

Кафедральный предмет "Исследование Операций"

ТЕОРИЯ К ЭКЗАМЕНУ "ТЕОРИЯ РИСКА"

Собрала Аладина Екатерина 311
Редактировала Племяшова Лидия 311

Москва
2023

Вопросы к экзамену по курсу «Теория риска», 6-й семестр

1	Случайные величины, их характеристики. Функция распределения, плотность распределения, функция выживания, масса вероятности, производящие функции. Уровень риска. Мода. Моменты	5
1.1	Функция распределения	5
1.2	Случайные величины	5
1.3	Функция выживания	6
1.4	Функция плотности распределения	6
1.5	Функция массы вероятности	7
1.6	Производящая функция	7
1.7	Уровень риска	7
1.8	Мода	8
1.9	Основные количественные характеристики распределений: моменты . . .	8
1.10	Основные количественные характеристики распределений: центральные моменты	8
2	Превышение ущерба. Цензурированные величины, их моменты. Ограничение убытка, моменты	10
2.1	Превышение ущерба	10
2.2	Превышение ущерба: моменты	10
2.3	Цензурированные переменные	10
2.4	Моменты цензурированной случайной величины	11
2.5	Ограничение убытка	11
2.6	Моменты ограниченной случайной величины	11
3	Хвосты распределений. Оценка тяжести хвоста (предельное соотношение, уровень риска, среднее превышение). Равновесные распределения	13
3.1	Хвосты распределений	13
3.1.1	Оценка через моменты	13
3.1.2	Предельное соотношение хвостов	14
3.1.3	Использование уровня риска для оценки хвоста	14
3.1.4	Использование среднего превышения для оценки тяжести хвоста .	14
3.2	Равновесные распределения	15
4	Меры риска. Согласованность меры риска. Процентиль. VaR, TVaR	17
4.1	Меры риска	17
4.2	Согласованность меры риска	17
4.3	Условия согласованности	18
4.4	Процентиль	18
4.5	Value at Risk (VaR)	19
4.6	Tail Value at Risk (TVaR)	19

5	Смеси распределений (конечные, переменные)	21
5.1	Параметрические и масштабируемые распределения	21
5.2	Конечные смеси распределений	21
5.3	Переменные смеси	21
5.4	Смешивание	21
6	Получение новых распределений (умножение на константу, возведение в степень, экспоненцирование)	23
6.1	Умножение на константу	23
6.2	Возведение в степень	23
6.3	Экспоненцирование	23
7	Склейка распределений	24
7.0.1	Склейка (splicing)	24
7.0.2	Свойства склейки	24
8	Экспоненциальное семейство распределений	25
8.1	Линейное экспоненциальное семейство	25
8.1.1	Математическое ожидание линейного экспоненциального семейства	25
8.1.2	Дисперсия линейного экспоненциального семейства	25
9	Дискретные распределения. Смесь пуассоновских распределений	27
9.1	Дискретные распределения	27
9.1.1	Распределение Пуассона	27
9.2	Отрицательное биномиальное распределение как смесь пуассоновских . .	28
9.3	Отрицательное биномиальное распределение	28
9.4	Смесь пуассоновских распределений	29
10	Классы $(a, b, 0)$, $(a, b, 1)$, соотношение между ними	30
10.1	Класс $(a, b, 0)$	30
10.2	Класс $(a, b, 1)$	30
10.3	Взаимосвязь классов $(a, b, 0)$ и $(a, b, 1)$	31
11	Условная и безусловная франшизы. Расчет на убыток и на платеж	32
11.1	Влияние модификации покрытия на частоту и тяжесть убытков	32
11.2	Безусловная франшиза	32
11.2.1	Функции для случайных величин с франшизой	32
11.3	Условная франшиза	33
11.3.1	Функции для случайных величин с франшизой	33
11.4	Ожидаемый убыток и ожидаемый платеж	34
11.5	Влияние инфляции на безусловную франшизу	34
12	Коэффициент устранения потерь	35
12.1	Моменты ограниченной случайной величины	35
12.2	Безусловная франшиза	35
12.3	Коэффициент устранения потерь	35

13 Модификации полиса (сострахование, франшиза, лимит). Влияние инфляции	36
13.1 Ограничения по полисам	36
13.2 Сострахование, франшизы и лимиты	36
13.3 Средние платежи и убытки	37
14 Влияние франшизы на частоту требований	38
15 Определение частоты убытков по частоте платежей	39
15.1 Определение частоты убытков по частоте платежей	39
16 Модель коллективного риска	40
16.1 Модель коллективного риска	40
16.2 Модель индивидуального риска	40
16.3 Моделирование суммарного убытка	40
17 Составное распределение, моменты	42
17.1 Составное (обобщенное) распределение	42
17.2 Составное распределение	42
17.3 Моменты составного распределения	43
18 Страхование стоп-лосс	44
18.1 Линейная интерполяция нетто-премии стоп-лосс	44
18.2 Дискретный случай	45
19 Обобщенное распределение Пуассона	46
20 Численное нахождение функции распределения совокупного ущерба	47
21 Дискретизация	48
22 Рекурсивный метод, рекурсия для распределения Пуассона	49
22.1 Рекурсивный метод	49
22.2 Рекурсия для распределения Пуассона	49
22.3 Проблемы переполнения и недостаточной точности	50
22.4 Проблемы переполнения и недостаточной точности	50
22.5 Численная устойчивость	50
22.6 Непрерывное распределение ущерба	51
23 Построение арифметического распределения (округление, согласование моментов)	52
23.1 Метод округления	52
23.2 Метод локального согласования моментов	52
23.3 Свойство метода моментов	53
24 Модель индивидуального риска	54

25	Функция полезности. Классы функций полезности	56
25.1	Функция полезности	56
25.2	Классы функций полезности	56
26	Отношение к риску. Коэффициент несклонности к риску	57
26.1	Отношение к риску	57
26.2	Коэффициент несклонности к риску	57
27	Неравенство Йенсена	58
27.1	Выпуклые функции	58
27.2	Неравенство Йенсена	58
28	Верхняя и нижняя границы премии	59
28.1	Верхняя и нижняя границы премии	59
28.2	Оценка максимальной премии	59
29	Сравнение случайных величин по степени рискованности	60
29.1	Конечная случайная величина	60
29.2	Непрерывные случайные величины	61
30	Премия стоп-лосс, ее свойства	62
30.1	Премия стоп-лосс	62
30.2	Оптимальность перестрахования стоп-лосс	62
30.3	Оптимальность пропорционального перестрахования	63
31	Невошедшее в билеты из презентаций	65
31.1	Свойства суммы случайных величин	65
31.2	Примеры согласованных и несогласованных мер	65
31.3	Модели уязвимости (Frailty models)	66
31.4	Усеченные и модифицированные распределения $(a, b, 0)$	66
31.5	Усеченные и модифицированные в нуле распределения	67
32	Информация по экзамену	68
32.1	Примерная теория с экзамена 2023 года	68
32.2	Примерная теория с 1 пересдачи 2023 года	68
32.3	Кратко о подготовке	68
33	Билет с зачета ФИИТа 4 курс	69
34	Задачи с экзамена 2023	70
35	Литература, указанная преподавателем	71

1 Случайные величины, их характеристики. Функция распределения, плотность распределения, функция выживания, масса вероятности, производящие функции. Уровень риска. Мода. Моменты

1.1 Функция распределения

Определение. Функция распределения $F_X(x)$ или $F(x)$ случайной величины X :

$$F_X(x) = P(X \leq x)$$

Свойства:

1. $0 \leq F(x) \leq 1$
2. $F(x)$ не убывает
3. $F(x)$ непрерывна справа
- 4.

$$\lim_{x \rightarrow -\infty} F(x) = 0$$

$$\lim_{x \rightarrow \infty} F(x) = 1$$

1.2 Случайные величины

Определение. Явления – это то, что можно наблюдать.

Определение. Эксперимент - это наблюдение за данным явлением при определенных условиях. Результат эксперимента называется исходом; событие - это набор из одного или нескольких возможных исходов. Стохастическое явление - это явление, для которого эксперимент имеет более одного возможного исхода.

Определение. Событие, связанное со стохастическим явлением, называется случайным.

Определение. Вероятность - это мера возможности наступления события, измеряемая от 0 до 1.

Определение. Случайная величина - присваивает числовое значение для каждого возможного исхода.

Определение. Носителем случайной величины называется множество ее возможных значений N_X : $P(X \notin N_X) = 0$

Определение. Случайная величина называется *дискретной* если ее носитель содержит не более счетного числа значений.

Определение. Случайная величина называется *непрерывной*, если функция распределения непрерывна и дифференцируема всюду за исключением, возможно, счетного числа точек.

Определение. Случайная величина называется *смешанной*, если она не является дискретной и является непрерывной всюду, за исключением не менее одного значения и не более чем счетного числа значений.

1.3 Функция выживания

Определение. Функцией выживания, обычно обозначаемой через $S_X(x)$ или $S(x)$, для случайной величины X называется дополнительная функция распределения:

$$S_X(x) = P(X > x) = 1 - F_X(x)$$

Свойства:

1. $0 \leq S(x) \leq 1$
2. $S(x)$ не возрастает
3. $S(x)$ непрерывна справа
- 4.

$$\lim_{x \rightarrow -\infty} S(x) = 1$$

$$\lim_{x \rightarrow \infty} S(x) = 0$$

Исторически сложилось описывать случайные величины, обозначающие деньги, с помощью функции распределения, а случайные величины, обозначающие время, с помощью функции выживания.

1.4 Функция плотности распределения

Определение. Функция плотности (вероятности) $f_X(x)$ (или $f(x)$) является производной функции распределения или, что эквивалентно, производной функции выживания с минусом:

$$f(x) = F'(x) = -S'(x)$$

Функция плотности определена только в тех точках, где производная существует.

Функция распределения и функция выживания получаются интегрированием функции плотности по соответствующему интервалу:

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

$$F(b) = \int_{-\infty}^b f(x) dx$$

$$S(b) = \int_b^{\infty} f(x) dx$$

1.5 Функция массы вероятности

Определение. Функция вероятности или функция массы вероятности обычно обозначается через $p_X(x)$ или $p(x)$ и показывает вероятность в определенной точке (когда она не равна 0):

$$p_X(x) = P(X = x)$$

Определяется для дискретных случайных величин. В этом случае $F(x) = \sum_{y \leq x} p(y)$,
 $S(x) = \sum_{y > x} p(y)$

1.6 Производящая функция

Определение. Производящая функция моментов случайной величины $XM_X(z) = E[e^{zX}]$ для всех z для которых математическое ожидание существует.

Определение. Производящая функция вероятностей $P_X(z) = E[z^X]$ для всех z для которых математическое ожидание существует.

$$M_X(z) = P_X(e^z) \text{ и } P_X(z) = M_X(\ln z)$$

Производящие функции моментов (п.ф.м.) используются для непрерывных случайных величин, а производящие функции вероятностей (п.ф.в.) для дискретных. Эти функции используются для вычисления моментов и вероятностей. Кроме того, существует взаимно однозначное соответствие между функцией распределения случайной величины и ее п.ф.м. и п.ф.в.

1.7 Уровень риска

Определение. Уровень (интенсивность) риска, также называемый интенсивностью смертности и интенсивностью отказов, обозначается через $h_X(x)$ или $h(x)$:

$$h_X(x) = \frac{f_X(x)}{S_X(x)}$$

Для интенсивности смертности используется обозначение $\mu(x)$, для интенсивности $\lambda(x)$.

$$h_X(x) = -\frac{S'(x)}{S(x)} = -\frac{d \ln S(X)}{dx}$$

$$S(b) = \exp\left(-\int_0^b h(x) dx\right)$$

Для смешанного распределения уровень риска определяется только на части носителя. В точках дискретной функции массы вероятности уровень риска не определяется.

1.8 Мода

Определение. Модой случайной величины называется наиболее вероятное значение.

Для дискретной переменной это значение с наибольшей вероятностью.

Для непрерывной величины это точка, в которой достигается наибольшее значение функции плотности.

Если имеется несколько локальных максимумов функции плотности, они все считаются модами.

1.9 Основные количественные характеристики распределений: моменты

Определение. k -й момент случайной величины это ожидаемое (среднее) значение k -й степени этой величины (если оно существует). Обозначается через $E[X^k]$ или μ'_k .

Первый момент математическое ожидание, часто обозначается через μ .

μ никак не связано с интенсивностью смертности $\mu(x)$

$$\mu'_k(x) = E[X^k] = \begin{cases} \int_{-\infty}^{\infty} x^k f(x) dx, & \text{если сл.в. непрерывна} \\ \sum_j x_j^k p(x_j), & \text{если сл.в. дискретна} \end{cases}$$

Все формулы приводятся для всюду непрерывных или всюду дискретных случайных величин. Для смешанных распределений соответствующие выражения получаются интегрированием по плотности распределения там, где переменная непрерывна, и суммированием по функции вероятности там, где переменная дискретна.

1.10 Основные количественные характеристики распределений: центральные моменты

Определение. k -й центральный момент случайной величины равен математическому ожиданию k -й степени отклонения величины от среднего $E[(X - \mu)^k] \stackrel{\text{def}}{=} \mu_k$

Определение. Второй центральный момент называется *дисперсией* σ^2 или $Var[X]$

Определение. Корень из дисперсии σ называется *стандартным отклонением*.

Определение. Отношение стандартного отклонения к среднему значению называется *коэффициентом вариации*: $\nu = \frac{\sigma}{\mu}$

Определение. Отношение третьего центрального момента к кубу стандартного отклонения называется *коэффициентом асимметрии*: $\gamma_1 = \frac{\mu_3}{\sigma^3}$

Определение. Отношение четвертого центрального момента к четвертой степени стандартного отклонения называется *эксцессом (куртозисом)*: $\gamma_2 = \frac{\mu_4}{\sigma^4}$

$$\mu_k(x) = E[(X - \mu)^k] = \begin{cases} \int_{-\infty}^{\infty} (x - \mu)^k f(x) dx, & \text{если сл.в. непрерывна} \\ \sum_j (x_j - \mu)^k p(x_j), & \text{если сл.в. дискретна} \end{cases}$$

Интеграл берется только по тем x , для которых $f(x) > 0$.

Экссесс больше 3 показывает, что при постоянном стандартном отклонении относительно нормального распределения большую вероятность имеют точки, удаленные от среднего, чем точки, расположенные рядом со средним.

Связь между вторыми моментами:

$$\mu_2 = \int_{-\infty}^{\infty} (x - \mu)^2 f(x) dx = \int_{-\infty}^{\infty} (x^2 - 2\mu x + \mu^2) f(x) dx = E[X^2] - 2\mu E[X] + \mu^2 = \mu'_2 - \mu^2$$

2 Превышение ущерба. Цензурированные величины, их моменты. Ограничение убытка, моменты

2.1 Превышение ущерба

Определение. Для заданного значения d такого, что $P(X > d) > 0$, превышение ущерба есть $Y^d = X - d$ для $X > d$

Ожидаемое значение $e_X(d) = e(d) = E[Y^d] = E[X - d | X > d]$ называется *функцией среднего превышения ущерба*. Другие названия: остаточное время жизни и полная ожидаемая продолжительность жизни.

Y^d также называется усеченной слева и сдвинутой переменной.

Усеченная слева - значит, что значения меньше d не рассматриваются. Сдвинутая, так как d вычитается из остальных значений.

Когда X является переменной, отвечающей платежам, среднее превышение ущерба равно ожидаемому количеству заплаченного при условии, что платежи производятся после вычета d .

Когда X отвечает возрасту смерти, среднее превышение ущерба является ожидаемым оставшимся временем до момента смерти при условии, что человек дожил до возраста d .

2.2 Превышение ущерба: моменты

k -й момент величины превышения ущерба:

$$e_X^k(d) = E[X^k] = \begin{cases} \frac{\int_d^\infty (x-d)^k f(x) dx}{1-F(d)} & \text{для непрерывной переменной} \\ \frac{\sum_{x_j > d} (x_j-d)^k p(x_j)}{1-F(d)} & \text{для дискретной переменной} \end{cases}$$

Интегрируя по частям первый момент, получаем более простое выражение:

$$e_X(d) = \frac{\int_d^\infty (x-d)f(x)dx}{1-F(d)} = \frac{-(x-d)S(d)|_d^\infty + \int_d^\infty S(x)dx}{S(d)} = \frac{\int_d^\infty S(x)dx}{S(d)}$$

2.3 Цензурированные переменные

Определение. Цензурированная слева и сдвинутая переменная

$$Y^L = (X - d)_+ = \begin{cases} 0, & X \leq d \\ X - d, & X > d \end{cases}$$

Левостороннее цензурирование означает, что значения меньше d не игнорируются, а приравниваются 0.

Стандартных обозначений для моментов таких переменных нет.

Различие между переменными превышения и цензурированными слева переменными для финансовых потоков состоит в том, что первое рассчитывается для платежа, а второе для убытка.

В расчете на платеж переменная существует только, если платежи были произведены. В расчете на ущерб переменная принимает нулевое значение всегда, когда ущерб не приводит к выплате.

2.4 Моменты цензурированной случайной величины

$$E[(X - d)_+]^k = \begin{cases} \int_d^\infty (x - d)^k f(x) dx & \text{для непрерывной переменной} \\ \sum_{x_j > d} (x_j - d)^k p(x_j) & \text{для дискретной переменной} \end{cases}$$

$$E[(X - d)_+]^k = e^k(d)(1 - F(d))$$

2.5 Ограничение убытка

Определение. Величина ограниченного убытка:

$$Y = X \cap u = \begin{cases} X, X \leq u \\ u, X > u \end{cases}$$

$E[X \cap u]$ - ограниченная ожидаемая величина.

Переменная также называется цензурированной справа, поскольку значения, большие u , устанавливаются равными u .

Этой ситуации в страховании соответствует лимит страхового покрытия, указанный в полисе (максимальная выплата)

$$(X - d)_+ + (X \cap d) = X$$

2.6 Моменты ограниченной случайной величины

$$E[(X \cap u)^k] = \begin{cases} \int_{-\infty}^u x^k f(x) dx + u^k [1 - F(u)] & \text{для непрерывной переменной} \\ \sum_{x_j \leq u} x_j^k p(x_j) + u^k [1 - F(u)] & \text{для дискретной переменной} \end{cases}$$

$$\begin{aligned} E[(X \cap u)^k] &= \int_{-\infty}^0 x^k f(x) dx + \int_0^u x^k f(x) dx + u^k [1 - F(u)] = \\ &= x^k F(x) \Big|_{-\infty}^0 - \int_{-\infty}^0 kx^{k-1} F(x) dx - x^k S(x) \Big|_0^u + \int_0^u kx^{k-1} S(x) dx + u^k S(u) = \end{aligned}$$

$$= - \int_{-\infty}^0 kx^{k-1}F(x)dx + \int_0^u kx^{k-1}S(x)dx$$

При $k=1$:

$$E[X \cap u] = - \int_{-\infty}^0 F(x)dx + \int_0^u S(x)dx$$

3 Хвосты распределений. Оценка тяжести хвоста (предельное соотношение, уровень риска, среднее превышение). Равновесные распределения

3.1 Хвосты распределений

Определение. Хвостом распределения (более точно, правым хвостом) называется часть распределения, соответствующая большим значениям случайной величины.

Большие возможные значения оказывают наибольшее влияние на общие потери. Случайные величины, у которых более высокие вероятности приписаны большим значениям, имеют более тяжелый хвост.

Вес хвоста может быть относительным понятием (модель А имеет более тяжелый хвост, чем модель В) или абсолютным (распределения классифицируются как имеющие тяжелый хвост).

Применяются разные способы классификации распределений по степени тяжести хвоста

3.1.1 Оценка через моменты

В общем случае существование всех моментов свидетельствует об относительно легком правом хвосте. У распределений с тяжелым хвостом интеграл, определяющий моменты, может расходиться. Отсутствие всех или части моментов говорит о тяжелом правом хвосте

Пример:

Для гамма-распределения

$$\mu'_k = \int_0^{\infty} x^k \frac{x^{\alpha-1} e^{-x/\theta}}{\Gamma(\alpha)\theta^{\alpha}} dx = \left\{y = \frac{x}{\theta}\right\} = \int_0^{\infty} (y\theta)^k \frac{(y\theta)^{\alpha-1} e^{-y}}{\Gamma(\alpha)\theta^{\alpha}} dy = \frac{\theta^k}{\Gamma(\alpha)} \Gamma(\alpha + k) < \infty \quad \forall k$$

Для распределения Парето

$$\mu'_k = \int_0^{\infty} x^k \frac{\alpha\theta^{\alpha}}{(x + \theta)^{\alpha+1}} dx = \{y = x + \theta\} = \int_{\theta}^{\infty} (y - \theta)^k \frac{\alpha\theta^{\alpha}}{y^{\alpha+1}} dy = \alpha\theta^{\alpha} \int_{\theta}^{\infty} \sum_{j=0}^k \binom{k}{j} y^{j-k-1} (-\theta)^{k-j} dy$$

интеграл сходится только если $\alpha + 1 - k > 1$, т.е. $k < \alpha$

Согласно этой классификации, распределение Парето имеет тяжелый хвост, а гамма распределение имеет легкий хвост.

Заметим, что если распределение не имеет всех своих положительных моментов, то у него нет п.ф.м.). Однако обратное неверно. Логнормальное распределение не имеет п.ф.м., несмотря на то, что все его положительные моменты конечны.

3.1.2 Предельное соотношение хвостов

Широко используемый признак того, что одно распределение имеет более тяжелый хвост, чем другое с тем же средним значением, состоит в том, что отношение двух функций выживания (с распределением с более тяжелым хвостом в числителе) стремится к бесконечности.

$$\lim_{x \rightarrow \infty} \frac{S_1(x)}{S_2(x)} = \lim_{x \rightarrow \infty} \frac{S'_1(x)}{S'_2(x)} = \lim_{x \rightarrow \infty} \frac{-f_1(x)}{-f_2(x)} = \lim_{x \rightarrow \infty} \frac{f_1(x)}{f_2(x)}$$

Это значит, что распределение в числителе приписывает заметно большие значения вероятностей большим значениям x .

3.1.3 Использование уровня риска для оценки хвоста

Уровень риска также содержит информацию о характере хвоста распределения. Распределения с убывающим (невозрастающим) уровнем риска имеют *тяжелый хвост*. Распределения с возрастающим (неубывающим) уровнем риска имеют *легкие хвосты*.

Экспоненциальное распределение имеет постоянный уровень риска, поэтому он является и возрастающим, и убывающим.

Для распределений с монотонным уровнем риска экспоненциальное распределение служит границей, разделяющие тяжелые и легкие хвосты.

Сравнение распределений также делается на основе уровня риска. Например, распределение имеет более легкий хвост, если уровень риска возрастает быстрее. Часто такая классификация и сравнения производятся только в правой части распределения, т.е. только для больших значений

3.1.4 Использование среднего превышения для оценки тяжести хвоста

Среднее превышение ущерба также дает информацию о тяжести хвоста. Если среднее превышение ущерба возрастает по d , распределение имеет тяжелый хвост. Если среднее превышение ущерба убывает по d , распределение имеет легкий хвост. Сравнение распределений осуществляется на основе того, является ли среднее превышение возрастающим или убывающим.

$$\frac{S(y+d)}{S(d)} = \frac{\exp[-\int_0^{y+d} h(x)dx]}{\exp[-\int_0^d h(x)dx]} = \exp[-\int_d^{y+d} h(x)dx] = \exp[-\int_0^y h(d+t)dt]$$

Поэтому, если уровень риска убывает, то для фиксированного $y \int_0^y h(d+t)dt$ также убывает по d , значит $\frac{S(y+d)}{S(d)}$ возрастающая функция по d . Но среднее превышение ущерба можно представить в виде:

$$e(d) = \frac{\int_0^\infty S(x)dx}{S(d)} = \int_0^\infty \frac{S(y+d)}{S(d)} dy$$

Таким образом, если уровень риска убывает, то среднее превышение ущерба $e(d)$ возрастает по d . И наоборот, если уровень риска возрастает, то среднее превышение является убывающей функцией. Обратное неверно.

При $d \rightarrow \infty$, $S(d)$, $\int_d^\infty S(x)dx$ стремятся к 0. Поэтому, применяя правило Лопиталя, получаем

$$\lim_{d \rightarrow \infty} e(d) = \lim_{d \rightarrow \infty} \frac{\int_d^\infty S(x)dx}{S(d)} = \lim_{d \rightarrow \infty} \frac{-S(d)}{-f(d)} = \lim_{d \rightarrow \infty} \frac{1}{h(d)}$$

Это предельное соотношение используется в случае, когда вид $F(x)$ достаточно сложный.

3.2 Равновесные распределения

Рассмотрим равновесные распределения (интегрированное хвостовое распределение). Для положительной случайной величины с $S(0) = 1$ при $d = 0$ $E[X] = \int_0^\infty S(x)dx$, т.е.

$1 = \int_0^\infty \frac{S(x)}{E[X]}dx$. Значит, $f_e(x) = \frac{S(x)}{E[X]}$ является плотностью распределения.

Функция выживания равна:

$$S_e[X] = \int_x^\infty f_e(t)dt = \frac{\int_x^\infty S(t)dt}{E[X]}, \quad x \geq 0$$

Уровень риска, соответствующий равновесному распределению

$$h_e(x) = \frac{f_e(x)}{S_e(x)} = \frac{S(x)}{\int_x^\infty S(t)dt} = \frac{1}{e(x)}$$

Таким образом, величина, обратная к среднему превышению, является уровнем риска. Этот факт можно использовать, чтобы показать что функция среднего превышения ущерба однозначно характеризует исходное распределение. Заметим, что

$$f_e(x) = h_e(x)S_e(x) = h_e(x)\exp\left(-\int_0^x h_e(t)dt\right)$$

или

$$S(x) = \frac{e(0)}{e(x)}\exp\left(-\int_0^x \frac{1}{e(t)}dt\right)$$

так как $e(0) = E[x]$.

Равновесное распределение позволяет получить и другие соотношения между уровнем риска, средним превышением и тяжестью хвоста. Предполагая, что $S(0) = 1$ и, значит, $e(0) = E[x]$, получаем $\int_x^\infty S(t)dt = e(0)S_e(x)$. С другой стороны, $\int_x^\infty S(t)dt = e(x)S(x)$, значит

$$\frac{e(x)}{e(0)} = \frac{S_e(x)}{S(x)}$$

Если среднее превышение возрастает (т.е. уровень риска убывает), то $e(x) \geq e(0)$, значит,

$$S_e(x) \geq S(x)$$

$$\int_0^\infty S_e(x)dx \geq \int_0^\infty S(x)dx$$

но $E[X] = \int_0^\infty S(x)dx$ и также

$$\int_0^\infty S_e(x)dx = \int_0^\infty x f_e(x)dx = \frac{1}{E[X]} \int_0^\infty x S(x)dx = \text{по частям} = \frac{E[X^2]}{2E[X]}$$

Значит, $\frac{E[X^2]}{2E[X]} \geq E[X]$ или $Var[X] \geq E^2[X]$.

Это значит, что квадрат коэффициента вариации и сам коэффициент вариации не менее 1, если $e(x) \geq e(0)$. Обратное неравенство приводит к верхней границе 1 для коэффициента вариации при $e(x) \leq e(0)$, что получается, если среднее превышение убывает или уровень риска возрастает. Таким образом, значение коэффициента вариации оказывается также связано с тяжестью хвоста.

4 Меры риска. Согласованность меры риска. Процентиль. VaR, TVaR

4.1 Меры риска

Вероятностные модели обеспечивают возможность описать подверженность риску. Уровень подверженности риску часть описывает одним числом или небольшим множеством чисел. Эти числа часто называются ключевыми индикаторами риска. Такие ключевые индикаторы информируют актуариев и других риск-менеджеров о степени, в которой компания подвержена определенным типам риска. В частности, VaR (Value at Risk, величина риска, стоимость под риском, денежная оценка риска) является квантилью распределения суммарных убытков. Риск-менеджеры имеют возможность учесть вероятность неблагоприятного исхода. Это может быть выражено посредством VaR для определенного уровня вероятности. VaR может быть использован для определения количества капитала, требующегося для того, чтобы противостоять неблагоприятным исходам.

VaR имеет определенные недостатки. Более информативной и более полезной мерой риска является TVaR (Tail Value at Risk, хвостовое значение риска). Оно появилось независимо во множестве областей и имеет различные названия, в том числе Conditional Value at Risk (CVaR), Average Value at Risk (AVaR), Conditional Tail Expectation (CTE) и Expected Shortfall (ES).

4.2 Согласованность меры риска

Мера риска является отображением множества случайных величины, представляющих убытки, ассоциированные с риском, на числовую прямую (множество действительных чисел).

Мера риска дает одно число, которое предназначено для количественной оценки подверженности риску. Например, стандартное отклонение или число, кратное стандартному отклонению, является мерой риска, так как обеспечивает оценку неопределенности. Это особенно хорошо подходит для стандартного нормального распределения. В финансовой сфере размер убытка, для которого вероятность превышения мала (например, 0.05%), является простой мерой риска.

Меры риска обозначаются через $\rho(X)$. Удобно рассматривать $\rho(X)$ как количество активов, требующихся для защиты от неблагоприятного исхода риска X .

Комбинирование рисков важно в изучении потребности в капитале страховой компании, когда рассматривается совокупная подверженности риску. Страховая компания может иметь разные страховые продукты, например, личное страхование жизни, страхование домовладельцев, автомобилей, групповое страхование жизни и здоровья. Когда меры риска применяются к отдельным типам страхования результат должен в некотором смысле соответствовать тому, что получается, когда мера риска применяется ко всей компании.

Далее будем рассматривать переменные X и Y в качестве переменных убытка для двух подразделений и $X + Y$ как случайную величину потерь для организации, образованной этими двумя подразделениями.

4.3 Условия согласованности

Будем считать, что если X и Y являются элементами множества, то cX и $X + Y$ также принадлежат этому множеству. **Определение.** Мера риска $\rho(X)$ называется согласованной, если для двух случайных величин X и Y , отвечающих потерям, выполнены следующие свойства:

1. полуаддитивность: $\rho(X + Y) \leq \rho(X) + \rho(Y)$

Полуаддитивность означает, что мера риска (и, значит, капитал, требуемый для его поддержки) для комбинации двух рисков не больше суммы мер каждого риска в отдельности. Она отражает тот факт, что от комбинирования рисков должно быть какое-то преимущество. Иначе компании будут иметь преимущества дезагрегации в более мелкие.

2. монотонность: если $X \leq Y$, то $\rho(X) \leq \rho(Y)$

Монотонность означает, что если один риск всегда приводит к большим потерям, чем другой, то и мера риска в первом случае должна быть больше.

3. положительная однородность: $\forall c > 0, \rho(cX) = c\rho(X)$

Положительная однородность означает, что мера риска не зависит от единицы измерения (валюты). Или, эквивалентно, удвоение подверженности риск требует удвоения капитала. Это разумно, так как удвоение не обеспечивает диверсификации.

4. трансляционная инвариантность: $\forall c, \rho(X + c) = \rho(X) + c$

Трансляционная инвариантность означает, что добавления риска нет там, где нет добавочной неопределенности. В частности, при $X = 0$, величина активов требующихся для определенного дохода, в точности равна величине этого дохода.

4.4 Процентиль

Определение. 100р *процентиль* случайной величины это любое значение π_p , такое, что

$$F(-\pi_p) \leq p \leq F(\pi_p)$$

Определение. 50 процентиль, $\pi_{0.5}$, называется *медианой*.

Квантиль принимает значения от 0 до 1, являясь по сути функцией, обратной к функции распределения.

Процентиль принимает значения от 0 до 100, выражая вероятность в процентах (0.7 квантиль и 70 процентиль это одно и то же).

Если функция распределения имеет значение p для единственного значения x то процентиль определяется однозначно.

Если функция распределения имеет скачок от значения меньше p до значения больше p то процентиль определяет в месте этого скачка.

Процентиль не определяется однозначно, если функция распределения является константой со значением p для некоторого диапазона случайных величин В этом случае любое значение из диапазона является процентилем.

4.5 Value at Risk (VaR)

Value at Risk (VaR) – стоимость под риском, стоимостная оценка риска. Это количество капитала, требующееся для обеспечения, с высокой степенью уверенности, того, что предприятие избежит технического разорения. Степень уверенности выбирается произвольно.

VaR широко распространен в финансовом риск менеджменте трейдинговых рисков на фиксированном (обычно довольно коротком) периоде. В таком случае нормальное распределение часто используется для описания выигрышей и убытков. Если распределения доходов и убытков ограничено нормальным распределением, VaR удовлетворяет условиям согласованности.

Однако нормальное распределение обычно не используется для описания страховых рисков, которые, как правило, являются асимметричными. Поэтому использование VaR затруднено из-за отсутствия полуаддитивности.

VaR для случайной величины X - это $100p$ процентиль распределения X , обозначаемая через $VaR_p[X] = \pi_p$

VaR часто называют квантильной мерой риска, так как когда страховой компании доступно такое количество капитала, то он может предотвратить $100p\%$ возможных исходов убытков.

Определение. Пусть X обозначает случайную величину ущерба. Value at Risk величины X на уровне $100p\%$ обозначаемая через $VaR_p[X]$ или π_p , является $100p$ процентилем распределения X :

$$VaR_p[X] = \inf_{x \geq 0} (x | F_X(x) \geq p), \quad 0 < p < 1$$

Для непрерывных распределений можно просто написать $VaR_p[X]$ для X как величину π_p , удовлетворяющую

$$P(X > \pi_p) = 1 - p$$

VaR не удовлетворяет одному из условий согласованности полуаддитивности.

4.6 Tail Value at Risk (TVaR)

Определение. Пусть X обозначает случайную величину убытка. Tail Value at Risk (хвостовое значение риска) X на уровне защиты $100p\%$, обозначаемое через $TVaR_p[X]$, равно среднему по всем значениям VaR выше уровня p , $0 < p < 1$, т.е.

$$TVaR_p[X] = \frac{\int_p^1 VaR_u[X] du}{1 - p}$$

TVaR является согласованной мерой риска. Альтернативная формула выглядит следующим образом:

$$TVaR_p[X] = VaR_p[X] + \frac{1 - F_X(VaR_p[X])}{1 - p} (E[X | X > VaR_p[X]] - VaR_p[X])$$

Если X непрерывна в $VaR_p[X]$, то $F_X(VaR_p[X]) = p$, и формула упрощается:

$$TVaR_p[X] = E[X|X > VaR_p[X]] = E[X|X > \pi_p] = \pi_p + E[X - \pi_p|X > \pi_p] = VaR_p[X] + e(\pi_p)$$

где $e(\pi_p)$ функция среднего превышения ущерба, вычисленная в 100 p процентилях.

VaR рассматривается как мера типа «все или ничего» в том смысле, что если происходит экстремальное событие в превышении порога VaR, то нет капитала, чтобы компенсировать убытки. Также VaR квантиль в TVaR определяет «плохие времена», т.е. ситуации, когда убытки превышают порог VaR. Тогда TVaR обеспечивает среднее превышение убытков в «плохие времена», т.е. когда порог VaR для «плохих времен» превышен.

На практике получение численных значений VaR или TVaR может быть сделано непосредственно на основе данных или из формул распределений. При оценке VaR на основе данных могут быть использованы методы получения квантилей на основе эмпирических распределений.

Так как TVaR является ожидаемым значением наблюдений, которые больше заданного порога, естественной будет оценка среднего значения наблюдений, превышающих порог. Однако следует предостеречь от использования этого подхода на практике до тех пор, пока не получено достаточно большое число наблюдений, превышающих пороговое значение.

В случаях, когда нет достаточного числа наблюдений, предпочтительнее получить модель распределения по всем наблюдениям или хотя бы по всем наблюдениям, превышающим относительно низкий порог.

Используя следующее соотношение, можно вычислить значения VaR и TVaR непосредственно из подходящего распределения.

$$\begin{aligned} TVaR_p[X] &= E[X|X > \pi_p] = \pi_p + \frac{\int_{\pi_p}^{\infty} (x - \pi_p) f(x) dx}{1 - p} = \\ &= \pi_p + \frac{\int_{-\infty}^{\infty} (x - \pi_p) f(x) dx - \int_{-\infty}^{\pi_p} (x - \pi_p) f(x) dx}{1 - p} = \\ &= \pi_p + \frac{E[x] - \int_{-\infty}^{\pi_p} x f(x) dx - \pi_p(1 - F(\pi_p))}{1 - p} = \\ &= \pi_p + \frac{E[x] - E[\min(X, \pi_p)]}{1 - p} = \pi_p + \frac{E[x] - E[X \cap \pi_p]}{1 - p} \end{aligned}$$

5 Смеси распределений (конечные, переменные)

5.1 Параметрические и масштабируемые распределения

Определение. Параметрическим распределением называется множество функций распределения, каждая из которых определяется указанием одного или более значений, называемых параметрами.

Число параметров фиксировано и конечно. Самое известное параметрическое распределение - нормальное с параметрами μ и σ^2 .

Определение. Параметрическое распределение называется масштабируемым, если случайная величина, подчиняющаяся этому распределению, умноженная на положительную константу, также подчиняется этому распределению.

5.2 Конечные смеси распределений

Определение. Пусть $X_1, X_2, \dots, X_k$ - множество случайных величин с соответствующими функциями распределений $F_{X_1}, F_{X_2}, \dots, F_{X_k}$. Случайная величина Y является k -точечной смесью случайных величин $X_1, X_2, \dots, X_k$, если ее функция распределения имеет вид

$$F_Y(y) = a_1 F_{X_1}(y) + a_2 F_{X_2}(y) + \dots + a_k F_{X_k}(y)$$

где $a_j > 0$ и $a_1 + a_2 + \dots + a_k = 1$

5.3 Переменные смеси

Если число распределений в смеси заранее не известно, то k становится параметром.

Определение. Смесью распределений с переменным числом компонент имеет функцию распределения следующего вида:

$$F(x) = \sum_{j=1}^K a_j F_j(x), \quad \sum_{j=1}^K a_j = 1, \quad a_j > 0, \quad j = 1, 2, \dots, K, \quad K = 1, 2, \dots$$

Такие модели называются полупараметрическими, они сложнее параметрических и проще непараметрических.

Если все компоненты имеют одно и то же параметрическое распределение (но разные значения параметров), результирующее распределение называется переменной смесью соответствующих распределений.

5.4 Смешивание

Концепция смешивания может быть расширена со смеси конечного числа случайных переменных на несчетное число. Далее функция плотности $f_\Lambda(\lambda)$ играет роль дискретных вероятностей a_j в k -точечной смеси.

Теорема. Пусть X имеет плотность распределения $f_{X|\Lambda}(x|\lambda)$ и функцию распределения $F_{X|\Lambda}(x|\lambda)$, где λ - параметр распределения X . Пусть λ реализация случайной

величины Λ с плотностью распределения $f_\Lambda(\lambda)$. Тогда безусловная функция плотности случайной величины X

$$f_X(x) = \int f_{X|\Lambda}(x|\lambda) f_\Lambda(\lambda) d\lambda$$

где интеграл берется по всем значениям λ , имеющим положительную вероятность. Функция смеси распределений

$$F_X(x) = \int_{-\infty}^x \int f_{X|\Lambda}(y|\lambda) f_\Lambda(\lambda) d\lambda dy = \int \int_{-\infty}^x f_{X|\Lambda}(y|\lambda) f_\Lambda(\lambda) dy d\lambda = \int F_{X|\Lambda}(x|\lambda) f_\Lambda(\lambda) d\lambda$$

Моменты смеси распределений:

$$E[X^k] = E[E[X^k|\Lambda]]$$

$$Var[X] = E[Var[X|\Lambda]] + Var[E[X|\Lambda]]$$

Как правило, смесь распределений, имеет тяжелые хвосты, поэтому такой способ является удобным при необходимости учитывать вероятности редких событий.

Если $f_{X|\Lambda}(x|\lambda)$ имеет убывающую функцию уровня риска для всех λ , тогда смесь распределений также имеет убывающую функцию уровня риска.

6 Получение новых распределений (умножение на константу, возведение в степень, экспоненцирование)

6.1 Умножение на константу

Теорема. Пусть X - непрерывная случайная величина с плотностью распределения $f_X(x)$ и функцией распределения $F_X(x)$. Пусть $Y = \theta X$, $\theta > 0$. Тогда $F_Y(y) = F_X(\frac{y}{\theta})$

$$f_Y(y) = \frac{1}{\theta} f_X\left(\frac{y}{\theta}\right)$$

θ - параметр масштаба.

6.2 Возведение в степень

Теорема. Пусть X - непрерывная случайная величина с плотностью распределения $f_X(x)$ и функцией распределения $F_X(x)$, $F_X(0) = 0$. Пусть $Y = X^{1/\tau}$. Тогда:

если $\tau > 0$, $F_Y(y) = F_X(y^\tau)$, $f_Y(y) = \tau y^{\tau-1} f_X(y^\tau)$, $y > 0$

если $\tau < 0$, $F_Y(y) = 1 - F_X(y^\tau)$, $f_Y(y) = -\tau y^{\tau-1} f_X(y^\tau)$.

Доказательство.

Если $\tau > 0$, $F_Y(y) = P(X \leq y^\tau) = F_X(y^\tau)$.

Если $\tau < 0$, $F_Y(y) = P(X \geq y^\tau) = 1 - F_X(y^\tau)$.

Так как обычно параметры считаются положительными, то при $\tau < 0$ после замены переменных получаем: $F_Y(y) = 1 - F_X(y^{-\tau})$, $f_Y(y) = \tau y^{-\tau-1} f_X(y^{-\tau})$

Определение. При возведении в степень, если $\tau > 0$, полученное распределение называется преобразованным; если $\tau = -1$, то обратным; если $\tau < 0$ (но не равно 1), то обратным преобразованным.

6.3 Экспоненцирование

Теорема. Пусть X - непрерывная случайная величина с плотностью распределения $f_X(x)$ и функцией распределения $F_X(x)$, $f_X(x) > 0$, $\forall x$. Пусть $Y = e^X$. Тогда для $y > 0$:

$$F_Y(y) = F_X(\ln(y)), \quad f_Y(y) = \frac{1}{y} f_X(\ln(y))$$

7 Склейка распределений

7.0.1 Склейка (splicing)

Еще один метод получения новых распределений склейка. Похож на смешивание в том смысле, что два или более отдельных процесса отвечают за порождение убытков.

При смешивании различные процессы действуют на подмножества x совокупности. При склейке процессы различаются в зависимости от величины потерь (на разных интервалах действуют разные модели).

Определение. k компонентное склеенное распределение имеет функцию плотности, которая представима в следующем виде:

$$f_X(x) = \begin{cases} a_1 f_1(x), & c_0 < x < c_1 \\ a_2 f_2(x), & c_1 < x < c_2 \\ \dots & \\ a_k f_k(x), & c_{k-1} < x < c_k \end{cases}$$

Для $j = 1, \dots, k$ $a_j > 0$ и $f_j(x)$ является функцией плотности, заданной на (c_{j-1}, c_j) , $a_1 + \dots + a_k = 1$.

7.0.2 Свойства склейки

Наличие функций плотности и коэффициентов гарантирует, что результат является функцией плотности.

Частая причина использования склейки параметрических моделей: поведение хвоста может быть несовместимым с поведением при небольших потерях.

Например, опыт (основанный на знаниях, выходящих за рамки того, что доступно на текущем, возможно, небольшом наборе данных) может указывать на то, что хвост соответствует распределению Парето, но в окрестности моды распределение больше соответствует логнормальному распределению.

Другой вариант. Имеется большой объем данных небольших значений, но ограниченный для больших значений. Тогда можно использовать эмпирическое распределение до определенной точки и параметрическую модель за пределами этого значения.

Другой способ построения склейки состоит в том, чтобы использовать стандартные распределения в диапазоне от c_0 до c_k .

Пусть $g_j(x)$ - j -я функция плотности. Тогда в Определении можно заменить $f_j(x)$ на $g_j(x)/[G(c_j) - G(c_{j-1})]$.

Такой способ позволяет сделать точки склейки параметрами, которые также могут быть оценены.

8 Экспоненциальное семейство распределений

8.1 Линейное экспоненциальное семейство

Определение. Случайная величина X (дискретная или непрерывная) имеет распределение из линейного экспоненциального семейства, если плотность ее распределения может быть параметризована с помощью параметра θ и выражена в виде

$$f(x; \theta) = \frac{p(x)e^{r(\theta)x}}{q(\theta)}$$

Функция $p(x)$ зависит только от x и функция $q(\theta)$ является нормировочной. Также носитель случайной величины не должен зависеть от θ . Параметр $r(\theta)$ называется каноническим параметром распределения.

Нормальное распределение входит в линейное экспоненциальное семейство. Гамма распределение принадлежит линейному экспоненциальному семейству

8.1.1 Математическое ожидание линейного экспоненциального семейства

Для этого сначала заметим, что: $\ln(f(x, \theta)) = \ln(p(x)) + r(\theta)x - \ln(q(\theta))$.

$$\frac{\partial}{\partial \theta} f(x, \theta) = [r'(\theta)x - \frac{q'(\theta)}{q(\theta)}] f(x, \theta)$$

$$\int \frac{\partial}{\partial \theta} f(x, \theta) dx = r'(\theta) \int x f(x, \theta) dx - \frac{q'(\theta)}{q(\theta)} \int f(x, \theta) dx$$

$$\frac{\partial}{\partial \theta} \int f(x, \theta) dx = r'(\theta) \int x f(x, \theta) dx - \frac{q'(\theta)}{q(\theta)} \int f(x, \theta) dx$$

$$\int f(x, \theta) dx = 1, \quad \int x f(x, \theta) dx = E[X], \quad E[x] = \mu(\theta) = \frac{q'(\theta)}{r'(\theta)q(\theta)}$$

8.1.2 Дисперсия линейного экспоненциального семейства

$$\frac{\partial}{\partial \theta} f(x, \theta) = [r'(\theta)x - \frac{q'(\theta)}{q(\theta)}] f(x, \theta)$$

$$\frac{\partial}{\partial \theta} f(x, \theta) = r'(\theta)[x - \mu(\theta)] f(x, \theta)$$

$$\frac{\partial^2}{\partial \theta^2} f(x, \theta) = r''(\theta)[x - \mu(\theta)] f(x, \theta) - r'(\theta)\mu'(\theta) f(x, \theta) + [r'(\theta)]^2 [x - \mu(\theta)]^2 f(x, \theta)$$

$$\int \frac{\partial^2}{\partial \theta^2} f(x, \theta) dx = \int r''(\theta)[x - \mu(\theta)] f(x, \theta) dx - \int r'(\theta)\mu'(\theta) f(x, \theta) dx + \int [r'(\theta)]^2 [x - \mu(\theta)]^2 f(x, \theta) dx$$

$$\frac{\partial^2}{\partial \theta^2} \int f(x, \theta) dx = r''(\theta) \int [x - \mu(\theta)] f(x, \theta) dx - r'(\theta)\mu'(\theta) \int f(x, \theta) dx + [r'(\theta)]^2 \int [x - \mu(\theta)]^2 f(x, \theta) dx$$

$$0 = -r'(\theta)\mu'(\theta) + [r'(\theta)]^2 \int [x - \mu(\theta)]^2 f(x, \theta) dx$$

$$Var[X] = v(\theta) = \frac{\mu'(\theta)}{r'(\theta)}$$

Линейное экспоненциальное семейство включает также дискретные распределения Пуассона, биномиальное и отрицательное биномиальное.

9 Дискретные распределения. Смесь пуассоновских распределений

9.1 Дискретные распределения

Функция вероятности p_k обозначает вероятность того, что произойдет ровно k событий таких как претензии или убытки). Пусть N - случайная величина, представляющая число таких убытков. Тогда

$$p_k = P(N = k), k = 0, 1, 2, \dots$$

Производящая функция вероятности дискретной случайной величины N с функцией вероятности p_k равна

$$P(z) = P_N(z) = E[z^N] = \sum_{k=0}^N p_k z^k$$

Как и производящая функция моментов, производящая функция вероятности порождает моменты:

$$P'(1) = E[N], P''(1) = E[N(N-1)]$$

$$P^m(z) = E\left[\frac{d^m}{dz^m} z^N\right] = E[N(N-1)\dots(N-m+1)z^{N-m}] = \sum_{k=m}^{\infty} k(k-1)\dots(k-m+1)z^{k-m}p_k$$

$$P^m(0) = m!p_m, p_m = \frac{P^m(0)}{m!}$$

9.1.1 Распределение Пуассона

Функция вероятности: $p(k) \equiv P(Y = k) = \frac{\lambda^k}{k!} e^{-\lambda}$

Производящая функция: $P(z) = e^{\lambda(z-1)}, \lambda > 0$

Математическое ожидание: $E[N] = P'(1) = \lambda$

$E[N(N-1)] = P''(1) = \lambda^2$

Дисперсия: $Var[N] = E[N(N-1)] + E[N] - E^2[N] = \lambda^2 + \lambda - \lambda^2 = \lambda$

Теорема. Пусть $N_1, \dots, N_n$ независимые случайные величины, имеющие распределения Пуассона с параметрами $\lambda_1, \dots, \lambda_n$. Тогда $N = N_1 + \dots + N_n$ имеет распределение Пуассона с параметром $\lambda_1 + \dots + \lambda_n$.

Доказательство.

Для суммы случайных величин, распределенных по закону Пуассона, имеем

$$P_N(z) = \prod_{j=1}^n P_{N_j}(z) = \prod_{j=1}^n e^{\lambda_j(z-1)} = e^{\sum_{j=1}^n \lambda_j(z-1)} = e^{\lambda(z-1)}$$

где $\lambda = \lambda_1 + \dots + \lambda_n$. Как и производящая функция моментов, производящая функция вероятности единственна, поэтому N имеет распределение Пуассона с параметром λ .

Теорема. Пусть число событий N - случайная величина, имеющая распределение Пуассона со средним λ . Предположим, что каждое событие может быть отнесено к одному из m типов с вероятностями $p_1, \dots, p_m$ независимо от всех других событий. Тогда

число событий $N_1, \dots, N_m$, соответствующих типам событий $1, \dots, m$ взаимно независимые случайные величины, имеющие распределения Пуассона со средними $\lambda p_1, \dots, \lambda p_m$, соответственно.

Доказательство.

Для фиксированного $N = n$ условное совместное распределение $(N_1, \dots, N_m)$ является полиномом с параметрами $(n, p_1, \dots, p_m)$. Также для фиксированного $N = n$ условное распределение N_j бином с параметрами (n, p_j) . Совместная функция вероятности $(N_1, \dots, N_m)$

$$\begin{aligned} P(N_1 = n_1, \dots, N_m = n_m) &= P(N_1 = n_1, \dots, N_m = n_m | N = n) * P(N = n) = \\ &= \frac{n!}{n_1! n_2! \dots n_m!} p_1^{n_1} \dots p_m^{n_m} \frac{e^{-\lambda} \lambda^n}{n!} = \prod_{j=1}^m e^{-\lambda} \frac{(\lambda p_j)^{n_j}}{n_j!} \end{aligned}$$

где $n = n_1 + n_2 + \dots + n_m$. Аналогично функция вероятности для N_j равна

$$\begin{aligned} P(N_j = n_j) &= \sum_{n=n_j}^{\infty} P(N_j = n_j | N = n) P(N = n) = \sum_{n=n_j}^{\infty} \binom{n}{n_j} p_j^{n_j} (1 - p_j)^{n-n_j} \frac{e^{-\lambda} \lambda^n}{n!} = \\ &= e^{-\lambda} \frac{(\lambda p_j)^{n_j}}{n_j!} \sum_{n=n_j}^{\infty} \frac{(\lambda(1 - p_j))^{n-n_j}}{(n - n_j)!} = e^{-\lambda} \frac{(\lambda p_j)^{n_j}}{n_j!} e^{\lambda(1-p_j)} = e^{-\lambda p_j} \frac{(\lambda p_j)^{n_j}}{n_j!} \end{aligned}$$

9.2 Отрицательное биномиальное распределение как смесь пуассоновских

Один из способов получить отрицательное биномиальное распределение состоит в смеси пуассоновских. Предположим, что риск имеет распределенное по закону Пуассона число претензий с известным λ . Можно считать λ реализацией случайной величины Λ . Обозначим плотность распределения Λ через $u(\lambda)$, причем Λ может непрерывным или дискретным, функцию распределения обозначим через $U(\lambda)$.

9.3 Отрицательное биномиальное распределение

Определение. Функция вероятности отрицательного биномиального распределения

$$P(N = k) = p_k = \binom{k+r-1}{k} \left(\frac{1}{1+\beta} \right)^r \left(\frac{\beta}{1+\beta} \right)^k, \quad k = 0, 1, 2, \dots, r > 0, \beta > 0$$

Биномиальные коэффициенты равны:

$$\binom{x}{k} = \frac{x(x-1)\dots(x-k+1)}{k!}$$

k является целым, x может быть любым действительным. При $x > k - 1$

$$\binom{x}{k} = \frac{\Gamma(x+1)}{\Gamma(k+1)\Gamma(x-k+1)}$$

Производящая функция вероятностей отрицательного биномиального распределения:

$$P(z) = [1 - \beta(z - 1)]^{-r}$$

Отсюда среднее и дисперсия отрицательного биномиального распределения:

$$E[N] = r\beta, \text{Var}[N] = r\beta(1 + \beta)$$

Геометрическое распределение является частным случаем отрицательного биномиального распределения с $r = 1$. Один из способов получить отрицательное биномиальное распределение состоит в смеси пуассоновских. Распределение Пуассона является предельным случаем отрицательного биномиального распределения.

9.4 Смесь пуассоновских распределений

Это то же смешивание распределений, которое было в непрерывном случае. В некоторых случаях это называется неопределенностью параметра. В Байесовском подходе распределение Λ называется априорным, а параметры его распределения называются гиперпараметрами. Роль распределения $u(\cdot)$ очень важна в доверительной теории. Когда параметр λ неизвестен, вероятность того, что будет ровно k претензий может быть записана как ожидаемое значение с той же вероятностью, но при условии $\Lambda = \lambda$, где мат. ожидание берется по распределению Λ . Согласно формуле полной вероятности,

$$p_k = P(N = k) = E[P(N = k|\Lambda)] = \int_0^\infty P(N = k|\Lambda = \lambda)u(\lambda)d\lambda = \int_0^\infty \frac{\lambda^k}{k!}e^{-\lambda}u(\lambda)d\lambda$$

Теперь предположим, что Λ имеет гамма распределение. Тогда

$$p_k = \int_0^\infty \frac{e^{-\lambda}\lambda^k}{k!} \frac{\lambda^{\alpha-1}e^{-\lambda/\theta}}{\theta^\alpha\Gamma(\alpha)}d\lambda = \frac{1}{k!} \frac{1}{\Gamma(\alpha)} \left(\frac{\theta}{1+\theta}\right)^k \left(\frac{1}{1+\theta}\right)^k \Gamma(k+\alpha) = \binom{k+\alpha-1}{k} \left(\frac{\theta}{1+\theta}\right)^k \left(\frac{1}{1+\theta}\right)^\alpha$$

Получили отрицательное биномиальное распределение.

10 Классы $(a, b, 0)$, $(a, b, 1)$, соотношение между ними

10.1 Класс $(a, b, 0)$

Определение. Пусть p_k функция вероятности дискретного распределения. Распределение входит в класс $(a, b, 0)$, если существуют такие a и b , что

$$p_k = (a + \frac{b}{k})p_{k-1}, k = 1, 2, \dots$$

Распределение Пуассона и отрицательное биномиальное распределение также принадлежат классу $(a, b, 0)$. Геометрическое распределение является частным случаем отрицательного биномиального, поэтому также принадлежит этому классу.

Рекурсивная формула может быть переписана в виде (при $p_{k1} > 0$)

$$k \frac{p_k}{p_{k-1}} = ak + b, k = 1, 2, \dots$$

Для Пуассона: $a = 0, b = \lambda, p_0 = e^{-\lambda}$.

Для биномиального: $a = -\frac{q}{1-q}, b = \frac{(m+1)q}{1-q}, p_0 = (1-q)^m$.

Для отрицательного биномиального: $a = \frac{\beta}{1+\beta}, b = \frac{(r-1)\beta}{1+\beta}, p_0 = (1+\beta)^{-r}$.

Для геометрического: $a = \frac{\beta}{1+\beta}, b = 0, p_0 = (1+\beta)^{-1}$.

Выражение слева линейно по k . В таблице видно, что наклон прямой $a = 0$ Для распределения Пуассона, $a < 0$ для биномиального распределения, $a > 0$ для отрицательного биномиального распределения. Это свойство позволяет графически определить, какое из распределений больше подойдет под данные.

Иногда распределения неадекватно описывают характеристики некоторых наборов данных, встречающихся на практике.

Например, хвост отрицательного биномиального распределения может оказаться недостаточно тяжелым или распределения в классе $(a, b, 0)$ не могут отразить форму набора данных в какой-либо другой части распределения. Рассмотрим проблему плохого соответствия в левом конце распределения. Для данных о количестве страховых случаев нулевая вероятность - это вероятность того, что претензий нет. Существуют также ситуации, которые порождают большие вероятности при нулевом значении.

10.2 Класс $(a, b, 1)$

Определение. Пусть p_k функция вероятности дискретного распределения. Распределение принадлежит классу $(a, b, 1)$ если существуют константы a и b такие, что $p_k = (a + \frac{b}{k})p_{k-1}, k = 2, 3, 4, \dots$

В отличие от $(a, b, 0)$ этот класс начинается рекурсию с p_1 , а не с p_0 . Распределения от $k = 1$ до $k = \infty$ имеют ту же форму, что и класс $(a, b, 0)$ в том смысле, что вероятности те же с точностью до коэффициента пропорциональности, так как $\sum_{k=1}^{\infty} p_k$ можно привести к любому числу из интервала $[0; 1]$. Оставшаяся вероятность относится к $k = 0$.

10.3 Взаимосвязь классов (а, b, 0) и (а, b, 1)

Пусть $P(z) = \sum_{k=0}^{\infty} p_k z^k$ - производящая функция вероятности для распределения из класса (а, b, 0). Пусть $P^M(z) = \sum_{k=0}^{\infty} p_k^M z^k$ обозначает производящую функцию вероятности соответствующего распределения из класса (а, b, 1), т.е. $p_k^M = c p_k$; $k = 1, 2, 3, \dots$ и p_0^M - произвольное число. Тогда

$$P^M(z) = p_0^M + \sum_{k=1}^{\infty} p_k^M z^k = p_0^M + \sum_{k=1}^{\infty} p_k^M z^k = p_0^M + c(P(z) - p_0)$$

Так как $P^M(1) = P(1) = 1$, то $1 = p_0^M + c(1 - p_0)$, т.е. $c = \frac{1-p_0^M}{1-p_0}$ или $p_0^M = 1 - c(1 - p_0)$

Это соотношение необходимо, чтобы гарантировать, что сумма p_k^M равна 1. Тогда имеем

$$P^M(z) = p_0^M + \frac{1 - p_0^M}{1 - p_0}(P(z) - p_0) = (1 - \frac{1 - p_0^M}{1 - p_0})1 + \frac{1 - p_0^M}{1 - p_0}P(z)$$

Это средневзвешенное производящих функций вероятности вырожденного распределения и соответствующего (а, b, 0) распределения.

Более того, $p_k^M = \frac{1-p_0^M}{1-p_0} p_k$, $k = 1, 2, \dots$

11 Условная и безусловная франшизы. Расчет на убыток и на платеж

11.1 Влияние модификации покрытия на частоту и тяжесть убытков

Рассмотрим страховые приложения, связанные с усеченными и модифицированными случайными величинами.

В некоторых случаях необходимо различать убыток (нулевая вероятность имеет место при потере без выплаты) и платеж (случайная величина не определяется, когда платежа нет).

Y^L убытки, Y^P платежи

Если различие несущественно (например, при установке максимального платежа), верхний индекс опускается.

11.2 Безусловная франшиза

Страховые полисы часто имеют франшизу на убыток d .

Когда убыток x меньше d , страховщик ничего не платит. Когда выше d , платеж составляет $x - d$.

Определение. Безусловная франшиза заменяет случайную величину либо на превышение убытка, либо на цензурированную слева и сдвинутую случайную величину. Результат зависит от того, применяется франшиза к платежу или убытку.

Случайная величина платежа с франшизой:

$$Y^P = \begin{cases} \text{не определена, } X \leq d \\ X - d; X > d \end{cases}$$

Случайная величина убытка с франшизой:

$$Y^L = \begin{cases} 0; X \leq d \\ X - d; X > d \end{cases}$$

11.2.1 Функции для случайных величин с франшизой

Заметим, что $Y^P = Y^L | Y^L > 0$. Функция плотности для случайной величины превышения платежа:

$$f_{Y^P}(y) = \frac{f_X(y + d)}{S_X(d)}, y > 0$$

Функция выживания:

$$S_{Y^P}(y) = \frac{S_X(y + d)}{S_X(d)}$$

Функция распределения:

$$F_{Y^P}(y) = \frac{F_X(y + d) - F_X(d)}{1 - F_X(d)}$$

Уровень риска:

$$h_{Y^P}(y) = \frac{f_X(y+d)}{S_X(y+d)} = h_X(y+d)$$

$$f_{Y^L}(y) = \begin{cases} F_X(d), & X = 0 \\ f_X(y+d), & \end{cases}$$

$$S_{Y^L}(y) = S_X(y+d)$$

$$F_{Y^L}(y) = F_X(y+d)$$

11.3 Условная франшиза

Определение. Условная франшиза изменяет безусловную франшизу, добавляя размер вычета в случае положительного платежа.

Цензурированность слева со сдвигом и превышение ущерба здесь не применимы.

При расчете на убыток:

$$Y^L = \begin{cases} 0, & X \leq d \\ X; & X > d \end{cases}$$

При расчете на платеж:

$$Y^P = \begin{cases} \text{не определена}, & X \leq d \\ X; & X > d \end{cases}$$

11.3.1 Функции для случайных величин с франшизой

При расчете на убыток:

$$f_{Y^L}(y) = \begin{cases} F_X(d), & y = 0 \\ f_X(y), & y > d \end{cases}$$

$$S_{Y^L}(y) = \begin{cases} S_X(d), & 0 \leq y \leq d \\ S_X(y), & y > d \end{cases}$$

$$F_{Y^L}(y) = \begin{cases} F_X(d), & 0 \leq y \leq d \\ F_X(y), & y > d \end{cases}$$

$$h_{Y^L}(y) = \begin{cases} 0, & 0 < y < d \\ h_X(y), & y > d \end{cases}$$

При расчете на платеж:

$$f_{Y^P}(y) = \frac{f_X(y)}{S_X(d)}, \quad y > d$$

$$S_{Y^P}(y) = \begin{cases} 1, & 0 \leq y \leq d \\ \frac{S_X(y)}{S_X(d)}, & y > d \end{cases}$$

$$F_{Y^P}(y) = \begin{cases} 0, & 0 \leq y \leq d \\ \frac{F_X(y) - F_X(d)}{1 - F_X(d)}, & y > d \end{cases}$$

$$h_{Y^P}(y) = \begin{cases} 0, & 0 < y < d \\ h_X(y), & y > d \end{cases}$$

11.4 Ожидаемый убыток и ожидаемый платеж

Теорема. Для безусловной франшизы:

ожидаемая стоимость убытка $E[X] - E[X \cap d]$

ожидаемая стоимость платежа $\frac{E[X] - E[X \cap d]}{1 - F(d)}$

Для условной франшизы:

ожидаемая стоимость убытка $E[X] - E[X \cap d] + d(1 - F(d))$

ожидаемая стоимость платежа $\frac{E[X] - E[X \cap d]}{1 - F(d)} + d$

11.5 Влияние инфляции на безусловную франшизу

Теорема. Для безусловной франшизы d после равномерной инфляции $1+r$ ожидаемая стоимость убытка составляет $(1+r)(E[X] - E[X \cap \frac{d}{1+r}])$

Если $F(\frac{d}{1+r}) < 1$ то ожидаемая стоимость платежа получается делением на $1 - F(\frac{d}{1+r})$

Доказательство.

После инфляции случайная величина становится равной $Y = (1+r)X$.

По Теореме 1 (Лекция 3): $f_Y(y) = \frac{1}{1+r}f_X(\frac{y}{1+r})$ и $F_Y(y) = F_X(\frac{y}{1+r})$

$$E[X \cap d] = \int_0^d y f_Y(y) dy + d(1 - F_Y(d)) = \int_0^d \frac{y}{1+r} f_X(\frac{y}{1+r}) dy + d(1 - F_X(\frac{d}{1+r})) =$$

$$= \int_0^{\frac{d}{1+r}} (1+r)x f_X(x) dx + d(1 - F_X(\frac{d}{1+r})) = (1+r) \left(\int_0^{\frac{d}{1+r}} x f_X(x) dx + \right.$$

$$\left. + \frac{d}{1+r} (1 - F_X(\frac{d}{1+r})) \right) = (1+r) E[X \cap \frac{d}{1+r}]$$

Так как $E[Y] = (1+r)E[X]$, то получаем утверждение теоремы. Для расчетов на платеж утверждение следует из соотношения между функциями распределения Y и X .

12 Коэффициент устранения потерь

12.1 Моменты ограниченной случайной величины

$$E[(X \cap u)^k] = \begin{cases} \int_{-\infty}^u x^k f(x) dx + u^k [1 - F(u)] & \text{для непрерывной переменной} \\ \sum_{x_j \leq u} x_j^k p(x_j) + u^k [1 - F(u)] & \text{для дискретной переменной} \end{cases}$$

$$E[(X \cap u)^k] = - \int_{-\infty}^0 kx^{k-1} F(x) dx + \int_0^u kx^{k-1} S(x) dx$$

При $k=1$:

$$E[X \cap u] = - \int_{-\infty}^0 F(x) dx + \int_0^u S(x) dx$$

12.2 Безусловная франшиза

Страховые полисы часто имеют франшизу на убыток d .

Когда убыток x меньше d , страховщик ничего не платит. Когда выше d , платеж составляет $x - d$.

Определение. Безусловная франшиза заменяет случайную величину либо на превышение убытка, либо на цензурированную слева и сдвинутую случайную величину. Результат зависит от того, применяется франшиза к платежу или убытку.

Случайная величина платежа с франшизой:

$$Y^P = \begin{cases} \text{не определена, } X \leq d \\ X - d; X > d \end{cases}$$

12.3 Коэффициент устранения потерь

Определение. Коэффициент устранения потерь (Loss Elimination Ratio) - это отношение уменьшения ожидаемого платежа при применении безусловной франшизы к ожидаемому платежу без франшизы.

Без франшизы ожидаемый платеж $E[X]$. С безусловной франшизой $E[X] - E[X \cap d]$. Таким образом, коэффициент устранения потерь равен:

$$\frac{E[X] - (E[X] - E[X \cap d])}{E[X]} = \frac{E[X \cap d]}{E[X]}$$

13 Модификации полиса (сострахование, франшиза, лимит). Влияние инфляции

Инфляция увеличивает затраты, но наличие франшизы усиливает эффект инфляции. Во-первых, некоторые события, ранее порождавшие убытки ниже франшизы, могут привести к платежам. Во-вторых, относительный эффект инфляции усиливается, так как франшиза вычитается после инфляции. Например, предположим что событие ранее породило убыток 600. При франшизе 500 платеж равен 100. Инфляция 10 процентов увеличит убытки до 660 и платеж до 160, т.е. 60 процентов увеличение стоимости для страховщика.

13.1 Ограничения по полисам

Противоположностью франшизе является ограничение по полису. Как правило, ограничение по полису указано в контракте, в котором убыток меньше u страховщик оплачивает полностью, а убыток больше u оплачивается только в размере u .

Ограничение приводит к цензурированной справа случайной величине. Она имеет смешанное распределение и функцию плотности вида (Y - случайная величина после применения ограничения)

$$F_Y(y) = \begin{cases} F_X(y), & y < u \\ 1, & y \geq u \end{cases}$$
$$f_Y(y) = \begin{cases} f_X(y), & y < u \\ 1 - F_X(u), & y = u \end{cases}$$

Эффект инфляции вычисляется следующим образом.

Теорема. Для ограничения по полису в размере u после применения равномерной инфляции $1 + r$ ожидаемая стоимость составит

$$(1 + r)E[X \cap \frac{u}{1 + r}]$$

Доказательство.

Ожидаемая стоимость есть $E[Y \cap u]$. Далее см. доказательство Теоремы 5.

13.2 Сострахование, франшизы и лимиты

Последнее широко встречающаяся модификация покрытия сострахование. Страховая компания платит долю убытков α , держатель полиса платит оставшуюся часть. Если сострахование является единственной модификацией, то переменная убытков X заменяется на переменную платежа $Y = \alpha X$. Если помимо сострахования применяются безусловная франшиза, лимит и инфляция, получается следующая случайная величина

в расчете на убыток:

$$Y^L = \begin{cases} 0, & X < \frac{d}{1+r} \\ \alpha((1+r)X - d), & \frac{d}{1+r} \leq X < \frac{u}{1+r} \\ \alpha(u - d), & X \geq \frac{u}{1+r} \end{cases}$$

13.3 Средние платежи и убытки

Важен порядок применения модификации. Лимит полиса равен $\alpha(u - d)$ это максимальная сумма, которая может быть выплачена. u - максимально возможный убыток, выше которого никакие выплаты не производятся. Переменная в расчете на платеж Y^P не определена для $X < \frac{d}{1+r}$.

Теорема. Для переменной в расчете на убыток

$$E[Y^L] = \alpha(1+r)(E[X \cap \frac{u}{1+r}] - E[X \cap \frac{d}{1+r}])$$

Ожидаемое значение переменной на платеж

$$E[Y^P] = \frac{E[Y^L]}{1 - F_X(\frac{d}{1+r})}$$

Теорема. Для переменной на убыток

$$E[(Y^L)^2] = \alpha^2(1+r)^2(E[(X \cap u^*)^2] - E[(X \cap d^*)^2] - 2d^*E[X \cap u^*] + 2d^*E[X \cap d^*])$$

где $u^* = \frac{u}{1+r}$ и $d^* = \frac{d}{1+r}$. Для второго момента переменной на платеж выражение делится на $1 - F_X(d^*)$.

Доказательство.

Из определения Y^L : $Y^L = \alpha(1+r)((X \cap u^*) - (X \cap d^*))$, и поэтому

$$\begin{aligned} \frac{(Y^L)^2}{(\alpha(1+r))^2} &= ((X \cap u^*) - (X \cap d^*))^2 = (X \cap u^*)^2 + (X \cap d^*)^2 - 2(X \cap u^*)(X \cap d^*) = \\ &= (X \cap u^*)^2 - (X \cap d^*)^2 - 2(X \cap d^*)((X \cap u^*) - (X \cap d^*)) \end{aligned}$$

Последнее слагаемое можно переписать в виде

$$2(X \cap d^*)((X \cap u^*) - (X \cap d^*)) = 2d^*((X \cap u^*) - (X \cap d^*))$$

Чтобы это получить, заметим, что когда $X < d^*$, обе части равны 0; когда $d^* \leq X < u^*$, обе части равны $2d^*(X - d^*)$ и когда $X \geq u^*$, обе части равны $2d^*(u^* - d^*)$. Взятие мат. ожидания доказывает равенство.

14 Влияние франшизы на частоту требований

Важная компонента анализа эффекта модификации полиса связана с изменением распределения частоты платежей, когда франшиза налагается или изменяется. Когда есть франшиза, число платежей становится меньше. Чем меньше франшиза, тем больше платежей. Можно оценить количественно этот эффект, если предположить, что внесение изменений в покрытие не влияет на процесс, который приводит к убыткам, или на тип физического лица, которое приобретет страховку. Например, те, кто покупает франшизу в размере 250 долларов на покрытие ущерба автомобилю, могут (правильно) считать, что вероятность попасть в аварию у них меньше, чем у тех, кто покупает полное покрытие. Аналогичным образом, работодатель может обнаружить, что уровень нетрудоспособности снижается, когда работникам впервые несколько лет работы предоставляются сокращенные пособия.

Пусть X_j размер исходного убытка j -го и нет никаких модификаций. Пусть N^L число убытков. Предположим, что покрытие модифицировано так, что v - вероятность того, что убыток приведет к платежу.

Например, если есть франшиза d , то $v = P(X > d)$.

Определим индикаторную случайную величину I_j , $I_j = 1$, если j -й убыток приводит к платежу и $I_j = 0$ иначе. Тогда I_j имеет распределение Бернулли с параметром v , и производящая функция вероятности I_j равна $P_{I_j}(z) = 1 - v + vz$.

Случайная величина $N^P = I_1 + \dots + I_{N^L}$ определяет число платежей.

Если $I_1, I_2, \dots$ попарно независимы и также независимы от N^L , то N^P имеет составное распределение, в котором N^L является первичным, а распределение Бернулли вторичным. Таким образом,

$$P_{N^P}(z) = P_{N^L}(P_{I_j}(z)) = P_{N^L}(1 + v(z - 1))$$

В важном частном случае, когда распределение N^L зависит от параметра θ так, что

$$P_{N^L}(z) = P_{N^L}(z, \theta) = B(\theta(1 - z))$$

где $B(z)$ не зависит от θ . Тогда

$$P_{N^P}(z) = B(\theta(1 - 1 - vz + v)) = B(v\theta(1 - z)) = P_{N^L}(z, v\theta)$$

Из этого результата следует, что N^L и N^P оба принадлежат одному и тому же параметрическому семейству, и только параметр θ меняется.

15 Определение частоты убытков по частоте платежей

15.1 Определение частоты убытков по частоте платежей

В приложениях может возникать ситуация, когда необходимо определить распределение N^L из распределения N^P . Например, данные могут быть собраны для числе платежей при применении франшизы, и из этих данных необходимо оценить N^P .

Мы можем захотеть узнать распределение платежей, если франшизу уберут. Рассуждая как и ранее,

$$P_{N^L}(z) = P_{N^P}(1 - v^{-1} + zv^{-1})$$

Из этого результата следует, что формулы, полученные ранее, выполнены при замене v на $1/v$.

Все распределения из классов $(a, b, 0)$ и $(a, b, 1)$ удовлетворяют указанным выше условиям. Для их распределений существуют таблицы пересчета параметров при переходе от N^L к N^P . Если N^L имеет составное распределение, можем записать $P_{N^L}(z) = P_1(P_2(z))$ и поэтому

$$P_{N^P}(z) = P_{N^L}(1 + v(z - 1)) = P_1(P_2(1 + v(z - 1)))$$

Таким образом, N^P также будет иметь составное распределение, в котором вторичное распределение изменено как показано выше. Если вторичное распределение имеет $(a, b, 0)$ распределение, то есть таблица.

16 Модель коллективного риска

16.1 Модель коллективного риска

Страховые компании существуют благодаря объединению рисков. При страховании большого количества людей индивидуальные риски объединяются в совокупный риск.

Существуют два варианта построения модели суммы платежа по всем искам, произошедшим за фиксированный период времени и касающихся определенного набора договоров страхования. Первый вариант – учесть все сделанные платежи и затем их сложить. В таком случае получаем совокупные убытки в виде суммы S случайного числа слагаемых N отдельных платежей $X_1, X_2, \dots, X_N$:

$$S = X_1 + X_2 + \dots + X_N, N = 0, 1, 2, \dots,$$

где $S = 0$ при $N = 0$.

Определение 1. Модель коллективного риска имеет вид S при условии, что X_j независимы и одинаково распределенные случайные величины, т.е.

1. При $N = n$ случайные величины $X_1, X_2, \dots, X_n$ одинаково распределенные и независимые случайные величины.

2. При $N = n$ общее распределение случайных величин $X_1, X_2, \dots, X_n$ не зависит от n .

3. Распределение N не зависит от $X_1, X_2, \dots$

16.2 Модель индивидуального риска

Второй вариант - приписать случайную величину каждому договору страхования.

Определение 2. Модель индивидуального риска представляет суммарный ущерб как сумму

$$X_1 + X_2 + \dots + X_n$$

фиксированного числа n договоров страхования. Убытки для n договоров равны $(X_1, X_2, \dots, X_n)$, где X_j предполагаются независимыми, но необязательно одинаково распределенными. Распределение X_j обычно имеет вероятность в нуле, соответствующую вероятности отсутствия убытка или платежа по договору.

Модель индивидуального риска используется с целью сложения убытков или платежей от фиксированного числа договоров страхования или множества страховых убытков группового страхования жизни или здоровья для группы n работников. Сотрудники могут иметь различное покрытие (выплаты по договору страхования кратны окладу) и различные уровни вероятности убытков (разный возраст и состояние здоровья).

В частном случае, когда X_j одинаково распределены, модель индивидуального риска становится частным случаем модели коллективного риска с вырожденным распределением N , т.е. $P(N = n) = 1$

16.3 Моделирование суммарного убытка

Распределение суммы S получается из распределения N и общего распределения X_j . При таком подходе частота и убытки моделируются отдельно. Информация об этих распределениях используется для того, чтобы получить информацию об S . Альтернативой

к этому подходу является просто сбор информации о S (например, суммарные убытки за каждый месяц в течение нескольких месяцев) и использование какой-либо модели распределений. Моделирование распределения N и распределений X_j отдельно имеет ряд преимуществ:

1. Ожидаемое число исков меняется по мере изменения числа полисов. Необходимо учитывать рост объема бизнеса при прогнозировании будущего числа исков на основе данных за прошлые годы.
2. Последствия инфляции (общей и/или отдельных выплат) отражаются на убытках, понесенных застрахованными лицами, и выплатах страховых компаний по искам. Такие эффекты часто маскируются, когда в страховых полисах предусмотрены франшизы и используются лимиты полисов, не зависящие от инфляции.
3. Влияние изменения индивидуальных франшиз и лимитов легче учитывать, изменяя распределения убытков и частоту предъявления претензий по отдельности.
4. Данные, являющиеся неоднородными точки зрения франшиз и лимитов, могут быть объединены для получения гипотетического распределения размера убытков. Такой подход полезен, когда объединяются данные за несколько лет, в течение которых менялись условия страхования.
5. Модели убытков страхователей, расходов по претензиям страховщиков и расходов по претензиям перестраховщиков могут быть взаимосогласованными. Это полезно при определении последствий передачи риска перестраховщикам.
6. Форма распределения S зависит от формы распределений N и X . Понимание относительных форм важно для моделирования. Например, если распределение убытков имеет гораздо более тяжелый хвост, чем распределение частоты, форма хвоста распределения совокупных убытков будет определяться распределением индивидуальных убытков и будет нечувствительна к выбору распределения частоты.

Раздельный анализ убытков и частоты позволят построить более точную и гибкую. Если N представляет фактическое число убытков для застрахованного, то X_j может представлять:

1. убытки для застрахованного
2. выплаты страховщика
3. выплаты перестраховщика
4. франшизы, выплаченные застрахованным

В каждом случае S имеет разную интерпретацию, и распределение убытков должно быть с ним согласовано.

N - случайная величина числа исков (число исков, иски), частота;

X_j - случайные величины индивидуальных (отдельных) потерь (просто потери, убытки). Как правило, X_j соответствуют платежам, однако там, где нет необходимости различать платежи и убытки, используется термин потери или убытки.

S - случайная величина суммарных (совокупных) убытков.

17 Составное распределение, моменты

17.1 Составное (обобщенное) распределение

Пусть S – суммарный убыток, связанный с множеством из N наблюдаемых исков $X_1, X_2, \dots, X_N$, которые являются независимыми и одинаково распределенными. Чтобы построить модель суммарного убытка, необходимо

1. На основании данных смоделировать распределение N .
2. Найти подходящее общее распределение для X_j .
3. Используя две эти модели, получить распределение S .

Пусть первые две задачи решены. Случайная сумма случайного числа слагаемых $S = X_1 + X_2 + \dots + X_N$ имеет функцию распределения (составное распределение)

$$F_S(x) = P(S \leq x) = \sum_{n=0}^{\infty} p_n P(S \leq x | N = n) = \sum_{n=0}^{\infty} p_n F_X^{*n}(x),$$

где $F_X(x) = P(X \leq x)$ функция распределения X_j , а $p_n = P(N = n)$, $F_X^{*n}(x)$ – свертка степени n функций распределения X .

Свертка степени n функций распределения X . $F_X^{*n}(x)$:

$$F_X^{*0}(x) = \begin{cases} 0, & x < 0 \\ 1, & x \geq 0 \end{cases}$$

$$F_X^{*k}(x) = \int_{-\infty}^{\infty} F_X^{*(k-1)}(x-y) dF_X(y), k = 1, 2, \dots$$

17.2 Составное распределение

Хвост распределения (функция выживания) для $x \geq 0$.

$$S_S(x) = 1 - F_S(x) = \sum_{n=1}^{\infty} p_n (1 - F_X^{*n}(x))$$

Если X непрерывная случайная величина с нулевой вероятностью неположительных значений, то

$$F_X^{*k}(x) = \int_0^x F_X^{*(k-1)}(x-y) f_X(y) dy, k = 2, 3, \dots$$

$F_X^{*1}(x) = F_X(x)$. Дифференцируя, получаем функцию плотности

$$f_X^{*k}(x) = \int_0^x f_X^{*(k-1)}(x-y) f_X(y) dy, k = 2, 3, \dots$$

Если X непрерывна, то S имеет функцию плотности для $x > 0$

$$f_S(x) = \sum_{n=1}^{\infty} p_n f_X^{*n}(x)$$

и дискретную точку массы $P(S = 0) = p_0$ в $x = 0$. Заметим, что $P(S = 0) \neq f_S(0) = \lim_{x \rightarrow 0+} f_S(x)$

Если X имеет счетное распределение,

$$F_X^{*k}(x) = \sum_{y=0}^x F_X^{*(k-1)}(x-y) f_X(y), x = 0, 1, \dots; k = 2, 3, \dots$$

Соответствующая функция вероятности

$$f_X^{*k}(x) = \sum_{y=0}^x f_X^{*(k-1)}(x-y) f_X(y), x = 0, 1, \dots; k = 2, 3, \dots$$

Будем полагать, что $f_X^{*0}(0) = 1$ и $f_X^{*0}(x) = 0$ при $x \neq 0$. Тогда S имеет дискретное распределение с функцией вероятности

$$f_S(x) = P(S = x) = \sum_{n=0}^{\infty} p_n f_X^{*n}(x), x = 0, 1, \dots$$

Производящая функция вероятности S в силу независимости $X_1, X_2, \dots, X_n$ для фиксированного n .

$$\begin{aligned} P_S(z) &= E[z^S] = E[z^0]P(N = 0) + \sum_{n=1}^{\infty} E[z^{X_1+X_2+\dots+X_n} | N = n]P(N = n) = \\ &= P(N = 0) + \sum_{n=1}^{\infty} E[\prod_{j=1}^n z^{X_j}]P(N = n) = \sum_{n=0}^{\infty} P(N = n)(P_X(z))^n = E[P_X(z)^N] = P_N[P_X(z)] \end{aligned}$$

Для производящей функции моментов аналогично $M_S(z) = P_N(M_X(z))$

17.3 Моменты составного распределения

Если $P_N(z) = P_1(P_2(z))$, т.е. N само по себе является составным распределением, $P_S(z) = P_1(P_2(P_X(z)))$.

Моменты S легко выражаются через моменты N и X_j :

$$E[S] = \mu'_{S1} = \mu'_{N1} \mu'_{X1} = E[N]E[X]$$

$$Var[S] = \mu_{S2} = \mu'_{N1} \mu_{X2} + \mu_{N2} (\mu'_{X1})^2 = E[N]Var[X] + Var[N]E^2[X]$$

$$E[(S - E[S])^3] = \mu_{S3} = \mu'_{N1} \mu_{X3} + 3\mu_{N2} \mu'_{X1} \mu_{X2} + \mu_{N3} (\mu'_{X1})^3$$

18 Страхование стоп-лосс

Частно к суммарным убыткам применяется франшиза. Когда убытки возникают у страхователя, это называется страховым покрытием, а когда убытки возникают у страховой компании, это называется перестраховочным покрытием. Последний вариант является распространенным методом для страховой компании защитить себя от неблагоприятного года (в отличие от защиты от единичного, очень крупного иска).

Определение 3. Страхование суммарных убытков с применением франшизы называется страхованием стоп-лосс. Ожидаемая стоимость такого страхования называется нетто-премией стоп-лосс и она равна $E[(S - d)_+]$, где d - франшиза.

Для совокупного распределения

$$E[(S - d)_+] = \int_d^{\infty} (1 - F_S(x)) dx$$

Если распределение является непрерывным при $x > d$, то

$$E[(S - d)_+] = \int_d^{\infty} (x - d) f_S(x) dx$$

Аналогично для дискретного случая

$$E[(S - d)_+] = \sum_{x>d}^{\infty} (x - d) f_S(x)$$

18.1 Линейная интерполяция нетто-премии стоп-лосс

Если есть интервал, в котором значения суммарного убытка отсутствуют, то можно воспользоваться следующей формулой.

Теорема 1. Предположим, что $P(a < S < b) = 0$. Тогда для $a \leq d \leq b$

$$E[(S - d)_+] = \frac{b - d}{b - a} E[(S - a)_+] + \frac{d - a}{b - a} E[(S - b)_+]$$

Таким образом, нетто-премия стоп лосс может быть вычислена с помощью линейной интерполяции.

Доказательство.

По предположению $F_S(x) = F_S(a)$, $a \leq x < b$. Тогда

$$\begin{aligned} E[(S - d)_+] &= \int_d^{\infty} (1 - F_S(x)) dx = \int_a^{\infty} (1 - F_S(x)) dx - \int_a^d (1 - F_S(x)) dx = E[(S - a)_+] - \\ &\quad - \int_a^d (1 - F_S(a)) dx = E[(S - a)_+] - (d - a)(1 - F_S(a)) \end{aligned}$$

Полагая $d = b$,

$$E[(S - b)_+] = E[(S - a)_+] - (b - a)(1 - F_S(a))$$

$$1 - F_S(a) = \frac{E[(S - a)_+] - E[(S - b)_+]}{b - a}$$

Подставляя в предыдущее равенство, получаем требуемое равенство.

18.2 Дискретный случай

Теорема 2. Пусть $x = kh$, $P(S = kh) = f_k \geq 0$ для фиксированного $h > 0$, $k = 0, 1, \dots$, и $P(S = x) = 0$ для всех остальных x . Тогда при условии $d = jh$, $j \in N \cup 0$,

$$E[(S - d)_+] = h \sum_{m=0}^{\infty} (1 - F_S((m + j)h))$$

Доказательство.

$$\begin{aligned} E[(S - d)_+] &= h \sum_{x>d} (x - d) f_S(x) = \sum_{k=j}^{\infty} (kh - jh) f_k = h \sum_{k=j}^{\infty} \sum_{m=0}^{k-j-1} f_k = \\ &= h \sum_{m=0}^{\infty} \sum_{k=m+j+1}^{\infty} f_k = h \sum_{m=0}^{\infty} (1 - F_S((m + j)h)) \end{aligned}$$

В дискретном случае имеет место следующая рекурсия (при равноотстоящий дискретных значениях).

Следствие. В условиях Теоремы 2 $E[(S - (j + 1)h)_+] = E[(S - jh)_+] - h(1 - F_S(jh))$. Этим свойством легко воспользоваться, так как при $d = 0$, $E[(S - 0)_+] = E[S] = E[N]E[X]$

19 Обобщенное распределение Пуассона

В договоре группового страхования у каждого участника есть составная модель Пуассона для совокупных убытков, и необходимо найти распределение общих совокупных убытков. Аналогичным образом можно оценивать совокупные убытки от нескольких независимых направлений бизнеса. В таких случаях совокупное распределение можно получить без вычисления каждого члена группы.

Теорема 3. Пусть S_j имеет сложное распределение Пуассона с параметром λ_j , и распределение ущерба имеет функцию распределения $F_j(x)$, $j = 1, 2, \dots, n$. Также предположим, что $S_1, S_2, \dots, S_n$ независимы. Тогда $S = S_1 + \dots + S_n$ имеет сложное распределение Пуассона с параметром $\lambda = \lambda_1 + \dots + \lambda_n$ и распределением ущерба

$$F(x) = \sum_{j=1}^n \frac{\lambda_j}{\lambda} F_j(x)$$

Доказательство.

Пусть $M_j(t)$ – производящая функция моментов $F_j(x)$ для $j = 1, 2, \dots, n$. Тогда S_j имеет производящую функцию моментов

$$M_{S_j}(t) = E[e^{tS_j}] = e^{\lambda_j(M_j(t)-1)}$$

и в силу независимости S_j , производящая функция моментов S равна

$$M_S(t) = \prod_{j=1}^n M_{S_j}(t) = \prod_{j=1}^n e^{\lambda_j(M_j(t)-1)} = \exp\left(\sum_{j=1}^n \lambda_j M_j(t) - \lambda\right) = \exp\left(\lambda \left(\sum_{j=1}^n \frac{\lambda_j}{\lambda} M_j(t) - 1\right)\right)$$

Так как $\sum_{j=1}^n \frac{\lambda_j}{\lambda} M_j(t)$ – производящая функция $F(x) = \sum_{j=1}^n \frac{\lambda_j}{\lambda} F_j(x)$, $M_S(t)$ – производящая функция моментов составного распределения Пуассона.

20 Численное нахождение функции распределения совокупного ущерба

Вычисление функции распределения $F(x) = \sum_{n=0}^{\infty} p_n F_X^{*n}(x)$ или соответствующей функции плотности в общем случае сложная задача.

1. Можно использовать аппроксимирующее распределение для того, чтобы избежать непосредственного вычисления суммы сверток с весами-вероятностями. Этот метод был использован в примере 4, где для оценки параметров распределений использовался метод моментов. Метод прост, однако имеет значительные недостатки: не всегда есть возможность определить качество аппроксимации, выбор различных аппроксимирующих распределений может привести к существенно различающимся результатам, особенно в правом хвосте распределения; чем больше моментов используется, тем точнее будет аппроксимация, однако после четвертого момента не останется подходящих распределений (так как большего числа параметров нет); приближенное распределение может не отражать особенности истинного распределения.

Например, когда распределение потерь непрерывно и существует максимально возможная претензия (при наличии лимита полиса), распределение ущерба может иметь массу в максимальной точке. Истинное распределение совокупных убытков является смешанным с пиками, кратными максимуму, соответствующими 1, 2, 3 и т.д. максимальными искам. Когда эти пики являются большими, они оказывают значительное влияние на вероятности в окрестности таких кратных значений. Подобные скачки в функции распределения совокупных убытков не могут быть заменены гладким аппроксимирующим распределением.

2. Второй метод состоит в прямом вычислении соответствующей функции плотности. Наиболее сложной (вычислительно) частью является оценка свертки степени n распределений ущерба для $n = 2, 3, 4, \dots$. Свертки численно оцениваются с использованием

$$F_X^{*k}(x) = \int_{0-}^x F_X^{*(k-1)}(x-y) dF_X(y)$$

(в общем случае интеграл берется по всей прямой). Этот интеграл записан в форме Лебега-Стилтьеса, так как возможны скачки функции распределения $F_X(x)$ в нуле и других точках. (Достаточно интерпретировать $\int g(y) dF_X(y)$ как интеграл $g(y)f_X(y)$ по y в тех точках, где X имеет непрерывное распределение, а затем добавить $g(y_i)P(X = y_i)$ в тех точках, где $P(X = y_i) > 0$)

Для вычисления интеграла обычно требуются методы численного интегрирования. Из-за того, что интегрируемую функцию распределения необходимо вычислить для всех возможных значений x , такой подход быстро становится вычислительно затратным. Когда распределение ущерба дискретно, вычисления сводятся к многочисленным умножениям и сложениям. Для непрерывных распределений ущерба простым способом избежать технических проблем является замена распределения ущерба дискретным распределением, определяемым в точках, кратных удобной денежной единицы (например, 1000).

21 Дискретизация

Денежную единицу можно сделать достаточно маленькой, чтобы учесть скачки при максимальных страховых суммах. Пик должен быть кратен денежной единице, чтобы он был расположен в точке дискретизации. При уменьшении единицы измерения дискретная функция распределения должна приближаться к истинной функции распределения. Самый простой подход - округлить все суммы до ближайшего значения, кратного денежной единице (например, до ближайшей 1000).

Когда распределение ущерба определено для неотрицательных целых чисел $0, 1, 2, \dots$, вычисление $f_X^{*k}(x)$ для целого x требует $x + 1$ умножений. Чтобы получить распределение от $x = 0$ до $x = n$ требуется порядка $O(n^3)$ умножений. Когда максимальное значение n , для которого рассчитывается совокупное распределение утверждений, велико, количество вычислений быстро становится большим даже для быстрых компьютеров. В реальных приложениях n может легко достигать 1000, т.е. требуется примерно 10^9 умножений.

Далее, если $P(X = 0) > 0$ и распределение частот неограниченно, требуется бесконечное число вычислений для получения одной вероятности. Причина этого в том, что $F_X^{*n}(x) > 0$ для всех n и всех x , и поэтому сумма, определяющая функцию распределения, содержит бесконечное число членов. Когда $P(X = 0) = 0$, мы имеем $F_X^{*n}(x) = 0$ для $n > x$, и поэтому сумма будет иметь не более $x + 1$ положительных слагаемых.

22 Рекурсивный метод, рекурсия для распределения Пуассона

22.1 Рекурсивный метод

Предположим, что распределение ущерба $f_X(x)$ определено в $x = 0, 1, 2, \dots, m$, представляя значения, кратные некоторой денежной единице. Число m - наибольший возможный платеж; может быть бесконечным. Предположим, что распределение частоты p_k является представителем класса $(a, b, 1)$ и поэтому

$$p_k = (a + \frac{b}{k})p_{k-1}, k = 2, 3, \dots$$

Теорема 4. Для класса $(a, b, 1)$

$$f_S(x) = \frac{(p_1 - (a + b)p_0)f_X(x) + \sum_{y=1}^{\min(x,m)} (a + \frac{by}{x})f_X(y)f_S(x-y)}{1 - af_X(0)}$$

Следствие. Для класса $(a, b, 0)$

$$f_S(x) = \frac{\sum_{y=1}^{\min(x,m)} (a + \frac{by}{x})f_X(y)f_S(x-y)}{1 - af_X(0)}$$

Когда распределение ущерба имеет максимальное возможное значение m , сумма в рекурсии ограничена не более чем m ненулевыми компонентами. В этом случае вычислительная сложность имеет порядок $O(x)$.

22.2 Рекурсия для распределения Пуассона

Заметим, что когда распределение ущерба не имеет вероятности в нуле, знаменатель равен 1.

Для распределения Пуассона

$$f_S(x) = \frac{\lambda}{x} \sum_{y=1}^{\min(x,m)} y f_X(y) f_S(x-y)$$

Ошибки вносятся в последующие значения при суммировании

Начальное значение для выполнения рекурсивной схемы $f_S(0) = P_N(f_X(0))$.

В случае распределения Пуассона

$$f_S(0) = e^{-\lambda(1-f_X(0))}$$

22.3 Проблемы переполнения и недостаточной точности

Рекурсия начинается с вычисления $P(S = 0) = P_N(f_X(0))$. Для больших страховых портфелей эта вероятность очень мала, так что не может быть представлена в машинной памяти. В таком случае она становится равной нулю, и рекурсивная формула перестает быть применимой. Есть несколько путей решения этой проблемы. Один из самых простых способов – начинать с произвольных значений $f_S(0), f_S(1), \dots, f_S(k)$, где k находится достаточно далеко слева в распределении, так что истинное значение $F_S(k)$ все еще незначительно. Обычно достаточно установить k в точку, лежащую на шесть стандартных отклонений левее среднего значения. Рекурсия используется для генерации значений распределения с этим набором начальных значений до тех пор, пока найденные значения не будут стабильно меньше $f_S(k)$. Затем «вероятности» суммируются и делятся на сумму так, чтобы «истинные» вероятности давали в сумме 1. Насколько маленьким должно быть значение k для конкретной задачи, определяется методом проб и ошибок.

22.4 Проблемы переполнения и недостаточной точности

Другим методом получения вероятностей, когда начальное значение слишком мало, является выполнение вычислений для субпортфеля. Например, для распределения Пуассона со средним λ найдем значений $\lambda^* = \frac{\lambda}{2^n}$ такое, что вероятность в нуле не равна машинному нулю, когда λ^* используется в качестве среднего распределения Пуассона.

Рекурсивная формула теперь используется для совокупного распределения с λ^* . Если $P_{(z)}$ является производящей функцией вероятности совокупных убытков со средним распределения Пуассона λ^* , тогда $P_S(z) = (P_*(z))^{2^n}$. Следовательно, можно последовательно получить распределения с производящими функциями вероятностей $(P_*(z))^2, (P_*(z))^4, (P_*(z))^8, \dots, (P_*(z))^{2^n}$ путем свертки результата на каждом этапе с самим собой. Этот подход требует n дополнительных сверток при выполнении вычислений, но не предполагает никаких приближений. Это может быть выполнено для любых частотных распределений, которые замкнуты относительно свертки. Для отрицательного биномиального распределения аналогичная процедура начинается с $r^* = \frac{r}{2^n}$. Для биномиального распределения параметр m должен быть целым. В этом случае $m^* = \lfloor \frac{m}{2^n} \rfloor$. Когда n сверток выполнено, необходимо выполнить рекурсию для параметра mm^{2^n} . Этот результат затем сворачивается с результатом n сверток. Для составных распределений частоты только первичное распределение должно быть замкнуто относительно свертки.

22.5 Численная устойчивость

Любая рекурсивная формула требует точного вычисления значений, поскольку каждое такое значение будет использоваться при вычислении последующих значений. Рекурсивные схемы подвержены риску распространения ошибок по всем последующим значениям и риску потенциального сбоя. В рекурсивной формуле ошибки вносятся путем округления на каждом этапе, поскольку компьютеры представляют числа с конечным числом значащих цифр. Вопрос устойчивости заключается в следующем: насколько

ко быстро растут ошибки в вычислениях по мере подстановки вычисленных значений в последовательные вычисления? Можно сделать некоторые общие выводы.

Ошибки вносятся в последующие значения при суммировании

$$\sum_{y=1}^x \left(a + \frac{by}{x}\right) f_X(y) f_S(x-y)$$

В правом хвосте распределения S эта сумма положительна (или, по крайней мере, неотрицательна), и последующие значения суммы будут уменьшаться.

Сумма останется положительной, даже с учетом ошибок округления, когда каждый из трех коэффициентов в каждом члене суммы будет положительным. В этом случае рекурсивная формула устойчива, так как приводит к относительным ошибкам, которые растут медленно.

Для распределений, основанных на распределении Пуассона и отрицательных биномиальных, множители в каждом члене всегда положительны. Однако для биномиального распределения сумма может иметь отрицательные члены, поскольку a отрицательно, b положительно, а y/x - положительная функция, не превышающая 1. В этом случае отрицательные члены могут привести к тому, что последовательные значения будут знакопеременными. Когда это происходит, бессмысленные результаты сразу становятся очевидными. Хотя на практике это случается нечасто, необходимо знать о такой возможности в моделях, основанных на биномиальном распределении.

22.6 Непрерывное распределение ущерба

Рекурсивный метод требует дискретного распределения ущерба, в то время как обычно используют непрерывное распределение. В этом случае аналогом рекурсии будет интегральное уравнение, решением которого является распределение совокупных убытков.

Теорема 5. Для распределения частоты из класса $(a, b, 1)$ и любого непрерывного распределения ущерба с положительными действительными значениями имеет место следующее интегральное уравнение:

$$f_S(x) = p_1 f_X(x) + \int_0^x \left(a + \frac{by}{x}\right) f_X(y) f_S(x-y) dy$$

Заметим, что начальное значение равно $p_1 f_X(x)$, а не $(p_1 - (a+b)p_0) f_X(x)$ как в дискретной рекурсивной формуле. Выражение в теореме 5 выполнено и для класса $(a, b, 0)$.

Интегральное уравнение – это интеграл Вольтерра второго типа. Численное решение можно найти в литературе. Рассмотрим альтернативный подход к непрерывному распределению ущерба, основанный на дискретизации.

23 Построение арифметического распределения (округление, согласование моментов)

Простейший способ получить дискретное распределение из непрерывного – заменить дискретными значениями вероятностей в точках, кратных подходящей единице измерения h (шаг, размах, интервал). Такое распределение называется арифметическим, так как оно определяется в неотрицательных целых точках. Для того, чтобы «арифметизировать» распределение, важно сохранить свойства исходного распределения – как локальные, в рамках диапазона, так и глобально – для всего распределения. Должна быть сохранена общая форма распределения и количественные характеристики (моменты).

23.1 Метод округления

Пусть f_j - вероятность значения jh , $j = 0, 1, 2, \dots$. Тогда положим

$$f_0 = P(X < \frac{h}{2}) = F_X(\frac{h}{2} - 0)$$

$$f_j = P(jh - \frac{h}{2} < X < jh + \frac{h}{2}) = F_X(jh + \frac{h}{2} - 0) - F_X(jh - \frac{h}{2} - 0), j = 1, 2, \dots$$

Этот метод собирает всю вероятность половины промежутка с каждой стороны jh и помещает ее в точку jh . Есть исключение для вероятности нулевого значения. По сути, все количество округляется до ближайшей подходящей денежной единицы, h , шага дискретизации.

Когда непрерывное распределение ущерба неограниченно, имеет смысл остановить процесс дискретизации в некоторой точке как только большая часть вероятности будет учтена. Если m – индекс последней точки, то $f_m = 1 - F_X((m - 0.5)h - 0)$. При таком методе вероятности всегда будут неотрицательными и в сумме будут равны 1.

23.2 Метод локального согласования моментов

В этом методе строится арифметическое распределение, у которого p моментов соответствуют истинному распределению ущерба. Рассмотрим произвольный интервал $[x_k, x_k + ph)$ длины ph . Поместим массы вероятностей $m_k^0, m_k^1, \dots, m_k^p$ в точках $x_k, x_k + h, \dots, x_k + ph$ так, что первые p моментов сохраняют свои значения. Получим систему из $p + 1$ уравнений:

$$\sum_{j=0}^p (x_k + jh)^r m_j^k = \int_{x_k-0}^{x_k+ph-0} x^r dF_X(X), r = 0, 1, 2, \dots, p$$

(дискретная вероятность в x_k включается, а дискретная вероятность в $x_k + ph$ исключается). Упорядочим интервалы так, чтобы $x_{k+1} = x_k + ph$ и чтобы концевые точки совпадали. Тогда значения в концевых точках складываются. При $x_0 = 0$ результирующее дискретное распределение имеет последовательные вероятности:

$$f_0 = m_0^0, f_1 = m_1^0, f_2 = m_2^0, \dots, f_p = m_p^0 + m_0^1, f_{p+1} = m_1^1, f_{p+2} = m_2^1, \dots$$

Суммируя равенства по всем положительным k с учетом $x_0 = 0$, видим, что первые p моментов сохраняются для всего распределения, и сумма вероятностей в точности равна 1. Остается решить систему уравнений.

Теорема 6. Решение системы имеет вид

$$m_j^k = \int_{x_k-0}^{x_k+ph-0} \prod_{i \neq j} \frac{x - x_k - ih}{(j - i)h} dF_X(x), j = 0, 1, \dots, p$$

Доказательство.

Интерполяционный многочлен Лагранжа для полинома $f(y)$ в точках $y_0, y_1, \dots, y_n$ равен

$$f(y) = \sum_{j=0}^n f(y_j) \prod_{i \neq j} \frac{y - y_i}{y_j - y_i}$$

Применяя эту формулу к полиному $f(y) = y^r$ в точках $x_k, x_k + h, \dots, x_k + ph$, получаем

$$x^r = \sum_{j=0}^p (x_k + jh)^r \prod_{i \neq j} \frac{x - x_k - ih}{(j - i)h}, r = 0, 1, \dots, p$$

Интегрируя по интервалу $[x_k; x_k + ph)$ по функции распределения ущерба, получаем

$$\int_{x_k-0}^{x_k+ph-0} x^r dF_X(x) = \sum_{j=0}^p (x_k + jh)^r m_j^k$$

где m_j^k определены выше. Таким образом, решение системы сохраняет первые p моментов.

23.3 Свойство метода моментов

Метод локального согласования моментов предложен еще в 1976 году (Gerber, Jones) и изучался для различных эмпирических и аналитических распределений ущерба.

При оценке влияния ошибки на совокупную нетто-премию стоп-лосс оказалось, что двух моментов обычно достаточно, и добавление третьего момента лишь незначительно влияет на точность. Более того, метод округления и согласования первого момента дают одинаковые одинаковые ошибки, тогда как второй момент позволяет получить существенное увеличение точности.

Причина выбора согласования нулевого и первого момента заключается в том, что результирующие вероятности всегда будут неотрицательными. При совпадении двух или более моментов это нельзя гарантировать.

Описанные методы качественно аналогичны численным методам, используемым для решения интегральных уравнений Вольтерры, разработанным в рамках численного анализа.

24 Модель индивидуального риска

(из презентации 2-3)

Существуют два варианта построения модели суммы платежа по всем искам, произошедшим за фиксированный период времени и касающихся определенного набора договоров страхования. Второй вариант - приписать случайную величину каждому договору страхования.

Определение 2. Модель индивидуального риска представляет суммарный ущерб как сумму

$$X_1 + X_2 + \dots + X_n$$

фиксированного числа n договоров страхования. Убытки для n договоров равны $(X_1, X_2, \dots, X_n)$, где X_j предполагаются независимыми, но необязательно одинаково распределенными. Распределение X_j обычно имеет вероятность в нуле, соответствующую вероятности отсутствия убытка или платежа по договору.

Модель индивидуального риска используется с целью сложения убытков или платежей от фиксированного числа договоров страхования или множества страховых убытков группового страхования жизни или здоровья для группы n работников. Сотрудники могут иметь различное покрытие (выплаты по договору страхования кратны окладу) и различные уровни вероятности убытков (разный возраст и состояние здоровья).

В частном случае, когда X_j одинаково распределены, модель индивидуального риска становится частным случаем модели коллективного риска с вырожденным распределением N , т.е. $P(N = n) = 1$

(презентация 5)

Модель.

Модель индивидуального риска представляет совокупный убыток как сумму фиксированного числа независимых (но не обязательно одинаково распределенных) случайных величин:

$$S = X_1 + X_2 + \dots + X_n.$$

Формула часто рассматривается как сумма убытков от n страховых контрактов, например, n человек, участвующих в полисе группового страхования. Модель индивидуального риска изначально была разработана для страхования жизни, в условиях которого есть вероятность смерти в течение года q_j , и в случае смерти j -го участника платится фиксированная сумма b_j . В таком случае распределение убытков по j -му полису есть

$$f_{X_j}(x) = \begin{cases} 1 - q_j, & x = 0 \\ q_j, & x = b_j \end{cases}$$

Среднее и дисперсия совокупных убытков равны

$$E[S] = \sum_{j=1}^n b_j q_j, \text{Var}[S] = \sum_{j=1}^n b_j^2 q_j (1 - q_j)$$

так как X_j предполагаются независимыми.

Тогда производящая функция вероятности совокупных убытков равна

$$P_S(z) = \prod_{j=1}^n (1 - q_j + q_j z^{b_j})$$

Когда все риски одинаковы, $q_j = q$ и $b_j = 1$, получаем $P_S(z) = (1 + q(z - 1))^n$, т.е. S имеет биномиальное распределение.

Модель индивидуального риска можно обобщить следующим образом. Пусть $X_j = I_j B_j$, где $I_1, \dots, I_n, B_1, \dots, B_n$ независимы. Случайные величины I_j являются индикаторными, они принимают значение 1 с вероятностью q_j и 0 с вероятностью $1 - q_j$.

Эти переменные показывают, производится ли платеж по j -му полису. Случайная величина B_j может иметь любое распределение и представляет величину платежа по j -му полису. В страховании жизни B_j вырождена, со всей вероятностью в значении b_j .

Производящая функция моментов

$$M_S(z) = \prod_{j=1}^n (1 - q_j + q_j M_{B_j}(z))$$

Если положить $\mu_j = E[B_j]$, $\sigma_j^2 = Var[B_j]$, тогда $E[S] = \sum_{j=1}^n q_j \mu_j$, $Var[S] = \sum_{j=1}^n (q_j \sigma_j^2 + q_j(1 - q_j) \mu_j^2)$.

25 Функция полезности. Классы функций полезности

25.1 Функция полезности

Таким образом, следует учитывать не фактические выплаченные суммы, а скорее «удовлетворение» или «полезность», которые приходят с обладанием определенным уровнем богатства (блага). Это объяснение лежит в основе концепции теории полезности.

Функция полезности $u: X \rightarrow R^1$

$u(x)$ возрастает; как правило, $u(x)$ вогнута (убывающая предельная полезность), в случае дифференцируемой функции $u'(x) > 0$, $u''(x) < 0$. В силу линейности мат. ожидания функция полезности определяется с точностью до линейного преобразования ($au + b, a > 0$). [в мат. экономике – с точностью до монотонного возрастающего преобразования $f(u(x)), f' > 0$].

Если в петербургском парадоксе взять функцию полезности $\ln x$ и учитывать ожидаемую полезность, а не ожидаемую стоимость фактических сумм, то окажется, что выигрыш равен

$$\sum_{n=1}^{\infty} \ln(2^n) \left(\frac{1}{2^n}\right) < \infty$$

25.2 Классы функций полезности

Линейная: $u(w) = w$

квадратичная: $u(w) = -(\alpha - w)^2, w \leq \alpha; u(w) = 0, w > \alpha$

логарифмическая: $u(w) = \ln(\alpha + w), w > -\alpha$

экспоненциальная: $u(w) = -\alpha e^{-\alpha w}, \alpha > 0$

степенная: $u(w) = w^c, w > 0, 0 < c \leq 1$

Все эти функции имеют неотрицательную и невозрастающую предельную полезность. Коэффициент несклонности к риску для линейной полезности равен 0, для экспоненциальной α . Для остальных может быть представлен в виде

$$(\gamma + \beta w)^{-1}$$

Заметим, что линейная полезность приводит к принципу эквивалентности и риск-нейтральности, так как если $E[w + P^- - X] = w, P^- = E[X]$.

26 Отношение к риску. Коэффициент несклонности к риску

26.1 Отношение к риску

Людей с вогнутыми функциями полезности называют *несклонными к риску* (*risk averse*), поскольку они предпочитают определенность неопределенности и, следовательно, извлекают пользу из страхования.

Людей с выпуклой функцией полезности называют *любителями риска* (*risk-seekers*), поскольку такие люди готовы платить за азартную игру даже при неблагоприятных шансах (казино).

Линейную функцию полезности приписывают *риск-нейтральным* лицам.

26.2 Коэффициент несклонности к риску

Коэффициент несклонности к риску для функции полезности $u(\cdot)$ при размере капитала w :

$$r(w) = -\frac{u''(w)}{u'(w)}$$

Тогда максимальная премия за риск X

$$P^+ \approx \mu + \frac{1}{2}r(w - \mu)\sigma^2$$

Коэффициент несклонности к риску действительно отражает степень неприятия риска: чем больше неприятие риска, тем большую премию готов заплатить страхователь.

Заметим, что $r(w)$ не меняется при линейном преобразовании $u(\cdot)$.

27 Неравенство Йенсена

27.1 Выпуклые функции

Определение 1. Функция g , определенная на интервале I вещественной прямой, называется выпуклой, если для всех x и y из I и $0 \leq \alpha \leq 1$

$$g(\alpha x + (1 - \alpha)y) \leq \alpha g(x) + (1 - \alpha)g(y)$$

Особенностью выпуклых функций, которая часто используется, является условие возрастающего наклона: для $x < y < z$ из I

$$\frac{g(y) - g(x)}{y - x} \leq \frac{g(z) - g(y)}{z - y}$$

Используется для не всюду дифференцируемых функций. Для дифференцируемых функций выпуклость определяется знаком второй производной.

27.2 Неравенство Йенсена

Теорема 1. (неравенство Йенсена) Пусть $g(\cdot)$ – выпуклая функция. Тогда

$$E[g(X)] \geq g(E[X])$$

Для вогнутой функции неравенство выполнено в другую сторону.

Доказательство.

Приведем доказательство в случае, когда g имеет непрерывную вторую производную. Используя разложение в ряд Тейлора с учетом $E[X] = \mu$, имеем

$$g(x) = g(\mu) + (x - \mu)g'(\mu) + \frac{(x - \mu)^2}{2}g''(\xi),$$

для некоторой ξ между μ и x .

Так как g выпуклая, то $g'' \geq 0$, поэтому $g(x) \geq g(\mu) + (x - \mu)g'(\mu)$. Беря мат. ожидание, получаем

$$E[g(X)] \geq g(\mu) + g'(\mu)E[X - \mu] = g(\mu)$$

28 Верхняя и нижняя границы премии

28.1 Верхняя и нижняя границы премии

Страховщик, имеющий функцию полезности $U(\cdot)$ и капитал W , застрахует убыток X за премию P , если

$$E[U(W + P - X)] \geq U(W)$$

т.е. $P \geq P^-$, где P^- - минимальная премия, допустимая для страховщика. Она находится из условия равновесия (принять риск за премию vs ничего не делать):

$$U(W) = E[U(W + P^- - X)]$$

Должно быть: $P^+ \geq P^-$.

В теории страховщики часто считаются риск-нейтральными, так что за риск X премия $E[X]$ достаточна. Тогда $E[U(W + E[X] - X)] = U(W)$ для любого риска X .

28.2 Оценка максимальной премии

Оценим максимальную премию P^+ за риск X для заданной функции полезности $u(\cdot)$.

Пусть μ и σ^2 - среднее и дисперсия X . Используя разложение $u(\cdot)$ в окрестности μ , получим

$$u(w - P^+) \approx u(w - \mu) + (\mu - P^+)u'(w - \mu)$$

$$u(w - X) \approx u(w - \mu) + (\mu - X)u'(w - \mu) + \frac{1}{2}(\mu - X)^2u''(w - \mu)$$

Беря мат. ожидание от обеих частей, получаем

$$E[u(w - X)] \approx u(w - \mu) + \frac{1}{2}\sigma^2u''(w - \mu)$$

Из принципа нулевой полезности

$$\frac{1}{2}\sigma^2u''(w - \mu) \approx (\mu - P^+)u'(w - \mu)$$

$$P^+ \approx \mu - \frac{1}{2}\sigma^2 \frac{u''(w - \mu)}{u'(w - \mu)}$$

29 Сравнение случайных величин по степени рискованности

Определение 2. Для $X, Y \geq 0$ с $E[X] = E[Y]$ X менее рискованно, чем Y , если для всех d

$$E[(X - d)_+] \leq E[(Y - d)_+]$$

Если случайные величины представляют собой убытки, то чистая премия за страхование с франшизой, всегда меньше для менее рискованного варианта.

Так как $E[X] = E[Y]$, то неравенство эквивалентно $E[X \cap d] \geq E[Y \cap d]$ для всех $d \geq 0$.

$$u_d(x) = (x - d)_+, v_d(x) = d - x \cap d$$

Если X менее рискованно, чем Y , то для всех $d \geq 0$

$$E[u_d(X)] \leq E[u_d(Y)], E[v_d(X)] \leq E[v_d(Y)]$$

и, следовательно, для любой кусочно-линейной выпуклой функции g

$$E[g(X)] \leq E[g(Y)]$$

Произвольную выпуклую функцию можно аппроксимировать кусочно-линейной функцией (выбрать точки и соединить отрезками).

Таким образом, из двух распределений с одинаковым средним значением в случае несклонности к риску будет выбрано менее рискованное.

Это оправдывает определение, данное Ротшильдом и Стиглицем, которое действительно отражает концепцию рискованности.

Не всегда легко решить, является ли одна случайная величина менее рискованной, чем другая, но есть определенные случаи, когда это можно проверить.

29.1 Конечная случайная величина

Теорема 3. Предположим, что X принимает значения $a_1 \leq a_2 \leq \dots \leq a_{N-1} \leq a_N$ (с вероятностью $1/N$) и Y принимает значения $b_1 \leq b_2 \leq \dots \leq b_{N-1} \leq b_N$ (с вероятностью $1/N$); $E[X] = E[Y]$.

X менее рискованно, чем Y тогда и только тогда, когда

$$\sum_{i=1}^k a_i \geq \sum_{i=1}^k b_i \text{ для } 1 \leq k \leq N$$

Доказательство.

Достаточность. Пусть неравенство выполнено. Тогда для некоторого d , $a_k \leq d < a_{k+1}$, имеем

$$E[X \cap d] = \frac{1}{N} \left(\sum_{i=1}^k a_i + (N-k)d \right) \geq \frac{1}{N} \left(\sum_{i=1}^k b_i + (N-k)d \right) \geq E[Y \cap d]$$

Последнее неравенство следует из того, что $b_i \cap d$ меньше или равно b_i , и d для всех i .

Необходимость. Обратно, пусть X менее рискованно чем Y . Фиксируем $k < N$, пусть $d = \max(a_k, b_k)$. Если $d = b_k$, то

$$\frac{1}{N} \sum_{i=1}^k (d - a_i) \leq E[v_d(X)] \leq E[v_d(Y)] = \frac{1}{N} \sum_{i=1}^k (d - b_i)$$

если $d = a_k$, то

$$\frac{1}{N} \sum_{i=k+1}^N (a_i - d) = E[u_d(X)] \leq E[u_d(Y)] \leq \frac{1}{N} \sum_{i=k+1}^N (b_i - d)$$

В любом случае требуемое условие следует из того, что предположение о равенстве мат. ожиданий означает, что

$$\sum_{i=1}^N a_i = \sum_{i=1}^N b_i$$

Замечание. Поскольку допускается повторение значений, приведенная Теорема 3 может быть применена ко всем конечным дискретным распределениям, где вероятности являются рациональными числами.

29.2 Непрерывные случайные величины

Теорема 4. (условие отсечения) Предположим, что для некоторой точки с $F_X(t) \leq F_Y(t)$ при $t < c$ и $F_X(t) \geq F_Y(t)$ при $t > c$. Тогда X менее рискованна чем Y .

Доказательство.

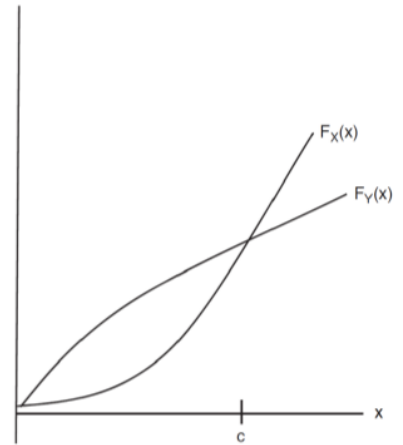
Если $d < c$,

$$E[X \cap d] = \int_0^d S_X(x) dx \geq \int_0^d S_Y(x) dx = E[Y \cap d]$$

Если $d > c$,

$$E[(X - d)_+] = \int_d^\infty S_X(x) dx \leq \int_d^\infty S_Y(x) dx = E[(Y - d)_+]$$

В силу непрерывности функции $g(d) = E[(X - d)_+]$, тот же результат получится при $d = c$. Во всех случаях по определению X менее рискованна чем Y .



30 Премия стоп-лосс, ее свойства

30.1 Премия стоп-лосс

Пусть размер убытка равен X (считаем, что $X \sim F_X(x)$, $f_X(x)$ – плотность распределения). Тогда размер выплаты перестраховщика $(X - d)_+ = \max(X - d, 0)$

Страховщик удерживает d , остальное оплачивает перестраховщик. С точки зрения страховщика рост убытка прекращается на уровне d . Покажем, что договор стоп-лосс является оптимальным с точки зрения дисперсии риска, который остается у страховщика. Это создает проблему – перестраховщики предлагают перестрахование стоп-лосс на менее выгодных условиях, чем другие виды перестрахования. Премией стоп-лосс назовем нетто-премию

$$\begin{aligned}\pi_X(d) &= E[(X - d)_+] \\ E[(X - d)_+] &= \int_d^\infty (x - d)f_X(x)dx = \{f_X(x) = F'_X(x) = -(1 - F_X(x))'\} = \\ &= -(x - d)[1 - F_X(x)]|_d^\infty + \int_d^\infty [1 - F_X(x)]dx = \int_d^\infty [1 - F_X(x)]dx\end{aligned}$$

Поэтому

$$\begin{aligned}\pi_X(d) &= \int_d^\infty [1 - F_X(x)]dx \\ \pi'_X(d) &= F_X(d) - 1\end{aligned}$$

Таким образом, $\pi_X(d)$ непрерывна и строго убывает.

Действительно, $\pi_X(0) = E[X]$, $\pi_X(\infty) = 0$.

Выплата, приходящаяся на долю перестраховщика при перестраховании стоп-лосс с собственным удержанием d для ущерба X , равна $(X - d)_+$.

30.2 Оптимальность перестрахования стоп-лосс

Теорема 5. Обозначим через $I(X)$ выплату по некоторому договору перестрахования, если ущерб равен X , $X \geq 0$. Пусть $0 \leq I(x) \leq x \forall x \geq 0$. Тогда выполнена импликация: $E[I(X)] = E[(X - d)_+] \Rightarrow \text{Var}[X - I(X)] \geq \text{Var}[X - (X - d)_+]$.

Доказательство.

Заметим, что $\forall I(\cdot)$ можно подобрать удержание d так, чтобы $E[I(X)] = E[(X - d)_+]$. В этом случае риски, оставляемые на собственном удержании, имеют вид

$$V(X) = X - I(X) \text{ и } W(X) = X - (X - d)_+$$

Поскольку $E[V(X)] = E[W(X)]$, то достаточно показать, что $E[(V(X) - d)^2] \geq E[(W(X) - d)^2]$. Для этого достаточно, чтобы неравенство

$$|V(X) - d| \geq |W(X) - d| \text{ выполнялось с вероятностью } 1$$

Если $X \geq d$, то $W(X) = d$ и неравенство очевидно. Если $X < d$, то $W(X) = X$ и

$$V(X) - d = X - I(X) - d \leq X - d = W(X) - d < 0 \Rightarrow |V(X) - d| \geq |W(X) - d|$$

Таким образом, перестрахование стоп-лосс минимизирует дисперсию риска, оставляемого на собственном удержании. В теореме решающую роль играет предположение, что при одинаковых ожидаемых выплатах премия при перестраховании стоп-лосс равна премии при других видах перестрахования. Поскольку дисперсия капитала перестраховщика при перестраховании стоп-лосс будет больше, чем при других видах перестрахования, на практике перестраховщик, будучи всегда не слишком склонным к риску, будет взимать за перестрахование стоп-лосс большую премию.

30.3 Оптимальность пропорционального перестрахования

Найдем ожидаемую прибыль страховщика (премия минус ожидаемое значение риска, оставленного на собственном удержании, минус перестраховочная премия):

$$A : (1 + \theta)E[X] - E[X - I(X)] - (1 + \lambda)E[I(X)] = \theta E[X] - \lambda E[I(X)]$$

$$B : (1 + \theta)E[X] - E[X - I(X)] - (E[I(X)] + \alpha \text{Var}[I(X)]) = \theta E[X] - \alpha \text{Var}[I(X)]$$

Легко видеть, что при обоих сценариях размер ожидаемой прибыли равен размеру исходной ожидаемой прибыли $\theta E[X]$ за вычетом ожидаемой прибыли перестраховщика. Значит, необходимо минимизировать ожидаемую прибыль перестраховщика.

Получаем задачу минимизации А:

$$\min E[I(X)]$$

$$\text{Var}[X - I(X)] = V$$

Для решения задачи А используем доказанную теорему об оптимальности перестрахования стоп-лосс.

$$\mu(d) = E[X - (X - d)_+] \text{ и } \sigma^2(d) = \text{Var}[X - (X - d)_+]$$

Функции $\mu(d)$ и $\sigma^2(d)$ непрерывно возрастают с ростом d и

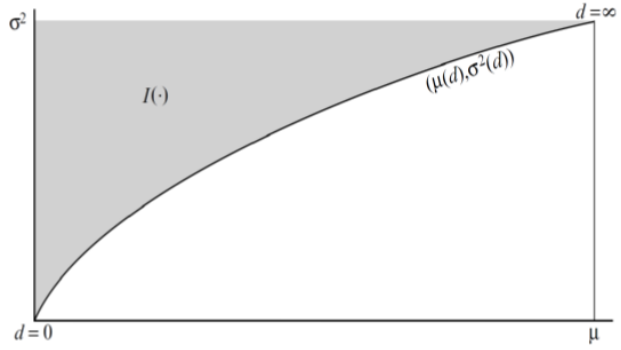
$$\mu(0) = \sigma^2(0) = 0, \mu(\infty) = E[X], \sigma^2(\infty) = \text{Var}[X]$$

(так как $\mu'(d) = 1 - F_X(d)$ и $(\sigma^2)'(d) = 2[1 - F_X(d)][d - \mu(d)]$,

Точки, изображающие ожидаемое значение и дисперсию риска, оставляемого на собственном удержании, могут располагаться только выше этой кривой на плоскости μ, σ^2 (по теореме об оптимальности стоп-лосс)

Это обусловлено тем, что дисперсия не может быть меньше, чем при перестраховании стоп-лосс с тем же ожидаемым значением. Отсюда следует, что такие точки могут располагаться лишь слева от кривой. И можно заключить, как и в теореме, что непропорциональное перестрахование стоп-лосс дает решение задачи А. Перестрахование стоп-лосс в данном случае является оптимальным по Парето: не существует решений с меньшей дисперсией и большим ожидаемым доходом.

На рисунке изображено множество точек $(\mu(d), \sigma^2(d))$, для $d \in [0, \infty)$ и некоторой случайной величины потерь.



31 Невошедшее в билеты из презентаций

31.1 Свойства суммы случайных величин

Рассмотрим страховой портфель. Общая сумма, выплаченная страховой компанией по всем полисам равна сумме платежей, поэтому далее рассмотрим свойства

$$S_k = X_1 + \dots + X_k$$

Центральная предельная теорема

$$E[S_k] = E[X_1] + \dots + E[X_k]$$

Если $X_1, \dots, X_k$ независимы, то $Var[S_k] = Var[X_1] + \dots + Var[X_k]$.

Если $X_1, X_2, \dots$ независимы, независимы, имеют конечные математическое ожидание и дисперсию, то

$$\lim_{k \rightarrow \infty} \frac{S_k - E[S_k]}{\sqrt{Var[S_k]}} \sim N(0, 1)$$

Сходимость рассматривается в смысле сходимости по распределению, т.е. случайные величины $\frac{S_k - E[S_k]}{\sqrt{Var[S_k]}}$ имеют распределение, пределом которого является стандартное нормальное распределение.

Определение. Производящая функция моментов случайной величины $M_X(z) = E[e^{zX}]$ для всех z для которых математическое ожидание существует. Производящая функция вероятностей $P_X(z) = E[z^X]$ для всех z для которых математическое ожидание существует.

$$M_X(z) = P_X(e^z) \text{ и } P_X(z) = M_X(\ln z)$$

Производящие функции моментов (п.ф.м.) используются для непрерывных случайных величин, а производящие функции вероятностей (п.ф.в.) для дискретных. Эти функции используются для вычисления моментов и вероятностей. Кроме того, существует взаимно однозначное соответствие между функцией распределения случайной величины и ее п.ф.м. и п.ф.в.

Теорема. Пусть $S_k = X_1 + \dots + X_k$, $X_1, \dots, X_k$ независимы. Тогда

$$M_{S_k}(z) = \prod_{j=1}^k M_{X_j}(z) \text{ и } P_{S_k}(z) = \prod_{j=1}^k P_{X_j}(z)$$

Доказательство: Так как математическое ожидание произведения независимых случайных величин равно произведению математического ожидания, то

$$M_{S_k}(z) = E[e^{zS_k}] = E[e^{z(X_1 + \dots + X_k)}] = \prod_{j=1}^k E[e^{zX_j}] = \prod_{j=1}^k M_{X_j}(z)$$

31.2 Примеры согласованных и несогласованных мер

Принцип стандартного отклонения.

Стандартное отклонение является мерой неопределенности распределения. Рассмотрим распределение убытков со средним μ и стандартным отклонением σ . Величина $\mu + k\sigma$, где k - константа, является мерой риска (часто называемой принципом стандартного отклонения). Коэффициент k обычно выбирается для обеспечения того, чтобы убытки превысили меру риска для какого-то распределения, например нормального, с заданной малой вероятностью.

Принцип стандартного отклонения не является согласованным – не удовлетворяет п. 2 (монотонности).

31.3 Модели уязвимости (Frailty models)

Введем случайную величину уязвимости $\Lambda > 0$ и определим условный (при заданном $\Lambda = \lambda$) уровень риска X :

$$h_{X|\Lambda}(x|\lambda) = \lambda a(x)$$

где $a(x)$ - заданная функция (зависит от конкретного приложения).

Уязвимость предназначена для количественной оценки неопределенности, связанной со степенью риска, которая действует мультипликативно. Тогда условная функция выживания $X|\Lambda$ равна

$$S_{X|\Lambda}(x|\lambda) = e^{-\int_0^x h_{X|\Lambda}(t|\lambda) dt} = e^{-\lambda A(x)}$$

где $A(x) = \int_0^x a(t) dt$

Чтобы найти смесь распределений (т.е. безусловное распределение X), определим производящую функцию моментов случайной величины уязвимости Λ как $M_\Lambda(z) = E[e^{z\Lambda}]$. Тогда

$$S_X(x) = E[e^{-\Lambda A(x)}] = M_\Lambda[-A(x)]$$

и $F_X(x) = 1 - S_X(x)$.

Тип используемой смеси определяется выбором $a(x)$ и, значит, $A(x)$. Наиболее важный подкласс моделей уязвимости это класс экспоненциальных смесей с $a(x) = 1$ и $A(x) = x$, так что $S_{X|\Lambda}(x|\lambda) = e^{-\lambda x}$, $x \geq 0$.

Смесь Вейбулла с $a(x) = \gamma x^{\gamma-1}$ и $A(x) = x^\gamma$.

Для оценки распределения уязвимости требуется выражение для производящей функции моментов $M_\Lambda(z)$.

Наиболее общий выбор гамма уязвимости, но и другие варианты, например, обратное распределение Гаусса, также используются.

31.4 Усеченные и модифицированные распределения (a, b, 0)

Необходимо различать ситуации, в которых $p_0 = 0$ и $p_0 > 0$. Первый подкласс называется усеченными (в нуле) распределениями. Ему принадлежат усеченное распределение Пуассона, усеченное биномиальное и усеченное отрицательное биномиальное.

Второй подкласс называется модифицированными в нуле распределениями, так как вероятность отличается от того, что есть в классе (a, b, 0). Эти распределения можно

рассматривать как смесь распределения $(a, b, 0)$ и вырожденного распределения с вероятностью, сосредоточенной в нуле. Покажем эту эквивалентность формально. Заметим, что все усеченные в нуле распределения можно рассматривать как модифицированные в нуле распределения с $p_0 = 0$. Будем для краткости использовать аббревиатуры ZT (zero truncated) и ZM (zero modified). Для того, чтобы различать вероятности будем использовать $p_k = P(N = k)$. Для усеченных в нуле p_k^T и для модифицированных p_k^M .

Пусть $P^T(z)$ обозначает производящую функцию вероятности усеченного в нуле распределения, соответствующего $(a, b, 0)$ производящей функции вероятности $P(z)$. Тогда полагая $P_0^M = 0$

$$p^T = \frac{P(z) - p_0}{1 - p_0}, p_k^T = \frac{p_k}{1 - p_0}, k = 1, 2, \dots$$

$$p_k^M = (1 - p_0^M)p_k^T, k = 1, 2, \dots$$

$$p^M(z) = p_0^M(1) + (1 - p_0^M)P^T(z)$$

31.5 Усеченные и модифицированные в нуле распределения

Пусть N^L зависит от параметров θ и α так, что

$$P_{N^L}(z) = P_{N^L}(z, \theta, \alpha) = \alpha + (1 - \alpha) \frac{\phi(\theta(1 - z)) - \phi(\theta)}{1 - \phi(\theta)}$$

Заметим, что $\alpha = P_{N^L}(0) = P(N^L = 0)$ так что является модифицированной вероятностью в нуле. Также имеет место, что если $\phi(\theta(1 - z))$ сама является производящей функцией вероятности, то написанная выше производящая функция является таковой для соответствующего модифицированного в нуле распределения. Однако необязательно, чтобы $\phi(\theta(1 - z))$ была производящей функцией вероятности $P_{N^L}(z)$ указанного вида. В частности, $\phi(z) = 1 + \ln(1 + z)$ приводит к ZM логарифмическому распределению, несмотря на то, что не существует распределения с такой $\phi(z)$ в качестве производящей функции вероятности. Аналогично, $\phi(z) = (1 + z)^{-r}$ для $1 < r < 0$ приводит к распределению ETNB.

Нетрудно видеть, что $P_{N^P}(z) = P_{N^L}(z, v\theta, \alpha^*)$, где $\alpha^* = P(N^P = 0) = P_{N^P}(0) = P_{N^L}(1 - v, \theta, \alpha)$. Вполне ожидаемо, что установление франшизы увеличивает α , так как периоды без платежей оказываются более характерными. В частности, если N^L является усеченной в нуле, то N^P будет модифицированной в нуле.

32 Информация по экзамену

32.1 Примерная теория с экзамена 2023 года

1. Определение среднего превышения, вывод зависимости среднего превышения от функции выживания.
2. Усеченная справа сл.в. График, смысл
3. Связь между $(a, v, 0)$ и $(a, v, 1)$
4. Функция выживания для платежи при условной и безусловной франшизах
5. Обобщенное распределение Пуассона
6. П.ф.в, доказательство того, что производящая суммы - произведение производящих
7. Метод округления, зачем он нужен
8. Полуаддитивность меры риска, примеры, смысл
9. Коэффициент несклонности к риску, его связь с вогнутостью функции полезности (вопросов было 10)

32.2 Примерная теория с 1 пересдачи 2023 года

1. Уровень риска. Определение, как выражается функция выживания.
2. Уровень риска, как связан с тяжестью хвоста.
3. Доказать, что биномиальное это $(a, b, 0)$.
4. Плотность равновесного распределения, доказать что тона является плотностью.
5. Обобщенное распределение Пуассона, функция распределения суммы.
6. Распределение с Пуассоном, сумма случайных величин.
7. Обратное, преобразованное, обратно преобразованное, показать на примере экспоненциального.
8. Метод округления, принцип, зачем нужен.
9. Неравенство Йенсена на примере производящей функции моментов экспоненциального распределения.
10. Коэффициент устранения потерь, суть, показать на примере распределения Парето.

32.3 Кратко о подготовке

На экзамене может спрашиваться практически какая угодно теория: от доказательств и кратких выводов до примеров рассмотренных в презентациях. Экзамен состоит из двух частей: теории и задач. На пересдаче не разрешалось использовать дополнительные материалы (в 2023 на экзамене на задачах можно было использовать интернет/записи). В обоих случаях по теории было 10 вопросов, а по задачам 4.

33 Билет с зачета ФИИТа 4 курс

ВАРИАНТ 1

1. Чему равно ожидаемое значение ограниченного убытка?
2. Неотрицательная случайная переменная имеет уровень риска $h(x) = A + e^{2x}, x \geq 0$. Также известно, что $S(0.4) = 0.5$. Найдите A .
3. Как связано среднее превышение ущерба с тяжестью хвоста?
4. Распределение Парето с параметрами $\theta, \alpha > 1$, имеет вид

$$F(x) = 1 - \left(\frac{\theta}{x + \theta} \right)^\alpha, x > 0.$$

Покажите, что $VaR_p[X] = \theta \left((1-p)^{-\frac{1}{\alpha}} - 1 \right)$ и $TVaR_p[X] = VaR_p[X] + \frac{VaR_p[X] + \theta}{\alpha - 1}$

5. Пусть X имеет распределение Парето. Определите обратное преобразованное распределение.
6. Что называется смесью распределений с переменным числом компонент?
7. Дайте определение условной франшизы.
8. Чему равна ожидаемая стоимость платежа в условиях безусловной и условной франшизы?
9. Чему равна ожидаемая стоимость страхования стоп-лосс в дискретном случае?
10. В чем состоит метод локального согласования моментов при дискретизации непрерывного распределения?

34 Задачи с экзамена 2023

ЗАДАЧИ

1. Размер убытков имеет распределение, являющееся смесью равномерного распределения на $[0, 10]$ с весом 0.5 и равномерного распределения на $[5, 13]$ с весом 0.5. Ограниченная ожидаемая величина на уровне a равна 6.11875. Определите a .
2. Годовые убытки имеют распределение Парето со средним 100 и дисперсией 20 000. Вычислите TVaR на уровне защиты 65%.
3. Случайная величина N имеет модифицированное в нуле распределение Пуассона. Известно, что $P\{N = 1\} = 0.25, P\{N = 2\} = 0.1$. Вычислите вероятность нулевого значения.
4. Страховщик применяет перестрахование превышения убытка. Ожидаемые убытки в год равны 10 000 000. Индивидуальные убытки имеют распределение Парето с

$$F(x) = 1 - \left(\frac{2000}{x + 2000} \right)^2.$$

Перестраховщик оплачивает превышение каждого ущерба выше 3000. Годовая премия перестраховщика равна 110% его ожидаемых убытков. Распределение частоты убытков не меняется. Определите величину премии.

35 Литература, указанная преподавателем

1. Кларк С.М., Харди М.Р., Макдоналд А.С., Вотерс Г.Р. Теория риска. Учебные материалы. М: 2008 ;
2. Фалин Г.И., Фалин А.И. Теория риска для актуариев в задачах. М.: Изд-во ф-та ВМиК, 2001. ;
3. Р. Каас, М. Гувертс, Ж. Дэнэ, М. Денут Современная актуарная теория риска М: Янус-К, 2007 (R. Kaas, M. Goovaerts, J. Dhaene, M. Denuit Modern Actuarial Risk Theory. Springer, 2008) ;
4. Klugman, S. A., Panjer, H. H., Willmot, G. E. Loss models. Hoboken, NJ: Wiley Interscience, 2004 ;
5. Promislow, S. Fundamentals of Actuarial Mathematics: Third Edition. Wiley, 2015 ;