

The County Fair Cyber Loss Distribution: Drawing Inferences from Insurance Prices

Daniel W. Woods^a, Tyler Moore^b, Andrew C. Simpson^a

^aDepartment of Computer Science, University of Oxford

^bTandy School of Computer Science, The University of Tulsa

Abstract

The actuarially fair insurance premium reflects the expected loss for each insured. Given the dearth of cyber security loss data, market premiums could shed light on the true magnitude of cyber losses despite noise from factors unrelated to losses. To that end, we extract cyber insurance pricing information from the regulatory filings of 26 insurers. We provide empirical observations on how premiums vary by coverage type, amount, policyholder type, and over time. A method using Particle Swarm Optimisation is introduced to iterate through candidate parameterised distributions with the goal of reducing error in predicting observed prices. We then aggregate the inferred loss models across 6,828 observed prices from all 26 insurers to derive the *County Fair Cyber Loss Distribution*. We demonstrate its value in decision support by applying it to a theoretical retail firm with annual revenue of \$50M. The results suggest that the expected cyber liability loss is \$428K, and that the firm faces a 2.3% chance of experiencing a cyber liability loss between \$100K and \$10M each year. The method could help organisations better manage cyber risk, regardless of whether they purchase insurance.

1 Introduction

Understanding how losses are distributed is an important step in assigning information security resources. The optimal investment level is intimately linked to both the likelihood and the impact of potential attacks. Put simply, “risks cannot be managed better until they can be measured better” [1].

How can organisations gain insights into their distribution of cyber losses? Information sources include mandatory breach notification laws, threat reports released by security vendors, self-reported survey data, and court dockets relating to cyber incidents. Potential insights are limited by challenges including the problem of denominators [2], reporting biases [3] and the tension between sample size and granularity.

The incentives for security vendors to misrepresent the threats faced have been widely recognised [4]. Self-reported survey results regarding cyber-crime costs are “dominated by a minority in the upper tail” and the effect of “refusal rates and small sample sizes” [3] are

magnified by the rarity of reported cyber attacks. Court dockets are detailed but can only provide insights into costs assigned by courts.

Repositories of data breach reports, such as Privacy Rights Clearinghouse¹, are drawn from the population of firms operating in the corresponding jurisdiction. This allows researchers to glean insights into the size of breaches and frequency in the aggregate. Estimating the frequency of incidents for an individual firm is complicated by not knowing the size of the population from which these reports are drawn [2].

Using the aforementioned data sources assumes homogeneity in the sample. This raises the question of how relevant the losses suffered by a \$300 billion energy firm are to a ‘mom-and-pop’ retail outlet. But, throwing out such data reduces an already limited sample size. A novel data source for cyber loss distributions could be valuable to complement existing information sources, especially given their limitations.

Hayek identified the price system as a mechanism to communicate information about “how to secure the best use of resources” [5], even when the “knowledge of the relevant facts is dispersed among many people”. This link between prices and dispersed information underlies faith in the efficient markets hypothesis stating that asset prices reflect all available information. Although, it has recently come under attack [6].

Prices can be operationalised for decision support. For example, prediction markets have been created to predict election outcomes [7], influenza outbreaks [8], scientific results [9], Oscar nominees and winners [10], and government policy outcomes [11]. In the security context, it has been suggested that exploit derivatives can be used to estimate the probability a software product will be compromised [12] and that cyber warranties can be used to estimate expected cyber losses with the corresponding control in place [13]. However, it is important to remember “markets incorporate falsehood as well as truth” [10]. Distributions implied from stock market prices failed “to anticipate” past crashes [14].

This paper introduces a method to infer loss distributions from cyber insurance prices, much like how probabilities are extracted from prediction markets. If the efficient markets hypothesis holds for insurance markets, the price of insurance for a single retail firm with a revenue of \$50M reflects all available information about retail firms with a revenue of \$50M. Admittedly, factors like operational costs and solvency risk add noise to the inferences.

Our aim is to identify the best-fit parameterised distributions for the observed prices. The contributions include a quantitative description of cyber insurance pricing, a novel technique for inferring loss distributions from insurance premiums, and the derivation of the *County Fair Cyber Loss Distribution*.

This method complements existing approaches to quantifying cyber losses in a number of ways. Insurance prices are provided based on firm characteristics like revenue and industry. This allows our method to provide granular estimates based on such firm characteristics. Insurance is offered for many cyber incidents that currently lack related data. This method provides implied distributions of dollar losses, whereas previous studies report point-estimates of dollar losses or distributions of the number of records breached.

Section 2 identifies research into the quantification of cyber losses at a firm level. Section 3 describes the pricing structure from all admitted insurers in California. We introduce the theory behind our inferences in Section 4. Section 5 applies the inference framework

¹<https://www.privacyrights.org/>

Study	Year	Data Source	Data Points	Frequency	Impact
[15]	2008	Identity Theft Resource Center	899		
[16]	2010	Open Security Foundation DatalossDB	956		Power law*
[17]	2016	Privacy Rights Clearinghouse (PRC)	2253		Lognormal*
[18]	2016	PRC & OPS DatalossDB &	8574		Pareto*
[19]	2017	Privacy Rights Clearinghouse	2266		Log-skew-normal*
[20]	2018	Privacy Rights Clearinghouse	600		Stochastic process*
[21]	2010	Research institute’s event log	23,000		Lognormal*
[22]	2014	Nordic bank’s event log	1,800		Lognormal*
[23]	2018	Proprietary data	53,308/2,216		
[24]	2018	Interviews	2200		μ
[25]	2015	Information security breaches survey	664	μ	μ
[26]	2018	Bank of Italy survey data	4,209		μ
[27]	2014	US court dockets	230		
[28]	2017	UK ICO regulatory actions	118		μ
[29]	2015	Operational risk database	994		μ, σ
[30]	2003	Event (DoS) window study	23		AR
[31]	2003	Event (Data Breach) window study	43		AR
[32]	2004	Event (Data Breach) window study	66		AR
[33]	2016	Proprietary data	12,585	μ	MR

Table 1: Research into the cost and frequency of cyber losses. * = Not financial impact, μ = mean/mode/median, σ = variation, MR = Multivariate regression, AR = Abnormal returns.

to evaluate how well these parameterised distributions explain observed data. We discuss limitations and compare our inferences to other empirical work related to loss events in Section 6. Finally, we offer conclusions in Section 7.

2 Related Work

This section identifies existing approaches and techniques to quantify cyber losses. Table 1 provides an overview of the studies we have considered, including the data source of the study. The frequency and impact columns concern whether the study provides insights for individual organisations, rather than statistics aggregated across multiple firms within an industry or economy.

Studying data breach repositories like Privacy Rights Clearinghouse provides insights into the size of breaches and aggregate frequency. Aggregate frequency is found to be stable over time [17, 18] and is distributed according to a “Poisson or negative binomial” [19]. Breach size was shown to be best described by a power law [16], Lognormal [17], Pareto [18] and log-skew-normal [19] distribution in successive publications. A recent paper has suggested both aggregate frequency and breach size are better modelled by stochastic processes [20].

Realistically, governments are the only entities who can realistically operationalise insights into aggregate frequency. Unfortunately they have a limited number of levers to pull in response. A notable exception is mandatory breach reporting laws, which have been shown to reduce identity theft by 6.1% [34]. Organisations have more risk management tools available, but how can they estimate potential losses?

Schroeder et al. [21] investigated 23,000 failures recorded on more than 20 different systems at a research institute. The results suggest time between failures is modelled by a Weibull distribution and “repair times are well modeled by a lognormal distribution”. Franke et al. [22] investigated 1800 incidents in a large Nordic bank and also found that the “log-normal distribution offers the best fit of IT service time to recovery”. These studies do not provide incident costs. It is not clear how much we can generalise from atypical organisations like the Los Alamos National Laboratory [21].

Verizon’s *Data Breach Investigations Report* [23] describes the relative frequency of different types of incidents by industry. The Ponemon Institute’s *Cost of a Data Breach* [24] estimates the average cost per record in a data breach. However, [23] does not provide absolute frequencies and [24] only surveys firms who have detected a breach. There are questions about the compatibility of commercial sponsors and scientific integrity, as evidenced by the \$1 trillion cyber crime figure².

Surveys commissioned by governments provide an alternative. One study of security investments [25] used a survey commissioned by the UK Government [35] providing point-estimates of both frequency and impact distributed according to a Bernoulli distribution. Piggy-backing security-related questions on the Bank of Italy’s annual survey provided a larger set of responses [26]. Neither set of losses was fitted to a distribution. Furthermore, self-reported surveys are complicated by response biases [3].

Court dockets provide rich information about legal cases. Romanosky et al. [27] investigated factors affecting the data breach litigation. Freedom of information requests have been used to understand regulatory actions in the UK [28]. But they only provide insights into cases contested in the courts, leading to small sample sizes.

Biener et al. [29] extracted 994 cyber incidents from an operational risk database. The mean and mode loss are \$41M and \$1.9M respectively, and “the distribution of the non-cyber risk sample is much heavier tailed than that of the cyber risk”. The mean loss figure of \$41M is 10 times larger than a similar figure in [33]. The types of attack included in their dataset depends on the “comprehensive set of keywords” comprising their definition of a cyber incident, so it is not clear what is included.

Event window studies provide insights into how events, such as denial of service attacks [30] and data breaches [31, 32], impact the stock market. By extracting information from the stock prices, event window studies represent an intellectual forefather of our proposed technique, which extracts information from insurance prices. However, event window studies are limited to publicly reported events. As a result, they tend to have small sample sizes and questionable relevance beyond listed companies.

The most comprehensive study quantifying cyber losses is based on a proprietary dataset with 15,000 incidents [33]. Attack frequencies for each industry are calculated using census data on the number of firms in each industry. Their multivariate regression, based on 265 observations, describes how total cost of incident varies according to an organisation’s revenue, industry, number of records and past incidents.

Table 1 summarises a research programme with strong evidence regarding how the number of records in a data breach is distributed and how frequently they occur across the United

²<https://www.forbes.com/sites/andygreenberg/2012/08/03/mcafee-explains-the-dubious-math-behind-its-u-#294dfc4525a3>

States. We have estimates [33] of breach frequency for an organisation, although this figure groups all firms in an industry regardless of their size³, as well as point estimates of financial cost of incidents via survey data [25, 26] and an insightful multivariate regression [33].

However, quantifying the costs of attacks beyond data breach and regulatory fines is a challenge. In particular, the studies described in Table 1 provide no insights into the potential cost of business email compromise, business interruption incidents, or ransomware attacks. A possible exception is [29] because it is unclear what is included under operational loss.

The next section describes the insurance pricing data we will analyse. It provides insights into policies covering business email compromise, business interruption and ransomware attacks, which have not been empirically studied in the literature so far. Section 4 introduces a method to derive distributions of dollar losses in terms of both frequency and impact, advancing the state-of-the-art in quantifying cyber losses.

3 Observed Prices

This section identifies publicly available cyber insurance data, which will provide context for the framework for inferences introduced in Section 4. Data collection is described in Section 3.1. Section 3.2 provides a quantitative description of cyber insurance prices, building on the overview provided in [36].

3.1 Data Collection

Insurers in the United States operating in admitted markets are required to “file their policies and rate schedules with the state insurance departments” [36]. Rate schedules describe the formulas and tables used to calculate the premium for a given applicant. These documents are made publicly available by the National Association of Insurance Commissioners (NAIC).

Documents are organised by state and made available in the SERFF Filing Access system. We chose to focus on the state of California, which is the largest in the US, rather than partially search multiple states in order to collect the widest range of policies. Even if the set of policies skews towards firms unique to California, the insights from a given policy should have general applicability as there are no differences between states “that would materially bias any results or conclusions” [36].

Our objective was to quantify how cyber insurance premiums are adjusted according to coverage type, limit, deductible, industry, revenue, and security infrastructure. This will allow us to generate sets of insurance prices for a hypothetical firm with specific characteristics, such as revenue or industry.

We extracted rate schedules related to cyber insurance by searching with keywords “cyber”, “security” and “privacy”, as in [36]. We downloaded only documents with either a “new program” or “rate” component in the filing type because these relate to pricing, whereas some filings concern the policy wording and have a “form” or “rule” filing type.

³Grouping corporations like Target with independent book stores under “retail trade” might explain the particularly low frequency for retail firms.

Revenue	Base rate	Deductible	Factor	Limit	Factor	Hazard Class	Factor
\$10m or less	\$1,914	\$10,000	1	\$500,000	0.809	1	.804
\$10m–\$20m	\$2,603	\$25,000	.95	\$1,000,000	1	2	1
\$20m–\$50m	\$3,502	\$50,000	.89	\$2,000,000	1.132	3	1.497
\$50m–\$100m	\$5,225	\$100,000	0.82	\$3,000,000	1.245	4	1.905

Table 2: A subset of the baserates and multiplicative factors found in a rate schedule, which can be downloaded from <https://tylermoore.utulsa.edu/cyberschedex.pdf>

$$\text{Price} = (\text{Base rate}) \times (\text{Deductible Factor}) \times (\text{Limit Factor}) \times (\text{Hazard Class Factor})$$

This resulted in 131 unique filings that were narrowed down to 26 rate schedules appropriate for our purposes. Romanosky et al. [36] identified three types of pricing: flat rate pricing, base rate modification, and information security pricing. There were 40 filings related to the first type and they were excluded because the inflexible pricing structure makes inferences difficult. Further, being sold alongside existing products means these endorsements may not be commercially viable alone.

A further five filings relate to the price of excess layers in which insurers offer additional layers of coverage to supplement an existing policy. These tend to be priced as a percentage of the original policy, which makes inferences difficult. We also excluded filings for policies with pricing structures for specific industries. For example, one policy was priced according to the number of doctors employed by a healthcare provider.

Some filings introduced new coverage areas without updating the original rates. We considered these to be additions to the original filing. A small number of insurers updated prices for the same coverage and these were considered to be new filings. This resulted in 26 unique filings, which correspond to sets of prices, along with meta data including when the pricing scheme was first filed.

A rate schedule does not provide prices directly. We had to read the document explaining how prices are calculated and extract the corresponding tables of multiplicative factors. The prices from the 26 rate schedules are determined by the product of a base rate (in USD) and many multiplicative factors, which increase or decrease the price.

After extracting the tables, we can calculate the price-limit-deductible triples for a hypothetical firm. The hypothetical firm would have characteristics, such as revenue or industry, corresponding to each multiplicative factor. Using the factors in Table 2, one of the triples for a retail firm with revenue \$50M would be (4964, 1000000, 25000) because $4964 = 5225 \times 1 \times 0.95 \times 1$ where 1, 0.95 and 1 are the factors for the limit of 1M, deductible of 25K, and hazard class of 2 respectively.

Generating the data set consists of taking all combinations of limits and deductibles offered by an insurer then computing the corresponding premium. If a rate schedule provided a choice of 8 deductible amounts and 15 limits, then there would be $8 \times 15 = 120$ triples. This leads to a total data set of 6,828 price-limit-deductible triples across all 26 rate schedules, with price varying based on the hypothetical firm’s characteristics. Section 4 describes how we make inferences from these triples.

The next subsection illustrates how the multiplicative factors vary with a change in coverage type, limit, deductible, and revenue. It also provides insights into how prices have changed over time.

Coverage	n	μ	σ
Cyber liability	26	1	0
Data breach/first party costs	15	0.95	0.97
Regulatory proceedings	11	0.23	0.17
Website multimedia	14	0.44	0.44
Business interruption	15	0.36	0.34
Contingent business interruption	4	0.06	0.05
Ransomware	15	0.15	0.13
Wire transfer	11	0.12	0.07
Notification Costs	7	0.21	0.08
Crisis management	9	0.1	0.11
PCI costs	5	0.24	0.15
Forensics	3	0.18	0.15
Data recovery	11	0.17	0.22

Table 3: Available endorsements to cyber liability coverage. μ is calculated by taking the mean of the price of each coverage as a fraction of the price for cyber liability.

3.2 Quantitative Analysis

We look at how prices are adjusted for coverage type, limit and deductible, and revenue and industry in turn. We then provide a longitudinal perspective on cyber insurance prices.

Pricing by coverage type Table 3 lists the coverage categories identified, along with frequency of occurrence (**n**), the average price as a fraction of the cyber liability premium (μ) and standard deviation (σ). Prices for different coverages are typically expressed as a fraction of the cyber liability premium, often in a range (e.g., 0.05–0.15), in which case we used the mid-point.

The prominence of cyber liability coverage is not surprising given most filings fall under insurance lines “related [to] corporate liability policies” [36]. Market entrants began offering coverage including first party, business interruption and ransomware in later years.

Data breach and first party costs vary across insurers in terms of what is covered, often including some combination of the last five entries in Table 3, which explains the high standard deviation. Further, policyholders can combine notification costs, public relations and forensics to build the equivalent of a comprehensive data breach policy. This may explain the seemingly small amount of insurers offering first-party data breach coverage. Although only three and five insurers offer a forensics and PCI costs endorsement respectively, coverage may still be offered under a general “first-party” costs endorsement.

Wire transfer fraud coverage is consistently priced at around 5–15% of cyber liability. Ransomware coverage tends to be somewhat more expensive. Business interruption coverage is even more expensive still. Variance in the price of multimedia coverage could result from it being a type of third-party cover, which is traditionally difficult to price.

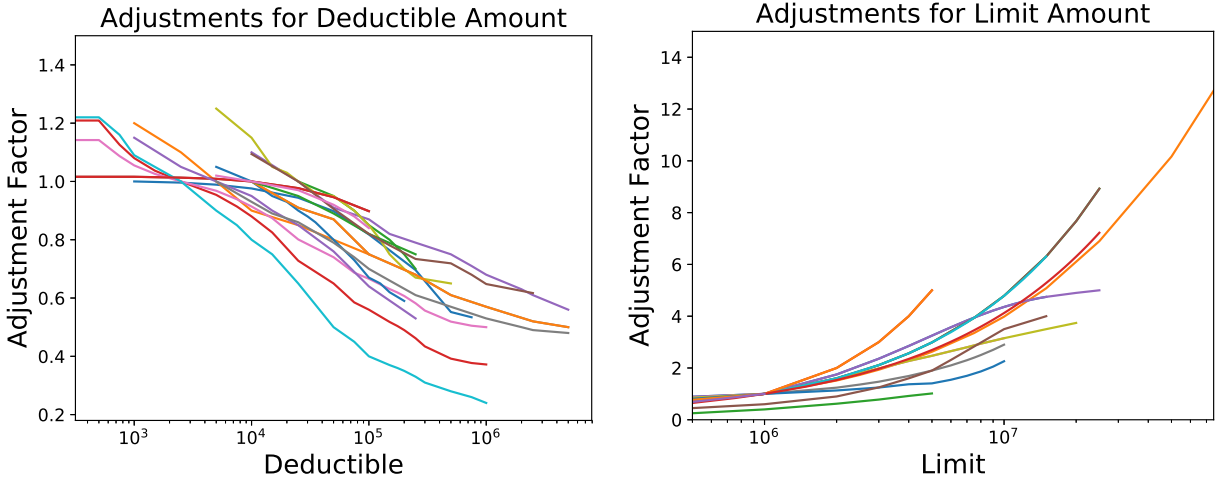


Figure 1: Insurers apply an adjustment factor depending on the selected deductible (left) and limit (right). Colours are randomly generated and correspond to a different pricing scheme.

Pricing by coverage amount Figure 1 (left) shows that increasing the deductible leads to a decrease in price as we would expect. The absolute adjustment depends on the specific insurer’s baseline, but we can see exponential increases in the deductible lead to linear decreases in the adjustment factor in general. Insurers use different baseline deductibles (corresponding to an adjustment factor of 1), while most insurers use \$1,000,000 as the baseline limit, hence why most of the lines in Figure 1 (right) pass through (1000000, 1.0).

The most obvious difference between insurers occurs when they adjust for higher limits. Remarkably, one insurer began including an option for a \$1 billion limit in 2017, which is 56 times the price of a \$1 million limit. Figure 1 (right) shows that some insurers offer sub-linear increases in price for an exponential increase in revenue, which suggests losses exceeding the limit become increasingly less likely.

Some insurers use the applicant’s revenue to set the adjustment for the deductible — when this was the case we used the adjustment for the smallest revenue band available for these figures. Other than these insurers, adjustments for the deductible and limit do not depend on the coverage type or firm characteristics. The importance of this will be established in Section 4 because it suggests the coverage type or industry has little effect on the shape of the distribution of losses.

Pricing by policyholder characteristics We cannot hope to describe in aggregate the range of adjustments for policyholder characteristics. Many of these adjustments are at the underwriter’s discretion, often ranging from 0.75 to 1.25. This is further complicated by insurers describing what they are adjusting for differently, which impedes cross-comparison. We would have to justify why adjustments for “privacy controls”, “encryption” and “encryption for network data, laptops and mobile devices” refer to the same firm underlying characteristic.

One might expect to be able to compare across adjustments for the policyholder’s in-

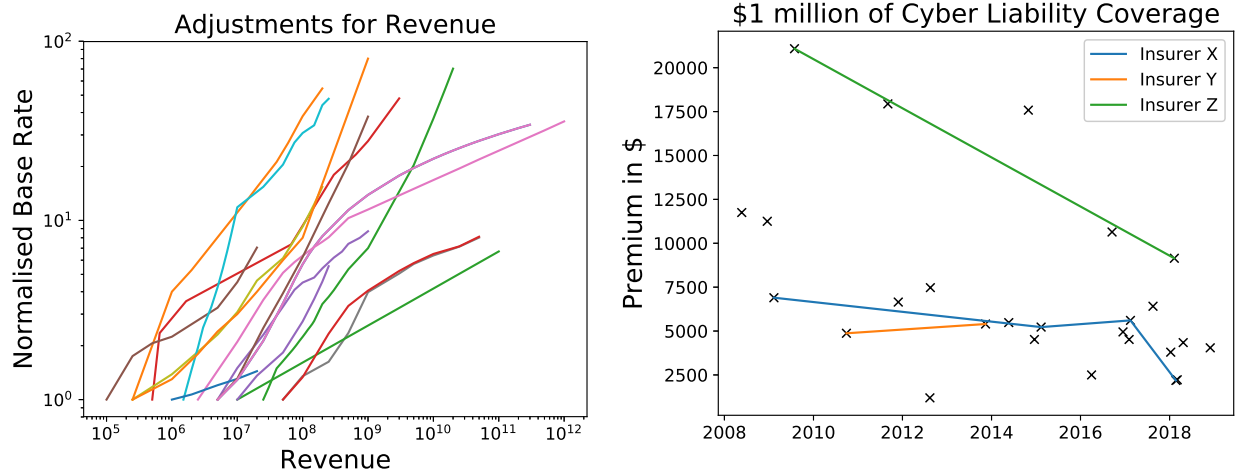


Figure 2: Insurers apply an adjustment factor depending on the policyholder’s revenue (left); cyber liability insurance premiums over time for selected insurers (right).

dustry, especially given there are popular standards for defining industry. As Romanosky et al. [36] observed, “there was no consistency regarding approach [to industry definitions], or any consensus on what the insurance industry would consider the most risky”. As a result, we cannot compare across these adjustments because there is no consistent definition.

Fortunately, we *can* examine how prices are adjusted for revenue. Larger companies are charged higher premiums by assigning tiered base rates according to revenue. Figure 2 (left) displays the adjustments that are made according to revenue. This also shows the maximum revenue that insurers are willing to price without further consultation. It should be noted that most rate schedules offer rates for larger organisations by request.

These adjustments cast further doubt over aggregate statistics regarding expected losses that group together companies regardless of revenues. Some insurers force large companies to pay 100 times the premium that smaller companies do, suggesting the expected loss is up to 100 times as large.

Pricing over time Base rates are typically given for \$1,000,000 of cyber liability coverage, with different rates depending on the organisation’s revenue band. Figure 2 (right) shows the base rate for a retailer with \$50 million of revenue, along with the date the policy was first filed with the regulator.

Three insurers updated their rate schedules with new prices, providing the only *apples to apples* comparison because the prices relate to the same underlying policy. Insurer X updated their prices four times with a notably sharp decrease in 2018. Meanwhile, Insurer Y had a minor increase in price. Insurer Z reduced their price by over 40% after 9 years. It is unclear whether prices for new and existing customers gradually fell during this period.

The trend in prices among market entrants is more difficult to assess due to subtle differences in the products being sold (although these differences are often overstated [36]). The cheapest policy was introduced in 2012 and the maximum limit was \$5M, which suggests it was targeted at smaller companies. We would need to control for changes in coverage to

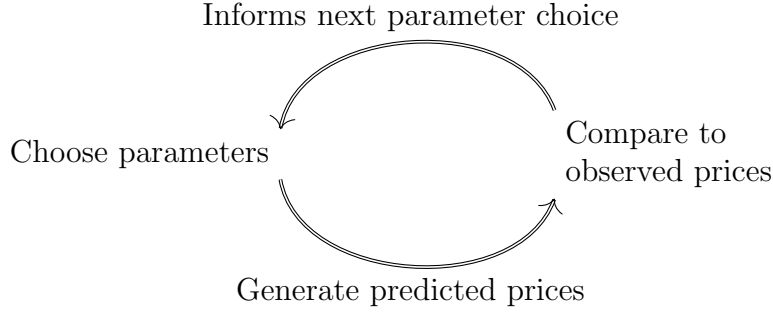


Figure 3: High level description of our iterative method for inferring loss distributions from insurance prices.

make reliable claims about how prices change over time⁴. Although, it is fair to say market entrants and price updates have become more frequent.

4 Method for Inferring Loss Distributions

We now consider how to make inferences from pricing data. Figure 3 describes our iterative method for inferring loss distributions. The method begins by choosing parameters for a given loss distribution. The parameterised loss distribution is used to generate the predicted price for an insurance policy with a given limit and deductible. The set of predicted prices is compared to the observed prices. The result of this comparison is used to improve the next parameter choice. The cycle repeats until some termination condition is met.

We describe how we generate predicted prices from a parameterised loss distribution in Section 4.1. A metric to compare predicted prices with observed prices is defined in Section 4.2. The heuristic governing parameter choices and the termination condition is described in Section 4.3. Finally, Section 4.4 explains how to translate between loss distributions for different coverage types and revenues.

4.1 Generating Price Predictions

This subsection assumes we already have a parameterised distribution for all loss events covered by a given insurance policy. This gives rise to the universe of possible losses $\Omega = \mathbb{R}_+$, which is distributed according to a random variable X . Each possible loss value x is a non-negative real number occurring with probability determined by the probability density function

$$f(x) : \mathbb{R}_+ \rightarrow \mathbb{R}_+$$

The probability that the loss amount is greater than a and less than b is given by

$$P(a < X < b) = \int_a^b f(x)dx \quad (1)$$

⁴Anecdotally, industry insiders suggest coverage has become broader. If this is true, then Figure 2 suggests prices are falling, which would be consistent with greater competition.

The first axiom of probability follows provided f is non-negative.
The second follows by normalising f so that

$$\begin{aligned} & P(\Omega) \\ &= \int_0^\infty f(x)dx \\ &= 1 \end{aligned} \tag{2}$$

The third follows from the definition of an integral.

For our purposes, an insurance contract is a promise that the insured will be indemnified for any losses greater than the deductible D up to a limit L . This gives rise to an indemnity function $I_{D,L}(x) : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ with

$$I_{D,L}(x) = \begin{cases} 0, & \text{for } x < D \\ x - D, & \text{for } D \leq x \leq L + D \\ L, & \text{for } L + D < x \end{cases} \tag{3}$$

which describes the size of the indemnity for a given loss x .

As a result, the insurer's expected loss is given by

$$\begin{aligned} E(X)_{D,L} &= \int_{-\infty}^\infty I_{D,L}(x)f(x)dx \\ &= \int_D^{L+D} (x - D)f(x)dx + \int_{L+D}^\infty Lf(x) \end{aligned} \tag{4}$$

Our method assumes insurance prices are determined by the value of $E(X)_{D,L}$. Algebraically we assume that a premium $p_{D,L}$ with limit L and deductible D is given by

$$p_{D,L} = \lambda E(X)_{D,L} \tag{5}$$

where $\lambda \geq 1$ is the loading factor. The loading factor determines how much the insurer sets aside for non-claims related expenses. An actuarially fair policy occurs if $\lambda = 1$.

4.2 Evaluating Loss Distributions via Predicted Prices

This subsection assumes the loss distribution has generated a set of predicted prices. We want to evaluate how the predicted prices compare to observed prices. We define the following cost functions.

Definition 1 (Observed Price Cost Functions) *For a set of premium-limit-deductible triples $S = \{(p_{D_1,L_1}, L_1, D_1), \dots, (p_{D_n,L_n}, L_n, D_n)\}$, denote the predicted price and observed price respectively as*

$$P_i = E(X)_{D_i,L_i}$$

and

$$O_i = p_{D_i,L_i}$$

Then define the following cost functions to evaluate a parameterised loss distribution X with an associated probability density function $f(x) : \mathbb{R} \rightarrow \mathbb{R}_+$ over the set of prices S :

$$C_T(f|S) = \frac{1}{n} \sum_{i=1}^n \frac{P_i - O_i}{O_i}$$

$$C_A(f|S) = \frac{1}{n} \sum_{i=1}^n \frac{|P_i - O_i|}{O_i}$$

The *absolute* function C_A sums the absolute differences between predicted and observed prices. A value of 0 occurs if and only if expected prices perfectly predict observed prices. Meanwhile, the *total* function C_T allows under-predictions to cancel out over-predictions in the aggregate. A score of zero can be achieved by balancing over- and under-predictions.

The parameterised distribution X with the lowest cost function value *best explains* the observed prices $(p_{D_1, L_1}, L_1, D_1), \dots, (p_{D_n, L_n}, L_n, D_n)$.

4.3 Choosing Parameters and Termination

Our method can be considered as an iterative optimisation problem. For a distribution with n parameters, this involves searching through $\mathbb{R}^n$ for the n -tuple of parameters minimising the cost function. The problem is non-convex with no way of distinguishing local minima from the globally optimal solution. Further, the objective function f is relatively costly because it consists of $2m$ integrals where m is the number of price-limit-deductible triples offered by the insurer. These integrals will be computed numerically when an analytic solution is not easily obtained (as will often be the case).

Many optimisation techniques could have been applied and may still prove fruitful in future work. The particle swarm optimisation (PSO) heuristic [37] consists of multiple candidate solutions moving around the search space. The direction and speed of movement of each candidate is determined by a weighting of each candidate's best known position and the best known position from the entire swarm. The search stops when the candidates converge or after a certain number of iterations. The process is repeated multiple times to decrease the chance of selecting an unsatisfactory solution.

PSO was chosen because it does not require calculating the gradient of the objective function, unlike gradient descent. It requires few assumptions about the search space, which is ideal for exploratory research. Further, multiple particles in the search space reduce the chance of terminating at a local minima. The Pyswarm [38] package was used for implementation.

4.4 Accounting for Multiplicative Pricing

This subsection provides a partial answer as to the effect of changing the coverage type, a characteristic about the firm, or even the profit loading factor λ . It could save re-running the optimisation analysis. We first provide the intuition behind the adjustment.

Suppose that we found that $f(x)$ best explains a set of prices $S = \{p_1, \dots, p_n\}$. Then let $aS = \{ap_1, \dots, ap_n\}$ be a set of prices for a new coverage type with $0 < a$, then intuitively

we might say that all losses are a times as likely. The following reasoning suggests $af(x)$ explains aS as well as $f(x)$ explains S .

$$C_A(af|aS) = \frac{1}{n} \sum_{i=1}^n \frac{|P_i - ap_i|}{ap_i} = \frac{1}{n} \sum_{i=1}^n \frac{|E(af)_{D_i, L_i} - ap_i|}{ap_i} \quad (6)$$

Then

$$\begin{aligned} E(af)_{D, L} &= \int_D^{L+D} (x - D)af(x)dx + \int_{L+D}^{\infty} Laf(x) \\ &= aE(f)_{D, L} \end{aligned} \quad (7)$$

substituted into (6) reveals

$$C_A(af|aS) = \frac{1}{n} \sum_{i=1}^n \frac{a|E(f)_{D_i, L_i} - p_i|}{ap_i} = C_A(f|S) \quad (8)$$

For example, if the price of business interruption is 36% of cyber liability across all limits and deductibles, then denoting the cyber liability prices as S_{CL} we have

$$C_A(0.36f(x)|0.36S_{CL}) = C_A(f(x)|S_{CL})$$

where $C_A(f|S_{CL})$ is the cost function value of $f(x)$, the cyber liability loss distribution.

Unfortunately, $af(x)$ is not a probability density function since

$$\int_0^{\infty} af(x)dx = a \int_0^{\infty} f(x)dx = a \neq 1$$

Providing $a < 1$, we can adjust $af(x)$ so that a loss of 0 occurs with probability $1 - a$. Constructing “mixed discrete/continuous distributions” has precedent in actuarial science [39]. This translates into a business interruption loss of 0 occurring with frequency 0.64. If $a > 1$, re-running the method on the new prices might be a better option.

5 Empirical Results

Our method begins by selecting a loss distribution with unknown parameters to be inferred. The polynomial distribution was chosen for analytic tractability. The Lognormal, Pareto, Burr and Gamma distributions were selected because they are commonly used by insurers [40]. Also, they were identified in Section 2. The Weibull distribution was included because it is mentioned in one of the documents analysed in Section 3.

We generated a set of 6,828 price-limit-deductible triples for cyber liability insurance for a hypothetical retail firm with a revenue of \$50M. We selected a multiplicative factor of 1 for all factors other than limit, deductible, industry, revenue and coverage type. Even if we assumed certain security controls were in place, it would be difficult to assign the corresponding multiplicative factors because the value depends on the underwriter’s subjective judgement.

Section 5.1 contains analysis of a set of prices from one insurer, enabling a more detailed description of the metrics introduced previously. In Section 5.2, we extend our analysis to consider all of the data collected in Section 3. The County Fair Cyber Loss Distribution is derived in Section 5.3.

Distribution	$f(x)$	Parameter Values	C_T	C_A
Polynomial	$-\frac{1+a}{c^{a+1}}x^a$	a=-1.369 c= 1.0	-0.261	0.371
Lognormal	$\frac{1}{\sigma\sqrt{2\pi}}e^{-(\log(x)-\mu)^2/2\sigma^2}$	$\mu = 1\text{e-}10, \sigma = 5.418$	-0.312	0.39
Pareto	$\frac{\alpha x_m^\alpha}{x^{\alpha+1}}$	$\alpha = 0.3692, x_m = 0.6712$	-0.266	0.37
Burr	$ck\frac{x^{c-1}}{(1+x^c)^{k+1}}$	c=1.219, k=0.303	-0.264	0.371
Gamma	$\frac{\beta^\alpha}{\Gamma(\alpha)}x^{\alpha-1}e^{-\beta x}$	$\alpha = 0.001579, \beta = 5.02\text{e-}08$	-0.058	0.185
Weibull	$\frac{k}{\lambda}(\frac{x}{\lambda})^{k-1}e^{-(\frac{x}{\lambda})^k}$	$k = 0.3039, \lambda = 0.0001574$	-0.263	0.41

Table 4: The distributions and parameters values minimising C_A across a set of 2,211 cyber liability prices offered by one insurer.

5.1 Analysis of One Insurer

Before considering the prices from all 26 rate schedules, we focus on the insurer with the most extensive pricing set. This pricing schedule contained 2,211 prices for cyber liability coverage, up to a limit of \$50,000,000.

Table 4 identifies parameter values that minimise C_A across all of the observed prices. The Gamma distribution does the best job of predicting these prices. The average error (C_A) is less than 20% of the observed premium and the net error (C_T) is less than 5%. We will achieve as low as 5% absolute error for some insurers in the next subsection.

The relationship between C_T and C_A reveals the balance of over- and under-predictions. The Lognormal is the worst offender for consistently under-predicting, as C_T is close in value to C_A . This is likely to occur when the distribution is not sufficiently heavy tailed to predict the prices. Indeed, Burnecki et al. [40] suggest the Pareto is more appropriate than Lognormal “where exceptionally large claims may occur” and the Burr distribution provides a more flexible heavy tailed distribution.

Both cost functions have an unbounded punishment for over-predictions but a bounded punishment for under-predictions (unless negative prices are predicted). This shifts all of the distributions towards under-prediction. We experimented with a cost function bounding the punishment for over-prediction at +1. For the best performing distributions, there was little difference in the optimal parameter values minimising C_A and the bounded cost function, and so we omitted this analysis. But it is important to note the parameter choices consistently tilt towards under-prediction.

5.2 Market Analysis

We ran the same analysis across all 6,828 of the insurance prices and show the results in Table 5, with each set of prices corresponding to a row. The number of prices (n) in a given

Set	Date	n	L_{max}	Poly	Lognrm	Pareto	Burr	Gamma	Weibull
1	2/17	39	2,000,000	0.322	0.31	0.188	0.324	0.174	0.334
2	9/16	55	5,000,000	0.1	0.101	0.276	0.287	0.183	0.385
3	11/11	56	5,000,000	0.112	0.102	0.22	0.22	0.082	0.282
4	12/16	396	15,000,000	0.174	0.172	0.313	0.317	0.544	0.465
5	8/12	74	5,000,000	0.194	0.182	0.171	0.168	0.156	0.182
6	5/14	458	25,000,000	0.17	0.163	0.333	0.342	0.277	0.516
7	7/09	168	25,000,000	0.327	0.327	0.389	0.391	0.21	0.384
8	4/16	90	25,000,000	0.19	0.237	0.19	0.19	0.156	0.305
9	2/09	132	1,000,0000	0.064	0.055	0.213	0.213	0.125	0.326
10	3/18	374	10,000,000	0.286	0.269	0.398	0.417	0.324	0.507
11	12/18	120	5,000,000	0.713	0.766	0.293	0.689	0.568	0.179
12	9/11	234	20,000,000	0.162	0.167	0.439	0.446	0.229	0.531
13	5/10	135	15,000,000	0.178	0.203	0.178	0.178	0.113	0.242
14	2/15	55	10,000,000	0.149	0.147	0.203	0.203	0.153	0.314
15	2/18	325	10,000,000	0.305	0.336	0.26	0.303	0.208	0.487
16	12/08	290	25,000,000	0.358	0.349	0.387	0.393	0.235	0.418
17	9/10	2,211	50,000,000	0.371	0.39	0.37	0.371	0.185	0.41
18	12/14	54	2,000,000	0.46	0.45	0.336	0.46	0.172	0.35
19	18/17	168	15,000,000	0.359	0.367	0.341	0.359	0.281	0.375
20	10/14	435	50,000,000	0.306	0.28	0.354	0.358	0.227	0.502
21	8/12	81	5,000,000	0.35	0.45	0.15	0.339	0.243	0.423
22	5/08	30	1,000,000	0.369	0.341	0.235	0.292	0.175	0.229
23	2/17	98	10,000,000	0.213	0.21	0.3	0.314	0.194	0.378
24	4/18	490	25,000,000	0.251	0.25	0.258	0.267	0.213	0.367
25	11/13	30	1,000,000	0.369	0.362	0.217	0.312	0.128	0.189
26	2/18	230	1,000,000,000	0.267	0.326	0.436	0.455	0.45	0.72
Mean				0.274	0.279	0.28	0.326	0.222	0.363
Variance				0.138	0.148	0.083	0.111	0.115	0.106

Table 5: The minimal C_A value across for each set of cyber liability prices offered by the insurers in our dataset. The best score is in bold text.

	Polynomial	Lognormal	Pareto	Burr	Gamma	Weibull
Polynomial	-	0.973	0.28	0.595	0.34	-0.017
Lognormal	0.981	-	0.258	0.607	0.389	0.048
Pareto	0.188	0.146	-	0.817	0.656	0.707
Burr	0.756	0.75	0.64	-	0.736	0.567
Gamma	0.422	0.458	0.449	0.685	-	0.668
Weibull	-0.194	-0.146	0.688	0.283	0.374	-

Table 6: Describing how the optimal C_A score for one distribution correlates with the optimal C_A score of another distribution across all of the rate schedules ($n = 26$). The top triangle shows the Spearman’s rank correlation coefficient and the bottom triangle shows the Pearson correlation coefficient.

set is determined by the number of limits and deductibles choices offered by the insurer. The maximum limit (L_{max}) ranges from \$1M to one thousand times that.

No loss distribution consistently outperformed all of the others. This could result from heterogeneity in the sets of observed prices. Factors might include the range of maximum limits and deductibles, or differences in the policy wording. Equally, heterogeneity may result from differing expectations among insurers about losses. In short, the best loss distribution is contingent on the set of prices it aims to predict. So what can we say?

Aggregating the scores suggests the Gamma distribution is the best candidate for predicting cyber liability prices, with the Burr and Weibull distributions performing poorly. However, for price set 11 is better explained by a Weibull parameterisation with a heavy tail. The actuarial model behind the 11th set of prices might expect a relatively more heavy tailed distribution of losses.

Seeing how these scores correlate with each other sheds some light. The scores in Table 6 tilting towards positive suggests that some sets of prices are easier to predict than others. The Lognormal’s performance on a given set of prices is a remarkably good indicator of the polynomial distribution’s performance. If the Weibull predicts a set of prices relatively well, we can expect the Polynomial/Lognormal to do relatively poorly. This provides more evidence of differing expectations regarding how heavy-tailed losses are.

5.3 The County Fair Cyber Loss Distribution

The previous subsection provides 26 potential candidates for the parameterised distribution of cyber losses. An argument could be made that the set of prices with the highest maximum limit is most useful, as this leads to the least extrapolation. Alternatively, we might want to select the distribution achieving the lowest C_A score, even though this discards all but 132 of our data points. Instead we take our lead from Francis Galton’s [41] method for estimating the size of an ox by aggregating guesses from attendees at a county fair. A similar method was applied to estimate the number of jelly beans in a jar [11].

The County Fair Cyber Loss Distribution (CFCLD) is derived by averaging the optimal parameterised loss distribution for each set of prices for a given firm and coverage type. The rest of this section will illustrate the CFCLD by considering cyber liability losses for a retail

	County Fair	Poly	Lognorm	Pareto	Gamma	Weibull
$P(\$0 < X < \$10K)$	0.9146	0.7894	0.9021	0.9968	0.9612	0.6374
$P(\$10K < X < \$50K)$	0.0386	0.1422	0.0581	0.0017	0.0124	0.0007
$P(\$50K < X < \$100K)$	0.0094	0.0257	0.0143	0.0004	0.0052	0.0003
$P(\$100K < X < \$250K)$	0.0089	0.0195	0.0121	0.0004	0.0065	0.0004
$P(\$250K < X < \$500K)$	0.0048	0.0084	0.0055	0.0002	0.0043	0.0003
$P(\$500K < X < \$1M)$	0.0035	0.0053	0.0034	0.0001	0.0035	0.0003
$P(\$1M < X < \$2.5M)$	0.0031	0.0041	0.0025	0.0001	0.0034	0.0004
$P(\$2.5M < X < \$5M)$	0.0015	0.0018	0.001	0.0001	0.0017	0.0003
$P(\$5M < X < \$10M)$	0.0009	0.0012	0.0005	0	0.0011	0.0003
$P(\$10M < X < \$50M)$	0.0007	0.0014	0.0005	0.0001	0.0006	0.0007
$P(\$50M < X < \$100M)$	0.0001	0.0003	0.0001	0	0	0.0003
$P(\$100M < X < \$500M)$	0.0001	0.0004	0	0	0	0.0007
$P(\$500M < X < \$1B)$	0	0.0001	0	0	0	0.0003
$P(\$1B < X < \$10B)$	0.0001	0.0001	0	0	0	0.0011
$E(X)_{min} (\$)$	107,328	242,592	28,360	17,449	23,660	1,3107,18
$E(X)_{mid} (\$)$	428,261	914,258	70,540	75,655	51,214	6,300,519
$E(X)_{max} (\$)$	749,194	1,585,924	112,719	133,860	78,768	11,290,318

Table 7: An overview of the distribution of losses for the County Fair Cyber Loss Distribution, as well as the average contribution from each of the distributions.

firm with revenue of \$50M across all 26 pricing schemes. We could easily derive the CFCLD for a different revenue, firm size or coverage type.

Table 7 displays the probability of different loss amounts according to this CFCLD. Treating 0 as part of a continuous distribution means these loss distributions underestimate the proportion of firms facing no losses. Interpreting smaller losses (such as those less than \$50K) as 0 may correct for this. This adjustment suggests an incident rate of 0.0468 per year since 95% of losses are less than \$50K.

The probability of a loss of between \$100K and \$250K is 0.0089. This falls to 0.0083 for losses between \$250,000 and \$1M. The probability of a loss of \$1M–\$10M is 0.0045, falling to 0.0009 for \$10–100M.

Inferences about losses beyond the maximum limit are essentially extrapolation. The indemnity payment for a loss exceeding the limit by a dollar is the same as for a loss exceeding the limit by a billion dollars. This leads to a question regarding losses exceeding \$50M as only one firm provides limits beyond this. For the CFCLD, this amounts to 1.4% of the distribution and it is not clear how it should be interpreted.

Table 7 also describes the average contribution to the CFCLD from each distribution. The contributions of each distribution are weighted according to how often each distribution led to the highest C_A score for a given set of prices. Consequently the Gamma, Lognormal, Polynomial, Pareto and Weibull distributions contribute 16, 5, 3, 1 and 1 respectively, with no contribution from the Burr distribution.

The implied Pareto distribution is notable for its vanishing tail and it corresponds to the set of prices including the remarkably cheap policy in Figure 2. The Weibull distribution (at least for this parameterisation) is a notable outlier and provides most of the support for the

CFCLD’s tail. In fact, the singular Weibull parameterisation contributes 99.2% of the losses greater than \$100M despite comprising one 26th of the CFCLD.

Table 7 also includes approximations of the expected loss of each distribution using the buckets in the table. The values of $E(X)_{min}$, $E(X)_{mid}$ and $E(X)_{max}$ are calculated by summing the $P(a < X < b)$ weighted by the minimum (a), midpoint ($\frac{b-a}{2}$), and maximum b of the range respectively. This presentation allows the reader to understand which sections of each distribution are contributing the most. For example, the over-extrapolated section of the Weibull distribution representing losses of between \$1 billion and 10 billion contributes 54% of $E(X)_{mid}$.

Readers might instead calculate expected losses by counting all losses above \$100M as \$100M. They might also decide to omit the contribution of the Weibull distribution. Extracting a concrete expected loss from these distributions is so challenging because any expected loss relies on some degree of extrapolation, unless the maximum limit exceeds plausible losses.

Section 4.4 provides a simple way of converting between different coverage types. Table 7 describes the cyber liability losses of a retail firm with revenue of \$50M. Section 3 showed that business interruption is priced at 38% of cyber liability. Section 4.4 suggests we can multiply all of the probabilities by 0.38 to estimate the distribution of business interruption incidents for the same hypothetical firm. This results in an expected loss of \$154K for business interruption events. The same process for ransomware and wire transfer incidents leads to expected losses of \$64K and \$51K respectively.

6 Discussion

We discuss how the results relates to other attempts to quantify cyber losses in Section 6.1. We critique our method in Section 6.2.

6.1 Quantifying Cyber Losses

The results allows us to comment on the shape of the distribution of losses and point-estimates of losses in dollars. Not identifying a single best distribution for cyber losses is in line with data breach studies, which have found no consensus; subsequent studies concluded breaches were best described by power law [16], Lognormal [17], Pareto [18] and log-skew-normal [19] distributions. However, the Gamma distribution being the most likely candidate differs from such studies.

Liability dollar losses could plausibly be less heavy tailed than distributions of the number of records in a data breach. Liability costs are assigned by courts, which are unlikely to follow the distributions found in data breaches. For example, an equivalent legal ruling to Yahoo!’s⁵ breach affecting up to 3 billion accounts is unlikely as it is an order of magnitude bigger than the next largest. Further evidence can be found in a study [29] of operational losses, which found that “the distribution of the non-cyber risk sample is much heavier tailed than that of the cyber risk”.

⁵<https://www.reuters.com/article/us-yahoo-cyber/yahoo-says-all-three-billion-accounts-hacked-in-2013>

Pricing different coverage types by multiplying the price of cyber liability coverage by a constant suggests each type of cyber loss is driven by a similarly shaped distribution. We have two separate studies [21, 22] that find the recovery time for IT systems to be well modelled by a Lognormal distribution, much like number of records in [17]. Another study investigating the time required to remove phishing websites has also been found to follow a lognormal distribution [42]. Although we cannot map these distributions to financial costs, this evidence supports the viability of multiplicative pricing. Given the different coverage types continue to be bought and sold, we can make estimates about previously under-studied types of cyber losses like ransomware.

The lower-bound, mid-point and upper-bound for the expected loss from the CFCLD are \$107K, \$428K and \$749K respectively. For comparison, one study of 921 event costs has mean of \$7.84M, median of \$250K and maximum of \$750M [33]. Another study found mean and median cyber losses of \$41M and \$1.9M respectively [29].

Why are the point estimates from the County Fair Loss Distribution smaller than related studies? Our inferences are based on the insurance prices for a retail firm with a revenue of \$50 million. Meanwhile we do not know the size or industry of the firms constituting the samples upon which [33, 29] are based. The firms likely have revenues exceeding \$50M. This highlights the problem of granularity in cyber security data.

6.2 Reflecting on the Method

If our aim is to uncover the true loss distribution, there are two types of error: insurer error and method error. Insurer error is driven by the insurer’s uncertainty around the loss distribution faced by an insured party. Method error is introduced by making flawed inferences from the observed prices. We can only seek to reduce method error, but it is worth discussing both.

Insurer Error Insurer error results from uncertainty in quantifying cyber risk (to be contrasted against certainty in quantifying the risk of a coin toss). This results in Knightian “unmeasurable” uncertainty [43] as evidenced by discussions in which insurers bemoan the lack of actuarial data in cyber insurance [44]. Even worse than random uncertainty, insurers may exhibit systemic bias resulting from tight professional networks. The resulting group-think could prevent aggregated information from mitigating random noise [45]. For example, the solicitors professional indemnity market was systemically under priced in 2010 [46]. Finally, the phenomenon of the underwriting cycle, in which prices across all lines of insurance rise and fall cyclically [47], illustrates how actual prices differ from the actuarially fair price. This introduces noise into our inferences.

Although the causes of cyber losses are new to insurers, costs are often realised like traditional insurance lines. For example, cyber liability is still assigned by courts and cyber business interruption is calculated via lost revenues. Insurers have amassed much experience quantifying realised losses; the industry dates to at least 14th Century Italy [48] when insurance could be purchased to cover voyages into uncharted territory. Professional qualifications should help with probability estimates given they are improved by even light-weight calibration training interventions [11]. Aggregating the inferred distributions is intended

to reduce random noise, much like how different political polls are combined into a *poll of polls* [10].

Prediction markets provide incentives for deviating from group-think. An analogous argument would go: if a line of insurance is systemically over-priced, then individual insurers have an incentive to defect by offering lower prices to increase market share. Figure 2 shows how more market entrants drove some insurers to reduce prices. Although losses should discipline insurers generous underwriting, competitive pressures can temporarily lead to systemic under pricing. Cass Sunstein’s aphorism “markets incorporate falsehood as well as truth” [10] summarises this discussion.

Method Error Even if the insurer priced risk perfectly, we might make flawed inferences based on those prices. Results are influenced by the sample of rate schedules and prices. Section 5 showed that the inclusion of rate schedule 11 more than doubled the expected loss from the County Fair Cyber Loss Distribution. For our purposes, choosing the largest data set was justified because any anomalies illustrate challenges for future work to overcome. Further, we presented the results so that the reader can extract loss estimates excluding the contribution of anomalous distributions.

Our model assumes risks are considered independently. In reality, insurers must consider how losses correlate with each other and implement costly solvency risk mitigation measures like reinsurance or holding more capital. Not considering insurer risk aversion or wealth endowment could lead to flawed inferences. These factors complicate extracting probabilities from prediction markets [49]. Our method is likely to face even greater challenges given the relative complexity of insurance markets as compared to the binary events considered by conventional prediction markets.

Given the novelty of the method, it is unclear whether the prediction errors are acceptable. For example, the best performing distribution (the Gamma) across all of the rate schedules has an average error of 22%. In our defence, there is heterogeneity in the expected losses of each insurer and the observed prices do not account for the underwriter’s subjective adjustment. Further, we are predicting as many as 2,200 data points with just two parameters and one distribution, which are commonly used by insurers [40]. More flexible distributions would improve predictions but increase the risk of over-fitting.

Making the best inference relies on solving an optimisation problem in a non-convex search space. The possibility of terminating at a local minima is ever present. Understanding the search space better and selecting the appropriate optimisation heuristic is a must for future work.

7 Conclusion

We provided empirical observations on how 26 cyber insurance providers in California vary premium by coverage type, amount, policyholder type and over time. The price of business interruption, wire transfer fraud and PCI costs as a percentage of cyber liability coverage are 36%, 12% and 24% respectively. Most insurers provide logarithmic increases in premium for a linear increase in the insured’s revenue. The data is inconclusive regarding whether prices are trending downwards, but the rate of market entrance has increased in recent years.

We introduced a method to infer loss distributions from insurance prices. The method uses particle swarm optimisation to iterate through candidate parameter values to identify the parameterised loss distribution which best explains the observed prices. The Gamma, Lognormal, Polynomial, Pareto and Weibull distributions best predicted 16, 5, 3, 1 and 1 respectively of the 26 sets of premiums.

The County Fair Cyber Loss Distribution aggregates each of the 26 parameterised distributions to provide estimates about cyber losses for a retail firm with a revenue of \$50M. The results suggest the expected loss resulting from cyber liability incidents is \$428K with a 0.0055 probability of a loss of between \$1M and \$10M. Expected losses for business interruption, ransomware and wire transfer incidents are \$154K, \$64K and \$51K respectively.

These estimates complement existing approaches to quantifying cyber losses by providing distributions of dollar losses, cost estimates for novel incident types, and granular insights for a specific revenue and industry. This first attempt at inferring losses from insurance prices can be improved by speaking to insurance professionals to understand how to construct a better sample of prices, including more flexible distributions to improve predictions, and by analysing performance on prices generated by a known distribution.

Acknowledgements

The authors would like to thank Rainer Böhme, Ulrik Franke, Steve Moyle, Jonathon Spring, and the anonymous reviewers for reading various versions of this paper and providing constructive feedback. The NAIC should also be commended for making rate schedules available to the public. The collaboration was made possible by a Fulbright Cybersecurity Scholar award from the US-UK Fulbright Commission. Daniel Woods' research is funded by the EPSRC.

References

- [1] Ross Anderson and Tyler Moore. The economics of information security. *Science*, 314(5799):610–613, 2006.
- [2] Eric Jardine. Mind the denominator: towards a more effective measurement system for cybersecurity. *Journal of Cyber Policy*, 3(1):116–139, 2018.
- [3] Dinei Florêncio and Cormac Herley. Sex, lies and cyber-crime surveys. In *Economics of information security and privacy III*, pages 35–53. Springer, 2013.
- [4] Ross Anderson, Chris Barton, Rainer Böhme, Richard Clayton, Michel JG Van Eeten, Michael Levi, Tyler Moore, and Stefan Savage. Measuring the cost of cybercrime. In *The economics of information security and privacy*, pages 265–300. Springer, 2013.
- [5] Friedrich August Hayek. The use of knowledge in society. *The American economic review*, 35(4):519–530, 1945.
- [6] Burton G Malkiel. The efficient market hypothesis and its critics. *Journal of economic perspectives*, 17(1):59–82, 2003.

- [7] David Rothschild. Forecasting elections: Comparing prediction markets, polls, and their biases. *Public Opinion Quarterly*, 73(5):895–916, 2009.
- [8] Philip M Polgreen, Forrest D Nelson, George R Neumann, and Robert A Weinstein. Use of prediction markets to forecast infectious disease activity. *Clinical Infectious Diseases*, 44(2):272–279, 2007.
- [9] Adam Mann. The power of prediction markets. *Nature News*, 538(7625):308, 2016.
- [10] Cass R Sunstein. *Infotopia: How many minds produce knowledge*. Oxford University Press, 2006.
- [11] Douglas W Hubbard. *How to measure anything: Finding the value of intangibles in business*. John Wiley & Sons, 2014.
- [12] Rainer Böhme. A comparison of market approaches to software vulnerability disclosure. In *International Conference on Emerging Trends in Information and Communication Security*, pages 298–311. Springer, 2006.
- [13] Daniel W Woods and Andrew C Simpson. Cyber-warranties as a quality signal for information security products. In *International Conference on Decision and Game Theory for Security*, pages 22–37. Springer, 2018.
- [14] Gordon Gemmill and Apostolos Saffekos. How useful are implied distributions? evidence from stock index options. *Journal of Derivatives*, 7(3):83–98, 2000.
- [15] Matthew Curtin and Lee T Ayres. Using science to combat data loss: Analyzing breaches by type and industry. *Journal of Law and Policy for the Information Society*, 4:569, 2008.
- [16] Thomas Maillart and Didier Sornette. Heavy-tailed distribution of cyber-risks. *The European Physical Journal B*, 75(3):357–364, 2010.
- [17] Benjamin Edwards, Steven Hofmeyr, and Stephanie Forrest. Hype and heavy tails: A closer look at data breaches. *Journal of Cybersecurity*, 2(1):3–14, 2016.
- [18] Spencer Wheatley, Thomas Maillart, and Didier Sornette. The extreme risk of personal data breaches and the erosion of privacy. *The European Physical Journal B*, 89(1):7, 2016.
- [19] Martin Eling and Nicola Loperfido. Data breaches: Goodness of fit, pricing, and risk measurement. *Insurance: Mathematics and Economics*, 75:126–136, 2017.
- [20] Maochao Xu, Kristin M Schweitzer, Raymond M Bateman, and Shouhuai Xu. Modeling and predicting cyber hacking breaches. *IEEE Transactions on Information Forensics and Security*, 13(11):2856–2871, 2018.
- [21] Bianca Schroeder and Garth Gibson. A large-scale study of failures in high-performance computing systems. *IEEE Transactions on Dependable and Secure Computing*, 7(4):337–350, 2010.

- [22] Ulrik Franke, Hannes Holm, and Johan König. The distribution of time to recovery of enterprise it services. *IEEE Transactions on Reliability*, 63(4):858–867, 2014.
- [23] Verizon LLC. 2018 Data Breach Investigations Report available at <https://enterprise.verizon.com/en-gb/resources/reports/dbir/>, 2018.
- [24] Ponemon Institute. Cost of a data breach study available at <https://www.ibm.com/security/data-breach>, 2018.
- [25] Chad D Heitzenrater and Andrew C Simpson. Policy, statistics and questions: Reflections on UK cyber security disclosures. *Journal of Cybersecurity*, 2(1):43–56, 2016.
- [26] Claudia Biancotti. The price of cyber (in) security: evidence from the italian private sector. In *Proceedings of The 17th Workshop on the Economics of Information Security (WEIS 2018)*, 2018.
- [27] Sasha Romanosky, David Hoffman, and Alessandro Acquisti. Empirical analysis of data breach litigation. *Journal of Empirical Legal Studies*, 11(1):74–104, 2014.
- [28] Aaron Ceross and Andrew Simpson. The use of data protection regulatory actions as a data source for privacy economics. In *International Conference on Computer Safety, Reliability, and Security*, pages 350–360. Springer, 2017.
- [29] Christian Biener, Martin Eling, and Jan Hendrik Wirfs. Insurability of cyber risk: An empirical analysis. *The Geneva Papers on Risk and Insurance-Issues and Practice*, 40(1):131–158, 2015.
- [30] Anat Hovav and John D’Arcy. The impact of denial-of-service attack announcements on the market value of firms. *Risk Management and Insurance Review*, 6(2):97–121, 2003.
- [31] Katherine Campbell, Lawrence A Gordon, Martin P Loeb, and Lei Zhou. The economic cost of publicly announced information security breaches: empirical evidence from the stock market. *Journal of Computer Security*, 11(3):431–448, 2003.
- [32] Huseyin Cavusoglu, Birendra Mishra, and Srinivasan Raghunathan. The effect of internet security breach announcements on market value: Capital market reactions for breached firms and internet security developers. *International Journal of Electronic Commerce*, 9(1):70–104, 2004.
- [33] Sasha Romanosky. Examining the costs and causes of cyber incidents. *Journal of Cybersecurity*, 2(2):121–135, 2016.
- [34] Sasha Romanosky, Rahul Telang, and Alessandro Acquisti. Do data breach disclosure laws reduce identity theft? *Journal of Policy Analysis and Management*, 30(2):256–286, 2011.
- [35] Department for Business, Innovation Skills. Information security breaches survey, 2015. [Online; accessed 27-July-2016].

- [36] Sasha Romanosky, Lillian Ablon, Andreas Kuehn, and Therese Jones. Content analysis of cyber insurance policies: How do carriers write policies and price cyber risk? In *Proceedings of The 16th Workshop on the Economics of Information Security (WEIS 2017)*, 2017.
- [37] James Kennedy. Particle swarm optimization. In *Encyclopedia of machine learning*, pages 760–766. Springer, 2011.
- [38] Lester James V. Miranda. Pyswarms, a research-toolkit for particle swarm optimization in python, 2017. 10.5281/zenodo.986300.
- [39] David Bahnemann. Distributions for actuaries. *CAS Monograph Series*, (2), 2015.
- [40] Krzysztof Burnecki, Grzegorz Kukla, and Rafał Weron. Property insurance loss distributions. *Physica A: Statistical Mechanics and its Applications*, 287(1-2):269–278, 2000.
- [41] Francis Galton. Vox populi (the wisdom of crowds). *Nature*, 75(7):450–451, 1907.
- [42] Tyler Moore and Richard Clayton. Examining the impact of website take-down on phishing. In Lorrie Faith Cranor, editor, *APWG eCrime Researchers Summit*, volume 269 of *ACM International Conference Proceeding Series*, pages 1–13. ACM, 2007.
- [43] Frank H Knight. *Risk, uncertainty and profit*. Courier Corporation, 2012.
- [44] Daniel W Woods and Andrew C. Simpson. Policy measures and cyber insurance: A framework. *Journal of Cyber Policy*, 2(2):209–226, 2017.
- [45] Scott E. Page. *The Difference: How the Power of Diversity Creates Better Groups, Firms, Schools, and Societies*. Princeton University Presss, 2007.
- [46] Daniel W Woods, Ioannis Agraftotis, Jason RC Nurse, and Sadie Creese. Mapping the coverage of security controls in cyber insurance proposal forms. *Journal of Internet Services and Applications*, 8(1):8, 2017.
- [47] Greg Niehaus and Andy Terry. Evidence on the time series properties of insurance premiums and causes of the underwriting cycle: new support for the capital market imperfection hypothesis. *Journal of Risk and Insurance*, pages 466–479, 1993.
- [48] Humbert O Nelli. The earliest insurance contract. a new discovery. *Journal of Risk and Insurance*, pages 215–220, 1972.
- [49] Charles F Manski. Interpreting the predictions of prediction markets. *economics letters*, 91(3):425–429, 2006.

Set	Poly	Lognrm	Pareto	Burr	Gamma	Weibull
1	-1.415, 6.39	0.6926, 5.1	0.329, 0.03881	0.2066, 1.576	0.007277, 1.147e-06	0.08108, 0.001172
2	-1.62, 564.9	5.643, 3.364	0.2652, 2.877	1.133, 0.2318	0.02423, 1.238e-06	0.06136, 0.002266
3	-1.66, 548.6	4.612, 3.692	0.3022, 2.511	0.6558, 0.4608	0.008544, 6.02e-07	0.5245, 0.001431
4	-1.655, 308.7	4.411, 3.593	0.3517, 2.501	1.2, 0.2931	0.02773, 3.106e-06	0.06657, 0.0005691
5	-1.286, 2.255	0.3091, 6.4	0.2558, 0.1918	0.1512, 1.494	0.008674, 9.961e-08	0.06174, 0.002714
6	-1.655, 340.9	4.694, 3.515	0.3517, 3.0	0.7278, 0.4834	0.008751, 5.452e-07	0.6507, 0.0005137
7	-1.396, 255.8	4.974, 4.838	0.2031, 0.1322	0.7801, 0.2604	0.01345, 1.018e-07	0.06178, 0.002765
8	-1.398, 1.19	1e-10, 5.126	0.3922, 0.3753	0.7442, 0.5269	0.001502, 9.592e-08	0.05443, 0.0002993
9	-1.722, 710.6	4.859, 3.475	0.3338, 2.073	0.8557, 0.3901	0.007281, 4.453e-07	0.3396, 0.0008387
10	-1.708, 494.9	4.722, 3.441	0.3449, 3.0	1.032, 0.3336	0.01308, 1.094e-06	0.05877, 0.0006117
11	-1.489, 1.0	1e-10, 4.215	0.3377, 5e-10	0.2231, 2.094	0.003038, 1.892e-07	0.1066, 0.001243
12	-1.796, 4023.0	7.622, 2.728	0.2625, 3.0	1.036, 0.2534	0.03272, 8.414e-07	0.0433, 0.001744
13	-1.367, 1.371	1e-10, 5.452	0.3571, 0.2254	0.5841, 0.6112	0.002467, 1.147e-07	0.491, 0.0004963
14	-1.706, 262.3	2.881, 3.949	0.3924, 2.878	1.165, 0.3368	0.008694, 1.168e-06	0.0536, 0.0004344
15	-1.42, 1.0	1e-10, 4.828	0.4142, 0.1765	0.3608, 1.159	0.001997, 2.887e-07	0.8208, 0.0002809
16	-1.468, 172.3	4.228, 4.426	0.2697, 0.01978	0.6569, 0.4132	0.01359, 2.172e-07	0.0687, 0.001438
17	-1.369, 1.0	1e-10, 5.418	0.3692, 0.6712	1.219, 0.303	0.001579, 5.02e-08	0.3039, 0.0001574
18	-1.438, 10.51	0.5467, 5.178	0.325, 0.03928	0.1881, 1.707	0.006347, 9.875e-07	0.08164, 0.001064
19	-1.307, 1.042	1e-07, 6.028	0.3072, 0.3652	0.6888, 0.446	0.004438, 1.012e-07	0.07249, 0.001049
20	-1.417, 82.3	3.978, 4.738	0.2743, 2.014	0.699, 0.3914	0.008581, 9.422e-08	0.04431, 0.0007566
21	-1.469, 1.0	1e-07, 4.339	0.4641, 0.04153	0.1861, 2.393	0.001529, 6.166e-07	0.6957, 0.0002709
22	-1.44, 124.6	2.485, 5.876	0.1988, 3.06e-05	0.1073, 1.118	0.01903, 9.225e-07	0.08372, 0.004
23	-1.709, 311.7	4.34, 3.406	0.3818, 3.0	0.6891, 0.5541	0.01925, 2.315e-06	0.0584, 0.0005017
24	-1.477, 20.38	1e-13, 5.427	0.3622, 0.1995	0.6908, 0.5254	0.002696, 1.385e-07	0.7393, 0.0004358
25	-1.439, 20.87	5.28e-07, 5.914	0.2942, 0.004664	1.043, 0.007721	0.01292, 2.127e-06	0.09196, 0.002301
26	-1.478, 104.4	4.655e-07, 6.201	0.3179, 9.962	0.8672, 0.3542	0.008126, 2.027e-07	0.03755, 0.0002793

Table 8: The parameter values corresponding to Table 5.