

2015 年度 CYBEX Working Group 活動報告

高橋 健志 (takeshi_takahashi@nict.go.jp)
門林 雄基 (youki-k@is.aist-nara.ac.jp)

宮本 大輔 (daisu-mi@nc.u-tokyo.ac.jp)
櫛山 寛章 (hiroa-ha@is.naist.jp)

2015 年 12 月 15 日

目 次

1	CYBEX WG 2015 年度の活動	1
2	ITU-T SG17 X.Cogent の標準化提案	1
2.1	視覚的要素	1
2.2	文章的要素	2
2.3	視覚障害者向けの情報	2
2.4	認知タスク分析	2
3	CVSS の各指標の自然言語解析による予測	3
4	今後の予定	3

1 CYBEX WG 2015 年度の活動

近年、様々なサイバーセキュリティ技術が開発されており、その普及は急務である。これまでの CYBEX WG は、様々なサイバーセキュリティ技術の先端的研究成果や運用によって得られた知見について、国際標準化活動などの様々なチャネルを通じ、多くのステークホルダーに啓蒙することを目的とした活動に従事してきた。

今年度の主な活動内容は以下の通りである。

- ITU-T SG17 X.Cogent の標準化提案
- CVSS の各指標の自然言語解析による予測

2 ITU-T SG17 X.Cogent の標準化提案

サイバー犯罪に用いられる技術、それに対抗するサイバーセキュリティ技術は著しく進歩している。その一方で、エンドユーザのセキュリティ技術への理解は追いついていない現状が指摘されている。そこで、ITU-T Study Group 17 Question 4 において、セキュリティ技術をエンドユーザに提示する画像や文章、色といった要素技術での工夫や、アクセシビリティについて調査を行い、X.Cogent という標準化提案を行った。

エンドユーザを狙ったサイバー攻撃では、エンドユーザが何を信用して良いか、何を信用してはいけないかを認識することが攻撃対策に重要である。例えば、フィッシングサイトのような攻撃の場合、正規サイトからコピーされたコンテンツを信用するのではなく、アドレスバーに表示されているセキュリティ情報や URL を信用し、これに基づいた真贋判定を行うべきである。また、標的型攻撃に用いられるメールであれば、メールの文言ではなく、メールヘッダやメールに付与されている証明書が信用に足ると考えられる。

しかし、こうした信用すべき情報は、エンドユーザにしばしば無視されており、注意をひくことができない。そこで、よりエンドユーザの注意をひき、理解を促すインタフェースの設計について考える。

2.1 視覚的要素

エンドユーザへの警告メッセージについて、ANSI Z545.4 Product safety signs and labels という標準は「危険」や「注意」はどのように表示すべきか定義さ

表 1: ウェブブラウザとアクセシビリティ

OS	ブラウザ	アクセシビリティAPI
Windows OS	Internet Explorer	MSAA and UIA
Windows OS	Firefox	MSAA and IAccessible2
Windows OS	Chrome	MSAA and IAccessible2
MacOS	Safari	AXAPI
Linux	(any)	ATK/AT-SPI

れている。例えば「危険」を表す文書では、赤色の背景色に白色のトライアングルを表示し、その中に赤色のエクスクラメーションマークを表示する。「注意」の場合は、黄色の背景色に黒字のトライアングルを表示し、その中に黄色のエクスクラメーションマークを表示する。

なお、過去には「目」のマークを警告に用いた場合、これをみた人々はより社会的に意識的に振る舞うことを見出している。このマークから人物の顔を想起することが、社会と協調活動を奨励するいわゆる「社会脳」を活性化させるという理論である。ただし、これについても懐疑的な意見はある。

2.2 文章の要素

警告メッセージの文言をどのように表記するかについて、「専門用語を避ける」「簡潔に記述する」「具体的なリスクを描写する」という点が重要視されている。例えば、Google Chrome においてサーバの証明書に問題があった場合の警告について、Chrome バージョン 36 では「不正な証明書が用いられていた」ことを正確に描写していたが、バージョン 37 以降は証明書という専門用語を用いることを避け、「この接続ではプライバシーが保護されない」という表記に置き換えた。また、警告メッセージの文字数を減らし、簡潔性を高めている。もちろん、簡潔性と正確性はトレードオフの関係にあるが、これを補うため、どのようなリスクがエンドユーザにあるのかを具体的に描写する手法を用いることにより、エンドユーザの理解を促している。

2.3 視覚障害者向けの情報

視覚障害者にとって、アクセスしているウェブサイトの証明書の有無を発見することは難しい。晴眼者にとってアドレスバーが緑色になったかどうかは理解しやすいが、この簡単なことすら今日のウェブブラウザのアクセシビリティでは達成困難である。

ウェブブラウザのベンダーは、表 1 に示すように OS の用意するアクセシビリティAPIを用い、ブラウザを作成している。これにより、視覚障害者が用いているスクリーンリーダは、ウェブブラウザが現在表示している DOM ツリーを取得可能となり、音声読み上げ機能により視覚障害者のコンテンツに対するアクセシビリティを高めている。

その一方で、フィッシングサイト対策ではコンテンツではなくアドレスバーに記載されている情報が求められている。国際標準 ISO/IEC 40500:2012 の定めるアクセシビリティのガイドラインでは、これらアドレスバーに記載されている情報に対するアクセシビリティが定義されていない。また、証明書及び証明書を発行する認証局の規定については CA/Browser Forum が Baseline Document を定めているが、証明書の内容をどのように表示するかについては言及されていない。

また、色情報について、警告では赤、黄色が使われる一方で、安全性を示す情報として緑色が用いられることが多い。緑色は可視光の波長においてちょうど中間となる値であり、認識する際に眼球の調整を必要としないという特徴がある。この一方で、赤緑色盲の状態の色覚異常者にとっては、警告と安全の区別がつきにくいこともある。緑色を他の色で代替可能とすべきという指摘はある。

2.4 認知タスク分析

信用に足ると考えられる情報を得るということは、タスク分析の学際領域であるとも考えられる。エンドユーザがインタフェースから情報を得て、「情報を得る」というタスクを行い、最終的に「信用していいかどうか意思決定する」というタスクを実施すると考えれば、エンドユーザがセキュリティ情報に対してどのようにユーザエクスペリエンスを獲得したかを分析することが可能である。

CYBEX WG では、このような領域に取り組んでいる SWAN WG と連携してこの領域の整理に取り組んでいる。具体的な内容については 2013 年度 SWAN Working Group 活動報告書を参照されたい。

表 2: CVSS の各指標の自然言語解析による予測

指標	カテゴリ	LSI	LDA	SLDA
AV	LOCAL	0.086	-	0.621
	ADJACENT NETWORK	0.067	-	-
	NETWORK	0.609	0.981	0.962
AC	HIGH	0.084	0.069	-
	MEDIUM	0.456	0.644	0.676
	LOW	0.124	0.009	0.718
AU	MULTIPLE INSTANCES	-	-	-
	SINGLE INSTANCE	-	-	0.451
	NONE	0.868	0.893	0.925
C	NONE	0.041	0.555	0.792
	PARTIAL	0.573	0.266	0.708
	COMPLETE	0.007	0.065	0.624
I	NONE	0.048	0.035	0.751
	PARTIAL	0.630	0.617	0.741
	COMPLETE	-	0.025	0.622
A	NONE	0.204	0.116	0.823
	PARTIAL	0.485	0.492	0.623
	COMPLETE	0.005	0.033	0.693

3 CVSS の各指標の自然言語解析による予測

システムの脆弱性の情報について、その重要度の高低を自動的に識別する手法について考える。近年、膨大なサイバー脅威が観測され、システム管理者は優先して対応すべき脅威が何であるかを認識することが難しくなりつつある。このため、システムの脆弱性の情報の重要さ、深刻さを評価するシステムである共通脆弱性評価システム (Common Vulnerability Scoring System, CVSS) が広く用いられている。CVSS は、脆弱性情報について、攻撃元 (Access Vector, AV)、攻撃の複雑さ (Attack Complexity, AC)、攻撃前の認証の要否 (Authentication, AU)、そして機密性 (C)、完全性 (I)、可用性 (A) への影響を三段階で評価し、脆弱性のスコアを数字で表現するための仕組みである。今日では、脆弱性が発見された場合、MITRE が脆弱性情報について共通脆弱性識別子 CVE における ID 番号及び内容を記述し、NIST がこの情報を専門家が分析し、CVSS を報告するという運用がなされている。

我々の研究グループは、脆弱性情報から脆弱性の評価指標を機械的に推測することは可能であると考えた。脆弱性の情報の多くは自然言語で記述されたテキストデータである一方で、脆弱性を評価する指標はカテゴリデータで表現されている。ここで自然言語処理技術を用いて文書を分類すれば、その結果、自然と脆弱性の評価指標の高低に沿うのではないかと予想した。

まず、2002 年 1 月～2014 年 9 月までに報告された CVE の概要に含まれる脆弱性のテキストデータを単語に分解し、ストップワード及びステミングにより余

分な単語を排除した。この上で、潜在意味解析 (Latent Semantic Indexing, LSI) 及び隠れディリクレ分布 (Latent Dirichlet Allocation, LDA) による文書分類を行った上で、CVSS の 6 種類の指標である AV, AC, AU, C, I, A ごとにある三段階どのカテゴリに分類されるかを調査し、文書分類の f_1 値を調べた。

結果を表 2 に示す。各指標ごとのカテゴリに分類された内容について、再現率及び適合率を調べて f_1 値を表示している。ハイフンは再現率または適合率が 0 となり f_1 値が計算できなかったことを示す。結果として、LSI, LDA とともに AV, AU についてはある程度の予想ができるものの、C, I, A についてはさほど正しい予想ができていとは言えない。これらの手法はテキストデータの単語数の増加にともなって次元数が増加し、ひいては計算量が増加するという問題に対し、次元削減手法により計算量の効率化を実現している。しかし、評価の予測に必要な重要な単語が小さな情報量しか持っておらず、これが削減されることにより文書の予測が正確に行われな可能性はある。我々は今回のデータはこのケースに該当したのではないかと考えた。そこで、LDA を教師あり学習に対応させた SLDA (Supervised LDA) により文書分類結果を同じく表 2 に試す。性能は大きく改善され、高確率で各指標が予測できることがわかった。

本研究のより詳細な内容については、文献 [1] を参照されたい。

4 今後の予定

今年度は標準化活動を通じ、あるいは標準化された技術の新たな利用方法を考えることにより、幅広いステークホルダーにセキュリティに関する知識を啓蒙していく研究を行った。来年度も継続し、CYBEX WG のみならず他の WG、あるいは全世界で創出される素晴らしいセキュリティ技術を、より多くのステークホルダーに広めていく活動を継続していきたい。

参考文献

- [1] Daisuke Miyamoto, Yasuhiro Yamamoto, Masaya Nakayama, “Text mining-based Approach for Estimating Vulnerability Score,” In *Proceedings*

*of the 4th International Workshop on Building
Analysis Datasets and Gathering Experience Re-
turns for Security (BADGERS), November 2015.*