

When the Wipers Win: How Practitioner Training and Tool Diversity Predict Success Against Anti-Forensic Techniques

Dr. Travis Eygabroad

Assistant Professor of Computer Science and Cybersecurity, St. Ambrose University, Davenport, IA, USA

Abstract

Anti-forensic techniques such as data wiping, encryption, and log tampering increasingly thwart digital investigations. This Year 1 survey of 83 practitioners examines whether formal cybersecurity credentials or the number of forensic platforms used predict perceptions of tool effectiveness and real-world anti-forensic encounter rates. We grouped training into “Trained” (CEH, EnCase Certified Examiner, CompTIA Security+, etc.) versus “Untrained,” and effectiveness ratings into “Effective” versus “Ineffective,” then applied Fisher’s Exact and χ^2 tests. A Kruskal–Wallis H test (with Mann–Whitney U follow-up) assessed ordinal ratings, and a negative-binomial GLM modeled yearly anti-forensic impact counts by training, role, tool diversity, and experience. None of the credential or tool-diversity predictors reached significance across analyses (all $p > .12$), suggesting that operational context and workflow integration—not résumé variables—drive both tool satisfaction and exposure to hiding techniques. Free-text responses identify practitioner priorities (e.g., threat-intel feeds, cross-tool hash sharing) that will guide Year 2 open-source enhancements.

Keywords: anti-forensics; digital forensics; training effectiveness; survey; negative-binomial GLM; Kruskal–Wallis.

1. Introduction

Anti-forensic techniques ranging from secure deletion and encryption to steganography and log tampering pose a growing challenge for digital forensic investigations. As adversaries employ sophisticated methods to obscure or destroy evidence, forensic practitioners must rely on both their formal training and the full range of tools at their disposal. Yet it remains unclear whether certifications such as CEH, EnCase Certified Examiner, or GCFA translate into higher perceived tool effectiveness or fewer real-world anti-forensic setbacks. Similarly, the value of using multiple forensic platforms (e.g., EnCase, FTK, Autopsy) in buffering against hidden evidence has not been quantified. This study addresses two hypotheses: H1 (Training effect): Practitioners holding one or more specialized credentials rate their tools as more effective than uncredentialed peers. H2 (Tool-diversity effect): Using three or more distinct forensic platforms predicts fewer annual “anti-forensic impacts.” We also pose an exploratory research question (RQ) regarding which anti-forensic methods most undermine effectiveness and whether this variation differs by agency type. By surveying 83 investigators across local, state, federal, and 2 private sectors, we establish a quantitative baseline to inform Year 2 laboratory experiments and targeted open-source countermeasures.

2. METHODS

We solicited 83 anonymous responses from digital forensic practitioners, who reported their primary training credential and rated their leading forensic software on a 5-point scale from “Extremely Effective” to “Not Effective.” For categorical analyses, we collapsed credentials into Trained (formal certifications) versus Untrained (none specified) and effectiveness into Effective (Extremely/Very/Moderately) versus Ineffective (Slightly/Not). We then conducted (1) Fisher’s Exact Test on the 2×2 table and a full χ^2 test on the 7×5 training-by-effectiveness matrix; (2) a Kruskal–Wallis H test (with Mann–Whitney U follow-up between CEH and GCFA) on the ordinal ratings; and (3) a negative-binomial generalized linear model predicting annual counts of anti-forensic impact incidents from training, role (local/state/federal/private), tool-diversity (number of platforms), and years of experience. Free-text “tool improvement” requests underwent thematic coding to rank practitioner-identified gaps. Of the 83 practitioners, 24 worked in local law enforcement, 9 in state agencies, 13 in federal agencies, 20 in private forensic firms, and 17 in other roles (contractors, academic labs). This distribution is shown in Table 1.

Table 1 Breakdown of Respondent Roles

Organization Type	Count	Percentage
Local law enforcement	24	28.9%
State law enforcement	9	10.8%
Federal law enforcement	13	15.7%
Private forensic analyst	20	24.1%
Other (Contractor/academic)	17	20.5%
Total	83	100%

3. RESULTS

3.1 Training vs. Perceived Effectiveness

In this study, we investigated whether the type of cybersecurity training someone has influences their perception of the effectiveness of their forensic tools. Participants (n = 83) listed their training credentials and rated their tools from “Extremely Effective” to “Not Effective” (Table 1). To facilitate a more precise analysis using Fisher’s Exact Test, training was categorized into two groups: trained (with a formal credential, such as CEH, EnCase, or CompTIA Security+) and untrained (with no specified training). Similarly, perceptions of effectiveness were grouped as either effective (extremely, very, or moderately effective) or ineffective (slightly or not effective at all) (Table 2).

Table 2 Training Credential x Perceived Tool-Effectiveness (N=83)

Training Type	Ext. Eff.	Very Eff.	Mod. Eff.	Sl. Eff	Not Eff.	Total
CEH certification	3	3	5	3	1	15
CISSP bootcamp	1	0	0	1	0	2
CompTIA Security+	4	6	3	5	3	21
EnCase Certified Examiner	6	4	4	3	1	18
GCFA	3	0	0	0	1	4
None specified	1	0	5	2	1	9

SANS FOR500	1	3	0	6	4	14
Column Total	19	16	17	20	11	83

A Fisher's Exact Test was used on the simplified 2×2 table to see if there was a significant relationship between training and perceived effectiveness. The result was $p = 0.2784$, which is not statistically significant, indicating that any difference observed between trained and untrained participants could be due to chance.

Additionally, a more detailed Chi-square test of independence was run using the full table of training types and five levels of tool effectiveness (Table 2). The test yielded a Chi-square statistic of $\chi^2 = 32.18$ with 24 degrees of freedom, and a p-value of .123. This, too, is not statistically significant and supports the same conclusion: training type does not appear to influence how effective users perceive their tools to be significantly.

In short, people with formal cybersecurity training were no more or less likely to rate their tools as effective than those without training. This suggests that other factors—like hands-on experience, tool familiarity, or real-world exposure to anti-forensic techniques—may have a greater influence on perceptions of effectiveness. Statistical tools, such as the Chi-square test, have been applied, for example, by Goonatilake (2007), who used a Chi-square test to detect abnormal activities in network traffic and developed a network intrusion detection system (Table 3).

Table 3 Chi-square Test of Independence (7x 5 matrix)

Test	χ^2	df	p-value
Training x Perceived Effectiveness	32.18	24	0.123

According to the study provided, the evaluation of different forensic training certifications, such as CEH or GCFA, is influenced by perceived effectiveness ratings (1 = Not effective to 5 = Extremely effective). Because our five-point effectiveness rating violates 4 normality, we used a Kruskal–Wallis test to compare median scores across training categories (Table 3). Datatab (2021). This was selected to compare the median scores according to the groups, and the results indicate that there are no significant differences ($H(6) = 9.27$, $p = 0.16$). No training groups differed significantly, so the medians all fell between ‘Moderately’ (3) and ‘Very’ (4) being effective.

To further investigate, we conducted post hoc Mann–Whitney U tests comparing specific credential groups. This analysis prioritized CEH ($n = 15$) versus GCFA ($n = 4$), as these represented the most significant credentialed classifications in Table 3, excluding “None specified,” and embody a critical contrast: CEH’s broad penetration-testing focus versus GCFA’s specialized forensic training. This test confirmed that there were no significant differences in their ratings ($U = 28.5$, $p = 0.42$), reinforcing that credential type does not predict effectiveness perceptions.

These findings align with prior research that has used non-parametric methods for ordinal ratings. Tambwe et al. (2023) applied the Kruskal–Wallis test to assess professional perspectives on confidentiality ($H = 18.581$, $p = 0.017$) and data availability ($H = 20.787$, $p = 0.008$) in construction risk management, mirroring our null result and suggesting that the operational context outweighs formal certifications. Barari Reykandeh & Shokri (2022) also used Kruskal–Wallis to demonstrate significant differences in cyber-attack

frequencies across institutions ($H = 14.70$, $p = 0.002$), validating the test's utility for skewed data in security contexts. Again, like Tambwe et al. (2023), who found that DRM “enhances integrity” ($MIS = 4.24$) through systematic risk approaches, these studies confirm that non-parametric methods are essential when comparing practitioner assessments with small samples or non-normal distributions, particularly in cybersecurity and data management domains.

The Kruskal–Wallis test was ideal for our subgroups and ordinal effectiveness ratings. The lack of significant differences ($H(6) = 9.27$, $p = 0.16$) suggests that tool perceptions depend more on case-specific factors, such as anti-forensics techniques, rather than credentials, aligning with practitioner-identified needs, such as cross-tool hash sharing (Table 4).

Table 4 Kruskal-Wallis Test of Median Effectiveness Ratings by Training Credential

H (KW)	df	p-value
9.27	6	0.16

3.3 Count-Model Findings

This study used a negative binomial generalized linear model (GLM) to predict annual anti-forensic impact counts among digital forensic investigators, incorporating predictors such as training type, role, tool diversity, and experience. The baseline profile, which included an investigator with CEH certification in a federal role, utilizing one tool platform, and with 0.5 years of experience, was associated with approximately 10.7 impacts per year ($IRR = 10.65$, $p < .001$).

No individual training category emerged as a statistically significant predictor ($p > .26$), though the GCFA credential had the highest estimated effect ($IRR \approx 1.97$). Tool diversity showed no meaningful association with impact frequency ($IRR = 1.01$, $p = 0.961$), and increased experience trended positively but was not statistically significant ($IRR = 1.02$ per year, $p = 0.157$). Non-

federal roles, including state and local law enforcement, showed lower impact rates, but none reached statistical significance. These results suggest that résumé-level variables may be weak proxies for exposure to advanced anti-forensic methods, which the complexity or nature of assigned cases may instead drive.

This finding parallels Leslie et al. (2018), who modeled cyber intrusion counts using a negative binomial framework. While they identified predictors such as DNS activity and policy violations, other expected factors, like host count and organizational visibility, were not significant. Both studies highlight the limitations of assumed indicators and underscore the need for contextual or behavioral predictors.

In contrast, Chapter 5 of *Artificial Intelligence and Sustainable Computing* (2022) adopted a machine learning approach, integrating Poisson processes into SVM kernels to detect DDoS attacks. While methodologically distinct, their use of statistical modeling to capture behavior patterns reflects a broader trend toward hybrid analytic frameworks. Unlike this study's focus on human and organizational factors, their model targets automated detection at the network level. Together, these studies suggest that conventional predictors often fall short in explaining cybersecurity outcomes, reinforcing the need for models that better account for operational context and adversarial complexity.

Table 5: Negative-Binomial GLM Predicting Annual Anti-Forensic Impact Counts

Predictor	IRR	p-value
Intercept	10.65	<0.001
CompTIA Security+	1.47	0.297
EnCase Certified Examiner	1.38	0.389
GCFA	1.97	0.268
None specified	0.98	0.968
SANS FOR500	1.23	0.603
Local Law Enforcement	0.87	0.689
State Law Enforcement	0.64	0.359
Private Forensic Analyst	0.68	0.319
Other role	0.54	0.114
Tool diversity (per additional platform)	1.01	0.961
Years of experience (per year)	1.02	0.157

Notes: Reference levels: training = CEH certification, role= federal law enforcement, tool diversity= 1 platform, experience= 0.5 years. IRR < 1 indicates a decreased in expected anti-forensic impact counts relative to the reference category; IRR > 1 indicates an increase.

4. PRACTITIONER-IDENTIFIED GAPS

Free-text responses yielded a ranked list of desired tool enhancements: automated threat-intel feeds, cross-tool hash sharing, log-tamper alerts, GPU-accelerated carving, and AI detection of encrypted containers (Table 6).

Table 6 Top Tool-Improvement Themes (Open-Ended Response, n= 83)

Rank	Suggested Improvement Theme	Mentions
1	Automated correlation with threat-intel feeds	24
2	Cross-tool hash-library sharing	18
3	Real-time alerts for log-manipulation attempts	15
4	Better GPU acceleration for large-scale carving	13
5	AI-driven detection of encrypted containers	13

5. DISCUSSION

Across categorical, ordinal, and count-model analyses, formal training credentials and tool diversity did not predict perceived tool effectiveness or annual anti-forensic impacts. This consistent null pattern highlights that operational context, real-world case complexity, and workflow integration likely dominate both satisfaction and success in countering anti-forensic tactics. Practitioner-identified priorities—live threat feeds, cross-tool interoperability, and tamper alerts—underscore the gap between résumé variables

and functional needs. Addressing these user-centered enhancements should take precedence over additional certification programs or broad tool inventories.

6. LIMITATIONS

We acknowledge several limitations of this Year 1 survey. First, our reliance on self-reported perceptions of tool effectiveness may introduce response and social desirability biases, rather than relying on objective performance metrics. Second, the sample sizes for some credential groups (e.g., GCFA, CISSP) were small, which reduces statistical power and may mask actual effects. Third, practitioners self-selected into the survey, so there may be systematic differences between respondents and the broader forensic community. Finally, our “None specified” training category likely includes both truly untrained individuals and those who simply omitted credential details, potentially leading to misclassification. Future work will address these limitations through controlled laboratory experiments, larger and more representative samples, and the integration of objective tool-performance data.

7. FUTURE WORK

In Years 2 through 5, our work will progress from survey findings to hands-on experimentation and the creation of tools. We will construct laboratory environments that simulate data wiping, encryption, steganography, and log tampering across several widely used forensic platforms. These controlled scenarios will let us move beyond perception and generate objective measures of detection and recovery performance. A key goal is to design and release open-source enhancements that leverage artificial intelligence techniques. Planned capabilities include models that highlight suspicious log activity, routines that recognize hidden or encrypted payloads in disk slack space, and interfaces that surface live threat intelligence indicators directly within everyday forensic workflows. All software, supporting notebooks, and pre-trained models will be shared under a permissive license, allowing practitioners and researchers to validate and extend our work. Each academic cycle, students will refine these tools, benchmark them against commercial alternatives, and gather feedback from partner laboratories. By Year 5, the project will provide a mature, freely available toolkit and a comprehensive synthesis of what we have learned, offering the forensic community practical guidance and empirically validated strategies for countering modern anti-forensic tactics.

8. CONCLUSION

Our Year 1 survey demonstrates that credential type and tool breadth alone are insufficient to explain forensic practitioners’ experiences with anti-forensic techniques. In Year 2, we will translate these findings into controlled laboratory experiments and develop targeted, open-source plugins that address the precise gaps identified by investigators. By iterating between field-informed requirements and lab-validated solutions, we aim to deliver robust counter-CAF methodologies that bolster forensic effectiveness regardless of formal credentials or organizational resources.

9. ACKNOWLEDGEMENTS

This research was generously supported by the Stoffel Fund for Excellence in Scientific Inquiry at St. Ambrose University. We also thank Dr. Kevin Lillis for his departmental support and guidance, as well as the digital forensics practitioners who took the time to complete our survey.

10. FUNDING

This work was supported by the Stoffel Fund for Excellence in Scientific Inquiry, St. Ambrose University (IRB 2425580).

11. ETHICS APPROVAL

This study was approved by the St. Ambrose University IRB (Protocol 2425580); all participants gave informed consent.

12. AUTHOR'S CONTRIBUTIONS

Travis Eygabroad conceived the study, designed the survey, performed and supervised students during statistical analysis, interpreted results, and drafted and revised manuscript. The author also acknowledges the following contributors for research assist: Luke Sproule (categorical findings, statistical analysis, editing,) Parker Trapkus (introduction and conclusion drafting,) Rachel Romilus (ordinal findings,) and Aleric Weber (count-model findings). All contributors provided support that does not meet the criteria for authorship.

13. CONFLICT OF INTEREST

The authors declare no competing interests.

13. AVAILABILITY OF DATA

The datasets generated and analyzed during the current study are not publicly available due to participant confidentiality and institutional review board (IRB) restrictions but are available from the corresponding author on reasonable request. De-identified survey data and statistical analysis code can be shared upon request in compliance with the St. Ambrose University IRB Protocol #2425580.

14. AUTHOR'S BIOGRAPHY

Dr. Travis Eygabroad serves as the Department Chair and Assistant Professor of Computer and Information Sciences at St. Ambrose University in Davenport, Iowa. His teaching and research focus on cybersecurity, digital forensics, and higher education leadership, with particular emphasis on anti-forensic methodologies, forensic tool performance, and technology integration in postsecondary education.

REFERENCES

1. Barari Reykandeh, K., & Shokri, S. A. (2022). Russian Digital Economy and Cybersecurity: An Overview of Recent Developments. *Journal of World Sociopolitical Studies*, 6(3), 439-498.
2. DATAtab Team. (2025). DATAtab: Online Statistics Calculator. <https://datatab.net/tutorial/anova>



3. Dubey, H. M., Mohan, H., Dubey, P., & Srivastava, L. (2022). Artificial Intelligence and Sustainable Computing. Springer Singapore.
4. Goonatilake, R., Herath, A., Herath, S., Herath, S., & Herath, J. (2007). Intrusion detection using the chi-square goodness-of-fit test for information assurance, network, forensics and software security. *Journal of Computing Sciences in Colleges*, 23(1), 255-263.
5. Leslie, N. O., Harang, R. E., Knachel, L. P., & Kott, A. (2018). Statistical models for the number of successful cyber intrusions. *Journal of Defense Modeling and Simulation*, 15(1), 49–63.
6. Tambwe, O. T., Aigbavboa, C. O., & Akinradewo, O. (2025). Benefits of construction data risks management in the construction industry. *Journal of Engineering, Design and Technology*, 23(2), 458-476.