

ORIGINAL ARTICLE

A statistical approach for assessing cyber risk via ordered response models

Silvia Facchinetti¹ | Silvia Angela Osmetti¹  | Claudia Tarantola²

¹Department of Statistical sciences, Università Cattolica del Sacro Cuore, Milan, Italy

²Department of Economics and Management, University of Pavia, Pavia, Italy

Correspondence

Silvia Angela Osmetti, Department of Statistical sciences, Università Cattolica del Sacro Cuore, Largo Gemelli 1, 20123, Milan, Italy.
Email: silvia.osmetti@unicatt.it

Funding information

COST Action FinAI, Grant/Award Number: CA19130

Abstract

Proper evaluation of the risk associated to a cyber attack is a crucial aspect for many companies. There is an increasing need to plan for and implement effective ways to address cyber security, data security, and privacy protection. Estimating the risk of a successful cyber attack is an important issue, since this type of threat is proliferating and thus poses increasing danger to companies and the customers who use their services. While quantitative loss data are rarely available, it is possible to obtain a qualitative evaluation on an ordinal scale of severity of cyber attacks from experts of the sector. Hence, it is natural to apply order response models for the analysis of cyber risk. In particular, we rely on cumulative link models. We explain the experts' assessment of the severity of a cyber attack as a function of a set of explanatory variables describing the characteristics of the attack under consideration. A measure of diffusion of the effects of the attacks obtained via the use of a network structure is also incorporated into the set of explanatory variables of the model. Along with the description of the methodology, we present a detailed analysis of a real data set that includes information on serious cyber attacks occurred worldwide in the period 2017–2018.

KEYWORDS

cumulative link model, cyber risk, marginal effect, social network analysis

1 | INTRODUCTION

Cyber risk can be defined as “any risk emerging from the use of Information and Communication Technologies systems (ICT) that compromises the Confidentiality, Integrity, and Availability (CIA) of data or services”; see, for example, Cebula and Young (2010), Edgar and Manz (2017), and Kopp et al. (2017) for a more detailed description. It is an operational risk¹ that arises from an external or internal attacker compromising a computer database or network, or from transactions on the internet (Mood, 2005). A cyber risk can be associated both to an event with a criminal intent (cyber attack) or without any criminal intent (i.e., IT outages due to a software update, or weather events that can lead

to business disruptions). In the following, when we refer to cyber risk, we mean risk due to a cyber attack.

In an increasingly digitized world, where organizations are affected by technological evolution, cyber attacks are multiplying rapidly. They have an impact on every class of business, and no industry can consider itself completely immune to the rising number of cyber attacks. According to World Economic Forum Global Risks Report (2019), cyber risks are consolidating their position alongside environmental risks in the high-impact/high-likelihood quadrant of the Global Risks Landscape, and according to the International Monetary Fund, they have become an increasing concern for policy makers (Bouveret, 2018).

To ensure that the CIA triad is enforced for information systems and that breaches are minimized, companies need to identify the key vulnerable assets that are exposed to cyber risk and to design, monitor, and improve information security. Chief technology officers generally invest a portion of their

¹ Operational risk has been defined, by the Basel Committee on Banking Supervision, as “the risk of a monetary loss caused by human resources, ICT, by organization processes or by external events.”

budget in IT security (i.e., investment in perimeter security elements, IT auditing, business continuity processes, and disaster recovery) that would automate and accelerate the threat defence (Kolfal et al., 2013).

Data are becoming more and more important in predicting future risks and disruptions to global activities, and in assessing vulnerabilities, see Choi and Lambert (2017). Institutions should be encouraged to collect data on cyber incidents in order to use statistical approaches to estimating the capital needed to cover losses due to the occurrence of cyber attacks. These attacks spread more quickly than other crimes and cause monetary losses, as well as impact the opportunity cost, market capitalization, and brand image; see, for example, Cavusoglu et al. (2004) and Hartwig and Wilkinson (2014).

Cyber security is a field with a growing amount of research in many contexts such as computer science, law, and business management, as well as technology sectors that did not originally involve the internet (e.g., smart grids and cars). We refer to Ramirez and Choucri (2016) and related references for a detailed cross-disciplinary review. Nevertheless, quantitative models for cyber risk measurement are still limited and, as discussed in Allodi & Massacci (2017), there is a lack of a shared framework for the quantification of risk that makes the adoption of comparable measures for risk mitigation very difficult. Relevant works include studies on the quantification of the impact of an attack—see, for example, Davis et al. (2009), Yayla and Hu (2011), and Romanosky (2016)—and measurement of attack likelihood (see Allodi & Massacci, 2017, and Cherdantseva et al., 2016, among others).

To our knowledge, there is a limited number of statistical papers addressing the problem of cyber risk assessment; this is due to the lack of quantitative data available, see, for example, Afful-Dadzie and Allen (2017). Indeed, cyber loss data are very difficult to obtain since these data are very sensitive. It is unlikely that an institution is willing to disclose them, since its reputation and, possibly, its security, may be at stake. Available public and commercial data sets exist, but they are incomplete, have different coverage, and use different definitions of cyber attacks, which makes the analysis of cyber losses difficult. To encourage data disclosure, the information is often collected on an ordinal scale, providing a classification of cyber risk in terms of its severity level. These ordinal measurements do not provide the actual magnitude of cyber attacks, but can be used to identify which types of attacks are the most dangerous.

Currently, there is no internationally recognized standard classification of cyber attack severity levels. It varies according to whom it is providing the classification. Usually, the victim of a cyber attack tends to minimize its gravity, in order to avoid pecuniary compensations and penalties or losses in terms of reputation. Customers of the affected company usually tend to exaggerate in order to get a higher compensation. Finally, national authorities tend to be quite conservative, due to political or social reasons. Here, we consider a classification of cyber attack severity levels provided by a pool of experts of the principal Italian authority in the field of cyber

security (Clusit²). They introduce an ordinal classification of cyber risk severity based on available information regarding the attack and their expertise.

Focusing on the quantification of cyber attack impact in terms of its severity level, we propose to model cyber risk via an ordered response model (a regression-type model with an ordinal response variable). In particular, we rely on cumulative link models that allow us to express the cumulative probabilities associated with the different severity levels as a nonlinear function of a suitable set of explanatory variables. In the analysis, we pay particular attention to the interpretation of the effect of each explanatory variable on the severity level of cyber attacks. For a review on how to interpret effects in cumulative link models, see Agresti and Tarantola (2018). We apply our model to a real data set that includes information on serious cyber attacks occurred worldwide in the last years.

The plan of the article is as follows. Section 2 contains the description of the examined cyber risk data. Specifically, in Section 2.1, the definition of the ordinal variable *Severity* is provided. This variable describes the severity levels of the examined cyber attack, and it will be used as a response variable in the examined cumulative link model. In Section 2.2, we describe the variables that will be used as explanatory ones in our model. In particular, in Section 2.2.1, we focus on the set of qualitative features of the examined attacks that may influence their severity. We further integrate this set of explanatory variables with a quantitative indicator of the level of systemic risk (variable *Closeness*). The definition, construction, and interpretation of the variable *Closeness* are provided in Section 2.2.2. The cumulative link model in its general formulation is presented in Section 3, while its formulation for the cyber risk data is described in Section 4 together with a discussion of the obtained results. The last section contains some concluding remarks.

2 | CYBER RISK DATA

We consider a data set collected by the researchers of the *Hackmanac* society³ and described by Clusit in its Report on ICT Security in Italy of 2019 (first semester; Antonielli et al., 2019). Clusit is the largest and most respected Italian association in the field of cyber security. It was created in 2000 at the University of Milan, and it includes important companies in different fields such as banking and insurance, public administration, health, telecommunications, and informatics, among others. Since 2012, Clusit publishes on biannual basis a report on “serious” cyber attacks that occurred worldwide in the previous year; in the period 2011–2018, it registered 8417 attacks.

² Associazione Italiana per la Sicurezza Informatica (<https://clusit.it>).

³ *Hackmanac* is a society based in Dubai that monitors the evolution of real global cyber threat by the analysis and classification of several open sources. It collaborates with Clusit performing the analysis for their half-yearly report. More details can be found at the webpage <https://hackmanac.com/>.

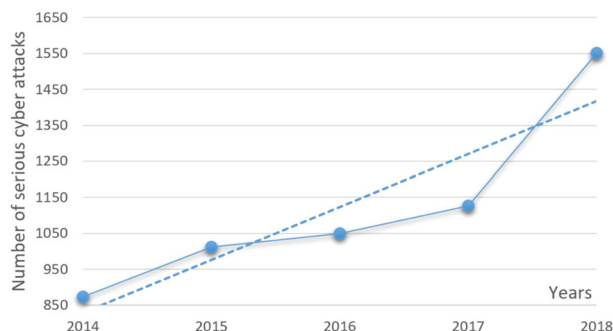


FIGURE 1 Number of “serious” attacks for years 2014–2018.

The classification criteria used by experts to identify cyber attack as “serious” have evolved during the years. For example, attacks that were considered “serious” in 2011–2013 are nowadays considered ordinary problems (i.e., the “defacements” of websites). According to the 2019 report, an attack is defined as “serious” if it had a significant impact on the victims in terms of economic losses, damages to reputation, and/or dissemination of sensitive data. This report shows that in the last few years, the number of cyber attacks has surged: 1552 attacks occurred in 2018, 1127 in 2017, 1050 in 2016, 1012 in 2015, and 873 in 2014, with a growth of approximately 45% between 2014 and 2018, as shown in Figure 1. The data examined in the report may represent only a partial situation that is less critical than the actual one, since many attacks may not be disclosed, or may be disclosed at a much later date.

In 2017, the experts of the *Hackmanac* developed for Clusit an ordinal classification of cyber risk assessment in terms of attack severity levels (medium, high, critical) on the base of their expertise. In the following, we refer to this ordinal classification as *Severity* variable. Since an ordinal classification of the severity level of each cyber attack has been introduced only in 2017, for our analysis, we focus on a sample of $n = 2679$ attacks occurred in the period 2017–2018.

The aspects that determine the risk assessment of each attack in terms of its severity level are multiple, and include the geopolitical, social, economic, image, and cost/opportunity impact on the victims. The geopolitical impact is considered relevant if governmental or national security institutions (or high representatives) are involved. The economic impact is measured in terms of the amount of estimated damages. It also includes the damage in terms of image and cost/opportunity impact on the victims. The economic impact is classified in three classes according to the amount of estimated damages (in dollars): tens/hundreds of thousands of USD, millions/tens of millions of USD, hundreds of millions/billions of USD. The evaluation in terms of social impact is based on the number of individuals involved: one individual/hundreds of individuals, thousands/hundreds of thousands of individuals, millions/hundreds of millions of individuals.

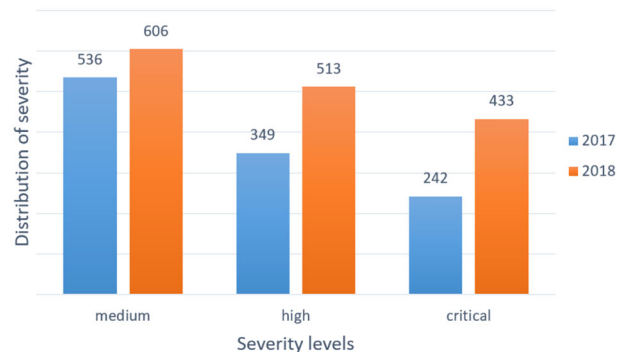


FIGURE 2 Distribution of variable *Severity* for the periods 2017 and 2018.

As stated in the Introduction, we use a cumulative link model (a particular type of regression model with an ordinal response variable) to explain variable *Severity* (Section 2.1) as a function of the following explanatory variables (Section 2.2):

- a set of qualitative variables (from the examined data set) describing peculiar characteristics of the observed cyber attacks (Section 2.2.1);
- an additional quantitative variable named *Closeness* indicating dependence, in terms of vulnerability, among the victims of cyber attacks (Section 2.2.2).

The general formulation of the cumulative link model and its properties are provided in Section 3, while its specific formulation with reference to cyber risk data is described in Section 4.

2.1 | Response variable *Severity*: The severity evaluation of a cyber attack

The expert evaluation of the severity level of a cyber attack can be considered a realization of a random ordinal variable S , named *Severity*, assuming K increasing values. In our context, S assumes $K = 3$ levels: 1 = medium severity, 2 = high severity, and 3 = critical severity.

Figure 2 shows the distribution of the severity level of cyber attacks for 2017 and 2018. We notice that the situation worsened from 2017 to 2018: While the number of medium-level attacks slightly increased (+13%), that of high-impact attacks significantly increased (+47%) and that of critical attacks almost doubled (+79%).

2.2 | Explanatory variables

2.2.1 | Qualitative explanatory variables

For each cyber attack, the data set reports information on the following qualitative variables: *Type of Attack*, *Attack*

TABLE 1 Description of the examined variables of the cyber risk data set.

Variable	Category	Frequency
<i>Type of Attack</i> (T- ...) (4 categories)	Cybercrime (T-Cyb)	77.98%
	Espionage/Sabotage (T-Esp)	12.39%
	Hacktivism (T-Hac)	5.23%
	Information Warfare (T-Inf)	4.40%
<i>Attack Technique</i> (A- ...) (5 categories)	0-day (A-Ody)	1.19%
	Multiple Threats/APT (A-Mul)	6.01%
	SQL Injection (A-SQL)	0.30%
	Trivial Threats (A-Tri)	66.93%
	Unknown (A-Unk)	25.57%
<i>Continent</i> (C- ...) (6 categories)	Africa (C-Afr)	1.08%
	America (C-Ame)	44.08%
	Asia (C-Asi)	10.90%
	Australia/Oceania (C-Aus)	1.75%
	Europe (C-Eur)	14.74%
	Multiple Continents (C-Mul)	27.44%
<i>Victim</i> (V- ...) (19 categories)	Automotive (V-Aut)	0.49%
	Banking/Finance (V-Ban)	10.19%
	Chemical/Medical (V-Che)	0.04%
	Critical Infrastructures (V-Cri)	3.62%
	Entertainment/News (V-Ent)	8.10%
	GDO/Retail (V-Gdo)	2.35%
	Gov-Mil-LEAs-Intelligence (V-Gmi)	16.09%
	Gov.Contractors/Consulting (V-Gco)	0.75%
	Health (V-Hea)	8.92%
	Hospitality (V-Hos)	2.95%
	Multiple Targets (V-Mul)	19.63%
	Online Services/Cloud (V-Onl)	8.36%
	Organization/ONG (V-Org)	0.97%
	Others (V-Oth)	2.61%
	Religion (V-Rel)	0.15%
	Research-Education (V-Res)	6.72%
	Security Industry (V-Sec)	0.52%
	Software/Hardware Vendor (V-Sof)	6.64%
	Telco (V-Tel)	0.90%

Technique, *Continent* (indicating where the attack took place), and *Victim* (indicating the sector affected by the cyber attack). A synthetic description of these variables is reported in Table 1. For more details on *Type of Attack* and *Attack Technique*, see the Appendix.

In our analysis, we consider only the qualitative variables *Type of Attack*, *Attack Technique*, and *Continent* as explanatory variable of the examined cumulative link model. Variable *Victim* (the sector affected by the attack) is not considered as an explanatory one, but it will be considered indirectly in the construction process of the variable *Closeness* as described in the following section.

2.2.2 | Explanatory variable *Closeness*: Definition, construction, and interpretation

We extend the set of explanatory variables with variable *Closeness*, a quantitative indicator of the level of systemic risk. Following the approach presented in Giudici (2018), we use social network analysis instruments to define this new variable starting from the information provided by variable *Victim*.

Variable *Closeness* is constructed via the following three steps procedure that will be described in detail in the following paragraphs.

- Step 1:** Calculation of the weekly “Criticality index” time series for each category of *Victim*.
- Step 2:** Construction of the “Network” among *Victim* categories.
- Step 3:** Calculation of the “Percentage Closeness Centrality Measure” for each category of *Victim* (*Closeness* values). Addition to the original data set of a column (*Closeness* values) containing for each category of *Victim* the corresponding value of the percentage closeness centrality measure.

We now proceed with a detailed description of the previous steps. In the following, for simplicity, we use the word victim to indicate a specific category of variable *Victim*.

Step 1: Calculation of the weekly “Criticality index.”

For each single victim, we consider the time series of the weekly Criticality indexes. The Criticality index proposed by Facchinetti and Osmetti (2018) and Facchinetti et al. (2019) is based on the relative cumulative frequencies $\hat{F}_k$ of cyber attacks suffered by a victim for $k = 1, \dots, K$ increasing levels of severity (in our analysis, $K = 3$ as shown in Section 2.1). For a specific victim v_l and for a specific week W , this index is calculated as

$$\hat{I}_W^{v_l} = 1 - \frac{\sum_{k=1}^K \hat{F}_k - 1}{K - 1}. \quad (1)$$

The index can be used to provide an indication of vulnerability of the victims. It assumes values in the interval $[0,1]$ (see footnote ⁴ for more detail). The higher the value of the Criticality index, the higher is the vulnerability of the corresponding victim.

Step 2: Network construction.

We first recall some basic notation and terminology of social network analysis that will be useful for the following part; see, for example, Wasserman and Faust (1996) for more details. A social network is represented by a weighted graph $G = (V, E)$, where the set of nodes $V = \{v_l; l = 1, \dots, |V|\}$ denotes the units under consideration, and the set of edges E indicates the connection between the different units. All edges are undirected, that is, if $(v_l, v_j) \in E$ also $(v_j, v_l) \in E$. We associate to each edge a measure of the strength of the association between the different units (w_{lj} weight of the edge). The width of each single edge is proportional to the corresponding weight. The weights can be stored in a matrix called the adjacency matrix of the weighted graph. If nodes v_l and v_j are adjacent in the graph (they are joined by an edge), we observe w_{lj} in position (l, j) of the adjacency matrix. A path connecting two nodes v_l and v_j is a set of nodes

$v_l = \tilde{v}_0, \dots, \tilde{v}_h, \dots, \tilde{v}_m = v_j$ such that for $h = 1, \dots, m$, $(\tilde{v}_{h-1}, \tilde{v}_h)$ is an edge of the graph.

In our study, the nodes represent the victims (sectors) affected by cyber attacks. Each pair of victims, v_l and v_j , is then connected by an edge in the network if the two corresponding criticality series of values present a nonzero partial correlation. The width of the edge (v_l, v_j) is proportional to the absolute partial correlation among the Criticality index time series.

In the left panel of Figure 3, we report the lower triangular part of the absolute partial correlation matrix (without the main diagonal), while in the right panel, we represent the obtained network structure. The grayscale and the width of the edges in the network show how strong such correlation is. High values of the absolute partial correlation indicate a strong dependence in terms of vulnerability of the victims. Low values imply that the vulnerabilities of two victims are not related. For the implementation, we used the R package “qgraph” (Epskamp et al., 2021).

Step 3: Variable *Closeness*.

To define variable *Closeness*, we rely on the closeness centrality measure for weighted graphs proposed by Opsahl (2010). This measure scores each node based on its distance to all the other nodes of the network, considering also the strength of the connection.

To proceed, first, we need to introduce the concepts of distance of two nodes in a path and the notion of shortest path among them. The distance between nodes v_l and v_j on a specific path $(v_l = \tilde{v}_0, \dots, \tilde{v}_h, \dots, \tilde{v}_m = v_j)$ is given by $d(l, j) = (\frac{1}{|w_{01}|} + \dots + \frac{1}{|w_{(h-1)h}|} + \dots + \frac{1}{|w_{(m-1)m}|})$, where $w_{(h-1)h}$ is the weight associated to the edge (v_{h-1}, v_h) for $h = 1, \dots, m$. The shortest path among nodes v_l and v_j is the one with the minimal distance $d^*(l, j)$ between the two nodes, that is, $d^*(l, j) = \min\{d(l, j)\}$.

The percentage closeness centrality measure c_l of node v_l is defined as

$$c_l = \frac{1}{\sum_{j \neq l} d^*(l, j)} \times 100\%. \quad (2)$$

Closeness is the variable assuming, for each victim l , values c_l calculated via Equation (2). *Closeness* values obtained for each single victim of our data set are reported in Table 2. The higher the value of *Closeness*, the higher the strength of the connection of a victim in the network.

If a victim is strongly correlated (positively or negatively) with many others, it is also highly interconnected to them in the network (high value of *Closeness*). This could indicate a connection of the victims in terms of vulnerability. Therefore, an increase in vulnerability of a victim can set alarm bells for the interconnected ones. On the other hand, we observe a low effect in case of weak correlation (correlation coefficient near zero) among the examined victim and other victims in the network.

⁴ $\hat{I}_W^{v_l} = 0$ when the cyber attacks are concentrated only on the lowest severity level or when the attacks do not occur for the specific victim in the examined week; $\hat{I}_W^{v_l} = 1$ when the attacks are concentrated only on the highest severity level. In all other situations, the index assumes any value between 0 and 1.

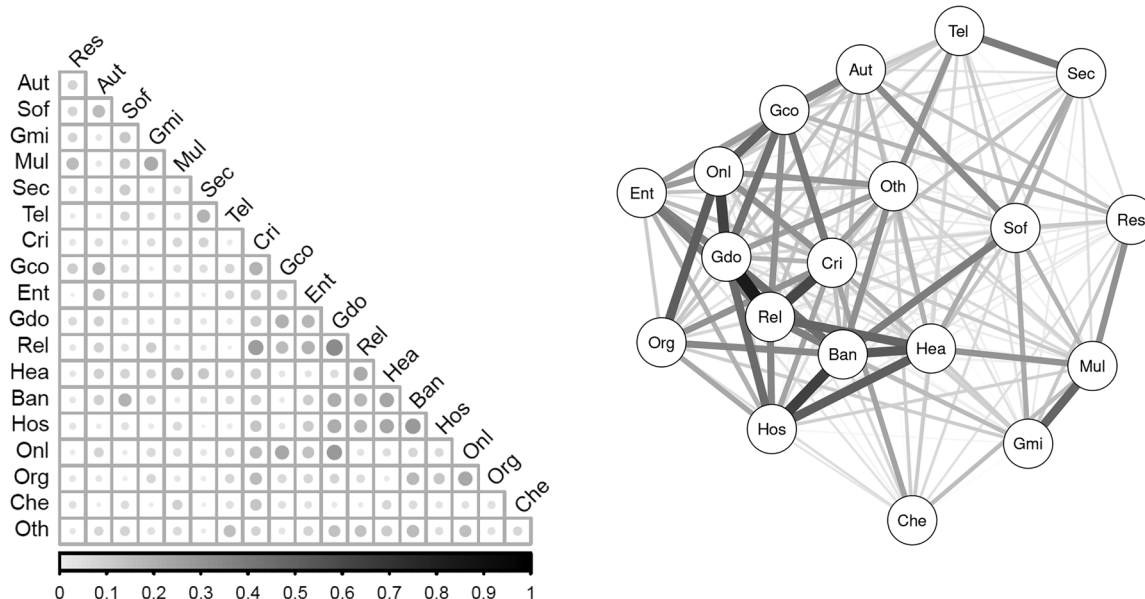


FIGURE 3 Lower triangular part of the absolute partial correlation matrix without the main diagonal (on the left) and the corresponding network structure (on the right).

TABLE 2 Closeness for the victims of cyber attacks.

Victim	Frequency of attack	Closeness
V-Aut	13	0.467
V-Ban	273	0.611
V-Che	1	0.380
V-Cri	97	0.578
V-Ent	217	0.492
V-Gdo	63	0.628
V-Gmi	431	0.414
V-Gco	20	0.559
V-Hea	239	0.620
V-Hos	79	0.564
V-Mul	526	0.447
V-Onl	224	0.558
V-Org	26	0.503
V-Oth	70	0.500
V-Rel	4	0.660
V-Res	180	0.369
V-Sec	14	0.396
V-Sof	178	0.520
V-Tel	24	0.377

3 | METHODOLOGICAL APPROACH: THE CUMULATIVE LINK MODEL

3.1 | Model specification

The cumulative link model is the most popular model for ordinal response data, see Agresti (2010). In this section, we

present its general formulation, and we refer to Section 4 for its application to cyber risk data.

Let R be a K -category ordinal response variable. For any i -th unit, $i = 1 \dots n$, a nonlinear link function (i.e. logit, probit, ...) is used to express the cumulative distribution of R_i as a function of a set of explanatory variables, that is,

$$\text{link}[P(R_i \leq r)] = \text{link}[F(r)] = \alpha_r - \mathbf{x}_i \boldsymbol{\beta}$$

$$r = 1, \dots, k, \dots, K-1; \quad i = 1, \dots, \{, \dots, n. \quad (3)$$

In Equation (3), r is the observed value of R_i , $F(\cdot)$ is the cumulative distribution function, α_r is the intercept, $\boldsymbol{\beta}$ is the (column) vector of the regression parameters, $\mathbf{x}_i$ is a (row) vector containing the values of the explanatory variables, and link is a suitable link function. The larger the value of $\mathbf{x}_i \boldsymbol{\beta}$, the higher the probability to obtain a higher level of R_i in terms of its ordinal scale.

The set of explanatory variables may consist of both quantitative and qualitative variables. The vector of explanatory variables $\mathbf{x}_i$ is obtained concatenating the following elements: $\mathbf{y}_i = [y_{i1}, \dots, y_{ih}, \dots, y_{ip}]$ a vector containing, for unit i , the observed values of p quantitative variables, and a q -dimensional factor vector $\mathbf{z}_i = [\mathbf{z}_{i1}, \dots, \mathbf{z}_{ih}, \dots, \mathbf{z}_{iq}]$ corresponding to the qualitative variables.

Indeed, qualitative variables are entered in the model as factors. For any qualitative variable h , we construct a set of dichotomous indicators (one for each level except for baseline one) $\mathbf{z}_{ih} = [z_{ih}^1, \dots, z_{ih}^\ell, \dots, z_{ih}^{(L_h-1)}]$, where $h = 1, \dots, q$ and L_h is the number of possible levels of the examined qualitative variable. The indicator z_{ih}^ℓ , $\ell = 1, \dots, (L_h - 1)$, assumes value one if the examined variable has level ℓ and zero otherwise. If all z_{ih}^ℓ are equal to zero, it means that in the model we are considering such a variable at its baseline level.

The (column) vector of the regression coefficients β can also be splitted in two parts, $\gamma = [\gamma_1, \dots, \gamma_h, \dots, \gamma_p]$ corresponding to the quantitative variables and $\lambda = [\lambda_1, \dots, \lambda_h, \dots, \lambda_q]$ to the qualitative ones, with $\lambda_h = [\lambda_h^1, \dots, \lambda_h^\ell, \dots, \lambda_h^{(L-1)}]$.

Note that β does not depend on r ; that is, the model assumes that the effect of the each explanatory variable is identical for all the considered cumulative probabilities. McCullagh (1998) named this assumption the proportional odds (PO) assumption. A discussion about the validity of the PO assumption in our context is presented in Section 4.2.

3.2 | Model estimates interpretation: Marginal effect measures

Even if the structure of model (3) resembles that of an ordinary linear model, the use of a nonlinear link function produces effects on the link scale that are not straightforwardly interpretable. For example, in a cumulative logit model, for the level ℓ of a generic qualitative variable indexed by h ($h = 1, \dots, q$), the corresponding parameter $-\lambda_h^\ell$ measures the difference between logits of cumulative probabilities of the examined category with respect to the baseline level; for a generic quantitative variable indexed by h ($h = 1, \dots, p$), $-\gamma_h$ measures the change in the cumulative logit per one-unit increase in the examined variable, adjusting for the other ones. Furthermore, the partial effect of one variable can be modified by the inclusion in the model of a new variable uncorrelated with it. On the other hand, the partial effect is identical in a standard linear regression model. See Agresti (2013) and Agresti and Tarantola (2018) for a critical discussion.

The effect of each explanatory variable on the response variable R_i can be more intuitively explained in terms of the so-called marginal effect (ME) measures; see Agresti and Tarantola (2018), Greene (2008), Long and Freese (2014), Long and Mustillo (2018), among others. The ME indicates how a change in a specific explanatory variable affects the response variable, holding constant the value of all the other explanatory variables. They are intuitive and can be calculated from any type of explanatory variable. The ME measures can be interpreted similarly as regression coefficients in a standard regression model. For quantitative continuous variables, they measure the instantaneous rate of change. For quantitative discrete variables or qualitative variables, they measure the discrete change with respect to the baseline level. They indicate how the probability of being in a particular level changes when we move from the examined level to the reference one. As an exemplification, the ME of a quantitative variable for level r of variable R_i is the partial derivative of $P(R_i = r)$ with respect to the examined variable, holding all the other variables constant. While the ME for level ℓ of a qualitative variable measures the change in the probability $P(R_i = r)$ when the level changes to the baseline one, and the other variables assume specific values. Naturally, these effects can be calculated for any value of r .

Depending on the way we fix the values of the other explanatory variables, we can obtain three different types of ME measures. The average marginal effect (AME), as the name suggests, is obtained by calculating the ME of a specific explanatory variable for each observation in the sample, and then averaging across all values. Alternatively, one can compute the marginal effects at representative (MER) values; that is, the ME measures are computed by choosing representative values of the other explanatory variables (i.e., values of particular interest for the considered problem). Finally, the marginal effects at the mean (MEM) are computed with all remaining explanatory variables held at their mean.

Among these summary measures, Long and Freese (2014) recommended the use of the AME since it can be interpreted as the sample average of the ME. Furthermore, AME measures are quite stable when we add an explanatory variable to the model that is uncorrelated with the variable whose effect we are describing, see Mood (2010). For an additional discussion on ME measures, see Agresti and Tarantola (2018) and Sun (2015) among others.

We use the standard function `polr` of the R package `Mass` to fit the model. For quantitative and binary explanatory variables, one can use the `ocAME` function by Agresti and Tarantola (2018), which supplies the AME using output from the `polr`. We use an extension of `ocAME` (`ocAME_CAT`) that deals with qualitative variables with more than two levels. This new function can be requested via e-mail from the authors.

3.3 | Link function

We now describe the principal types of link functions used for cumulative link models. In `polr` standard link functions are examined (*logit*, *probit*, *log-log*, and *cloglog*). It is possible to extend `polr` to incorporate more flexible and asymmetric link functions, such as the *bgeva* link function proposed by Calabrese and Osmetti (2013) and Calabrese et al. (2016). The *bgeva* link function is defined as the quantile function of a generalized extreme value (GEV) random variable, with a tail parameter $\tau \in \mathfrak{R}$, which regulates the shape of the function. Depending on the value of τ , several special cases can be recovered; for example, when $\tau \rightarrow 0$, the log-log link function is obtained; see Agresti (2013) and Calabrese and Osmetti (2013). This link function is particularly suitable for unbalanced samples; see Andreeva et al. (2015). A cumulative *bgeva* model for ordinal outcomes could be obtained by using the *bgeva* link function in Equation (3). The list of the examined link functions is reported in Table 3.

4 | EMPIRICAL RESULTS

We now apply a cumulative link model to our data to identify the variables that most influence the severity of a cyber attack. We consider the ordinal variable *Severity S* (Section 2.1), with $K = 3$ levels, as response variable and variables *Type*

TABLE 3 Link functions.

Link function	Form
logit	$\log\left[\frac{F(r)}{1-F(r)}\right]$
probit	$\Phi^{-1}[F(r)]$
log-log	$\log[-\log(F(r))]$
cloglog	$\log[-\log(1-F(r))]$
cauchy	$\tan\left[\pi\left(F(r) - \frac{1}{2}\right)\right]$
bgeva	$\frac{[-\ln(F(r))]^{-\tau}-1}{\tau}$

of Attack, Attack Technique, Continent, and Closeness (Section 2.2) as explanatory variables. In the following, S_i denotes the random variable *Severity* for the i -th cyber attack.

The cumulative link model in 3 becomes

$$\begin{aligned} \text{link}[P(S_i \leq s)] &= \alpha_s - \lambda y_i - \mathbf{z}_{iT} \boldsymbol{\gamma}_T - \mathbf{z}_{iA} \boldsymbol{\gamma}_A - \mathbf{z}_{iC} \boldsymbol{\gamma}_C \\ s &= 1, 2 \quad i = 1, \dots, \{, \dots, n, \end{aligned} \quad (4)$$

with

- y_i the observed *Closeness* value for unit i ,
- $\mathbf{z}_{iT} = [z_{iT}^{(T-Esp)}, z_{iT}^{(T-Hac)}, z_{iT}^{(T-Inf)}]$ i -th factor vector for *Type of Attack*,
- $\mathbf{z}_{iA} = [z_{iA}^{(A-0dy)}, z_{iA}^{(A-Mul)}, z_{iA}^{(A-Tri)}, z_{iA}^{(A-Unk)}]$ i -th factor vector for *Attack Technique*,
- $\mathbf{z}_{iC} = [z_{iC}^{(C-Afr)}, z_{iC}^{(C-Ame)}, z_{iC}^{(C-Asi)}, z_{iC}^{(C-Aus)}, z_{iC}^{(C-Eur)}]$ i -th factor vector for *Continent*,
- λ parameter for *Closeness*,
- $\boldsymbol{\gamma}_T = [\gamma_T^{(T-Esp)}, \gamma_T^{(T-Hac)}, \gamma_T^{(T-Inf)}]$ parameter vector for *Type of Attack*,
- $\boldsymbol{\gamma}_A = [\gamma_A^{(A-0dy)}, \gamma_A^{(A-Mul)}, \gamma_A^{(A-Tri)}, \gamma_A^{(A-Unk)}]$ parameter vector for *Attack Technique*,
- $\boldsymbol{\gamma}_C = [\gamma_C^{(C-Afr)}, \gamma_C^{(C-Ame)}, \gamma_C^{(C-Asi)}, \gamma_C^{(C-Aus)}, \gamma_C^{(C-Eur)}]$ parameter vector for *Continent*.

For example, if for unit i , we observe Hacktivism (T-Hac) as a *Type of Attack*, 0-day (A-0dy) as *Attack Technique*, and America (C-Ame) as *Continent*, Equation (4) becomes

$$\begin{aligned} \text{link}[P(S_i \leq s)] &= \alpha_s - \lambda y_i - z_{iT}^{(T-Hac)} \gamma_T^{(T-Hac)} - \\ & z_{iA}^{(A-0dy)} \gamma_A^{(A-0dy)} - z_{iC}^{(C-Ame)} \gamma_C^{(C-Ame)} \quad s = 1, 2. \end{aligned} \quad (5)$$

Note that if all elements of a factor vector are equal to zero, it means that we are considering the corresponding variable at the baseline level. As an exemplification, let us consider the qualitative variable *Type of Attack*. As described in Table 1, it has five different levels: Cybercrime (T-Cyb), Espionage/Sabotage (T-Esp), Hacktivism (T-Hac), and Information Warfare (T-Inf), with Cybercrime (T-Cyb) as baseline level. If $\mathbf{z}_{iT} = [z_{iT}^{(T-Esp)} = 0, z_{iT}^{(T-Hac)} = 0, z_{iT}^{(T-Inf)} = 0]$, it means that in the model we are considering Cybercrime(the baseline level) as type of attack.

We now discuss the main results obtained fitting the model to our data and using the link functions in Table 3. As shown in Table 4, *logit*, *probit*, and *log-log* cumulative link models outperform the others, presenting similar values of residual deviance (Res. dev.) and Akaike information criteria (AIC). They also provide coherent parameter estimations. Hence, in the following, we only present the results for the logit model, see Table 5. The baseline levels for variables *Type of Attack*, *Attack Technique*, and *Continent* are Cybercrime, SQL Injection, and Multiple continents, respectively; similar ML parameter estimates occur for different choices of the baseline level.

We notice that the p -values corresponding to Hacktivism, 0-day and Africa are high. Hence, these categories have a negligible impact on the severity level. The quantitative variable *Closeness* is significant, it affects positively the severity level and may improve the predictive performance of the model.

We remind that variable *Closeness* (defined in Section 2.2.2) is based on the values of the Criticality index that measures the vulnerability of each victim to suffer from a cyber attack. *Closeness* is an indicator of the connection of the victims in the network in terms of vulnerability, and it is a measure of the systemic risk and its possible effect on the probability distribution of *Severity*. In terms of real-world implications, the information that a cyber attack hits a victim strongly connected in the network could generate an alarm signal for the other connected victims. Victims close in terms of vulnerability should work together to prevent critical attacks and hopefully collaborate for the development of common security protocols to protect against cyber attacks.

In order to better explain the role of variable *Closeness*, in Figure 4, we depict the changes in the probability distribution for the extreme severity levels when *Closeness* varies from its minimum to its maximum value. In each graph of the figure, we consider one single explanatory variable at the time, and we fix the others at the baseline level. We notice a similar behavior for each explanatory variable: The probability decreases for the medium severity and increases for the critical severity.

For a better interpretation of the effect of each explanatory variable, we now present the AME measures for those explanatory variables significant at level 0.05 in Table 5. In a cumulative link model that contains solely main effects (as the one examined here), only the highest and lowest probabilities change monotonically as an explanatory variable increases. Therefore, in the following, we present the ME uniquely for extreme categories of the response variable. Furthermore, extreme categories (in our context medium vs. critical severity) often represent noteworthy states and are of special interest.

The AMEs for extreme categories of the response variable (medium vs. critical severity) are reported in Table 6, together with the standard errors, the test statistics, and the corresponding p -values. The standard errors are obtained via the Delta-method approach. We recall that for qualitative variables with more than two possible values, the ME

TABLE 4 Accuracy measures for the model with alternative link functions.

Link	logit	probit	log-log	cloglog	cauchy	bgeva ($\tau = -0.25$)
Res. dev.	5538.333	5541.049	5540.789	5565.718	5577.085	5733.813
AIC	5562.333	5565.049	5564.789	5589.718	5601.085	5757.813

Abbreviations: Res. dev., residual deviance; AIC, Akaike information criteria.

FIGURE 4 Extreme category probability changes as a function of the range of *Closeness*.

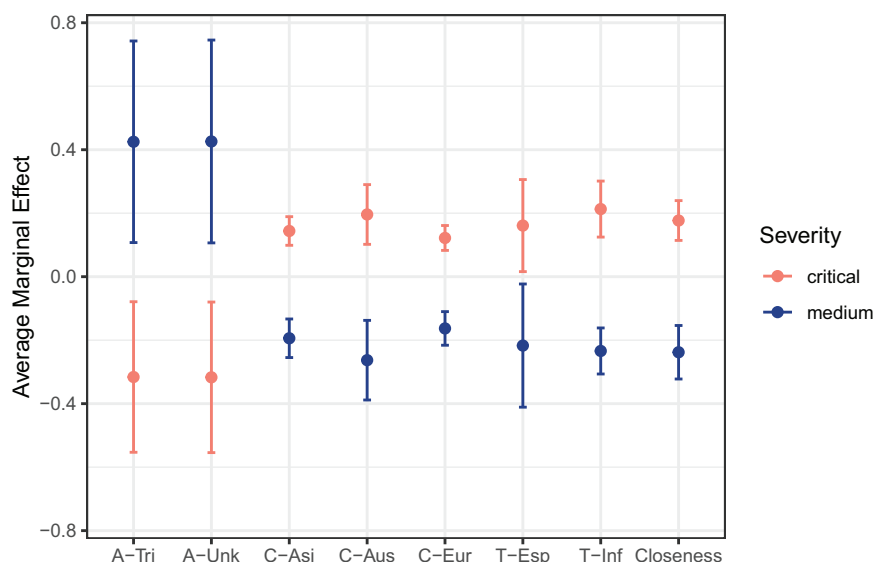


FIGURE 5 95% confidence intervals for the average marginal effects (AMEs) significant at level 0.05.

TABLE 5 The cumulative logit model fitted to the cyber risk data.

Coefficients:	Value	Std. error	z Value	p Value
T-Espionage/Sabotage	1.879	0.140	13.431	0.000
T-Hacktivism	0.252	0.162	1.558	0.119
T-Information Warfare	1.083	0.199	5.441	0.000
A-0-day	-0.568	0.854	-0.665	0.506
A-Multiple Threats	-1.470	0.761	-1.932	0.053
A-Trivial Threats	-1.933	0.740	-2.613	0.009
A-Unknown	-1.935	0.743	-2.606	0.009
C-Africa	0.367	0.377	0.974	0.330
C-America	0.756	0.097	7.749	0.000
C-Asia	0.883	0.143	6.161	0.000
C-Australia/Oceania	1.197	0.292	4.093	0.000
C-Europe	0.743	0.125	5.928	0.000
Closeness	0.985	0.452	2.179	0.029

measures show the difference in the estimated probabilities for cases in one category relative to the baseline one. In Figure 5, we plot the 95% confidence intervals only for the significant effects reported in Table 6.

We now briefly comment on the results presented in Table 6 and Figure 5. First, we consider the variable *Type of Attack*; Espionage/Sabotage and Information Warfare are types of attacks that are more likely than Cybercrime (the baseline level) to generate a critical severity. More precisely, for Espionage/Sabotage (Information Warfare), the estimated probability to generate a critical severity is on average 0.213 (0.177) higher than for Cybercrime. For medium severity, we observe an inverted effect. In accordance with the opinion of the Hackmanac experts, although Cybercrime attacks represent a huge percentage of the total number of cyber attacks (77.98%), in terms of gravity they are nowadays classified

as minor risks. Espionage/Sabotage and Information Warfare cause more severe damage because they are used mainly to steal important geopolitical and economical information. Furthermore, the available countermeasures are currently particularly ineffective; for more details, see Antonielli et al. (2019, pp. 13–14).

Next, we consider the variable *Attack Technique*. We notice that all AME measures for critical severity level are negatives. Trivial Threats and Unknown attack technique have a negative impact with respect to the baseline level (SQL Injection). Hence, even if SQL Injection is a low-frequency attack technique, its impact on critical severity level is quite relevant. This can find an explanation in the peculiarity of this type of attack. The SQL injection attack consists of reading and modifying sensitive data, performing unauthorized operations as an administrator on a database, retrieving the contents of a given system, and in some cases commanding the operating system. All these issues can cause serious problems to the victim. On the contrary, Trivial Treats and Unknown attack technique affect positively medium severity level. This means that attackers can rely on the effectiveness of more trivial and unknown attacks to achieve the majority of attacks (together 92.5% of the total number) even if of medium severity.

We now consider the variable *Continent*, for which the baseline level is Multiple Continents. Critical severity attacks directed at individual continents show a positive effect (the ME measures are positive), while attacks of medium severity are effective against multiple continents (the ME measures are negative). It turns out that attacks targeted against an individual continent are more effective than the ones directed to more continents in which the effect is dispersed. A possible explanation can be found in the peculiarities of the different continents. Note that, both the AME measures of America are not significant. A possible explanation could be that when an attack to Multiple Continents

TABLE 6 Average marginal effects (AMEs) for the cumulative logit model fitted to the cyber risk data.ocAME_CAT(*logit.m*) # new function available from the authors

\$ME.1 (medium severity)

	Effect	Std. error	z Value	p Value
T-Espionage/Sabotage	−0.234	0.037	−6.420	0.000
T-Information Warfare	−0.238	0.043	−5.521	0.000
A-Trivial Threats	0.425	0.162	2.619	0.009
A-Unknown	0.426	0.163	2.612	0.009
C-America	−0.083	0.086	−0.968	0.333
C-Asia	−0.194	0.031	−6.293	0.000
C-Australia/Oceania	−0.263	0.064	−4.136	0.000
C-Europe	−0.163	0.027	−6.055	0.000
Closeness	−0.217	0.099	−2.186	0.029
\$ME.3 (critical severity)				
T-Espionage/Sabotage	0.213	0.045	4.696	0.000
T-Information Warfare	0.177	0.032	5.493	0.000
A-Trivial Threats	−0.316	0.121	−2.621	0.009
A-Unknown	−0.317	0.121	−2.614	0.009
C-America	0.050	0.056	0.895	0.371
C-Asia	0.144	0.023	6.186	0.000
C-Australia/Oceania	0.196	0.048	4.097	0.000
C-Europe	0.122	0.020	5.936	0.000
Closeness	0.161	0.074	2.180	0.029

TABLE 7 Measures of correct classification obtained by *PROC LOGISTIC* in SAS.

Percent concordant	65.2	<i>d</i>	0.333
Percent discordant	31.9	<i>CI</i>	0.667
Percent tied	2.9	γ	0.342
Pairs	2337104	τ	0.217

(the baseline level) is performed, America is often affected.

Finally, a unit change in the quantitative variable *Closeness* would produce an average rate of change in the estimated probability equal to −0.217 for medium severity attacks, while it is equal to 0.161 for the critical ones. This confirms our previous findings regarding this variable.

4.1 | Predictive power of the model

In order to assess the predictive power of the model, we calculate the mostly used measures of correct classification reported in Table 7. They are based on the number of concordant (n_C) and discordant (n_D) pairs between the observed and the predicted values. A pair of observations is said to be concordant (discordant) if the observation with the larger (lower) observed response value also has the higher (lower) predicted response value; Harrel et al. (1996). If a pair of observations

with different responses is neither concordant or discordant, it is a tie. In Table 7 are also reported the percentage of concordant, discordant, and tied.

Somers' *d* statistics is the difference between the proportions of concordant and discordant pairs over the untied pairs $d = (n_C - n_D)/t$, where t is the number of pairs with different responses. It assumes values in $[-1, 1]$. A normalized version of this index in $[0, 1]$ is the concordance index (*CI*), that estimates the probability that the predictions and the outcomes are concordant $CI = (d + 1)/2$. *CI* is equal to 0.667. Therefore, we notice that for roughly the 67% of the untied pairs on the severity level, the observations with the higher severity are also associated to higher predicted values. This highlights a high concordance between observed and predicted values. Also, Goodman–Kruskal's γ and Kendall's τ indicate the presence of a positive association.

In order to evaluate the goodness of fit of the model, we considered the Efficient Score test and Wald test. Both indicate that the model gives a significant improvement over the baseline intercept-only model (p -values < 0.0001).

As suggested in Agresti (2010, pp. 69–70), since our model contains a continuous predictor, we compute the Lipsitz global goodness-of-fit test (Lipsitz et al., 1996). It is an alternative goodness-of-fit test for ordinal response logistic regression models that generalizes the Hosmer–Lemeshow test for binary logistic regression (Fagerland & Hosmer, 2016). It involves partitioning the data into g groups based on the cardinality $c(S)$ of the response variable. Lipsitz et al.

(1996) propose that g is chosen such that $6 \leq g < n/5c(S)$, where n is the sample size. This method is available in the R package `generalhoslem`. In our context, the Lipsitz test shows that there is no evidence of lack of fit for low value of g ($g = 6$, p -value = 0.1001).

4.2 | Discussion on the PO assumption

The score test by Peterson & Harrell (1990) is commonly used to check the PO assumption. This test evaluates if the effects are the same for each cumulative logit against the alternative hypothesis of different effects (see Agresti, 2010, section 3.5.5). It compares the null hypothesis of a model with PO (i.e., the coefficients do not depend on the level of the response variable), to the alternative one, a model with a separate set of coefficients for each threshold.

However, when the number of explanatory variables is large or the sample size is large, this test tends to reject the null hypothesis even when the PO assumption is reasonable; see Allison (1999), Btamt (1990), and O'Connell (2006). Furthermore, it tends to reject the null hypothesis also when few observations fall in one of the outcome categories (Peterson & Harrell, 1990) or when some explanatory variables are continuous (Allison, 1999). For these reasons, following O'Connell (2006), we computed the score test for each qualitative explanatory variable separately. We rejected the null hypothesis of PO assumption only for the categories Unknown and Espionage/Sabotage.

Finally, we analyzed how the odds ratios (OR) computed for these variables vary at the different thresholds. In particular, we dichotomized the three levels of the variable *Severity* and we computed three separate binary logistic regression models with all the explanatory variables. For both categories, we then compared the estimated ORs across all the severity levels with the cumulative OR. Furthermore, we compared the estimated values of the coefficients. Espionage/Sabotage presents broadly similar ORs across all the severity levels, with an average close to the cumulative OR. Moreover, the coefficients are broadly consistent in magnitude across all the severity levels. Hence, for Espionage/Sabotage, it seems reasonable to apply a PO cumulative logit model. Only Unknown attack technique presents a different picture, with ORs and coefficients varying across the severity levels.

The above discussion suggests that the ordinal PO model is a fair summary of the patterns in the data in relation to the severity levels. Naturally, it is possible to consider a more complicated model by relaxing the PO assumption; see Agresti (2010) and Peterson & Harrell (1990).

5 | CONCLUSIONS

In this article, we presented how cumulative link models can be a useful instrument for cyber risk assessment. These types

of models require only ordinal data for the response variable, that in our context describes the severity levels of a cyber attack, and not the actual losses they produce. This can protect the privacy of the cyber victims, can induce a wider disclosure of cyber risk data, and makes the model easily repeatable. We have applied our model to a real data set that includes information on “serious” cyber attacks that occurred worldwide from 2017 to 2018. We have also included in our model a quantitative variable named *Closeness* as explanatory variable, which could provide important insights about the existence of significant relations, in terms of vulnerability, among victims of attacks.

Moreover, we have considered different link functions such as *logit* and *bgeva*. Since the parameters of ordinal response models are not as simple to interpret as slopes and correlations for ordinary linear regression, we introduce alternative measures such as ME for evaluating the effect of each explanatory variable on the cyber risk level.

The empirical analysis suggests that Espionage/Sabotage and Information Warfare are the type of attacks that are more likely than the others to generate a critical severity attack. Among the most commonly used attack techniques, SQL Injection shows a quite relevant impact on critical severity level of attacks. Also, the quantitative variable *Closeness* has a significant impact on the attack severity level. This indicates that victims close in terms of vulnerability should collaborate for the development of common security protocols to protect against cyber attacks.

The proposed model can be used by practitioners and regulators for setting properly cyber security policy taking into account the specific characteristic of the observed attacks. More specifically, they could study the intersection of the threat landscape with the attack surface of a specific organization, in order to manage appropriate countermeasures and prioritize interventions on the basis of the estimated risk, considering also the available resources, the size of the organization, and its culture on cyber security issue.

In conclusion, we remark that there is no internationally recognized standard classification of the gravity of a cyber attack. In our application, the classification of the severity levels has been provided by cyber security experts. A different classification could lead to different results. This highlights the necessity to introduce a standardized classification of cyber risk levels that can be adopted worldwide.

ACKNOWLEDGMENTS

We thank the Associate Editor and the two referees for helpful comments and suggestions. This work acknowledges research support by COST Action CA19130 “Fintech and Artificial Intelligence in Finance - Towards a transparent financial industry” (FinAI), supported by COST (European Cooperation in Science and Technology) and the PRIN project “Fin4Green—Finance for a Sustainable, Green and Resilient Society Quantitative approaches for a robust assessment and management of risks related to sustainable investing”. We also thank the experts of the Hackmanac Project and Clusit for sharing the data set.

ORCID

Silvia Angela Osmetti  <https://orcid.org/0000-0002-0321-5731>

REFERENCES

- Afful-Dadzie, A., & Allen, T. T. (2017). Data-driven cyber-vulnerability maintenance policies. *Journal of Quality Technology*, 46, 234–250.
- Agresti, A. (2010). *Analysis of ordinal categorical data* (2nd ed.). J. Wiley & Sons.
- Agresti, A. (2013). *Categorical data analysis* (3rd ed.). Wiley.
- Agresti, A., & Tarantola, C. (2018). Simple ways to interpret effects in modeling ordinal categorical data. *Statistica Neerlandica*, 72, 210–223.
- Allison, P. D. (1999). *Logistic regression using the SAS system: Theory and application*. SAS Institute.
- Allodi, L., & Massacci, F. (2017). Security events and vulnerability data for cybersecurity risk estimation. *Risk Analysis*, 37, 1606–1627.
- Andreeva, G., Calabrese, R., & Osmetti, S. A. (2015). A comparative analysis of the UK and Italian small businesses using generalised extreme value models. *European Journal of Operational Research*, 249, 506–516.
- Antonielli, A., Arsene, L., Barletta, V. S., Butti, G., Bechelli, L., Benedetti, D., Biggio, B., Bologna, P., Carboni, D., Carboni, D., Ciardi, N., Conte, G., D'Agostino, R., Digregorio, P., Dinardo, L., Dozio, L., Dragoni, G., Faggioli, G., Fontana, R., ... Zapparoli Manzoni, A. (2019). *Rapporto Clusit 2019 sulla Sicurezza ICT in Italia*. Clusit.
- Bouveret, A. (2018). *Cyber risk for the financial sector: A framework for quantitative assessment*. IMF Working Paper WP/18/143, 1–27.
- Brant, R. (1990). Assessing proportionality in the proportional odds model for ordinal logistic regression. *Biometrics*, 46, 1171–1178.
- Calabrese, R., & Osmetti, S. A. (2013). Modelling SME loan defaults as rare events: The generalized extreme value regression model. *Journal of Applied Statistics*, 40, 1172–1188.
- Calabrese, R., Marra, G., & Osmetti, S. A. (2016). Bankruptcy prediction of small and medium enterprises using a flexible binary generalized extreme value model. *Journal of the Operational Research Society*, 67, 604–615.
- Cavusoglu, H., Mishra, B., & Raghunathan, S. (2004). The effect of Internet security breach announcements on market value: Capital market reaction for breached firms and Internet security developers. *International Journal of Electronic Commerce*, 9, 69–105.
- Cebula, J. J., & Young, L. R. (2010). *A taxonomy of operational cyber security risks*. Technical Note CMU/SEI-2010-TN-028, Software Engineering Institute, Carnegie Mellon University, 1–34.
- Cherdantseva, Y., Burnap, P., Blyth, A., Eden, P., Jones, K., Soulsby, H., & Stoddart, K. (2016). A review of cyber security risk assessment methods for SCADA systems. *Computers & Security*, 56, 1–27.
- Choi, T. M., & Lambert, J. H. (2017). Advances in risk analysis with big data. *Risk Analysis*, 37, 1435–1442.
- Davis, G., Garcia, A., & Zhang, W. (2009). Empirical analysis of the effects of cyber security incidents. *Risk Analysis*, 29, 1304–1316.
- Edgar, T. W., & Manz, D. O. (2017). *Research methods for cyber security*. Elsevier.
- Epskamp, S., Costantini, G., Haslbeck, J., Isvoranu, A., Cramer, A. O. J., Waldorp, L. J., Schmittmann, V. D., & Borsboom, D. (2021). *qgraph: Graph plotting methods, psychometric data visualization and graphical model estimation*. CRAN Repository. R package version 1.6.9.
- Facchinetti, S., Giudici, P., & Osmetti, S. A. (2020). Cyber risk measurement with ordinal data. *Statistical Methods and Application*, 29, 173–185. <https://doi.org/10.1007/s10260-019-00470-0>
- Facchinetti, S., & Osmetti, S. A. (2018). A risk index for ordinal variables and its statistical properties: A priority of intervention indicator in quality control framework. *Quality and Reliability Engineering International*, 34, 265–275.
- Fagerland, M. W., & Hosmer, D. W. (2016). Tests for goodness of fit in ordinal logistic regression models. *Journal of Statistical Computation and Simulation*, 86, 3398–3418.
- Giudici, P. (2018). Financial data science. *Statistics and Probability Letters*, 136, 160–164.
- Greene, W. (2008). *Econometric analysis* (6th ed.). Pearson Prentice Hall.
- Harrell, F. E., Lee, K. L., & Mark, D. B. (1996). Multivariable prognostic models: issues in developing models, evaluating assumptions and adequacy, and measuring and reducing errors. *Statistics in Medicine*, 15, 361–387.
- Hartwig, R. P., & Wilkinson, C. (2014). *Cyber risks: The growing threat*. Insurance Information Institute, p. 1–27.
- Kolfal, B., Patterson, R. A., & Yeo, M. L. (2013). Market impact on IT security spending. *Decision Sciences*, 44, 517–556.
- Kopp, E., Kaffenberger, L., & Wilson, C. (2017). *Cyber risk, market failures, and financial stability*. IMF Working Paper WP/17/185, 1–35.
- Lipsitz, S. R., Fitzmaurice, G. M., & Molenberghs, G. (1996). Goodness-of-fit tests for ordinal response regression models. *Journal of the Royal Statistical Society (Series C)*, 45, 175–190.
- Long, J. S., & Freese, J. (2014). *Regression models for categorical dependent variables using Stata*. Stata Press.
- Long, J. S., & Mustillo, A. S. (2018). Using predictions and marginal effects to compare groups in regression models for binary outcomes. *Sociological Methods & Research*, 50(3), 1284–1320. <https://doi.org/10.1177/0049124118799374>
- McCullagh, P. (1980). Regression models for ordinal data (with discussion). *Journal of the Royal Statistical Society, Series B*, 42, 109–142.
- Mood, C. (2010). Logistic regression: Why we cannot do what we think we can do, and what we can do about it. *European Sociological Review*, 26, 67–82.
- Moore, R. (2005). *Cybercrime: Investigating high-technology computer crime*. Anderson Publishing.
- O'Connell, A. (2006). *Logistic regression models for ordinal response variables*. Sage Publications.
- Opsahl, T., Agneessens, F., & Skvoretz, J. (2010). Node centrality in weighted networks: Generalizing degree and shortest paths. *Social Networks*, 32, 245–251.
- Peterson, B., & Harrell, F. E. Jr. (1990). Partial proportional odds models for ordinal response variables. *Journal of the Royal Statistical Society (Series C)*, 39, 205–217.
- Ramirez, R., & Choucri, N. (2016). Improving interdisciplinary communication with standardized cyber security terminology: A literature review. *IEEE Access*, 4, 2216–2243.
- Romanosky, S. (2016). Examining the costs and causes of cyber incidents. *Journal of Cybersecurity*, 2, 121–135.
- Sun, C. (2015). *Empirical research in economics: Growing up with R*. Pine Square LLC.
- The Global Risks Report 2019, 14th Edition (2019), World Economic Forum, Geneva, https://www3.weforum.org/docs/WEF_Global_Risks_Report_2019.pdf
- Wasserman, S., & Faust, K. (1996). *Social network analysis: Methods and applications*. Cambridge University Press.
- Yayla, A. A., & Hu, Q. (2011). The impact of information security events on the stock value of firms: The effect of contingency factors. *Journal of Information Technology*, 26, 60–77.

How to cite this article: Facchinetti, S., Osmetti, S. A., & Tarantola, C. (2024). A statistical approach for assessing cyber risk via ordered response models. *Risk Analysis*, 44, 425–438. <https://doi.org/10.1111/risa.14186>

APPENDIX A

In Tables A1 and A2, a detailed description of the explanatory variables *Attack Technique* and *Type of Attack* is reported.

For sake of parsimony and to facilitate the interpretation of the model, we decided to reduce the 10 levels of variable *Attack Technique* defined in the data set grouping together Account Cracking, DDoS, Malware, Phishing/Social Engineering, Phone Hacking, and Vulnerabilities in a new category named Trivial Threats. This encompasses less complex and sophisticated attack techniques: attackers can rely on the effectiveness of “simple” malware

and on more trivial and easy attack techniques to achieve their goals. The fact that the Trivial Threats category represents about 67% of the total implies that attackers can make successful attacks with relative simplicity and at very low costs. Therefore, the variable *Attack Technique* considered in our model presents five levels: A-SQL = SQL Injection, A-Ody = 0-day, A-Mul = Multiple Threats/APT, A-Tri = Trivial Threats, and A-Unk = Unknown.

Note that the model is robust to the choice of the aggregation of the more simple attack techniques in Trivial Threats.

TABLE A1 *Attack technique description.*

0-day	Flaw in software, hardware or firmware that is unknown to the party responsible for patching. 0-day refers to the fact that the developers have zero days to fix the problem that has just been exposed.
Account Cracking	Identification of valid login credentials by trying different values for user names and/or passwords.
DDoS	Multiple compromised computer systems attack a target and cause a denial of service for users of the targeted resource.
Malware	Software which is specifically designed to disrupt, damage, or gain authorized access to a computer system.
Multiple Threats/APT	Multiple Threats refers: attacks based on more the one threats. APT (Advanced Persistent Threat): prolonged and targeted cyber attack in which an intruder gains access to a network and remains undetected for a period of time.
Phishing/Social Engineering	Phishing: fraudulent attempt to obtain sensitive information such as usernames, passwords and credit card details. Social Engineering: use of psychological manipulation to trick users into making security mistakes or giving away sensitive information.
Phone Hacking	Get access to an individual's cellular phone through a variety of methods.
SQL Injection	Code injection technique, used to attack data-driven applications, in which nefarious SQL statements are inserted into an entry field for execution.
Unknown	Most of the used techniques are Data Breach kind, for which the consequences may be known, but almost never the execution methods.
Vulnerabilities	Weakness in the computational logic found in software and hardware components that, when exploited, results in a negative impact to the CIA.

TABLE A2 *Type of attack description.*

Cybercrime	Use of a computer or internet as an instrument to further illegal ends, such as committing fraud, stealing identities, or violating privacy.
Espionage/Sabotage	Espionage: act or process of learning secret information through clandestine means. Sabotage: deliberate and malicious actions aimed at weakening an enemy through disruption the of the normal processes and functions.
Hacktivism	Act of hacking or breaking into a computer system for a politically or socially motivated purpose.
Information Warfare	Use of information technology as an active weapon of war.