

Link-Spam-Resistant Domain Ranking: A Comparative Evaluation of PageRank, HITS, and TrustRank Under Simulated Link-Farm Attacks Using Common Crawl Web Graphs

Romelyn J. Banaybanay¹, Reagan B. Ricafort²

¹Initao College, Initao, Misamis Oriental, Philippines

²AMA University, Makati City, Philippines

DOI: <https://doi.org/10.51583/IJLTEMAS.2026.150600016>

Received: 17 June 2026; Accepted: 22 June 2026; Published: 03 July 2026

ABSTRACT

Link-based ranking is still an important part of web information retrieval, but it remains vulnerable to coordinated link-farm attacks. This study examined the resistance of PageRank, HITS, and TrustRank to link spam using simulated attacks on domain-level web graphs extracted from Common Crawl. A quantitative, comparative, and controlled experimental design was used. Five independent base graphs were prepared, with each graph containing 100,000 legitimate domain nodes. Four link-farm structures were tested: star-centered, bounded-density clique, bipartite, and mixed. The attacks were applied at four intensity levels: 1%, 3%, 5%, and 10% of the legitimate graph size. Across all conditions, the experiment produced 480 attacked graph instances and 1,440 verified algorithm-level runs. The algorithms were compared using spam target promotion, spam score share, top-1000 spam infiltration, legitimate-domain displacement, ranking stability, execution time, peak memory use, and iteration count. The results showed that PageRank was the most vulnerable algorithm. It had the highest spam target promotion, highest spam score share, measurable top-1000 spam infiltration, greater legitimate-domain displacement, and lowest ranking stability. HITS showed the strongest overall resistance and computational efficiency. It recorded near-zero spam score share, zero top-1000 infiltration, zero legitimate-domain displacement, highest stability, fastest runtime, and fewest iterations. TrustRank also blocked top-1000 infiltration and legitimate-domain displacement, although it allowed higher target promotion than HITS. The bounded-density clique structure caused the most damage to PageRank, while higher attack intensity increased spam impact in most cases. The findings show that link-spam resistance depends on the ranking algorithm, attack structure, attack intensity, and evaluation metric used. A single ranking score or attack model is not enough to evaluate resistance to link-farm manipulation.

Keywords: Web information retrieval, link spam, PageRank, HITS, TrustRank

INTRODUCTION

Web information retrieval needs ranking algorithms to decide which pages or domains appear first in search results. Link-analysis methods use hyperlinks as signs of authority, popularity, or trust. PageRank, HITS, and TrustRank offer three related but different ways to read link structure. PageRank estimates importance through a random-walk model. HITS separates hub and authority scores. TrustRank limits rank propagation by starting from trusted seed nodes. Recent work on sparse, dense, contextual, and late-interaction retrieval shows that ranking remains an active research area [5, 6, 8, 15, 25]. Studies on exact lexical matching, dense passage retrieval, sequence-to-sequence ranking, and reproducible IR tools also support the need for reliable ranking methods [9, 14, 16, 21, 23]. Graph-based recommendation, knowledge graph, link-prediction, and graph neural network studies show that graph structure remains useful for ranking and relational analysis [7, 11, 12, 24, 28].

Link-analysis methods are useful, but they are not safe from manipulation. Link spam happens when artificial links are created to raise the rank of a target page or domain. A link farm is one example. It uses a group of spam nodes to pass ranking influence to one or more targets. This behavior can promote low-quality or harmful

domains and push legitimate sources downward. Recent work on spam and web-spam detection supports the need to test ranking systems under adversarial conditions [1, 2, 13, 20, 26]. Fraud and graph-anomaly studies also show that coordinated or camouflaged graph behavior can distort system outputs [3, 17, 18, 22].

This study responds to that problem by comparing PageRank, HITS, and TrustRank under controlled link-farm attacks. It uses domain-level Common Crawl graphs and synthetic attack structures instead of surveys or private search-engine logs. The aim is not to claim that the simulated attacks represent all real spam behavior. Rather, the aim is to test how each algorithm behaves under clear, repeatable, and measurable attack settings.

The main research question was: How do PageRank, HITS, and TrustRank differ in their resistance to simulated link-farm attacks across attack structures and intensities using domain-level Common Crawl web graphs? The study measured differences in spam target promotion, spam score share, top-k spam infiltration, legitimate-domain displacement, ranking stability, execution time, memory use, and convergence behavior.

The following null hypotheses were tested:

H01: There is no significant difference among PageRank, HITS, and TrustRank in spam target promotion.

H02: Link-farm structure does not significantly affect top-k spam infiltration.

H03: Attack intensity does not significantly affect legitimate-domain displacement.

H04: There is no significant algorithm-by-structure interaction for spam target promotion.

H05: There is no significant algorithm-by-structure-by-intensity interaction for ranking stability.

H06: There is no significant difference among PageRank, HITS, and TrustRank in execution time.

H07: There is no significant difference among PageRank, HITS, and TrustRank in peak memory consumption.

METHODOLOGY

Research Design

The study used a quantitative, comparative, and controlled experimental design. It compared PageRank, HITS, and TrustRank under simulated link-farm attacks. The study did not involve human participants, surveys, interviews, questionnaires, or confidential records. The independent variables were ranking algorithm, link-farm structure, and attack intensity. The dependent variables were spam target promotion, spam score share, top-1000 spam infiltration, legitimate-domain displacement, Spearman ranking stability, rank-biased overlap, execution time, peak memory consumption, and iteration count.

Data Source and Unit of Analysis

Domain-level web graph data from Common Crawl were used. In the graph, each node represented a pay-level domain and each directed edge represented a hyperlink relationship between two domains. An edge from domain u to domain v meant that at least one page from u linked to a page from v . The pay-level domain was the unit of analysis.

Graph Sampling

Five independent base graphs were prepared. Each base graph contained 100,000 legitimate domain nodes. Direction-aware random-walk-with-restart sampling preserved connected web-graph neighborhoods while keeping the experiment manageable. After sampling, self-loops and duplicate directed edges were removed. The cleaned graphs were used as the base testbeds.

TrustRank Seed Selection

TrustRank needed trusted seed domains. Archived Tranco domain rankings were used as a reproducible proxy for trusted seeds. For each base graph, the highest-ranked Tranco domains that also appeared in the sampled graph were selected. Synthetic spam nodes were not allowed in the seed set.

Ranking Algorithms

PageRank was run using a random-surfer model with a damping factor of 0.85. HITS was run using authority and hub updates over the directed adjacency matrix, and authority scores served as the main ranking output. TrustRank was run as seed-personalized PageRank. All algorithms used sparse matrix representations in the same computing environment. This workflow followed common scientific computing and reproducible IR tool practices [10, 19, 27].

Link-Farm Attack Model

Each attack included one target spam domain and a set of supporting spam domains. Four link-farm structures were tested: star-centered, bounded-density clique, bipartite, and mixed. Each attack used a fixed spam-originated edge budget of 10m, where m was the number of synthetic spam nodes. External ingress links were added from legitimate non-seed domains.

Attack Intensities and Replication

Four attack intensities were tested: 1%, 3%, 5%, and 10% of the 100,000 legitimate nodes in each base graph. Each structure-intensity-base graph condition was repeated six times using independent attack realizations. The final experiment produced 480 attacked graph instances and 1,440 algorithm-level runs.

Statistical Analysis

Descriptive statistics were calculated for each algorithm, structure, and intensity. Friedman tests were used for paired nonparametric comparisons among algorithms. Pairwise Wilcoxon signed-rank tests with Holm adjustment were used for post hoc analysis. Structure and intensity effects were examined by algorithm. Interaction effects were tested using rank-based interaction ANOVA, following current procedures for multifactor rank-based analysis [4]. The statistical workflow used Python-based scientific computing and reproducible IR tools [10, 19, 27].

RESULTS

The experiment was completed successfully. It produced 480 attacked graph instances and 1,440 algorithm-level metric records. All algorithm runs converged. The final master dataset had no missing values and no duplicate experimental rows. It was balanced across algorithms, structures, intensities, and base graphs.

Across all attack conditions, PageRank had the highest mean target promotion at 0.930980. TrustRank followed at 0.753470, while HITS had the lowest mean target promotion at 0.186051. PageRank also had the highest spam score share, measurable top-1000 spam infiltration, greater legitimate-domain displacement, and the lowest ranking stability.

Table 1. Main descriptive comparison of ranking algorithms.

Metric	HITS	PageRank	TrustRank
Mean Target Promotion	0.186051	0.930980	0.753470
Mean Spam Score Share	0.000003	0.032597	0.000058

Mean Infiltration@1000	0.000000	0.007198	0.000000
Mean P95 Legitimate Displacement	0.000000	9.173438	0.000000
Mean Spearman Stability	0.998953	0.906140	0.990426
Mean Rank-Biased Overlap	1.000000	0.972289	0.999985
Mean Execution Time (seconds)	5.005167	7.988084	7.727799
Mean Peak Memory (GiB)	0.729034	0.729571	0.731206
Mean Iterations	31.800000	98.964583	96.952083

Friedman tests showed significant algorithm effects for all tested metrics: target promotion, spam score share, top-1000 spam infiltration, P95 legitimate displacement, Spearman stability, RBO, execution time, peak memory consumption, and iterations. Pairwise Wilcoxon tests with Holm adjustment showed that HITS had lower target promotion than PageRank and TrustRank. PageRank had higher target promotion than TrustRank.

Link-farm structure affected the algorithms in different ways. For PageRank, the bounded-density clique was the most damaging structure. It had the highest spam score share, highest top-1000 infiltration, highest P95 displacement, lowest Spearman stability, and lowest RBO. For HITS, the star-centered structure produced the highest target promotion, but HITS still had zero top-1000 infiltration and zero P95 displacement across all structures. For TrustRank, the bipartite structure had the highest target promotion, while the clique structure had the highest spam score share and lowest stability.

Higher attack intensity generally increased spam impact. For PageRank, higher intensity led to stronger spam score share, higher infiltration, greater displacement, and lower ranking stability. HITS and TrustRank also showed increases in target promotion and spam score share, but both kept zero top-1000 infiltration and zero P95 legitimate displacement at all intensity levels.

The Algorithm x Structure interaction for target promotion was significant, $F = 114.918537$, $p = 2.814419e-118$, partial eta squared = 0.327479. For ranking stability, the Algorithm x Structure x Intensity interaction was significant when Spearman stability was used, $F = 65.500370$, $p = 1.949972e-170$, partial eta squared = 0.460189. The same three-way interaction was not significant when RBO was used, $F = 0.873179$, $p = 0.612166$, partial eta squared = 0.011237.

Table 2. Final hypothesis-decision summary.

Hypothesis	Metric	Decision	Interpretation
H01	Target Promotion	Reject H01	Algorithms differed significantly; HITS had the lowest target promotion, PageRank the highest, Trust Rank intermediate.
H02	Infiltration@1000	Partially reject H02	Structure affected spam infiltration under PageRank; HITS and Trust Rank produced tied zero infiltration.
H03	P95 Legitimate Displacement	Partially reject H03	Intensity affected displacement under PageRank; HITS and Trust Rank produced tied zero displacement.
H04	Target Promotion	Reject H04	Algorithm x Structure interaction was significant: $F=114.918537$,

			$p=2.814419e-118$, partial eta squared=0.327479.
H05	Spearman Stability and RBO	Partially reject H05	Three-way interaction was significant for Spearman Stability but not significant for RBO.
H06	Execution Time Seconds	Reject H06	Execution time differed significantly among algorithms; HITS was fastest.
H07	Peak Process RSS GiB	Reject H07	Peak memory differed significantly, although practical differences were small.

DISCUSSION

PageRank, HITS, and TrustRank responded differently to the simulated link-farm attacks. PageRank was the most vulnerable. HITS showed the strongest overall resistance. TrustRank showed strong but selective resistance. These results reflect the different assumptions of the algorithms. PageRank spreads importance through links. HITS separates hub and authority behavior. TrustRank limits influence through trusted seed personalization. Related graph-based studies also show that recommendation, graph convolution, knowledge representation, and graph learning methods depend on relational structure [7, 11, 12, 28, 30]. Link-prediction research supports the same role of graph structure in relational inference and ranking behavior [24].

PageRank was vulnerable because it depends on link-based score propagation. When spam nodes are coordinated, they can pass ranking influence to a target. In this study, the clique structure caused the most damage because spam nodes reinforced one another before passing influence outward. This led to higher spam score share, higher infiltration, greater displacement, and lower stability. The finding is consistent with fraud and graph-anomaly studies showing that coordinated graph behavior can distort detection and ranking outputs [3, 17, 18, 22]. It also supports web-spam studies that recommend combining link-based signals with content-based and anomaly-based defenses [1, 20, 26].

HITS performed best overall. It stopped top-1000 spam infiltration and legitimate-domain displacement across all structures and intensities. It also showed the highest stability and the best computational efficiency. The spam target still gained some rank position under some structures, especially star-centered attacks, but the effect did not spread into wider ranking pollution.

TrustRank also showed strong resistance. It stopped top-1000 infiltration and legitimate-domain displacement across all structures and intensities. This result supports the use of seed-personalized ranking to limit spam diffusion. Still, TrustRank allowed higher target promotion than HITS. This means that trust personalization reduced broad spam spread but did not fully stop the target spam domain from improving its rank. This point agrees with graph-based research that shows the need to manage trust, links, and relational structure carefully [12, 24, 28, 30].

The results also show that attack structure matters. PageRank was most affected by clique structures. HITS had its highest target promotion under star-centered attacks. TrustRank had different weaknesses under bipartite and clique structures. These differences show that one attack model is not enough to evaluate link-spam resistance. Recent studies on graph anomaly detection, web-spam detection, and imbalanced graph learning also support structure-aware testing of adversarial graph behavior [18, 20, 26, 29].

Overall, link-spam resistance is not defined by one metric alone. PageRank was weak against coordinated reinforcement. HITS was strongest overall. TrustRank contained broad spam spread but was less effective than

HITS in suppressing target promotion. These findings support multi-algorithm, multi-metric, and structure-aware evaluation in web information retrieval.

CONCLUSION AND RECOMMENDATIONS

Conclusion

This study tested the link-spam resistance of PageRank, HITS, and TrustRank under simulated link-farm attacks using domain-level Common Crawl web graphs. The experiment produced 480 attacked graph instances and 1,440 verified algorithm-level runs. The results showed clear differences among the three algorithms.

PageRank was the most vulnerable algorithm. It had the highest spam target promotion, highest spam score share, measurable top-1000 spam infiltration, greater legitimate-domain displacement, and lowest ranking stability. The bounded-density clique structure and higher attack intensities caused the most damage to PageRank.

HITS showed the strongest overall resistance and computational efficiency. It had the lowest target promotion, lowest spam score share, zero top-1000 infiltration, zero legitimate-domain displacement, highest ranking stability, fastest execution time, and fewest iterations. TrustRank also performed well in preventing top-1000 infiltration and legitimate-domain displacement, but it allowed higher target promotion than HITS.

The study concludes that link-spam resistance depends on algorithm design, attack structure, attack intensity, and the metric used to measure spam success. HITS was strongest overall. TrustRank limited broad spam spread. PageRank was most vulnerable to link-farm manipulation.

Recommendations

Search systems should not rely on PageRank alone when link manipulation is possible. PageRank should be combined with trust signals, spam-pattern detection, link-anomaly detection, content-quality signals, and temporal link-growth monitoring. Spam and web-spam studies support the use of both content-based and link-based anti-spam defenses [1, 13, 20, 26]. Fraud and anomaly detection studies also support graph-aware defenses against coordinated behavior [3, 17, 18, 22].

Systems that need stronger resistance to link-farm attacks should include HITS-like authority and hub analysis as part of a broader ranking model. HITS should still be combined with other anti-spam methods because target promotion can still occur under some structures. Graph-based recommendation, graph convolution, knowledge graph, and graph neural network studies support the use of relational structure in ranking and detection models [7, 11, 12, 28, 30]. Link-prediction research also supports the value of modeling relationships among connected entities [24].

TrustRank or seed-personalized ranking should be used when the goal is to limit spam influence from untrusted graph regions. Careful seed selection, periodic seed validation, and added spam-detection features are still needed to reduce target-level gains. This recommendation fits graph-based research showing that link structure, knowledge representation, and graph neural methods depend on the quality of relational signals [12, 24, 28, 30].

Future evaluations should use several metrics, link-farm structures, and attack intensities. Future work should extend the experiment to page-level graphs, labeled spam datasets, temporal spam patterns, and hybrid ranking models. Future studies should compare link-based ranking with sparse, contextual, dense, and exact-matching retrieval approaches under adversarial settings [5, 6, 8, 9, 14]. Late-interaction ranking should also be tested in spam-resistant retrieval settings [15, 25]. Reproducible IR toolkits, sequence-to-sequence ranking, optimized dense retrieval, and imbalanced graph learning should be included in future comparisons [16, 21, 23, 29].

LIMITATIONS

This study used simulated link-farm attacks instead of naturally labeled spam networks. The findings therefore measure resistance to the attack models used in the experiment. This matters because real spam behavior can include email, social, content, review, and web-spam patterns that were not fully represented by the simulated link farms [1, 2, 13, 20, 26]. Real adversarial behavior can also include fraud and graph-anomaly patterns that were not directly modeled [3, 17, 18, 22].

The study used domain-level graphs instead of page-level graphs. Domain-level aggregation supports large-scale analysis, but it removes page-level details such as anchor text, topical relevance, internal site structure, and link placement. These details matter because lexical expansion, contextual matching, dense passage retrieval, and late-interaction retrieval often work below the domain level [5, 6, 9, 14, 15]. Sequence-to-sequence ranking, optimized dense retrieval, and lightweight late-interaction models also rely on document-level or text-level signals [21, 23, 25].

The sampled graphs were controlled testbeds and should not be treated as complete representations of the web. Random-walk sampling may give more weight to well-connected graph regions. Future replication should record sampling settings, software versions, and computing tools to strengthen reproducibility [10, 16, 19, 27].

TrustRank used Tranco-based proxy seed domains. These seeds were reproducible, but popularity is not the same as trustworthiness. Future work should compare different seed-selection methods and test whether trusted seeds remain stable across time, domain categories, and graph regions [12, 18, 24, 28, 30].

Data and Reproducibility Statement

This study used public domain-level web graph data from the Common Crawl Host- and Domain-Level Web Graphs October-November-December 2024 release. The domain-level graph was used to build five independent legitimate-domain base graphs, each with 100,000 nodes. The original Common Crawl graph was not used as a labeled spam dataset.

TrustRank seed domains were selected from archived Tranco ranking data as a reproducible proxy for trusted domains. For each base graph, the highest-ranked Tranco domains found in the sampled graph were used as seed nodes. Synthetic spam nodes were excluded from the seed set.

The experiment used controlled synthetic link-farm attacks so that testing could be repeated. Four attack structures were simulated: star-centered, bounded-density clique, bipartite, and mixed. Four attack intensities were tested: 1%, 3%, 5%, and 10% of the legitimate graph size. Each structure-intensity-base graph condition was repeated using six independent attack realizations.

The reproducibility package contains the dataset source record, sampling notes, experimental pseudocode, statistical output files, result screenshots, and appendix tables. These files document the data source, sampling process, attack design, ranking workflow, statistical analysis, and reported results. Full reproduction requires access to the cited public datasets and implementation of the documented experimental pipeline.

REFERENCES

1. Akinyelu, A. A. (2021). Advances in spam detection for email spam, web spam, social network spam, and review spam: ML-based and nature-inspired-based techniques. *Journal of Computer Security*, 29(5), 473-529. <https://doi.org/10.3233/JCS-210022>
2. Alom, Z., Carminati, B., & Ferrari, E. (2020). A deep learning model for Twitter spam detection. *Online Social Networks and Media*, 18, Article 100079. <https://doi.org/10.1016/j.osnem.2020.100079>
3. Dou, Y., Liu, Z., Sun, L., Deng, Y., Peng, H., & Yu, P. S. (2020). Enhancing graph neural network-based fraud detectors against camouflaged fraudsters. *Proceedings of the 29th ACM International Conference on Information and Knowledge Management*, 315-324. <https://doi.org/10.1145/3340531.3411903>

4. Elkin, L. A., Kay, M., Higgins, J. J., & Wobbrock, J. O. (2021). An aligned rank transform procedure for multifactor contrast tests. *Proceedings of the 34th Annual ACM Symposium on User Interface Software and Technology*, 754-768. <https://doi.org/10.1145/3472749.3474784>
5. Formal, T., Lassance, C., Piwowarski, B., & Clinchant, S. (2022). From distillation to hard negative sampling: Making sparse neural IR models more effective. *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval*, 2353-2359. <https://doi.org/10.1145/3477495.3531857>
6. Formal, T., Piwowarski, B., & Clinchant, S. (2021). SPLADE: Sparse lexical and expansion model for first stage ranking. *Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval*, 2288-2292. <https://doi.org/10.1145/3404835.3463098>
7. Gao, C., Zheng, Y., Li, N., Li, Y., Qin, Y., Piao, J., Quan, Y., Chang, J., Jin, D., He, X., & Li, Y. (2023). A survey of graph neural networks for recommender systems: Challenges, methods, and directions. *ACM Transactions on Recommender Systems*, 1(1), Article 3. <https://doi.org/10.1145/3568022>
8. Gao, L., & Callan, J. (2021). Condenser: A pre-training architecture for dense retrieval. *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, 981-993. <https://doi.org/10.18653/v1/2021.emnlp-main.75>
9. Gao, L., Dai, Z., & Callan, J. (2021). COIL: Revisit exact lexical match in information retrieval with contextualized inverted list. *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 3030-3042. <https://doi.org/10.18653/v1/2021.naacl-main.241>
10. Harris, C. R., Millman, K. J., van der Walt, S. J., Gommers, R., Virtanen, P., Cournapeau, D., Wieser, E., Taylor, J., Berg, S., Smith, N. J., Kern, R., Picus, M., Hoyer, S., van Kerkwijk, M. H., Brett, M., Haldane, A., del Río, J. F., Wiebe, M., Peterson, P., Gérard-Marchant, P., Sheppard, K., Reddy, T., Weckesser, W., Abbasi, H., Gohlke, C., & Oliphant, T. E. (2020). Array programming with NumPy. *Nature*, 585, 357-362. <https://doi.org/10.1038/s41586-020-2649-2>
11. He, X., Deng, K., Wang, X., Li, Y., Zhang, Y., & Wang, M. (2020). LightGCN: Simplifying and powering graph convolution network for recommendation. *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval*, 639-648. <https://doi.org/10.1145/3397271.3401063>
12. Ji, S., Pan, S., Cambria, E., Marttinen, P., & Yu, P. S. (2021). A survey on knowledge graphs: Representation, acquisition, and applications. *IEEE Transactions on Neural Networks and Learning Systems*, 33(2), 494-514. <https://doi.org/10.1109/TNNLS.2021.3070843>
13. Kaddoura, S., Chandrasekaran, G., Popescu, D. E., & Duraisamy, J. H. (2022). A systematic literature review on spam content detection and classification. *PeerJ Computer Science*, 8, Article e830. <https://doi.org/10.7717/peerj-cs.830>
14. Karpukhin, V., Oguz, B., Min, S., Lewis, P., Wu, L., Edunov, S., Chen, D., & Yih, W. (2020). Dense passage retrieval for open-domain question answering. *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing*, 6769-6781. <https://doi.org/10.18653/v1/2020.emnlp-main.550>
15. Khattab, O., & Zaharia, M. (2020). ColBERT: Efficient and effective passage search via contextualized late interaction over BERT. *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval*, 39-48. <https://doi.org/10.1145/3397271.3401075>
16. Lin, J., Ma, X., Lin, S.-C., Yang, J.-H., Pradeep, R., & Nogueira, R. (2021). Pyserini: A Python toolkit for reproducible information retrieval research with sparse and dense representations. *Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval*, 2356-2362. <https://doi.org/10.1145/3404835.3463238>
17. Liu, Z., Dou, Y., Yu, P. S., Deng, Y., & Peng, H. (2020). Alleviating the inconsistency problem of applying graph neural network to fraud detection. *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval*, 1569-1572. <https://doi.org/10.1145/3397271.3401253>
18. Ma, X., Wu, J., Xue, S., Yang, J., Zhou, C., Sheng, Q. Z., Xiong, H., & Akoglu, L. (2023). A comprehensive survey on graph anomaly detection with deep learning. *IEEE Transactions on Knowledge and Data Engineering*, 35(12), 12012-12038. <https://doi.org/10.1109/TKDE.2021.3118815>

19. Macdonald, C., & Tonellotto, N. (2020). Declarative experimentation in information retrieval using PyTerrier. *Proceedings of the 2020 ACM SIGIR International Conference on the Theory of Information Retrieval*, 161–168. <https://doi.org/10.1145/3409256.3409829>
20. Makkar, A., & Kumar, N. (2020). An efficient deep learning-based scheme for web spam detection in IoT environment. *Future Generation Computer Systems*, 108, 467-487. <https://doi.org/10.1016/j.future.2020.03.004>
21. Nogueira, R., Jiang, Z., & Lin, J. (2020). Document ranking with a pretrained sequence-to-sequence model. *Findings of the Association for Computational Linguistics: EMNLP 2020*, 708-718. <https://doi.org/10.18653/v1/2020.findings-emnlp.63>
22. Pang, G., Shen, C., Cao, L., & van den Hengel, A. (2021). Deep learning for anomaly detection: A review. *ACM Computing Surveys*, 54(2), Article 38, 1-38. <https://doi.org/10.1145/3439950>
23. Qu, Y., Ding, Y., Liu, J., Liu, K., Ren, R., Zhao, W. X., Dong, D., Wu, H., & Wang, H. (2021). RocketQA: An optimized training approach to dense passage retrieval for open-domain question answering. *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 5835-5847. <https://doi.org/10.18653/v1/2021.naacl-main.466>
24. Rossi, A., Barbosa, D., Firmani, D., Matinata, A., & Merialdo, P. (2021). Knowledge graph embedding for link prediction: A comparative analysis. *ACM Transactions on Knowledge Discovery from Data*, 15(2), Article 14. <https://doi.org/10.1145/3424672>
25. Santhanam, K., Khattab, O., Saad-Falcon, J., Potts, C., & Zaharia, M. (2022). ColBERTv2: Efficient and effective retrieval via lightweight late interaction. *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 3715-3734. <https://doi.org/10.18653/v1/2022.naacl-main.272>
26. Shahzad, A., Nawari, N. M., Gillani, S. M. Z. R., & Khan, A. (2021). An improved framework for content- and link-based web-spam detection: A combined approach. *Complexity*, 2021, Article 6625739. <https://doi.org/10.1155/2021/6625739>
27. Virtanen, P., Gommers, R., Oliphant, T. E., Haberland, M., Reddy, T., Cournapeau, D., Burovski, E., Peterson, P., Weckesser, W., Bright, J., van der Walt, S. J., Brett, M., Wilson, J., Millman, K. J., Mayorov, N., Nelson, A. R. J., Jones, E., Kern, R., Larson, E., Carey, C. J., et al. (2020). SciPy 1.0: Fundamental algorithms for scientific computing in Python. *Nature Methods*, 17(3), 261-272. <https://doi.org/10.1038/s41592-019-0686-2>
28. Wu, Z., Pan, S., Chen, F., Long, G., Zhang, C., & Yu, P. S. (2021). A comprehensive survey on graph neural networks. *IEEE Transactions on Neural Networks and Learning Systems*, 32(1), 4-24. <https://doi.org/10.1109/TNNLS.2020.2978386>
29. Zhao, T., Zhang, X., & Wang, S. (2021). GraphSMOTE: Imbalanced node classification on graphs with graph neural networks. *Proceedings of the 14th ACM International Conference on Web Search and Data Mining*, 833-841. <https://doi.org/10.1145/3437963.3441720>
30. Zhou, J., Cui, G., Hu, S., Zhang, Z., Yang, C., Liu, Z., Wang, L., Li, C., & Sun, M. (2020). Graph neural networks: A review of methods and applications. *AI Open*, 1, 57-81. <https://doi.org/10.1016/j.aiopen.2021.01.001>

APPENDICES

Appendix A. Experimental Design Summary

Appendix Table A1. Experimental design summary.

Component	Verified value
Base graphs	5 independent graphs
Legitimate nodes per base graph	100,000 domains
Algorithms	PageRank, HITS, TrustRank
Structures	Star, Clique, Bipartite, Mixed
Attack intensities	1%, 3%, 5%, 10%
Attack realizations	6 per structure-intensity-base graph condition
Attacked graph instances	480
Algorithm-level runs	1,440
Primary data source	Common Crawl domain-level web graph
Trust seed proxy	Archived Tranco ranking

Appendix B. Direction Table by Algorithm and Structure

Appendix Table B1. Mean direction by algorithm and link-farm structure.

Algorithm	Structure	Target Promotion	Spam Score Share	Infil@1000	P95 Disp.	Spearman	RBO	Time	Memory	Iterations
HITS	Bipartite	0.138682	0.000003	0.000000	0.000000	0.999011	1.000000	4.937690	0.729367	31.800000
HITS	Clique	0.141947	0.000003	0.000000	0.000000	0.998585	1.000000	5.000123	0.729063	31.800000
HITS	Mixed	0.217248	0.000003	0.000000	0.000000	0.998532	1.000000	4.949230	0.729363	31.800000
HITS	Star	0.246326	0.000003	0.000000	0.000000	0.999684	1.000000	5.133626	0.728343	31.800000
PageRank	Bipartite	0.931403	0.025555	0.001000	1.000000	0.934824	0.958214	7.835396	0.729907	95.600000
PageRank	Clique	0.931355	0.071694	0.022917	22.716667	0.817822	0.946760	8.697891	0.729601	109.058333
PageRank	Mixed	0.930872	0.020865	0.003875	9.676250	0.886683	0.989270	7.725024	0.729908	95.600000
PageRank	Star	0.930290	0.012274	0.001000	3.300833	0.985231	0.994913	7.694027	0.728868	95.600000
TrustRank	Bipartite	0.820064	0.000043	0.000000	0.000000	0.995908	0.999992	7.775917	0.731467	96.800000

TrustRank	Clique	0.802913	0.000130	0.000000	0.000000	0.976711	0.999992	7.714501	0.731224	97.408333
TrustRank	Mixed	0.715707	0.000037	0.000000	0.000000	0.990708	0.999978	7.670509	0.731453	96.800000
TrustRank	Star	0.675196	0.000021	0.000000	0.000000	0.998379	0.999976	7.750268	0.730681	96.800000

Appendix C. Direction Table by Algorithm and Attack Intensity

Appendix Table C1. Mean direction by algorithm and attack intensity.

Algorithm	Intensity %	Target Promotion	Spam Score Share	Infil@100	P95 Disp.	Spearman	RBO	Time	Memory	Iterations
HITS	1	0.162504	5.14E-07	0.000000	0.000000	0.999841	1.000000	5.162207	0.720385	31.800000
HITS	3	0.169351	1.59E-06	0.000000	0.000000	0.999459	1.000000	4.858365	0.722302	31.800000
HITS	5	0.185284	2.65E-06	0.000000	0.000000	0.999002	1.000000	5.008369	0.725999	31.800000
HITS	10	0.227063	5.45E-06	0.000000	0.000000	0.997511	1.000000	4.991728	0.747450	31.800000
PageRank	1	0.945636	7.39E-03	0.002883	3.266667	0.976512	0.995539	8.084670	0.720805	98.058333
PageRank	3	0.938031	2.16E-02	0.004333	5.925000	0.934219	0.984497	7.983587	0.722729	99.050000
PageRank	5	0.929841	3.52E-02	0.006292	8.584167	0.896510	0.971081	7.972144	0.726570	99.300000
PageRank	10	0.910412	6.62E-02	0.015283	18.917917	0.817318	0.938040	7.9111936	0.748180	99.450000
TrustRank	1	0.677634	1.44E-05	0.000000	0.000000	0.997806	0.999994	7.854634	0.722520	96.833333
TrustRank	3	0.744481	3.42E-05	0.000000	0.000000	0.994031	0.999989	7.601424	0.725146	96.908333
TrustRank	5	0.782194	5.90E-05	0.000000	0.000000	0.990004	0.999984	7.668601	0.728852	96.966667
TrustRank	10	0.809570	1.23E-04	0.000000	0.000000	0.979865	0.999971	7.786537	0.748306	97.100000

Appendix D. Key Interaction Effects

Appendix Table D1. Key H04 and H05 rank-based interaction effects.

Test	Metric	Effect	F	p-value	Partial eta ²
H04: Algorithm x Structure	Target Promotion	Algorithm x Structure	114.918537	2.814419e-118	0.327479
H05: Algorithm x Structure x Intensity	Spearman Stability	Algorithm x Structure	1551.453659	0.000000e+00	0.870648
H05: Algorithm x Structure x Intensity	Spearman Stability	Algorithm x Structure x Intensity	65.500370	1.949972e-170	0.460189
H05: Algorithm x Structure x Intensity	RBO	Algorithm x Structure	67.681206	6.032906e-74	0.226980
H05: Algorithm x Structure x Intensity	RBO	Algorithm x Structure x Intensity	0.873179	0.612166	0.011237