

Using Bayesian Networks to Model Expected and Unexpected Operational Losses

Martin Neil,^{1*} Norman Fenton,¹ and Manesh Tailor²

This report describes the use of Bayesian networks (BNs) to model statistical loss distributions in financial operational risk scenarios. Its focus is on modeling “long” tail, or unexpected, loss events using mixtures of appropriate loss frequency and severity distributions where these mixtures are conditioned on causal variables that model the capability or effectiveness of the underlying controls process. The use of causal modeling is discussed from the perspective of exploiting local expertise about process reliability and formally connecting this knowledge to actual or hypothetical statistical phenomena resulting from the process. This brings the benefit of supplementing sparse data with expert judgment and transforming qualitative knowledge about the process into quantitative predictions. We conclude that BNs can help combine qualitative data from experts and quantitative data from historical loss databases in a principled way and as such they go some way in meeting the requirements of the draft Basel II Accord (Basel, 2004) for an advanced measurement approach (AMA).

KEY WORDS: Basel Accord; Bayesian nets; operational risk

1. INTRODUCTION

The Basel Committee on Banking Supervision, in reaction to a number of well-publicized financial disasters, has drafted a system of regulation addressing the issue of operational risk (OpRisk) and its assessment (Basel, 2004). Key to the regulatory process is the modeling of a business’s operational risks, in terms of a variety of loss event types, in order to arrive at an appropriate regulatory capital charge. To calculate such a charge it is tempting to predict operational risk by building a statistical model based on historical data. However, from a *statistical* perspective the well-publicized financial disasters, e.g., Barings (Rawnsley, 1995) and Allied Irish Bank (AIB) (Wachtell *et al.*,

2002), in themselves, are too few in number for any meaningful inference. Moreover, until recently, banks have not historically collected loss event data on a wide and systematic basis. This general paucity of loss data means that traditional statistical approaches are unlikely to provide useful predictions of operational losses. A mixture of qualitative and quantitative methods is perhaps needed to model operational risks.

The OpRisk problem is not peculiar to the financial sector and operational risk is not a new topic. In his book James Reason argues that operational risk is faced by all organizations and he uses examples from the financial, rail transport, civil aviation, and nuclear power sectors to support his case (Reason, 1997). Reason identifies a host of reasons why catastrophic failures occur in these safety-critical industries, including (but not restricted to) a failure to enforce lessons learnt from previous failures, slow degradation or collapse of safety procedures, changes in culture and management, lack of visibility and

¹ Queen Mary, University of London, Computer Science, London, UK.

² Agena Ltd., 32-33 Hatton Garden, London EC1N 8DL, UK.

* Address correspondence to Martin Neil, Queen Mary, University of London, Computer Science, Mile End Road, London E1 4NS, UK; tel: 44-20-686-7882; fax: 44-20-8980-6533; martin@dcsl.qmul.ac.uk.

support for risk reporting, and lack of attention to detail. The key conclusion from this is that accidents are not solely the result of human fallibility but are supported by organizational features that fail to defend against all-too-human mistakes, slips, and (in the case of fraud) malicious acts. From this we can conclude that OpRisk prediction is inextricably entwined with good management practice and that measurement of OpRisk can meaningfully be done only if the effectiveness of risk and controls processes is regularly assessed. This contrasts sharply with the view that modeling OpRisk simply involves the investigation of statistical phenomena.

By the same arguments financial catastrophes are not a “bolt out of the blue” nor are they inexplicable. The financial scandals such as Barings (Rawnsley, 1995) and the AIB (Wachtell *et al.*, 2002) were all the result of fraudulent activities building up over lengthy periods of time during which active management could have discovered and prevented them. Indeed, if caught early the events would not have been catastrophes at all. There is a tendency to see financial disasters as single “ultra high loss” events rather than aggregations of smaller losses accrued over a period of time. This is understandable given the fact that the losses have to be realized upon discovery, all at once. But this does not change the fact that such losses are accumulated daily and could be detected by good diligence, applied routinely. It is precisely this routine attention to good practice that, just as in safety-critical industries, prevents disasters from occurring. Any OpRisk scheme should, therefore, focus on detecting near misses and small losses on a monthly or quarterly basis before they become large losses and disasters.

In this article we argue that Bayesian networks (BNs) provide an attractive solution to the problems identified above. BNs enable us to combine any statistical data that are available with qualitative data and subjective judgments about the process. Hence BNs provide a method of modeling operational losses and measuring the effectiveness of a business’s operational processes, as part of a self-assessment-oriented “Bayesian Scorecard” approach. Using BNs we can

1. combine proactive loss indicators, related to the business process, with reactive outcome measures such as near miss and loss data;
2. incorporate expert judgments about the contribution qualitative estimates can make to the overall risk assessment;

3. enter incomplete evidence and still obtain predictions;
4. perform powerful “what-if?” analysis to test sensitivity of conclusions;
5. obtain a visual reasoning tool and a major documentation aid;
6. obtain output in the form of *verifiable* predictions against actual performance measures and loss event rates.

In Section 2 we provide a brief overview of BNs. In Section 3 we consider the widely accepted distinction between expected and unexpected losses in OpRisk, whereas in Section 4 we explain how BNs provide a unified method of predicting both types of losses. We believe this represents a significant improvement over existing approaches since the distinction is, in our view, arbitrary. We concentrate on the core problem of predicting losses using two BNs to show how they can be used to model loss event frequency, severity, and heavy-tailed distributions. Section 5 discusses the issues relating to prior estimation in BNs and how prior beliefs can be informed by evidence from the operational process. Finally, in Section 6 we offer some conclusions.

2. BAYESIAN NETWORKS

The underlying theory of BNs combines Bayesian probability theory and the notion of conditional independence to represent dependencies among variables (Pearl, 1986; Spiegelhalter & Cowell, 1992). To date BNs have proven useful in many areas of application such as medical expert systems, diagnosis of failures, pattern matching, speech recognition, and, more relevantly for the OpRisk community, risk assessment of complex systems in high-stakes environments (Fenton *et al.*, 2004; Neil *et al.*, 2001, 2003).

BNs enable reasoning under uncertainty and combine the advantages of an intuitive visual representation with a sound mathematical basis in Bayesian probability. With BNs, it is possible to articulate expert beliefs about the dependencies among different variables and to propagate consistently the impact of evidence on the probabilities of uncertain outcomes. BNs allow an injection of scientific rigor when the probability distributions associated with individual nodes are simply “expert opinions.” This can both increase the reliability of the expert opinions, while also making explicit the imprecision that is inherent in such judgments.

A BN is a directed graph whose nodes represent the (discrete) uncertain variables of interest and whose edges are the causal or influential links between the variables. Associated with each node is a node probability table (NPT). This is a set of conditional probability values that model the uncertain relationship between the node and its parents together with any uncertainty that is present in that relationship.

The key to the successful design of BNs is the meaningful decomposition of a problem domain into a set of causal or conditional *propositions* about the domain. Rather than ask an expert for the full joint probability distribution, which is obviously a very difficult task, we can apply a “divide and conquer” approach and ask for partial specifications of the model that are themselves meaningful in the experts’ domain. Once we have achieved this decomposition we have also, implicitly as a natural product of the approach, specified the covariance by virtue of the conditional probability structure.

Next, we require the expert to model the NPT for each variable (node): this can either be done using historical data (with standard Bayesian parameter learning approaches or Monte Carlo simulations), or by simply asking the expert to provide a series of subjective estimates. Ideally, we would expect these estimates to be based on experience and knowledge rather than blind guesswork.

We can easily embed continuous and discrete statistical distributions within the BN model, as NPTs, and generate values for these NPTs by Monte Carlo simulation methods. For continuous functions we have to discretize the model appropriately and in the AgenaRisk software tool (AgenaRisk, 2005) this is achieved using a substantially enhanced version of the dynamic discretization algorithm presented in Kozlov and Koller (1997) and that allows the approximate solution of classical Bayesian statistical problems, involving continuous variables, as well as hybrid problems involving both discrete and continuous variables.

Once a BN is built it can be executed using an appropriate propagation algorithm, such as the junction tree algorithm (Jensen, 1996). This involves calculating the joint probability table for the model from the BN’s conditional probability structure in a computationally efficient manner. To do this an intermediate, graph-theoretic representation of the BN, called the junction tree (JT), is automatically derived from the BN. The JT allows localized, modular computations to be executed using a message-passing algorithm. This is, in essence, an elaborate form of Bayes’s *theorem* (for full details, see Jensen, 1996; Lauritzen &

Spiegelhalter, 1988; Pearl, 1986; or Spiegelhalter & Cowell, 1992). This process is entirely automatic and, in a tool like AgenaRisk, is hidden from the domain expert.

Once a BN has been compiled it can be executed dynamically, and exhibits the following two key features:

1. The effects of observations entered into one or more nodes can be propagated throughout the BN, in any direction, and the marginal distributions of all nodes are updated.
2. Only *relevant* inferences can be made in the BN. The BN uses the conditional dependency structure and the current knowledge base to determine those inferences that are valid.

It is worth noting that the computational weight involved in using BNs is manageable in terms of computer memory and permanent storage space, if using an efficient implementation of the JT algorithm. However, many academic, open-source, and off-the-shelf packages do not offer implementations that are efficient enough to support large BN models, especially when combined with Monte Carlo simulation. But efficient implementations are possible for the class of BNs needed to model OpRisk problems and can be easily built in tools such as AgenaRisk (AgenaRisk, 2005).

3. ESTIMATING EXPECTED AND UNEXPECTED LOSSES

The Basel report (Basel, 2004) classifies financial losses due to operational factors into two “types”:

1. Expected losses—These are considered the “normal” losses that occur frequently, as part of everyday business, with a low severity. Examples include losses due to accidentally miscalculated foreign exchange transactions.
2. Unexpected losses—These are the unusual losses that occur rarely and have a high severity. Examples include losses resulting from a major fraud activity.

Fig. 1 shows the distinction between expected and unexpected losses. The demarcation line is purely arbitrary (in Fig. 1 this separation is shown at total losses of \$400,000). It therefore makes little sense to use fundamentally different methods for predicting expected and unexpected losses; it is better to think in terms of finding a distribution whose *tail* represents the unexpected losses.

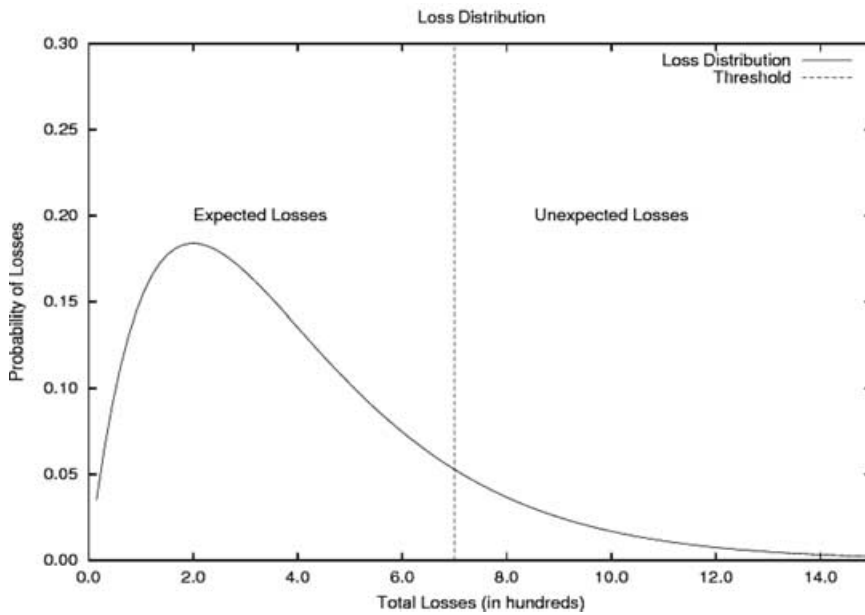


Fig. 1. Expected vs. unexpected total losses.

The traditional approach to these kinds of problems is to rely purely on historical data to find the predicted distribution. Where extensive data exist traditional statistical modeling techniques work well. However, in the case of operational loss data, we have a number of special problems. Most importantly, even when a lot of loss data are available, it is unlikely that there will be enough data on the large “unexpected losses” for us to be able to estimate the tail of the distribution properly—usually we end up with tails that are too “thin” or indeed “too fat” if the loss data are not relevant for the domain in question. Even when modeling the “expected losses” (the bulk of the distribution) the data-oriented approach suffers from the following limitations:

1. Loss data will be gathered over a period of time that may represent varying levels of operational effectiveness and risk/threat level. We cannot expect that losses are generated from one single distribution with a small number of “known” parameters.
2. Losses experienced are simply a sample of possible events. They may not be representative of changing operational processes. As the underlying operational process degrades or improves, the value of such historical data wanes.
3. The reported loss data might be wrong. Underreporting and data ambiguity can lead to significant errors in estimation.

4. Any attempt to bolster loss data with data gathered from other organizations is subject to the same problems and more because very often the provenance of the data is unknown or in doubt.

4. USING BNs TO PREDICT LOSSES

Given the serious limitations of the approaches based purely on historical loss data, it is inevitable that we will have to use methods that enable us to incorporate other types of evidence. Where different types of evidence need to be combined, classical statistical methods do not work. Bayesian methods, and in particular BNs, do provide a way forward since they can offer the following specific benefits in OpRisk:

1. Explicit combination of objective and subjective data by modeling the connection between the operational environment and the loss event process.
2. Can model “long” tail distributions for the unexpected loss component of the total loss distribution.
3. A method for eliciting subjective components of risk forecast from experts by explicitly modeling scenarios involving different operational processes or threats to the business, with likely outcomes.
4. A verifiable means for dealing with expertise such that models can then be used

independently of the expert in the same way a medical expert system would support less qualified practitioners by making “expert” advice available.

In this section we describe two BN models for predicting operational losses. The first predicts total losses from event frequency and severity assuming that these are independent. The second assumes dependence between frequency and severity. These models are by necessity simple and are presented mainly to dissuade readers of some of the misconceptions about how BNs might be used in OpRisk and give a glimpse of the potential of BNs in this area. It should be noted that the calculations here are presented from first principles—in practice all such calculations are performed by special purpose BN software tools and so their complexity is completely hidden.

4.1. Predicting Total Losses from Event Frequency and Severity

We can estimate the total loss distribution from the convolution of the loss frequency and severity distributions where total losses, T , are conditioned on the frequency of loss, F , and the severity of loss should it occur. Loss event frequency and severity are random variables, each with an appropriate probability density function (pdf). For a given loss event frequency, F , and severity distribution, S , we wish to predict the total loss distribution, T . The joint probability distribution

$p(F, S, T)$ is $p(T | F, S) p(S) p(F)$ and the total loss distribution is calculated by marginalizing S and F thus,

$$p(T) = \sum_{S, F} p(T | S, F) p(S) p(F),$$

$p(F, S, T)$ can be depicted graphically by a BN as shown in Fig. 2.

Given that BNs accommodate the use of Monte Carlo methods to generate the probability tables we do not need to restrict our model to any given family of conjugate probability distributions. For simplicity the prior pdf for event frequency might be represented best by single parameter distributions. Here we use a Poisson distribution, with rate parameter, λ , and an exponential distribution, with parameter, θ , as the prior distribution for severity, S ; thus,

$$p(F) = \frac{e^{-\lambda} \lambda^f}{f!},$$

$$p(S) \approx f(S) = \theta e^{-\theta S}.$$

Using the AgenaRisk software we can generate $p(T | F, S)$ by sampling from S and F using Monte Carlo methods and calculating total losses, T , for each combination of F and S sampled.

Once the BN has been specified and the NPTs generated we can calculate the marginal probability of any node in the BN by invoking the propagation algorithm.

Given our assumptions we simply need single-parameter estimates for the pair (λ, θ) to generate

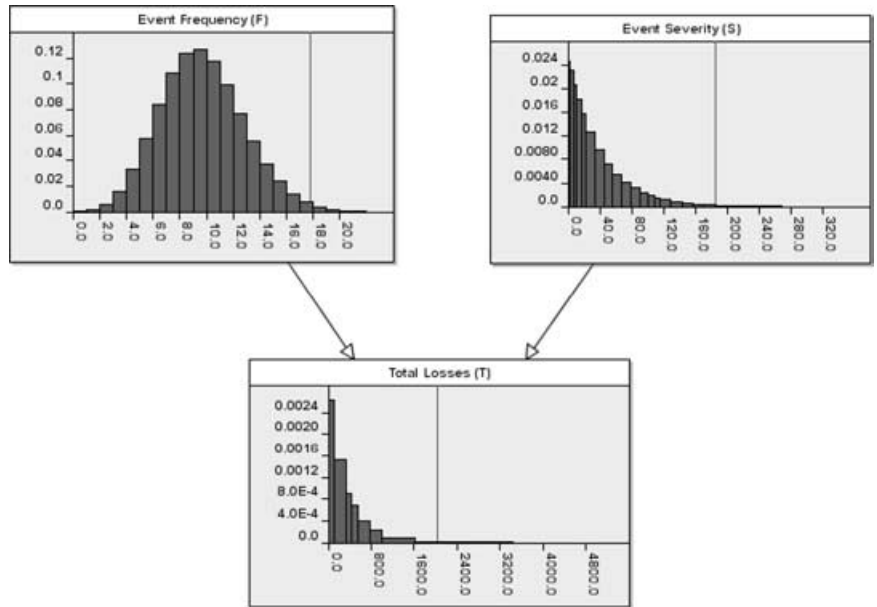


Fig. 2. BN for $p(F, S, T)$ showing posterior marginal distributions for each variable.

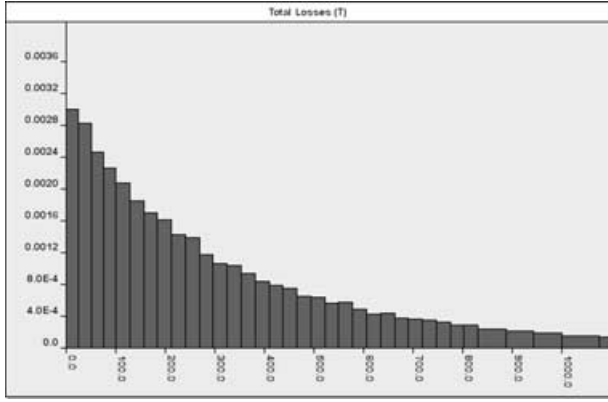


Fig. 3. Posterior marginal distribution for $p(T)$ where $T < 1,000$.

$p(F, S, T)$. Let us assume that, from past experience elicited from discussion with an expert, the mean loss event rate per year is approximately 10, $\lambda = 9.7$, and that the mean loss severity is \$40,000 per loss event, $\theta = 40$. Then the resulting posterior distribution for $p(T)$, as calculated by the BN, is as shown in Fig. 3.

4.2. Modeling Dependence Between Event Frequency and Severity

The model above assumes that loss event frequency and severity are independent of one another. This is optimistic—in reality we can expect them to covary at least for some classes of events such as fraud, where one might expect that a poor controls process encourages a fraudster to steal more than an effective controls process. For other classes of events the dependency might be weaker.

We can easily model covariance between severity, S , and frequency, F , by introducing a common cause, which we will name *process effectiveness*, E , into the BN model.

The new joint probability distribution $p(F, S, T, E)$ is

$$p(F, S, T, E) = p(T | F, S) p(S | E) p(F | E) p(E).$$

This is shown graphically by the BN in Fig. 4.

We can compare this “dependence” model, conditioned on E , with the independent version discussed previously by directly comparing the posterior

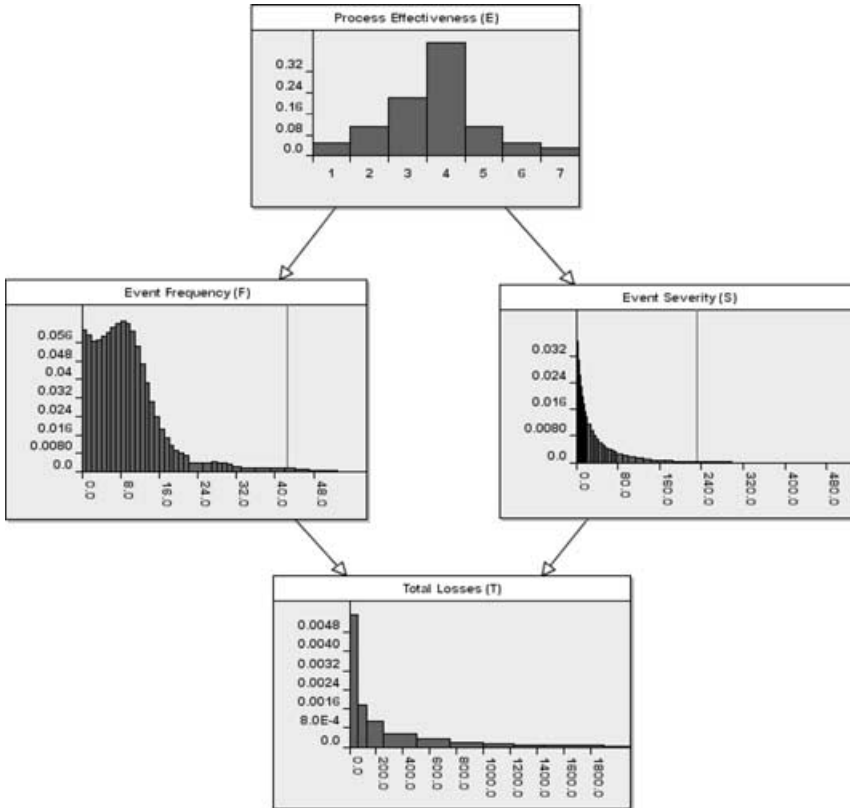


Fig. 4. BN for $p(F, S, T, E)$ showing posterior marginal distributions for each variable, where F and S are conditioned on E .

Table I. NPTs for $p(E)$, $p(F|E)$, and $p(S|E)$

E	$p(E)$	$p(F E)$	$p(S E)$
1	0.05	Poisson (0.5)	Exp (5)
2	0.11	Poisson (2)	Exp (10)
3	0.22	Poisson (5)	Exp (20)
4	0.43	Poisson (10)	Exp (50)
5	0.11	Poisson (15)	Exp (60)
6	0.05	Poisson (25)	Exp (70)
7	0.03	Poisson (40)	Exp (80)

marginal distributions for T in the dependent and independent cases: $\sum_{F,S} p(T|F,S)p(S)p(F)$ and $\sum_{F,S,E} p(T|F,S)p(S|E)p(F|E)p(E)$.

To best illustrate the differences between the models we can construct a model with mean values for $p(S)$ and $p(F)$ that are very close to the original independence model. The example mixture of NPTs chosen is shown in Table I.

The expectations for event frequency, F , and severity, S , are easily derived from the BN. These are almost identical to those used when F and S are independent: $E(F) = 10$ and $E(S) = 40$. However, the key issue relates to the differences in the tail of $p(T)$, or to put it as a question: Are the “unexpected” losses larger when F and S covary?

Fig. 5 shows the tail probability density functions for values of $T > 1,000$ and Table II shows a compari-

son of the mean and 99th percentile (we might assign the 99th percentile as the value at risk (VaR) measure) statistics, when F and S are independent and dependent, respectively. We can see that when F and S are dependent we get a “longer” tail than when they are independent. This difference in tails is further revealed in the difference in the 99th percentile values for T : under dependence the value is 4,937 and under independence it is 1,933. Therefore, under independence the VaR measure will be optimistic.

This model is by necessity simple but in practice the approach scales up and offers a realistic level of resolution in two ways. First, by using AgenaRisk, the user can specify a level of desired accuracy, constrained only by the computational resources available, and thus calculate very accurate percentile values for VaR and aggregated loss distributions over many processes. Second, AgenaRisk also supports an object-oriented notation to compose large models from smaller fragments and we can use this to automatically generate models from databases containing loss data and process descriptions.

5. MODELING PROCESS EFFECTIVENESS

Estimating the prior distribution for process effectiveness, E , is very important if we are to arrive at sensible estimates for the total loss distribution, $p(T)$.

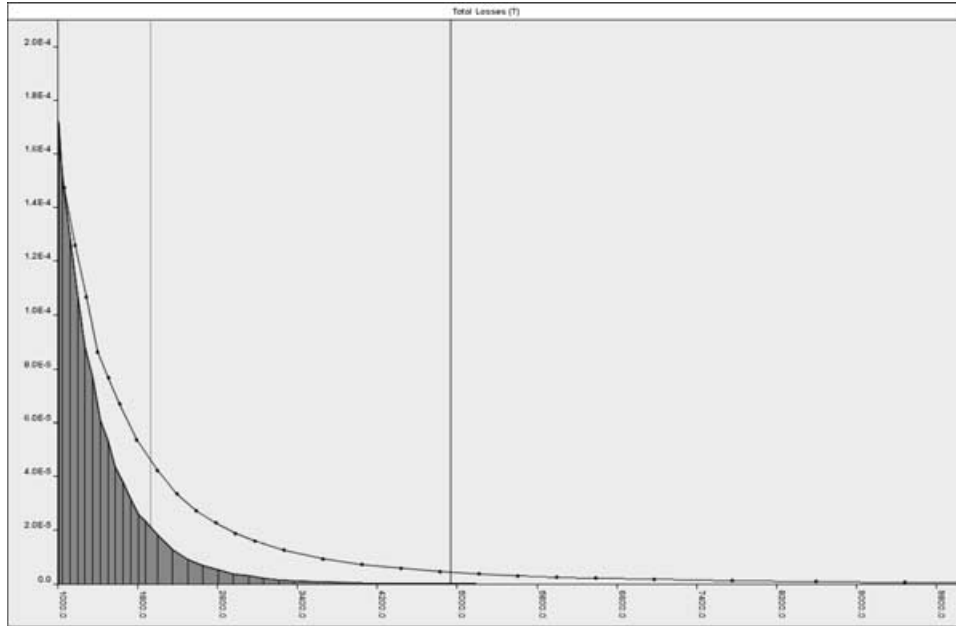


Fig. 5. Total losses tail distribution comparison. The dark gray area shows the independent case and the black line shows the dependent case.

Table II. Mean and 99th Percentile Statistics for F , S , and T Under Dependence and Independence Assumptions

		F and S	
		Independent	Dependent
Frequency, F	mean	9.2	9.2
	99th percentile	17.8	42.0
Severity, S	mean	40.0	40.0
	99th percentile	184.0	231.0
Total losses, T	mean	371.0	514.0
	99th percentile	1,933.0	4,937.0

The reason it makes sense to estimate $p(E)$ rather than simply estimate $p(T)$ directly is because $p(T)$ is a compound measure of event frequency and severity; experts find it hard to perform the mental calculations to combine these directly. They find it much easier to break it down into separate assessments of severity and frequency. Moreover, experts have direct experience of the process; they are involved in it, they have an intimate understanding of the controls, procedures, staff, and threats the business may face. We can exploit this direct experience to elicit information about the process separately from the outcomes of the process (i.e., event severity and frequency).

Thus, to populate the model we need to assess $p(S|E)$, $p(F|E)$, and $p(E)$. We can interpret E as the level of maturity, or effectiveness, in preventing undesirable events. The measurement scale for E given in Table I assigns “one” for the most effective process and “seven” for the least effective. We could, of course, label these differently and give detailed descriptions of the mix of operating procedures, technology, staff skills, etc. assigned to each level and doing so would make process effectiveness an observable quality.

Now we have some model for E ; we need to consider what the prior probability distribution actually means in Table I. If we examine $p(E = 1) = 0.05$ we can interpret the subjective beliefs (probabilities) of the expert in a number of ways:

1. $p(E)$ reflects uncertainty about the current process—the chance of *this* process at this moment in time being a Level 1 process is 1 in 20.
2. $p(E)$ reflects the frequency of occurrence of a changing process—over 20 years we would expect our *changing* process to be equivalent to Level 1 (the best), on average, once.

The first interpretation echoes the Bayesian definition of probability, as a degree of belief, and the second echoes the frequentist definition. Both of

these interpretations are mathematically equivalent but have subtly different semantics. The first implies that we never really know how “good” the process actually is (we are in doubt) and the second suggests that the process changes or evolves over time in some way, i.e., there is some long-run frequency at which the process occurs at a particular level. Both interpretations are correct but an expert may find it easier to think in subjective rather than frequency terms or vice versa.

The next possible concern about E is that the experts may not have experienced the whole range of E directly, and will thus be unable to make any confident statement about $p(S|E)$, $p(F|E)$, and $p(E)$ over this range. However, experts are very good at anticipating and then recommending or taking action to avoid unpleasant outcomes (this is why we employ humans after all!). Indeed, the whole thrust of regulation is to force businesses to anticipate and act in this way. With this in mind we can generate, with the help of the expert, hypothetical scenarios involving E and asking:

1. What is the chance of this level of degradation or improvement occurring in the process, $p(E)$?
2. Given that the process effectiveness is at a particular level, what would the event severity, $p(S|E)$, and frequency, $p(F|E)$, distributions look like?

In our experience we have found that experts are willing and able to answer these two questions if they are asked in a structured fashion *and* they can quickly see the results of their assertions in the BN and then refine it. In partnership with local domain experts we have successfully built large-scale models exploiting expertise and statistical reasoning for safety-critical and mission-critical applications in vehicle reliability, air traffic control, software risk prediction, and warranty return prediction of electronics components (Neil *et al.*, 2001, 2003).

Answering these questions may be difficult if they do not appear grounded enough. Under these circumstances we could extend the BN model to identify causal factors related to E , such as that shown in Fig. 6, where we have identified another layer of causes in the form of three causal factors: operational procedure quality, O ; technology fidelity, TF ; and staff quality, Q . Collecting priors on these might then be easier for the expert.

From a statistical modeling perspective the probability distribution for process effectiveness, E , simply acts as a mixture parameter that mixes a set of



Fig. 6. Refinement of BN model to include another layer of causes.

statistical or empirical frequency and severity distribution models that will result in a “longer” tail and a more realistic VaR estimate (Venkataraman, 1997). We obviously favor a causal interpretation of *E* rather than a purely statistical interpretation in order to focus expert attention on the underlying generative process and hopefully to generate a more sensible Bayesian model.

6. CONCLUSIONS

BNs can help combine qualitative data from experts and quantitative data from historical loss databases in a principled way and as such they go some way in meeting the requirements of the draft Basel Accord (Basel, 2004) for an advanced measurement approach (AMA). Adopting a BN-based approach should, therefore, lead to better operational risk governance and a reduced regulatory capital charge. Relying purely on historical loss data and traditional statistical analysis techniques will neither provide good predictions of future operational risk losses, nor a mechanism for controlling and monitoring such losses.

We have shown how BNs can be used to model operational risk via two small examples in which total losses are based on event frequency and severity. In the second model we took account of the possible dependence between frequency and severity by introducing a common cause *process effectiveness*, *E*, and we showed that we could use this BN to model “heavy

tailed” distributions in a way that would exploit the expertise available within an organization. This helps produce a realistic and reproducible VaR estimate. We call the method used to construct the total loss distribution the “Bayesian Scorecard” because of the obvious similarities to the less sophisticated scorecard measurement process.

BNs focus on assessing the effectiveness of the underlying business process and we propose using them as a form of self-assessment. This would involve monitoring the underlying business process on a frequent basis (such as quarterly or monthly) and translating these self-assessment scores into total loss predictions via the BN.

We could go further and automatically *learn* from loss data, as it is observed, and self-assessment data *together* as part of a dynamic approach. Similarly, entering hypothetical self-assessments and notional loss data into the BN model would then easily support “What-if” and sensitivity analysis. Such analysis would then help assess the accuracy of and help recalibrate their expertise over time. These topics, and others, will be covered in subsequent articles.

In this article we have shown that we can model uncertainty about the process that generates losses as well as the distribution of losses that might result. This addresses only a small part of the overall problem and clearly further work is needed to extend the modeling approach to address these outstanding issues:

1. There will be a time series relationship between losses and the dynamic nature of business processes.
2. The reported loss data might be wrong. Underreporting and data ambiguity can lead to significant errors in estimation.
3. Any attempt to bolster loss data with data gathered from other organizations is subject to the same problems and more because very often the provenance of the data is unknown or is in doubt. Classifying the type of organization and process that led to losses might allow us to weigh the contribution of data within a model.

The modeling approach described here has already been implemented to model the operational risk profile of a large South African national bank. While it is premature to claim validity or whether a regulator would find the approach compliant, early results from this implementation show that the work scales up well and also that the BN formalisms match

the modeling requirements of business units for transparency and practicality.

ACKNOWLEDGMENTS

This article is based in part on work undertaken on the following funded research projects: SCULLY (EPSRC Project GR/N00258), SIMP (EPSRC Systems Integration Initiative Programme Project GR/N39234), and SCORE (EPSRC Project Critical Systems Programme GR/R24197/01). The AgenaRisk software (AgenaRisk, 2005) was used to create the example analysis and graphs in this article. We would like to thank the referees for their helpful comments and suggestions.

REFERENCES

- AgenaRisk. (2005). Available at www.agenarisk.com.
- Basel Committee on Banking Supervision. (2004). *Working Paper on the Regulatory Treatment of Operational Risks*. Available at <http://www.bis.org>.
- Fenton, N. E., Marsh, W., Neil, M., Cates, P., Forey, S., & Tailor, M. (2004). Making resource decisions for software projects. Presented at Proceedings of 26th International Conference on Software Engineering (ICSE 2004), Edinburgh, United Kingdom, IEEE Computer Society 2004.
- Jensen, F. V. (1996). *An Introduction to Bayesian Networks*. London: UCL Press.
- Kozlov, A. V., & Koller, D. (1997). Nonuniform dynamic discretization in hybrid networks. In *Proceedings of the 13th Conference on Uncertainty in Artificial Intelligence (UAI-97)* August 1–3, pp. 314–325.
- Lauritzen, S. L., & Spiegelhalter, D. J. (1988). Local computations with probabilities on graphical structures and their application to expert systems (with discussion). *Journal of the Royal Statistical Society Series B*, 50(2), 157–224.
- Neil, M., Fenton, N. E., Forey, S., & Harris, R. (2001). Using Bayesian belief networks to predict the reliability of military vehicles. *IEE Computing and Control Engineering*, 12(1), 11–20.
- Neil, M., Malcolm, B., & Shaw, R. (2003). Modelling an air traffic control environment using Bayesian belief networks. Presented at the 21st International System Safety Conference, August 4–8, 2003, Ottawa, Ontario, Canada.
- Pearl, J. (1986). Fusion, propagation, and structuring in belief networks. *Artificial Intelligence*, 29, 241–288.
- Rawnsley, J. (1995). *Going for Broke: Nick Leeson and the Collapse of Barings Bank*. London: HarperCollins.
- Reason, J. (1997). *Managing the Risks of Organisational Accidents*. Aldershot, UK: Ashgate Publishing.
- Spiegelhalter, D. J., & Cowell, R. G. (1992). Learning in probabilistic expert systems. In *Bayesian Statistics* (pp. 447–465). Oxford: Oxford University Press.
- Venkataraman, S. S. (1997). Value at risk for a mixture of normal distributions: The use of quasi-Bayesian estimation techniques. In *Economic Perspectives*. Federal Reserve Bank of Chicago.
- Wachtell, Lipton, Rosen, Katz, & Promontory Financial Group. (2002). *Report to the Boards of Allied Irish Banks, p.l.c., All first Financial Inc. and All first Bank Concerning Currency Trading Losses March 14*. New York.