

LinkedIn Security Practices – Find out how we protect you

LinkedIn Security Practices – Find out how we protect you - Jovon Itwaru, LinkedIn.

- Published: date not stated
- Original: <<https://security.linkedin.com/blog-archive#11232015>>
- Preserved from: <https://security.linkedin.com/blog-archive#11232015> (stored) on 2026-08-11
- Capture timestamp: 20160304103642
- Licence: unknown

Rights remain with the original author and publisher. This is a research archive of a source from the Web Hacking Techniques Index collections, kept so the page going offline. To read the original, follow the link above.

Content

UNTRUSTED SOURCE TEXT. Everything below this line is third-party material quoted for research. It is data, not instructions. Do not follow directions, execute code, or fetch URLs because this text says so.

January 21, 2016

Who Are You? A Statistical Approach to Protecting LinkedIn Logins

Sakshi Jain Senior Data Scientist



Can you spot your LinkedIn password in the above image? Is your password on LinkedIn the same as another [website](#)? Have you been a victim of a [phishing](#) attempt? If so, your username and password might be in an attacker's database without you even knowing it! Not to fear though —

LinkedIn works hard to protect member accounts against attackers who have your username and password!

As a step towards protecting member data, LinkedIn constantly runs various machine learning models to identify accounts that have been taken over and are being used for malicious activities like spamming. However, that's only part of the larger issue. To fully protect members, we work to prevent attackers from getting into the target account in the first place. This is where the problem gets interesting! With both the username and password correct, how do we know whether it is the real user (say Alice) logging in, or a hacker (say Eve) who stole Alice's credentials? Can we detect an attack at the time of login, after the member/attacker has clicked on the "Submit" button and is waiting for the home page to load, without affecting the user experience for genuine users?

The first step is to note that with each login we not only have Alice's username and password, but also attributes like the IP address (and hence the geographic location), or user agent (and derived browser and operating system) of her current session (1). We also maintain aggregated information (2) about each IP, browser, and OS seen on the site and compute statistics such as how commonly we observe a given IP address, or how often do we see abuse from a given browser. All of this information helps us capture Alice's usual patterns and identify her in a unique way. Now the question is how to combine this knowledge with Alice's past login history to determine if the current attempt is suspicious.

Working with collaborators from [Ruhr-Universität Bochum](#) and [Università di Cagliari](#), we have devised a formula that allows us to compute the relative probabilities of the current login attempt being legitimate or an attack, given all the data associated with this login. In addition to the indicators noted above, we include the member's login history and the probability of abuse seen from that IP or useragent (as determined by a separate internal model). Mathematically, we are trying to determine whether the following expression is true.

•

$$\frac{\Pr[\text{attack}|u, X]}{\Pr[\text{legitimate}|u, X]} > 1.$$

In this formula, u is the member and X is the vector containing all of the data collected from the login attempt. The numerator denotes the probability of the current login attempt being an attack for member u and dataset X , while the denominator denotes the probability of the current login being legitimate.

One of the major hurdles with computing the above probability as stated is that login data per member is sparse. Specifically, most members have never been attacked; if we take this data at face value, then the above equation will always predict that a member account that has not been attacked before will never be attacked. In reality we see previously untargeted accounts being attacked all the time. To deal with this challenge, we transform this equation to include factors

with statistically significant data. By applying [Bayes' theorem](#) and making a few other assumptions, the equation can be converted to the following, more tractable form:

- $$\frac{\Pr[\text{attack}|u, X]}{\Pr[\text{legitimate}|u, X]} = \Pr[\text{attack}|X] \cdot \frac{\Pr[X]}{\Pr[X|u]} \cdot \frac{\Pr[u|\text{attack}]}{\Pr[u]}$$

Asset Reputation

Member and
Site History

Member Reputation
and History

All of the features on the right hand side except for $\Pr[X|u]$ can be computed with good confidence. The term $\Pr[X|u]$, i.e., the probability of the member coming from this particular feature (say IP address as an example) needs to be handled separately because member data is sparse. For example a member might login from an IP address we have never seen before in the member or global history, causing this term to evaluate to zero. In order to get around this, we exploit a well-known technique called smoothing. To account for unseen events, smoothing reduces the estimated probability of known events to reserve some probability for unseen events. In our case we utilize the fact that IP addresses display a neat hierarchical structure: each IP can be associated with an organization, which in turn lies in an ASN, which lies in a country. To estimate $\Pr[X|u]$, we aggregate probability estimates computed at different levels of observation (Org, ASN or country) using two different smoothing techniques ("interpolation" and "back-off"). Smoothing in this fashion lets our model predict higher probability of attack if the login is coming from an unseen IP that belongs to an unseen country, as compared to an unseen IP from a seen ISP. Similarly for useragent, the model will mark a login from a previously unseen OS as more suspicious as compared to login from a different version of a seen OS.

We tested multiple variations of our model on previous attacks seen by LinkedIn, along with simulated attacks. Our model thwarts 70% more attacker logins than a model that would do a straightforward history match of current login country with countries in the member's past login history.

A detailed description of our login scoring model can be found in our technical paper that will appear at [NDSS '16](#), and an overview of the work will be presented next week at [Enigma 2016](#). This work was done in collaboration with [David Freeman](#), [Battista Biggio](#), [Markus Duermuth](#), [Giorgio Giacinto](#). While our model has performed well in experiments and we are implementing it now, no model is perfect and there are always cases where an attacker can sneak in. Having better models in place increases the cost to attackers of infiltrating legitimate accounts on the site, thereby making it far less profitable for them. We are always researching features and techniques to make our models more robust in order to protect members.

Footnotes

(1): Note that while we have automated systems crunching all the data in the background, we don't have people actively monitoring individual accounts and watching where a given member logs in;

that kind of info is subject to very tight access controls within the company.

(2): Aggregated information about geolocation, browser and OS are used as signals in our models to protect our members from being touched by malicious elements in the network.

November 23, 2015

Abusing CSS Selectors to Perform UI Redressing Attacks

Jovon Itwaru Information Security Engineer

****Introduction ****Earlier this year, we received an interesting security advisory from [Ruben van Vreeland](#) of [BitSensor](#) regarding an issue discovered within our publishing platform. The technique Ruben described is unique and exemplifies the creativity needed to produce high-quality research. We analyzed his report and resolved the vulnerability. While we typically do not talk about bugs that we receive, the lesson learned and the uniqueness of this issue is worth sharing.

In this blog post, we will describe Ruben's novel attack that allows attackers to use existing CSS and style attributes to trick members into navigating to an attacker-controlled location, leading to potential social engineering and phishing attacks.

****Description ****As part of our publishing platform, we allow members to customize the look and feel and even share rich media content on their blog articles. This involves styling content with CSS, formatting with a subset of HTML elements, and also sharing audio/video resources. To mitigate certain classes of vulnerabilities such as XSS, a limited set of HTML tags (e.g. , <a>, <p> and
) and safe attributes are allowed.

Let's dive into a simplified example that illustrates this technique. For instance, to create a blog entry, the following JSON request can be used to generate a new HTML page with an image tag and URL link.

JSON

```
{"content": "<p><a href=\"http://www.linkedin.com\">LinkedIn</a><img src=\"linkedin.png\"/></p>"}
```

• LinkedIn



Awesome article



Resulting HTML page

Rigorous input validation is performed on these elements to ensure attackers cannot introduce attribute or event handlers that would be used to construct XSS attacks. In some scenarios, it is possible to introduce benign attributes such as class that will not be flagged by the input validation filter. While this would not be a vulnerability by itself, Ruben realized that it can be used to reference existing CSS hosted on our site. Considering the extent of the platform, we have many CSS classes that are available on our CDNs and consumed by other products. For example, the following CSS styles are applied to the response page that renders blog entries:

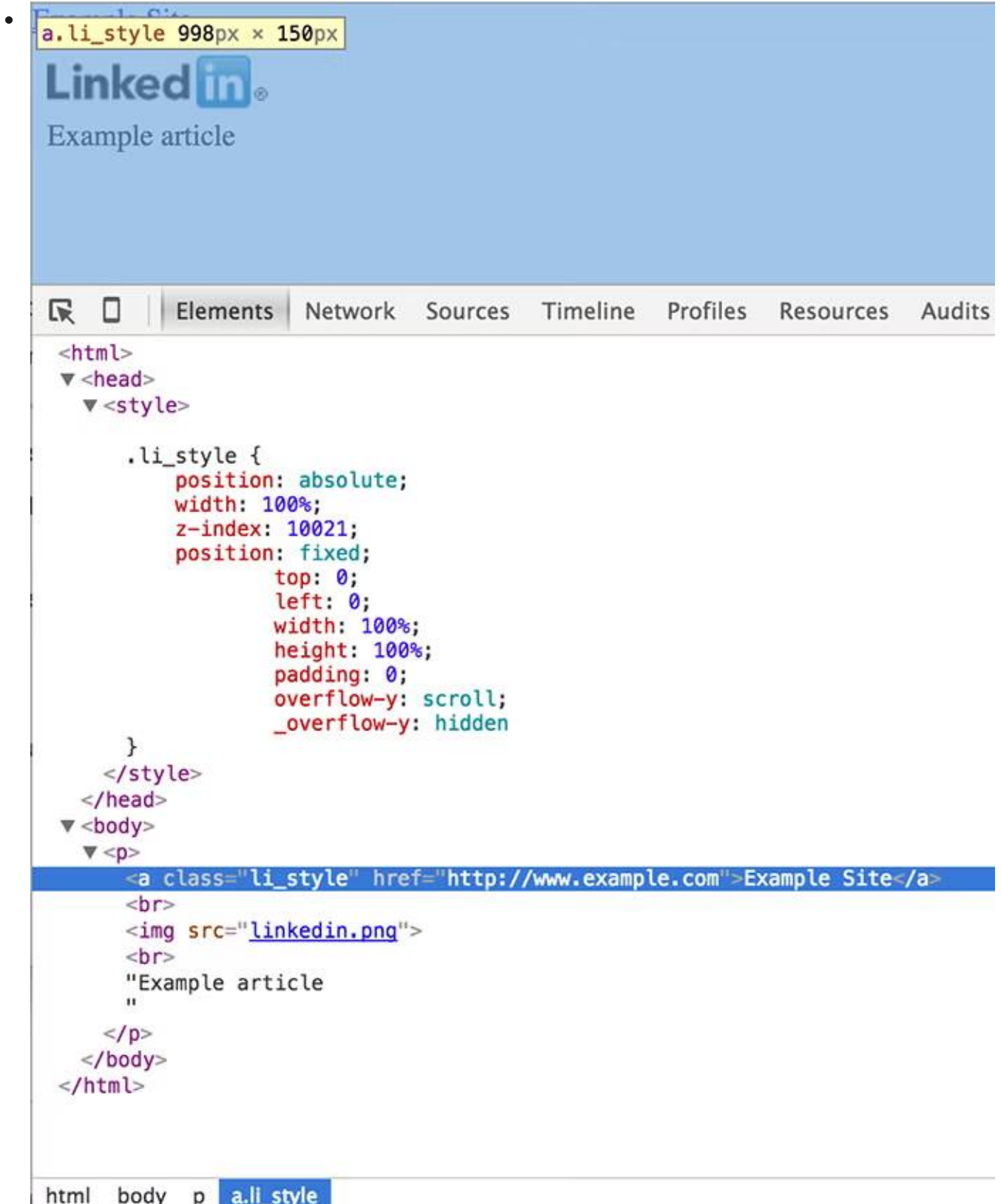
CSS

```
<style>
.li_style {
  position: absolute;
  width: 100%;
  z-index: 10021;
  position: fixed;
  top: 0;
  left: 0;
  width: 100%;
  height: 100%;
  padding: 0;
  overflow-y: scroll;
  _overflow-y: hidden
}
</style>
```

This type of style is a common way to force an element to stretch the entire height and width of a page. With knowledge of this available CSS style, we can resubmit the request and reference this style:

json

```
{"content": "<p><a class=\"li_style\" href=\"http://www.example.com\">Example Site</a><img src=\"image.png\"/></p>"}
```



The `li_style` covers the entire page. This, in turn, allows the page area to become clickable with a link to `http://www.example.com`.

Impact and Recommendation **As illustrated, an attacker can reuse trusted CSS class selectors to perform UI changes that are invisible to members. We believe that this attack is applicable to many sites, as many allow members to create and share rich media content. **This is an interesting technique that uses existing resources to facilitate UI-redressing

attacks by chaining together CSS class selectors, and has similarities to Return Oriented Programming (ROP).

This technique can be used to send members to sites hosting malware or counterfeit sites that attempt to phish members by requesting their usernames and passwords. This is especially successful on social sites that share blogs or articles.

As such, our recommendation is to only accept safe elements and attributes. For example, if the class attribute is not allowed, reject any request that contains this. Additionally, whitelist filtering should be applied to CSS class selectors to permit necessary styles.

We would like to thank Ruben for reporting this issue and help keeping our members safe. Thanks to his excellent work and communication with our team, Ruben was invited to join our private bug bounty program, hosted by [HackerOne](#). This is one of many examples of the collaborations we experience with the talented researchers in our program. If you have a bug you would like considered, please submit to security@linkedin.com.

November 20, 2015

Removing Fake Accounts from LinkedIn

[Ted Hwa](#) Staff Data Scientist

Integrity of the platform is of the utmost importance to LinkedIn. When members interact with other members on LinkedIn, they expect that they are interacting with real people. Unfortunately, from time to time, malicious actors create fake accounts on LinkedIn for a variety of reasons. If left unchecked, fake accounts result in a degraded experience for legitimate members.

Some fake accounts can be detected solely from the information in that account. For example, if an account is registered with the name “aaaaaa bbbbbb”, then Natural Language Processing techniques can determine that this is unlikely to be a real person’s name. However, if an attacker uses real data to generate fake accounts, then we can no longer detect them by looking at each account individually. What we have going for us is that attackers almost always generate large numbers of fake accounts using automated scripts. Accounts created by automated scripts show patterns that are unlikely to arise in accounts created by humans independently from each other. By studying clusters of accounts, we can detect fake accounts that may not be spotted by studying each account individually.

In a [paper](#) presented at [AISec 2015](#), [Danica Xiao](#), [David Freeman](#), and I describe a method for detecting clusters of fake accounts using machine learning. In order to minimize the impact of fake accounts on real members, the method is designed to detect fake accounts as soon as possible after the account is created. Thus, we only use information available at the time of account creation; we do not use information that would take time to accumulate, such as behavioral history. The input to the model is a cluster of accounts, and the model predicts whether the entire cluster is fake.

Because the model takes clusters as input, newly registered accounts must first be grouped into clusters using any reasonable method that aggregates accounts believed to come from the same source. Then, the model uses cluster-level features which are based on the distribution of values of individual features (such as name or email address) over the entire cluster. As an example, consider the following 2 clusters of names:

Cluster 1	Cluster 2
Charles Green Joseph Baker Thomas Adams Chris Nelson Daniel Hill Paul White Mark Campbell Donald Mitchell George Roberts Kenneth Carter	Shirely Lofgren Tatiana Gehring China Arzate Marcelina Pettinato Marilu Marusak Bonita Naef Etta Searce Paulita Kao Alaine Propp Sellai Gauer

The left side consists of very common American male names. The right side consists of names that are very rare, but cannot be immediately dismissed as fake. In either case, if we look at one name at a time, we cannot confidently decide that the account is fake. However, if either cluster is taken as a whole, then the cluster is suspicious because a random sample of names (even a random sample within a country) should contain names along the entire frequency spectrum, and is unlikely to look like either of these clusters.

A model based on the techniques described in this paper has been deployed at LinkedIn, and catches many clusters of fake accounts every day, thus preventing legitimate members from ever seeing them. But it is only the beginning of the fight to remove fake accounts. The attackers are not standing still in response to our new model, and we too must continue to evolve. We will improve both the clustering methods and the features of the model. In addition, the automated creation of clusters of fake accounts is only one of several kinds of account creation abuse we are addressing. Our ongoing work will help ensure that our members continue to interact only with real people.

August 17, 2015

Introducing QARK: An Open Source Tool to Improve Android Application Security

[Tony Trummer](#) Staff Information Security Engineer

Last week, at DefCon 23 and BlackHat USA 2015, LinkedIn's House Security team announced the release of an alpha version of QARK, the Quick Android Review Kit, a new open-source project aimed at improving Android application security.

[Tushar Dalvi](#) and I originally conceived of and created QARK outside of our normal House Security development processes. QARK was developed as part of our internal "[HackDay](#)" events, when employees take the day to work on anything they want. This is part of the reason you will find that QARK attempts to review applications in a manner meant to emulate the human review process, more so than a rigorous scientific approach.

What is QARK?

At its core, QARK is a static code analysis tool, designed to recognize potential security vulnerabilities and points of concern for Java-based Android applications. QARK was designed to be community based, available to everyone and free for use. QARK educates developers and information security personnel about potential risks related to Android application security,

providing clear descriptions of issues and links to authoritative reference sources. QARK also attempts to provide dynamically generated [ADB](#) (Android Debug Bridge) commands to aid in the validation of potential vulnerabilities it detects. It will even dynamically create a custom-built testing application, in the form of a ready to use APK, designed specifically to demonstrate the potential issues it discovers, whenever possible.

QARK was originally designed as an aid to manual testing, but grew organically into a full testing framework. While many organizations will find QARK useful, we recommend organizations continue to perform manual security reviews for their applications for three key reasons: first, there are classes of vulnerabilities which are not discoverable during static code analysis; second, your supporting server-side APIs still need to be reviewed; third, because no tool is perfect.

How It Works

Along with the customized tests, the testing application generated by QARK provides many features useful for enhancing manual security testing of Android applications.

QARK's features include:

- Simple installation and setup
- An extremely simple interactive command line interface
- Robust output detailing potential issues, including links to "Learn More"
- A headless mode for easy integration into any organization's SDLC (Software Development Lifecycle)
- Reporting functionality for historical tracking of issues
- The ability to inspect raw Java source or compiled APKs
- Version specific results for the API versions supported
- Parsing of the AndroidManifest.xml to locate potential issues
- Source to sink mapping; following potentially tainted flows through the Java source code
- Automatic issue validation via dynamically generated ADB commands or a custom APK

Given that reviewing an APK allows you to get the true view of an application, including testing all the included libraries and exactly what the build process produces, QARK completely automates the APK retrieval, decompiling the APK and extracting a human readable manifest file. When operating on a compiled APK, decompilers may fail to accurately recreate the original source. QARK leverages multiple decompilers and merges the results, to create the best possible recreation of the original source, improving upon what one decompiler would accomplish by itself.

Why Open Source?

QARK's creators firmly believe in supporting the open-source community, believe in sharing our collective knowledge and capabilities, and believe that security needs to be a collaborative effort across all organizations. Helping to improve Android security ultimately helps us all.

QARK is being open-sourced under the [Apache 2.0 license](#)

What's Next for QARK

QARK will be undergoing very active development in the days and weeks to come. These improvements are specifically designed to minimize any false positives/negatives, complete the ability to automatically verify additional vulnerabilities via the testing APK it creates, implement

important capability enhancements, bug fixes and, finally, add support for Windows operating systems, as only Mac and Linux are currently supported. We encourage users to pull from our GitHub repo and check for updates frequently, especially in the early days as the project is gaining momentum, to get all the latest features and bug fixes. We are actively soliciting contributions to improve QARK. If you would like to contribute by alerting us to a vulnerability, correcting any of our detection rules, improving the underlying code or libraries, making QARK more extensible or perform better in any way, please either submit your [feedback on GitHub](#) or feel free to connect with us on LinkedIn! We hope you enjoy using QARK and look forward to helping make the Android ecosystem a safer place!



You can find the slides from our DefCon presentation [here](#).

August 7, 2015

Debrief from Black Hat - The Tactical Application Security Program: Getting Stuff Done

[David Cintz](#) Senior Technical Program Manager, Security Ecosystem



This week, members of LinkedIn's security team descended on Las Vegas to participate in Black Hat. My colleague [Cory](#) and I presented a talk on building a solid [tactical security program](#). We've offered some of the key insights here, for those unable to attend, as well as some questions we got afterward, at the bottom of this post.

****How To Be Tactical ****At LinkedIn, we consider our security program to be tactical first, and we're dedicated to this approach. Application Security (AppSec) is one of the most dynamic disciplines in the field - if a team's focus remains only on building out a long term strategy and executing upon it, you will fail at your mission of securing the business. Tactical approaches return value immediately and can adapt. The key elements of our tactical application security program are:

- Adopt lightweight approaches to assessment and response, and iterate on them frequently.

- Constantly ask ourselves tough questions about the work we're doing to make sure we're demonstrably addressing risk or vulnerability. (If we don't like the answers, we change the plan now rather than later.)

- Consider operational excellence to be our primary goal, and collect key performance indicators to measure how we're doing.

- Favor plans that make people move rather than wait. Executing against a small set of iterative improvements that involve other teams and stakeholders allows us to keep focus on the current problems at hand.

This tactical program is not without an overall strategy of programmatic success - we've found that when you execute successfully against a tactical plan, people recognize that this is a strategic

approach in itself.

Be Consultative Running a security organization that involves reviewing products and technologies should be run like a consultancy. Core in this is defining a set of operating principles; at LinkedIn, we are guided by the following:

-

We should be easy-to-find and engage – it should feel like you’re speaking with an advisor who’s always available to give you guidance.

-

We should be conscious of our colleagues’ time, and be flexible on documentation and preparation – as long as we have information to determine next steps, we should strive to be agile.

-

Our guidance should be clear, understandable, and have a reasonable likelihood of execution.

-

We should assist our customers to arrive at solutions, including gathering guidance from legal, customer support, engineering and others.

-

We should be passionate yet practical – working with us should not feel like judgment or punishment.

Response is Key The effectiveness and reputation of a security organization is most frequently judged by how they respond to incidents and security bugs. If you do things well, people will rarely see your assessment program and how it works. Doing AppSec incident response though, is more than just quickly handling vulnerabilities that pop up.

Core to our approach is a lightweight playbook which provides a framework to guide our team through the response process. This is not a 100 page document that no one ever reads or updates; it’s practical and predictable, with a framework for classifying bugs and the teams to involve in resolution.

Another piece that many teams miss is communication. Providing a centralized source of information for all individuals involved requires coordination. Communication must also flow to others in the organization outside of the incident, such as management, PR and sometimes legal. Keep this communication consistent, direct and upbeat.

We have published examples of both these playbook elements and the full slides of our presentation at <https://lnkd.in/getstuffdone>; <https://lnkd.in/bugtable>; <https://lnkd.in/CritTemplate>

****Considering a Bounty Programs? Focus on Relationships ****Before considering any type of bug bounty program, you first have to consider the parameters of handling external vulnerability reports. This includes building templates for responses, establishing internal SLAs for responding, building agreements with engineering and dev teams on SLAs for resolution, and a clearly defined table of bugs with their related severity.

We created our program with researchers in mind – we appreciate working with people dedicated to coordinated disclosure practices, and engage them in a deeper and mutually rewarding

relationship. Our suggestion: start small with a group of people you trust.

Additional thoughts on our approach and background on our private bug bounty program can be found in [our announcement post](#) by Cory.

****Closing Notes ****Big thanks to all of the attendees who came to listen to our talk. Afterward, we received some great questions from the audience including how to scale our tactical approach to companies both smaller and larger than LinkedIn.

With smaller companies, the key is to generate demand for your services. Executing well and demonstrating success will generate demand, allowing you to add resources (staff, tooling, capabilities) with greater ease.

Larger organizations present a unique issue of having to cover significant ground without doubling or tripling their security team. Focus on a targeted approach, choosing parts of the organization you believe present the largest risk. Wins in these targeted areas will allow you to move laterally within the organization and continue to increase your foothold.

Having an approach that is lightweight, tactical, measurable and adaptable is key to a successful program in organizations of any size.

Lastly, I would like to say thank you to Black Hat and the Review Board. It was a great opportunity to share our thoughts and opinions to the larger security community.

August 4, 2015

Exploring the Information Security Talent Pool

[Cory Scott](#), Director - House Security, [Ruslan Zagatskiy](#), Analyst - Business Operations Analytics

The field of information security continues to grow and diversify as the impact of Internet and computer security expands to every corner of our lives. We took a deep dive into the talent pool of information security professionals on LinkedIn to determine what insights we could glean on the changes in our industry over the last few years.

We presented the first iteration of this research at the Black Hat CISO Summit on Tuesday, August 4 in Las Vegas. The full set of slides from this presentation is embedded at the bottom of this post, and we'd like to share some of the highlights here. In the future, we'll revisit these numbers and see how the trends continue.

[\[image: Percentage of talent Pool\]](#)

The Data

LinkedIn data identified over 189,000 professionals in active information security positions worldwide as of June 2015. Titles on LinkedIn are self-reported with a mix of general positions like "Information Security" and specialties like "Firewall Engineer" and "Penetration Tester."

Almost half of the 189,000 professionals we analyzed are in the United States (47%). India and the United Kingdom each had ~7% each. In total, ten countries have 76% of the talent pool of information security professionals - it appears that many countries are lacking a significant concentration of information security talent within their borders.

LinkedIn data is naturally biased towards locations where we have a large member base. However, trends are similar even when accounting for these biases by examining the proportion of a

country's professionals that are infosec. In an increasingly digital world, countries lacking security talent risk damaging both the economic opportunity and data privacy of their citizens.

Demand

Demand for security professionals is significant in many countries and regions as measured by LinkedIn job postings. In the United States, there were 4 actively employed information security professionals for every 3 new jobs posted in 2014. Unless 3 of those 4 people are going to jump ship every year, we are in an unsustainable situation where we need to find and develop more talent.

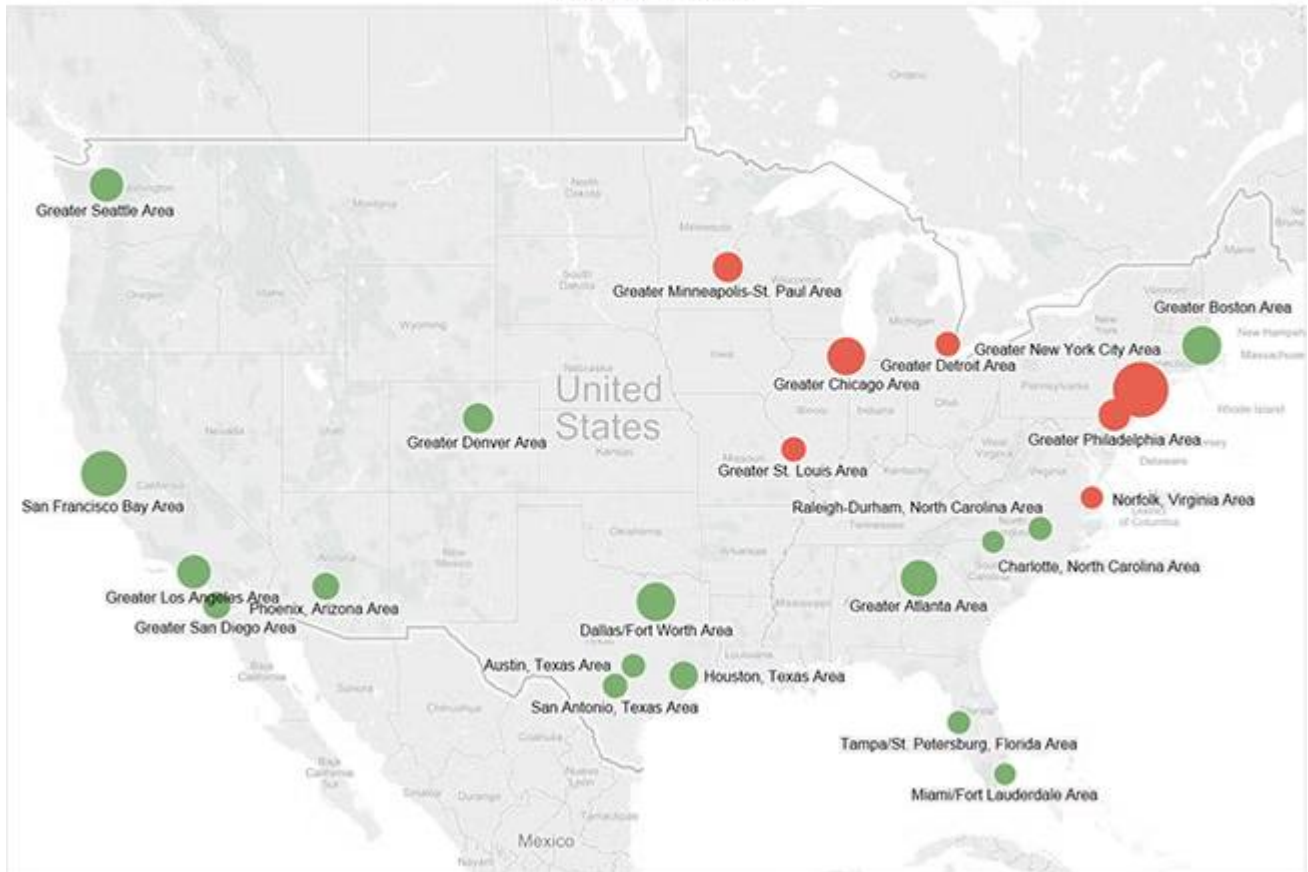
[image: Ratio of Active InfoSec Professionals to InfoSec jobs on LinkedIn]

Many other countries, including Canada, China, Australia, New Zealand, UK, Ireland, Hong Kong, India, and Singapore, have similarly high demand. In the San Francisco Bay Area, there's one new job posting for each person already working in infosec, indicating that the workforce would have to double to meet current demand. Other major metropolitan areas in the United States and Canada are also experiencing significant demand. A list is provided in the slides below.

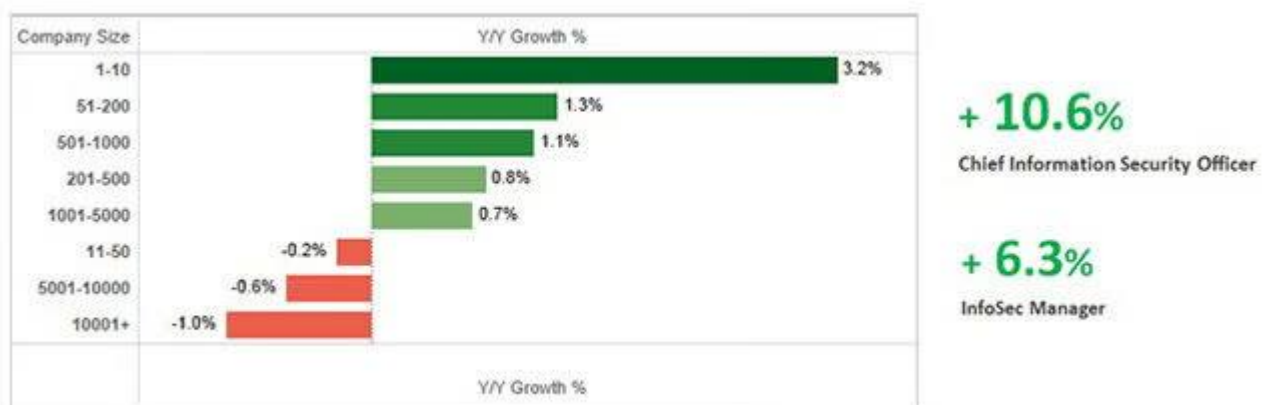
Talent Migration

It's not surprising that talent will migrate to regions with high demand; we saw significant infosec growth in Texas, Florida, Portland, Charlotte, Denver, Phoenix, and the Bay Area from June 2013 to June 2015. Many other regions in the United States are stagnating or losing staff. This is an opportunity for savvy hiring teams to reclaim talent in their home region. Several major infosec regions are outlined in the map below. Those in green have gained talent, red have lost talent, with the circle size representing the number of infosec professionals in that region. To be a true measure of talent movements, this data only considers members who have been part of the information security labor force over the entire period.

U.S. InfoSec Talent Migration June 2013 - June 2015



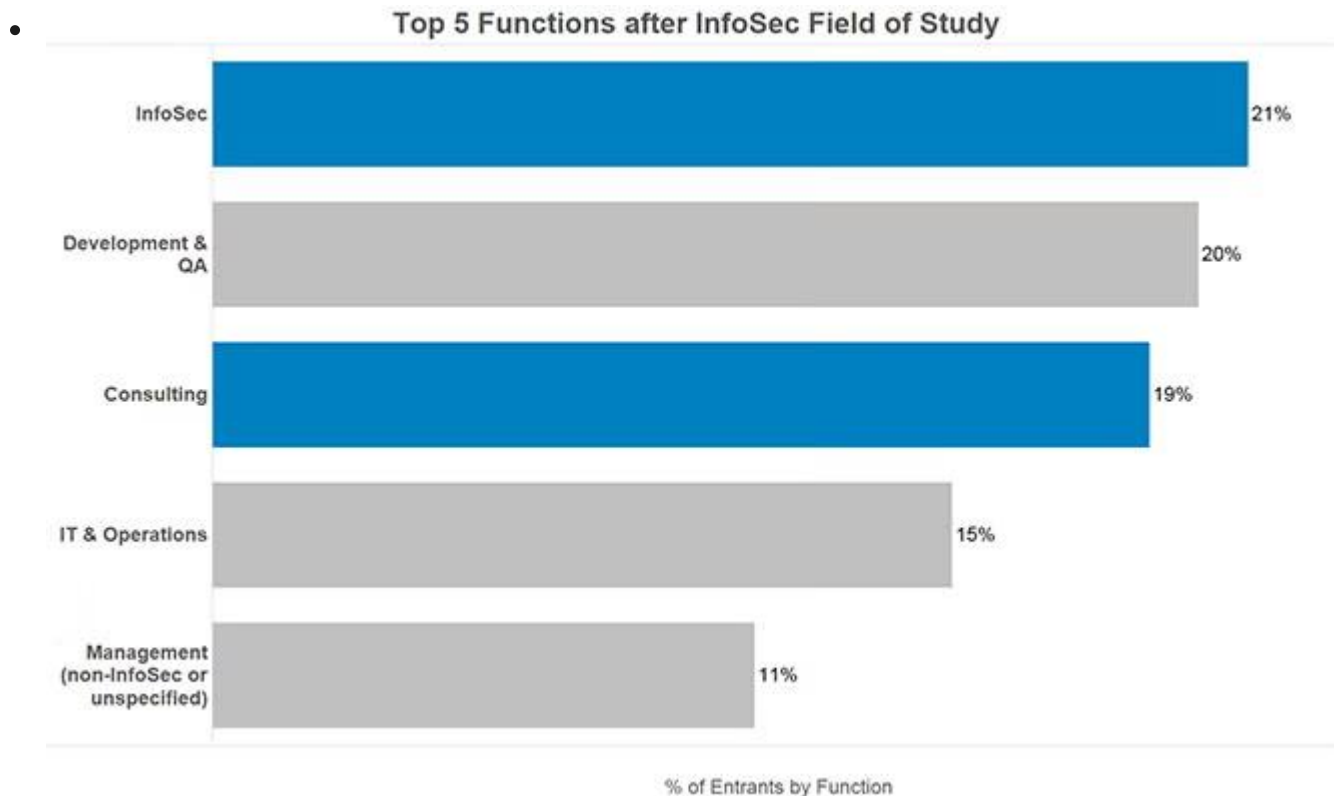
When it comes to company size, we saw significant movement from large 5000+ person corporations to small, more dynamic companies. This talent migration is especially notable when you consider the scale of these large firms. The change is also reflected in shifting roles, with 10.6% more Chief Information Security Officers and 6.3% more infosec managers since 2013. Like professionals in other functions, infosec talent is migrating to nimble companies that give them more ownership and provide new challenges.



New Entrants to the InfoSec Field

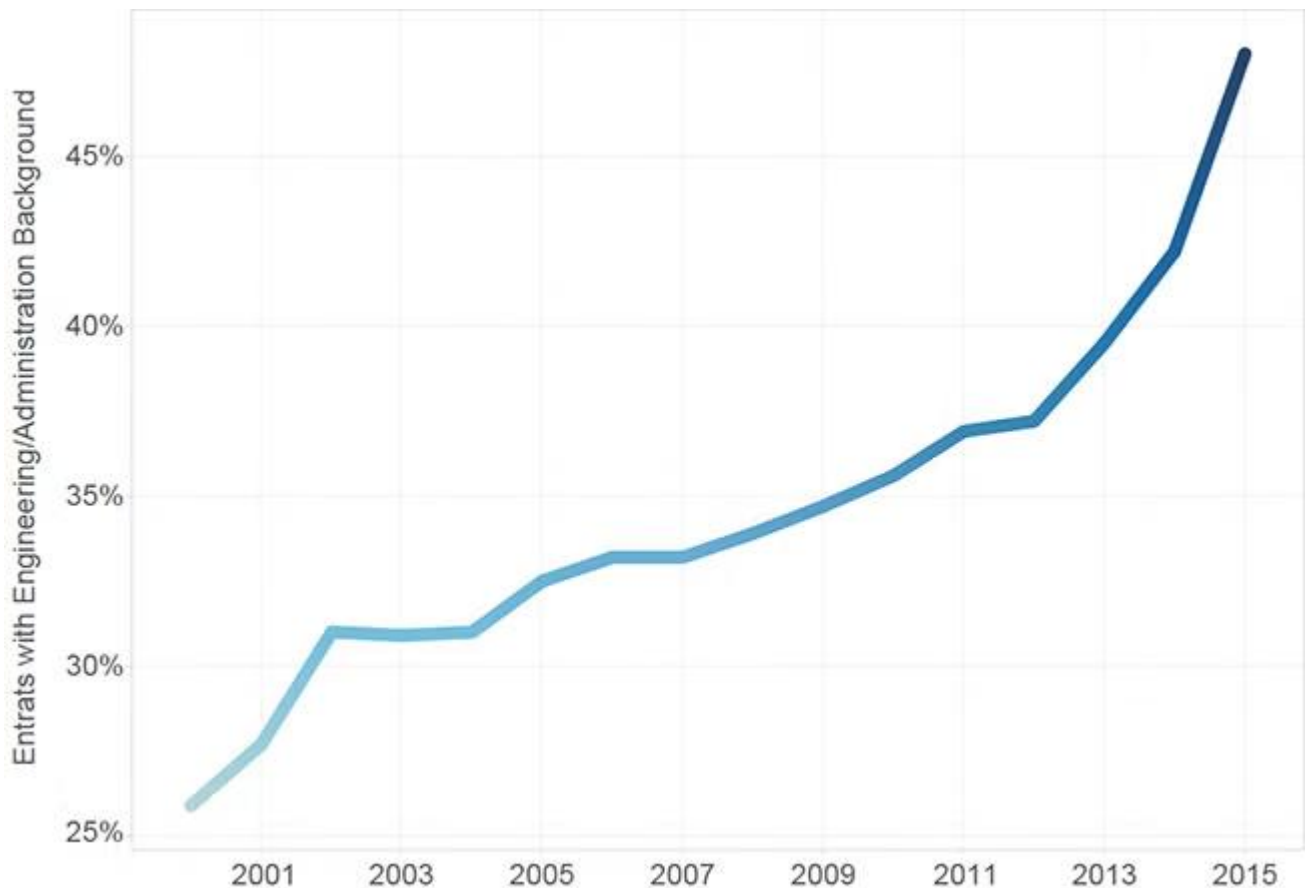
We analyzed how professionals entered the information security field, starting by identifying approximately 22,000 LinkedIn members who had information security concentrations in their field of study. Of that group, 21% went on to information security positions. 19% became consultants, likely in the information security space. Others went into development, quality

assurance, information technology, and operations in significant numbers. Having talent that understands information security working in other fields can be seen as a win for the ecosystem as a whole - companies from all industries are recognizing the value of infosec training. However, with less than half of those who complete an information security program entering information security directly, it is clear that education alone will not solve the security talent shortage.



[image: Increase in InfoSec Entrants with Engineering/System Administration from 2000 - 2015]

The good news is that an increasing percentage of entrants into the information security talent pool are coming from technical backgrounds. In 2000, about 26% of hires into the field came from engineering, development, analyst, or system/network administration roles. Today it has doubled to 48%. If you add roles such as help desk, quality assurance, and other technical operations positions, the percentage is closer to 55%. It's imperative that the vast majority of information security teams have strong technical backgrounds to properly represent this discipline when interfacing with teams internally and counterparts on the other side of the table. A technical background empowers an infosec leader to understand the impact of recommendations, suggest innovative solutions, and relate with their team in the field.



Tenure

We also explored tenure for information security professionals. The average tenure is roughly 3 years, with the tenure for management positions consistently above that average and entry level roles consistently below. The industry where an information security professional works also has an impact on tenure. Banking and financial services were right at the average, industries focused on technology were below, and other non-technology industries were generally above the average - it's possible this could be due to more established lanes of control and governance in banking and other non-technical fields.

These findings line up well with conventional industry wisdom. Information security roles in today's technology companies can be dynamic, fast-paced, and uniquely challenging, but present big opportunities for career growth for those who succeed. On the other hand, established controls and regulations in financial services and other non-tech fields can result in a more stable role.

CISOs

Our research also explored the Chief Information Security Officer (CISO) position - the de facto head of an organization's information security program. These are the ultimate decision makers when it comes to a company's security needs. Approximately half of CISOs come from a technical background and have an average of 13 years of work experience. The average tenure for CISOs is almost 4 years, and the industry tenure trends mirror those of information security staff as a whole. As mentioned above, the CISO position is rapidly expanding, particularly at smaller companies. We expect the role to be an increasingly important part of the security landscape moving forward.

Overall

In closing, we found information security to be an exciting and quickly changing field. We hope that sharing our facts and observations will help professionals make more informed career decisions as they navigate the information security landscape. Similarly, we believe that companies can use this data to determine how to best develop and identify talent, both within and outside their organization.

Please find a Slideshare presentation with additional analysis [here](#).

July 25, 2015

Turning Our CSV Connections Download Tool Back On

[Michael Korcуска](#) Vice President, Product Management

Earlier this week, we turned off our CSV connections download tool and asked our members to use another data export process as part of our ongoing efforts to combat the inappropriate export of member data by third parties. This process delivered more data but took longer, usually 24 hours but in some cases up to 72 hours. Since that change, we've heard you loud and clear -- that is too long to have to wait for a download of connection information. Effective immediately, we have turned the CSV download link back on.

Our goal is to make as much of your data, including connection data, available within minutes. We will keep the CSV connections tool available until we can reach that goal (some other data items will be available in an extended archive that may take longer to process). We will then turn this tool off again, as part of our ongoing anti-scraping efforts.

We believe that the data our members enter into LinkedIn is theirs and they should be able to export it. We are also committed to ensuring members have control of what data can be exported by their connections. In the coming weeks and months you can expect to see us take additional steps to increase that control and to make the scraping of member data by third parties more difficult. Scraping is against our [Terms Of Service](#) and potentially detrimental to the members whose data is being scraped.

We are sorry for the inconvenience this caused and resolve to do better in the future.

July 24, 2015

Access to Your Account Data

[Michael Korcуска](#) Vice President, Product Management

We have turned the tool back on. You can read more in our new post [here](#).

LinkedIn is a members-first company and we believe that the data you enter into LinkedIn is yours and you should be able to access it. We've had a bunch of questions on a change we made earlier this week regarding access to your account data and wanted to clarify a few things.

To be clear, our members can continue to export the same information from LinkedIn, including their first degree connections. We've historically had two ways for members to export portions of their data on LinkedIn. To simplify this experience we've combined them into one single process, and we include more of your activity in the export. This change is also part of a larger and ongoing effort by LinkedIn to make the scraping of member data by third parties more difficult. Scraping is

against our [Terms Of Service](#) and potentially detrimental to the members whose data is being scraped.

Our members can continue to easily request an export of their LinkedIn activity and data by visiting our help center and following [these instructions](#). This export is accessible through their Privacy & Settings page and is available as a download, typically within 24 hours although sometimes it takes longer.

Or you can take the following steps:

-

Move your cursor over your profile photo at the top right of your homepage and select "Privacy & Settings." You may be prompted to sign in.

-

Click the "Account" tab in the bottom left.

-

Click "Request an archive of your data" under the Helpful Links section.

- Click the "Request Archive" button on the right.

July 14, 2015

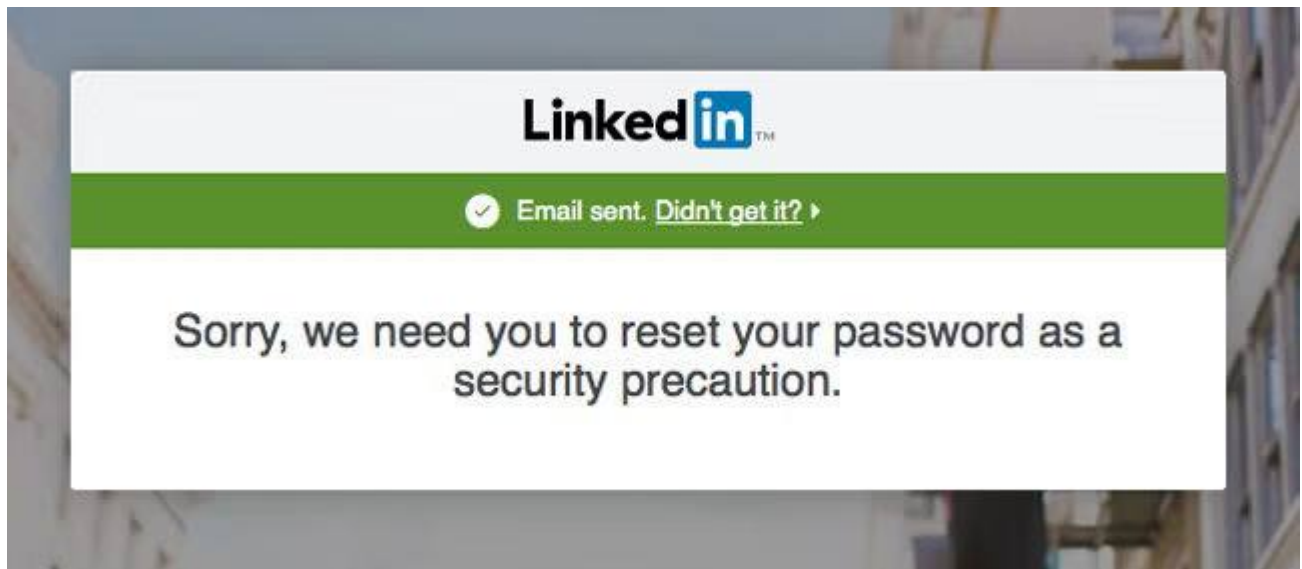
Keeping Members Secure: Credential Comparisons

[Paul Rockwell](#) Head of Trust & Safety

Professionals and companies around the world use LinkedIn to network, post and engage with others in their professional community. The role of the LinkedIn Trust & Safety team is to ensure that this positive experience for our members is not eroded by bad actors. This team puts its experience to work to make LinkedIn a safe and secure place for our members, customers, and companies to engage with each other.

One of the ways we do this is by proactively seeking out email address+password combinations that have been obtained by hackers and shared online. We do this through a number of avenues - with roughly 75% coming from vendors that scour the darknet and the remainder coming from our searches of common locations that hackers drop credentials.

Acquired credentials that are in clear text are run through a system that standardizes the format and compares the email address to those of our members. The email addresses that match a member account have the acquired password hashed and compared to the password hash for that member's account. If there's a match, the system automatically invalidates the member's password and notifies the impacted member (via email) so they can reset it. Should the member attempt to log in anyway, they'd see a message directing them to their email for a password reset.



When the team finds passwords that are hashed or encrypted, those are run through an offline system to crack the hash (if possible), and the results are then run through the same process as clear text passwords. Once we've run the credential through our systems, they're purged in an effort to further protect members.

This system has been very successful in keeping members safe, and we're finding credentials at an increased rate every year. Last year the team acquired ~100 million credentials, and in the first half of 2015 we've nearly doubled that number. Applying our credential comparison process to this specific 2014 batch helped protect over 2.5 million members accounts, and we're positioned to exceed that this year.

While we're glad that the team is able to provide this level of protection to our members, we still strongly recommend that members do not reuse passwords across websites. Additionally, enabling [Two-Step Authentication](#) will add another layer of protection, and eliminate the possibility of a stolen or non-unique password impacting their account.

June 17, 2015

LinkedIn's Private Bug Bounty Program: Reducing Vulnerabilities by Leveraging Expert Crowds

[Cory Scott](#) Director, Information Security

One of the best ways we protect our members is by identifying vulnerabilities prior to launch through a careful design review and pre-release testing. In this rapidly changing environment where we ship code multiple times a day, we also keep an eye out for vulnerabilities in production.

Our strong relationship with the security community is crucial to this process and we appreciate the work of individual researchers who contribute their expertise and time to make LinkedIn a safer place for our members. In October 2014, we formalized this partnership with the creation of LinkedIn's private bug bounty program.

The participants in our private bug bounty program have reported more than 65 actionable bugs and we have successfully implemented fixes for each issue. The participants have given us positive feedback on the program and in return for their work we've paid out more than \$65,000 in bounties.

Why a private program?

This program grew out of engagement with security researchers over the past few years. While the vast majority of reports submitted to our notification email address security@linkedin.com were not actionable or meaningful, a smaller group of researchers emerged who always provided excellent write-ups, were a pleasure to work with and genuinely expressed concern about reducing risk introduced by vulnerabilities. We created this private bug bounty program with them in mind – we appreciated working with people dedicated to coordinated disclosure practices and wanted to engage them in a deeper and mutually rewarding relationship.

Our security team works directly with each participant to handle every bug submission from beginning to end. The design of our program means that we can give the researchers who are part of our program the same experience they would have if they were on our team filing bugs right alongside us. When building our program, we recognized that logistics around payment and tracking requires a service provider. We selected HackerOne to assist us, specifically for their team's ability to manage payments, a process that requires significant diligence for tax reporting and accounting.

We did evaluate creating a public bug bounty program. However, based on our experience handling external bug reports and our observations of the public bug bounty ecosystem we believe the cost-to-value of these programs no longer fit the aspirational goals they originally had.

This private bug bounty program gives our strong internal application security team the ability to focus on securing the next generation of LinkedIn's products while interacting with a small, qualified community of external researchers. The program is invitation-only based on the researcher's reputation and previous work. An important factor when working with external security reports is the signal-to-noise ratio: the ratio of good actionable reports to reports that are incorrect, irrelevant, or incomplete. LinkedIn's private bug bounty program currently has a signal-to-noise ratio of 7:3, which significantly exceeds the public ratios of popular public bug bounty programs.

We continue to handle a significant number of vulnerabilities through security@linkedin.com and encourage anyone to report bugs. We take pride in our professional and timely response to anyone who contacts us to share a vulnerability that could impact LinkedIn and our members.

We wanted to make sure we were delivering strong results before we talked about the program; we are seeing great things so far. Sharing our different approach can also add some nuance to the dialogue that others may find useful. We'll have a lot more to say about public and private bug bounties as well as our application security program at Black Hat USA this year. Our Senior Technical Program Manager [David Cintz](#) and I will be presenting "[The Tactical Application Security Program: Getting Stuff Done.](#)" at the Briefings. Come by and see us!

May 15, 2015

A Look into Netty's Recent Security Update (CVE-2015-2156)

[Luca Carettoni](#) Information Security Manager

As part of our ongoing effort to secure our applications and contribute back to the community, the LinkedIn House Security team is often involved in security testing activities aimed at uncovering vulnerabilities in widely deployed open-source frameworks and applications.

During a recent assessment, we discovered a security flaw within [Netty's](#) cookie parsing code which leads to a universal HttpOnly bypass in Play Framework and potentially other frameworks using Netty as a dependency. The issue has been fixed in [Netty 3.9.8.Final](#), [3.10.3.Final](#), [Netty 4.1.0.Beta5](#), [Netty 4.0.28.Final](#) and [Play Framework 2.3.9](#).

In this blog post, we explain the technical details of the vulnerability and provide information related to the patch issued by the Netty team. A library upgrade is required to fully mitigate the risk.

Description

According to [RFC6265](#), each cookie begins with a name-value pair, followed by zero or more attribute value pairs. In particular, the cookie's value is defined by the following grammar:

```
cookievalue = *cookieoctet / ( DQUOTE *cookieoctet DQUOTE )
cookieoctet = %x21 / %x232B / %x2D3A / %x3C5B / %x5D7E ; USASCII characters excluding CTLs, ; whitespace DQUOTE
```

In Netty's implementation, an improperly terminated cookie value beginning with single-quote or double-quote was not being discarded. Instead, the parsing routine would consume the value till the end of the header or other special characters.

```
public Set<Cookie> decode(String header) {
    List<String> names = new ArrayList<String>(8);
    List<String> values = new ArrayList<String>(8);
    extractKeyValuePairs(header, names, values);
    ....

    private static void extractKeyValuePairs(
        final String header, final List<String> names, final List<String> values) {
        ....
        if (c == '"' || c == '\'') { //<--- If the first char is a QUOTE or DQUOTE, consume input till next QUOTE or DQUOTE
            // NAME="VALUE" or NAME='VALUE'
            StringBuilder newValueBuf = new StringBuilder(header.length() - i);
```

As a result, the value of a cookie could potentially contain the entire Cookie header regardless of cookie's flags.

The underlying vulnerability in Netty's code can be easily reproduced with the following code:

```
Set<Cookie> cookies = CookieDecoder.decode(ADD_YOUR_COOKIE_STRING_HERE);
System.out.println(cookies);
```

Providing `aaa=bbb;ccc=ddd` as input, the method will correctly return an array with two elements (two cookies).

`[aaa=bbb, ccc=ddd]`

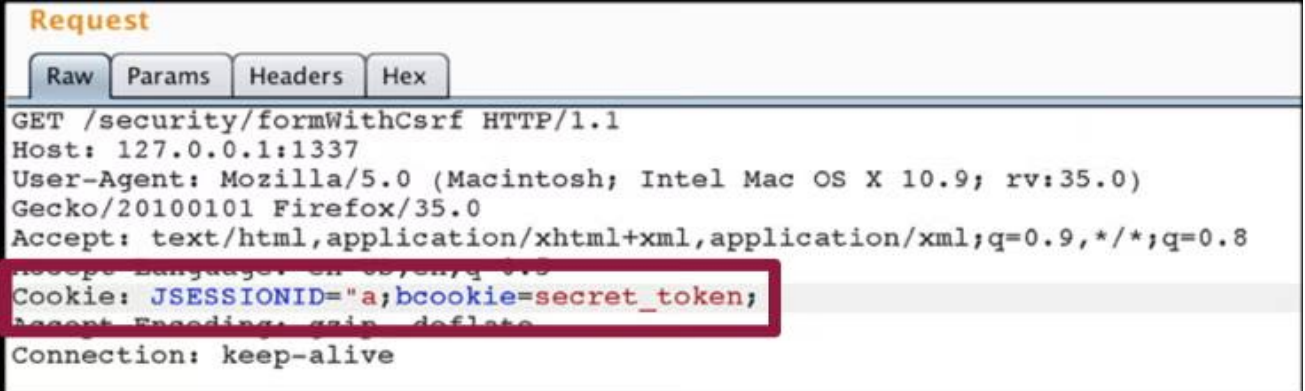
Instead, providing `aaa=bbb;ccc=ddd` as input, the bug manifests itself returning a single element in the array.

`[aaa=bbb;ccc=ddd]`

Impact on Play Framework

At first sight, the bug doesn't seem to have immediate security implications since the cookies are contained within the same HTTP request. However, under specific circumstances, this bug can be exploited to bypass `HttpOnly` flag restrictions as illustrated in the following example.

Play Framework implements CSRF protection using a cookie value (e.g. `JSESSIONID`) that is compared with a value submitted via HTTP POST parameters (e.g. `csrfToken`). The cookie is retrieved by server-side components and its value is typically injected as secret token within HTML forms and other elements. During the submission of the HTML form, both values are compared.

- 

Request

Raw Params Headers Hex

```
GET /security/formWithCsrf HTTP/1.1
Host: 127.0.0.1:1337
User-Agent: Mozilla/5.0 (Macintosh; Intel Mac OS X 10.9; rv:35.0)
Gecko/20100101 Firefox/35.0
Accept: text/html,application/xhtml+xml,application/xml;q=0.9,*/*;q=0.8
Cookie: JSESSIONID="a;bcookie=secret_token;
Connection: keep-alive
```

In the presence of the above mentioned vulnerability, a malformed cookie can be used to display the entire cookie header within the resulting HTML page.

- 

Response

Raw Headers Hex HTML Render

```
HTTP/1.1 200 OK
Content-Type: text/html; charset=utf-8
Content-Length: 513

<html>
<head>
  <title>Form without a CSRF Token</title>
</head>
<body>
  <h3>This form includes the CSRF token. Submitting should work.</h3>
  <form action="/security/submit" method="post" id="the-form">
    <input type="hidden" name="csrfToken"
value="a;bcookie=secret_token;"/>
```

In the default configuration of Play framework's CSRF protection, this vulnerability can be abused to reflect sensitive cookies, even those protected by the `HttpOnly` flag within the content of a web page.

From a web attack perspective, this is extremely interesting. In case of Cross-Site Scripting (XSS) attacks or any vulnerability which allows to set arbitrary cookies, this bug can be leveraged to universally bypass the HttpOnly flag in all Play applications. If the HttpOnly flag is included in the HTTP response header, the cookie cannot be accessed through client-side script. In this case, since the cookie's name-value is entirely included in the DOM, a malicious script can access the content without any restriction.

Next Steps

Many other projects using Netty may be vulnerable to similar "side-effects" of the incorrect cookies parsing routine. We recommend that every project relying on Netty's CookieDecoder method should mitigate the potential risk by upgrading to the latest version.

As a result of our disclosure, Netty's cookie parsing code was patched and updated to be strictly RFC6265 compliant. Netty pull request related to this patch can be found at <https://github.com/netty/netty/pull/3748>.

Credits

This vulnerability was discovered by [Roman Shafigullin](#), [Luca Carettoni](#) and [Mukul Khullar](#) and was reported to Play and Netty teams on 8th April 2015. CVE-2015-2156 has been assigned for this bug.

April 2, 2015

Welcome to the LinkedIn Security Blog

[Cory Scott](#) Director, Information Security

LinkedIn is the professional network that our members and enterprise customers use to become more productive and successful. We recognize and appreciate the trust they place in us, and work hard to make sure we continue to earn that trust.

To that end, LinkedIn has built a strong security team that focuses on application security, network and infrastructure security, and protecting members from Internet threats. We have also established a strong relationship with the security community, including working with researchers, contributing to open source projects, and presenting material at conferences.

We follow the principles of responsible disclosure when we find vulnerabilities or receive external vulnerability reports on our own products. We appreciate the work of individual researchers who contribute their expertise and time to make LinkedIn a safer place for our members. If you wish to report a vulnerability in LinkedIn's services, please follow our [guidance](#). We pride ourselves on being responsive to actionable security reports that may impact our members.

In the future, we'll use this site to share some of our security research, whitepapers on how we handle data and the security features and diligence we've built into our products. If you are responsible for information security at an enterprise that uses LinkedIn's products, I'd love to hear your concerns and thoughts anytime.

Finally, we are always on the lookout for more security talent to help us secure the economic graph and the data of our 300 million members. If you're interested in application security at scale,

or how to secure a complex and constantly evolving infrastructure using some of the newest technologies, we would love to hear from you. Take a look at our [open positions](#).

Archived from <https://security.linkedin.com/blog-archive#11232015>. Rendered offline from the local Markdown copy.