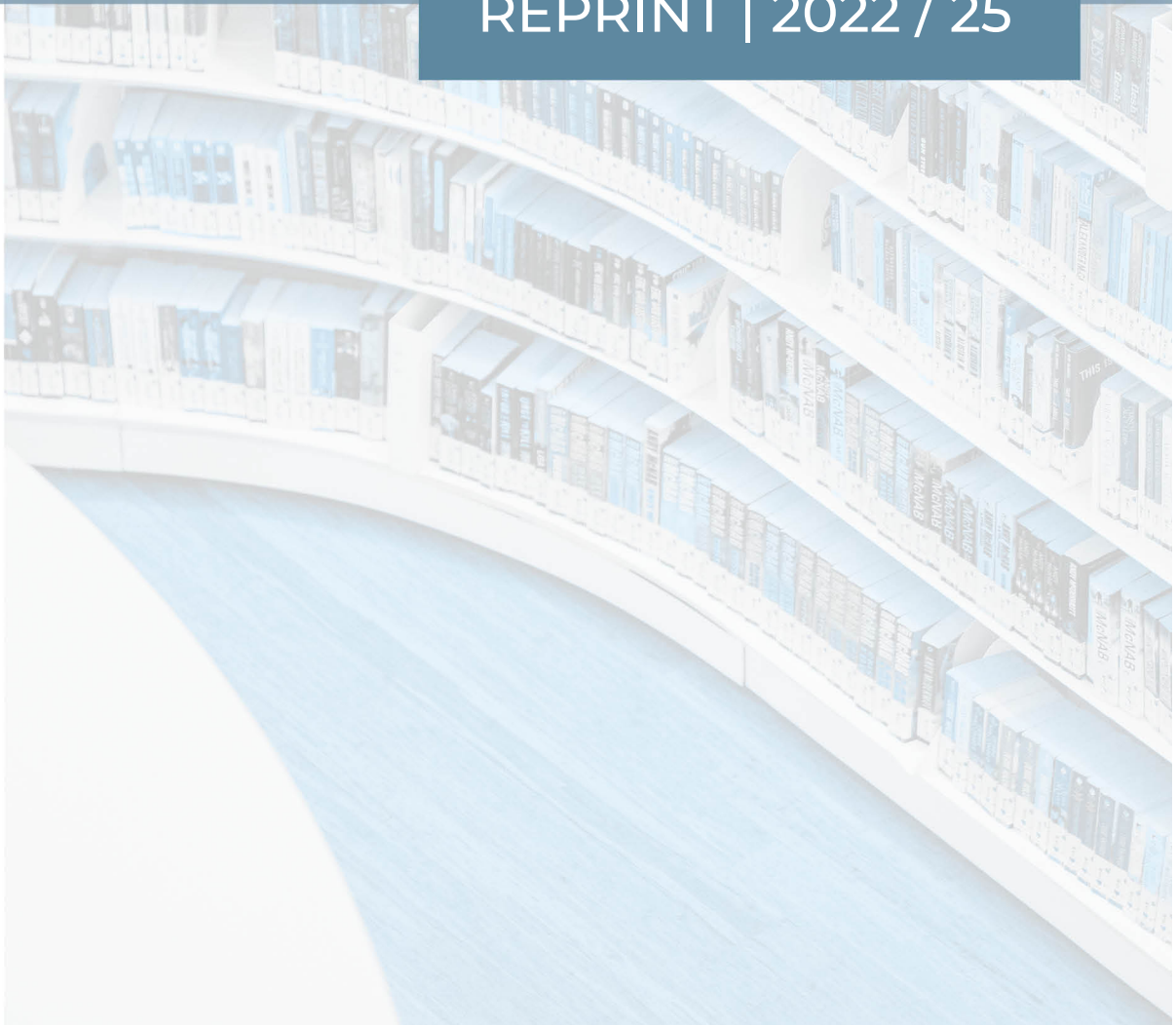


# CONDITIONAL TAIL EXPECTATION DECOMPOSITION AND CONDITIONAL MEAN RISK SHARING FOR DEPENDENT AND CONDITIONALLY INDEPENDENT LOSSES

Michel Denuit, Christian Y. Robert

REPRINT | 2022 / 25



## **ISBA**

Voie du Roman Pays 20 - L1.04.01

B-1348 Louvain-la-Neuve

Email : [lidam-library@uclouvain.be](mailto:lidam-library@uclouvain.be)

<https://uclouvain.be/en/research-institutes/lidam/isba/publication.html>

CONDITIONAL TAIL EXPECTATION  
DECOMPOSITION AND CONDITIONAL MEAN  
RISK SHARING FOR DEPENDENT AND  
CONDITIONALLY INDEPENDENT RISKS

MICHEL DENUIT

Institute of Statistics, Biostatistics and Actuarial Science - ISBA  
Louvain Institute of Data Analysis and Modeling - LIDAM  
UCLouvain  
Louvain-la-Neuve, Belgium  
`michel.denuit@uclouvain.be`

CHRISTIAN Y. ROBERT

Laboratory in Finance and Insurance - LFI  
CREST - Center for Research in Economics and Statistics  
ENSAE  
Paris, France  
`chrobert@ensae.fr`

June 8, 2020

### **Abstract**

Conditional tail expectations are often used in risk measurement and capital allocation. Conditional mean risk sharing appears to be effective in collaborative insurance, to distribute total losses among participants. This paper develops analytical results for risk allocation among different, correlated units based on conditional tail expectations and conditional mean risk sharing. Results available in the literature for independent risks are extended to correlated ones, in a unified way. The approach is applied to mixture models with correlated latent factors that are often used in insurance studies.

**Keywords:** Weighted distributions, size-biased transform, mixture models.

# 1 Introduction and motivation

The paper aims at contributing to the rich literature on risk allocation and risk sharing that are both core topics in actuarial science. In the former case, actuaries often need to allocate the available amount of capital across various entities, such as business lines, territories or even individual products. Costs of holding capital can then be distributed among entities so that it is fairly supported by policyholders, on the one hand, and the respective performances of each entity can be assessed, on the other hand. There are many ways to allocate capital, most of them being based on risk measures such as Conditional Tail Expectation (CTE) considered in this paper.

Risk sharing is a related, though distinct problem. In that context, economic agents hand their individual losses over to a pool and agree on some rule as to how the total pooled loss has to be divided among them. In the present paper, we focus on the conditional mean risk sharing, as defined by Denuit and Dhaene (2012). According to this rule, each participant to an insurance pool contributes the conditional expectation of the loss brought to the pool, given the total loss experienced by the entire pool.

In both cases, we wish to determine how to split the risk measure of the total risk, or the total risk itself, among individual entities. Precisely, the contribution of this paper is as follows. Representations for capital allocation based on CTE and conditional mean risk sharing derived in the literature for independent risks are extended here to correlated ones. To this end, we use a multivariate weighted distribution considered by Arratia et al. (2019) in the context of size-biasing sums of random variables. Several examples are discussed, including Liouville and infinitely divisible multivariate distributions. Mixture models where correlation arises from latent variables are carefully studied because of their wide applicability to insurance studies.

The remainder of this paper is organized as follows. Section 2 establishes the decomposition of CTE among correlated business lines whereas Section 3 considers the conditional mean risk sharing of correlated losses. As an application, Section 4 deals with multivariate mixture models.

## 2 General CTE decomposition formulas

### 2.1 CTE and size-biasing

Consider an insurance loss modeled as a non-negative random variable  $X$  with distribution function  $F_X$  and expected value  $E[X]$  such that  $0 < E[X] < \infty$ . Several risk measures for  $X$  are based on the function

$$t \mapsto E[X|X > t] = \frac{1}{P[X > t]} \int_t^\infty x dF_X(x)$$

where  $t$  is such that  $P[X > t] > 0$ . The Conditional Tail Expectation (CTE) is obtained when  $t$  corresponds to some Value-at-Risk (or quantile) of  $X$ . Such risk measures are intimately related to the size-biased transform, as shown next. Consider a (measurable) function  $g$  and

write

$$\begin{aligned}
\mathbb{E}[Xg(X)] &= \int_0^\infty xg(x)dF_X(x) \\
&= \mathbb{E}[X] \int_0^\infty g(x) \frac{x dF_X(x)}{\mathbb{E}[X]} \\
&= \mathbb{E}[X] \mathbb{E}[g(\tilde{X})]
\end{aligned} \tag{2.1}$$

where the non-negative random variable  $\tilde{X}$  with distribution function

$$\mathbb{P}[\tilde{X} \leq t] = \frac{1}{\mathbb{E}[X]} \int_0^t x dF_X(x),$$

is called the size-biased version of  $X$ . In some sense,  $\tilde{X}$  is larger compared to  $X$ , so that  $\tilde{X}$  represents a worse loss than  $X$ . In order to see why this is true, it suffices to notice that  $\tilde{X}$  is distributed as  $\max\{X, Z\}$  where the random variable  $Z$  is independent of  $X$  and has distribution function

$$\mathbb{P}[Z \leq t] = \frac{\mathbb{P}[\tilde{X} \leq t]}{\mathbb{P}[X \leq t]} = \frac{\mathbb{E}[X|X \leq t]}{\mathbb{E}[X]}. \tag{2.2}$$

Notice that the size-biased version of any constant  $c$  leaves it unchanged, that is,  $\tilde{c} = c$ . This shows that no unjustified loading is induced by the size-biased transform.

Now, inserting the function

$$g(x) = \mathbb{I}[x > t] = \begin{cases} 1 & \text{if } x > t \\ 0 & \text{otherwise} \end{cases} \tag{2.3}$$

in identity (2.1), we see that the representation

$$\mathbb{E}[X|X > t] = \mathbb{E}[X] \frac{\mathbb{P}[\tilde{X} > t]}{\mathbb{P}[X > t]} \tag{2.4}$$

is valid for any threshold  $t$ . The ratio  $\mathbb{P}[\tilde{X} > t]/\mathbb{P}[X > t]$  appearing in (2.4) exceeds unity so that we recover the classical inequality  $\mathbb{E}[X|X > t] \geq \mathbb{E}[X]$  for all  $t$  from elementary probability.

The size-biased transform can be traced back to the late 1960s in the statistical literature. It is an example of weighted distribution. Initially developed in order to unify various sampling distributions when the chance of being recorded by an observer varies, weighted distributions are closely related to weighted risk measures and weighted capital allocation rules. See Furman and Zitikis (2009). Among these weighted distributions, the size-biased, or length-biased one corresponds to the identity weight function. It refers to the situation where larger observations are more likely to be recorded. Hence, the available data are of bigger size compared to the actual population values. Translated to an actuarial context, this means that claim amounts are made larger before performing actuarial calculations, which generates a safety loading. The size-biased transform has proven to be useful in the study of risk measures after the pioneering work by Furman and Landsman (2005, 2008) and Furman and Zitikis (2008a,b) in connection to (2.4).

## 2.2 Independent losses

Consider  $n$  insurance losses modeled as non-negative, mutually independent random variables  $X_1, X_2, \dots, X_n$ , each with  $0 < E[X_i] < \infty$ . Define the aggregate loss  $S$  as  $S = \sum_{i=1}^n X_i$ . Clearly,

$$E[S|S > t] = \sum_{i=1}^n E[X_i|S > t]$$

so that the total  $E[S|S > t]$  can be decomposed into the sum of individual contributions  $E[X_i|S > t]$  that are useful in risk management to allocate the CTE across lines of business, for instance (taking for  $t$  a quantile of  $S$ ). It is known from Furman and Landsman (2005) that the representation

$$E[X_i|S > t] = E[X_i] \frac{P[S - X_i + \tilde{X}_i > t]}{P[S > t]} \quad (2.5)$$

holds true, where the size-biased version  $\tilde{X}_i$  of  $X_i$  is assumed to be independent of  $X_1, X_2, \dots, X_n$ . In words, the contribution  $E[X_i|S > t]$  of risk  $X_i$  to  $E[S|S > t]$  is equal to its expected value  $E[X_i]$  increased by the ratio of the excess probabilities  $P[S - X_i + \tilde{X}_i > t]$  and  $P[S > t]$ , where

$$S - X_i + \tilde{X}_i = \sum_{j \neq i} X_j + \tilde{X}_i$$

is the sum of all risks  $X_j$ , except the  $i$ th one which is replaced with its size-biased version  $\tilde{X}_i$ . Hence,  $S - X_i + \tilde{X}_i$  tends to be larger compared to  $S$ .

Identities (2.4)-(2.5) have been exploited in a number of papers (including those cited in the preceding section) to derive many useful properties for risk measures, mainly for continuous losses assumed to be mutually independent. In the next section, we extend these results to general losses, possibly correlated.

## 2.3 Dependent losses

Let us now extend the representation (2.5) recalled in the preceding section to correlated risks. To this end, consider a collection of  $n$  insurance losses, modeled as  $n$  non-negative, possibly correlated random variables  $X_1, X_2, \dots, X_n$ . Hence, we consider the random vector, or insurance portfolio  $\mathbf{X} = (X_1, \dots, X_n)$  with joint distribution function  $F_{\mathbf{X}}$ . As before, define the aggregate loss  $S$  as  $S = \sum_{i=1}^n X_i$ .

For each  $i \in \{1, 2, \dots, n\}$ , define the random vector  $\mathbf{Y}^{[i]} = (Y_1^{[i]}, \dots, Y_n^{[i]})$  with joint distribution function  $F_{\mathbf{Y}^{[i]}}$  given by

$$\begin{aligned} F_{\mathbf{Y}^{[i]}}(y_1, \dots, y_n) &= \int_0^{y_1} \dots \int_0^{y_n} \frac{x_i dF_{\mathbf{X}}(x_1, \dots, x_n)}{E[X_i]} \\ &= \frac{E[X_i | X_1 \leq y_1, \dots, X_n \leq y_n]}{E[X_i]} F_{\mathbf{X}}(y_1, \dots, y_n). \end{aligned} \quad (2.6)$$

The random vector  $\mathbf{Y}^{[i]}$  with distribution function (2.6) plays a central role in the extensions of the results presented in the preceding section to correlated risks. Distributions (2.6)

have been considered by Arratia et al. (2019) in relation to size-biasing sums of random variables. See Section 2.4 below for details. The distribution function (2.6) is a multivariate weighted version of the distribution function for  $\mathbf{X}$ . Multivariate weighted distributions have been reviewed by Navarro et al. (2006). The weight function  $w$  corresponding to (2.6) is  $w(x_1, \dots, x_n) = x_i/E[X_i]$  that has been considered in Jain and Nanda (1995).

We see from (2.6) that  $Y_i^{[i]}$  is distributed as  $\tilde{X}_i$ , marginally. The marginal distribution of  $Y_j^{[i]}$  for  $j \neq i$  is given by

$$\begin{aligned} P[Y_j^{[i]} \leq y_j] &= \int_0^{y_j} \int_0^\infty \frac{x_i dF_{(X_i, X_j)}(x_i, x_j)}{E[X_i]} \\ &= \frac{E[X_i I[X_j \leq y_j]]}{E[X_i]} \\ &= P[X_j \leq y_j] + \frac{\text{Cov}[X_i, I[X_j \leq y_j]]}{E[X_i]} \end{aligned} \quad (2.7)$$

$$= \frac{E[X_i | X_j \leq y_j]}{E[X_i]} P[X_j \leq y_j], \quad (2.8)$$

where  $I[\cdot]$  is the indicator function, equal to 1 if the condition appearing between the brackets is satisfied, and to 0 otherwise. Interestingly, some positive dependence is needed among the components of  $\mathbf{X}$  to ensure that the components of the random vector  $\mathbf{Y}^{[i]}$  are “larger” compared to those of  $\mathbf{X}$ . If  $X_i$  and  $X_j$  are negatively related then it can be expected that  $X_i$  and  $I[X_j \leq y_j]$  are positively correlated (since  $x_j \mapsto I[x_j \leq y_j]$  is a decreasing function) so that we see from (2.7) that the inequality  $P[Y_j^{[i]} \leq y_j] \geq P[X_j \leq y_j]$  holds true. If this is the case for all  $y_j$  then  $Y_j^{[i]}$  tends to be smaller compared to  $X_j$ , in the sense that  $Y_j^{[i]}$  is dominated by  $X_j$  in (first-order) stochastic dominance. This is for instance the case when  $X_i$  is negatively expectation dependent on  $X_j$ , that is, when the inequality

$$E[X_i | X_j \leq y_j] \geq E[X_i]$$

holds true for all  $y_j$ . In words, this means that the knowledge that  $X_j$  is small, that is,  $X_j$  falls below the threshold  $y_j$ , makes  $X_i$  larger on average. We then see from (2.8) that the inequality  $P[Y_j^{[i]} \leq y_j] \geq P[X_j \leq y_j]$  is valid for all  $y_j$ , so that  $Y_j^{[i]}$  is smaller than  $X_j$  in (first-order) stochastic dominance, as mentioned previously. Switching from  $\mathbf{X}$  to  $\mathbf{Y}^{[i]}$  is thus not necessarily a conservative strategy in this case. To prevent such a phenomenon to occur, some positive dependence is needed among  $X_1, X_2, \dots, X_n$ .

Considering (2.6), in order to ensure that  $F_{\mathbf{X}}$  dominates  $F_{\mathbf{Y}^{[i]}}$ , so that  $\mathbf{Y}^{[i]}$  is larger than  $\mathbf{X}$  (in the sense of the lower orthant order, that is, the joint distribution function  $F_{\mathbf{X}}$  dominates  $F_{\mathbf{Y}^{[i]}}$  everywhere), the inequality

$$E[X_i | X_1 \leq y_1, \dots, X_n \leq y_n] \leq E[X_i]$$

has to be valid for all  $y_1, \dots, y_n$ . This condition expresses some positive relationship between the individual risks  $X_1, \dots, X_n$ . It appears to be similar to the one imposed by Guo et al. (2016); see also Denuit and Mesfioui (2017).



We know from Proposition 6.1(v) in Navarro et al. (2006) that provided  $\mathbf{X}$  possesses positively associated components,  $\mathbf{Y}^{[i]}$  is larger than  $\mathbf{X}$  in multivariate stochastic dominance. Notice that

$$dF_{\mathbf{Y}^{[i]}}(y_1, \dots, y_n) = \frac{y_i}{E[X_i]} dF_{\mathbf{X}}(y_1, \dots, y_n)$$

so that the ratio  $dF_{\mathbf{Y}^{[i]}}/dF_{\mathbf{X}}$  is non-decreasing. This corresponds to condition (1.8) in Cohen and Sackrowitz (1995). Together with association, that condition is known to imply multivariate (first-order) stochastic dominance. This is formally stated in the next result, of which we provide the reader with an elementary proof borrowed from Shaked and Shanthikumar (2007, Theorem 6.B.8).

**Property 2.1.** *Assume that the risks  $X_1, \dots, X_n$  are associated random variables, that is, the covariance  $\text{Cov}[g_1(\mathbf{X}), g_2(\mathbf{X})]$  is non-negative for all non-decreasing functions  $g_1$  and  $g_2$ . Then, the random vector  $\mathbf{Y}^{[i]}$  with distribution function (2.6) is larger than  $\mathbf{X}$  in the sense of multivariate first-order stochastic dominance, that is, the inequality  $E[g(\mathbf{X})] \leq E[g(\mathbf{Y}^{[i]})]$  holds true for every non-decreasing function  $g$  such that the expectations exist.*

*Proof.* Considering a non-decreasing function  $g$ ,

$$\begin{aligned} E[g(\mathbf{Y}^{[i]})] &= \int_0^\infty \dots \int_0^\infty g(\mathbf{y}) dF_{\mathbf{Y}^{[i]}}(\mathbf{y}) \\ &= \int_0^\infty \dots \int_0^\infty g(\mathbf{y}) \frac{dF_{\mathbf{Y}^{[i]}}(\mathbf{y})}{dF_{\mathbf{X}}(\mathbf{y})} dF_{\mathbf{X}}(\mathbf{y}) \\ &\geq \int_0^\infty \dots \int_0^\infty g(\mathbf{y}) dF_{\mathbf{X}}(\mathbf{y}) \int_0^\infty \dots \int_0^\infty \frac{dF_{\mathbf{Y}^{[i]}}(\mathbf{y})}{dF_{\mathbf{X}}(\mathbf{y})} dF_{\mathbf{X}}(\mathbf{y}) \\ &= E[g(\mathbf{X})] \end{aligned}$$

where the inequality comes from the covariance inequality defining association and the non-decreasingness of the ratio  $dF_{\mathbf{Y}^{[i]}}(\mathbf{y})/dF_{\mathbf{X}}(\mathbf{y})$ . This ends the proof.  $\square$

Property 2.1 considers risks  $X_1, \dots, X_n$  that are associated random variables. Association is a positive dependence concept reviewed in Section 7.2.3 of Denuit et al. (2005). It is known to increase many risk measures compared to the independent case.

The next result gives a general representation formula for  $E[X_i|S > s]$  based on the random vector  $\mathbf{Y}^{[i]}$ . As a corollary, we obtain the existing result for independent risks recalled in Section 2.2.

**Proposition 2.2.** *Consider insurance losses  $X_1, \dots, X_n$  with joint distribution function  $F_{\mathbf{X}}$  and expected values  $E[X_i]$  such that  $0 < E[X_i] < \infty$ . Let  $\mathbf{Y}^{[i]} = (Y_1^{[i]}, \dots, Y_n^{[i]})$  be a random vector with joint distribution function  $F_{\mathbf{Y}^{[i]}}$  given by (2.6). Then for any  $s \geq 0$  and for any  $i \in \{1, 2, \dots, n\}$ , we have*

$$E[X_i|S > s] = E[X_i] \frac{P[Y_1^{[i]} + \dots + Y_n^{[i]} > s]}{P[X_1 + \dots + X_n > s]}.$$

*Proof.* Following the lines leading to (2.1), we obtain

$$\begin{aligned}
\mathbb{E}[X_i g(S)] &= \int_0^\infty \cdots \int_0^\infty x_i g(x_1 + \dots + x_n) dF_{\mathbf{X}}(x_1, \dots, x_n) \\
&= \mathbb{E}[X_i] \int_0^\infty \cdots \int_0^\infty g(x_1 + \dots + x_n) \frac{x_i dF_{\mathbf{X}}(x_1, \dots, x_n)}{\mathbb{E}[X_i]} \\
&= \mathbb{E}[X_i] \mathbb{E}[g(Y_1^{[i]} + \dots + Y_n^{[i]})].
\end{aligned} \tag{2.9}$$

Now, let us apply (2.9) to the function  $g$  given by

$$g(x_1, \dots, x_n) = \mathbb{I} \left[ \sum_{j=1}^n x_j > s \right]$$

to obtain the identity

$$\mathbb{P} \left[ Y_1^{[i]} + \dots + Y_n^{[i]} > s \right] = \frac{\mathbb{E}[X_i \mathbb{I}[S > s]]}{\mathbb{E}[X_i]}.$$

This ends the proof.  $\square$

**Corollary 2.3.** *Consider independent risks  $X_1, X_2, \dots, X_n$ . Let  $\tilde{X}_1, \tilde{X}_2, \dots, \tilde{X}_n$  be their corresponding size-biased versions, assumed to be independent and independent of  $X_1, X_2, \dots, X_n$ . If  $X_1, \dots, X_n$  are mutually independent then  $Y_1^{[i]}, \dots, Y_n^{[i]}$  are also independent random variables. Furthermore, (2.8) shows that  $Y_j^{[i]}$  is distributed as  $X_j$  for  $j \neq i$ , and  $Y_i^{[i]} \stackrel{d}{=} \tilde{X}_i$ , where “ $\stackrel{d}{=}$ ” means “is distributed as”. The distributional equality*

$$Y_1^{[i]} + \dots + Y_n^{[i]} \stackrel{d}{=} S - X_i + \tilde{X}_i$$

*thus holds true when  $X_1, \dots, X_n$  are mutually independent and we recover (2.5) as a particular case.*

## 2.4 Size-biased transform of a sum of dependent risks

Arratia et al. (2019, Section 2.4) studied the effect of size-biasing on a sum of random variables. Their result can also be obtained with the help of the CTE decomposition derived in Proposition 2.2. This can be shown as follows. On the one hand, the identity appearing in equation (2.4) allows us to write

$$\mathbb{E}[S | S > t] = \mathbb{E}[S] \frac{\mathbb{P}[\tilde{S} > t]}{\mathbb{P}[S > t]}$$

for any threshold  $t$ . On the other hand, Proposition 2.2 allows us to write

$$\begin{aligned}
\mathbb{E}[S | S > s] &= \sum_{i=1}^n \mathbb{E}[X_i | S > s] \\
&= \sum_{i=1}^n \mathbb{E}[X_i] \frac{\mathbb{P}[Y_1^{[i]} + \dots + Y_n^{[i]} > s]}{\mathbb{P}[X_1 + \dots + X_n > s]}.
\end{aligned}$$

Combining these identities, we get

$$P[\tilde{S} > t] = \sum_{i=1}^n \frac{E[X_i]}{E[S]} P[Y_1^{[i]} + \dots + Y_n^{[i]} > t].$$

This shows that  $\tilde{S}$  obeys a discrete mixture with mixing distribution assigning probability  $\frac{E[X_i]}{E[S]}$  on index  $i$ . Precisely, define the random variable  $K$  valued in  $\{1, 2, \dots, n\}$ , independent of  $\mathbf{X}$  and of  $\mathbf{Y}^{[1]}, \dots, \mathbf{Y}^{[n]}$  such that  $P[K = k] = \frac{E[X_k]}{E[S]}$  for  $k \in \{1, 2, \dots, n\}$ . Then,

$$\tilde{S} \stackrel{d}{=} \sum_{i=1}^n Y_i^{[K]} = \begin{cases} Y_1^{[1]} + \dots + Y_n^{[1]} & \text{with probability } \frac{E[X_1]}{E[S]} \\ Y_1^{[2]} + \dots + Y_n^{[2]} & \text{with probability } \frac{E[X_2]}{E[S]} \\ \vdots & \\ Y_1^{[n]} + \dots + Y_n^{[n]} & \text{with probability } \frac{E[X_n]}{E[S]} \end{cases}.$$

This shows that size-biasing a sum of correlated risks is equivalent to selecting at random (according to the discrete distribution assigning probability  $\frac{E[X_k]}{E[S]}$  to index  $k$ ) one element among  $\{\mathbf{Y}^{[1]}, \dots, \mathbf{Y}^{[n]}\}$ . Thus, the elements  $\mathbf{Y}^{[i]}$  corresponding to risks  $X_i$  with larger expected values are more likely to be selected.

As a particular case, consider independent risks  $X_1, \dots, X_n$ . Let  $\tilde{X}_1, \dots, \tilde{X}_n$  be their corresponding size-biased versions, assumed to be independent, and independent of  $X_1, \dots, X_n$ . Let the random variable  $K$  valued in  $\{1, 2, \dots, n\}$  be defined as above, and independent of  $X_1, \dots, X_n$  and of  $\tilde{X}_1, \dots, \tilde{X}_n$ . We then have  $\tilde{S} \stackrel{d}{=} S - X_K + \tilde{X}_K = \sum_{j \neq K} X_j + \tilde{X}_K$ .

### 3 Conditional mean risk sharing

#### 3.1 General representation formula

Consider the random vector, or insurance portfolio  $\mathbf{X} = (X_1, \dots, X_n)$  with joint distribution function  $F_{\mathbf{X}}$ . The next result gives a general representation formula for  $E[X_i|S = s]$  in terms of the random vectors  $\mathbf{Y}^{[i]}$ .

**Proposition 3.1.** *Consider insurance losses  $X_1, \dots, X_n$  with joint distribution function  $F_{\mathbf{X}}$  and expected values  $E[X_i]$  such that  $0 < E[X_i] < \infty$ . Let  $\mathbf{Y}^{[i]} = (Y_1^{[i]}, \dots, Y_n^{[i]})$  be a random vector with joint distribution function  $F_{\mathbf{Y}^{[i]}}$  given by (2.6). The following results hold:*

- (i) *if  $X_1, \dots, X_n$  are continuous random variables with joint probability density function  $f_{\mathbf{X}}$  then for any  $s \geq 0$ ,*

$$E[X_i|S = s] = E[X_i] \frac{f_{Y_1^{[i]} + \dots + Y_n^{[i]}}(s)}{f_{X_1 + \dots + X_n}(s)}.$$

- (ii) *if  $X_1, \dots, X_n$  are valued in  $\{0, 1, 2, \dots\}$  then for any  $s \in \{0, 1, 2, \dots\}$ ,*

$$E[X_i|S = s] = E[X_i] \frac{P[Y_1^{[i]} + \dots + Y_n^{[i]} = s]}{P[X_1 + \dots + X_n = s]}.$$

(iii) if  $X_1, \dots, X_n$  are zero-augmented random variables with positive probability masses at the origin and probability density functions over  $(0, \infty)$  then  $E[X_i|S = 0] = 0$  and for any  $s > 0$ , the representation in (i) holds true.

*Proof.* Let us apply (2.9) to the function  $g$  given by

$$g(x_1, \dots, x_n) = \mathbf{I} \left[ \sum_{j=1}^n x_j \leq s \right]$$

to get the identity

$$\mathbf{P} \left[ Y_1^{[i]} + \dots + Y_n^{[i]} \leq s \right] = \frac{E[X_i \mathbf{I}[S \leq s]]}{E[X_i]} \Leftrightarrow E[X_i \mathbf{I}[S \leq s]] = E[X_i] \mathbf{P} \left[ Y_1^{[i]} + \dots + Y_n^{[i]} \leq s \right].$$

To establish the validity of statement (i), notice that we can also write

$$E[X_i \mathbf{I}[S \leq s]] = \int_0^s E[X_i | S = t] f_S(t) dt.$$

Taking the derivative of these expressions with respect to  $s$  gives the identity stated in (i). Turning to (ii), we can write

$$\begin{aligned} E[X_i \mathbf{I}[S = s]] &= E[X_i \mathbf{I}[S \leq s]] - E[X_i \mathbf{I}[S \leq s-1]] \\ &= E[X_i] \mathbf{P} \left[ Y_1^{[i]} + \dots + Y_n^{[i]} = s \right], \end{aligned}$$

as announced. This ends the proof.  $\square$

It is worth to stress that the representation in Proposition 3.1(i) is related to Bartlett's formula for conditional expectations  $E[U|V = v]$  in terms of characteristic functions (Zabell, 1979). Essentially, the densities are replaced with their inverse transform to obtain Bartlett's formula.

When  $X_1, \dots, X_n$  are mutually independent, we recover as a particular case the result obtained by Denuit (2019, Proposition 2.3), as formally stated next.

**Corollary 3.2.** *Consider independent risks  $X_1, X_2, \dots, X_n$ . Let  $\tilde{X}_1, \tilde{X}_2, \dots, \tilde{X}_n$  be their corresponding size-biased versions, assumed to be independent and independent of  $X_1, X_2, \dots, X_n$ . If  $X_1, \dots, X_n$  are mutually independent then  $Y_1^{[i]}, \dots, Y_n^{[i]}$  are also independent random variables. Furthermore,  $Y_j^{[i]}$  is distributed as  $X_j$  for  $j \neq i$ , while  $Y_i^{[i]} \stackrel{d}{=} \tilde{X}_i$ . The distributional equality*

$$Y_1^{[i]} + \dots + Y_n^{[i]} \stackrel{d}{=} S - X_i + \tilde{X}_i$$

*thus holds true when  $X_1, \dots, X_n$  are mutually independent and*

(i) *if  $X_1, X_2, \dots, X_n$  are absolutely continuous random variables with respective probability density functions  $f_{X_1}, f_{X_2}, \dots, f_{X_n}$ , then for any  $s > 0$*

$$E[X_i | S = s] = E[X_i] \frac{f_{S-X_i+\tilde{X}_i}(s)}{f_S(s)}. \quad (3.1)$$

(ii) if  $X_1, X_2, \dots, X_n$  are valued in  $\{0, 1, 2, \dots\}$ , then for any  $s \in \{0, 1, 2, \dots\}$

$$\mathbb{E}[X_i|S = s] = \mathbb{E}[X_i] \frac{\mathbb{P}[S - X_i + \tilde{X}_i = s]}{\mathbb{P}[S = s]}. \quad (3.2)$$

(iii) if  $X_1, \dots, X_n$  are zero-augmented random variables with positive probability masses at the origin and probability density functions over  $(0, \infty)$  then  $\mathbb{E}[X_i|S = 0] = 0$  and for any  $s > 0$ , (3.1) holds true.

## 3.2 Linear case

### 3.2.1 Characterization

With the help of multivariate Laplace transforms, Furman et al. (2018) characterized the case where the conditional expectations are linear in  $s$ , that is, the situation where the identity  $\mathbb{E}[X_i|S = s] = \frac{\mathbb{E}[X_i]}{\mathbb{E}[S]}s$  holds true. Considering Theorem 3.2 in that paper, denote as  $X = X_i$  and  $Y = \sum_{j \neq i} X_j$ . Let  $L(u, v) = \mathbb{E}[\exp(-uX - vY)]$  be the Laplace transform of the random couple  $(X, Y)$ . Then,

$$\mathbb{E}[X_i|S = s] = \frac{\mathbb{E}[X_i]}{\mathbb{E}[S]}s \Leftrightarrow \left. \frac{\frac{d}{du}L(u, v)}{\frac{d}{dv}L(u, v)} \right|_{(u,v)=(t,t)} = \frac{\mathbb{E}[X_i]}{\mathbb{E}\left[\sum_{j \neq i} X_j\right]} \text{ for all } t. \quad (3.3)$$

This further constrains the distribution of the random vectors  $\mathbf{Y}^{[i]}$ , as shown next. Considering Proposition 3.1(i), we see that under (3.3), we must have

$$\frac{\mathbb{E}[X_i]}{\mathbb{E}[S]}s = \mathbb{E}[X_i] \frac{f_{Y_1^{[i]} + \dots + Y_n^{[i]}}(s)}{f_S(s)} \Rightarrow f_{\tilde{S}}(s) = f_{Y_1^{[i]} + \dots + Y_n^{[i]}}(s). \quad (3.4)$$

Thus, (3.3) implies that all sums  $Y_1^{[i]} + \dots + Y_n^{[i]}$  are identically distributed, as  $\tilde{S}$ .

### 3.2.2 Inverted Dirichlet distributions

Assume that  $\mathbf{X}$  obeys the Inverted Dirichlet( $a_1, \dots, a_{n+1}$ ) distribution, for some positive parameters  $a_1, \dots, a_{n+1}$ , that is,  $\mathbf{X}$  possesses the joint probability density function

$$f_{\mathbf{X}}(x_1, \dots, x_n) = \frac{\Gamma(\sum_{i=1}^{n+1} a_i)}{\prod_{i=1}^{n+1} \Gamma(a_i)} \prod_{i=1}^n x_i^{a_i-1} \left(1 + \sum_{i=1}^n x_i\right)^{-\sum_{i=1}^{n+1} a_i} \quad (3.5)$$

for positive  $x_1, x_2, \dots, x_n$ . The random vector  $\mathbf{X}$  admits the representation

$$\mathbf{X} \stackrel{d}{=} \frac{1}{Z_{n+1}} (Z_1, \dots, Z_n)$$

where  $Z_1, \dots, Z_n, Z_{n+1}$  are independent random variables obeying Chi-Square distributions with  $2a_j$  degrees of freedom,  $j \in \{1, \dots, n, n+1\}$ .

If  $a_{n+1} > 1$ , we have  $E[X_i] = \frac{a_i}{a_{n+1}-1}$  and the joint probability density function of  $\mathbf{Y}^{[i]}$  is obtained from

$$\begin{aligned} \frac{x_i}{E[X_i]} f_{\mathbf{X}}(x_1, \dots, x_n) &= \frac{\Gamma\left(\sum_{j=1}^{n+1} a_j\right)}{\Gamma(a_i + 1) \Gamma(a_{n+1} - 1) \prod_{i=1, \dots, n, j \neq i} \Gamma(a_j)} \\ &\quad x_i^{a_i} \prod_{i=1, \dots, n, j \neq i} x_j^{a_j - 1} \left(1 + \sum_{j=1}^n x_j\right)^{-\sum_{j=1}^{n+1} a_j}. \end{aligned}$$

We deduce that  $\mathbf{Y}^{[i]}$  follows the Inverted Dirichlet( $a_1, \dots, a_i + 1, \dots, a_{n+1} - 1$ ) distribution. Since

$$S = \sum_{j=1}^n X_j = \frac{\sum_{j=1}^n Z_j}{Z_{n+1}},$$

the probability density function of the sum  $S$  is given by

$$f_S(s) = \frac{\Gamma\left(\sum_{i=1}^{n+1} a_i\right)}{\Gamma\left(\sum_{i=1}^n a_i\right) \Gamma(a_{n+1})} s^{\sum_{i=1}^n a_i} (1+s)^{-\sum_{i=1}^{n+1} a_i}.$$

For the same reason, we also get

$$f_{\sum_{j=1}^n Y_j^{[i]}}(s) = \frac{\Gamma\left(\sum_{i=1}^{n+1} a_i\right)}{\Gamma\left(\sum_{i=1}^n a_i + 1\right) \Gamma(a_{n+1} - 1)} s^{\sum_{i=1}^n a_i + 1} (1+s)^{-\sum_{i=1}^{n+1} a_i}.$$

It follows that from Proposition 3.1(i) that

$$\begin{aligned} E[X_i | S = s] &= E[X_i] \frac{f_{\sum_{j=1}^n Y_j^{[i]}}(s)}{f_S(s)} \\ &= \frac{a_i}{a_{n+1} - 1} \frac{\Gamma\left(\sum_{j=1}^n a_j\right)}{\Gamma\left(\sum_{j=1}^n a_j + 1\right)} \frac{\Gamma(a_{n+1})}{\Gamma(a_{n+1} - 1)} s \\ &= \frac{a_i}{\sum_{j=1}^n a_j} s. \end{aligned}$$

The conditional expectation is thus linear in  $s$ , with a slope equal to the ratio of  $a_i$  over the sum of  $a_1, \dots, a_n$ . As it will be explained below for the whole class of Liouville distributions, this result can also be deduced from (3.3).

### 3.2.3 Liouville distributions

The result established in the preceding section is just a particular case of a more general statement applying to the whole family of Liouville distributions. Recall that an absolutely continuous random vector  $\mathbf{X}$  obeys a multivariate Liouville distribution if its joint probability density function is proportional to

$$g\left(\sum_{j=1}^n x_j\right) \prod_{i=1}^n x_i^{a_i - 1}$$

where  $x_i > 0$ ,  $a_i > 0$  and the function  $g$  is positive, continuous and appropriately integrable. Gupta and Richards (1987) have used the notation  $L_n[g; a_1, \dots, a_n]$  to refer to this distribution.

We only consider Liouville distributions of the first kind for which the support is non-compact. In this case

$$\mathbf{X} \stackrel{d}{=} R\mathbf{Z},$$

where  $R = \sum_{j=1}^n X_j$  obeys the univariate Liouville distribution  $L_1[g; a]$  with  $a = \sum_{i=1}^n a_i$  and  $\mathbf{Z}$  is independent of  $R$  and possesses the Dirichlet distribution with joint probability density function

$$f_{\mathbf{Z}}(z_1, \dots, z_n) = \frac{\Gamma(\sum_{i=1}^n a_i)}{\prod_{i=1}^n \Gamma(a_i)} \prod_{i=1}^{n-1} z_i^{a_i-1} \left(1 - \sum_{i=1}^{n-1} z_i\right)^{a_n-1}, \quad z_i > 0, \quad \sum_{i=1}^n z_i = 1.$$

We recover as a particular case the Inverted Dirichlet distribution considered in the preceding section with  $g(t) = (1+t)^{-\sum_{i=1}^{n+1} a_i}$  for  $t > 0$ ,  $a_{n+1} > 0$ . The corresponding probability density function is proportional to

$$\frac{\prod_{i=1}^n x_i^{a_i-1}}{(1 + \sum_{i=1}^n x_i)^{\sum_{i=1}^{n+1} a_i}}$$

which corresponds to (3.5). As another example, the multivariate correlated Gamma distribution is obtained with  $g(t) = t^{c-1}e^{-bt}$  for  $t > 0$ ,  $c > 0$ ,  $b > 0$ . The corresponding probability density function is proportional to

$$\left(\sum_{i=1}^n x_i\right)^{c-1} \prod_{i=1}^n (e^{-bx_i} x_i^{a_i-1}).$$

Let us now derive the distribution of  $\mathbf{Y}^{[i]}$ . Since  $x_i f_{\mathbf{X}}(x_1, \dots, x_n)$  is proportional to

$$g\left(\sum_{j=1}^n x_j\right) x_i^{a_i} \prod_{j=1, \dots, n, j \neq i} x_j^{a_j-1},$$

we deduce that  $\mathbf{Y}^{[i]}$  follows the  $L_n[g; a_1, \dots, a_i + 1, \dots, a_n]$  distribution as soon as  $E[X_i] < \infty$ . We have  $f_S(s) \propto g(s) s^{a-1}$  where “ $\propto$ ” means “is proportional to”, and therefore,  $f_{\sum_{j=1}^n Y_j^{[i]}}(s) \propto g(s) s^a$ . It follows that  $E[X_i | S_n = s] \propto s$  which implies

$$E[X_i | S_n = s] = \frac{E[X_i]}{E[S_n]} s.$$

Since  $E[X_i] = a_i E[R] / a$ , we finally get

$$E[X_i | S_n = s] = \frac{a_i}{\sum_{j=1}^n a_j} s.$$

This result can also be obtained as a consequence of (3.3). To this end, consider

$$X = X_i = RZ_i \text{ and } Y = \sum_{j \neq i} X_j = R \sum_{j \neq i} Z_j = R(1 - Z_i).$$

Then,

$$L(u, v) = \mathbb{E} \left[ \exp \left( -R \left( uZ_i + v \sum_{j \neq i} Z_j \right) \right) \right]$$

and the ratio appearing in (3.3) writes

$$\begin{aligned} \frac{\frac{d}{du} L(u, v)}{\frac{d}{dv} L(u, v)} \Big|_{(u, v) = (t, t)} &= \frac{\mathbb{E}[RZ_i \exp(-tR)]}{\sum_{j \neq i} \mathbb{E}[RZ_j \exp(-tR)]} \\ &= \frac{\mathbb{E}[Z_i]}{\sum_{j \neq i} \mathbb{E}[Z_j]} \\ &= \frac{\mathbb{E}[X_i]}{\mathbb{E} \left[ \sum_{j \neq i} X_j \right]}. \end{aligned}$$

This confirms that the conditional mean risk sharing is linear in the case of Liouville distributions.

In Section 4, we consider conditionally independent risks. This construction is widely used in actuarial models and allows us to derive extensions of the results obtained earlier in the independent case. This is also another approach to deal with Liouville distributions which are correlated by the common factor  $R$  that plays the role of the latent variable  $\Lambda$  inducing correlation between the conditionally independent risks.

### 3.2.4 Infinitely divisible distributions

According to Corollary 2.5 in Horn and Steutel (1978), a positive random vector  $\mathbf{X}$  with infinitely divisible distribution on  $\mathbb{R}_+^n$  can be characterized by its Laplace transform of the form

$$\mathbb{E}[e^{-\langle \mathbf{t}, \mathbf{X} \rangle}] = \exp \left( \int_{\mathbb{R}_+^n} \frac{e^{-\langle \mathbf{t}, \mathbf{x} \rangle} - 1}{\|\mathbf{x}\|} \nu(d\mathbf{x}) \right)$$

where  $\langle \mathbf{x}, \mathbf{y} \rangle$  stands for the inner product  $\sum_{i=1}^n x_i y_i$ ,  $\|\mathbf{x}\| = \sqrt{\langle \mathbf{x}, \mathbf{x} \rangle}$  for the Euclidean norm and where  $\nu$  is a measure such that  $\nu(\{0\}) = 0$  and  $\int_{\mathbf{y} > \mathbf{x}} \frac{\nu(d\mathbf{y})}{\|\mathbf{y}\|} < \infty$  for all  $\mathbf{x}$  in the interior of  $\mathbb{R}_+^n$ . The expectation of  $X_i$  is given by

$$\mathbb{E}[X_i] = \int_{\mathbb{R}_+^n} \frac{x_i}{\|\mathbf{x}\|} \nu(d\mathbf{x}).$$

The random vector  $\mathbf{Y}^{[i]}$  possesses the Laplace transform

$$\begin{aligned} \mathbb{E}[e^{-\langle \mathbf{t}, \mathbf{Y}^{[i]} \rangle}] &= \frac{\mathbb{E}[X_i e^{-\langle \mathbf{t}, \mathbf{X} \rangle}]}{\mathbb{E}[X_i]} = -\frac{1}{\mathbb{E}[X_i]} \frac{\partial}{\partial t_i} \mathbb{E}[e^{-\langle \mathbf{t}, \mathbf{X} \rangle}] \\ &= \frac{1}{\int_{\mathbb{R}_+^n} (x_i / \|\mathbf{x}\|) \nu(d\mathbf{x})} \int_{\mathbb{R}_+^n} \frac{x_i}{\|\mathbf{x}\|} e^{-\langle \mathbf{t}, \mathbf{x} \rangle} \nu(d\mathbf{x}) \int_{[0, \infty)} e^{itx} \frac{\nu(dx)}{\nu([0, \infty))} \mathbb{E}[e^{-\langle \mathbf{t}, \mathbf{X} \rangle}] \\ &= \mathbb{E}[e^{-\langle \mathbf{t}, \mathbf{Z}^{[i]} \rangle}] \mathbb{E}[e^{-\langle \mathbf{t}, \mathbf{X} \rangle}] \end{aligned}$$



where  $\mathbf{Z}^{[i]}$  is a random variable with probability density function proportional to  $(x_i / \|\mathbf{x}\|) \nu(d\mathbf{x})$ , and therefore  $\mathbf{Y}^{[i]} \stackrel{d}{=} \mathbf{X} + \mathbf{Z}^{[i]}$  where  $\mathbf{X}$  and  $\mathbf{Z}^{[i]}$  are independent. See also item (b) in Theorem 2.4 by Horn and Steutel (1978).

Assume that there exist positive constants  $\lambda_{ij}$  such that for all  $i \neq j$

$$\frac{\int_{\mathbb{R}_+^n} \frac{x_i}{\|\mathbf{x}\|} e^{-\langle \mathbf{t}, \mathbf{x} \rangle} \nu(d\mathbf{x})}{\int_{\mathbb{R}_+^n} \frac{x_j}{\|\mathbf{x}\|} e^{-\langle \mathbf{t}, \mathbf{x} \rangle} \nu(d\mathbf{x})} = \lambda_{ij}, \text{ for all } \mathbf{t}.$$

In this case, we have

$$\lambda_{ij} = \frac{\int_{\mathbb{R}_+^n} (x_i / \|\mathbf{x}\|) \nu(d\mathbf{x})}{\int_{\mathbb{R}_+^n} (x_j / \|\mathbf{x}\|) \nu(d\mathbf{x})}$$

and

$$\mathbb{E}[e^{-\langle \mathbf{t}, \mathbf{Z}^{[1]} \rangle}] = \dots = \mathbb{E}[e^{-\langle \mathbf{t}, \mathbf{Z}^{[n]} \rangle}]$$

as well as

$$\mathbb{E}[e^{-\langle \mathbf{t}, \mathbf{Y}^{[1]} \rangle}] = \dots = \mathbb{E}[e^{-\langle \mathbf{t}, \mathbf{Y}^{[n]} \rangle}].$$

Note that

$$\mathbb{E}[e^{-t\langle \mathbf{1}, \mathbf{Y}^{[i]} \rangle}] = \mathbb{E}[e^{-t\langle \mathbf{1}, \mathbf{Z}^{[i]} \rangle}] \mathbb{E}[e^{-t\langle \mathbf{1}, \mathbf{X} \rangle}] = \mathbb{E}[e^{-t\langle \mathbf{1}, \mathbf{Z}^{[i]} \rangle}] \mathbb{E}[e^{-tS}]$$

and that

$$\begin{aligned} \mathbb{E}[e^{-t\tilde{S}}] &= \frac{\mathbb{E}[S e^{-tS}]}{\mathbb{E}[S]} = -\frac{1}{\mathbb{E}[S]} \frac{\partial}{\partial t} \mathbb{E}[e^{-tS}] = -\frac{1}{\mathbb{E}[S]} \frac{\partial}{\partial t} \mathbb{E}[e^{-t\langle \mathbf{1}, \mathbf{X} \rangle}] \\ &= \frac{1}{\mathbb{E}[S]} \int_{\mathbb{R}_+^n} \frac{\langle \mathbf{1}, \mathbf{x} \rangle}{\|\mathbf{x}\|} e^{-t\langle \mathbf{1}, \mathbf{x} \rangle} \nu(d\mathbf{x}) \exp \left( \int_{\mathbb{R}_+^n} \frac{e^{-t\langle \mathbf{1}, \mathbf{x} \rangle} - 1}{\|\mathbf{x}\|} \nu(d\mathbf{x}) \right) \\ &= \sum_{i=1}^n \frac{\mathbb{E}[X_i]}{\mathbb{E}[S]} \mathbb{E}[e^{-t\langle \mathbf{1}, \mathbf{Z}^{[i]} \rangle}] \mathbb{E}[e^{-t\langle \mathbf{1}, \mathbf{X} \rangle}] \\ &= \sum_{i=1}^n \frac{\mathbb{E}[X_i]}{\mathbb{E}[S]} \mathbb{E}[e^{-t\langle \mathbf{1}, \mathbf{Z}^{[i]} \rangle}] \mathbb{E}[e^{-tS}] \\ &= \mathbb{E}[e^{-t\langle \mathbf{1}, \mathbf{Y}^{[i]} \rangle}] \text{ for } i = 1, \dots, n, \end{aligned}$$

in accordance with (3.4).

It finally follows that

$$\mathbb{E}[X_i | S = s] = \mathbb{E}[X_i] \frac{f_{\tilde{S}}(s)}{f_S(s)} = \frac{\mathbb{E}[X_i]}{\mathbb{E}[S]} s = \frac{\int_{\mathbb{R}_+^n} (x_i / \|\mathbf{x}\|) \nu(d\mathbf{x})}{\sum_{j=1}^n \int_{\mathbb{R}_+^n} (x_j / \|\mathbf{x}\|) \nu(d\mathbf{x})} s.$$

Again, this result can also be obtained with the help of (3.3).

## 4 Application to dependence by mixture

### 4.1 Size-biased transform of a mixture

Arratia et al. (2019, Lemma 2.6) studied the size-biased transform of mixtures. This section summarizes their findings, that appear to be useful for studying dependence induced by correlated latent factors in the next sections. We provide the reader with an elementary reasoning, for convenience and because this result will also be central to the extension to size-biasing conditionally independent risks.

Consider a family of non-negative random variables  $\{X(\lambda), \lambda \geq 0\}$  indexed by a single, non-negative parameter  $\lambda$ . Let  $\Lambda$  be a mixing parameter with distribution function  $F_\Lambda$ . The corresponding mixture  $X(\Lambda)$  has distribution function

$$P[X(\Lambda) \leq x] = \int_0^\infty P[X(\lambda) \leq x] dF_\Lambda(\lambda).$$

Define  $\tilde{X}(\lambda)$  to be the size-biased version of  $X(\lambda)$ , with distribution function

$$P[\tilde{X}(\lambda) \leq z] = \frac{E[X(\lambda)I[X(\lambda) \leq z]]}{E[X(\lambda)]}.$$

Then, the size-biased version  $\widetilde{X(\Lambda)}$  of the mixture  $X(\Lambda)$  corresponds to the mixture of the non-negative random variables  $\{\tilde{X}(\lambda), \lambda \geq 0\}$  with mixing parameter  $\Lambda^*$  distributed according to

$$dF_{\Lambda^*}(\lambda) = \frac{E[X(\lambda)]}{E[X(\Lambda)]} dF_\Lambda(\lambda). \quad (4.1)$$

This result can be rewritten as

$$\widetilde{X(\Lambda)} \stackrel{d}{=} \tilde{X}(\Lambda^*) \quad (4.2)$$

and is easily obtained as follows:

$$\begin{aligned} P[\widetilde{X(\Lambda)} \leq x] &= \frac{1}{E[X(\Lambda)]} \int_0^x t dF_{X(\Lambda)}(t) \\ &= \frac{1}{E[X(\Lambda)]} \int_0^\infty \int_0^x t dF_{X(\lambda)}(t) dF_\Lambda(\lambda) \\ &= \frac{1}{E[X(\Lambda)]} \int_0^\infty E[X(\lambda)] F_{\tilde{X}(\lambda)}(t) dF_\Lambda(\lambda) \\ &= \int_0^\infty F_{\tilde{X}(\lambda)}(t) dF_{\Lambda^*}(\lambda) \end{aligned}$$

which shows that the announced distributional equality (4.2) is indeed valid.

In the particular case where  $E[X(\lambda)]$  is proportional to  $\lambda$ , that is, if the identity  $E[X(\lambda)] = a\lambda$  holds for some positive constant  $a$ , we obtain

$$dF_{\Lambda^*}(\lambda) = \frac{\lambda}{E[\Lambda]} dF_\Lambda(\lambda),$$

so that  $\Lambda^* \stackrel{d}{=} \tilde{\Lambda}$  and  $\widetilde{X(\Lambda)} \stackrel{d}{=} \tilde{X}(\tilde{\Lambda})$ .

## 4.2 Common mixture model

### 4.2.1 Definition

The common mixture model (in the terminology of Wang, 1998) consists in conditionally independent risks and forms the basis of the credibility approach. The kind of dependence induced by this construction is comprehensively studied in Denuit et al. (2005, Chapter 7). The intuition behind this modeling approach is as follows: an external mechanism, described by the positive random variable  $\Lambda$ , influences several risks  $X_1, X_2, \dots, X_n$  with  $X_i \stackrel{d}{=} X_i(\Lambda)$ . Given the environmental parameter  $\Lambda$ , the individual risks are independent. Formally, the joint distribution function of the portfolio vector  $\mathbf{X}$  can be written as

$$\begin{aligned} F_{\mathbf{X}}(t_1, \dots, t_n) &= \mathbb{E}[\mathbb{P}[X_1 \leq t_1, \dots, X_n \leq t_n | \Lambda]] \\ &= \int_0^\infty \left( \prod_{i=1}^n \mathbb{P}[X_i(\lambda) \leq t_i] \right) dF_\Lambda(\lambda). \end{aligned} \quad (4.3)$$

Notice that this construction is rather general and covers for instance the case of the common shock model, where each risk  $X_i$  is obtained as the sum of two independent random variables, with the second one common to all risks. This common shock then plays the role of  $\Lambda$  in the common mixture model (4.3). The multivariate Gamma distribution proposed in Furman and Landsman (2005) is built in this way, defining  $X_i = \beta_i Y_0 + Y_i$  for some  $\beta_i > 0$ , where  $Y_0, Y_1, \dots, Y_n$  are independent, Gamma-distributed random variables. Here,  $Y_0$  plays the role of  $\Lambda$ . This is also the approach followed by Furman and Landsman (2010) to derive a multivariate extension of the Tweedie distribution. The multivariate Pareto II distribution used by Asimit et al. (2013) is also obtained in this way, specifying  $X_i = \mu_i + Y_i/Y_0$  where  $Y_0, Y_1, \dots, Y_n$  are independent random variables with  $Y_1, \dots, Y_n$  obeying the same Negative Exponential distribution and  $Y_0$  following the Gamma distribution. The multivariate Pareto distribution proposed by Asimit et al. (2010) corresponds to  $X_i = \min\{\sigma_i Y_0 + \mu_i, Y_i\}$  with  $\sigma_i > 0$ , where  $Y_0, Y_1, \dots, Y_n$  are independent, Pareto-distributed random variables. In credit risk modeling, time-to-defaults  $X_i$  are sometimes assumed to be subject to a competing risk mechanism with  $X_i = \min\{Y_i, Y_0\}$  where  $Y_0, Y_1, \dots, Y_n$  are independent, positive random variables. The common factor  $Y_0$  impacting all times-to-default accounts for a systematic shock. See, e.g., Giesecke (2003).

### 4.2.2 Association

It is reasonable to expect that the risks  $X_i$  resulting from this construction are positively correlated provided they all move in the same direction with  $\Lambda$ . This makes the random vector  $\mathbf{Y}^{[i]}$  with distribution function (2.6) larger so that size-biasing induces a safety margin in this case. This is formally stated in the next result.

**Property 4.1.** *Consider the risks  $X_1, X_2, \dots, X_n$  with joint distribution function (4.3). Assume that*

$$\mathbb{P}[X_i(\lambda) \leq t] \geq \mathbb{P}[X_i(\lambda + \delta) \leq t] \text{ for all } t, \lambda > 0, \delta > 0, \text{ and } i = 1, 2, \dots, n.$$

Then, the random vector  $\mathbf{Y}^{[i]}$  with distribution function (2.6) is larger than  $\mathbf{X}$  in the sense of multivariate first-order stochastic dominance, that is, the inequality  $\mathbb{E}[g(\mathbf{X})] \leq \mathbb{E}[g(\mathbf{Y}^{[i]})]$  holds true for every non-decreasing function  $g$  such that the expectations exist.

*Proof.* We know from Property 7.2.16 in Denuit et al. (2005) that the risks  $X_1, \dots, X_n$  are associated random variables, that is, the covariance  $\text{Cov}[g_1(\mathbf{X}), g_2(\mathbf{X})]$  is non-negative for all non-decreasing functions  $g_1$  and  $g_2$ . The announced result then follows from Property 2.1.  $\square$

### 4.2.3 Representation formula

When the risks  $X_1, X_2, \dots, X_n$  are conditionally independent, given  $\Lambda$ , with joint distribution function (4.3), it turns out that each  $\mathbf{Y}^{[i]}$  also obeys a common mixture model with a change in the  $i$ th conditional distribution and in the mixing distribution. This is formally stated next.

**Proposition 4.2.** *Consider the risks  $X_1, X_2, \dots, X_n$  with joint distribution function (4.3). Then, each  $\mathbf{Y}^{[i]}$  also possesses a joint distribution function of the form (4.3), that is,*

$$F_{\mathbf{Y}^{[i]}}(t_1, \dots, t_n) = \int_0^\infty \left( \prod_{j \neq i} \mathbb{P}[X_j(\lambda) \leq t_j] \right) \mathbb{P}[\tilde{X}_i(\lambda) \leq t_i] dF_{\Lambda_i^*}(\lambda)$$

where  $\Lambda_i^*$  is distributed according to (4.1), that is,

$$dF_{\Lambda_i^*}(\lambda) = \frac{\mathbb{E}[X_i(\lambda)]}{\mathbb{E}[X_i(\Lambda)]} dF_{\Lambda}(\lambda).$$

*Proof.* It suffices to write

$$\begin{aligned} F_{\mathbf{Y}^{[i]}}(y_1, \dots, y_n) &= \frac{\mathbb{E}[X_i \mathbb{I}[X_1 \leq y_1, \dots, X_n \leq y_n]]}{\mathbb{E}[X_i]} \\ &= \frac{\mathbb{E} \left[ \mathbb{E} \left[ X_i \prod_{j=1}^n \mathbb{I}[X_j \leq y_j] \middle| \Lambda \right] \right]}{\mathbb{E}[X_i]} \\ &= \mathbb{E} \left[ \frac{\mathbb{E}[X_i \mathbb{I}[X_i \leq y_i] | \Lambda]}{\mathbb{E}[X_i]} \prod_{j \neq i} \mathbb{P}[X_j \leq y_j | \Lambda] \right]. \end{aligned}$$

This ends the proof.  $\square$

### 4.2.4 Individual model of risk theory

Assume that  $\Lambda \in [0, 1]$  and that each risk  $X_i$  is of the form  $X_i = I_i C_i$ . Given  $\Lambda = \lambda$ ,  $I_i = I_i(\lambda)$  is Bernoulli( $\lambda$ ) distributed, that is,

$$\mathbb{P}[I_i = 1 | \lambda] = 1 - \mathbb{P}[I_i = 0 | \lambda] = \lambda.$$

We assume that the costs  $C_i$  are independent of  $\Lambda$  with respective means  $E[C_i] = \mu_i$ . Then,  $X_i = X_i(\lambda)$  is such that  $E[X_i(\lambda)] = \lambda\mu_i$  for every  $i = 1, 2, \dots, n$ . Given  $\Lambda$ , all the random variables  $I_1, \dots, I_n, C_1, \dots, C_n$  are assumed to be independent.

Each  $\mathbf{Y}^{[i]}$  obeys a mixture of  $(X_1(\lambda), \dots, \tilde{C}_i, \dots, X_n(\lambda))$  with mixing parameter  $\Lambda_i^*$  distributed according to

$$dF_{\Lambda_i^*}(\lambda) = \frac{\lambda\mu_i}{E[\Lambda]\mu_i} dF_{\Lambda}(\lambda) \Leftrightarrow \Lambda_i^* \stackrel{d}{=} \tilde{\Lambda} \text{ for all } i.$$

Thus, the following representations hold true:

$$\begin{aligned} E[X_i|S > t] &= E[\Lambda]\mu_i \frac{P\left[\tilde{C}_i + \sum_{j \neq i} X_j(\tilde{\Lambda}) > t\right]}{P\left[\sum_{j=1}^n X_j(\Lambda) > t\right]} \\ E[X_i|S = t] &= E[\Lambda]\mu_i \frac{f_{\tilde{C}_i + \sum_{j \neq i} X_j(\tilde{\Lambda})}(t)}{f_{\sum_{j=1}^n X_j(\Lambda)}(t)}. \end{aligned}$$

**Example 4.3.** Assume that given  $\Lambda \in (0, 1)$ , each  $I_i$  is Bernoulli distributed with the same mean  $\Lambda$  where  $\Lambda$  follows the Beta( $\eta, \beta$ ) distribution. The costs  $C_i$  are independent of  $\Lambda$  and obey Gamma( $\alpha_i, \tau$ ) distributions. Given  $\Lambda$ , all the random variables  $I_1, \dots, I_n, C_1, \dots, C_n$  are assumed to be independent.

It is easy to see that  $\tilde{C}_i$  follows the Gamma( $\alpha_i + 1, \tau$ ) distribution. Moreover,

$$\mathbf{Y}^{[i]} \stackrel{d}{=} (X_1(\tilde{\Lambda}), \dots, \tilde{C}_i, \dots, X_n(\tilde{\Lambda}))$$

where  $\tilde{\Lambda}$  follows the Beta( $\eta + 1, \beta$ ) distribution. Since

$$\begin{aligned} & f_{\tilde{C}_i + \sum_{j \neq i} X_j(\tilde{\Lambda})|\tilde{\Lambda}=\lambda}(s) \\ &= \exp(-s\tau) \sum_{j=0}^{n-1} \lambda^j (1-\lambda)^{n-1-j} \sum_{\sum_{l \neq i} k_l = j: k_l \in \{0,1\}} \frac{\tau}{\Gamma(\alpha_i + \sum_{l \neq i} k_l \alpha_l + 1)} (\tau s)^{\alpha_i + \sum_{l \neq i} k_l \alpha_l} \end{aligned}$$

we then have

$$\begin{aligned} & f_{\tilde{C}_i + \sum_{j \neq i} X_j(\tilde{\Lambda})}(s) \\ &= \exp(-s\tau) \sum_{j=0}^{n-1} \frac{B(\eta + 1 + j, \beta + n - 1 - j)}{B(\eta + 1, \beta)} \sum_{\sum_{l \neq i} k_l = j: k_l \in \{0,1\}} \frac{\tau}{\Gamma(\alpha_i + \sum_{l \neq i} k_l \alpha_l + 1)} (\tau s)^{\alpha_i + \sum_{l \neq i} k_l \alpha_l}. \end{aligned}$$

If  $\alpha_1 = \dots = \alpha_n = \alpha$ , we thus obtain

$$\begin{aligned} & f_{\tilde{C}_i + \sum_{j \neq i} X_j(\tilde{\Lambda})}(s) \\ &= \exp(-s\tau) \sum_{j=0}^{n-1} \frac{B(\eta + 1 + j, \beta + n - 1 - j)}{B(\eta + 1, \beta) B(j + 1, n)} \frac{\Gamma(n + 1 + j)}{\Gamma(n - j)} \frac{\tau}{\Gamma(j\alpha + 1)} (\tau s)^{j\alpha + 1} \end{aligned}$$

and  $E[X_i|S = s] = \frac{\alpha_i}{\alpha} s$  holds true in this case, as expected.

#### 4.2.5 Compound mixed Poisson model

Let us now assume that given  $\Lambda = \lambda$  each  $X_i$  is a compound Poisson sum. Precisely, assume that  $X_i = \sum_{k=1}^{N_i} C_{ik}$  where the claim severities  $C_{ik}$  are positive, continuous, and distributed as  $C_i$ , all these random variables being independent,  $i = 1, 2, \dots, n$ , and independent of  $(N_1, \dots, N_n)$ . We denote as  $\mu_i = E[C_i]$  the mean claim severity for  $X_i$ . Each  $N_i$  obeys the  $\text{Poisson}(a_i \Lambda)$  distribution given  $\Lambda$ , for some positive constants  $a_1, \dots, a_n$ . Claim severities thus remain independent but the random vector of claim frequencies  $(N_1, \dots, N_n)$  now obeys the common mixture model. In this case, all components  $X_i$  increase in  $\Lambda$  in the first-order stochastic dominance (see Chapter 7 in Denuit et al., 2005). This construction is a particular case of the model considered in Kim et al. (2019).

Here, we see that

$$dF_{\Lambda_i^*}(\lambda) = \frac{a_i \lambda \mu_i}{a_i E[\Lambda] \mu_i} dF_{\Lambda}(\lambda) = \frac{\lambda}{E[\Lambda]} dF_{\Lambda}(\lambda)$$

so that  $\Lambda_1^*, \dots, \Lambda_n^*$  are identically distributed and  $\Lambda_i^* \stackrel{d}{=} \tilde{\Lambda}$  for all  $i$ . Since  $\tilde{X}_i(\lambda) \stackrel{d}{=} X_i(\lambda) + \tilde{C}_i$  where the random variable  $\tilde{C}_i$  is independent of  $X_i(\lambda)$ , we get

$$F_{\mathbf{Y}^{[i]}}(t_1, \dots, t_n) = \int_0^\infty \left( \prod_{j \neq i} P[X_j(\lambda) \leq t_j] \right) P[X_i(\lambda) + \tilde{C}_i \leq t_i] dF_{\tilde{\Lambda}}(\lambda).$$

Thus, the following representations hold true:

$$\begin{aligned} E[X_i | S > t] &= a_i E[\Lambda] \mu_i \frac{P\left[\sum_{j=1}^n X_j(\tilde{\Lambda}) + \tilde{C}_i > t\right]}{P\left[\sum_{j=1}^n X_j(\Lambda) > t\right]} \\ E[X_i | S = t] &= a_i E[\Lambda] \mu_i \frac{f_{\sum_{j=1}^n X_j(\tilde{\Lambda}) + \tilde{C}_i}(t)}{f_{\sum_{j=1}^n X_j(\Lambda)}(t)}. \end{aligned}$$

Now, assume that severities also depend on  $\Lambda$ , that is,  $C_{ik} = C_{ik}(\Lambda)$  with  $E[C_{ik}(\lambda)] = \lambda \mu_i$ . Given  $\Lambda$ , all the random variables are supposed to be independent. Then,  $\mathbf{Y}^{[i]}$  is a mixture of

$$(X_1(\lambda), \dots, X_i(\lambda) + \tilde{C}_i(\lambda), \dots, X_n(\lambda))$$

with mixing parameter  $\Lambda_i^*$  distributed as

$$dF_{\Lambda_i^*}(\lambda) = \frac{a_i \lambda^2 \mu_i}{a_i E[\Lambda^2] \mu_i} dF_{\Lambda}(\lambda) = \frac{\lambda^2}{E[\Lambda^2]} dF_{\Lambda}(\lambda).$$

Thus, the mixing parameters  $\Lambda_i^*$  are identically distributed for all  $i$ , that is,  $\Lambda_i^* \stackrel{d}{=} \Lambda^*$  for all  $i$ . We have the following representations:

$$\begin{aligned} E[X_i | S > t] &= a_i E[\Lambda^2] \mu_i \frac{P\left[\sum_{j=1}^n X_j(\Lambda^*) + \tilde{C}_i(\Lambda^*) > t\right]}{P\left[\sum_{j=1}^n X_j(\Lambda) > t\right]} \\ E[X_i | S = t] &= a_i E[\Lambda^2] \mu_i \frac{f_{\sum_{j=1}^n X_j(\Lambda^*) + \tilde{C}_i(\Lambda^*)}(t)}{f_{\sum_{j=1}^n X_j(\Lambda)}(t)}. \end{aligned}$$

Some families of distributions are closed in the sense that if  $\Lambda$  obeys a distribution belonging to the family then this is also the case for  $\Lambda^*$ . Gamma distributions satisfy this property.

### 4.3 General mixture models

#### 4.3.1 Definition

Let us now extend the common mixture model by letting each risk  $X_i$  depend on its own latent variable  $\Theta_i$ , that is,

$$F_{\mathbf{X}}(t_1, \dots, t_n) = \int_0^\infty \dots \int_0^\infty \left( \prod_{i=1}^n P[X_i(\theta_i) \leq t_i] \right) dF_{\Theta}(\theta_1, \dots, \theta_n). \quad (4.4)$$

If  $\Theta_1 = \dots = \Theta_n = \Lambda$  then we are back to the common mixture model (4.3) studied in the preceding section.

Mixture models (4.4) are also widely used in insurance studies. In multiperil insurance modeling for instance, policyholders may own several insurance products (up to  $n$ , say). Given  $\Theta$ , the number of claims  $N_i$  recorded for product  $i$  is often assumed to be Poisson distributed with parameter  $\lambda_i \Theta_i$ , where  $\lambda_i$  is the a priori expected claim frequency for product  $i$ , based on policyholder's specific risk profile. Denuit and Lu (2020) assumed that  $\Theta$  obeys the Wishart distribution with dimension  $n$ . As another example, Denuit et al. (2015) introduced the Max-factor individual risk model for credit portfolios. Given  $\Theta \in [0, 1]^n$ , default indicators  $I_i$  are Bernoulli distributed with

$$P[I_i = 1 | \Theta_i] = 1 - P[I_i = 0 | \Theta_i] = \Theta_i,$$

with

$$\Theta_i = F_{\Psi} \left( \max\{\nu_i + \sigma_i \Psi_i, \mu_0 + \sigma_0 \Psi_0\} \right)$$

where  $F_{\Psi}$  is a suitable distribution function,  $\nu_i \in \mathbb{R}$  and  $\sigma_i \geq 0$ . There is thus a competition between the risk-specific effect  $\nu_i + \sigma_i \Psi_i$  and the global effect  $\mu_0 + \sigma_0 \Psi_0$  and only the largest one impacts on the occurrences of losses.

#### 4.3.2 Association

As it was the case for common mixture models, it is reasonable to expect that the risks  $X_i$  resulting from this construction are positively correlated provided they all move in the same direction with  $\theta_i$  and the components of  $\Theta$  are positively related. This makes the random vector  $\mathbf{Y}^{[i]}$  with distribution function (2.6) larger so that size-biasing induces a safety margin in this case. This is formally stated in the next result.

**Property 4.4.** *Consider the risks  $X_1, X_2, \dots, X_n$  with joint distribution function (4.4). Assume that*

$$P[X_i(\theta_i) \leq t] \geq P[X_i(\theta_i + \delta) \leq t] \text{ for all } t, \theta_i > 0, \delta > 0, \text{ and } i = 1, 2, \dots, n,$$

*and  $\Theta$  is associated. Then, the random vector  $\mathbf{Y}^{[i]}$  with distribution function (2.6) is larger than  $\mathbf{X}$  in the sense of multivariate first-order stochastic dominance, that is, the inequality  $E[g(\mathbf{X})] \leq E[g(\mathbf{Y}^{[i]})]$  holds true for every non-decreasing function  $g$  such that the expectations exist.*

*Proof.* We know from Property 7.2.17 in Denuit et al. (2005) that the risks  $X_1, \dots, X_n$  are associated random variables. The announced result then follows from Property 2.1.  $\square$

### 4.3.3 Representation formula

The next result shows that each  $\mathbf{Y}^{[i]}$  obeys a common mixture model with a change in the  $i$ th conditional distribution and in the mixing distribution.

**Proposition 4.5.** *Consider the risks  $X_1, X_2, \dots, X_n$  with joint distribution function (4.4). Then, each  $\mathbf{Y}^{[i]}$  also possesses a joint distribution function of the form (4.4), that is,*

$$F_{\mathbf{Y}^{[i]}}(t_1, \dots, t_n) = \int_0^\infty \dots \int_0^\infty \left( \prod_{j \neq i} P[X_j(\theta_j) \leq t_j] \right) P[\tilde{X}_i(\theta_i) \leq t_i] dF_{\Theta^{*[i]}}(\theta_1, \dots, \theta_n)$$

where the mixing parameter  $\Theta^{*[i]}$  has joint distribution

$$dF_{\Theta^{*[i]}}(\theta_1, \dots, \theta_n) = \frac{E[X_i(\theta_i)]}{E[X_i]} dF_{\Theta}(\theta_1, \dots, \theta_n).$$

*Proof.* It suffices to write

$$\begin{aligned} F_{\mathbf{Y}^{[i]}}(y_1, \dots, y_n) &= \frac{E[X_i I[X_1 \leq y_1, \dots, X_n \leq y_n]]}{E[X_i]} \\ &= \frac{E \left[ E \left[ X_i \prod_{j=1}^n I[X_j \leq y_j] \middle| \Theta \right] \right]}{E[X_i]} \\ &= E \left[ \frac{E[X_i I[X_i \leq y_i] | \Theta_i]}{E[X_i]} \prod_{j \neq i} P[X_j \leq y_j | \Theta_j] \right]. \end{aligned}$$

This ends the proof.  $\square$

### 4.3.4 Individual model of risk theory

Assume that each risk  $X_i$  is of the form  $X_i = I_i C_i$  and that  $\Theta \in [0, 1]^n$ . Given  $\Theta = \theta$ ,  $I_i = I_i(\theta_i)$  obeys the Bernoulli( $\theta_i$ ) distribution, that is,

$$P[I_i = 1 | \theta_i] = 1 - P[I_i = 0 | \theta_i] = \theta_i$$

for every  $i = 1, 2, \dots, n$ . Given  $\Theta$ , all the random variables  $I_j$  and  $C_j$  are assumed to be independent. Then,  $\mathbf{Y}^{[i]}$  is a mixture of

$$(X_1(\theta_1), \dots, \tilde{C}_i, \dots, X_n(\theta_n))$$

with mixing parameter  $\Theta^{[i]}$  defined in Proposition 4.5.



### 4.3.5 Compound mixed Poisson model

Assume that  $X_i = \sum_{k=1}^{N_i} C_{ik}$  where the claim severities  $C_{ik}$  are positive, continuous, and distributed as  $C_i$ . Given  $\Theta = \theta$ , each  $N_i$  obeys the  $\text{Poisson}(a_i\theta_i)$  distribution, for some positive constants  $a_1, \dots, a_n$ , and severities are such that  $E[C_{ik}(\theta_i)] = \theta_i\mu_i$ . Given  $\Theta$ , all the random variables are assumed to be independent. Then,  $\mathbf{Y}^{[i]}$  is a mixture of

$$(X_1(\theta_1), \dots, X_i(\theta_i) + \tilde{C}_i(\theta_i), \dots, X_n(\theta_n))$$

with mixing parameter  $\Theta^{\star[i]}$  distributed as

$$dF_{\Theta^{\star[i]}}(\theta_1, \dots, \theta_n) = \frac{a_i\theta_i^2\mu_i}{a_iE[\Theta_i^2]\mu_i}dF_{\Theta}(\theta_1, \dots, \theta_n) = \frac{\theta_i^2}{E[\Theta_i^2]}dF_{\Theta}(\theta_1, \dots, \theta_n).$$

Sometimes,  $\Theta$  and  $\Theta^{\star[i]}$  obey distributions in the same family. This is for instance the case when  $\Theta$  follows a multivariate Liouville distribution.

If the severities  $C_{ik}$  are independent of  $\Theta$  then we see that  $\mathbf{Y}^{[i]}$  is a mixture of

$$(X_1(\theta_1), \dots, X_i(\theta_i) + \tilde{C}_i, \dots, X_n(\theta_n))$$

with mixing parameter  $\Theta^{\star[i]}$ . Since

$$dF_{\Theta^{\star[i]}}(\theta_1, \dots, \theta_n) = \frac{a_i\theta_i\mu_i}{a_iE[\Theta_i]\mu_i}dF_{\Theta}(\theta_1, \dots, \theta_n) = \frac{\theta_i}{E[\Theta_i]}dF_{\Theta}(\theta_1, \dots, \theta_n),$$

we see that  $\Theta^{\star[i]} \stackrel{d}{=} \Theta^{[i]}$  for all  $i = 1, 2, \dots, n$ . Here also, if  $\Theta$  obeys a Liouville distribution then this is also the case for  $\Theta^{[i]}$ .

## References

- Arratia, R., Goldstein, L., Kochman, F. (2019). Size bias for one and all. *Probability Surveys* 16, 1-61.
- Asimit, A.V., Furman, E., Vernic, R. (2010). On a multivariate Pareto distribution. *Insurance: Mathematics and Economics* 46, 308-316.
- Asimit, A.V., Vernic, R., Zitikis, R. (2013). Evaluating risk measures and capital allocations based on multi-losses driven by a heavy-tailed background risk: The multivariate Pareto-II model. *Risks* 1, 14-33.
- Cohen, A., Sackrowitz, H.B. (1995). On stochastic ordering of random vectors. *Journal of Applied Probability* 32, 960-965.
- Denuit, M. (2019). Size-biased transform and conditional mean risk sharing, with application to P2P insurance and tontines. *ASTIN Bulletin* 49, 591-617.
- Denuit, M., Dhaene, J. (2012). Convex order and comonotonic conditional mean risk sharing. *Insurance: Mathematics and Economics* 51, 265-270.

- Denuit, M., Dhaene, J., Goovaerts, M.J., Kaas, R. (2005). Actuarial Theory for Dependent Risks: Measures, Orders and Models. Wiley, New York.
- Denuit, M., Kiriliouk, A., Segers, J. (2015). Max-factor individual risk models with application to credit portfolios. Insurance: Mathematics and Economics 62, 162-172.
- Denuit, M., Lu, Y. (2020). Wishart-Gamma random effects for multivariate experience ratemaking, frequency-severity experience ratemaking and micro loss-reserving. Submitted for publication.
- Denuit, M., Mesfioui, M. (2017). Preserving the Rothschild-Stiglitz type increase in risk with background risk: A characterization. Insurance: Mathematics and Economics 72, 1-5.
- Furman, E., Kuznetsov, A., Zitikis, R. (2018). Weighted risk capital allocations in the presence of systematic risk. Insurance: Mathematics and Economics 79, 75-81.
- Furman, E., Landsman, Z. (2005). Risk capital decomposition for a multivariate dependent gamma portfolio. Insurance: Mathematics and Economics 37, 635-649.
- Furman, E., Landsman, Z. (2008). Economic capital allocations for non-negative portfolios of dependent risks. ASTIN Bulletin 38, 601-619.
- Furman, E., Landsman, Z. (2010). Multivariate Tweedie distributions and some related capital-at-risk analyses. Insurance: Mathematics and Economics 46, 351-361.
- Furman, E., Zitikis, R. (2008a). Weighted risk capital allocations. Insurance: Mathematics and Economics 43, 263-269.
- Furman, E., Zitikis, R. (2008b). Weighted premium calculation principles. Insurance: Mathematics and Economics 42, 459-465.
- Furman, E., Zitikis, R. (2009). Weighted pricing functionals with applications to insurance: An overview. North American Actuarial Journal 13, 483-496.
- Giesecke, K. (2003). A simple exponential model for dependent defaults. Journal of Fixed Income 13, 74-83.
- Guo, X., Li, J., Liu, D., Wang, J. (2016). Preserving the Rothschild-Stiglitz type of increasing risk with background risk. Insurance: Mathematics and Economics 70, 144-149.
- Gupta, R.D., Richards, D.S.P. (1987). Multivariate Liouville distributions. Journal of Multivariate Analysis 23, 233-256.
- Horn, R.A., Steutel, F.W. (1978). On multivariate infinitely divisible distributions. Stochastic Processes and Their Applications 6, 139-151.
- Jain, K., Nanda, A.K. (1995). On multivariate weighted distributions. Communications in Statistics-Theory and Methods 24, 2517-2539.

- Kim, J. H., Jang, J., Pyun, C. (2019). Capital allocation for a sum of dependent compound mixed Poisson variables: A recursive algorithm approach. *North American Actuarial Journal* 23, 82-97.
- Navarro, J., Ruiz, J.M., Del Aguila, Y. (2006). Multivariate weighted distributions: A review and some extensions. *Statistics* 40, 51-64.
- Shaked, M., Shanthikumar, J.G. (2007). *Stochastic Orders*. Springer, New York.
- Wang, S. (1998). Aggregation of correlated risk portfolios: Models and algorithms. *Proceedings of the Casualty Actuarial Society*, 848-939.
- Zabell, S. (1979). Continuous versions of regular conditional distributions. *The Annals of Probability* 7, 159-165.