

Adapting Banking Models to Insurer ERM

Gary G. Venter

Presented at
Enterprise Risk Management Symposium
Society of Actuaries

Chicago, IL

April 23–26, 2006

Copyright 2006 by the Society of Actuaries.

All rights reserved by the Society of Actuaries. Permission is granted to make brief excerpts for a published review. Permission is also granted to make limited numbers of copies of items in this monograph for personal, internal, classroom or other instructional use, on condition that the foregoing copyright notice is used so as to give reasonable notice of the Society's copyright. This consent for free limited copying without prior consent of the Society does not extend to making copies for general distribution, for advertising or promotional purposes, for inclusion in new collective works or for resale.

Abstract

Insurance company issues that do not necessarily arise in banks are identified and utilized to develop insurer ERM modeling approaches that start from the work done for banks but respond to insurance-specific matters. Some of the approaches discussed could also be used to refine banking models.

1. Introduction

Banking enterprise risk management (ERM) models provide a framework for financially consistent risk assessment, and serve as useful examples for insurance companies' ERM modeling. Some of the methods used in the bank models are reviewed here and directions that might improve their application to insurer modeling are suggested.

Central to the banking approach has been having a comprehensive enterprise-wide risk measure (sometimes called "economic capital"), identifying the types of risk that contribute to that (asset risk, credit risk, etc.), allocating the overall risk to operating units and using the allocated risk in performance measurement and evaluation of strategic alternatives. This approach can provide the discipline to identify and try to quantify the key sources of risk a company faces, and to reach across disparate divisions to get a common risk framework and language and a consistent decision-making process.

Insurers like banks are financial institutions facing risk, but there are some important differences that can complicate modeling. Bank risk is probably more subject to standardization. For instance, standardized home loans are packaged and traded on financial markets, but this has proved to be difficult for books of automobile and homeowners' insurance policies. While both face financial risk, the basic business of banks is financing, while that of insurers is risk. This has traditionally made insurers more invested in risk modeling and perhaps more concerned about the details of capturing the risk, and more willing to use more detailed modeling approaches. Thus while banking models make a good starting point, insurers will often want to extend such models further.

Possible extensions include alternative risk measures, allocation of risk, correlation, modeling specific to insurance risk such as pricing and reserving, and models for particular risk categories like asset risk and operational risk.

2. Risk Measures

ERM modeling can give a comprehensive picture of company risk by using a common risk measure for different sources of risk, which tend to be managed by different departments.

2.1 Economic Capital

Banking derived models often focus on a single risk measure, scale it in monetary units, and call it “economic capital.” For instance the 99.95th percentile of the negative of annual profits might be the chosen risk measure. This is typically positive since there is usually more than a 1/2000 probability of incurring an annual loss. Thus it can be interpreted as the amount of capital needed to have a probability of 99.95 percent of staying solvent. That might seem like a reasonable target, so that amount of capital could reasonably be identified as economic capital.

Usually the risk measure is selected so that economic capital is a bit less than actual capital. Unless the company is planning to raise more capital, it is awkward to say it has economic capital of \$10 billion and actual capital of \$9 billion, even though that is allowed by the terminology. In practice the rule for selecting economic capital often appears to be to set it at the largest round percentile of negative annual profit that is below actual capital. For well financed companies, the percentiles that would allow insolvency probabilities of 1/2000 or 1/3333 are typical, while for weaker companies this might be 1/1000 or more.

This terminology could be misleading for insurers, however. Insurer management is accustomed to thinking in terms of risk measures, so selecting one and calling it economic capital can obfuscate the fact that the discussion is about a risk measure. It also creates a somewhat artificial distinction with actual capital. If the actual probability of insolvency is 1/3481, for example, singling out the capital that would increase it to 1/3333 does not seem to be particularly useful.

It is useful to have risk measures, however. In fact, insurers tend to want to see an analysis with more than one measure of risk. This is in part to see if conclusions are consistent with different risk measures, and in part because there is as yet no consensus on what risk measure is most representative of optimum strategy.

2.2 Tail-Based Risk Measures

The probability of insolvency is a long-standing risk measure in insurance literature, and the capital required to prevent insolvency at a selected probability level is closely related. Both of these are tail-based risk measures, that is, risk measures that look only at the adverse scenarios. In contrast, the standard deviation is a common risk measure whose calculation requires the entire distribution of probabilities of outcomes. Because capital is needed to respond to adverse outcomes, it is natural to view risk as a tail phenomenon. However, there are other risk measures discussed below that are also

worth considering.

Probably the most commonly referenced tail measure is value-at-risk, or VaR. This is just a selected percentile of a distribution. For example, the 99.95th percentile of the negative of profit would be the VaR at probability level $= 0.9995$ for the random variable negative of profit. This is what was being called economic capital in the example above.

While popular and sometimes easy to calculate, VaR has serious limitations as a risk measure. Since it focuses entirely on a single probability level, it does not give much information about the distribution, or even the possibilities for other adverse outcomes. No probability level stands out as being particularly relevant, so the selection of level is arbitrary. It is also difficult to analyze VaR into component contributions from risk sources or business units. A loss at a selected probability level might have a completely different breakout into components than a loss at a nearby probability level from a simulation model. Also VaR is often misrepresented as a return time. That is the VaR at 90 percent is said to be the loss that recurs at a 10-year period. But if so, after 10 of those periods, that loss would recur 10 times, and the 1-in-100 loss would never happen.

For probability levels in the vicinity of insolvency, like 1/3333, VaR is actually quite difficult to measure for a large insurer. The individual events that could contribute to such an annual loss are difficult to quantify in themselves, and this is compounded by aggregation of several events—large loss reserve movements, management improprieties, etc.—that all become more critical in the extreme cases. It appears to be well beyond the current state of the art of modeling to know the insolvency probability of a large insurer with much accuracy.

It is more reasonable to quantify impairment probabilities, such as the probability of losing more than 1/5th of capital in one year. This could be done for a few impairment levels, to see what percent of capital could be lost say with probability 25 percent, 10 percent, 4 percent, etc. These could even be turned around to express capital as multiples of losses at different probability levels. E.g., capital could be 20 times the 90th percentile loss and 5 times the 96th percentile loss. This in effect scales the impairment percentiles to the level of capital, so they are comparable to capital needed to prevent insolvency, but much more quantifiable. These could be called economic capital A, economic capital B, etc., but it is more straightforward to just identify them as different risk measures.

Some of the shortcomings of VaR are overcome by tail value at risk (TVaR), also called conditional tail expectation (CTE). This is defined at a probability level α as the average loss when the loss is above that probability level.¹ E.g., for $\alpha = 0.9$, TVaR at 90 percent is the average loss beyond the 90th percentile. This does not look just at a single percentile of the loss distribution, but in fact looks at all loss scenarios beyond that level. It is more meaningfully analyzed into components, in that changing the percentile slightly will not change most of the contributions to TVaR. It also can be reasonably interpreted as the loss at a certain return time. If you think that every 10 years you get a random draw of the losses beyond the 90th percentile, then one in 10 of those will be the 99th percentile loss, on the average. So the 1-in-10 loss can be represented as the average of the loss beyond the 90th percentile, which is the TVaR at 90 percent.

TVaR is thus a good way to represent impairment probabilities. Capital can be expressed as multiples of TVaR at a few probability levels, which give risk measures scaled to capital but within the ability of models to quantify reasonably, and which can be interpreted as multiples of capital lost at various return times.

TVaR still has the problem of arbitrary selection of probability level, however. This is one reason for selecting several levels. One meaningful level is the probability of a drop in surplus. TVaR at this level is the average drop in surplus when there is a drop in surplus. Expressing capital as a multiple of this tells you the percentage of capital lost on the average when capital is lost.

Another problem of TVaR is that it is linear in large losses. This is not usually thought to be appropriate for risk measures. E.g., a loss twice as large is usually regarded as more than twice as bad, unless you are risk-neutral. A way around this is to use weighted TVaR (WTVaR). This is TVaR calculated with adjusted probabilities that make the worst cases more likely than they actually are. Even the ground-up mean under transformed probabilities might make a reasonable risk measure. This is discussed further below when alternatives to tail-based measures are considered. It might be possible to find a probability transform whose mean represents the market value of the risk undertaken. In that case WTVaR would represent the value of the tail risk.

Sometimes ignoring losses above the insolvency threshold is recommended when calculating TVaR, under the reasoning that shareholders do not care how big the loss is once the company is ruined. On the other hand, the policyholders and regulators would care about this, and they could probably make it worthwhile for the shareholders to care about it as well.

¹ Often the average at or beyond that level is used, but for a normal insurance company that would be the same thing.

The most popular tail-based measure in academic literature is the insolvency put option. This is the market value of the right of the shareholders not to pay once capital is exhausted. It reflects both the probability and severity of insolvency, and the market response to this. The put is typically valued from the capital market perspective, not the insurance buyer perspective, and is regarded as part of the market value of the insurer. There is evidence that the insurance buyer demands a discount in insurance pricing well beyond the capital market value of this option. This makes sense when you think of the buyers as non-diversified in insurance purchases, and the shareholders as diversified in shares owned.

The difficulties of modeling the probability of insolvency extend to modeling the value of the insolvency put. Perhaps some degree of quantification is possible through observing security prices, but it does not appear to be possible through detailed modeling of the insurance risk at the current level of modeling expertise.

2.3 Transformed Probability Risk Measures

A way to consider the entire distribution of profit amounts but still give the greatest emphasis to adverse outcomes is to define a risk measure by the expected value of the profit under transformed probabilities. There are a number of probability transforms in the actuarial literature, such as the Esscher transform and the Wang transform. The latter has been shown (Wang 2004) to approximate market prices for both commercial bonds and catastrophe bonds with nearly the same parameters. The general idea is to make the probabilities of the bad events higher and the probabilities of the better scenarios lower. The transformed mean profit then decreases to something less than the actual mean profit. A candidate for a reasonable transform is one that makes the expected profit the risk-free rate. Then the transform could be considered to account for the risk in the portfolio.

In an ERM context, a consistent risk measure is desired for different sources of risk, like insurance losses and investment assets. This could possibly be done with the Wang transform with fixed parameters. Another way to get consistency is to use different transforms for different risk sources, but derive them all from the same rule, like minimizing the distance from the original probabilities while making the adjusted return the risk-free rate. Then the distance measure would provide the consistency. In the theory of incomplete markets, two popular distance measures are quadratic distance and information distance. Minimizing quadratic distance gives the minimum martingale transform, and minimizing information distance gives the minimum entropy martingale transform. The latter seems to be stronger in theory and in practice, e.g., see

Venter, Barnett and Owen (2004). This transform has in fact been shown to be a generalized version of the Esscher transform, and has been computed for many stochastic processes, including processes for asset prices and insurance losses.

For a compound Poisson process with frequency λ and severity distribution $g(y)$, the entropy transformed distributions, with parameter c , are:

$$Ee^{cy/EY}$$

$$g(y) = g(y)e^{cy/EY}/Ee^{cy/EY}.$$

This gives a risk loading factor on expected losses of $1 + \lambda = E[Ye^{cy/EY}]/EY$. This can be used to solve for c if the loading is known. If the c is allowed to differ by line of business, then by-line profit targets would be an input rather than an outcome of the modeling process. This may be realistic in practice. The consistency of risk measurement of losses and investments would again be achieved by also using the entropy transform of asset prices as the risk measure for assets.

Covariance with another variable can be expressed as a particular kind of transformed probability—see Venter (1991)—so the capital asset pricing model (CAPM) and related asset pricing models (as well as Black-Scholes) are special cases of this approach. Thus it should be possible to express the market value of a risk as its mean under a probability transform. That would be a good risk measure.

3. Allocation of Risk

Once the risk of the firm has been quantified with one or more risk measures, the contributions of the various business units to the overall risk are often sought. When the risk measure is called “economic capital” this process is called “allocation of economic capital.” However it is more like allocation of risk than allocation of capital. In fact, estimating the contributions of the business units to the overall risk is conceptually a bit different from allocation of the risk to business units. It is more like seeing where the risk comes from rather than sending it out. This is more like decomposition of the risk measure than allocation of it.

A typical method of allocating a risk measure is first calculating the risk measure separately on each business unit then spreading the company overall risk measure proportionally. Analyzing the risk measure into component contributions is also a two-step process. To do this the risk measure first has to be definable as an average of company results under certain conditions. Many but not all risk measures can be so defined. Then the contribution from each business unit is the average of the business unit results under the same conditions.

This can be easily done for a risk measure like TVaR, which is the average loss when the loss is above a selected probability level. Then each business unit's contribution to TVaR is the business unit's average loss when the company loss exceeds that threshold. Similarly, the components of VaR from the business units are the units' average losses when the company loss is at that probability level. Thus, under risk decomposition, VaR is additive—the business units' contributions to VaR add up to company VaR. However these contributions are unstable when simulation is used to compute the distribution, because there is only one simulation exactly at that probability level for each run of the simulation model.

The variance $E[(Y - EY)(Y - EY)]$ is the average squared deviation of a variable from its mean. Denoting Y , the total company negative profit, as the sum of the business units' negative profits $X_1 + \dots + X_n$, allows expressing the contribution of the j^{th} unit to the variance as $E[(X_j - EX_j)(Y - EY)]$. This is the definition of the covariance of X_j with Y , and these covariances add up over the X s to the variance. Because of this terminology, the contributions of business units to other risk measures as described above are called co-measures, like co-TVaR, co-VaR, etc.

One of the goals of decomposition of a risk measure is to be able to measure risk-adjusted profitability of business units. This works particularly well if the risk measure is proportional to the market value of the risk. Then dividing the profit of a unit by its contribution to the risk measure would give a constant times the ratio of the profit to the market value of the risk. Thus business units with higher ratios would have more profit relative to the value of the risk they are taking.

Although there are various theories of how to measure the market value of risk, at this point this is not a settled question. Thus it makes sense to compare several risk measures, in hopes market value will be close to proportional to one of them, and that the indicated strategic directions will not be too different among them.

A desirable feature of a risk decomposition methodology would be that it is a marginal method. That is, the change in overall company risk due to a small change in a business unit's volume should all be attributed to that business unit. This links to the standard financial theory of pricing proportionally to marginal cost. It also leads to consistent strategic implications. Under a marginal methodology, if a business unit with an above average ratio of profit to risk increases its volume, then the overall company ratio of profit to risk will increase.

So how can a marginal risk attribution method be developed? This is easier when the business units can change volume in a homogeneous fashion. An example would be

business units that buy quota share reinsurance and so can change their volume uniformly just by changing the quota share percentages. For such a company, if the risk measure is also homogeneous, then there is a marginal risk attribution method. The risk measure is homogeneous if changing the volume by a factor changes the risk measure by the same factor. Under these conditions, the marginal attribution is just the change in the company risk measure due to a small change in the volume of the business unit. Then this is a co-measure as well and the marginal attributions sum up to the company risk measure. This is a direct consequence of a theorem of Euler about derivatives of homogeneous functions. Common risk measures expressed in monetary units, like standard deviation, TVaR, etc., are homogeneous, but other measures are not, such as variance (square dollars) and probability of insolvency (unitless).

For many companies and business units, growth in exposure units can approximate homogeneous growth, so the same procedure would apply. However it can happen that exposure units come in large enough lumps compared to overall volume that adding one changes the shape of the distribution, so the changes in risk measure will not add up to the overall risk even for some homogeneous risk measures. However even in this case, transformed probability risk measures will still be marginal and additive.

Thus some transformed probability risk measure is likely to exist that is proportional to the market value of the risk being measured, and it will have an attribution to business unit that is marginal.

3.1 Capital Consumption

An alternative method of comparing profitability of business units was suggested by Merton and Perold (1993), and this has a lot in common with what Mango (2003) calls capital consumption. The general idea is to compute the value of the right of the business unit to call upon the capital of the firm, and check to see if the value of the profits the unit is generating exceeds this. If it does, the unit is adding value to the firm.

The value of the right to access capital and the value of the profit stream if positive can be computed using the theory of pricing of contingent claims. However these are not simple options. If the unit does require a capital call, its timing is not fixed in advance, and probably a sequence of cash flows would be needed over time. The timing of the realization of profit is not pre-determined either. The theory of pricing in incomplete markets could be applied, however. This would again involve a probability transform of the profit stochastic process.

4. Correlation

Risk sources, such as assets and insurance losses, and business units, such as workers' compensation and commercial property, are likely to have some degree of correlation of results. Historical results for establishing these correlations tend to be sparse at best, so modeling and judgment need to be used to come up with reasonable values. The worst company outcomes can be when several things go bad at once, so extreme results can be sensitive to correlation assumptions, and it is useful to test these sensitivities.

One approach to modeling correlation is to build up a correlation matrix across all the risk categories, and then simulate scenarios based on the matrix. One problem with this is that standard simulation approaches tend to build in assumptions that come down to simulating from the normal copula. This procedure in effect builds in the correlation structure of the multivariate normal distribution that has the assumed correlation matrix.

The main problem with this is that the multivariate normal, and so anything built up from the normal copula, is independent in the extremes. That is, even though there is correlation overall, the correlation among tail events is weaker, and tends to zero if you go far enough out into the extreme tail. This is not appropriate for many risks insurers face. For instance, catastrophe losses across lines of business tend to be most strongly correlated in extreme events, which is just the opposite behavior.

The t-copula is closely related to the normal copula, but has an extra parameter that allows inclusion of virtually any degree of extreme tail correlation. The simulation procedure is also similar, just with an extra step to include the extra parameter by application of common shocks that hit all correlated variables.

A modeler that is using the t-copula will usually say so. If the modeler refers to the Cholesky decomposition, or the @risk package, which use machinery needed for both the t and normal copulas, they may be using the t-copula. If they don't say so, they probably are just using the normal copula.

5. Modeling Insurance Risk

Insurance loss processes are somewhat dissimilar to what bank models have traditionally dealt with, although models of bank operational risk are starting to include losses due to fire, liability, fidelity, etc. However actuaries and such have been model-

ing these losses for centuries.² Typically aggregate loss models are built up from claim count and size distributions. A comprehensive review of such models is beyond the scope here, but a few key points are addressed.

5.1 Parameter Uncertainty

The basis of insurance is diversification of risk. However writing more independent exposure units does not diversify the risk of systematic underpricing. Not knowing the distributions exposed to could thus be a large part of the risk of running an insurance enterprise. When distributions are fit to historical data and then projected into the future, parameter risk arises from estimation risk and projection risk. These are the risks of getting the historical distributions wrong and of errors in projection of the historical data to future periods. Getting these wrong is not necessarily due to insufficient expertise or diligence—it is a built-in risk when fitting and projecting distributions. There are statistical methods available for quantifying such risks, and these should be used when modeling risk.

Catastrophe risk models are in a different category. Rather than modeling loss patterns themselves, these models use historical data on geophysical processes to model events that drive losses. Parameter uncertainty is more difficult to quantify here. Sensitivity testing by randomly changing assumptions can lead to very wide dispersion of estimates; it is usually thought that the parameters have to be harmonized by judgment to keep model results within reasonable ranges. Differences between subsequent models by a single vendor and among models of different vendors have sometimes been large. So-called “secondary uncertainty” within catastrophe models does not capture all of the parameter uncertainty issues. This is an area that requires continued attention in loss modeling.

5.2 Loss Reserve Adequacy

Modeling the distributional aspects of loss reserves is also a large topic in itself. The recent CAS working party white paper on this topic goes into a lot of detail. For strategic planning and risk quantification purposes, perhaps it is most reasonable to project the payouts and risk to payouts for each proposed business unit over the entire payment horizon. This gives a series of risk elements for each unit, and perhaps these can be discounted back to present to get a single risk measure. What this can miss, however, is the possibility of correlation with the runoff risk of existing liabilities. This runoff risk over time including the correlations is part of the company risk. Perhaps these

² For example, in 1693, Halley, of comet fame, created a mortality table based on records from Breslau.

could also be quantified as a series over time and also present-valued.

If business units have had comparable and fairly steady growth over time, the one-year risk of development of the entire book of reserves for each unit could be an approximation of the discounted risk of all future development for the unit's new underwriting year. These assumptions could be fairly imprecise, however, so the more reliable approach would usually be to model out the series of future risks for the underwriting year. In either approach, the risk of all future development of existing reserves is a separate but correlated issue.

The development risk is probably more acute on an incurred basis than a paid basis. If there is an acceleration of payments, or an increase in liability patterns, the entire book of reserves could be hit at once, even though payments are stretched out over time. Models need to be able to incorporate this risk, which can sometimes be understated by assumptions about independence of development fluctuations. Incorporating calendar-year effects (sometimes called "superimposed inflation") into the reserve model can help recognize this risk.

6. Other Risk Categories

Asset risk is somewhat different for insurers than banks due to different investment portfolios, so different modeling approaches may be appropriate. Operational risk is actually defined somewhat differently in these arenas, with much of insurance risk considered operational for banks. Finally, credit risk for insurers is chiefly reinsurance recoverable risk, while for banks it is loan default risk. These are linked with major catastrophe events and economic recessions, respectively, so again require different approaches.

6.1 Asset Risk

Insurers' assets are invested in bonds, stocks and various real-estate related instruments. Pretty good bond models are available, but modeling gets more difficult for the other items.

For government bonds, it is important to use arbitrage-free models. This is not because there are never any arbitrage opportunities in the bond market, as in fact there might be some from time to time. Rather it is important because having arbitrage strategies possible in a set of simulated scenarios will show these as the optimal strategies, and will show other strategies that are close to these as close to optimal. The trou-

ble with that is that arbitrage opportunities in the market do not last for very long, and will not still be there after running the model. Thus having them in the model will create inappropriate priorities.

A useful standard for evaluating bond models is their ability to capture actual aspects of the market. Examples of bond models and a testing methodology are discussed in Venter (2004). Some features of bond markets are:

1. Very high autocorrelation of interest rates
2. Higher volatility with higher interest rates
3. Long-term mean reversion
4. Short-term mean reversion to a temporary mean
5. Stochastic volatility
6. Lower yield spreads with higher short-term rates
7. Stochastic yield spreads

Bond models that capture these effects are known.

For shares, the models are still developing. Models like Black-Scholes assume geometric Brownian motion, but this is not realistic. Option prices, for instance, are higher for longer term and far out-of-the-money options, relative to shorter term and at-the-money options, than this model would predict. The next level of sophistication is Levy processes, which allow for jumps and heavier tails, and which can get the right relationship between at-the-money and out-of-the money option prices. It appears, however, that if these models are calibrated to short-term volatility levels they actually have too much volatility for longer terms. There might be some way to correct for this with mean reversion. Another possibility that shows some promise is regime-change models, which postulate a few or a continuum of geometric Brownian motion processes and a mechanism for shifting among them. Probably there are better proprietary models that never get published, as information is valuable in share trading. This is an evolving area and selecting the best model involves choosing which compromises to make.

Detailed models of real-estate instruments are even less developed in the public domain. There are also a variety of such instruments with subtly different risks. Perhaps at this point simply simulating from a selected distribution is all that can be done, even though this would miss correlations with other assets.

6.2 Operational Risk

Banks and insurers tend to classify risk categories somewhat differently. Insurers typically use a fourfold system:

- **Strategic Risk**—changes in customer priorities, shifts in brand power, new technology, legal & regulatory changes
- **Operational Risk**—succession planning, HR issues, governance, audit and control, product failure, supply chain, IT
- **Financial Risk**—volatility in interest rates, exchange rates, equity, credit risk and liquidity
- **Hazard Risk**—non-financial asset impairment, such as natural hazards, employee actions, legal liability, product recall and integrity and business interruption.

Banks tend to use a different breakout. For instance, operational risk is defined as “the risk of direct or indirect loss resulting from inadequate or failed internal processes, people and systems or from external events.” This at first appears to include just about anything, and in fact it is usually further qualified in the fine print to be somewhat narrower than it sounds here. Nonetheless, it usually does include things that insurance covers, like fire, theft, liability, employee dishonesty, etc., all of which would be classified as hazard risk in the insurance list.

Some aspects of operational risk in the insurance list can be quantified. Pension funding issues would probably be in this category, for example. There has been some quantification of IT failure risk, and perhaps some of the others as well. Other aspects of operational risk do not fit well into the paradigm of quantification and funding for risk, however. It is important to identify these risks, but then managing them with operational controls is probably more important than quantifying and funding for them.

For instance, what is the probability that the incentive compensation plan will lead to inappropriate managerial behavior and decision-making? This is a key issue in the agency theory approach to financial risk. Rather than quantifying and funding for this risk, studying the plan and understanding its incentives, and if necessary adjusting it, would be more useful. Or in the area of reputational risk, firms have been severely impaired due to damage to their reputation and image from off-hours behavior of key employees. It is not clear that adding extra capital would help in either of these examples. Thus for much of operational risk, the primary role of ERM would seem to be identification and management of such risks rather than quantification and funding for

them.

6.3 Credit Risk

For banks credit risk is central to the ERM exercise. It is important for insurers as well, chiefly in the area of reinsurance recoveries. As an asset class these tend to be comparable in size to the major categories of investment assets, and also comparable to total capital. Non-recovery is a risk, but also reluctant recovery has to be considered. Slowdowns in payments and coverage disputes would be included in the latter. These are much more likely for reinsurers who have ceased underwriting, even though they are still solvent. These reinsurers tend to have much less of an incentive to maintain a good reputation in the market and more incentive to retain assets than do ongoing writers.

Regulators and rating agencies have charges for reinsurance credit risk, and that is a starting point for quantification. These models classify reinsurers by strength and have different charges depending on that. Long and short term recoveries could also have different charges. Also too much emphasis on any one reinsurer is penalized.

These charges are aimed at expected credit shortfalls but are not distributional. Studies of distributional aspects would thus be useful. Also correlations with other risk sources should be incorporated. For instance, reinsurers could fail or withdraw from the market after major catastrophes. Large drops in asset markets could have similar impacts.

7. Summary

Banking ERM models have useful concepts that can serve as a starting point for insurer ERM models. Refinements are possible in several areas, however.

References

- Mango, D.F. 2003. Capital consumption: An alternative methodology for pricing reinsurance. *Casualty Actuarial Society Forum*. Winter: 351-378.
Online at <http://www.casact.org/pubs/forum/03wforum/03wf351.pdf> .
- Merton, R., and Perold, A. 1993. Theory of risk capital in financial firms. *Journal of Applied Corporate Finance* Fall: 16-32.
- Venter, G. 1991. Premium calculation implications of reinsurance without arbitrage. *ASTIN Bulletin* 21:2.
- — —. 2004. Testing distributions of stochastically generated yield curves. *ASTIN Bulletin*
- Venter, G., Barnett, J., and Owen, M. 2004. Market value of risk transfer: Catastrophe reinsurance case. *International AFIR Colloquium*,
http://afir2004.soa.org/afir_papers.htm
- Wang, S. 2004. Cat bond pricing using probability transforms. *Geneva Papers: Etudes et Dossiers*, special issue on “Insurance and the State of the Art in Cat Bond Pricing,” No. 278, January, 19-29, also available at
http://www.rmi.gsu.edu/Faculty/work_papers/Wang/catbondpaperbywang.pdf