

Cross-Domain Fault Localization: A Case for a Graph Digest Approach

William D. Fischer
Naval Postgraduate School
Email: wdfisch1@nps.edu

Geoffrey G. Xie
Naval Postgraduate School
Email: xie@nps.edu

Joel D. Young
Naval Postgraduate School
Email: jdyoung@nps.edu

Abstract—Prior research has focused on intra-domain fault localization leaving the cross-domain problem largely undressed. Faults often have widespread effects, which if correlated, could significantly improve fault localization. Past efforts rely on probing techniques or assume hierarchical domain structures; however, administrators are often unwilling to share network structure and state and domains are organized and connected in complex ways. We present an inference-graph-digest based formulation of the problem. The formulation not only explicitly models the inference accuracy and privacy requirements for discussing and reasoning over cross-domain problems, but also facilitates the re-use of existing fault localization algorithms while enforcing domain privacy policies. We demonstrate our formulation by deriving a cross-domain version of SHRINK, a recent probabilistic fault localization strategy.

I. INTRODUCTION

Cross-domain, multi-domain, and inter-domain fault localization are synonymous terms that describe determining the root cause of a network failure whose effects propagate across administrative domains. When data is required from more than one domain to isolate a fault, a cross-domain solution is needed. A study of routing instability found that all parties pointed to another party as the cause in about 10% of the problems [1].

Faults often have widespread effects, which if correlated, can significantly increase fault localization accuracy. We define *inference gain* to be the increase in inference accuracy achieved by correlating additional evidence. Cross-domain network failures can not always be localized without a coordinated effort between domains. As an example consider a scenario in which an operator makes a typo in the A record for a web service in an authoritative DNS server. The domain administrator may not be able to isolate the fault quickly, and may not even be aware that a problem exists for a period of time. While the fault remains unabated and potentially unnoticed, there may be observations external to the domain that could help detect and localize the fault.

Privacy, scalability, and interoperability issues hinder efforts to achieve accurate cross-domain fault localization. While prior work has stated the importance of these issues [1]–[4], to our knowledge there has been no formal definition of requirements addressing them. Network domain managers are often unwilling or not permitted to share detailed internal network architectures and quality-of-service issues with outside agencies, running face-first into the need to share data

to successfully troubleshoot networking issues. Automated techniques for finding faults across a large number of domains face serious computational issues and exact computation using belief networks is NP-hard [5]. Interoperability in network management and fault isolation techniques is a perennial problem: Different modeling techniques and tools using different algorithms will be employed in various domains. Conflict of information formats and semantics may arise between domains, with each domain’s model assigning a different value to the same parameter.

In this paper we characterize the problem space for cross-domain fault localization and propose an inference graph digest approach. A cross-domain approach must achieve acceptable accuracy while satisfying privacy concerns. We illustrate an approach whereby domain managers using causal graphs to model fault propagation can use a function to create a digest representation of their network state and dependencies to participate in a collaborative effort to localize a cross-domain fault by sharing observations. By addressing privacy, scalability, and interoperability issues with our graph digest approach, we attack the obstacles that prevent collaborative cross-domain fault localization. Although the discussion in this paper focuses on the formulation, and demonstrates the utility of the formulation by creating a causal graph digest for a probabilistic inference method, the concepts discussed in creating a graph digest could be extended to include other domain representations and inference methods.

In Section II we describe related work including Bayesian approaches to localize intra-domain faults and approaches addressing the cross-layer fault localization problem. In Section III we discuss the challenges and tradeoffs associated with using a graph digest and formulate an approach to perform cross-layer fault localization. Section IV illustrates the graph digest approach for a cross domain implementation of SHRINK. We describe our assumptions and initial efforts to define a graph digest.

II. NETWORK FAULT LOCALIZATION

As characterized by Steinder and Sethi, fault localization is the second step in fault diagnosis following *fault detection* and preceding *testing* [6]–[8]. Network administrators use fault localization techniques to discover best hypotheses explaining the observations detected in the fault detection step. Myriad

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE OCT 2008		2. REPORT TYPE N/A		3. DATES COVERED -	
4. TITLE AND SUBTITLE Cross-Domain Fault Localization: A Case for a Graph Digest Approach				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Naval Postgraduate School Operations Research Department Monterey, CA 93943				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release, distribution unlimited					
13. SUPPLEMENTARY NOTES in Proc. IEEE Third Workshop on Internet Network Management (INM'08) , Orlando, FL, USA, October 2008, The original document contains color images.					
14. ABSTRACT Prior research has focused on intra-domain fault localization leaving the cross-domain problem largely unaddressed. Faults often have widespread effects, which if correlated, could significantly improve fault localization. Past efforts rely on probing techniques or assume hierarchical domain structures; however, administrators are often unwilling to share network structure and state and domains are organized and connected in complex ways. We present an inference-graph-digest based formulation of the problem. The formulation not only explicitly models the inference accuracy and privacy requirements for discussing and reasoning over cross-domain problems, but also facilitates the re-use of existing fault localization algorithms while enforcing domain privacy policies. We demonstrate our formulation by deriving a cross-domain version of SHRINK, a recent probabilistic fault localization strategy.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT SAR	18. NUMBER OF PAGES 6	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

techniques have been developed for fault localization, including: rule-based systems, model-based systems, case-based systems, neural networks, decision trees, model traversing techniques, code-based techniques, Bayesian networks, dependency graphs, causal graphs, and phrase structured grammars [6].

A current trend attacks the problem by modeling network dependencies in a directed acyclic graph having root causes as parentless nodes, observations as childless nodes, and dependencies represented as directed edges in the graphs with uncertainties captured in conditional probability distributions associated with each node [5], [9], [10]. This graph structure is also known as a causal (or causality) graph [6]. Approaches typically perform probabilistic (Bayesian) inference on bipartite causal graphs [5], [10].

Root causes to network failure are also known as shared risk groups (SRGs) [3]. SRGs typically represent hardware components that can fail, impacting service for a set of dependent services or communication channels. In the graphical model, edges depicted from an SRG to each node directly influenced by the state of the SRG represent conditional dependencies. In bipartite causal graphs the edges only connect from SRG nodes to observation nodes allowing faster probabilistic inference as compared with general Bayesian networks.

The SCORE [9], SHRINK [5], and Sherlock [11] approaches form the state of the art for leveraging causal graphs for fault localization. SCORE uses a set covering approach for finding the best explanation (set of failed SRGs) for observed outages based on a bipartite graph. SHRINK enhances the model to allow probabilistic inference by attaching edge weights that are combined using the noisy-OR [12] model to form conditional probability tables for each observation node. Sherlock further extends these approaches with a multilevel causal graph.

Very little research is published on cross-domain fault localization. Probing and monitoring techniques can be leveraged to assist with collection of information about network state and structure. A cross-domain fault localization approach is presented by Steinder and Sethi [13] for hierarchically organized networks. This approach locates the source of an end-to-end service failure through distributed coordination between the domains along the path of the failure. In addition to an existing domain hierarchy, the approach relies on full knowledge of each end-to-end data path at the domain level.

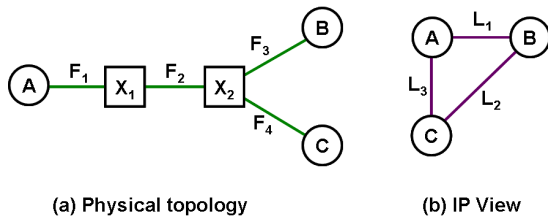


Fig. 1. Example network.

SHRINK [5] performs Bayesian inference on a bipartite causal graph. The SHRINK model assumes independent failures of root cause nodes and that no more than three SRGs

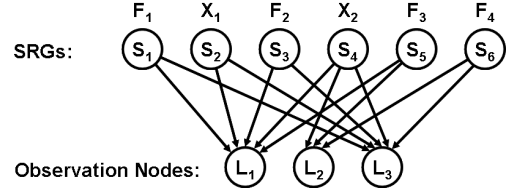


Fig. 2. Causal Graph. Noisy edges not depicted.

will fail simultaneously in a large network based on the extremely low likelihood of four or more simultaneous failures. Noisy-OR is used to calculate the conditional probability table for a node with multiple parents. The SHRINK algorithm is defined as follows. Let $\langle S_1, \dots, S_n \rangle$ denote a hypothesis vector, where $S_i = 1$ if a failure of SRG S_i is assumed, and $S_i = 0$ otherwise. Let $\langle L_1, \dots, L_m \rangle$ denote an observation vector, where $L_j = 1$ if a failure of L_j is observed, and $L_j = 0$ otherwise. Given a particular observation vector, the SHRINK algorithm searches through all hypothesis vectors with no more than three assumed failures, and returns those maximizing the posterior probability:

$$\underset{\langle S_1, \dots, S_n \rangle}{\operatorname{argmax}} \Pr(\langle S_1, \dots, S_n \rangle \mid \langle L_1, \dots, L_m \rangle) \quad (1)$$

Consider the simple scenario depicted in Figs. 1, and 2. Fig. 1(a) depicts the network physical topology, in which IP routers A, B, and C are connected across fibers $F_1 - F_4$ and optical cross-connects X_1 and X_2 . Each IP router has a link to each other router as shown in Fig. 1(b). If any of the optical components, fibers, or optical cross-connects fail, the IP routers will detect link failures. The prior SRG failure probabilities are 10^{-4} and 10^{-6} for the fibers and the cross-connect respectively.

The causal graph (Fig. 2) has six optical components mapped to SRGs $S_1 - S_6$. To account for potential database and observation errors a noise value (10^{-4}) is subtracted from the conditional probability of each edge, and noisy edges with this value are added to form a complete bipartite graph. E.g., $\Pr(L_1|S_1)$ is 0.9999 while $\Pr(L_2|S_1) = 10^{-4}$.

Suppose L_1 and L_2 are down, and L_3 is up. Intuitively, the cause is most likely the failure of fiber link F_3 . As described above, SHRINK only considers hypothesis vectors with at most three total assumed failures. For this six SRG example SHRINK searches through $\sum_{k=0}^3 \binom{6}{k} = 42$ hypotheses, with hypothesis vector $\langle 0, 0, 0, 0, 1, 0 \rangle$ maximizing the posterior probability for the given observations. SHRINK correctly identifies SRG S_5 (i.e., the failure of fiber link F_3), to be the root cause.

III. FORMULATION OF GRAPH DIGEST APPROACH

In this section, we present a formulation for cross-domain fault localization based on information-preserving transformations of intra-domain inference graphs. We propose a set of criteria to explicitly define the two primary requirements of cross-domain fault localization: *preservation of inference gain* and *protection of privacy*. Finally, we discuss the main

technical challenges for deriving practical algorithms from the proposed formulation.

A. General Framework

As discussed above, recent intra-domain approaches use graphical models to model dependencies between aspects of network operation, particularly the causal relationships between hardware failures and observed anomalies. These models (also called inference graphs), enable inference algorithms to determine those failure scenarios best explaining observed anomalies. In practice, faults often propagate across network domain boundaries, depriving intra-domain algorithms of critical information required for accurate inference. We address the problem by sharing summarized intra-domain models (called *graph digests*) between domains. A graph digest captures cross-domain dependencies while hiding internal details.

Our approach is based on, and designed to address, the problems that arise from the following assumptions: domains are administratively separated, domain managers are unwilling to reveal their internal network structures and associated inference graphs, and finally domain managers *are* willing to collaborate to localize faults if their domain's internal details are hidden. We believe these assumptions are fundamental constraints all general cross-domain approaches must address.

A cross-domain inference model based on graph digests is defined as follows. Consider n network domains:

- G_i is the inference graph for the i_{th} domain.
- f is (ideally) a one-way function on G_i implementing a privacy policy. $f(G_i)$ is called the *inference graph digest*, or simply *digest*, for G_i .
- $\mathcal{G}^j = \left(\biguplus_{i \neq j}^n f(G_i) \right) \uplus G_j$, where j is a domain performing cross-domain inference and $\uplus$ is a model-specific union. $\mathcal{G}^j$ is the cross-domain model integrating the digests from all the other domains with domain j 's undigested graph. Now, domain j may use an existing algorithm such as SHRINK to perform inference over $\mathcal{G}^j$.

Before a practical graph digest approach can be implemented, interoperability standards must be developed. Domains using different inference methods can potentially use a digest approach if standards are implemented and adhered to. Items to be standardized include data types and attributes as well as cross-domain management structures such as centralized, distributed, iterative, etc. We define a *shared attribute* as a physical entity or logical concept modeled in two or more fault propagation causal graphs, and that has the same semantics in each graph. In order to create a domain digest to connect to another domain's fault propagation causal graph, shared attributes must be identified and agreed upon.

B. Modeling Preservation of Inference Gain

The function f above is useless if the digest it produces is not useful for inference. A digest function (transformation) is *inference preserving* if it maintains enough structure to allow successful inference. Ideally, we achieve the same inference gain using digests versus undigested graphs.

Let B_u and B_d be the best hypotheses produced using undigested graphs and graph digests respectively. B_u and B_d are sets of potential causes. percentage of elements in a hypothesis that are consistent with the observations, and *coverage ratio* c that measuring the percentage of all of the observations that a hypothesis can explain:¹

$$h = \frac{|B_d \cap B_u|}{|B_d|}, \quad c = \frac{|B_d \cap B_u|}{|B_u|}.$$

It is clear that $1 \geq h, c \geq 0$. The ratios of false positives and false negatives are $1 - h$ and $1 - c$ respectively, both relative to B_u . The ratios h and c can each be easily optimized at the expense of the other, which may be overcome by computing the harmonic mean of the two values.² We propose to use the harmonic mean α as the criterion to measure how well a digest model preserves inference gain:

$$\alpha = \begin{cases} 0 & \text{if } h = c = 0 \\ \frac{2 \cdot h \cdot c}{h + c} & \text{otherwise} \end{cases} \quad (2)$$

Ideally α is one, when both h and c equal one. The definition of α can be generalized for the case where more than one best explanations are derived by the inference algorithm.

C. Modeling Protection of Privacy

We define a *sensitive property* as a detail the domain manager considers private. Ideally, a graph digest should not help to reveal any sensitive properties. Specific sensitive properties will vary between domains and may include bottlenecks, customer information, peering agreements, the number of failed components, and many other characteristics. Furthermore, a set of digests from a domain collected over time should not aid in deriving the sensitive properties from the original graph.

Shannon says that "perfect secrecy" is defined to be when the *a priori* probability is equal to the *posterior* probability for message traffic deciphering by an adversary [16]. The same concept can be applied as a criterion for inference graph privacy. One has to assume that an adversary has some domain knowledge, has passive access to externally observable information, and can infer some level of knowledge about a distribution over time. Using an information theoretic approach, the relative entropy (Kullback Leibler distance) [17], between the probability distribution of the sensitive property without a digest and the probability distribution after receiving a digest measures the privacy loss due to sharing a digest. Let s represent a sensitive property in a domain conditioned by the adversary's knowledge, where $s|d$ represents the property further conditioned by a digest. Let X represent the set of possible values for s . The relative entropy equation is:

$$KL(s|d, s) = \sum_{x \in X} Pr(s = x|d) \log_2 \frac{Pr(s = x|d)}{Pr(s = x)} \quad (3)$$

Ideally this distance will equal zero for each sensitive property in a domain, meaning that the information about a

¹In AI [14] these ratios are known as *precision* and *recall* respectively

²The harmonic mean of precision and recall is also known as F Score [15]

sensitive property is unchanged after receiving a digest. Even if the entropy is reduced for a sensitive property, the entropy of $s|d$ may remain sufficiently high to protect the privacy of the property. Ultimately, the resultant entropy of $s|d$ and not the amount of entropy lost, indicates the level of privacy protection for a sensitive property.

Unfortunately, deriving accurate probability distributions about a sensitive property in a domain, particularly from an adversary's perspective, may not be possible. An ontology of sensitive properties with privacy protection implementation methods is needed. If prior and posterior probability distributions can be derived for a sensitive property, we suggest that the relative entropy Eq. (3) be used to evaluate the privacy protection for the property. For sensitive properties that are not conducive to evaluation with probability distributions, techniques to protect properties against obvious attacks should be implemented. As an example, reachability information may be difficult to hide with a relative entropy approach, but has a well established method for evaluation: the transitive closure of an adjacency matrix determines the reachability between any two components in a network. A digest must prevent this obvious attack method by denying adjacency information that could be used to establish reachability. A function f could be augmented to break all paths between nodes that would reveal the sensitive property (or add nodes to hide a lack of reachability). We evaluate privacy against the transitive closure attack for reachability in Section IV.C.

IV. ILLUSTRATION OF DIGEST APPROACH

In this section, we demonstrate the utility of our formulation by showing how a digest for a bipartite causal graph can be created to enable the use of SHRINK for cross-domain fault localization. We selected SHRINK to illustrate the ideas presented in Section III due to its robustness and simplicity. We first present an algorithm to create a digest for inference by SHRINK. Next, we describe a hypothetical cross-domain scenario to illustrate the possible steps of creating such a digest. Finally, we provide a brief analysis of inference preservation and privacy protection for the scenario.

A. Bipartite causal graph digest creation

A main challenge for implementing the digest approach is to find digest creation algorithms that to be practical, must be general while at the same time meeting a wide range of domain specific information hiding needs. We present one such algorithm, *createBipartiteDigest* (Fig. 3), that although a bit naive, embodies the concepts presented in Section III. We crafted *createBipartiteDigest* to establish a cross-domain extension to SHRINK, and thus the assumptions previously presented for SHRINK in Section II apply to this algorithm as well. The algorithm executes in a sequential four step process: (1) pruning (lines 3,4), (2) partial evaluation (lines 5-8), (3) aggregation (lines 10-12), then (4) renaming (lines 13,14). The pruning step removes any SRG nodes that have no non-noisy edges to observation nodes reporting failure. Except for highly unlikely cases, these nodes will have a very low score and will

createBipartiteDigest(G)

- 1: Add node L_{new} to G
- 2: **for** all $SRG S_i \in G$
- 3: **if** (for all edges $(S_i, L_j) \in G$, L_j is up)
- 4: **then** Prune S_i and its edges (S_i, L_j)
- 5: **else**
- 6: Collect edges $(S_i, L_j) \in G$ such that L_j is up
- 7: Add edge (S_i, L_{new}) using Eq. (4) on collected edges
- 8: Prune collected edges (S_i, L_j)
- 9: Remove all isolated observation nodes L_i
- 10: **for** all $SRG S_x, S_y \in G$
- 11: **if** S_x and S_y are indistinguishable
- 12: Aggregate S_x and S_y into S'_x such that $S'_x = S_x \cup S_y$
- 13: Rename all SRGs that are not shared attributes
- 14: Rename all Observation nodes other than L_{new}

Fig. 3. Algorithm for computing a digest from a bipartite causal graph G .

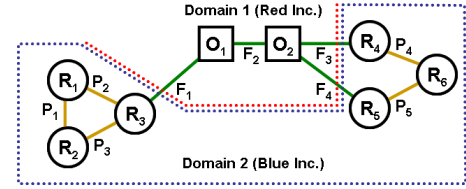


Fig. 4. Physical Topology of considered scenario.

not appear on a list of best explanations. The partial evaluation step uses noisy-OR to combine all edges from an SRG S_i that point to k observation nodes $L_1, L_2, \dots, L_k$ reporting liveness into a single node L_{new} . The noisy-OR equation to compute the new edge weight is:

$$Pr(L_{new}|S_i) = 1 - \prod_{j=1}^k 1 - Pr(L_j|S_i) \quad (4)$$

The aggregation step of the algorithm combines SRGs that have the same prior probabilities and edges. These SRGs will have identical scores on a list of best explanations. Aggregation of SRGs means that one SRG represents a set of SRGs. The final step, renaming, is simply assigning a new label on each node in the resultant graph, except for shared attributes and L_{new} .

B. Cross-domain scenario

Three months ago Blue Inc. (Domain 2) started a lease for three optical circuits across the optical mesh provided by Red Inc. (Domain 1), a large provider with many customers. The physical view of the overlap between Blue and Red is depicted in Fig. 4 and the view of provisioning in Fig. 5. Redundant components in the mesh (not portrayed) provided by Red were

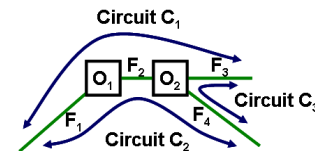
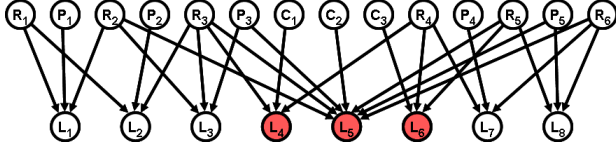
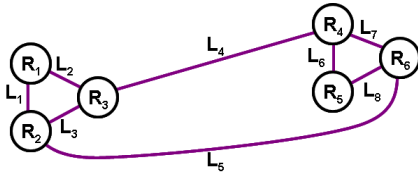


Fig. 5. View of leased circuits provisioned to Domain 2.



tested to ensure that Blue could transit the mesh in the event of a failover. Two months ago Green Inc. subscribed to a number of circuits across the Red mesh. Through oversight, an older image was used to add the circuits for Green Inc. to the tables in the backup Optical Digital Cross Connect (ODCX) O_2 . The backup ODCX O_2 has no records honoring the leased circuits for Blue. This morning, O_2 went offline and the backup component came online, severing connectivity for Blue across the mesh.

Conducting inference in isolation, neither Red nor Blue is able to isolate the problem. From the perspective of Red, all tools show a healthy network status and the circuits show liveness. From the perspective of Blue, no traffic can cross the leased circuits.

Fig. 6 portrays the IP link connectivity for Blue. The administrators at Blue use the SHRINK algorithm for fault localization, and Fig. 7 reflects the graph with the three failed IP links highlighted. The optical circuits are shared attributes, which Blue has modeled as SRGs $C_1 - C_3$. Not knowing how the optical mesh is configured, Blue has assigned a uniform prior probability of failure at 10^{-5} for each SRG in the graph. Inference for the best explanation returns $\{R_4\}$, $\{R_5\}$, $\{C_1\}$, $\{C_2\}$, and $\{C_3\}$ as equally likely.

Red agrees to perform probabilistic inference on a combined inference graph using a digest from Blue. Blue has one sensitive property to hide from Red: the internal reachability between R_4 and R_5 . Blue is interested in hiding whether

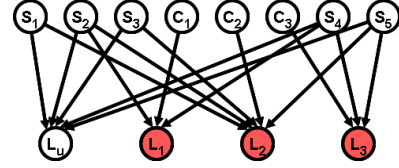
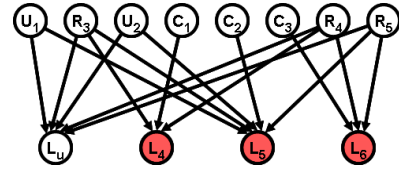
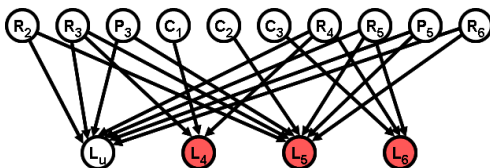
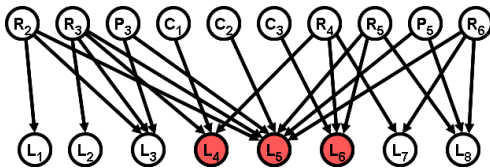


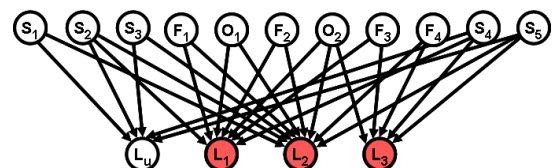
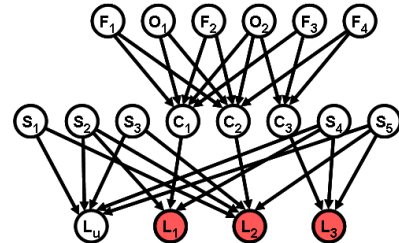
Fig. 10. Domain 2 reflecting aggregation of R_2, R_6 into U_1 , and P_3, P_5 into U_2 .

they are able to transmit data between R_4 and R_5 if Circuits $C_1 - C_3$ fail. The digest construction proceeds with pruning (Fig. 8), partial evaluation (Fig. 9), aggregation (Fig. 10), and renaming (Fig. 11).

Red creates a union of their graph with the digest from Blue. Red also uses SHRINK for inference, so the combined graph must be converted to a bipartite graph. The provided circuits C_1 , C_2 , and C_3 are logical in nature and although they are needed to connect the graphs, they are not needed in the final graph for inference. Each edge in the Red portion of the graph that is directed to one of the circuit nodes is redirected to each observation node that the circuit node is directed to, and all of the shared attribute nodes are then pruned from the graph. The final causal graph is shown in Fig. 13. Red uniformly assigns 10^{-5} prior probability to each SRG. SHRINK inference run on the graph returns $B_d = \{O_2\}$ as the best explanation, with a score significantly higher than the other hypotheses. Convinced that the lost connectivity for Blue is most likely caused by O_2 , Red proceeds to diagnose the ODCX and restore service to Red.

C. Analysis of criteria

Running SHRINK on the union of the Red graph with the undigested Blue graph returns $B_u = \{O_2\}$ as the best explana-



tion with a score significantly higher than the other hypotheses. The hit ratio h and coverage ratio c are both 1.0, resulting in a perfect inference preservation score ($\alpha = 1.0$). To further illustrate the α criterion, suppose instead the inference results received were $B_d = \{O_2, F_3, S_4\}$ and $B_u = \{O_2, F_1\}$. For this scenario $h = 0.33$, $c = 0.5$, and α would be 0.4.

To evaluate privacy protection, we first need a technique to create an adjacency matrix using the SHRINK inference graph. Since the privacy hiding goal is based on hiding reachability between hardware components, we populate the elements of the adjacency matrix with the SRGs. For each SRG, all of the parents of the observation nodes that the SRG has an edge to are considered adjacent to the SRG. For example, R_1 in Fig. 7 is adjacent to R_1 , P_1 , and R_2 via L_1 ; and R_1 , P_2 , and R_3 via L_2 . In building the adjacency matrix for the digest, we can't use L_u to establish adjacency since the parent nodes of L_u don't necessarily reach each other. We can't use any observation nodes for which $C_1 - C_3$ is a parent since Blue wants to hide reachability between R_4 and R_5 in the event of failure across the optical mesh. Since there are no observation nodes available to establish adjacency, privacy protection for the sensitive property is trivially satisfied.

Suppose that the digest for Blue included the IP links L_7 and L_8 as down observation nodes. By the algorithm (Fig. 3), these nodes and all SRGs that can affect them will appear in the digest. Constructing the adjacency matrix as above and computing the transitive closure will clearly reveal the sensitive property. The algorithm may be refined so that the SRG reachable by both R_4 and R_5 in the closure that has the lowest inference result using the undigested graph can be pruned from the digest. This process can be repeated and checked until R_4 and R_5 are no longer reachable in the closure matrix. Clearly accuracy may suffer from such a heuristic, further highlighting the tension between accuracy and privacy.

V. CONCLUSION

Network faults often have observable effects in multiple domains. This paper demonstrated that cross-domain fault localization, by correlating the observations from different domains, has the potential to significantly increase the accuracy of network fault localization. It also articulated the main challenges to realize the inference gain, particularly the privacy consideration. The main contribution is an inference-graph-digest based formulation of the problem. The formulation not only explicitly models the inference accuracy and information hiding requirements, but also facilitates the re-use of existing fault localization algorithms without compromising each domain's information hiding policy.

To move forward certainly requires a fundamental understanding of the issues that is beyond the formulation and scenario described in this paper. Is the graph digest approach applicable to a wide range of network scenarios? What about scenarios involving more than two domains? Does there exist a general, yet easily calculable metric for quantifying the highly domain-specific information hiding policy? How should observation errors and graph model inaccuracies be detected

and controlled? We believe these and other similar questions constitute a new area of networking research which may have a major impact on how we trouble-shoot network faults.

ACKNOWLEDGMENTS

We thank the anonymous reviewers and our shepherd Keisuke Ishibashi for their constructive comments. Srikanth Kandula provided valuable input for our background research. This research was partially sponsored by the NSF under the grants CNS-0520210 and CNS-0721574. Views and conclusions contained in this document are those of the authors and should not be interpreted as representing the official policies, either expressed or implied, of NSF, or the U.S. government.

REFERENCES

- [1] D. G. Thaler and C. V. Ravishanker, "An architecture for inter-domain troubleshooting," *Journal of Network and Systems Management*, vol. 12, no. 2, pp. 155–189, 2004.
- [2] X. Huang, S. Zou, W. Wang, and S. Cheng, "Mdfm: Multi-domain fault management for internet services," *Management of Multimedia Networks and Services: 8th International Conference on Management of Multimedia Networks and Services, MMNS 2005, Barcelona, Spain, October 24-26, 2005: Proceedings*, 2005.
- [3] D. Larrabeiti, R. Romeral, I. Soto, M. Uruena, T. Cinkler, J. Szigeti, and J. Tapolcai, "Multi-domain issues of resilience," *Transparent Optical Networks, 2005, Proceedings of 2005 7th International Conference*, vol. 1, 2005.
- [4] P. Reynolds, J. L. Wiener, J. C. Mogul, M. K. Aguilera, and A. Vahdat, "Wap5: black-box performance debugging for wide-area systems," *Proceedings of the 15th international conference on World Wide Web*, pp. 347–356, 2006.
- [5] S. Kandula, D. Katabi, and J. P. Vasseur, "Shrink: a tool for failure diagnosis in ip networks," *Applications, Technologies, Architectures, and Protocols for Computer Communication*, pp. 173–178, 2005.
- [6] M. Steinder and A. S. Sethi, "A survey of fault localization techniques in computer networks," *Science of Computer Programming*, vol. 53, no. 2, pp. 165–194, 2004.
- [7] I. Katzela, A. Bouloutas, and S. Calo, "Centralized vs. distributed fault localization," *Proceedings of the fourth international symposium on Integrated network management IV table of contents*, pp. 250–261, 1995.
- [8] M. Y. Chen, A. Accardi, E. Kiciman, J. Lloyd, D. Patterson, A. Fox, and E. Brewer, "Path-based failure and evolution management," *Proceedings of the 1st conference on Symposium on Networked Systems Design and Implementation-Volume 1 table of contents*, pp. 23–23, 2004.
- [9] R. R. Kompella, A. Greenberg, J. Rexford, A. C. Snoeren, and J. Yates, "Cross-layer visibility as a service," *Proc. of fourth workshop on Hot Topics in Networks (HotNet-IV)*, 2005.
- [10] M. Steinder and A. Sethi, "Probabilistic fault diagnosis in communication systems through incremental hypothesis updating," *Computer Networks*, vol. 45, no. 4, pp. 537–562, 2004.
- [11] P. Bahl, R. Chandra, A. Greenberg, D. A. Maltz, and M. Zhang, "Towards highly reliable enterprise network services via inference of multi-level dependencies," *Proceedings of the 2007 conference on Applications, technologies, architectures, and protocols for computer communications*, pp. 13–24, 2007.
- [12] J. Pearl, *Probabilistic Reasoning in Intelligent Systems: Networks of Plausible Inference*. Morgan Kaufmann, 1988.
- [13] M. Steinder and A. S. Sethi, "Multidomain diagnosis of end-to-end service failures in hierarchically routed networks," *IEEE Transactions on Parallel and Distributed Systems*, vol. 19, no. 1, pp. 126–144, 2008.
- [14] S. J. Russel and P. Norvig, *Artificial intelligence*. Prentice-Hall, 2003.
- [15] R. Baeza-Yates and B. Ribeiro-Neto, *Modern information retrieval*. Addison-Wesley Longman, 1999.
- [16] C. E. Shannon, "Communication theory of secrecy systems," *Bell System Technical Journal*, vol. 28, no. 4, pp. 656–715, 1949.
- [17] T. M. Cover, J. A. Thomas, J. Wiley, and W. InterScience, *Elements of Information Theory*. Wiley-Interscience New York, 2006.