

Getting started with risk quantification using quantrr

John Benninghoff
Security Differently



Hi, I'm John Benninghoff. I started my consulting company, Security Differently, with a goal of making cybersecurity less scary, and as much a part of technology engineering as safety is part of mechanical or structural engineering.

This talk is a demo of a new tool. While the tool is new, the concepts are based on a lot of past work, my peer David Grimmer's work starting risk quantification at our last company, and my own analysis and research. It's also the story of my ongoing journey to better model and estimate risk.

I'll have a QR code at the end for you to download the slides with notes and links to all the references.



What's my story? Me on upper left, wife Jolene and our dog Gertie. Started in security after attending SANS Network Security 1998. 20 years later, MSc in safety science (managing risk and systems change, 2018-2021). More recently, I worked in Site Reliability Engineering, starting in 2020, and spoke at SREcon earlier this year!

SANS: <https://www.sans.org>

TCD: <https://psychology.tcd.ie/postgraduate/msc-riskandchange/>, image:

https://commons.wikimedia.org/wiki/File:Trinity_college_library.jpg

SREcon: <https://www.usenix.org/srecon>

I ♥ Interruptions!

I'm known to be guilty of interrupting. As this is an interactive talk, please feel free to jump in and ask questions!

Why CRQ?

What is the value of (Cyber) Risk Quantification?

Before we launch into the demo, I'll present a recap of the history of cyber risk quantification, and my motivations for building a new tool.

What is the value of (Cyber) Risk Quantification? I see risk quantification as a tool to help inform and improve organizational decisions, primarily investments. I've seen CRQ used effectively to get funding to reduce risk.

Blunt End and Sharp End



In safety, we talk about the blunt end and the sharp end: the sharp end of the organization is the people who do the work, and the blunt end is leadership. There's a gap in understanding; the executives at the blunt end don't and can't have as complete an understanding as those at the sharp end. (Also true for different practitioners). Risk quantification helps bridge the gap by informing the blunt end.

Image: Figure 3 from: Cook, R., Woods, D., & Miller, C. A. (1998). A Tale of Two Stories: Contrasting Views of Patient Safety.
https://www.researchgate.net/publication/245102691_A_Tale_of_Two_Stories_Contrasting_Views_of_Patient_Safety

Risk Matrices Considered Harmful

		Severity				
		Very low	Low	Medium	High	Very high
Probability	Very high					
	High					
	Medium					
	Low					
	Very Low					

Risk matrices are a popular decision tool, but we have good evidence that they lead to poorer decisions than not using them at all: they are qualitative (how do you compare a “red” risk against an investment of \$1M with an expected \$500K return?) and don’t communicate uncertainty (typically, you pick only one box, not multiple).

<https://safetyofwork.com/episodes/ep8-do-risk-matrices-help-us-make-better-decisions>

Image: https://commons.wikimedia.org/wiki/File:Risk_heat_map.png (David Vose)

https://www.linkedin.com/posts/davidvoserisk_inmemoriam-davidvose-riskanalysis-activity-7221501204109930497-wWkC

Cyber Risk Quantification (CRQ)

- Metrology: all measurements are estimates
- *How to Measure Anything*: expert estimation
- Factor Analysis of Information Risk (FAIR)

How did CRQ evolve? Metrology, the study of measurement, teaches us that all measurements are estimates, with a certain level of accuracy and precision. The work of Doug Hubbard, as captured in the book *How to Measure Anything*, asks experts to estimate a range of outcomes, which captures uncertainty. Jack Jones adapted the work of Hubbard and others to create FAIR, the leading methodology for quantifying cyber risk.

<https://en.wikipedia.org/wiki/Metrology>

https://www.goodreads.com/book/show/444653.How_to_Measure_Anything,

<https://hubbardresearch.com>

https://en.wikipedia.org/wiki/Factor_analysis_of_information_risk,

<https://www.opengroup.org/open-fair>, <https://www.fairinstitute.org/what-is-fair>

Why do we need another RQ tool?

Commercial

- [Safe Security](#)
- [RiskRecon](#)
- [Ostrich](#)
- [Riskconnect](#)
- [Derive](#)
- [Onspring](#)
- [Axio](#)

• [SecurityScorecard](#)

- [Black Kite](#)
- [ProcessUnity](#)
- [Kovrr](#)
- [Erambra](#)
- [ThreatConnect](#)
- [CyberSaint](#)
- [Alfahive](#)

• [Monaco Risk](#)

- [Vivo Security](#)

Limited Use

- [FAIR-U](#)
- [Open FAIR™](#)
- [Guesstimate](#)

Free!

- [tidyrisk*](#)

So many tools... but most are commercial products or limited-use. tidyrisk is great (but no longer in development)!

I wanted to create something simpler that wasn't based on FAIR.

A tool for people who have...



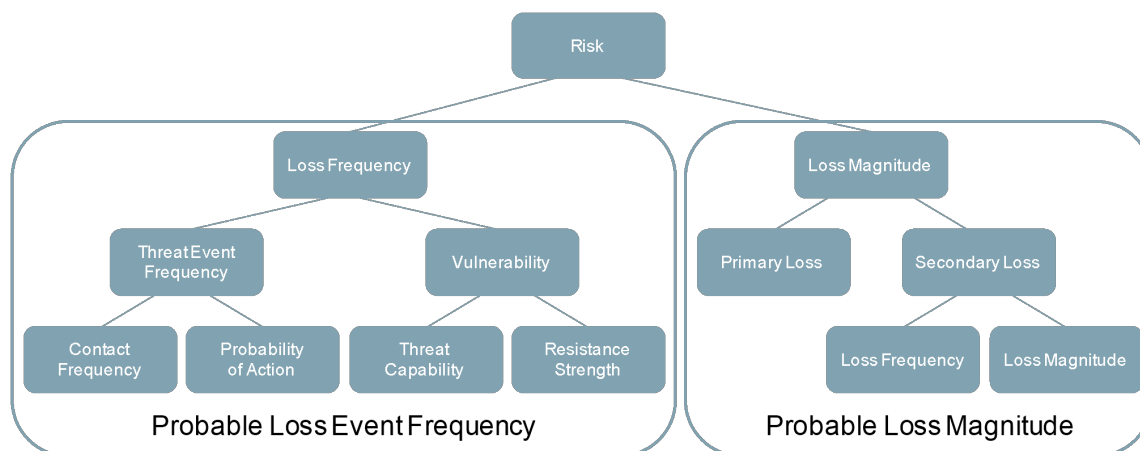
Limited Time



No Budget

The demo is for people with limited time and essentially no budget; the one person trying to start Risk Quant who understands enough math to explain a basic model.

Limitations of FAIR – “What are we missing?”



More factors isn't necessarily better; there is evidence that *fewer* factors give better estimates, and in my opinion, more factors serves the Analyst, but not the experts; it's easier for the experts to simply estimate frequency and magnitude directly, and use fast-thinking for the many factors that contribute to each. Additionally, by focusing only on cybersecurity risk, FAIR excludes non-security risks that may be larger. By asking "What are we missing?" the analyst is mining for knowledge of hidden risks.

* I lost track of the reference on the benefit of fewer factors; it was from a talk Miles Edmunson gave at Secure360, where he spoke about using Monte Carlo for risk estimation (without knowledge of FAIR)

Image: https://pubs.opengroup.org/security/openfair-process-guide/#_Toc503856057

Expert Elicitation and Calibration

Do you know what you know?

A Confidence Calibration Exercise

After answering each of the true/false questions below, indicate how confident you are in your answer using the corresponding slider. A value of 50% means you have no idea what the right answer is (the same probability as a random guess between the two choices); a value of 100% means you are completely confident in your answer.

English is the top language used on the Internet

☐ True ☐ False Probability of correctness: 50%

The ferrule connects the bristles to the handle on a paint brush

☐ True ☐ False Probability of correctness: 50%

Neutrons are the particles in an atom that have positive charge

☐ True ☐ False Probability of correctness: 50%



The accuracy of expert estimates can be affected by under or overconfidence (typically over). Expert estimates can be calibrated to increase the accuracy by reducing precision – a 90% confidence range. Key tools are calibration exercises (to test your calibration) and the equivalent bets method (imagine a wheel with 9 green wedges and 1 red one; you win 90% of the time and lose 10%. Do you want to bet \$100 on your estimate, spin the wheel, or is there no difference?)

<https://www.fairinstitute.org/blog/calibrated-estimation-for-fair-cyber-risk-quantitative-analysis-explained-in-3-to-4-minutes>

<https://www.tonym-v.com/blog/2019/10/2/improve-your-estimations-with-the-equivalent-bet-test>

<http://confidence.success-equation.com>

<https://perfectlyconfident.com>

Images:

<http://confidence.success-equation.com>

https://commons.wikimedia.org/wiki/File:Wheel_of_Fortune_-_Season_26_-_Round_4.svg

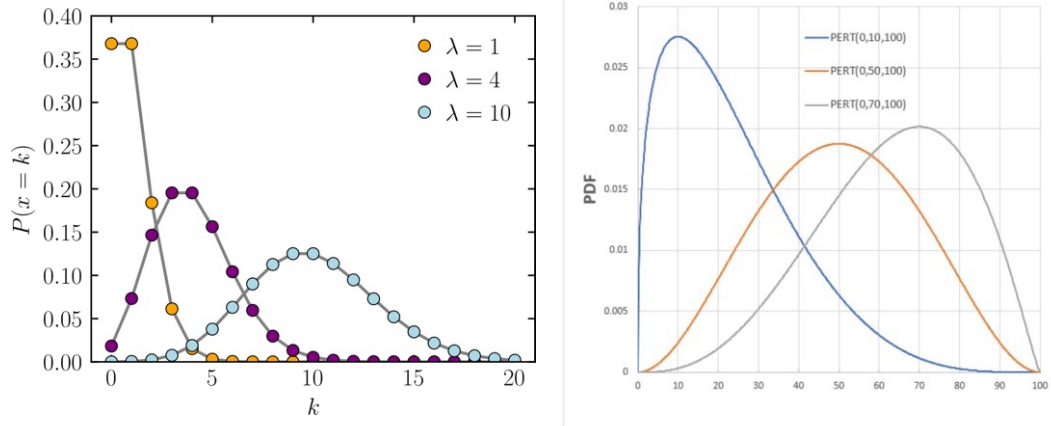
Demo

Show me, don't tell me

The demo is meant to be interactive – please ask questions!

Spreadsheet > Report > Code Walkthrough.

Limitations of FAIR – probability distributions



While OpenFAIR doesn't prescribe a specific distribution, historically, FAIR uses BetaPERT. There are 2 issues with this: first, frequency is better modeled with a discrete distribution, like Poisson.

The Open Group. (2021). *Risk Analysis (O-RA), Version 2.0.1*.

<https://publications.opengroup.org/c20a>

Images: https://commons.wikimedia.org/wiki/File:PERT_pdf_examples.jpg,

<https://commons.wikimedia.org/wiki/File:Log-normal-pdfs.png>

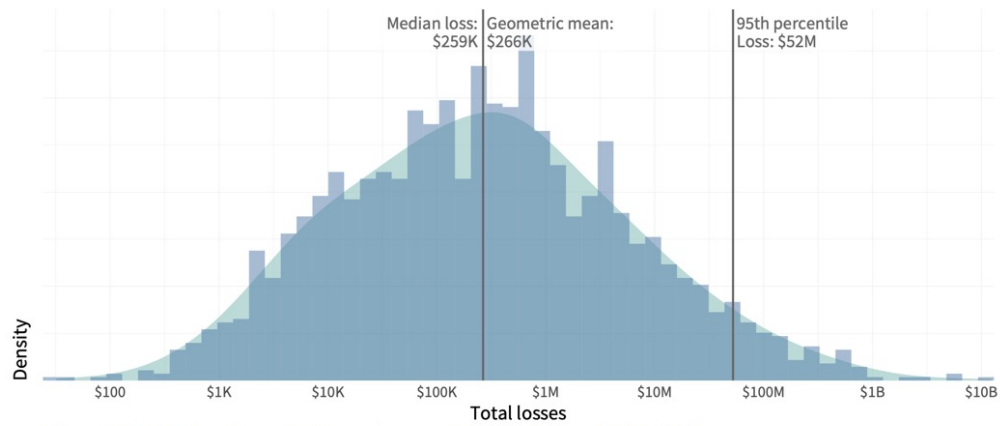
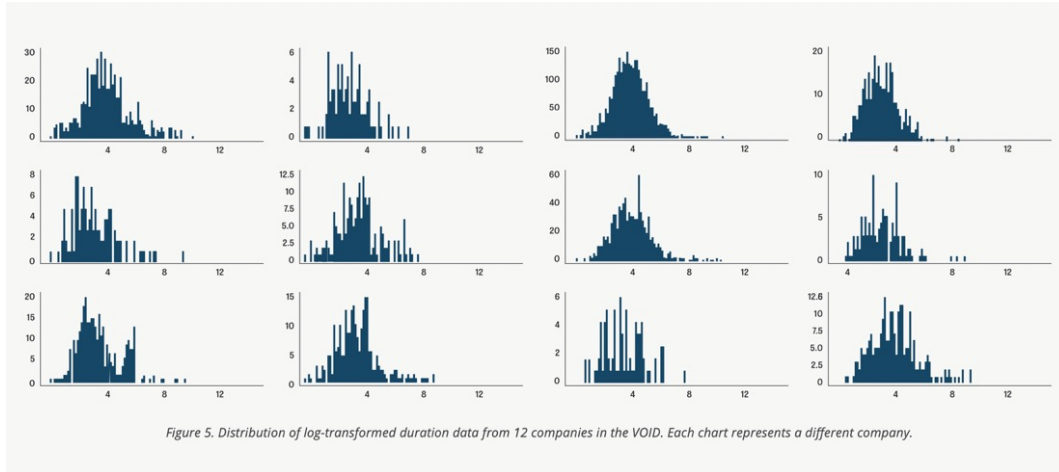


Figure 6: Distribution of reported losses for security incidents from 2012 to 2021

The Cyentia IRIS reports show that loss distribution is log-normal.

Cyentia Institute. (2022). *Information Risk Insights Study (IRIS) 2022*.
https://www.cyentia.com/wp-content/uploads/IRIS-2022_Cyentia.pdf



Technology outage duration times are also generally log-normal. In my own work, I found that outages of a specific type did fit log-normal quite well, some of these may be mixing multiple types and are multi-modal as a result.

Nash, C. (2022). *The VOID Report 2022*. <https://www.thevoid.community/report>

Practical Takeaways



Scoping and
Scheduling



Outlier Experts



Modeling and
Communication

Some common themes I've heard from those who have done risk quantification:

1. Scoping and Scheduling: the main challenges are scoping the risks, finding the experts, and scheduling time for the interviews. The time spent and running the models is comparatively easy.
2. Outlier Experts: sometimes an expert is far different from the rest. Methods of weighting expert opinion don't improve the estimate, and typically one expert won't change the story much.
3. Modeling and Communication: in practice, the model is less important – its primary value is in facilitating discussion, discovery, and bringing knowledge from front-line workers to management.

Get Involved!

- Try out the report
- Direct feedback
- Open an issue
- Submit a pull request



I'd love to make quantrr a community-supported project, but I need your help – if you're interested, try it out, you can send me feedback directly, open an issue on GitHub, or submit code via a pull request.

<https://github.com/jabenninghoff/quantrr/>

Slides, Connect & Resources



Resources:

jabenninghoff.github.io/quantrr/

Connect:

linkedin.com/in/jbenninghoff/

Website:

jbenninghoff.com
security-differently.com



Scan the QR code for slides and more! Questions?