuition is that ASes’ strategies are based in part on known features *and* are reflected in visible links, which are the partial outcomes of their strategies. These features and known links provide information about the likelihood of a link existing between two ASes. For example, ASes that peer with Google are more likely to also peer with Amazon than ASes that do not peer with Google (Fig. 1), as they made a decision to connect to a large cloud/content provider and so may connect to other cloud/content providers for similar reasons. Additionally, if an AS that peers with Google is *also* known to host many end-users and have predominantly inbound traffic, it increases our confidence that Amazon, which serves a lot of outbound traffic to end-users, may peer with it. We validate these intuitions in the context of cloud providers, whose peering we measure via traceroutes out from the cloud (§2). The challenge lies in combining topological knowledge obtained from visible links with meta-information about ASes to capture factors driving peering strategies and infer which invisible links are likely to exist.

Our solution - metAScritic. We introduce a novel system called Metro AS Critic, shortened to metAScritic,⁴ which we designed to measure and predict interconnections between ASes within a metropolitan area (metro). Since connectivity in a metro is influenced by factors that make two ASes mutually appealing, we frame

the problem of discovering invisible interconnections as a recommendation problem, where the goal is to determine likely peering partners based on their properties. metAScritic combines knowledge about AS features and existing connectivity from measurements to predict which invisible interconnections exist (§3). To be clear, metAScritic does not infer the business relationship between ASes (e.g., customer-to-provider or peer-to-peer), nor does it provide peering recommendations for an AS. Rather, it uses a recommender system architecture to deduce interconnections likely to exist.

To mitigate the biases in the public measurements that undermine the efficacy of the recommender system, metAScritic conducts targeted measurements to uncover links of ASes with few known interconnections to capture their peering strategies. By adopting this method, metAScritic does not just perform measurements; it does so with the intent of aiding the inferential process.

The paper addresses the challenge of validation in the presence of limited ground truth by creating a validation dataset that combines multiple datasets (§4). The results show that metAScritic has an average precision of 0.89 and a recall of 0.85 on all validation data considered and for different training and testing splits.

We introduce two frameworks for using metAScritic. The first progressively adds links from the highest to lower confidence scores, allowing the construction of extended topologies with inferred links at a target confidence level, while the second estimates network properties such as existing paths or estimated catchments as random variables, derived from the likelihood of each inferred link’s existence (§5). We use the links derived from metAScritic in three case studies: (i) we analyze the measured and inferred links, (ii) quantify the impact of the additional links on hijack predictions, and (iii) examine their effect on Internet flattening (§6).

2 Overview of approach

Confronted with the reality that a significant part of the Internet’s topology is beyond the reach of current measurement techniques [77, 118], we develop an inferential method. Inferential approaches introduce a trade-off: they can significantly expand our visibility of the topology, at the cost of adding uncertainty. For the inferences to be useful despite their uncertainty, our approach quantifies the uncertainty associated with each inferred link. This quantification lets us adjust the trade-off between false positives (selecting only highly likely links) and false negatives (including lower-likelihood links) depending on the use case or account for this uncertainty when applying our inferred topology in practical scenarios. Such a compromise in precision is worthwhile when it enables a deeper understanding of the Internet’s structure. For instance, inferential techniques based on incomplete topology measurements are state of the art for IP-to-AS mapping [101] and for characterizing AS relationships [84]. We aim to close the gaps in network topology by adopting a more flexible definition of topology discovery and enable research that has been hindered by incomplete data.

To further elucidate the core principles underpinning metAScritic, we draw an analogy to a Tinder-like matchmaking system [33]. Just as Tinder uses a variety of factors—such as shared interests and geographic proximity—to recommend matches, metAScritic utilizes business strategies, traffic patterns, and other network characteristics to predict the likelihood of a peering link. In this analogy,

¹Here we consider a vantage point to be either a BGP monitor in RouteViews [107] or RIPE RIS [4] or a probe from RIPE Atlas [130] or Ark [24]

²As an illustrative example, the London Internet Exchange (LINX) has disconnected Russian telecoms companies Megafon and Rostelecom following the Russian invasion of Ukraine [108].

³For instance the IPv6 peering disputes that have led Hurricane Electric and Cogent to have neither a peering nor a transit interconnection [89].

⁴“Metro” denotes the level of granularity at which we conduct our recommendations, while “metAScritic” alludes to the website metacritic.com, which compiles reviews for films, television shows, music albums, and video games, subsequently, settings in which recommender systems were developed that we adapt to our problem.

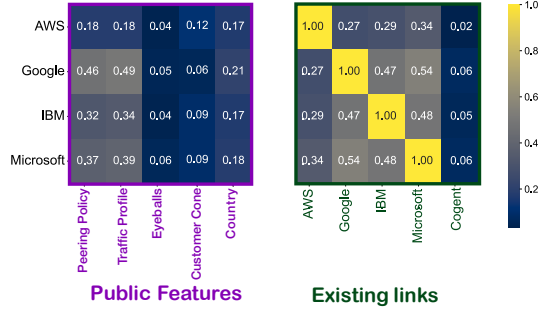


Figure 1: Correlation matrices between the peering relationships of AWS, Google, IBM, Microsoft and (left) five key publicly available features: Peering Policy (e.g., Open vs Selective), Traffic Profile (e.g., Inbound vs Outbound), Number of Eyeballs, Size of the Customer Cone, and Country of Registration; and (right) the existence of other peering links with other networks: AWS, Google, IBM, Microsoft and Cogent. Peering Policies and Traffic Profiles provide the most valuable information for predicting connections with most Cloud Providers. Additionally, identifying a peering link with another cloud provider is more informative than identifying a connection with Cogent, a Tier 1.

ASes are like users, an AS’s peering strategy resembles a user’s preferences, and the establishment of a peering link is akin to a match. Our confidence in these predictions ranges from -1 (indicating no link likely exists) to 1 (indicating a high likelihood of a link), reflecting the certainty of our inferences.

Identifying the right granularity for peering link inference.

The vast majority of peering links are established at peering facilities. However, metAScritic focuses on inferring AS links at the metro level rather than at the facility level, as the metro level strikes a practical balance between accuracy and feasibility: it captures most of the possible interdomain connections while avoiding the complexities of pinpointing interconnections to a specific peering facility, which are more challenging to map with existing methods [70], whereas mapping traceroutes and interconnections to metros is well-supported by existing techniques [28, 48, 95, 137]. The metro granularity is valuable for applications like latency prediction [97], Point of Presence (PoP) mapping [149], network troubleshooting [150], and outage detection [69]. Our metro-level insights can also be aggregated for AS-level analysis, which has been used for monitoring BGP hijacks [143], making predictions about anycast catchments [102, 140–142, 148, 158], improving overlay networks [76], and interpreting geopolitical tussles on the Internet [57, 134]. Most large interconnection facilities are concentrated in a small number of metros. According to iGDB [11], only 226 metros host more than 50 ASes worldwide, and the top 10 metros account for 21% of all (Metro, AS) pairs. This concentration suggests that targeting a limited number of metros will suffice to allow metAScritic to capture a significant portion of missing links in the Internet’s topology. While our framework is optimized for metro granularity, we believe it can be adapted to other granularities in the future.

Understanding the motivations of peering decisions with data. Peering decisions among ASes stem from multiple underlying factors. These factors will be reflected in the AS’s established peering interconnections with other ASes. Therefore, the visible peers of an AS provide information about what may be its invisible peers. Some AS attributes that reflect or influence its peering practices are publicly available. For instance, ASes can disclose whether their traffic is predominantly inbound or outbound and their criteria for peering—whether open to any requesting AS or more selective [122]. ASes that host many users typically have heavy inbound traffic [23] and are likely to connect with cloud/content providers that host most of the user-facing services on the Internet. ASes with selective peering policies seek partnerships that offer manifest value, finding that establishing peering relationships with cloud providers yields benefits in view of the considerable volume of traffic these providers generate. Additionally, an AS peering with one major cloud provider is likely to connect with others, driven by similar customer demands and performance needs across many hypergiants’ services.

We corroborated these insights using a dataset of measured (and validated) peering links from cloud providers [15], and, for a comparison with other large networks that have different practices, supplemented this data with interdomain links from Tier 1 ASes, according to AS Rank [25].⁵ Figure 1 depicts the correlation between various AS properties (columns) and whether the AS peers with each of four cloud providers (rows). The properties on the left are various AS features (Peering Policy, Traffic Profile, Eyeballs, Customer Cone, Country) and, on the right, the presence of peering links with other cloud providers and Cogent (AWS, Google, IBM, Microsoft, Cogent) where the color indicates whether the variable is binary, categorical, or continuous. For categorical features, we use the correlation ratio [139], while for continuous and binary features, we compute Pearson’s correlation coefficient. The left side of the figure highlights the correlations between an AS’s Peering Policy (e.g., Open, Selective) and Traffic Profile (e.g., Inbound, Outbound) and it having a peering link with a given cloud provider. For example, a correlation ratio of 0.37 between Microsoft and Peering Policy indicates that an AS’s Peering Policy is moderately predictive of its likelihood to peer with Microsoft. In particular, ASes with selective peering policies are more likely to peer with Microsoft than ASes with other policies. Other policies, such as open policies, are less correlated, resulting in a moderate positive total correlation. Overall, we find that ASes with a selective peering policy and heavy inbound traffic are the most likely to be connected to cloud providers. The right side shows moderately strong correlations (0.27-0.54) between connections with one cloud provider and the existence of a link with other cloud providers: an AS known to peer with Microsoft is more likely to peer with Google than an AS known to not peer with Microsoft. On the other hand, we see no significant correlation with Cogent (0.02-0.06), nor with any other Tier 1 ASes (results not shown, but correlations fall between 0.01 and 0.12). Furthermore, these correlations offer complementary insights, such that combining data on Peering Policy, Traffic

⁵Past research has shown that the peering connectivity of these Tier 1 ASes is fully visible via the publicly available BGP collectors [120].

Profile, and the existence of connections with other cloud providers provides more information than any single feature or link alone.

Modeling the factors that shape peering decisions. The structured nature of peering decisions, where ASes with aligned business or technical incentives are more likely to form links, suggests that the connectivity matrix of ASes is likely to be shaped by the properties of the ASes. The limited number of underlying motivations—whether it is traffic patterns, geographic proximity, or customer demands—means that the set of plausible strategies is likely small, resulting in a connectivity matrix that can be efficiently summarized with few dimensions, making it low-rank. Low-rankedness implies that, with the right model, we can accurately predict invisible peering links using the factors embodied in the visible ones. We believe this low-rankedness holds for two main reasons:

First, the establishment of IXPs encourages multilateral peering relationships, often resulting in dense mesh networks among ASes that use an IXP’s route server to exchange routes. The full connectivity between this subset of ASes can be represented as a connectivity matrix composed entirely of 1s, a matrix of rank 1. While some ASes may choose not to connect to the route server or restrict peering with specific members, we confirmed that low-rank property holds in the 18 IXPs for which we found public peering matrices (e.g., [92]). The rank ranges between 3.7% (NAPAFRICA [111]) and 26% (INEX Dublin [81]) of the total connectivity matrix dimension, with an average of 12.6%.⁶

Second, the similarity of business models among ASes, such as ISPs, CDNs, Cloud Providers, and enterprise networks, pushes these entities to form peering connections with a similar set of peers, as we demonstrate in the previous example (Fig. 1). By modeling factors that drive the formation of visible links and assuming these same factors also influence the formation of unseen links, we can effectively infer the invisible links. This transferability assumes that (1) there is no fundamental difference between pairs of ASes that peer where available vantage points have measured the link and pairs where they have not (i.e., the decision to host a vantage point does not stem from properties of an AS that would also impact its (or its provider’s) peering strategy)⁷, except for (2) properties that are observable and can be incorporated into the model. For example, in Iran, both the probability of an AS hosting a vantage point and its peering practices are influenced by the country’s regulatory environment, and so our model will include an AS’s country of registration to capture this effect. We discuss these assumptions in greater detail in Section 5.3.

Therefore, we conceptualize our model as encompassing two types of properties of an AS: (i) aspects of the identity of AS (e.g., business, economic, political, historical factors), approximated by features of the AS, and (ii) its currently measured peering links, which reflect outcomes of its peering policy. Just as Tinder combines user demographics (e.g., age, location) with interaction history (e.g., likes, matches) to recommend connections, our model integrates AS characteristics and PeeringDB profiles (akin to user demographics) with known interconnections (interaction history) to infer unseen connections (suggested matches). Just as the likelihood of a Tinder

user liking a profile can be predicted in part from the similarity of the profile to others the user has liked and from how other users similar to the user have responded to the profile, known interconnections for two ASes inform how likely they are to interconnect. Section 4 demonstrates that available measurements and AS properties capture these latent factors in ways that machine learning techniques can use for inferences.⁸ To increase the reliability of each inferred link, we include a reflexive component that evaluates the model’s confidence in its predictions. This per-link assessment enables users to manage the trade-off between false positives and false negatives and to incorporate modeling of uncertainty into use cases that can accommodate it (§5).

Finding the right representation. The previous paragraph highlighted the potential of a model to reveal which links are more likely to exist but did not specify its characteristics. We now unveil metAScritic’s primary insight: using matrix completion theory to pinpoint what constitutes a good representation.

We transform the observed interconnections between ASes found in publicly available traceroutes into entries within a connectivity matrix. These data points not only confirm the existence of interconnections between ASes but sometimes even hint at a lack thereof (called “non-existence”) (§3.4). Using both this positive and negative data, metAScritic initializes ‘ratings’ to train a recommender system (§3.1). In some instances, findings about a link in one metro can be extrapolated to another, a process that enriches our representation (§3.3.2).

Improving representation through targeted measurements. Public measurements are often skewed, heavily featuring ASes that host or are near vantage points [34], while leaving others underrepresented. This phenomenon is similar to a version of Tinder’s recommendation system being overly trained on a narrow demographic, making it less effective at catering to a broader, more diverse user base. To counteract this skew, we augment our dataset with targeted measurements to uncover links in underrepresented sections of the Internet in public datasets and move toward a more balanced (and better) representation (§3.3).

The literature on matrix completion highlights that, given certain conditions—particularly when the entries in the matrix are revealed randomly or distributed uniformly across the matrix [6, 29, 133]—the completion can be achieved with high accuracy. Based on these conditions, we crafted a test based on the matrix’s effective rank [38] to estimate the quality of our representation and, consequently, our capacity to discern peering strategies. Our approach involves incrementally issuing more traceroutes to raise the rank until the test reveals that enough links have been measured, stopping once additional data ceases to affect the results (§3.2). In Tinder’s scenario, this approach is similar to progressively showing users random profiles and observing how well the matchmaking system predicts their preferences as it gets factored in their swipes. This process is stopped when it becomes clear that the system is capturing irrelevant data, like the users’ current mood, rather than their long-term preferences. Unlike Tinder, metAScritic can dynamically adjust the number of traceroutes performed to allow us to refine predictions

⁶We found the peering matrices by searching for the keyword “IXP Manager and Peering Matrix” on Google and crawling all the available matrices.

⁷Prior work has shown that ASes hosting RIPE Atlas showed little bias when compared to those that do not, across a broad range of metrics [145].

⁸Our model focuses on predicting associations rather than uncovering causal relationships, which would require a more advanced level of causal inference, as described by Judea Pearl [18].

incrementally, something that Tinder’s recommendation system cannot do. We integrate this approach into metAScritic, developing an iterative algorithm to ensure we gather enough ratings for precise identification of the latent factors.

Translating inferences to interconnections. The ultimate objective of metAScritic is to map the ratings back to AS links in a metro. While inferential approaches might raise reservations for some applications, we demonstrate that higher inferred ratings consistently map to greater accuracy for inferred links. This property allows users to customize thresholds to balance precision and recall based on their specific requirements (§5.1). Additionally, we implement techniques originating from explainable machine learning to offer insight as to why links were inferred and to help users gauge the reliability of the inferred links (§5.2).

3 Design

VAR.	DESCRIPTION
m	Metro
ℓ_{ijm}	Link between ASes i and j at metro m
$\mathbb{M}_m$	Initial connectivity matrix at m ; $\mathbb{M}_{ijm} \in [-1, 1]$
$\mathbb{C}_m$	Completed connectivity matrix at metro m ; $\mathbb{C}_{ijm} \in \{0, 1\}$
$\mathbb{E}_m$	Balanced estimated connectivity matrix at m ; $\mathbb{E}_{ijm} \in [-1, 1]$
$\mathbb{P}_m$	Probability of uncovering entry in $\mathbb{E}_m$; $\mathbb{P}_{ijm} \in [0, 1]$
$\mathbb{T}_m$	True (latent) connectivity matrix at m ; $\mathbb{T}_{ijm} \in \{0, 1\}$
λ	Threshold on $\mathbb{E}_{ijm}$ to decide link (non)existence
ϵ	Fraction (probability) of exploration measurements
n	Number of ASes in $\mathbb{C}_m$
r	Effective rank of $\mathbb{E}_m$

Table 1: Notation used in the paper.

Given the impossibility of measuring all links in the Internet, metAScritic uses matrix completion on a partial view of Internet connectivity to infer a more complete view (§3.1), where the accuracy of the inferences in the more complete view depends on the input partial view. We describe a principled approach to identify if the partial view contains enough information to enable accurate matrix completion (§3.2) and a system to iteratively collect traceroute measurements and extract information to complement the partial view until it contains enough information for a reliable completion (§§ 3.3 and 3.4). Our notations are given in Table 1.

3.1 Inferring interconnections via completion

metAScritic builds the final binary connectivity matrix $\mathbb{C}_m$ using matrix completion on the estimated connectivity matrix $\mathbb{E}_m$ built from traceroute measurements. We detail now how we employ complementary features to improve completion, pick the recommender system, and make a binary decision about the existence of a link.

Combining $\mathbb{E}_m$ and AS features. We complement information from traceroute measurements in $\mathbb{E}_m$ with AS features describing the profile of an AS and its peering strategy. For example, the traffic profile [122] or the number of users served by an AS [13] are factors that could drive the peering decisions of the organization maintaining the AS. A complete description of the features can be found in Appendix C. When building our recommendation system, we balance AS features, which inform peering strategies, with known peering link data, which show the results of these

strategies but not the reasoning behind them. Too much focus on AS attributes might overlook complex patterns observable in measurements, while underestimating them could ignore the impact of factors like economics or geopolitics [96, 103, 134]. We address this challenge by making the weights applied to $\mathbb{E}_m$ and AS features hyperparameters of the model, configured during training.

Picking the recommender system. Despite the ubiquitousness of neural networks in modern recommender systems, we opt for a simpler linear model trained with ALS [83]. This choice is driven by three main rationales. First, deep learning models are less interpretable, which complicates extracting insights, making them less trustworthy [74]. Second, the performance gains are negligible compared to our simpler approach (Appx. E.2). Third, neural network architectures do not provide a metric to assess whether enough links have been measured to perform good completion, whereas our linear approach does via the effective rank (§3.2).

Recovering the links. To map link existence estimates in $\mathbb{E}_m$ —continuous values in $[-1, 1]$ —into a binary determination of whether a link exists, we use a threshold λ . Varying λ allows us to be conservative or aggressive, trading off between precision and recall (§5.1). By default, we search to pinpoint the λ that maximizes the F-score, resulting in a configuration that combines good accuracy and recall. Section 5.1 discusses how varying λ affects the inferences and how different use cases benefit from choosing different thresholds.

3.2 Assessing if completion is trustworthy

Accurate completion of a matrix with an effective rank r requires at least r entries in each row and column [133], where the effective rank of a matrix determines the smallest number of dimensions required to reconstruct the full matrix within a small error margin. Applied to our case, knowing the effective rank of an AS connectivity matrix would tell us whether we have measured sufficient links to enable an accurate completion.

In practice, the actual effective rank is unknown, so we estimate it in an iterative fashion. We initialize the target effective rank $r_1 = 1$. In each iteration i , we randomly remove 3 entries per row, which we use to evaluate the accuracy of completion at the currently estimated effective rank. For each AS with fewer entries than our current rank estimate, we conduct targeted traceroutes until we reach our effective rank or hit a limit of successive traceroutes that fail to reveal the (non)-existence of any new links. We provide more details on the theory behind our work in Appendix B).

After these measurements, we compute the matrix completion’s mean squared error (MSE) using the removed entries from the rows with more than r_i entries. Rows with fewer than r_i entries are temporarily set aside to ensure precise rank estimation but get completed during the final completion (§3.1). We increment our rank estimate r_i by 1 and repeat the process. If the MSE does not improve over several iterations, we conclude the estimation, setting the effective rank to the one with the lowest MSE. This iterative process helps us zero in on the matrix’s true effective rank, with the MSE decreasing initially and then stabilizing as we approach

this rank. The stopping condition saves measurement budget and is justified experimentally (§4.2).⁹

Our methodology requires conducting sufficient measurements to discover r entries from as many rows as possible. Its efficacy is contingent on the presence of well-positioned vantage points. However, we do not possess a formal definition of what qualifies as ‘sufficient’ vantage points. As we ignore rows with fewer than r_i rows after a few batches for the prediction of the rank, limited vantage points could result in underestimating ranks and unreliable matrix completions. The impact of available vantage points is detailed based on our empirical findings in Section 4.4.

3.3 Tracerouting toward an accurate completion

A reliable completion at a given estimated effective rank requires at least that many known entries per AS. This section describes how metAScritic issues traceroutes from RIPE Atlas [130] to achieve that property. Although it is apparent which ASes need more information, it is not obvious which traceroutes to issue. A traceroute to or from a “deficient” AS may only traverse one of the AS’s known interconnections, and other traceroutes may not traverse the AS. Given the large number of vantage points and potential targets relative to RIPE Atlas rate limits [41], the challenge is selecting traceroutes that will efficiently uncover connectivity.

metAScritic maintains a matrix $\mathbb{P}_m$, with the same dimensions as $\mathbb{E}_m$, where each entry $\mathbb{P}_{ijm}$ represents the current estimate of the probability that it can strategically select a traceroute to fill $\mathbb{E}_{ijm}$. We first describe how we select the traceroutes given $\mathbb{P}_m$ (§3.3.1) and then how we derive the entries in $\mathbb{P}_m$ (§3.3.2).

3.3.1 Exploitation and exploration. To reduce the time it takes to collect traceroute measurements and better integrate with RIPE Atlas, metAScritic issues traceroute measurements in batches. We select the traceroutes in a batch iteratively, updating the number of known entries in each row of $\mathbb{E}_m$ assuming that previously selected measurements in the batch will be successful.

metAScritic uses two processes to select traceroutes. *Exploitation* traceroutes employ $\mathbb{P}_m$ to select traceroutes likely to fill entries in $\mathbb{E}_m$ that will uncover its rank. More precisely, we select exploitation traceroutes greedily by choosing the row i in $\mathbb{E}_m$ with the least number of filled entries (*i.e.*, the AS with the smallest number of existing or non-existing links) but with at least one $\mathbb{P}_{ijm} > 0.1$, breaking ties at random. We then select the entry $\mathbb{P}_{ijm}$ in that row with the highest estimated probability of success.

Exploration traceroutes do not use $\mathbb{P}_m$ and are a mechanism to counteract errors in probability estimations in $\mathbb{P}_m$. We select exploration traceroutes by choosing the row i and column j in $\mathbb{E}_m$ such that the sum of the number of entries in row i and column j is minimized and then picking the traceroute with the highest probability of determining whether or not that link exists. If there is no traceroute that can measure that entry, we keep iterating until we find a pair (i, j) with a possible traceroute. We use a parameter ϵ to control the probability of issuing an exploitation or exploration traceroute. Through empirical analysis detailed in Section 4, we

find that our estimated probabilities of success in $\mathbb{P}_m$ strongly correlate with filling entries in $\mathbb{E}_m$, and that a split of 0.9 exploitation and 0.1 exploration yields good results in practice. The fraction of exploration traceroutes could be increased to compensate for scenarios where $\mathbb{P}_m$ is less accurate (*e.g.*, when the entries in $\mathbb{P}_m$ do not correlate with the output of a batch). In each batch, we limit the number of exploration traceroutes per row to 1 to prevent the exploration from focusing on one particular row. Throughout the construction of $\mathbb{E}_m$, we limit the number of exploration traceroutes per entry in $\mathbb{E}_m$ to 1 to prevent repeated tries of measurements that are unlikely to succeed.

3.3.2 Estimating the probability of a traceroute to be informative. For each link ℓ_{ijm} , our goal is to identify (vantage point, target) pairs whose traceroutes have a high likelihood of showing the existence or the non-existence of ℓ_{ijm} . If a traceroute does provide such information, we say it is *informative*. We compute $\mathbb{P}_m$ by looking at whether previous traceroutes from the same source or to the same destination traversed the AS i , AS j , or metro m .

Categorizing vantage points and targets. For a possible link ℓ_{ijm} , we categorize vantage points and targets based on their geographical and topological relationship to the link. Those geographically and topologically closer to m and AS i are more likely to yield useful information about the existence of ℓ_{ijm} , but in some cases, distant vantage points may be the only ones available. For each vantage point, we categorize whether it is in the metro m , the same country (but not in the same metro), the same continent (but not in the same country), or elsewhere. We also categorize whether it is in AS i , i ’s customer cone, or outside i ’s customer cone. Vantage points are then categorized into one of the 12 categories based on the cross-product of geographical and topological locations.

We reuse these geographic and topological categories for targets with one change. Given that any IP address can be used as a target, whereas sources are limited to available vantage points, we do not consider targets announced by ASes outside AS j ’s customer cone, which are very unlikely to uncover unknown connectivity for j [118]. Instead, we add a category for targets that belong to AS j and are adjacent (in an existing traceroute) to an IP address belonging to an IXP in metro m . We break ties across all targets for a given strategy using the responsiveness in the ISI hitlist [55].

A pair composed of a category for a vantage point and a category for a target is called a *measurement strategy*, *e.g.*, measuring from a vantage point in AS i ’s metro m to a target in j in the same country. For each link ℓ_{ijm} , the set of (vantage point, target) pairs are grouped by measurement strategy (144 in total). Note that a (vantage point, target) pair can belong to different strategies for different links.

Initial Estimation of $\mathbb{P}_m$. For each strategy, we select a set of random (vantage point, target) pairs, if available, to conduct traceroutes and bootstrap its initial probability of success. Appendix D.6 describes how we apply hierarchical modeling to transfer knowledge across metros, employing the probability of success of each strategy at other metros to initialize their probability of success at a new metro m using less than 20% as many traceroutes. metAScritic maintains an estimated probability of collecting informative traceroutes using each strategy for each link ℓ_{ijm} . The probabilities of running an informative measurement when using a strategy are

⁹Even if a larger effective rank $r_i > r$ were to yield a lower MSE, it might not be preferable as a larger effective rank requires more traceroutes per row for accurate completion and may not generalize when applied to rows in $\mathbb{E}_m$ that have far fewer entries than r_i .

initialized with identical values across all links but evolve independently after initialization. In particular, we increase the estimated success probability for strategies based on the number of available vantage points and potential targets for a specific link. This adjustment favors measurements using strategies with a larger pool of (vantage point, target) pairs. The value of $\mathbb{P}_{ijm}$ is set to the highest estimated probability of success across all strategies for ℓ_{ijm} .

Updating $\mathbb{P}_m$. Each batch of traceroutes prompts an update in the count of both informative and total traceroutes that result in a recalibration of the probability of success as new measurements are collected. Should a traceroute yield no valuable insights about the target link ℓ_{ijm} (§3.4), we reduce the corresponding strategy’s success probability. This conservative approach prevents excessive measurements of any single link ℓ_{ijm} using one strategy, thereby enhancing diversity in the selection process. The reduction factor strikes a balance between repeating measurements on elusive links and avoiding undue penalties for unsuccessful attempts. This approach, along with our partition of (vantage point, target) pairs into strategies, allows $\mathbb{P}_m$ to quickly identify links ℓ_{ijm} and ASes i for which we lack good vantage points, steering exploitation traceroutes towards fruitful measurements and supporting effective use of the limited probing budget.

Choosing specific vantage points. We often find that multiple vantage points are available for a given target link ℓ_{ijm} and strategy selected using $\mathbb{P}_m$. This is particularly the case for strategies involving large geographical areas and the customer cone of large transit providers. However, not all vantage points within a strategy have the same probability to provide information about a link ℓ_{ijm} . metAScritic records and evaluates the performance of each vantage point in detecting links associated with AS i . Vantage points are scored based on the highest fraction of previous measurements that have confirmed the existence or non-existence of links in an AS i . When choosing vantage points, we adopt a biased random selection method that favors those with higher scores. Traceroute targets within a strategy are chosen randomly.

3.4 From traceroutes to links/non-links

We identify that a link exists whenever a traceroute traverses two consecutive hops belonging to different ASes. We map traceroutes into AS-paths using BDRMAPIT [101] and geolocate interconnections between ASes using a combination of IXP prefix databases [78, 82, 122], rDNS reverse engineering [95], and RTT-based constraints [70, 114] (Appendix D for details).

However, identifying that a link does not exist is more challenging and is not a problem that has been addressed previously, to the best of our knowledge. A link ℓ_{ijm} at metro m may exist and not be observed because, for example, (i) we lack a nearby vantage point to observe ℓ_{ijm} ; (ii) intradomain routing policies of ASes i and j steer traffic through another interconnection at a different metro (for load balancing or cost reasons); or even (iii) violations of valley-free routing [66]. To navigate these challenges, we only conclude ℓ_{ijm} does not exist if (i) AS i and AS j have *consistent routing policies* at metro m and (ii) there exists a traceroute originating from a *well-positioned vantage point* between ASes i and j that traverses a transit provider at m , two notions we detail now.

We say AS i consistently routes towards AS j in a location (metro, country, or continent) if AS i always uses the same type of link (direct interconnection or via an intermediate transit) when forwarding traffic towards destinations in AS j located at that location. We define an AS i as having a *consistent routing policy* if it has consistent routing toward all ASes j in the metro. Identifying ASes with *consistent routing policies* is detailed in Appendix D.5. We found that *inconsistent routing policies* are more common among major CDNs, cloud providers, and transit providers.

We say a vantage point is *well-positioned* to identify links at a given metro m for AS i if it has never issued a measurement or has previously issued a measurement that traversed an interface belonging to AS i located at m . This ensures the vantage point’s capability to assess connectivity for AS i in m . If a traceroute from this vantage point observes an intermediate transit provider between two ASes with *consistent routing* i and j , it strongly suggests that a direct interconnection ℓ_{ijm} does not exist. Identification of these well-placed vantage points is performed by analyzing already run traceroutes for each (AS, metro) pair.

Transferring links between locations. For some links ℓ_{ijm} , the scarcity of *well-positioned* vantage points constrains our ability to measure them. To work around this, we introduce the idea of geographic transferability, which builds on the observation that a link ℓ_{ijm} has a higher likelihood of existing if ASes i and j have other interconnections close to metro m . This idea is supported by the common practice among ASes to exchange routes with peers in shared locations to enhance redundancy, load-balancing, and resilience [115]. Appendix E.4 studies in detail the validity of our assumption and reveals that between 42–65% of measured interconnections exist across all locations where the two ASes are collocated. Geographic transferability significantly increases data retention in our model, capturing more information than if we were only considering the metro area.

To account for the uncertainty in this transferability of information from one metro to another, we assign different values to $\mathbb{E}_{ijm}$ for each AS pair i and j , based on the geographic location of their interconnection. Specifically, we use values of 1, 0.7, 0.4, and 0.1 to correspond to an interconnection’s presence in the same metro m , the same country, the same continent, or a different location, respectively. Similarly, for non-existent links, if both ASes i and j have consistent routing policies at the given granularity and only transit links are seen across all the paths that traversed them, we consider the transit link geographically closest to m . In this case, a direct interconnection between ASes i and j has a higher likelihood of existing the further away the transit is from m . This accounts for the possibility that the vantage point is far from m (i.e., not very well placed, low $\mathbb{P}_{ijm}$), and both ASes i and j are unwilling to carry the traffic to m . We fill $\mathbb{E}_{ijm}$ with a value of -1, -0.7, -0.4, and -0.1, depending on where the closest transit link is geolocated (metro, country, continent, or elsewhere). If we have both interconnections and transit links between ASes i and j , we fill $\mathbb{E}_{ijm}$ with the biggest absolute value. Although this approach may lead to occasional inaccuracies (i.e., mistakenly inferring connections in one metro when they only exist in another), the overall increase in the amount of data we can include in the matrix $\mathbb{E}_m$ for matrix completion outweighs these errors.

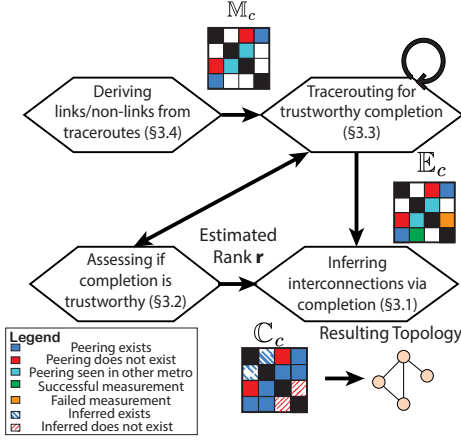


Figure 2: Interaction between metAScritic's modules.

3.5 Piecing the steps together

Having described the components of metAScritic, we now describe how they work together, as depicted in Figure 2. metAScritic takes a metro as input and outputs its inferences as to interconnections between ASes within that metro. To do so, metAScritic constructs a connectivity matrix (§3.4). One of metAScritic's key insights is that we can estimate the rank of the true matrix $\mathbb{T}_m$ on the fly. This step allows metAScritic to identify entries to fill in order to infer the remaining entries using matrix completion. To fill those entries, metAScritic relies on an algorithm that performs targeted traceroutes in rounds to estimate the effective rank of the matrix and fill in enough connectivity entries for an accurate completion (§3.2). These targeted traceroutes are selected based on a probability matrix estimating their likelihood to measure an interconnection's existence or non-existence (§3.3). After measuring enough entries, we finally perform our matrix completion $\mathbb{C}_m$ and translate the inferred ratings into binary links by fixing a threshold (§3.1).

4 Evaluation

In this section, we evaluate the effectiveness of each of the components that make up metAScritic by answering the following questions: Can metAScritic (i) identify AS interconnections with high recall and precision (§4.1), (ii) precisely estimate the effective rank of the connectivity matrix (§4.2), (iii) estimate the likelihood of a traceroute uncovering (or ruling out) a targeted link (§4.3), (iv) identify connectivity for ASes with limited vantage point without excessively inferring links where confidence is low (§4.4)?

Overview: metAScritic has an average Area Under the Precision-Recall Curve (AUPRC) score over six metros of 0.91 on different testing datasets. It achieves this performance (§4.1) while using only 0.3% of the number of measurements that an exhaustive approach would demand (Appx. E.3). metAScritic is capable of finding the true underlying effective rank of a generated connectivity matrix in a controlled environment (Appx. E.5) and is likely to find a better effective rank than alternative approaches in the wild (§4.2). metAScritic can effectively uncover links between ASes that lack vantage points while still adopting a conservative approach when data is too sparse (§4.4). metAScritic can predict whether a traceroute will be informative with high accuracy (§4.3).

4.1 End-to-end performance

We evaluate metAScritic on 6 metros with validation datasets obtained from multiple sources. We chose these 6 metros because they are geographically distributed across the world and together amount to approximately 11% of all the $(\text{Metro}, \text{AS}) \times \text{space}$ pair according to iGDB. We use the public traceroutes from RIPE Atlas and CAIDA Ark from the first week of February 2023 and targeted traceroutes from RIPE Atlas from February 2023 to August 2023 to obtain $\mathbb{E}_m$. This amounts to a total of more than 3 billion traceroutes, among which 1.8 billion traverse ASes hosted in one of the metros.

Precision and recall: For each metro, we initialize the connectivity matrix $\mathbb{E}_m$ with the public traceroutes and then run the subsequent steps of metAScritic except the completion. We evaluate the performance of the completion using precision-recall curves for each metro under two splits: (i) A *stratified* split, which removes 20% of the entries per row. (ii) A *completely-out* split, which removes all the entries of random rows removed until 20% of the entries of $\mathbb{E}_m$ are removed. This split simulates the case where there are ASes with no good vantage points and targets.

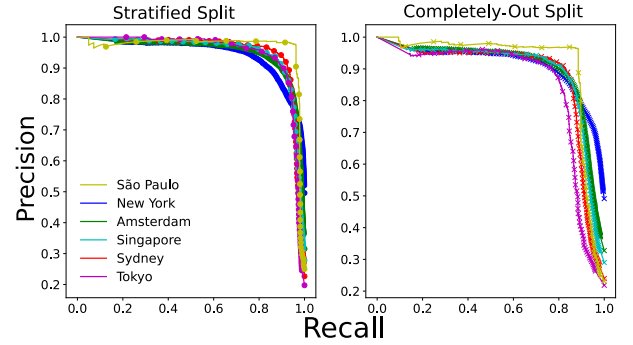


Figure 3: Precision-Recall curves for metAScritic across six metros for two splits. The high AUPRC scores (0.85-0.96) highlight metAScritic's capacity to accurately differentiate between existing and non-existing links.

In Figure 3, we present the precision-recall curves for six metros. The curves' proximity to the graph's top right corner indicates metAScritic high recall and precision. The high AUPRC, ranging from 0.85 to 0.96, demonstrates metAScritic's ability to identify true positives while maintaining a high level of precision. The worst performance is achieved with the completely-out split, with an average AUPRC equal to 0.88. ROC Curves are also shown in Appendix E.2 with an average Area-Under-Curve (AUC) of 0.98.

The completely-out split has worse results largely because most bad inferences are for ASes with fewer known existent and non-existent links than the estimated rank, justifying the need for supplemental targeted measurements. This observation also suggests two important properties of metAScritic: greater reliability in inferences for ASes with existing entries exceeding the estimated rank prior to completion and a point of diminishing returns when revealing more adjacencies for a specific AS. As detailed in Appendix E.8, ASes with fewer known links than the estimated rank are more

prone to misclassification. Conversely, 93.1% of cases with more entries than the estimated rank achieved a recall above 0.9.

Validation on external datasets. External validation plays a vital role because relying solely on cross-validation is susceptible to limitations of the input data—a system displaying high recall and precision with flawed data may be fundamentally ineffective.

A common hurdle is the scarcity of reliable ground truth data. To navigate this, we have combined six different external datasets, each shedding light on aspects of AS-level topology from various angles. The first dataset contains the peering links of two ASes: those of Vultr, a cloud provider, inferred from AS-paths observed on complete routing tables they export to the PEERING Testbed [138] across 30 PoPs around the globe, including 5 of our 6 metros; those of Google, obtained by replicating prior work’s methodology which consists of using a VM in every region [15]. The second dataset contains inferred peering links from Albakour *et al.* to infer IP aliases belonging to the same router [9]. In both instances, we geolocate the interconnections (Appx. D) and only consider those occurring within one of our metros. The third dataset contains links inferred from public BGP AS-paths tagged with *BGP Communities* known to tag interconnection locations, following the procedure outlined in prior work [67]. The fourth dataset corresponds to all the interconnections located at a single metro by *iGDB* [11]. The fifth dataset contains peering links obtained from measurements using *Looking Glass* servers from 12 different ASes used in recent papers [71, 162]. The sixth dataset contains the peering links obtained using IXP-manager from IXPs that openly share bilateral and multilateral peering matrices [1]. Our differentiation of bilateral peering and multilateral peering is supported by prior research findings, which showed that the majority of routes announced to route servers often lead to destinations outside the continent, typically involving three or more AS hops [125]. These routes are thus less likely to be discovered in our measurements. Additionally, another study has noted that multilateral peering accounts for only a minor portion of total traffic [129]. None of the validation datasets constitute ground truth, but those from Vultr and Google are the closest, allowing us to evaluate both precision and recall. The other validation datasets only evaluate the recall, as they only provide information about a set of existing links. The setup of metAScritic is the same as before, except that we perform the completion with all the entries in $\mathbb{E}_m$, as opposed to only 80% as in the cross-validation datasets. On the datasets from Vultr and Google, we have an average precision of 0.88 and recall of 0.91. On our other datasets, where we can only compute the recall, metAScritic has an even higher recall than on our testing dataset, with an average of 0.89 for the alias links, 0.98 for the BGP communities, 0.92 for the iGDB links, 0.90 for the looking glasses and 0.99 for the bilateral links and 0.68 for the multi-lateral links. More details on each metro can be found in Table 4.

4.2 Can metAScritic select the right entries to measure to improve accuracy?

The efficacy of the completion module is intrinsically tied to the quality of the traceroutes provided as input (§3.1). As delineated in Section 3.2, capturing an accurate representation hinges upon recovering the true effective rank. To this end, we engineered an iterative measurement selection mechanism. However, in the real

Strategies	Precision	Recall	Estimated Rank
Greedy	0.71	0.64	20
IXP-mapped	0.83	0.79	30
Random	0.67	0.65	18
Only Exploration	0.64	0.61	23
Only Exploitation	0.85	0.87	31
metAScritic ($\epsilon = 0.1$)	0.93	0.96	35

Table 2: Comparison of several targeted measurement strategies. metAScritic performs best.

world, there exists no way to know the true effective rank of the real connectivity matrix $\mathbb{T}_m$ *unless* one has direct access to the matrix. Instead, we compare the accuracy of metAScritic’s algorithm for selecting which link ℓ_{ijm} to measure against five alternatives: (1) a strategy that picks links ℓ_{ijm} at *random*; (2) a strategy that uses only *exploration* measurements, selecting ℓ_{ijm} for ASes i and j with the fewest filled entries regardless of the probability of success in $\mathbb{P}_{ijm}$; (3) a strategy that uses only *exploitative* measurements, selecting the AS i with the fewest filled entries and then the entry with the highest probability of success in $\mathbb{P}_{ijm}$; (4) *greedy*, where the links ℓ_{ijm} with the highest probability of being informative $\mathbb{P}_{ijm}$ are selected first, and (5) the algorithm described in IXP-mapped [17], a prior traceroute selection technique that is designed to uncover the links in an IXP. Once a link ℓ_{ijm} is selected, all strategies use metAScritic’s source and target ranking for that entry. We limit our analysis to Sydney as it requires a significant number of measurements.

To do a head-to-head comparison, we run metAScritic’s rank estimation approach, compute the number of measurements it issued, and set that as the budget to be used by all the other strategies and IXP-mapped. For the other techniques, the effective rank r of the resulting matrix is carried out post-hoc by adjusting it as a hyperparameter [56]. For each possible effective rank r , we run matrix completion and compute the F-score of the inferred links against the extensive measurements collected at Sydney (Appx. E.3). We set r to the value that maximizes the F-score.

Table 2 shows that metAScritic performs better than all the other approaches, with a precision of 0.93 and recall of 0.96. Techniques that focus on exploiting the entries (IXP-mapped and Only Exploitation) result in an average precision of 0.84 and recall of 0.83. The Pure Greedy, Random, and Only Exploration techniques result in the worst performance with precision ranging from 0.64 to 0.71 and recall ranging from 0.61 to 0.65. The estimated rank is also higher for metAScritic, indicating that its measured entries are capturing more complexity about the underlying structure of the connectivity matrix compared to the other techniques. Appendix E.5 presents a complementary evaluation showing that metAScritic also achieves better performance than the alternate strategies in a controlled scenario. In the controlled scenario, we simulated a complete topology (hence, with a known matrix rank) and progressively removed edges to mimic our partial topologies. This controlled experiment confirmed that metAScritic not only outperforms alternative strategies but also recovers the correct effective rank.

4.3 Can metAScritic estimate the probability of a traceroute to be informative?

Section 4.2 showed that metAScritic does a good job at determining the entries that need to be measured, allowing accurate completion. Recall that $\mathbb{P}_m$ estimates the probability of an entry to be filled in $\mathbb{E}_m$ if we choose to issue a traceroute for it. We now evaluate whether

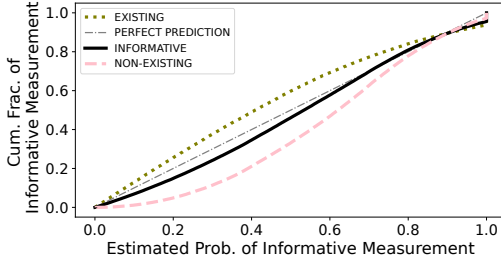


Figure 4: CDF of the estimated probabilities, for different sets of traceroutes. Overall, the set of informative traceroutes is quite close to the perfect prediction line.

our estimated probabilities in $\mathbb{P}_m$ are accurate by looking at whether the traceroutes that we select fit their estimated probabilities.

Results: We extract all the targeted traceroutes from the end-to-end evaluation (§4.1) and their estimated probability of measuring the existence or the non-existence of a link (*i.e.*, being informative) before issuing them, and look at whether they were actually informative or not. Figure 4 shows the CDF of how different sets of traceroutes performed: the ones that found existing links, those that indicated non-existing links, the ones that gave us information (the combination of the previous two), and those that did not yield useful information. Ideally, we would like to observe a ‘perfect prediction line’ that matches a straight line in the graph. Our analysis reveals the informative measurements closely match this ideal, with a Kolmogorov-Smirnov distance of 0.04, indicating that our approach to categorizing sources and destinations by topology and geography effectively predicts the likelihood of traceroutes being informative. It is also important to note that these estimations do not need absolute accuracy; instead, they must act as good tie-breakers to select between different measurement strategies.

4.4 Can metAScritic identify connectivity for ASes lacking vantage points?

A critical aspect of metAScritic is its ability to infer connectivity for ASes even when there are limited or no vantage points available. This section evaluates metAScritic’s performance under such conditions, focusing on how well it can uncover invisible links that traditional measurement-based approaches might miss. We explore two key dimensions: first, how well metAScritic performs in identifying connections for ASes lacking vantage points, and second, how metAScritic handles connectivity inference across different cities with varying vantage point coverage. Together, these evaluations help us understand how metAScritic balances its inferences and moderates the transfer of knowledge from observable connections to the hidden parts of the topology, reducing inference confidence where data is insufficient.

Identifying connectivity for ASes lacking vantage points: To assess metAScritic’s performance for ASes where there are limited or no vantage points available, we explore two key questions: (i) does metAScritic assign lower confidence ratings for links where measurements provide less information (*e.g.*, ASes lacking well-placed vantage points), and (ii) can metAScritic still make high-confidence inferences in these challenging conditions?

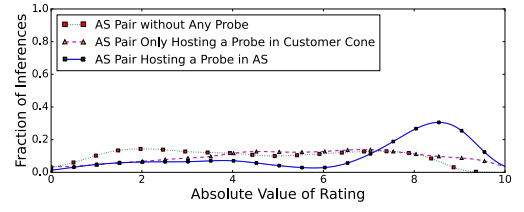


Figure 5: Relationship between probe coverage and the absolute value of inferred rating. When we have vantage points at one end of an AS pair, inferred ratings are generally higher than when we do not. However, even when we have no vantage points, we can still infer some links with high confidence, which would likely remain invisible if we were only using measurements.

Figure 5 illustrates the relationship between the presence of probes in the AS or its customer cones and the absolute value of rating. It shows the distribution of inferred ratings across all metros and AS pairs for three main categories: AS pairs with a vantage point at either end, AS pairs with a vantage point within the customer cone at either end but not in the ASes themselves, and AS pairs without any vantage point in their customer cone. A higher rating means higher confidence. Unsurprisingly, AS pairs without vantage points are the most common to receive the lowest confidence scores, but metAScritic can still infer some links with high confidence. These links would be impossible to detect with measurement-only methods. In contrast, ASes hosting vantage points tend to produce higher confidence inferences.

Probe coverage in different cities: The availability of vantage points within ASes (or within its customers) in a given metro or country impact metAScritic’s ability to make accurate inferences. In particular, if most ASes hosted in a metro only have poorly positioned vantage points, metAScritic may exhaust all available measurements and struggle to accurately infer connectivity. Although we lack a formal methodology to determine *a priori* the minimum number or placement of vantage points needed, we discuss some insights from running metAScritic across multiple metros. Figure 6 shows the distribution of the best available vantage points (§3.3.2) for every AS across all the 120 metros that host more than 50 ASes, according to iGDB [11]. In iGDB, these cities collectively amount to more than 95% of the potential city-level peering links on the Internet. The metros are ordered by the total percentage of ASes covered by at least one type of vantage point. We have highlighted in bold the metros where we ran metAScritic. São Paulo stands out as a metro where metAScritic struggled to infer connectivity accurately, with 86% of ASes with no probes in any category. This presents a challenge in a metro where the largest IXP in the world carries the most traffic volume [113], creating an expectation for rich connectivity. With such limited vantage point visibility, the accuracy of metAScritic’s inferences is more challenging to guarantee. In contrast, most European metros have much broader probe coverage, especially when focusing on probes located in the same metro or country as the AS or its customer, which gives us more data points to learn the underlying structure of the peering matrix in the metro. This increased visibility makes it more likely that our

matrix completion will be accurate, leading to more reliable inferences from metAScritic. These examples reflect a broader pattern of geographical disparities. Most metros with more than 50 ASes have at least one type of vantage point in at least half of their ASes, but all African and Latin American metros have fewer than 60% of their ASes covered by any probes. The relatively low number of African and Latin American metros in the figure highlights that the peering ecosystem is thin according to iGDB [11].

These insights guide us in predicting where metAScritic will likely perform better or worse. Specifically, based on the ranks computed in our study (Tbl. 4), we estimate that having approximately 5% to 15% of edges measurable is necessary to accurately infer metro-wide properties with higher confidence where the percentages are derived by dividing the matrix dimension by the estimated rank. Below this threshold, confidence decreases, and the results should be interpreted with greater caution.

5 Using the data

Researchers and operators are accustomed to links *measured* through traceroute or BGP paths, and so using links *inferred* by our system will require a shift in how links are used and how results are interpreted, to account for the uncertainty in inferences. Although using these links will introduce new challenges, it is important to remember that even links identified through traceroutes are, in essence, also inferential. For instance, Arnold *et al.* reported a false positive rate of 11-15% in their inferred links [15], and Marder *et al.* reported an error rate of 1.2-8.9% in their mapping technique [101]. Yet, these inaccuracies have not diminished the value of traceroutes in mapping AS-level connectivity, and over time, practitioners have learned to understand and often work around errors introduced by this tool [16, 99, 152]. In this section, we discuss how the research and operational communities might harness the insights from metAScritic. Even in cases where researchers are skeptical about inferential techniques, our system can be utilized to conduct targeted measurements to uncover peering links in a metro.

5.1 Using an adapted threshold

Bounding analysis by sweeping through thresholds: When translating the ratings inferred from metAScritic to binary values that capture whether two ASes are interconnected, the chosen threshold significantly affects the resulting topology. As the threshold decreases, more links are added but their precision diminishes, a phenomenon we explore in Figure 15. By analyzing the topology at various thresholds, we can increase our confidence in the relevance of the findings. Take, for instance, studies exploring the vulnerability of national Internet traffic to external observation and interference [49, 57, 90, 134]. Since these studies rely solely on BGP collectors for collecting topology data, they are likely to overstate the influence of a few key transit providers in controlling the nation’s Internet routes, when compared to topologies that factor in concealed peering links. These studies could be improved by reassessing their findings as they sweep through thresholds, including more peering links. The robustness of their findings across diverse thresholds enhances their reliability.

Enabling probabilistic reasoning: Another strategy is to assign each link a probability of existing based on its precision at a given

threshold. This approach enables probabilistic analysis and estimation of Internet properties as random variables. For example, RIPE could couple a probabilistic topology with our analysis of how vantage points are likely to measure the existence of a link (§3.3.2), to identify the best locations to deploy new RIPE Atlas vantage points [14, 124]. Depending on RIPE objectives, the best locations could be those predicted to (i) uncover the most previously unseen links, (ii) those predicted to remove the most uncertainty from the topology or (iii) a combination of those two criteria.

Balancing precision and recall: In Appendix F.3, we explore the trade-off between precision and recall as threshold values are adjusted. Our analysis shows that inferred edges with a 0.9 threshold have a 97-99% likelihood of being accurate. These high-confidence inferred edges, found in just six metro areas, reveal over 226K previously unseen peering links —equivalent to 0.7× the total peering links in the *entire* CAIDA AS relationship dataset [94].

5.2 Explain the outputs

Although metAScritic achieves high precision for higher thresholds (§5.1), it is important to recognize that its outputs are *still* not entirely error-free. This is similarly true for traceroute-inferred links. However, a fundamental distinction exists: Errors in traceroute-inferred interconnections can be explained. They mainly stem from non-responsive, invisible routers [31, 47, 99], off-path third-party addresses [80, 93, 159], or rare cases of “lying devices” using unauthorized IP addresses [91, 135]. In contrast, the inaccuracies in inferential approaches do not have clearly defined origins. Nonetheless, to enable some interpretation of the output, we use Shapley values, which quantify the contribution of each feature to individual predictions [106]. In our context, to compute the Shapley values of the features, we consider all possible combinations of features, or “coalitions,” and compute how the inclusion or exclusion of a feature affects the system’s recommendation. The Shapley value for each feature is then calculated as the average of its marginal contributions across these coalitions. In reality, the number of coalitions increases exponentially, so we use the *SHAP* library to approximate it [128]. This approach offers a comprehensive understanding of each feature’s influence on the recommendation and provides visibility on how and why metAScritic makes its decisions.

Appendix F.2.1 describes our comprehensive analysis across all features. We summarize the findings here: The model primarily relies on the number of existing and non-existing links from traceroutes. Additionally, information about the shared footprint of ASes, their customer cone size, and individual AS attributes (captured through one-hot encoding) are the most informative. Conversely, features from PeeringDB contribute minimally to metAScritic’s decisions, which is consistent with prior research [67]. This does not imply that these factors do not correlate with the existence of a peering link but rather that the model prioritizes other features. Appendix F.2.2 provides an example of how Shapley values can be used to understand metAScritic’s inferences for a specific link.

5.3 Limitations

Measurement limits: Our measurements rely on two assumptions. First, we assume that traceroutes can reliably identify peering links between ASes and in a specific metro with high confidence. In

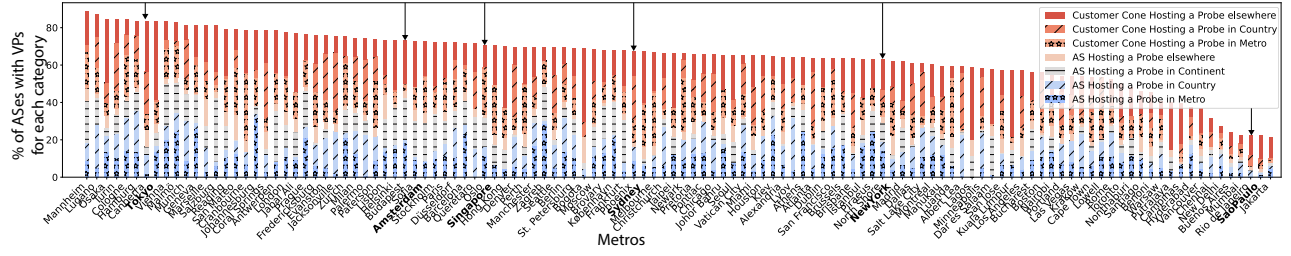


Figure 6: Distribution of the best vantage points present in RIPE Atlas for the metros housing more than 50 ASes. Arrows ↓ indicate Metros where metAScritic has been run. While coverage is generally close across EU and NA metros, there exist some disparities for certain African and Latin American metros, which may affect inference accuracy in these regions.

reality, many factors can influence traceroute probes, as discussed in Section 5.3, potentially affecting the input data to our recommender system and, by extension, our final results. Second, we assume that we have enough traceroutes to infer whether the observed routing behavior also holds consistently across unobserved paths at different geographic levels (§3.4). However, missed inconsistent measurements—whether due to temporal changes or destination-specific variations—could lead to incorrectly labeling an AS as consistent, which may result in wrongly inferring that certain links do not exist when they could.

Rank limits: metAScritic operates under the assumption that the true rank of the peering matrix in a city can be identified with enough measurements (§3.2). While this would be trivially true if we had vantage points everywhere and an unlimited measurement budget, there are cases where available vantage points have no visibility into whether a certain link exists, making accurate rank estimation impossible to guarantee and causing our algorithm that estimates ranks to potentially fail. The algorithm assumes convergence, such that adding more measurements no longer changes the best-performing rank for the completion. While this assumption works well in controlled scenarios (Appx. E.5), we lack formal proof of its general validity. Proving such guarantees is challenging, especially since the result may only hold for some matrices. Understanding these limits remains an open problem.

Inferential limits: metAScritic rests on two key assumptions: (1) The decision to host a vantage point does not influence the AS’s peering strategy, except for (2) observable properties that we can account for in the model. As metAScritic ingests more data about these observable properties, it becomes increasingly adept at identifying how these properties influence peering decisions. This learning process enables the system to refine its predictions by accounting for the impact of these properties on the outcome. In causal inference, this property is known as conditional ignorability, or unconfoundedness [121]. Simply put, given a treatment T (i.e., the presence of a vantage point), a covariate X (i.e., AS features, whether invisible or observable), and the potential outcomes (i.e., the existence of a peering link), unconfoundedness assumes that, conditional on X , the potential outcomes are independent of the treatment T . However, these assumptions can break down when certain latent factors influence all observable variables in ways we cannot untangle. For example, a political factor, such as government policy, may affect both the likelihood of an AS hosting a vantage point and its peering strategy. If this political factor is unobserved

and unaccounted for, any resulting analysis may incorrectly attribute changes in the AS’s peering behavior solely to the presence of a vantage point. In reality, the unobserved political factor could be influencing both the treatment (i.e., hosting a vantage point) and the outcome (i.e., peering connectivity in the country), leading to biased inferences (e.g., Iran [134] or Venezuela [32]). We leave as future work the task of more rigorously defining the conditions under which we can run metAScritic.

6 Use cases

We use the expanded topologies generated by metAScritic to revisit Internet topology properties and models. These topologies can also inform projects impacted by the incompleteness of the AS-level topology. Appendix F.1 details many applications where our new inferred links could help enhance the topology.

Internet topology: In Appendix G, we study the effects of measured and inferred links on the topology. Our analysis reveals a drastic increase in the number of links for hypergiants and content providers—quadrupling and almost doubling, respectively, compared to public BGP data. In contrast, the increase in links for Tier-1/2 and stub networks is comparatively smaller (less than $\times 1.3$) as these entities mostly connect via customer-provider links that are already better captured in the BGP public view.

Predicting the impact of route leaks and prefix hijacks: This is challenging due to limited visibility of existing links and opaque routing policies and filtering practices such as RPKI and Peer-Lock [65, 105, 116, 127]. Although we cannot model these opaque policies, we show that metAScritic’s inferences support more accurate predictions, helping operators to better evaluate these threats.

We predict the impact of BGP hijacks under three topologies. We start with CAIDA’s AS-relationship database built using *public BGP* information as a baseline [94], then add links discovered with active *measurements*, and finally add link *inferences* from metAScritic. We assume each AS follows Gao-Rexford routing policies [58], with AS-relationships taken from CAIDA’s database and assuming all metAScritic’s links are peer-to-peer [120]. We propagate all paths that are tied for best according to the Gao-Rexford model so an AS can choose multiple equally preferred nest routes, potentially including one that is hijacked and one that is legitimate.

We compare predictions against ground truth obtained from PEERING [138]. For every pair of metros where we have run metAScritic, we emulate different hijacks by making competing

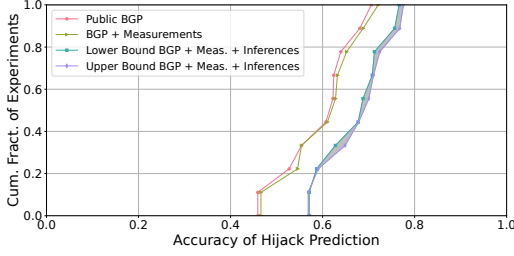


Figure 7: CDF of the accuracy of hijack prediction (fraction of ASes correctly predicted to be hijacked or not hijacked) across 90 BGP announcement configurations for 3 different AS topologies. Incorporating metAScritic’s inferred links leads to better predictions compared to using only BGP data or BGP data with measured links.

prefix announcements to nonoverlapping subsets of PEERING providers at each metro [144]. We deploy a total of 90 BGP announcements across all metro pairs and nonoverlapping subsets of providers. We use Verfploeter [44] to identify which ASes choose routes toward legitimate and hijacking announcements. Finally, we compare the (ground truth) results against predictions for the equivalent BGP announcement configurations.

Figure 7 shows the accuracy of predictions relative to ground truth across all three topologies. We consider the prediction correct if any of the best paths are hijacked/not hijacked in the same way as the actual route. We find that metAScritic’s inferences provide an average increase of 25% of the prediction accuracy compared to the public BGP data. The shaded area represents the range of values obtained by varying the threshold from 0.3 to 1 and demonstrates that the link generation threshold λ has little impact on the prediction accuracy, a positive sign that suggests that the improvement for this application is agnostic to metAScritic’s configuration.

Internet flattening: We evaluate the impact on metAScritic’s measurements and inferences on some flattening metrics, *i.e.*, for each AS with newly measured or inferred links; we compare the fraction of best paths going through providers and the AS-path length across three different topologies in each metro where we have run metAScritic: one with only BGP data, one with metAScritic’s measured links and one with metAScritic’s measured and inferred links. We also consider a *global* extended topology with links discovered across all 6 metros. We compute the best paths using the Gao-Rexford routing model. Table 3 shows an overview of the average decrease in AS-path lengths and the fraction of paths through providers for all metros in our evaluation. We report results for all ASes as well as for ASes registered in the country of the metro where we ran metAScritic. The results show that metAScritic has a significant impact on flattening metrics, particularly at the country granularity, with 20% of the path length reduced on average and a decrease of 21% of the paths that cross a transit link.

7 Related Work

Missing links of the AS topology: Recovering the AS connectivity has traditionally depended on BGP collector data [35, 54, 119], with limited visibility of peering paths due to valley-free routing [40]. This issue is further exacerbated by the Internet flattening

METRO	Fraction of shorter paths				Fraction of provider paths					
	All ASes		Country		All ASes		Country			
	+ M	+ Inf	+ M	+ Inf	BGP	+ M	+ Inf	BGP	+ M	+ Inf
Amsterdam	.050	.082	.124	.246	.830	.804	.772	.780	.759	.666
NewYork	.072	.144	.100	.187	.862	.789	.708	.896	.800	.697
Santiago	.002	.006	.072	.233	.932	.931	.928	.887	.858	.757
Singapore	.019	.045	.068	.243	.893	.885	.859	.849	.836	.727
Sydney	.009	.015	.096	.174	.945	.940	.934	.866	.832	.768
Tokyo	.007	.017	.105	.247	.873	.869	.854	.863	.823	.711
Global	.076	.115	—	—	.893	.873	.835	—	—	—

Table 3: Fraction of shorter AS-paths and of provider paths for different topologies. “+M” indicates addition of active measurements and “+Inf” indicates the addition of metAScritic inferences. Adding metAScritic’s inferred links not only increases the fraction of shorter paths but also significantly reduces reliance on provider paths.

[45, 64, 87]. To improve visibility, researchers have turned to alternative sources like IRR records [77], IXP datasets [98], active measurements [79, 146], and Bayesian statistics [131]. Research in this area typically focuses on contrasting publicly available Internet topologies with the more comprehensive views held by large organizations [7, 118] or on mapping the Internet’s invisible links in specific setup [15, 36, 77]. Our work aligns with the latter efforts but extends them by deploying more complex measurement techniques and applying inferential methods. Unlike recent studies that apply machine learning to AS topology without constraining the set of predicted links to plausible locations - potentially leading to the inference of unrealistic connections- our approach focuses on the metro-level connectivity, employs traceroutes for data collection, and runs targeted measurement for higher accuracy [61, 164].

Low-rankedness and completion in the Internet: Prior research has underscored the low-rank characteristics of various network matrices, such as traffic [132], latency [160], and those in network tomography [39, 51, 100]. These characteristics have been leveraged for matrix completion in several contexts [52, 75, 161]. Our research is the first to showcase the effectively low-rank structure of the connectivity matrix at the metro level and to incorporate the concept of adding measurements to enhance the “learning” process. This approach, while theoretically explored by Ruchansky *et al.* [133], is extended by accounting for the stochastic nature of measurements and determining the correct rank.

Pinning down interconnections to cities: Augustin *et al.* designed a methodology to specifically uncover AS connectivity within an IXP [17], while Motamedi *et al.* [109] and Giotsas *et al.* [70] designed techniques to pinpoint interconnections to facility/colocations. Our measurement strategies build from theirs, but we update ours to pin down the interconnections in the metro.

8 Conclusion

We presented metAScritic, a system that is able to infer the connectivity matrix between ASes within a metro, revealing significantly more links than prior work. We have shown that metAScritic has high precision and recall on various datasets, provided that it can accurately estimate the effective rank of the connectivity matrix. metAScritic can benefit any techniques and studies relying on an AS topology.

9 Acknowledgements

We would like to thank the anonymous reviewers for their constructive feedback and our shepherd for guiding us through the final submission process. This paper was partially funded by NSF grant CNS-2148275.

References

- [1] IXP Manager. <https://www.ixpmanager.org/>. Accessed: 2024-01-26.
- [2] Internet Exchange Points. URL https://www.pch.net/services/internet_exchange_points.
- [3] PCH Peering Policy. URL <https://www.pch.net/about/peering>.
- [4] Routing Information Service (RIS). <https://www.ripe.net/analyse/internet-measurements/routing-information-service-ris/>. Accessed: 2024-01-26.
- [5] Internet Exchange Points (ixps), Nov 2022. URL <https://www.internetsociety.org/issues/ixps/>.
- [6] Anish Agarwal, Munther Dahleh, Devavrat Shah, and Dennis Shen. Causal matrix completion. In *The Thirty Sixth Annual Conference on Learning Theory*, pages 3821–3826. PMLR, 2023.
- [7] Bernhard Ager, Nikolaos Chatzis, Anja Feldmann, Nadi Sarrar, Steve Uhlig, and Walter Willinger. Anatomy of a large european IXP. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, ACM SIGCOMM, page 163–174, New York, NY, USA, 2012. Association for Computing Machinery. ISBN 9781450314190. doi: 10.1145/2342356.2342393. URL <https://doi.org/10.1145/2342356.2342393>.
- [8] Apache Airflow. Apache airflow. URL <https://airflow.apache.org>, 2020.
- [9] Taha Albakour, Oliver Gasser, and Georgios Smaragdakis. Pushing Alias Resolution to the Limit. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, page 584–590, New York, NY, USA, 2023. Association for Computing Machinery. ISBN 9798400703829. doi: 10.1145/3618257.3624840. URL <https://doi.org/10.1145/3618257.3624840>.
- [10] Abdullah Alomar, Pouya Hamadian, Arash Nasr-Esfahany, Anish Agarwal, Mohammad Alizadeh, and Devavrat Shah. CausalSim: A causal framework for unbiased Trace-Driven simulation. In *USENIX Symposium on Networked Systems Design and Implementation*, USENIX NSDI, pages 1115–1147, Boston, MA, April 2023. USENIX Association. ISBN 978-1-939133-33-5. URL <https://www.usenix.org/conference/nsdi23/presentation/alomar>.
- [11] Scott Anderson, Loqman Salamatian, Zachary S Bischof, Alberto Dainotti, and Paul Barford. iGDB: connecting the physical and logical layers of the internet. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 433–448, 2022.
- [12] Ruwaifa Anwar, Haseeb Niaz, David Choffnes, Ítalo Cunha, Phillipa Gill, and Ethan Katz-Bassett. Investigating Interdomain Routing Policies in the Wild. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, page 71–77, New York, NY, USA, 2015. Association for Computing Machinery. ISBN 9781450338486. doi: 10.1145/2815675.2815712. URL <https://doi.org/10.1145/2815675.2815712>.
- [13] APNIC. Visible ASNs: Customer Populations (Est.). <https://stats.labs.apnic.net/asn/>.
- [14] Malte Appel, Emile Aben, and Romain Fontugne. Metis: Better Atlas Vantage Point Selection for Everyone. In Roy A Ensafi, Andra Lutu, Anna Sperotto, and Roland van Rijswijk-Deij, editors, *Network Traffic Measurement and Analysis Conference*, TMA. IFIP, 2022. ISBN 978-3-903176-47-8. URL <http://dblp.uni-trier.de/db/conf/tma/tma2022.html#AppelAF22>.
- [15] Todd Arnold, Jia He, Weifan Jiang, Matt Calder, Ítalo Cunha, Vasileios Giotas, and Ethan Katz-Bassett. Cloud provider connectivity in the flat internet. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 230–246, 2020.
- [16] Brice Augustin, Xavier Cuvellier, Benjamin Orgogozo, Fabien Viger, Timur Friedman, Matthieu Latapy, Clémence Magnien, and Renata Teixeira. Avoiding traceroute anomalies with Paris traceroute. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, page 153–158, New York, NY, USA, 2006. Association for Computing Machinery. ISBN 1595935614. doi: 10.1145/1177080.1177100. URL <https://doi.org/10.1145/1177080.1177100>.
- [17] Brice Augustin, Balachander Krishnamurthy, and Walter Willinger. IXPs: mapped? In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, 2009.
- [18] Elias Bareinboim, Juan Correa, Duligur Ibeling, and Thomas Icard. On Pearl’s Hierarchy and the Foundations of Causal Inference (1st edition). pages 507–556, 2022.
- [19] Srinadh Bhojanapalli and Prateek Jain. Universal Matrix Completion. In *International Conference on Machine Learning*, pages 1881–1889. PMLR, 2014.
- [20] Zachary S Bischof, Kennedy Pitcher, Esteban Carisimo, Amanda Meng, Rafael Bezerra Nunes, Ramakrishna Padmanabhan, Margaret E Roberts, Alex C Snoren, and Alberto Dainotti. Destination Unreachable: Characterizing Internet Outages and Shutdowns. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, ACM SIGCOMM, pages 608–621, 2023.
- [21] Boston University School of Public Health. Confidence Intervals. https://sphweb.bumc.bu.edu/otlt/mph-modules/bs/bs704_confidence_intervals/bs704_confidence_intervals_print.html. Accessed: 02-02-2024.
- [22] Timm Böttger, Gianni Antichi, Eder L Fernandes, Roberto di Lallo, Marc Bruyere, Steve Uhlig, and Ignacio Castro. The Elusive Internet Flattening: 10 Years of IXP Growth. *arXiv e-prints*, 2018.
- [23] Timm Böttger, Felix Cuadrado, and Steve Uhlig. Looking for hypergiants in peeringDB. *ACM SIGCOMM Computer Communication Review*, 48(3):13–19, sep 2018. ISSN 0146-4833. doi: 10.1145/3276799.3276801. URL <https://doi.org/10.1145/3276799.3276801>.
- [24] CAIDA. The CAIDA IPv4 Routed /24 Topology Dataset - 2023/10, . https://www.caida.org/catalog/datasets/ipv4_routed_24_topology_dataset/.
- [25] CAIDA. The CAIDA UCSD AS Classification Dataset, 2022–2023, . <https://www.caida.org/catalog/datasets/as-classification>.
- [26] CAIDA. The CAIDA UCSD AS to Organization Mapping Dataset, 2023/02. https://www.caida.org/data/as_organizations.xml, .
- [27] CAIDA, UCSD. The CAIDA UCSD BGP Community Dictionary. <https://www.caida.org/catalog/datasets/bgp-communities/>. [Accessed: 09/01/2024].
- [28] Matt Calder, Xun Fan, Zi Hu, Ethan Katz-Bassett, John Heidemann, and Ramesh Govindan. Mapping the expansion of Google’s serving infrastructure. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 313–326, 2013.
- [29] Emmanuel Candès and Benjamin Recht. Exact matrix completion via convex optimization. *Commun. ACM*, 55(6):111–119, jun 2012. ISSN 0001-0782. doi: 10.1145/2184319.2184343. URL <https://doi.org/10.1145/2184319.2184343>.
- [30] Emmanuel J Candes and Yaniv Plan. Matrix completion with noise. *Proceedings of the IEEE*, 98(6):925–936, 2010.
- [31] Esteban Carisimo, Caleb J. Wang, Mia Weaver, Fabián E. Bustamante, and Paul Barford. A Hop Away from Everywhere: A View of the Intercontinental Long-haul Infrastructure. *Proceedings of the ACM on Measurement and Analysis of Computing Systems*, 7(3), 2023. doi: 10.1145/3626778. URL <https://doi.org/10.1145/3626778>.
- [32] Esteban Carisimo, Rashna Kumar, Caleb J. Wang, Santiago Klein, and Fabián E. Bustamante. Ten years of the Venezuelan crisis - An Internet perspective. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, ACM SIGCOMM, page 521–539, New York, NY, USA, 2024. Association for Computing Machinery. ISBN 9798400706141. doi: 10.1145/3651890.3672218. URL <https://doi.org/10.1145/3651890.3672218>.
- [33] Tinder Help Center. Powering Tinder: The Method Behind Our Matching, 2023. URL <https://www.help.tinder.com/hc/en-us/articles/7606685697037-Powering-Tinder-The-Method-Behind-Our-Matching>. Accessed: 2024-09-07.
- [34] RIPE Network Coordination Centre. RIPE Atlas Anchors. Web Page, 2023. URL <https://www.ripe.net/measurements-tools/ripe-atlas/anchors>. Accessed: July 31, 2023.
- [35] Hyunseok Chang, Ramesh Govindan, Sugih Jamin, Scott J Shenker, and Walter Willinger. Towards capturing representative AS-level Internet topologies. In *Proceedings of the ACM on Measurement and Analysis of Computing Systems*, ACM SIGMETRICS, pages 280–281, 2002.
- [36] Kai Chen, David R. Choffnes, Rahul Potharaju, Yan Chen, Fabian E. Bustamante, Dan Pei, and Yao Zhao. Where the sidewalk ends: extending the internet as graph using traceroutes from p2p users. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, page 217–228, New York, NY, USA, 2009. Association for Computing Machinery. ISBN 9781605586366. doi: 10.1145/1658939.1658964. URL <https://doi.org/10.1145/1658939.1658964>.
- [37] Yi-Ching Chiu, Brandon Schlinker, Abhishek Balaji Radhakrishnan, Ethan Katz-Bassett, and Ramesh Govindan. Are We One Hop Away from a Better Internet? In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 523–529, 2015.
- [38] David B. Chua, Eric D. Kolaczyk, and Mark Crovella. Network Kriging. *IEEE Journal on Selected Areas in Communications*, 24:2263–2272, 2005.
- [39] Mark Coates, Alfred Hero, Robert Nowak, and Bin Yu. Internet tomography. *IEEE Signal processing magazine*, 19(3):47–65, 2002.
- [40] Rami Cohen and Danny Raz. The Internet Dark Matter-on the Missing Links in the AS Connectivity Map. In *INFOCOM*, IEEE INFOCOM. IEEE, 2006.
- [41] Ítalo Cunha, Pietro Marchetta, Matt Calder, Yi-Ching Chiu, Bruno V. A. Machado, Antonio Pescapè, Vasileios Giotas, Harsha V. Madhyastha, and Ethan Katz-Bassett. Sibyl: A Practical Internet Route Oracle. In *USENIX Symposium on Networked Systems Design and Implementation*, USENIX NSDI, pages 325–344, Santa Clara, CA, March 2016. USENIX Association. ISBN 978-1-931971-29-4. URL <https://www.usenix.org/conference/nsdi16/technical-sessions/presentation/cunha>.
- [42] Alberto Dainotti, Claudio Squarcella, Emile Aben, Kimberly C Claffy, Marco Chiesa, Michele Russo, and Antonio Pescapè. Analysis of country-wide internet outages caused by censorship. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 1–18, 2011.

- [43] Omar Darwich, Hugo Rimlinger, Milo Dreyfus, Matthieu Gouel, and Kevin Vermeulen. Replication: Towards a publicly available internet scale ip geolocation dataset. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 1–15, 2023.
- [44] Wouter B. de Vries, Ricardo de O. Schmidt, Wes Hardaker, John Heidemann, Pieter-Tjerk de Boer, and Aiko Pras. Broad and load-aware anycast mapping with verfloeter. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, page 477–488, New York, NY, USA, 2017. Association for Computing Machinery. ISBN 9781450351188. doi: 10.1145/3131365.3131371. URL <https://doi.org/10.1145/3131365.3131371>.
- [45] Amogh Dhamdhere and Constantine Dovrolis. The internet is Flat: Modeling the Transition From a Transit Hierarchy to a Peering Mesh. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, pages 1–12, 2010.
- [46] Amogh Dhamdhere and Constantine Dovrolis. Twelve Years in the Evolution of the Internet Ecosystem. *IEEE/ACM Transactions on Networking*, 19(5):1420–1433, 2011.
- [47] Benoît Donnet, Matthew Luckie, Pascal Mérindol, and Jean-Jacques Pansiot. Revealing MPLS tunnels obscured from traceroute. *ACM SIGCOMM Computer Communication Review*, 42(2):87–93, mar 2012. ISSN 0146-4833. doi: 10.1145/2185376.2185388. URL <https://doi.org/10.1145/2185376.2185388>.
- [48] Ben Du, Massimo Candela, Bradley Huffaker, Alex Snoeren, and kc claffy. RIPE Imap Active Geolocation: Mechanism and Performance Evaluation. *ACM SIGCOMM Computer Communication Review*, 50, Apr 2020.
- [49] Anne Edmundson, Roya Ensafi, Nick Feamster, and Jennifer Rexford. Nation-state Hegemony in Internet Routing. In *Proceedings of the ACM SIGCAS Conference on Computing and Sustainable Societies*, ACM SIGCAS, pages 1–11, 2018.
- [50] Stanley C. Eisenstat and Ilse C. F. Ipsen. Three Absolute Perturbation Bounds for Matrix Eigenvalues Imply Relative Bounds. *SIAM Journal on Matrix Analysis and Applications*, 20(1):149–158, 1998. doi: 10.1137/S0895479897323282. URL <https://doi.org/10.1137/S0895479897323282>.
- [51] Brian Eriksson, Gautam Dasarathy, Paul Barford, and Robert Nowak. Efficient Network Tomography for Internet Topology Discovery. *IEEE/ACM Transactions on Networking*, 20(3):931–943, 2011.
- [52] Brian Eriksson, Laura Balzano, and Robert Nowak. High-rank matrix completion. In *Artificial Intelligence and Statistics*, pages 373–381. PMLR, 2012.
- [53] Ksenia Ermoshina, Benjamin Loveluck, and Francesca Musiani. A market of Black Boxes: The Political economy of Internet surveillance and censorship in Russia. *Journal of Information Technology and Politics*, 19(1):p. 18–33, 2022. doi: 10.1080/19331681.2021.1905972. URL <https://hal.archives-ouvertes.fr/hal-03190007>.
- [54] Michalis Faloutsos, Petros Faloutsos, and Christos Faloutsos. On Power-Law Relationships of the Internet Topology. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, volume 29 of *ACM SIGCOMM*, pages 251–262. ACM New York, NY, USA, 1999.
- [55] Xun Fan and John Heidemann. Selecting Representative IP Addresses for Internet Topology Studies. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 411–423, Melbourne, Australia, November 2010. ACM. doi: <http://dx.doi.org/10.1145/1879141.1879195>. URL <http://www.isi.edu/%7Ejohnh/PAPERS/Fan10a.html>.
- [56] Bruno G Galuzzi, Ilaria Giordani, Antonio Candelieri, Riccardo Perego, and Francesco Archetti. Hyperparameter optimization for recommender systems through Bayesian optimization. *Computational Management Science*, 17:495–515, 2020.
- [57] Alexander Gamero-Garrido, Esteban Carisimo, Shuai Hao, Bradley Huffaker, Alex C Snoeren, and Alberto Dainotti. Quantifying Nations’ Exposure to Traffic Observation and Selective Tampering. In *Passive and Active Measurement*, PAM, pages 645–674. Springer, 2022.
- [58] Lixin Gao. Stable Internet routing without global coordination. *IEEE/ACM Transactions on networking*, 9(6):681–692, 2001.
- [59] Andrew Gelman. Multilevel (Hierarchical) Modeling: What It Can and Cannot Do. *Technometrics*, 48(3):432–435, 2006. doi: 10.1198/004017005000000661. URL <https://doi.org/10.1198/004017005000000661>.
- [60] Andrew Gelman, John B Carlin, Hal S Stern, David B Dunson, Aki Vehtari, and Donald B Rubin. *Bayesian data analysis*. CRC press, 2013.
- [61] Dimitrios P. Giakatos, Sofia Kostoglou, Pavlos Sermpezis, and Athena Vakali. Benchmarking Graph Neural Networks for Internet Routing Data. In *Proceedings of the International Workshop on Graph Neural Networking*, ACM GNNet, page 1–6, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450399333. doi: 10.1145/3565473.3569187. URL <https://doi.org/10.1145/3565473.3569187>.
- [62] Pietro Giardina, Enrico Gregori, Alessandro Improta, Alessandro Pischedda, and Luca Sani. Isolario: a Do-ut-des Approach to Improve the Appeal of BGP Route Collecting.
- [63] Petros Gligis, Matt Calder, Lefteris Manassakis, George Nomikos, Vasileios Kotronis, Xenofontas Dimitropoulos, Ethan Katz-Bassett, and Georgios Smaragdakis. Seven years in the life of Hypergiants’ off-nets. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, ACM SIGCOMM, pages 516–533, 2021.
- [64] Phillipa Gill, Martin Arlitt, Zongpeng Li, and Anirban Mahanti. The flattening internet topology: Natural evolution, unsightly barnacles or contrived collapse? In *Passive and Active Network Measurement*, PAM, pages 1–10. Springer, 2008.
- [65] Phillipa Gill, Michael Schapira, and Sharon Goldberg. Modeling on quicksand: dealing with the scarcity of ground truth in interdomain routing data. *ACM SIGCOMM Computer Communication Review*, 42(1):40–46, jan 2012. ISSN 0146-4833. doi: 10.1145/2096149.2096155. URL <https://doi.org/10.1145/2096149.2096155>.
- [66] Phillipa Gill, Michael Schapira, and Sharon Goldberg. A Survey of Interdomain Routing Policies. *ACM SIGCOMM Computer Communication Review*, 44(1):28–34, 2013.
- [67] Vasileios Giotsas, Shi Zhou, Matthew Luckie, and kc claffy. Inferring multilateral peering. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, pages 247–258, 12 2013. doi: 10.1145/2535372.2535390.
- [68] Vasileios Giotsas, Matthew Luckie, Bradley Huffaker, and kc claffy. Inferring Complex AS Relationships. pages 23–29, 11 2014. doi: 10.1145/2663716.2663743.
- [69] Vasileios Giotsas, Christoph Dietzel, Georgios Smaragdakis, Anja Feldmann, Arthur Berger, and Emile Aben. Detecting Peering Infrastructure Outages in the Wild. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, ACM SIGCOMM, pages 446–459, 2017.
- [70] Vasilis Giotsas, Giorgios Smaragdakis, Bradley Huffaker, Matthew Luckie, and kc claffy. Mapping Peering Interconnections to a Facility. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, pages 37:1–37:13, New York, NY, USA, 2015. ACM. ISBN 978-1-4503-3412-9. doi: 10.1145/2716281.2836122. URL <http://doi.acm.org/10.1145/2716281.2836122>.
- [71] Vasilis Giotsas, Amogh Dhamdhere, and kc claffy. Periscope: Unifying Looking Glass Querying. In *Passive and Active Network Measurement*, PAM, Mar 2016. doi: https://doi.org/10.1007/978-3-319-30505-9_14.
- [72] Matthieu Gouel, Kevin Vermeulen, Maxime Mouchet, Justin P Rohrer, Olivier Fourmaux, and Timur Friedman. Zeph & Iris map the internet: A resilient reinforcement learning approach to distributed IP route tracing. *ACM SIGCOMM Computer Communication Review*, 52(1):2–9, 2022.
- [73] Enrico Gregori, Alessandro Improta, Luciano Lenzi, Lorenzo Rossi, and Luca Sani. On the incompleteness of the AS-level graph: a novel methodology for BGP route collector placement. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 253–264, 2012.
- [74] Arpit Gupta, Chris Mac-Stoker, and Walter Willinger. An effort to democratize networking research in the era of ai/ml. In *Proceedings of the ACM Workshop on Hot Topics in Networks*, ACM HotNets, pages 93–100, 2019.
- [75] Gonca Gürsun and Mark Crovella. On Traffic Matrix Completion in the Internet. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 399–412, 2012.
- [76] Junghee Han, David Watson, and Farnam Jahanian. Topology Aware Overlay Networks. In *Proceedings IEEE 24th Annual Joint Conference of the IEEE Computer and Communications Societies.*, volume 4 of *IEEE INFOCOM*, pages 2554–2565 vol. 4, 2005. doi: 10.1109/INFCOM.2005.1498540.
- [77] Yihua He, Georgios Siganos, Michalis Faloutsos, and Srikanth Krishnamurthy. Lord of the Links: A Framework for Discovering Missing Links in the Internet Topology. *IEEE/ACM Transactions On Networking*, 17(2):391–404, 2008.
- [78] Packet Clearing House. Internet Exchange Directory. <https://www.pch.net/ixp/dir>.
- [79] Bradley Huffaker, Daniel Plummer, David Moore, and kc claffy. Topology discovery by active probing. In *Proceedings 2002 Symposium on Applications and the Internet (SAINT) Workshops*, pages 90–96, 2002. doi: 10.1109/SAINTW.2002.994558.
- [80] Young Hyun, Andre Broido, et al. On third-party addresses in traceroute paths. In *Passive and Active Network Measurement*, PAM. Cooperative Association for Internet Data Analysis (CAIDA), 2003.
- [81] INEX. INEX Peering Matrix. <https://www.inex.ie/ixp/peering-matrix>, 2024. Accessed: 2024-08-28.
- [82] Euro IX. IXP Database (IXPDB). <https://ixpdb.euro-ix.net/en/>.
- [83] Prateek Jain, Praneeth Netrapalli, and Sujay Sanghavi. Low-Rank Matrix Completion Using Alternating Minimization. In *Proceedings of the Annual ACM Symposium on Theory of Computing*, ACM STOC, page 665–674, New York, NY, USA, 2013. Association for Computing Machinery. ISBN 9781450320290. doi: 10.1145/2488608.2488693. URL <https://doi.org/10.1145/2488608.2488693>.
- [84] Yuchen Jin, Colin Scott, Amogh Dhamdhere, Vasileios Giotsas, Arvind Krishnamurthy, and Scott Shenker. Stable and Practical AS Relationship Inference with ProbLink. In *USENIX Symposium on Networked Systems Design and Implementation*, USENIX NSDI, pages 581–598, 2019.
- [85] Olga Klopp, Karim Lounici, and Alexandre B Tsybakov. Robust matrix completion. *Probability Theory and Related Fields*, 169:523–564, 2017.
- [86] Thomas Krenc, Robert Beverly, and Georgios Smaragdakis. AS-level BGP community usage classification. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, page 577–592, New York, NY, USA, 2021. Association

- for Computing Machinery. ISBN 9781450391290. doi: 10.1145/3487552.3487865. URL <https://doi.org/10.1145/3487552.3487865>.
- [87] Craig Labovitz, Scott Iekel-Johnson, Danny McPherson, Jon Oberheide, and Farnam Jahanian. Internet inter-domain traffic. 40(4):75–86, 2010.
- [88] Geoff Huston APNIC Labs. How Big is That Network?, 10 2014. URL <https://labs.apnic.net/index.php/2014/10/02/how-big-is-that-network/>.
- [89] Qrator Labs. National Internet Segments' Reliability Survey. <https://qratorlabs.medium.com/national-internet-segments-reliability-survey-df8129b47743>, August 2018.
- [90] Kirtus G Leyba, Benjamin Edwards, Cynthia Freeman, Jedidiah R Crandall, and Stephanie Forrest. Borders and gateways: measuring and analyzing national as chokepoints. In *Proceedings of the 2nd ACM SIGCAS Conference on Computing and Sustainable Societies*, pages 184–194, 2019.
- [91] Ioana Livadariu, Ahmed Elmokashfi, Amogh Dhamdhere, and kc claffy. A first look at IPv4 transfer markets. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, page 7–12, New York, NY, USA, 2013. Association for Computing Machinery. ISBN 9781450321013. doi: 10.1145/2535372.2535416. URL <https://doi.org/10.1145/2535372.2535416>.
- [92] LONAP. LONAP Peering Matrix. <https://portal.lonap.net/peering-matrix>, 2024. Accessed: 2024-08-28.
- [93] Matthew Luckie and Kc Claffy. A second look at detecting third-party addresses in traceroute traces with the IP timestamp option. In *Passive and Active Measurement*, PAM, pages 46–55. Springer, 2014.
- [94] Matthew Luckie, Bradley Huffaker, kc claffy, Amogh Dhamdhere, and Vasilis Giotas. AS Relationships, Customer Cones, and Validation. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, 2013.
- [95] Matthew Luckie, Bradley Huffaker, Alexander Marder, Zachary Bischof, Marianne Fletcher, and K Claffy. Learning to Extract Geographic Information from Internet Router Hostnames. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, page 440–453, New York, NY, USA, 2021. Association for Computing Machinery. ISBN 9781450390989. doi: 10.1145/3485983.3494869. URL <https://doi.org/10.1145/3485983.3494869>.
- [96] Richard TB Ma, Dah Ming Chiu, John CS Lui, Vishal Misra, and Dan Rubenstein. Internet Economics: The use of Shapley value for ISP settlement. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, pages 1–12, 2007.
- [97] Harsha V. Madhyastha, Ethan Katz-Bassett, Thomas Anderson, Arvind Krishnamurthy, and Arun Venkataramani. IPlane Nano: Path Prediction for Peer-to-Peer Applications. In *USENIX Symposium on Networked Systems Design and Implementation*, USENIX NSDI, page 137–152, USA, 2009. USENIX Association.
- [98] Priya Mahadevan, Dmitri Krioukov, Marina Fomenkov, Xenofontas Dimitropoulos, kc claffy, and Amin Vahdat. The Internet AS-Level Topology: Three Data Sources and One Definitive Metric. *ACM SIGCOMM Computer Communication Review*, 36(1):17–26, jan 2006. ISSN 0146-4833. doi: 10.1145/1111322.1111328. URL <https://doi.org/10.1145/1111322.1111328>.
- [99] Pietro Marchetta, Antonio Montieri, Valerio Persico, Antonio Pescapè, Italo Cunha, and Ethan Katz-Bassett. How and how much traceroute confuses our understanding of network paths. pages 1–7, 06 2016. doi: 10.1109/LANMAN.2016.7548847.
- [100] Morteza Mardani and Georgios B Giannakis. Estimating traffic and anomaly maps via network tomography. *IEEE/ACM transactions on networking*, 24(3): 1533–1547, 2015.
- [101] Alexander Marder, Matthew Luckie, Amogh Dhamdhere, Bradley Huffaker, KC Claffy, and Jonathan M Smith. Pushing the boundaries with bdrmapIT: mapping router ownership at Internet scale. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 56–69, 2018.
- [102] Michael Markovitch, Sharad Agarwal, Rodrigo Fonseca, Ryan Beckett, Chuanji Zhang, Irena Atov, and Somesh Chaturmohta. TIPSy: predicting where traffic will ingress a WAN. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, ACM SIGCOMM, page 233–249, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450394208. doi: 10.1145/3544216.3544234. URL <https://doi.org/10.1145/3544216.3544234>.
- [103] Ashwin Jacob Mathew. *Where in the world is the internet? Locating political power in internet infrastructure*. University of California, Berkeley, 2014.
- [104] Fabrizio Mazzola, Pedro Marcos, and Marinho Barcellos. Light, Camera, Actions: Characterizing the Usage of IXPs' Action BGP Communities. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, page 196–203, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450395083. doi: 10.1145/3555050.3569143. URL <https://doi.org/10.1145/3555050.3569143>.
- [105] Benjamin Tyler McDaniel. Interdomain Route Leak Mitigation: A Pragmatic Approach. 2021.
- [106] Luke Merrick and Ankur Taly. The explanation game: Explaining machine learning models using shapley values. In *Machine Learning and Knowledge Extraction*, 4th IFIP TC 5, TC 12, WG 8.4, WG 8.9, WG 12.9 International Cross-Domain Conference, CD-MAKE 2020, Dublin, Ireland, August 25–28, 2020, *Proceedings* 4, pages 17–38. Springer, 2020.
- [107] David Meyer. University of oregon route views project. <http://www.antc.uoregon.edu/route-views/>, 1997.
- [108] Sebastian Moss. London Internet Exchange disconnects Megafon and Rostelecom. *Data Center Dynamics*, March 2022. URL <https://www.datacenterdynamics.com/en/news/london-internet-exchange-disconnects-megafon-and-rostelecom/>. Accessed: 2024-09-01.
- [109] Reza Motamedi, Bahador Yeganeh, Balakrishnan Chandrasekaran, Reza Rejaie, Bruce Maggs, and Walter Willinger. On Mapping the Interconnections in Today's Internet. *IEEE/ACM Transactions on Networking*, PP:1–15, 10 2019. doi: 10.1109/TNET.2019.2940369.
- [110] Wolfgang Mühlbauer, Anja Feldmann, Olaf Maennel, Matthew Roughan, and Steve Uhlig. Building an AS-Topology Model That Captures Route Diversity. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, ACM SIGCOMM, page 195–206, New York, NY, USA, 2006. Association for Computing Machinery. ISBN 1595933085. doi: 10.1145/1159913.1159937. URL <https://doi.org/10.1145/1159913.1159937>.
- [111] NAFAfrica. NAFAfrica Peering Matrix. <https://ix.nap.africa/peering-matrix>, 2024. Accessed: 2024-08-28.
- [112] Luong Trung Nguyen, Junhan Kim, and Byonghyo Shim. Low-rank matrix completion: A contemporary survey. *IEEE Access*, 7:94215–94237, 2019.
- [113] NIC.br. In a new record, IX.br surpasses 3.1 Tbit/s of peak Internet traffic exchange. <https://nic.br/noticia/releases/in-a-new-record-ix-br-surpasses-31-tbit-s-of-peak-internet-traffic-exchange/>, 2024. Accessed: 2024-01-30.
- [114] George Nomikos, Vasileios Kotronis, Pavlos Sermpezis, Petros Gigos, Lefteris Manassakis, Christoph Dietzel, Stavros Konstantaras, Xenofontas Dimitropoulos, and Vasileios Giotas. O Peer, Where Art Thou? Uncovering Remote Peering Interconnections at IXPs. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, page 265–278, 10 2018. doi: 10.1145/3278532.3278556.
- [115] William B. Norton. The Art of Peering: The Peering Playbook. 2010. URL <https://api.semanticscholar.org/CorpusID:167173508>.
- [116] NTT. *Peerlock Manual*, 2023. URL https://instituut.net/~job/peerlock_manual.pdf. Accessed: 2023-02-01.
- [117] University of Oregon Route Views Project. University of Oregon Route Views Project. <http://www.routeviews.org/>. Accessed: 2024-08-26.
- [118] Ricardo Oliveira, Dan Pei, Walter Willinger, Beichuan Zhang, and Lixia Zhang. The (In)Completeness of the Observed Internet AS-level Structure. *IEEE/ACM Trans. Netw.*, 18:109–122, 02 2010. doi: 10.1109/TNET.2009.2020798.
- [119] Ricardo V Oliveira, Beichuan Zhang, and Lixia Zhang. Observing the evolution of Internet AS topology. In *Proceedings of the 2007 conference on Applications, technologies, architectures, and protocols for computer communications*, ACM SIGCOMM, pages 313–324, 2007.
- [120] Ricardo V Oliveira, Dan Pei, Walter Willinger, Beichuan Zhang, and Lixia Zhang. In search of the elusive ground truth: the Internet's AS-level connectivity structure. *Proceedings of the ACM on Measurement and Analysis of Computing Systems*, 36(1):217–228, 2008.
- [121] Judea Pearl. Causal inference in statistics: An overview. 2009.
- [122] PeeringDB. PeeringDB. <http://www.peeringdb.com>.
- [123] Gergana Petrova. RACI funding in 2021 - planning and projects, Nov 2021. URL https://labs.ripe.net/author/gergana_petrova/raci-funding-in-2021-planning-and-projects/.
- [124] Gergana Petrova. RACI Funding in 2021: Planning and Projects. https://labs.ripe.net/author/gergana_petrova/raci-funding-in-2021-planning-and-projects/, 2021. Accessed: 2023-02-01.
- [125] Lars Prehn, Franziska Lichtblau, Christoph Dietzel, and Anja Feldmann. Peering Only? Analyzing The Reachability Benefits of Joining Large IXPs Today. In *Passive and Active Measurement*, PAM, page 338–366, Berlin, Heidelberg, 2022. Springer-Verlag. ISBN 978-3-030-98784-8. doi: 10.1007/978-3-030-98785-5_15. URL https://doi.org/10.1007/978-3-030-98785-5_15.
- [126] Steffen Rendle, Walid Krichene, Li Zhang, and John Anderson. Neural collaborative filtering vs. matrix factorization revisited. In *Proceedings of the ACM Conference on Recommender Systems*, ACM RecSys, pages 240–248, 2020.
- [127] Andreas Reuter, Randy Bush, Italo Cunha, Ethan Katz-Bassett, Thomas C Schmidt, and Matthias Wählisch. Towards a rigorous methodology for measuring adoption of RPKI route validation and filtering. *ACM SIGCOMM Computer Communication Review*, 48(1):19–27, 2018.
- [128] Marco Tulio Ribeiro, Sameer Singh, and Carlos Guestrin. "Why Should I Trust You?": Explaining the Predictions of Any Classifier. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, KDD '16, page 1135–1144, New York, NY, USA, 2016. Association for Computing Machinery. ISBN 9781450342322. doi: 10.1145/2939672.2939778. URL <https://doi.org/10.1145/2939672.2939778>.
- [129] Philipp Richter, Georgios Smaragdakis, Anja Feldmann, Nikolaos Chatzis, Jan Boettger, and Walter Willinger. Peering at Peerings: On the Role of IXP Route Servers. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, page 31–44, New York, NY, USA, 2014. Association for Computing Machinery.

- ISBN 9781450332132. doi: 10.1145/2663716.2663757. URL <https://doi.org/10.1145/2663716.2663757>.
- [130] RIPE NCC. RIPE Atlas. <https://atlas.ripe.net/>.
- [131] Matthew Roughan, Jonathan Tuke, and Olaf Maennel. Bigfoot, Sasquatch, the Yeti and other missing links: What we don't know about the AS graph. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 325–330, 10 2008. doi: 10.1145/1452520.1452558.
- [132] Matthew Roughan, Yin Zhang, Walter Willinger, and Lili Qiu. Spatio-temporal compressive sensing and internet traffic matrices (extended version). *IEEE/ACM Transactions on Networking*, 20(3):662–676, 2011.
- [133] Natali Ruchansky, Mark Crovella, and Evimaria Terzi. Matrix completion with queries. In *Proceedings of the ACM SIGKDD international conference on knowledge discovery and data mining*, ACM KDD, pages 1025–1034, 2015.
- [134] Loqman Salamati, Frédéric Douzet, Kavé Salamati, and Kévin Limonier. The geopolitics behind the routes data travel: a case study of Iran. *Journal of Cybersecurity*, 7(1):tyab018, 2021.
- [135] Loqman Salamati, Todd Arnold, Ítalo Cunha, Jiangchen Zhu, Yunfan Zhang, Ethan Katz-Bassett, and Matt Calder. Who Squats IPv4 Addresses? *ACM SIGCOMM Computer Communication Review*, 53(1):48–72, apr 2023. ISSN 0146-4833. doi: 10.1145/3594255.3594260. URL <https://doi.org/10.1145/3594255.3594260>.
- [136] Loqman Salamati, Calvin Ardi, Vasileios Giotas, Matt Calder, Ethan Katz-Bassett, and Todd Arnold. What's in the Dataset? Unboxing the APNIC per AS User Population Dataset. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, ACM, 2024. ISBN 979-8-4007-0592-2/24/11. doi: 10.1145/3646547.3688411.
- [137] Quirin Scheitle, Oliver Gasser, Patrick Sattler, and Georg Carle. HLOC: Hints-based geolocation leveraging multiple measurement frameworks. In *Network Traffic Measurement and Analysis Conference*, TMA, pages 1–9. IEEE, 2017.
- [138] Brandon Schlinder, Todd Arnold, Ítalo Cunha, and Ethan Katz-Bassett. PEERING: Virtualizing BGP at the Edge for Research. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, page 51–67, New York, NY, USA, 2019. Association for Computing Machinery. ISBN 9781450369985. doi: 10.1145/3359989.3365414. URL <https://doi.org/10.1145/3359989.3365414>.
- [139] Patrick Schober, Christa Boer, and Lothar A Schwarte. Correlation Coefficients: Appropriate Use and Interpretation. *Anesthesia & Analgesia*, 126(5):1763–1768, 2018.
- [140] Kyle Schomp and Rami Al-Dalky. Partitioning the Internet using Anycast catchments. *ACM SIGCOMM Computer Communication Review*, 50(4):3–9, oct 2020. ISSN 0146-4833. doi: 10.1145/3431832.3431834. URL <https://doi.org/10.1145/3431832.3431834>.
- [141] Pavlos Sermpezis and Vasileios Kotronis. Inferring Catchment in Internet Routing. *Proc. ACM Meas. Anal. Comput. Syst.*, 3(2), jun 2019. doi: 10.1145/3341617.3326145. URL <https://doi.org/10.1145/3341617.3326145>.
- [142] Pavlos Sermpezis and Vasileios Kotronis. Inferring Catchment in Internet Routing. 47(1):19–20, dec 2019. ISSN 0163-5999. doi: 10.1145/3376930.3376943. URL <https://doi.org/10.1145/3376930.3376943>.
- [143] Pavlos Sermpezis, Vasileios Kotronis, Petros Gigis, Xenofontas Dimitropoulos, Danilo Cicalese, Alistair King, and Alberto Dainotti. ARTEMIS: Neutralizing BGP hijacking within a minute. *IEEE/ACM Transactions on Networking*, 26(6):2471–2486, 2018.
- [144] Pavlos Sermpezis, Vasileios Kotronis, Konstantinos Arakadakis, and Athena Vakali. Estimating the impact of BGP prefix hijacking. In *2021 IFIP Networking Conference (IFIP Networking)*, pages 1–10. IEEE, 2021.
- [145] Pavlos Sermpezis, Lars Prehn, Sofia Kostoglou, Marcel Flores, Athena Vakali, and Emile Aben. Bias in Internet Measurement Platforms. In *Network Traffic Measurement and Analysis Conference*, TMA, pages 1–10. IEEE, 2023.
- [146] Yuval Shavitt and Eran Shir. DIMES: Let the Internet measure itself. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, volume 35 of *ACM SIGCOMM*, pages 71–74. ACM New York, NY, USA, 2005.
- [147] Brivaldo A Silva Jr, Paulo Mol, Osvaldo Fonseca, Ítalo Cunha, Ronaldo A Ferreira, and Ethan Katz-Bassett. Automatic inference of BGP location communities. *Proceedings of the ACM on Measurement and Analysis of Computing Systems*, 6(1):1–23, 2022.
- [148] Rachee Singh, David Tench, Phillipa Gill, and Andrew McGregor. PredictRoute: A Network Path Prediction Toolkit. *Proc. ACM Meas. Anal. Comput. Syst.*, 5(2), jun 2021. doi: 10.1145/3460090. URL <https://doi.org/10.1145/3460090>.
- [149] Neil Spring, Ratul Mahajan, and David Wetherall. Measuring ISP topologies with Rocketfuel. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, volume 32 of *ACM SIGCOMM*, pages 133–145. ACM New York, NY, USA, 2002.
- [150] Neil Spring, Ratul Mahajan, and Thomas Anderson. Quantifying the Causes of Path Inflation. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, ACM SIGCOMM, page 113–124, New York, NY, USA, 2003. Association for Computing Machinery. ISBN 1581137354. doi: 10.1145/863955.863970. URL <https://doi.org/10.1145/863955.863970>.
- [151] Richard Steenberg. A practical guide to (correctly) troubleshooting with ... - Nanog. URL https://archive.nanog.org/meetings/nanog47/presentations/Sunday/RAS_Traceroute_N47_Sun.pdf.
- [152] Richard A Steenberg. A practical guide to (correctly) troubleshooting with traceroute. Nanog, 2009.
- [153] Gilbert Strang. *Introduction to linear algebra*. SIAM, 2022.
- [154] Bill Toulas. Web Page. URL <https://www.bleepingcomputer.com/news/security/all-dutch-govt-networks-to-use-rpki-to-prevent-bgp-hijacking/>.
- [155] Kevin Vermeulen, Justin P Rohrer, Robert Beverly, Olivier Fourmaux, and Timur Friedman. Diamond-Miner: Comprehensive Discovery of the Internet's Topology Diamonds. In *USENIX Symposium on Networked Systems Design and Implementation*, USENIX NSDI, pages 479–493, 2020.
- [156] Kevin Vermeulen, Ege Gurmericililer, Ítalo Cunha, David Choffnes, and Ethan Katz-Bassett. Internet scale reverse traceroute. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 694–715, 2022.
- [157] Bill Woodcock, Marco Frigino, and Packet Clearing House. 2016 Survey of Internet Carrier Interconnection Agreements. Packet Clearing House, Nov 2016.
- [158] Xiao Zhang, Tanmoy Sen, Zheyuan Zhang, Tim April, Balakrishnan Chandrasekaran, David Choffnes, Bruce M. Maggs, Haiying Shen, Ramesh K. Sitaraman, and Xiaowei Yang. AnyOpt: predicting and optimizing IP Anycast performance. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*, ACM SIGCOMM, page 447–462, New York, NY, USA, 2021. Association for Computing Machinery. ISBN 9781450383837. doi: 10.1145/3452296.3472935. URL <https://doi.org/10.1145/3452296.3472935>.
- [159] Yu Zhang, Ricardo Oliveira, Hongli Zhang, and Lixia Zhang. Quantifying the pitfalls of traceroute in AS connectivity inference. In *Passive and Active Measurement*, PAM, pages 91–100. Springer, 2010.
- [160] Rui Zhu, Bang Liu, Di Niu, Zongpeng Li, and Hong Vicky Zhao. Network latency estimation for personal devices: A matrix completion approach. *IEEE/ACM Transactions on Networking*, 25(2):724–737, 2016.
- [161] Rui Zhu, Bang Liu, Di Niu, Zongpeng Li, and Hong Vicky Zhao. Network Latency Estimation for Personal Devices: A Matrix Completion Approach. *IEEE/ACM Transactions on Networking*, 25(2):724–737, 2017. doi: 10.1109/TNET.2016.2612695.
- [162] Shuying Zhuang, Jessie Hui Wang, Jilong Wang, Zujang Pan, Tianhao Wu, Fenghua Li, and Zhiyong Zhang. Discovering Obscure Looking Glass Sites on the Web to Facilitate Internet Measurement Research. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, page 426–439, New York, NY, USA, 2021. Association for Computing Machinery. ISBN 9781450390989. doi: 10.1145/3485983.3494857. URL <https://doi.org/10.1145/3485983.3494857>.
- [163] Shuying Zhuang, Jessie Hui Wang, Jilong Wang, Zujang Pan, Tianhao Wu, Fenghua Li, and Zhiyong Zhang. Discovering Obscure Looking Glass Sites on the Web to Facilitate Internet Measurement Research. In *Proceedings of the ACM International Conference on Emerging Networking Experiments and Technologies*, ACM CoNEXT, page 426–439, New York, NY, USA, 2021. Association for Computing Machinery. ISBN 9781450390989. doi: 10.1145/3485983.3494857. URL <https://doi.org/10.1145/3485983.3494857>.
- [164] Shuying Zhuang, Jessie Hui Wang, Jilong Wang, Changqing An, Yuedong Xu, and Tianhao Wu. Predicting Unseen Links Using Learning-based Matrix Completion. In *NOMS 2022-2022 IEEE/IFIP Network Operations and Management Symposium*, pages 1–9. IEEE, 2022.
- [165] Maya Ziv, Liz Izhikevich, Kimberly Ruth, Katherine Izhikevich, and Zakir Durumeric. ASdb: a System for Classifying Owners of Autonomous Systems. In *Proceedings of the ACM Internet Measurement Conference*, ACM IMC, pages 703–719, 2021.

A Ethics

This paper raises no ethical concerns.

B Theoretical Formulation

We formalize the insights we have discussed in Section 2 and discuss some of the properties that we can guarantee via metAScritic.

B.1 Matrix Representation

Given a set of ASes, we represent the geo-annotated AS-level topology as a bipartite graph $\mathcal{B} = (\mathbf{M}, \mathbf{AS})$, where $\mathbf{M}$ are the metros considered in the dataset, and $\mathbf{AS}$ are the ASes in the topology. For every metro $m \in \mathbf{M}$, we build an edge to an $\mathbf{AS}_i$ in $\mathcal{B}$ if $\mathbf{AS}_i$

is present in any facility in m . Each $m \in \mathbf{M}$ admits a connectivity matrix $\mathbb{M}_m$ that represents the ASes that agree to exchange traffic in the given metro.

Effective low-rankedness of the true connectivity matrix:

Our work draws upon the insight that connectivity matrices often exhibit a low effective rank, an idea first introduced in Internet datasets by Chua et al. [38]. Low-rankedness of a matrix implies that only a small number of row vectors can approximate a matrix within a small error margin. In the context of connectivity matrices, a low effective rank suggests that knowing few AS interconnections is sufficient to capture the latent relationships that drive the connections between all ASes. We believe that the low-rankedness of the connectivity matrix in a metro holds for two main reasons:

The role of IXPs: The establishment of IXPs fosters multilateral peering relationships. These facilitate the emergence of peering connections via a centralized route server. This architecture typically results in the formation of a dense mesh network among all ASes utilizing the route server for traffic exchange. Nonetheless, deviations from this full mesh are observed when ASes choose not to connect to the IXP route server or deliberately avoid to exchange traffic with certain members of the traffic using BGP communities or IXP services that forward the BGP announcements on behalf of their members. A recent work published in 2022 has shown that Do not announce and Announce only to actions communities were rarely seen in the fabric [104].

The convergence of business models among ASes: ISPs, CDNs, Cloud Providers, Enterprise, Campus networks often align in their business strategies. This leads to the formation of peering connections with a similar set of entities, suggesting that the set of existing links for a given type of AS is likely smaller than the entire set of ASes.

Completion with Oracle: The problem of filling in missing entries in a partially observed matrix has a long history in inferential literature and is called matrix completion (see the survey [112] and references therein for a detailed description). Given the support set

$$\Omega_m := \{(i, j) \mid \text{AS } i \text{ has a peering link (or not) with AS } j \text{ in } v\}$$

whose cardinality is μ , we can define the projection operator $\mathcal{P}_\Omega$ by

$$\mathcal{P}_\Omega[X]_{(i,j)} = \begin{cases} m_{ij} & \text{if } (i, j) \in \Omega \\ 0 & \text{if } (i, j) \notin \Omega. \end{cases}$$

Then matrix completion problem asks if it is possible to use the μ -sparse observed ratings $Y \in \mathbb{R}^{i \times j}$ to recover X given

$$Y \approx \mathcal{P}_\Omega[X].$$

However, the problem we describe in Section 2 deviates from standard matrix completion for two reasons: (i) the pattern of missing entries is not random, and (ii) the likelihood of discovering new entries through additional measurements is not uniform, so classical approaches to uncover new entries are unsuitable for full matrix recovery. These observations align with challenges described in recent works applying causal techniques for predicting network properties [6, 10].

We draw a parallel between the ability to add new measurements to recover entries and the existence of an oracle that answers our queries. For each query sent to the oracle, there is a probability that

they will answer it. This probability expresses various underlying and unknown factors that influence the result of our query. These factors include, for example, the location from which measurements are taken, the chosen destination, the targeted ISP's responsiveness to ICMP probes, or the network's load at the time of measurement. Our objective is to achieve a high-accuracy reconstruction of the underlying matrix while keeping the number of queries to the oracle as low as possible. More specifically, for a given metro c , we devise an algorithm to reconstruct the symmetric matrix $\mathbb{T}_m = (t_{i,j})$. In this matrix, each element has a specific probability, $\Pi = (p_{i,j})$, of being disclosed by the oracle.

We first assume that the true rank r of the matrix is known. The spectral theorem for symmetric matrices states that $\mathbb{T}_m$ can be expressed as the product of a diagonal matrix Δ of dimension $n \times n$ and of a matrix X (of dimension $r \times n$) and its transpose X^T , such that $\mathbb{T}_m = X\Delta X^T$ [153]. Each element of the matrix, $t_{i,k}$, can be computed by the formula $\sum_{j=1}^n x_{i,j} \cdot x_{k,j}$. Consequently, the number of independent variables, or the degree of freedom, in $\mathbb{T}_m$ is equal to $r(n-r)/2$. As an example, let us assume that $\mathbb{T}_m$ is a rank 1 matrix, such that each entry is represented as $t_{i,j} = x_i \cdot x_j$. If all the entries for $i = j$ are observed, then for every i , $x_i = \pm \sqrt{t_{i,i}}$, we can then easily complete the rest of the matrix. The practical limitations in the number of measurements that can be issued imply that we need to adopt a strategy that involves the minimum number of measurements to ensure at least r filled entries in every row and column. This strategy, when framed as an optimization problem, seeks to minimize the expected number of measurements, guided by $\Pi(v)$, while constraining that every row and column possesses at least one filled entry.

Approximating the connectivity matrix using MLE: Our matrix is not exactly low-rank in practice. Instead, each entry in the matrix $\mathbb{T}_m$ is the product of a row of X and a column of Y , plus some unknown noise. This noise, assumed to be independently decomposable for each AS, reflects biases like measurement accessibility, network position, and business role. Therefore, each entry can be represented as $t_{i,k} = \sum_{j=1}^n x_{i,j} \cdot y_{j,k} + \epsilon_i + \epsilon_k$. Our final goal is to recover a low-rank approximation matrix $\mathbb{C}_m$ such that $|\mathbb{C}_m - \mathbb{T}_m| \approx 0$.

Extracting the underlying factors in this noisy scenario is more challenging. We compute the maximum likelihood estimator (MLE) to reduce noise for each estimated $t_{i,j}$ to address this. In this model, the value of each entry diminishes as more entries in the same row and column are revealed. This means that rows and columns with fewer known elements are more prone to noise, affecting the accuracy of our predictions. Our method balances this by uncovering new elements that simulate a perfectly random sample, aligning the matrix structure with the conditions needed for traditional matrix completion methods [19, 29, 85].

The algorithm we developed targets the least observed parts of the matrix, querying information efficiently while minimizing costs. It also avoids assuming a known rank by iteratively approximating the effective rank. While proving convergence to the true rank of $\mathbb{T}_m$ or its optimality remains an area for further research, empirical evaluations show promising results with real-world data and confirm its ability to identify the true rank in controlled simulations. These findings are detailed in Section 4.2 and Appendix E.5.

C Datasets

Our system takes support from many datasets. We describe them in detail in this section:

BGP Data: We use CAIDA’s AS relationship dataset to estimate the customer-cone of an AS [94]. The AS relationship is obtained by aggregating BGP announcements from BGP speakers at cooperating ASes and inferring the most likely relationship using the heuristics described in [94]. In addition, we collect BGP announcements from route servers at IXPs located in all the metros that we considered. Finally, we gather BGP routes and their associated BGP communities from all RIPE RIS [4] and University of Oregon Routeviews collectors [117] over a week of data.

Geographic Data: We leverage iGDB to bridge logical and geographical datasets [11]. iGDB extracts AS interconnection information indicating the geographic presence of participating networks for more than 25K networks that share information regarding their physical footprints and peering policies in Hurricane Electric, PeeringDB [122] and PCH [78]. To geolocate IP addresses, we also update the IP Table in iGDB using RIPE IP Map [48], HOIHO [95] and RTT constraints to geolocate IP addresses [43].

IXP Prefixes: We gather the prefixes corresponding to IXPs and the metro-level location of the IXPs by extracting information from PeeringDB, PCH [78], Hurricane Electric, and EuroIX. EuroIX collects data directly from IXPs through a recurring automated process, providing a comprehensive public source of IXP-related data and is known to be the most reliable source of information from operators. When extracting this information, we leverage data originating from sources that use a variety of data formats and naming conventions. To address this challenge we use heuristics from prior work [11] which looked at the overlap between the EuroIX, PeeringDB, Hurricane Electric, and PCH. For example, we pair up different naming conventions for infrastructures by looking at overlapping prefixes present in distinct datasets. When sources disagree, we prioritize Euro-IX then PeeringDB, PCH, and finally Hurricane Electric, since Euro-IX updates its database on a daily basis by directly and automatically gathering data from IXPs. In contrast, PCH and PeeringDB’s data are compiled manually by the ASes and the IXPs themselves, which is more error-prone and at the risk of being outdated. The exact strategy for data collection of Hurricane Electric is unknown and therefore harder to assess.

Number of Eyeballs: We regularly crawl APNIC’s eyeball estimate [13] and adjoin that information for every AS. The specific details of the technique are unknown; a brief description is provided in [88] suggesting that the technique is based on the utilization of Google’s advertisement delivery placement system to provide a rough estimate of the number of customers hosted in a network. This dataset has been demonstrated to be highly close to CDN traffic volume in all the countries considered in our study [136].

Looking Glasses: We have devised a methodology that utilizes Looking Glasses to validate a subset of our inferences. We replicate the pipeline outlined in [163] to run measurements from Looking Glasses. This technique leverages Looking Glasses that can be found from sources such as PeeringDB, Traceroute.org, BGP4.as, and BGPLookingglass.com. We also use Periscope, a publicly accessible tool that unifies LGs that automate the use of LG capabilities [71].

Peering DB AS Outbound/Inbound and AS Type Profile: We also crawl information about outbound properties of the different ASes [122]. For every specific IXP, we also fetch the data regarding the policy of member ASes when available..

RIPE Historical Measurements and Arkipelago: We build a system to gather, process and transform publicly available measurements from RIPE Atlas [130] and Arkipelago probes [24]. The data used in this paper come from the first week of February 2023.

D Methodology Details

D.1 Mapping IP to Org:

We use the state-of-the-art technique *bdrmapit* [101] to map the traceroute paths to their AS paths and keep only the traceroutes that cross at least two ASes in A , either consecutively or with intermediate ASes separating them. As a convention, we translate all the sibling ASes [26] to the lowest ASN managed by the organization.

D.2 Geolocating interconnections in v

To recover the matrix of interconnections $\mathbb{M}(v)$, we need to geolocate, for each interconnection between two ASes of A , the interfaces of the two border routers. First, we examine if the IP address assigned to any border routers belongs to an IXP prefix in v [78, 82, 122]. If so, we consider the interconnection to be located in v . Next, we examine the Round-Trip Time (RTT) obtained from the source of the traceroute. If the source is located in v and the RTT to both border routers is less than 3 milliseconds (as established in previous research [114]), we consider the interconnection to be in v . If the source is not in v , we perform ping measurements from 10 probes in v . If the minimum RTT for the two border routers from any of the probes is less than 3 milliseconds, we consider the interconnection to be located in v , in accordance with prior studies [114]. Another scenario that arises is when the two border routers are not located in the same metro, either because one is remotely peering with the other or due to inflated latency or unresponsive interfaces. In these cases, we adopt as a convention that the interconnection happens simultaneously in both metros, as it is difficult to determine which border router is remotely connected.

D.3 Classifying ASes

We classify ASes into the following classes, in order: *Tier-1/2* ASes are taken from Wikipedia (as currently used by CAIDA’s AS-relationship inference algorithm [68, 94] and other recent work [15, 147]); *Hypergiants* are taken from Gigis’ 2021 dataset [63]; *Large ISPs* are inferred as the ASes with the most users in each country covering 80% of the population using APNIC’s AS-population estimates [13]; we consider ASes as *Content* and *Enterprise* if they self-report as such on PeeringDB [122]; *Stub* networks are defined as those without any customer in CAIDA’s AS-relationship database [94]; and all the remaining ASes are classified as *Transit*. In the matrix completion part of metAScritic, we include all categories derived by ASdb as one-hot encoded variables [165].

D.4 Recommender system architecture

Matrix completion problems have been extensively studied in the context of recommender systems and generally fall into three primary categories: (i) Collaborative Filtering, which depends solely on matrix entries, (ii) Content-Based Filtering, where the emphasis lies on the side features of each entry, (iii) Hybrid models, which integrate both approaches. AS properties influence their roles as peering agents, as various organizations select peers according to their business incentives and their role in the network, making hybrid models the most suitable for our problem. By integrating features from different ASes into the prediction process, and we can refine our estimates about matrix entries. We split our features into numerical and category features. For categorical features, we utilize one-hot encoding and assume that there are n_1 numerical features and n_2 one-hot encoded features. To incorporate these features into our model, we augment the matrix by appending the feature associated with each AS to every row and column, expanding the matrix dimension from $|A| \times |A|$ to $(|A| + |n_1| + |n_2|) \times (|A| + |n_1| + |n_2|)$. Most recommender systems use a regularizer to avoid overfitting. We select the regularizer by performing a hyperparameter tuning following the strategy described in [56].

Alternating Least Square: The technique we use in this paper is called Alternating Least Squares (ALS). ALS aims to factorize a given matrix M into two matrices P and Q of rank r by optimizing $\arg\min |M - P^T Q|^2 + \lambda(|P|^2 + |Q|^2)$. Here, $P \in \mathbb{R}^{r \times n}$ and $Q \in \mathbb{R}^{n \times r}$. The ALS algorithm works by alternating between fixing one of the matrices, P or Q , and solving for the other. By fixing one of the matrices, the optimization problem becomes equivalent to solving multiple regression problems, and the optimal value for the fixed matrix can be computed explicitly. This process is repeated back and forth until convergence, providing an efficient way to solve the matrix completion problem.

D.5 Consistent Routing

For a given granularity, an AS exhibits consistent routing behavior if its connections with other ASes are consistently either peering links or discarded links. To catch ASes with inconsistent routing, we introduce a consistency binary matrix Γ , where rows and columns correspond to the ASes at location m (so same dimension as $\mathbb{E}_m$). A value of 1 signifies that we have observed at least one traceroute featuring a direct link between the two ASes and at least one traceroute with an intermediate link utilizing a provider at the considered granularity. In contrast, a value of 0 indicates that we consistently observed direct links or intermediate links with a provider for the given ASes. For example, if we observed two traceroutes connecting AS 1 and AS 2 directly in New York and Seattle and another traceroute passing through AS3, a provider of AS 1, in Toronto, then $\Gamma_{AS1,AS2} = 0$ at both country and metro granularities. However, the value would be 1 at the continent granularity since measurements have resulted in contradictory inferences for two metros within the same continent. In practice, our goal is to utilize inferences that span interconnections in other metros when an AS demonstrates consistent routing behavior at a particular granularity. Upon examining the consistency binary matrix, we discover that a small number of rows account for all inconsistent routing behaviors, with CDNs, Cloud Providers and a few large Transit Providers being

the common factors. This finding implies that inconsistency can be attributed to a few organizations. By iteratively eliminating the rows and columns associated with the ASes exhibiting the highest number of inconsistent rows until convergence, we derive a submatrix $\Gamma_{*,j}$ containing only consistent ASes. For this subset of ASes, we can rely on measurements from other metros to infer properties about the connectivity matrix in the metro.

D.6 Refining the probabilities via hierarchical modelling

To quickly predict the probability of the measurement strategies in an unseen metro, we build a hierarchical Bayesian estimator that captures the relationships between the strategies in different metros [60]. Prior works have shown that hierarchical models give more accurate predictions than no-pooling (i.e., considering every metro to be completely independent) and complete-pooling (i.e., collapsing all the metros together) estimation, especially when predicting group averages [59]. Hierarchical models recognize general patterns across metros while still allowing for unique metro-specific factors to exist. For example, we observed that in South America and South-Eastern Asia, probes hosted in the customer cone and the same metro were twice as likely to run informative measurements as in Europe because of the smaller pool of available upstreams. To fit our model, we began by taking 1K measurements for each strategy across 5 metros. We limited ourselves to 5 metros since the initiation step involves a large number of measurements which is practically challenging to conduct. Our model is structured in three tiers: the source-destination level, the metro level, and the AS level. For new metros, we use a maximum likelihood estimator for each (source, destination) category, building upon the model derived from previous metros. This approach allows us to recover $\Pi(v)$ with far fewer measurements (on average $6 \times$ fewer measurements) than if we had to rerun extensive measurements to recover the converging probabilities for every strategy in each metro under consideration. This step will allow us to scale measurement campaigns to more and more metros in the future.

E Additional Evaluation

E.1 Detailed table with the performance of metAScritic

Table 4 shows the detailed results on the validation dataset and different splits for the 6 metros.

E.2 ROC curves and comparison with other classifiers

When it comes to evaluating classification models, ROC and Precision-Recall curves are often used conjointly. The ROC Curves display the trade-off between the true positive and false positive rate of the classifier and are favored in scenarios where the dataset exhibits a balanced class distribution – that is, when the number of true positives and negatives are comparable. Since different metros result in different degrees of class imbalance, the ROC curves provide a complementary picture of metAScritic’s performance than the one shown in Figure 3. In Figure 8, we plot ROC curves for each metro of our dataset using the stratified split. We note

Does metAScritic identify AS interconnections efficiently and with high accuracy and coverage?							
Metro		Amsterdam	New York	Sao Paulo	Singapore	Sydney	Tokyo
Country		NL	US	BR	SG	AU	JP
Number of ASes		1470	748	1574	660	666	367
Estimated Rank by metAScritic		59	51	32	44	35	26
Train/Test Split and Validation with External Datasets							
Stratified	Recall	0.88	0.82	0.94	0.93	0.89	0.90
	Precision	0.85	0.84	0.96	0.90	0.96	0.86
Random	Recall	0.85	0.81	0.91	0.91	0.86	0.82
	Precision	0.84	0.78	0.89	0.91	0.87	0.81
Completely Out	Recall	0.79	0.81	0.87	0.78	0.81	0.85
	Precision	0.75	0.74	0.79	0.67	0.82	0.82
BGP Community	Recall	0.96	0.95	1.0	0.90	1.0	1.0
iGDB Geographic Hint	Recall	0.79	1.0	0.94	1.0	0.51 ¹⁰	0.61
Looking Glass	Recall	0.87	0.86	0.99	0.91	0.96	0.95
Bilateral IXP	Recall	0.97			1.0	1.0	1.0
Multilateral IXP	Recall	0.53			0.67	0.81	0.71
IP Aliasing	Recall	0.94	1.0	0.82	0.93	1.0	1.0
Ground Truth (Vultur and Google)	Recall	0.94	0.85	0.88	0.97	0.84	0.92
	Precision	0.82	0.89	0.78	0.93	0.95	0.84
How many measurements does metAScritic need to recover the connectivity matrix?							
Expected # of Measurements to Complete via Measurements		100M	43M	210M	34M	7M	6M
# of Measurements Run by metAScritic		330K	150K	275K	180K	150K	100K
metAScritic Eval. on Extensive Measurements	Recall					0.96	0.92
	Precision					0.93	0.93
No Targeted Measurements Eval. on Extensive Measurements	Recall					0.65	0.73
	Precision					0.56	0.63

Table 4: Performance of metAScritic across six different metros under varying configurations. On the upper part, we describe metrics related to precision and recall on diverse validation datasets and on the bottom part, we focus on the efficiency of metAScritic in terms of measurement reduction. The Estimated Rank by metAScritic corresponds to the rank estimated by the end of the process described in Section 3.2. This presentation of data allows for a holistic understanding of metAScritic’s performance across diverse scenarios and validates its robustness in inferring the missing interconnections, drastically reducing measurements while maintaining high precision and recall.

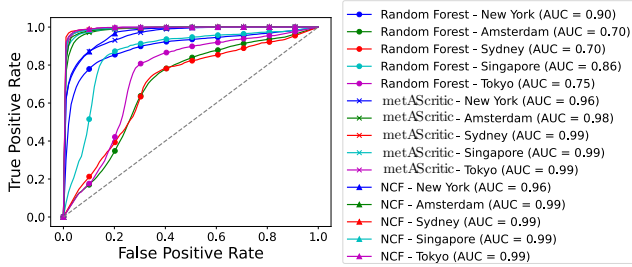


Figure 8: Comparison of ROC Curves for metAScritic, Random Forest, and Neural Collaborative Filtering (NCF). This graph demonstrates metAScritic’s efficiency in achieving a high True Positive Rate (TPR) while maintaining a low False Positive Rate (FPR). Notably, metAScritic’s performance parallels the more complex NCF architecture.

that for every metro considered, our AUC is even higher than our AUPC (between 0.96 and 0.99). Additionally, we extend this analysis by comparing our classifier’s performance with two alternative classifiers: (1) a random forest and (2) a recommender system trained with neural collaborative filtering [126]. The random forest serves as a baseline classifier and only builds on available public features to predict the existence of a peering link. Neural collaborative filtering moves away from decomposing the user-item interaction into two lower-dimensional matrices. Instead, it employs a neural network to create non-linear mappings of users and items. This approach enables the discovery of more complex patterns beyond linear relationships, albeit at the expense of

opacity and complexity in training. For our comparison, we opted for a simple design and used a multi-layer perceptron. We note that the Area-Under-Curve (AUC) is very similar for both metAScritic and the NCF. Furthermore, a decision tree that is agnostic to the global patterns in the connectivity matrix underperforms matrix completion approaches. These two observations further reinforce the appeal of our technique.

E.3 Efficiency in number of traceroutes

We compare metAScritic with two sets of *extensive measurements* collected for Tokyo and Sydney. In these datasets, we issue five traceroutes targeting each entry of the connectivity matrix (ranking source and targets using metAScritic’s approach in §3.3.2). This approach will not scale in many metros and requires launching millions of traceroutes, which would take months to complete at normal probing rates allowed by RIPE Atlas.¹¹ We run metAScritic on the $\mathbb{E}_m$ matrix obtained using its standard rank estimation algorithm, and evaluate the precision and recall relative to the entries of the connectivity matrix of the extensive measurements. Finally, we also compute precision and recall for a technique that performs matrix completion directly on the publicly available traceroutes

We observe that metAScritic performs almost as well as the extensive approach, evidenced by a marginal reduction of 0.06 in recall and 0.07 in precision compared to the extensive measurement. Notably, metAScritic achieves this with approximately 50 times fewer measurements than with the other approach. Finally, the bottom two rows of Table 4 show that the approach with only the public

¹¹Our RIPE Atlas account has a higher probing rate limit, but even then we limit ourselves to two metros.

measurements yields considerably worse results, with an average dip of 0.25 in recall and 0.34 in precision compared to the extensive campaign.

The theoretical literature on matrix completion suggests that, given the ability to recover any entries in an unknown matrix, the number of measurements needed is $O(nr \log(n))$, where n is the matrix dimension and r is the effective rank [30]. This corresponds to a scenario where one can determine the existence or non-existence of any link in the matrix with *one* traceroute, which is unrealistic in our scenario due to budget restrictions, vantage point placement, and unknown routing policies. Even under these challenges, comparing the number of measurements issued by metAScritic and the limit above (Tbl. 4) shows that metAScritic approaches this theoretical result while achieving high precision and recall.

E.4 How often can we transfer a peering link observed in one metro to another one?

Our method, as outlined in Section 3.4, relies on extrapolating the information about the existence and non-existence of a peering link from one metro to another. For example, in Dhaka, this inference accounts for up to 44% of our data entries, significantly enriching the initial matrix before its completion.

To validate this approach, we conducted the following study. Starting with ASes in Amsterdam with consistent routing defined in Section 3.4, we geolocated all the metros where these ASes interconnect in our measurements. We infer their geolocation by running our own implementation of HLOC, a technique that combines ping measurements with geographic hints observed in rDNS [137]. If by the end of our measurement campaign, the lowest observed latency from all of the utilized probes is smaller than 3 ms, we assume that the IP address is collocated in the same metrocity as the probe. Using the iGDB database [11], we determined for each AS pair with a known peering link, the set of metros where both ASes were collocated. For metros without measured interconnections between the two AS, we checked for available probes to discover the interconnection. In an effort to avoid unnecessary measurements, we only considered probes hosted in the metro and the AS or a direct customer or in the country and the AS (if no other interconnections in the country exist) as those were the instances, where the probability of identifying the existence or non-existence, were maximized in $\mathbb{P}_m$. If such probes were present, we consider that this link should be “measurable” and we conduct measurements from the best available probe-destination pair to uncover links in the metro. These measurements can lead to several possible outcomes: (1) The traceroute confirmed an interconnection between the two ASes in the targeted metro. (2) The traceroute identified an interconnection between the two ASes, but in a different metro. (3) A new, previously unobserved interconnection between the two ASes was discovered, though not in the intended metro. (4) The measurement provided no useful data (e.g., failure to cross either of the two ASes, unresponsive traceroute, etc.). (5) The packet routing occurred via a transit, indicating inconsistent routing in one of the ASes.

Outcome (1) confirms our hypothesis of transferability, whereas (2-3) suggests it might be incorrect (although the specific link could

be a backup in face of failures or heavy load but we adopt a conservative stance). (4-5) do not impact transferability as they imply either the link was initially unmeasurable or one of the peers had inconsistent routing (not relevant for transferability analysis).

As we were running our analysis, we noticed some instances of interconnections that we could not associate to any facilities from iGDB. This discrepancy could stem from their absence from the various databases collecting physical presence of ASes or from incorrect geolocation. For simplicity purposes, we decided to drop those instances and leave it as future work to complete and study in detail the missing facilities from iGDB.

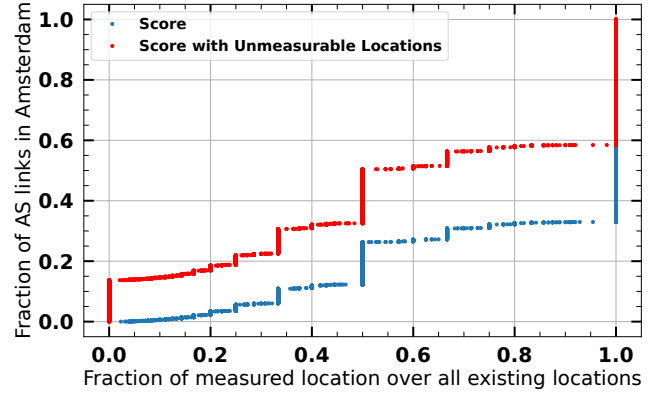


Figure 9: On the x-axis, we plot the ratio of metros where two ASes are interconnecting over the metros where they are both present according to the data from iGDB. In particular, we notice that 90% of the AS peering connect in at least 50% of all the metros where they have an overlapping presence.

In Figure 9, we evaluated the accuracy of our transferability hypothesis by computing two metrics. The first metric, depicted in red, adopts a conservative approach by calculating the score as the ratio of observed interdomain links to the total links for each AS pair. The second metric, shown in blue, takes a balanced approach, where the score is the ratio of observed interdomain links to the number of links that can be measured for each AS pair. In the conservative scenario, we presumed every unmeasurable link did not exist, while in the balanced scenario, these links were not taken into account. Our findings indicate that between 42-65% of interconnections are present across all locations, and a significant 70-90% are found in at least half of the locations. These results reinforce our belief that our measurements are in line with our expectations and confirm the metro-transferability for ASes with consistent routing.

E.5 Can metAScritic find the effective rank in a controlled environment?

The goal is to retrieve the (hopefully low) effective rank of a matrix with our algorithm that orchestrates the selection of traceroutes to issue, round by round (§3.2). To achieve this goal, we need three things: a matrix of low effective rank, a probability matrix to select the traceroutes to issue, and an initial set of visible entries preceding the algorithm execution, which correspond to the entries obtained

from the public traceroutes. To make our environment resemble real data, we use the connectivity matrix and the probability matrix inferred in Amsterdam from the run of metAScritic in Section 4.1. The probability matrix will be used as is, and we extract the set of indexes (i, j) of entries that were unknown before running any targeted measurements, that we call the mask of the connectivity matrix. This mask will serve to select the initial missing entries of the generated matrix. Finally, we extract the effective rank r of the connectivity matrix of Amsterdam, to generate a matrix of similar effective rank. To obtain this matrix, we start by generating a symmetric matrix of rank r with random values, where we add a Gaussian noise of standard deviation δ to every entry of our generated matrix. This technique generates an effective rank of r , as it provably results in at most r eigenvalues larger than δ [50]. Now that we have our generated matrix of low effective rank r , we use the mask to hide entries and to mimic the visible entries from the matrix before executing the algorithm. In our setup, we use a matrix of dimension $n = 1470$, an effective rank of $r = 59$ and we start at $r_0 = 1$, copying the data in Amsterdam from Table 4.

With this controlled environment, we then execute our algorithm, estimating the rank and running the targeted measurements in batches. To simulate the outcome of a traceroute, we draw a value from a uniform distribution and assume that a traceroute is informative if the value is smaller than the associated entry in the probability matrix. If it is the case, we consider that it unveils the actual value in the generated matrix for the computation of the MSE that serves to know when to stop running more traceroutes (§3.2). To be clear, the value does not matter here. What matters is the capacity of the algorithm to hit the true rank of the generated matrix and stop when it has hit it.

Figure 10 shows the mean square error of the different approaches over batches of 3,500 measurements, the setup that mirror the one uses in our end to end evaluation. We draw a black line at the iteration number where our estimated rank converges. metAScritic is the only approach with a decreasing RMSE, corresponding at some point to the true underlying rank. After the true underlying rank has been found, the RMSE starts to increase. In this simulation we let more than three rounds where the RMSE increases, but in practice metAScritic would have stopped and found the true effective rank. The other approaches exhibit a stable RMSE. These approaches fall short in providing a reliable mechanism to infer the matrix’s effective rank.

E.6 Does metAScritic pick good measurements?

In the previous experiment, we demonstrated the efficacy of our algorithm when operating under ideal conditions (*i.e.*, where the probability distribution is accurately inferred, and noise adheres to well-behaved theoretical properties). We now shift our focus to the performance of our algorithm when faced with real-world data. In this context, we both lack access to the true underlying rank and cannot compute the *true* RMSE. As a substitute, we evaluate the algorithm based on the total number of adjacencies discovered and the number of ASes exceeding the expected rank, as determined by a metAScritic run (Fig. 11).

We begin by examining the performance of the metAScritic selection scheme alongside three baseline strategies: random, greedy,

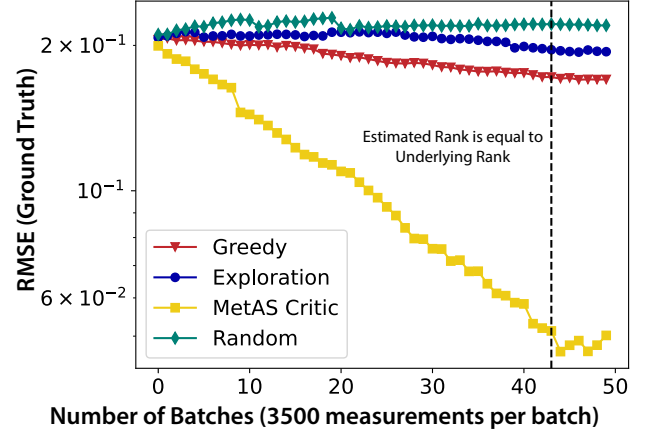


Figure 10: RMSE evolution across batches of targeted measurements on a matrix with a known effective rank in a controlled environment for different approaches. The black line indicates the effective rank at which convergence happens. metAScritic RMSE rapidly converges to the true underlying effective rank compared to other approaches.

and explorative. Additionally, we include the fully exploitative approach, in which our algorithm forgoes the explorative step (*i.e.*, setting ϵ to 0), and the metAScritic variant with a fixed exploratory step value of 0.1. We also compare our results with IXP-mapped [17], a technique closely related to our work, designed to uncover interconnections between members hosted within an IXP. IXP-mapped is predicated on selecting probes situated either in the source AS or in customer ASes within a maximum distance of two AS hops from the source AS. This technique prioritizes vantage points that (1) have previously been observed traversing the IXP fabric and (2) are geographically closer to the IXP. Target selection prioritizes responsive IP addresses in the destination AS. metAScritic with an ϵ exploratory step covers less of the visible entries compared to the fully greedy, fully exploitative and IXP-mapped but results in 12% more entries above the threshold compared to other approaches. In particular, we notice that the fully exploitative approach results in significantly fewer ASes above the threshold, underscoring the importance of random exploration to accommodate unexpectedly valuable measurements. Our greedy approach discovers more links than IXP-mapped, demonstrating the added value of our finer grain separation of probes and destination. Interestingly, metAScritic discovers entries nearly as efficiently as the Greedy version, although those entries are significantly more informative.

E.7 Can metAScritic exploit traceroutes to rule out the existence of links?

We study the accuracy of our methodology for determining the non-existence of a link. We only assign negative values in the connectivity matrix under certain conditions on the routing consistency of an AS and the position of the sources (§3.4). We evaluate our model versus simpler approaches (ordered from the most to the least conservative): (1) a *0-negative* approach, where we never introduce a negative value in the connectivity matrix (*i.e.*, non-existence

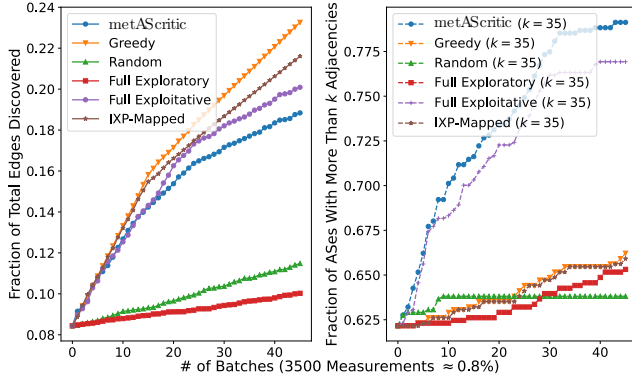


Figure 11: Fraction of the AS edges that were recovered or discovered per batch of 3500 measurements (left). Number of edges where we recovered more than k adjacencies, indicating that we could recover them if the underlying rank of the system is smaller than $r = 35$ (right).

is never inferred from measurements); (2) a *routing-inconsistency oblivious* approach where we put a negative value in $\mathbb{E}_{ijm}$ when the probe is well positioned (*i.e.*, non-existence is only inferred from specific probes that have previously observed a traceroute going through AS i in the metro). This approach will add more negative values in $\mathbb{E}_m$ than metAScritic as it ignores inconsistent routing (§3.4); (3) a *full negative* approach where we also remove the criteria of the probe being well positioned, resulting in even more negative values in $\mathbb{E}_m$. We then compute the precision and the recall of these different approaches on the Vultr and Google validation datasets on all overlapping metros, as those are the only datasets that provide a way of evaluating the accuracy of our non-existence inferences. We cannot evaluate the precision and recall on a train/test split as those metrics depend on the input data and could result in very high accuracy despite the starting matrix being incorrect.

Before the completion, the 0-negative approach fills 64.3% less entries in $\mathbb{E}_m$ than metAScritic, as it excludes any measurements that indicate the non-existence of a link. Conversely, the full negative and inconsistency oblivious approaches wrongly label 27% and 19% of the links that exist as non-existent, respectively. Analyzing these links in more detail, we find that these misclassifications are due to negative inferences made from traceroutes that would have been considered as inconsistent routing by metAScritic. We find that metAScritic has the best precision and recall in all metros. This empirical evidence underscores the importance of integrating the concept of routing consistency and advocates for a more nuanced model than that suggested by Gao-Rexford for deducing link non-existence—a finding aligned with other recent works discussing interdomain routing policies [12, 66, 110].

E.8 What is the relationship between the number of visible entries and the probability of incorrectly inferring peers?

In Figure 12, we discern a pivotal relationship between the number of visible entries and the prediction accuracy. Specifically, rows

with a count of measurements below the estimated rank exhibit an average error increase of 134% compared to those with more entries than the estimated rank. This finding highlights the significance of maintaining a balanced matrix to enhance the reliability of predictions.

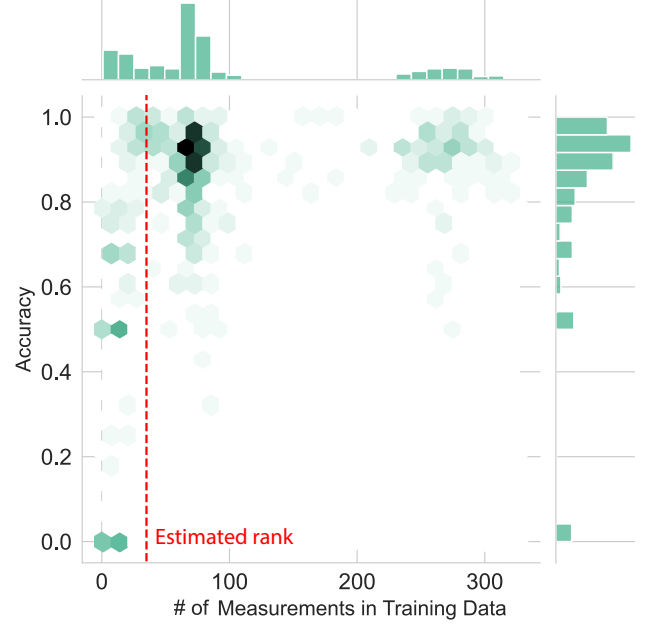


Figure 12: Relationship between the number of measured entries and the accuracy. We notice that as the number of entries increases, the accuracy increases. In particular, as the number of entries goes beyond the threshold, the accuracy gets close to 1.

F More on Usability

F.1 Building a taxonomy of applications

We delineate several pathways for how both the research and operational communities might use the outputs from our system. Here, we separate applications into three groups based on the nature of the datasets utilized: *conservative*, *loose*, and *balanced*. In practice, all these use-cases could be studied through the prism of the frameworks we described in Section 5 but they all possess a more natural setup to be considered. This list does not intend to be exhaustive but rather describe a few applications that can be potentially improved by our dataset.

F.1.1 Conservative topology. Designed to exclusively incorporate high-confidence links or only the measured links, the conservative topology could be utilized when assessing resiliency of the network. In particular, an improved understanding of the AS-level topology enables various stakeholders to predict how the network might react to local failures (*e.g.*, disruptions at interconnection facilities or BGP sessions going down). By identifying weak points, they can strategically guide improvements in the infrastructure and mitigate the risks associated with such failures globally. Similarly, a

comprehensive AS-level map is instrumental in delineating the Internet’s potential “attack surface” from the perspective of a provider. This would facilitate the prediction, detection and mitigation of large-scale security threats, such as BGP hijacking and fail-over.

F.1.2 Loose topology. The loose topology incorporates all inferred links regardless of their confidence level and can be used for applications that revolve around coverage and compliance. In particular, in recent years, there has been interest in improving vantage points deployment across the Internet [73, 145]. Our methodology could help guiding the placement of new vantage points to fill gaps in our topological understanding, allowing to both verify that those fractions of missing topology that we inferred are correct and provide more visibility onto the Internet. In parallel, governments have shown their interests in more actively monitor and track compliance across the Internet [20, 42, 53, 134]. The government of the Netherlands have for example required all of its services to be hosted on RPKI compliant ASes [154]. Our dataset can be used for targeted auditing of network fraction to ensure legal compliance.

F.1.3 Balanced topology. Lastly, the Balanced topology, created by selecting the threshold that maximize the F -score, lies between the conservative and loose models and supports applications that require guarantees over a spectrum of settings. This view would enable the generation of “what-if” scenarios, allowing for the prediction of user impacts based on new BGP announcements and the estimation of catchment areas. This dataset could serve as a reliable ground for designing and evaluating new routing protocols, ensuring expected behaviors align with the different topology versions at hand. For network planning, ISPs and businesses can harness insights from this balanced approach to strategically orchestrate their networks and peering relationships, thereby enhancing both performance and service reliability. The balanced topology could be utilized to reverse-engineer network routing policies can potentially lead to insights into operators’ traffic routing practices, thus stimulating a reevaluation of the Gao-Rexford model’s validity. Ultimately, it can pave the way for recognizing more faithfully quasi-routers and improve our understanding of the Internet [110].

F.2 Using Shapley to Understand the Inferences.

F.2.1 How to use Shapley to interpret what metAScritic learns. Figure 13 illustrates how the Shapley beeswarm plot can be used to interpret the contributions of different features to individual ratings. The beeswarm plot visually represents the impact of each feature across all predictions. A positive or negative Shapley (SHAP) value means that the feature is pushing for the existence or the non-existence of a link, while the color of each dot represents the feature’s actual value in the data. . For instance, if red dots are consistently positioned further to the right, like for the “# of Existing Links” feature, it means higher values of that feature push for the existence of the link. By examining the Shapley values of all features, we can determine whether the most influential features in the model’s predictions are contextually relevant, which helps assess whether the recommender system of metAScritic is effectively capturing meaningful network properties. We discuss the findings below:

- (1) **Most important features:** The most influential features, as indicated by their position on the x -axis and their impact on model output, are the number of existing links (# of Existing Links) and the number of non-existing links (# of Non-Existing Links). These features have the most significant positive or negative SHAP values, indicating that they strongly drive the model’s predictions.
- (2) **Geographic features:** Features related to geography and topology, such as Overlapping City, Overlapping Country, and Overlapping Facility, also play a crucial role in the model’s decision-making. The positive or negative SHAP values suggest that the presence or absence of overlap in the different granularity impacts the model’s confidence in the inferred links.
- (3) **AS specific characteristics:** Characteristics like Eyeballs, ASN, # in Customer Cone, Peering Policy and Traffic Profile also contribute to the model’s inferences, although their influence is less pronounced compared to the link-based features. These characteristics indicate how the model learns specific patterns associated with different ASes and their impact on the final predictions.
- (4) **IXP Overlap:** While overlapping IXP has minimal impact on the model’s output, this does not imply that IXP overlap carries no information. Instead, it suggests that our model relies on other features to make its inferences, potentially because an overlap in an IXP might be too coarse to accurately capture the nuances of peering connectivity.

F.2.2 How to use Shapley to understand an inference . We present an example of how to relate metAScritic’s predicted rating to the networking behavior of the two ASes involved. Figure 14 shows the most important features that were used to infer the link between AS 42 (PCH) and AS 138466 (Datamossa) in Sydney. Looking at the feature with the highest contribution, we see that the number of non-existing links for AS 42 and the number of existing links for AS 138466, before completion, which is 0 and 199 are the most informative features. They both contribute positively in favor of the existence of the link. Given PCH’s well-documented open peering policy [3], and the important number of links for AS 138466 suggesting a similar policy, the reliability of metAScritic’s output is reinforced. In that case, there is no reason to not trust metAScritic’s output This kind of reasoning can help a user to decide whether they want to trust the result returned by metAScritic.

F.3 Precision-Recall trade-off

Increasing λ to increase confidence: In Figure 15, we compute the recall and precision of links inferred in the six metros where we ran metAScritic. For each threshold, we use the recall and precision results across metros to estimate the 95th percentile confidence interval of the recall and precision at that threshold [21]. We observe a consistent and monotonic trend: increasing the threshold raises the precision at the expense of the recall, with a threshold of 0.3 maximizing the F -score. Edges inferred at a 0.9 threshold have a 97-99% likelihood of being correct. These inferred edges, only within 6 metros, represent over 226K unseen peering links in the topology, or 0.7× the peering links currently found in the *whole* CAIDA AS relationship dataset [94].

	Tier	Hypergiant	LargeISP	Content	Enterprise	Stub	Transit	ASes	Links/AS	% Increase
Tier	665 + 123	475 + 295	5379 + 1557	5841 + 1713	2340 + 550	31704 + 1623	22876 + 4717	51	209.8	15.3
Hypergiant	475 + 295	5 + 131	838 + 1797	397 + 1886	171 + 626	1176 + 2265	2949 + 5491	21	601.0	209.8
LargeISP	5379 + 1557	838 + 1797	11135 + 5536	8371 + 8060	2885 + 2098	42025 + 9927	47448 + 26870	2446	25.1	47.5
Content	5841 + 1713	397 + 1886	8371 + 8060	4096 + 5199	2993 + 3144	13441 + 9682	41559 + 28357	2041	31.0	78.3
Enterprise	2340 + 550	171 + 626	2885 + 2098	2993 + 3144	528 + 537	4742 + 3102	15220 + 8242	1426	13.2	64.1
Stub	31704 + 1623	1176 + 2265	42025 + 9927	13441 + 9682	4742 + 3102	6162 + 6794	122770 + 29680	61249	1.1	30.6
Transit	22876 + 4717	2949 + 5491	47448 + 26870	41559 + 28357	15220 + 8242	122770 + 29680	98712 + 40658	9573	19.3	41.0
Links Added	10578	12491	55845	58041	18299	63073	144015			

Table 5: Number of additional measured and inferred links by metAscritic compared to the public BGP view for different AS class pairs.

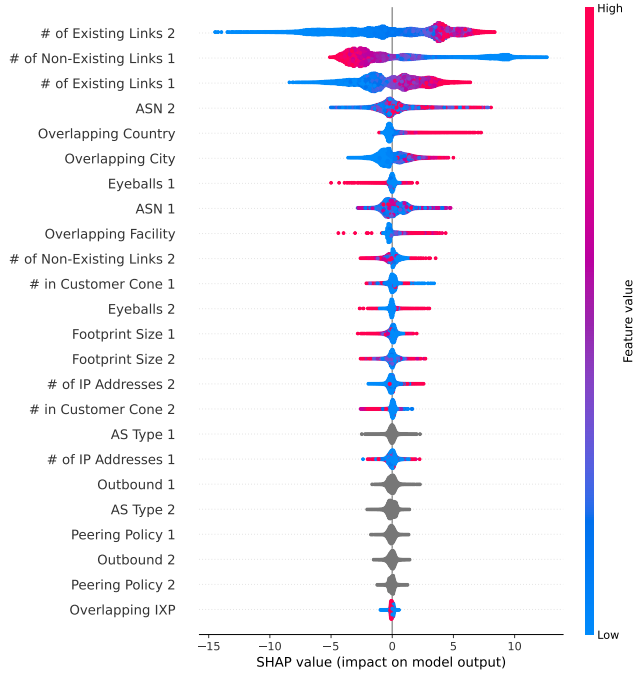


Figure 13: Shapley Summary Plot depicting the individual contributions of each feature to the ratings predicted by metAscritic in Sydney. Each point represents the impact of a specific feature on a particular rating, with the color denoting the magnitude of the feature.

F.4 Implementation

We implement our system in Python in around 800K lines. Our system first relies on EasymapIT¹² to fetch all of the inferred links from RIPE Atlas and CAIDA Ark within a week of interest. We have built several scripts to automate the collection of all the data described in Appendix C. We augment this pipeline to only consider traceroute crossing at least two ASes in the metro of interest. We also map every IP addresses observed in our dataset to their associated rDNS. At the end of this process, we get a set of traceroutes that cross at least 2 ASes of interest. We then implement a script to

¹²<https://github.com/cunha/easymapit>

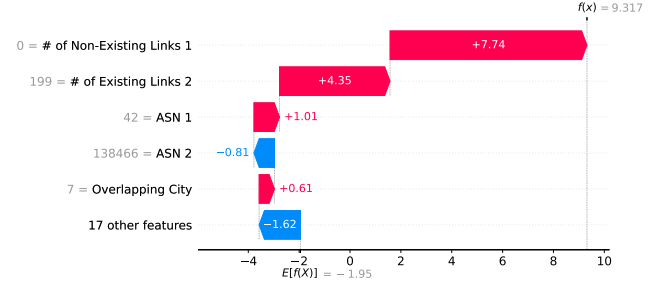


Figure 14: Shapley Force Plot showing the effect of the different features in the metAscritic's decision for the existence of the link between AS 42 (PCH) and AS 138466 (Datamossa).

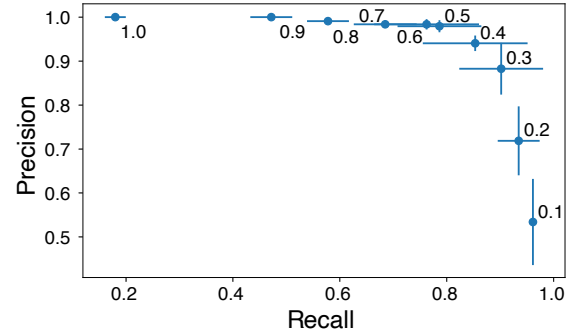


Figure 15: Scatter plot highlighting the relationship between the recall and precision and the threshold in each of the metros considered. The cross around the point corresponds to the 95% percentile confidence interval around the points.

geolocate all the interconnections that do not possess informative rDNS entries.

We leverage Apache Airflow [8] to architect the different functionalities of the system. For the measurement, we build a script, on top of RIPE API, that implements the algorithm that identify the right set of traceroutes to run. At the end of this process, we pass an Airflow query that starts the process of fetching all the measurements and geolocating the eventual unknown interconnections. Airflow then monitors when this process ends and either launches

the completion if all the viable candidates have been probed. Otherwise, Airflow goes back to the measurement phase. By the end of this process, Airflow activates a function that decides whether the final recommendation is ready to be launched. Otherwise, the system gets back to the measurement phase. The system run-time is conditioned by the number of measurements that can be performed simultaneously in a batch. We implemented the whole system in an AWS m6a.4xlarge EC2 VM.

G Impact on Internet topology

We characterize the links measured and inferred by metAScritic, which could benefit any system or tool that uses topology information. Table 5 shows the number of links in the BGP public view and the number of additional links inferred by metAScritic combining our measurements across all six metros. We classify ASes by type (see Appendix D.3 for details). On the rightmost column, we see that metAScritic infers a significant number of additional links, particularly for hypergiants and content providers, which peer aggressively and are underrepresented in the BGP public view. Conversely, metAScritic infers relatively fewer links for Tier-1/2 and stub networks, as these have mostly customer-provider links better captured in the BGP public view [118].

Figure 16 shows the number of links measured and inferred for each metro in our evaluation. We find that metAScritic needs a small number of measurements (patterned areas) compared to the total number of inferred links, which is aligned with our goals of overcoming practical limitations imposed by limited probing budgets and lack of vantage points to observe all links. We order metros on the x axis by their number of links. For each metro, we classify links as *existing* if they have been measured or inferred in one of the previous metros (to the left). The small number of *existing* links indicates that collecting measurements from diverse locations significantly contributes to complete our view of the topology and motivates metAScritic’s localized approach. For the remaining *new* links, we also identify the ones where *both* ASes are present at a previous metro, and have thus already been probed by metAScritic in previous iterations. Finally, links found between previously-probed ASes could occur either because ASes have different connectivity at different locations or failure to measure or infer a link in one of the locations. In either case, probing at different locations has benefits and improves our view of connectivity. These findings may be explained by route diversity due to varying routing policies across different continents as well as vantage point availability and placement.

H Validation datasets

Drawing from previous work, we assemble an extensive validation dataset. In this section, we rank the observed links based on our confidence in their accuracy.

Confirmation from operators: The most reliable method for validating our inferences involves obtaining direct confirmation from the organizations involved. However, operators often lack incentives to disclose this information, and those who are willing to validate our inferences are usually operators already engaged with the research community. These operators are also more likely

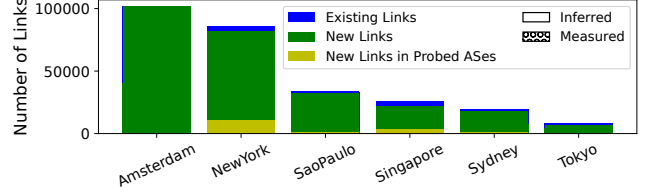


Figure 16: Number of links measured and inferred for each metro in our evaluation. We order metros by number of ASes, and classify a link as *existing* if it has been measured or inferred in a metro with more ASes (to the left). Among the *new* links, we differentiate those belonging to ASes previously *unprobed* (green).

to host RIPE Atlas probes, which already provide them with high visibility through measurements alone.

BGP communities: We use prior techniques based on BGP communities to identify where announcements have been learned. Using CAIDA’s manually labeled dataset of communities [27, 86], we can identify the interconnections geolocated in a specific location from publicly available BGP feeds. For each observed path from a feed, we match the communities with ASes present in the path. If we find a community associated with geographic information, we infer that the AS stamped the path with the community when it received the path from the previous AS in the path. For instance, if we observe [AS1, AS2, AS3] and notice a community stamped from AS2 in location v , we infer that AS2 and AS3 are connected in location v . This technique is reliable but limited in coverage, as most ASes remove communities as they cross their network, and communities lack standardized meaning, requiring considerable manual labor to decipher their meaning.

Openly accessible looking glasses: Another validation source previously employed is LGs. Although they can provide unique visibility into an operator’s view of the Internet, we found that only a small subset of LGs actually provide a complete view of their routing table. Most LGs only allow querying the routing table for a given prefix instead. We devise a strategy that queries all of the prefixes in v announced by the AS that we have inferred to be peers of the AS hosting the LGs to confirm our inference. Unfortunately, our experience confirms that LGs are not intended for large-scale measurement, making the process of extensive probing very complicated in practice. Even under optimal conditions, where we can query them as much as we want, they can still miss many less-preferred paths from other border routers. As a result, it is challenging to know precisely when an inferred link is incorrect unless the LG is hosted in the exact metro of interest. For all of those reasons, LGs are useful tools for assessing true positive rate (*i.e.*, correctly inferred links), but weaker when considering false negative (*i.e.*, incorrectly discarded links).

Extensive measurements: While the other datasets allow to study metAScritic in its capacity to recover links that are completely invisible from any vantage points, one of the key assumption on which relies metAScritic is that conducting exhaustive measurements from all vantage points generates redundancy, and that the

same information could be obtained using significantly fewer measurements with our completion mechanism. To validate this hypothesis, we carried out an extensive measurement campaign for two metros, running over a few million traceroutes between the majority of probes and available destinations. Due to the impracticality of executing such a large volume of measurements, we limited our scope to two metros. Next, we compared the links inferred by metAScritic to those measured through the comprehensive measurement campaign. Our hope is that metAScritic reduces drastically the number of measurements without sacrificing the accuracy of the inferences.

iGDB: Another source of validation relies on the geographic footprint inferred from iGDB. In particular, we query the database to find all the ASes whose geographic footprint only overlaps in v . For this set, we deduce that the interconnections occur in v . This technique assumes the database is complete, which is difficult to verify. We observe that, in general, ASes that do provide information are likely to provide correct information, and the sources on which iGDB builds are likely to be audited and updated regularly (e.g., IXP presences are likely to be up to date on IXP's websites, certain Cloud Providers requesting for up-to-date PeeringDB information...). This set of geolocated interconnections is limited by the completeness of iGDB.

Train/Test split: We explore several data split strategies to assess the performance of our completion across diverse scenarios. Our focus lies on three distinct partitioning methodologies that exemplify different facets of our algorithm's performance. The first split we consider is a **stratified** split where we discard 20% of all the entries from each row. This partitioning serves as the standard splitting considered in our work as for real-world measurements, most rows will possess a few visible entries, due to the symmetry of interconnections resulting in sources with probes hosted in the (AS, metro) to successfully run informative measurements from their end. Another crucial scenario to consider is the *completely-out* split, where random rows are fully discarded until 20% of the visible entries of $\mathbb{M}(v)$ have been discarded. This scenario aligns with extreme instances where we fail to obtain any informative measurement due to the lack of good sources or targets. When investigating such cases, the only source of usable information for predicting the interconnection is side information pertaining to the ASes. We finally consider a *random* split which serves as the baseline methodology considered in other works in standard recommender system literature.