

LOSS ELIMINATION RATIO OF TOTAL MOTOR VEHICLE INSURANCE LOSSES USING ORDINARY DEDUCTIBLE

BERLIAN SETIAWATY*, MUTIARA MAHARANI, RUHIYAT, WINDIANI ERLIANA

*Departemen Matematika,
Institut Pertanian Bogor, Kampus IPB Darmaga, Bogor 16680
email : berlianse@apps.ipb.ac.id, mutiara-dee98@apps.ipb.ac.id,
ruhiyat-mat@apps.ipb.ac.id, windi@apps.ipb.ac.id*

Diterima 4 Juni 2024 Direvisi 2 Oktober 2024 Dipublikasikan 31 Januari 2025

Abstrak. Penelitian ini membahas total kerugian dari asuransi kendaraan bermotor jika diterapkan *ordinary deductible*. Data yang digunakan dalam penelitian ini adalah data Motorcycle Insurance dari *package insuranceData R*. Pemodelan total kerugian dilakukan dengan memodelkan banyak klaim menggunakan sebaran binomial negatif dan besar klaim menggunakan sebaran *beta-prime*. Sebaran total kerugian merupakan sebaran *compound* dengan sebaran primernya binomial negatif dan sebaran sekundernya *beta-prime*. Sebaran *compound* ini sulit diperoleh bentuk analitiknya sehingga digunakan simulasi Monte Carlo. Dengan simulasi Monte Carlo dapat dihitung nilai harapan kerugian agregat tanpa *ordinary deductible* maupun dengan *ordinary deductible*, sehingga diperoleh *loss elimination ratio* dari penggunaan *ordinary deductible*. Dari perhitungan yang dilakukan dapat disimpulkan semakin besar nilai *ordinary deductible*, semakin meningkat pula *loss elimination ratio*-nya.

Kata Kunci: loss elimination ratio, ordinary deductible, simulasi Monte Carlo, total kerugian

Abstract. This research discusses the total loss from motor vehicle insurance if the *ordinary deductible* is applied. The data used in this research is Motorcycle Insurance data in the *insuranceData R* package. Total loss modeling was done by modeling the frequency of claims using a negative binomial distribution and the severity of claims using a *beta-prime* distribution. The total loss distribution is a compound distribution, with the primary distribution being negative binomial and the secondary distribution being *beta-prime*. The distribution of this compound is difficult to obtain in analytical form, so the Monte Carlo simulation method was used. With the Monte Carlo simulation, the expected value of aggregate losses without *ordinary deductible* or with *ordinary deductible* can be calculated to obtain a *loss elimination ratio*. From the calculations, it can be concluded that the greater the deductible value, the greater the *loss elimination ratio*.

Keywords: loss elimination ratio, ordinary deductible, Monte Carlo simulation, total loss

*penulis korespondensi

1. Pendahuluan

Asuransi adalah bentuk perjanjian antara kedua belah pihak, yaitu pihak tertanggung dan penanggung asuransi, di mana pihak tertanggung membayar premi kepada penanggung asuransi untuk mendapatkan bentuk ganti rugi atas risiko finansial yang dapat terjadi secara tak terduga. Pihak asuransi akan memberikan jaminan terhadap risiko atau kerusakan kepada pihak tertanggung (nasabah) untuk memenuhi hak pemegang polis sesuai yang tertera dalam polis yang disebut klaim. Pihak tertanggung melakukan suatu atau serangkaian pembayaran yang ditetapkan sebagai biaya pengalihan risiko dari pemegang polis kepada penyedia asuransi yang disebut premi. Besarnya premi asuransi untuk masing-masing tertanggung dapat berbeda. Hal ini ditentukan oleh beberapa faktor seperti: besarnya uang pertanggungan, jangka waktu asuransi, kebijakan perusahaan asuransi, dan lain-lain.

Untuk merancang polis asuransi perusahaan harus memprediksi total kerugian yang dialami oleh pemegang polis yang nantinya harus ditanggung oleh perusahaan asuransi dalam suatu periode waktu tertentu. Untuk memodelkan total kerugian menggunakan *collective risk model*, menurut [1] diperlukan informasi tentang banyak dan besar klaim. Banyak klaim dimodelkan oleh peubah acak diskret dan besar klaim dimodelkan oleh peubah acak kontinu tak negatif. Sebaran total kerugian merupakan sebaran *compound* dari sebaran banyak klaim dan sebaran besar klaim.

Untuk mengurangi total kerugian yang harus ditanggung oleh perusahaan asuransi, bisa diterapkan beberapa strategi. Salah satunya adalah menerapkan kebijakan *ordinary deductible*. Menurut [2], *ordinary deductible* atau biasa juga disebut dengan *deductible* atau *own risk* adalah jumlah tertentu yang harus dibayarkan oleh pemegang polis jika terjadi klaim. Tujuan penggunaan *ordinary deductible* adalah untuk mencegah klaim yang ukurannya kecil tapi sering terjadi, sehingga total kerugiannya bisa menjadi besar. Kebijakan *ordinary deductible* akan mendorong pemegang polis untuk lebih berhati-hati agar tidak mengalami kerugian karena mereka harus ikut menanggung bagian dari kerugian tersebut. *Ordinary deductible* sering digunakan pada asuransi kendaraan bermotor.

Penetapan *ordinary deductible* yang terlalu rendah akan memberatkan perusahaan asuransi, di lain pihak *ordinary deductible* yang terlalu tinggi tidak menarik bagi calon peserta asuransi. Pada penelitian ini akan diamati perilaku *loss elimination ratio* (LER) yang disebabkan oleh penerapan *ordinary deductible*. Data yang digunakan dalam penelitian ini adalah data Motorcycle Insurance yang merupakan salah satu data yang tersedia pada *package insuranceData R*.

Data dari *package insuranceData R* ini bersifat *open* dan banyak dipakai untuk berbagai penelitian, di antaranya oleh [3] dan [4]. Penelitian [3] menggunakan berbagai algoritma *machine learning* seperti *Support Vector Machines* (SVM), *Decision Trees*, *Random Forests*, dan *Boosting* untuk memprediksi klaim asuransi kendaraan bermotor. Penelitian ini menemukan bahwa variabel seperti penjualan mobil baru dan kondisi cuaca sangat berpengaruh dalam memprediksi klaim asuransi. Dalam [4], data tersebut digunakan untuk menguji model regresi baru dalam masalah frekuensi dan keparahan klaim. Model seperti *Generalized Linear Models* (GLM),

Generalized Linear Mixed Models (GLMM), dan model campuran non-linear juga sering diuji menggunakan data ini.

Pada penelitian ini dianalisis kaitan antara penggunaan *deductible* dan kemampuannya mengurangi total kerugian yang dihitung dengan LER pada asuransi kendaraan bermotor. Studi tentang total kerugian telah dilakukan antara lain oleh [1], [5] dan [6]. Mohamed dkk. [6] dalam artikelnya membahas cara mengaproksimasi sebaran total kerugian menggunakan simulasi. Klugman [1] membahas aspek teoritis LER dari total kerugian jika dilakukan modifikasi *coverage* secara umum, sedangkan Gee dan Frees [5] membahas hubungan antara LER total kerugian dengan penggunaan *ordinary deductible* pada asuransi properti.

Untuk menghitung LER diperlukan sebaran dari total kerugian yang merupakan sebaran *compound*. Bentuk analitik sebaran ini biasanya tidak mudah diperoleh, oleh sebab itu dilakukan simulasi Monte Carlo yang idenya diperoleh dari [6]. Dari simulasi Monte Carlo diperoleh nilai harapan total kerugian tanpa *ordinary deductible* maupun dengan *ordinary deductible*. Dari kedua nilai harapan dapat dihitung LER.

2. Landasan Teori

2.1. Sebaran binomial dan sebaran beta-prime

Peubah acak N menyebar binomial negatif dengan parameter $r, \beta > 0$, dilambangkan $N \sim NB(r, \beta)$, jika memiliki fungsi massa peluang:

$$p_k = P(N = k) = \frac{r(r+1) \cdots (r+k-1)\beta^k}{k!(1+\beta)^{r+k}}, \quad k = 0, 1, 2, \dots \quad (2.1)$$

Nilai harapan dan ragam dari sebaran binomial negatif menurut [7] dinyatakan sebagai berikut.

$$E[N] = r\beta, \quad (2.2)$$

$$Var(N) = r\beta(1+\beta). \quad (2.3)$$

Peubah acak X menyebar *beta-prime* dengan parameter $\alpha, \beta > 0$, dilambangkan $X \sim B'(\alpha, \beta)$ jika mempunyai fungsi kepekatan peluang sebagai berikut [8].

$$f(x) = \frac{\Gamma(\alpha+\beta)}{\Gamma(\alpha)\Gamma(\beta)} \cdot \frac{x^{\alpha-1}}{(1+x)^{\alpha+\beta}}, \quad x > 0. \quad (2.4)$$

Nilai harapan dan ragam dari sebaran *beta-prime* menurut [8] adalah

$$E[X] = \frac{\alpha}{\beta-1}, \quad (2.5)$$

$$Var(X) = \frac{\alpha(\alpha+\beta-1)}{(\beta-1)^2(\beta-2)}. \quad (2.6)$$

2.2. Metode Maximum Likelihood

Menurut [1], metode *maximum likelihood* adalah suatu metode statistik yang digunakan untuk menduga nilai parameter dari suatu sebaran berdasarkan data yang

diamati. Misalkan $X_1, X_2, \dots, X_n$ adalah peubah acak saling bebas dan identik yang memiliki fungsi kepekatan peluang $f(x; \theta)$, dengan θ adalah parameter yang tidak diketahui dan akan diduga. Definisikan fungsi *likelihood* sebagai:

$$L(\theta; x_1, x_2, \dots, x_n) = \prod_{i=1}^n f(x_i; \theta),$$

dan fungsi *log-likelihood* sebagai:

$$l(\theta; x_1, x_2, \dots, x_n) = \log(L(\theta; x_1, x_2, \dots, x_n)) = \sum_{i=1}^n \log f(x_i; \theta).$$

Notasi $\hat{\theta} = \hat{\theta}(x_1, x_2, \dots, x_n)$ adalah dugaan *maximum likelihood* apabila:

$$\hat{\theta} = \arg \max l(\theta; x_1, x_2, \dots, x_n).$$

Menurut Montgomery [9], $\hat{\theta}$ merupakan penyelesaian dari persamaan

$$\frac{\partial l(\theta; x_1, x_2, \dots, x_n)}{\partial \theta} = 0.$$

2.3. Uji khi-kuadrat

Menurut [1], uji khi-kuadrat merupakan uji statistik yang digunakan untuk mengevaluasi sejauh mana data yang diamati akan sesuai dengan sebaran atau model yang diharapkan. Pada pengujian ini, nilai sampel acak $x_1, x_2, \dots, x_n$ akan diuji apakah berasal dari fungsi sebaran $F(x)$.

Hipotesisnya sebagai berikut:

H_0 : data menyebar $F(x)$,

H_1 : data tidak menyebar $F(x)$.

Statistik ujinya adalah:

$$d = \sum_{j=1}^k (o_j - e_j)^2 / e_j,$$

di mana o_j adalah banyaknya pengamatan dan e_j adalah nilai harapan pada kategori ke- j , dengan $j = 1, 2, \dots, k$ serta k adalah banyaknya kategori. Pada uji suai khi-kuadrat, $D = \sum_{j=1}^k (O_j - E_j)^2 / E_j$ menyebar khi-kuadrat dengan derajat kebebasan $k - 1$.

Nilai *p-value* dihitung menggunakan:

$$p = Pr(D > d).$$

Kriteria keputusan untuk tingkat signifikansi α , jika $p > \alpha$, maka tidak ada bukti untuk menolak H_0 sehingga kita dapat memodelkan data dengan sebaran $F(x)$.

2.4. Uji Kolmogorov-Smirnov

Uji Kolmogorov-Smirnov dilakukan untuk mengetahui apakah data mengikuti sebaran tertentu. Misalkan $x_{(1)}, x_{(2)}, \dots, x_{(n)}$ adalah sampel acak terurut dari X yang diduga menyebar $F(x)$ dan fungsi sebaran empiris dari X adalah $F_n(x)$.

Hipotesis uji Kolmogorov-Smirnov adalah:

$$H_0 : \text{data menyebar } F(x),$$

$$H_1 : \text{data tidak menyebar } F(x).$$

Statistik ujinya adalah:

$$d = \sup |F_n(x) - F(x)|.$$

Menurut Feller [10], $D_n = \sup |F_n(X) - F(X)|$ tidak memiliki sebaran eksak, tapi untuk $n \rightarrow \infty$, berlaku:

$$Pr(D_n \leq xn^{-1/2}) \rightarrow L(x), \quad x \geq 0,$$

di mana $L(x)$ adalah fungsi sebaran kumulatif, dengan:

$$L(x) = 1 - 2 \sum_{k=1}^{\infty} (-1)^{k-1} e^{-k^2 x^2} = (2\pi)^{1/2} x^{-1} \sum_{k=1}^{\infty} e^{-(2k-1)^2 \pi^2 / 8x^2}, \quad x \geq 0.$$

Menurut Marsaglia dkk. [11], p -value dihitung menggunakan

$$p = Pr(D_n > d) \approx 1 - L(d\sqrt{n}) \approx 1 - K(n, d),$$

dengan $K(n, d)$ adalah aproksimasi dari $L(d\sqrt{n})$.

Kriteria keputusan untuk tingkat signifikansi α , jika $p > \alpha$, maka tidak ada bukti untuk menolak H_0 sehingga kita dapat memodelkan data dengan sebaran $F(x)$.

2.5. Ordinary deductible, total kerugian dan loss elimination ratio

Misalkan X adalah peubah acak kontinu tak negatif yang menyatakan besar klaim. *Ordinary deductible* d adalah jumlah yang harus dibayarkan oleh pemilik polis jika terjadi klaim. Dalam hal ini, menurut [2] ganti rugi yang harus dibayarkan oleh perusahaan asuransi menjadi:

$$Y = (X - d)_+ = \begin{cases} 0, & X < d, \\ X - d, & X \geq d. \end{cases} \quad (2.7)$$

Jika N adalah *counting random variable* yang menyatakan banyak klaim dan Y_i adalah peubah acak yang menyatakan besar klaim ke- i dengan *ordinary deductible* d , maka total kerugian yang harus ditanggung perusahaan mengikuti *collective risk model* [1] menjadi:

$$S_d = \sum_{i=1}^N Y_i = \sum_{i=1}^N (X_i - d)_+. \quad (2.8)$$

Jika $X_i, i = 1, 2, \dots$ saling bebas dan identik serta bebas terhadap N , maka sebaran S_d merupakan sebaran *compound* dengan sebaran primer N dan sebaran sekunder Y . Sebaran ini biasanya sulit ditentukan secara analitik.

Untuk menghitung rasio pengurangan kerugian yang harus ditanggung perusahaan karena kebijakan *ordinary deductible* sebesar d , dalam [1] didefinisikan *loss elimination ratio* (LER) sebagai berikut.

$$LER = \frac{E[S_0 - S_d]}{E[S_0]} = 1 - \frac{E[S_d]}{E[S_0]}, \quad (2.9)$$

dengan:

$$S_0 = \sum_{i=1}^N X_i.$$

Berdasarkan sifat sebaran *compound* [12], diperoleh:

$$E[S_0] = E[N]E[X], \quad (2.10)$$

$$Var(S_0) = E[N]Var(X) + Var(N)(E[X])^2. \quad (2.11)$$

3. Metode Penelitian

3.1. Data

Data yang digunakan dalam penelitian ini adalah data Motorcycle Insurance yang ada di dalam *package* insuranceData R. Data ini berasal dari perusahaan asuransi Swedia yang berisi data klaim asuransi kendaraan bermotor selama 1994 – 1998. Terdapat sembilan variabel dalam Motorcycle Insurance, yaitu: *agarald*, *kon*, *zon*, *mcklass*, *fordald*, *bonuskl*, *duration*, *anstkad* dan *skadkost* yang berturut-turut menyatakan usia pemilik kendaraan, zona tempat tinggal pemilik kendaraan, klasifikasi mesin kendaraan, usia kendaraan, kelas bonus asuransi, lama waktu mengikuti asuransi kendaraan bermotor, banyak klaim, dan besar klaim. Karena penelitian ini fokus pada total kerugian maka hanya dua variabel yang diteliti, yaitu *anstkad* dan *skadkost*. Variabel *anstkad* merupakan data banyaknya klaim (N) dan *skadkost* data besarnya klaim (X).

3.2. Prosedur penelitian

Langkah-langkah penelitian untuk mengaproksimasi total kerugian dengan *ordinary deductible* dan menghitung LER menggunakan simulasi yaitu sebagai berikut:

- (1) Mendeskripsikan data yang diperoleh dari Motorcycle Insurance untuk mendapatkan gambaran umum tentang data dari peubah acak N dan X .
- (2) Menentukan jenis sebaran yang tepat untuk N dan X .
- (3) Menentukan nilai penduga parameter sebaran banyak klaim dan besar klaim dengan metode *maximum likelihood*.
- (4) Menguji kesesuaian sebaran yang digunakan untuk memodelkan banyak klaim menggunakan uji khi-kuadrat dan besar klaim dengan menggunakan uji Kolmogorov Smirnov.
- (5) Membangun sebaran total kerugian tanpa *ordinary deductible*.
- (6) Membangun sebaran total kerugian dengan *ordinary deductible* menggunakan simulasi Monte Carlo.

(7) Menghitung *Loss Elimination Ratio* (LER).

Seluruh perhitungan dilakukan menggunakan bahasa pemrograman R.

4. Hasil dan Pembahasan

4.1. Sebaran banyak klaim N

Data banyak klaim N dari Motorcycle Insurance tahun 1994-1998 yang digunakan dalam penelitian ini dapat dilihat pada Tabel 1.

Tabel 1. Frekuensi banyak klaim

N	Frekuensi	Peluang Empirik
0	63716	0.989625
1	641	0.009956
2	27	0.000419
Total	64384	1.000000

Berdasarkan Tabel 1, terdapat total 64384 polis dengan hampir 99% polis tidak pernah mengajukan klaim, hampir 1% polis mengajukan klaim satu kali dan sangat sedikit sekali polis yang mengajukan klaim dua kali. Tidak ada polis yang mengajukan klaim tiga kali atau lebih.

Mean dan *varians* data berturut-turut adalah 0.01079461 dan 0.01151698. Karena nilai rata-rata lebih kecil dari ragamnya, diduga $N \sim NB(r, \beta)$. Menggunakan metode *maximum likelihood* dan R sebagai alat bantu diperoleh penduga untuk parameternya adalah sebagai berikut.

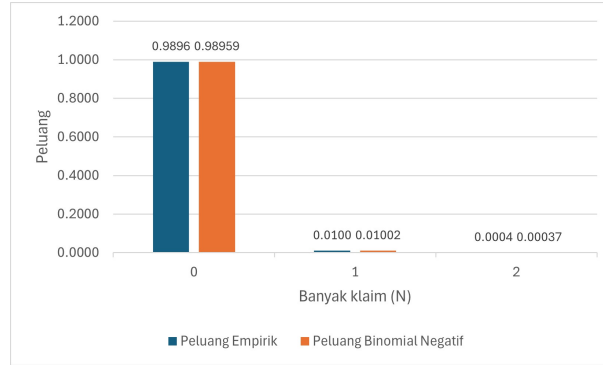
$$\hat{r} = 0.16130, \quad \hat{\beta} = 0.06691. \quad (4.1)$$

Perbandingan peluang sebaran $NB(\hat{r}, \hat{\beta})$ yang dihitung menggunakan Persamaan 2.1 dan peluang empiriknya yang diperoleh dari Tabel 1 ditunjukkan oleh Gambar 1.

Berdasarkan Gambar 1, peluang binomial negatif sangat mirip dengan peluang empirik untuk $N = 0, 1$ dan 2. Hal ini mengindikasikan bahwa sebaran binomial negatif mungkin cocok dalam memodelkan data banyak klaim. Perbandingan antara *mean* dan *varians* empirik dengan binomial negatif yang diperoleh dari Persamaan (2.2) dan Persamaan (2.3) ditunjukkan oleh Tabel 2 berikut.

Tabel 2. Perbandingan *mean* dan *varians* empirik dan binomial negatif

N	Empirik	$NB(\hat{r}, \hat{\beta})$
$E[N]$	0.01079461	0.01079464
$Var(N)$	0.01151698	0.01151701

Gambar 1. Perbandingan peluang $NB(\hat{r}, \hat{\beta})$ dengan peluang empirik

Berdasarkan Tabel 2, *mean* dan varians binomial negatif juga sangat mirip dengan *mean* dan varians empirik. Hal ini menguatkan dugaan bahwa data banyak klaim dapat dimodelkan dengan sebaran binomial negatif.

Untuk menguji dugaan tersebut, digunakan uji khi-kuadrat dengan hipotesis:

$$H_0 : \text{Data banyak klaim menyebar } NB(\hat{r}, \hat{\beta}),$$

$$H_1 : \text{Data banyak klaim tidak menyebar } NB(\hat{r}, \hat{\beta}).$$

Menggunakan pemrograman R diperoleh Tabel 3.

Tabel 3. Uji Khi-kudrat untuk banyak klaim

Peubah	Sebaran	Nilai Dugaan Parameter	<i>p-value</i>
Banyak Klaim N	Binomial Negatif	$\hat{r} = 0.16310$ $\hat{\beta} = 0.006691$	0.8272832

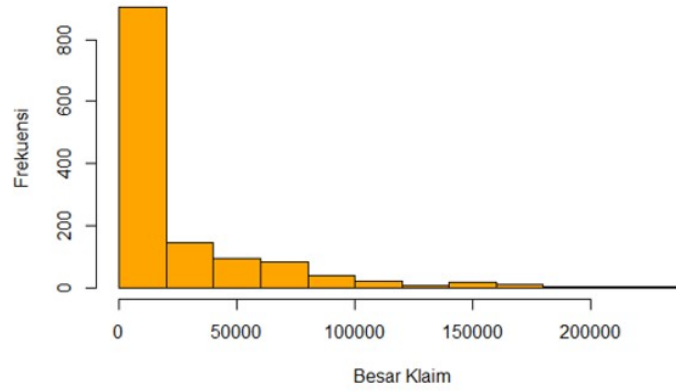
Berdasarkan Tabel 3, *p-value* = 0.8272832 lebih besar dari $\alpha = 0.05$, sehingga dapat disimpulkan bahwa tidak ada bukti untuk menolak H_0 atau data banyak klaim mengikuti sebaran binomial negatif.

4.2. Sebaran besar klaim X

Data besar klaim X (dalam Dolar AS) memiliki nilai minimum sebesar 16, nilai maksimum sebesar 224281.4, *mean* 24481.65, dan varians 1216487602. Histogram besar klaim dari data Motorcycle Insurance pada *package* InsuranceData R ditunjukkan oleh Gambar 2.

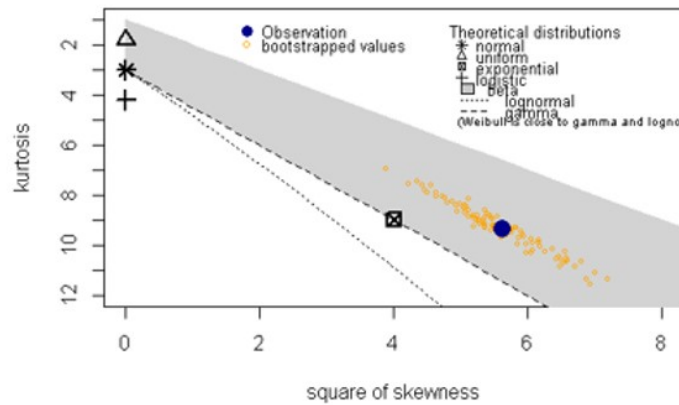
Dengan hanya bermodalkan histogram pada Gambar 2, kita belum dapat memprediksi secara spesifik sebaran apa yang dapat memodelkan data besar klaim. Oleh sebab itu, untuk mencari sebaran yang sesuai untuk X , dengan bantuan R dilakukan analisis menggunakan Cullen and Frey *graph* yang ditunjukkan oleh Gambar 3.

Cullen and Frey *graph* adalah grafik yang menganalisis jenis-jenis sebaran melalui hubungan antara kuadrat *skewness* dan kurtosisnya. Plot kuadrat *skew-*



Gambar 2. Histogram frekuensi besar klaim

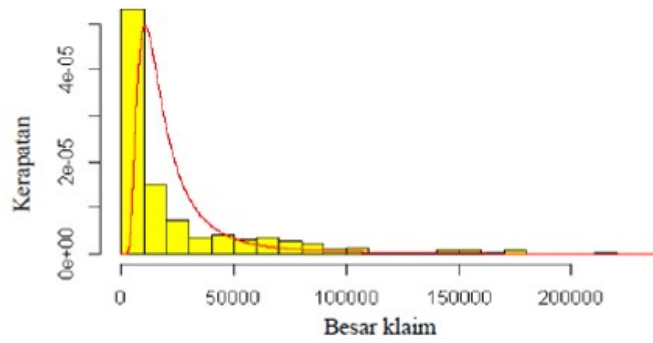
ness dan kurtosis data pada grafik ditandai dengan warna biru. Titik-titik kuning menunjukkan nilai *bootstrapped* yang dihasilkan untuk data. *Bootstrapping* adalah metode *resampling* yang digunakan untuk memperkirakan sebaran statistik dari sampel. Nilai *bootstrap* yang digunakan dalam penelitian ini sebanyak 1000.

Gambar 3. Plot Cullen Frey *graph* untuk X

Dari plot pada Gambar 3, terlihat bahwa titik biru dan kuning terletak di daerah yang berwarna abu-abu yang merupakan daerah keluarga sebaran beta. Sehingga diambil kesimpulan, sebaran yang mungkin cocok adalah sebaran beta. Sebaran beta yang digunakan adalah sebaran *beta-prime* $B'(\alpha, \beta)$. Menggunakan metode *maximum likelihood* diperoleh dugaan parameternya adalah:

$$\hat{\alpha} = 36544.0051, \quad \hat{\beta} = 2.4927. \quad (4.2)$$

Secara visual, Gambar 4 memperlihatkan perbandingan antara histogram peluang dan fungsi kepekatan peluang $B'(\hat{\alpha}, \hat{\beta})$ yang diberikan oleh formula pada Persamaan 2.4. Terlihat bahwa bentuk grafik fungsi kepekatan peluangnya cukup menyerupai bentuk histogram peluang dari data besar klaim. Hal ini mengindikasikan bahwa sebaran *beta-prime* mungkin sesuai dalam memodelkan data besar klaim.



Gambar 4. Perbandingan fungsi kepekatan peluang $B'(\hat{\alpha}, \hat{\beta})$ dengan histogram peluang empirik

Menggunakan formula *mean* dan *varians* pada Persamaan 2.5 dan 2.6, diperoleh perbandingan *mean* dan *varians* empirik dengan sebaran *beta-prime* pada Tabel 4.

Tabel 4. Perbandingan *mean* dan *varians* empirik dan *beta-prime*

X	Empirik	$B'(\hat{\alpha}, \hat{\beta})$
$E[X]$	24481.65	24481.81
$Var(X)$	1216487602	1216528806

Berdasarkan Tabel 4, *mean* dan *varians beta-prime* juga sangat dekat dengan *mean* dan *varians* empirik. Hal ini menguatkan dugaan bahwa data besar klaim dapat dimodelkan dengan sebaran *beta-prime*.

Untuk memastikan bahwa sebaran *beta-prime* adalah sebaran yang sesuai untuk data banyak klaim, dilakukan uji Kolmogorov-Smirnov dengan hipotesis sebagai berikut.

$$H_0 : \text{Data menyebar } B'(\hat{\alpha}, \hat{\beta}),$$

$$H_1 : \text{Data tidak menyebar } B'(\hat{\alpha}, \hat{\beta}).$$

Menggunakan R diperoleh hasil uji Kolmogorov Smirnov pada Tabel 5.

Berdasarkan Tabel 5, *p-value* = 0.6499 lebih besar dari $\alpha = 0.05$, sehingga dapat disimpulkan bahwa tidak ada bukti untuk menolak H_0 atau data besar klaim mengikuti sebaran *beta-prime*.

Tabel 5. Uji Kolmogorof-Smirnov untuk data besar klaim

Peubah	Sebaran	Nilai Dugaan Parameter	p -value
Besar Klaim X	$Beta\text{-prime}$	$\hat{\alpha} = 36544.0051$ $\hat{\beta} = 2.4927$	0.6499

4.3. Total kerugian

Untuk menentukan sebaran total kerugian secara analitik sulit dilakukan, oleh sebab itu dilakukan simulasi Monte Carlo. Simulasi Monte Carlo untuk menentukan total kerugian tanpa menggunakan *ordinary deductible* (S_0) maupun menggunakan *ordinary deductible* d (S_d) dilakukan dengan mengikuti algoritme berikut.

- (1) Tetapkan $d = 0, 1000, 1500, 2000, \dots, 15000$.
- (2) Lakukan untuk $k = 1, 2, 3, \dots, 21000$.
 - (a) Bangkitkan sebuah data banyak klaim N_k dari sebaran $NB(0.1613066, 0.06691)$.
 - (b) Bangkitkan sebanyak N_k data besar klaim $X_{k,1}, X_{k,2}, \dots, X_{k,N_k}$ dari sebaran $B'(36544.01, 2.4927)$.
 - (c) Hitung $S_d^k = \sum_{i=1}^{N_k} Y_{k,i} = \sum_{i=1}^{N_k} (X_{k,i} - d)_+$.
- (3) Hitung $S_d^{sim} = \sum_{k=1}^{21000} \frac{S_d^k}{21000}$.

Karena $E[X] = 24481.65$ dan *ordinary deductible* diterapkan pada X , maka nilai *ordinary deductible* d pada penelitian ini diambil kelipatan 500, mulai dari 1000 sampai dengan 15000.

Untuk menghitung *mean* dan *varians* dari total kerugian tanpa *ordinary deductible* yaitu S_0 , dapat dilakukan secara empirik menggunakan data pada Tabel 1 secara langsung tanpa pemodelan data, secara analitik menggunakan Persamaan 2.10 dan 2.11 dan parameter-parameter besaran yang telah diduga sebelumnya, serta secara numerik menggunakan simulasi Monte Carlo dengan algoritma di atas. *Mean* dan *varians* dari S_0 ditunjukkan oleh Tabel 6 berikut.

Tabel 6. *Mean* dan *varians* dari total kerugian tanpa *ordinary deductible*

S_0	Empirik	Analitik	Simulasi
$E[S_0]$	264.2706	264.2723	266.2606
$Var(S_0)$	20034279	20034814	19982533

Tabel 6 menunjukkan perbandingan *mean* dan *varians* S_0 secara empirik, analitik, dan simulasi yang cukup akurat. Hal ini mengindikasikan bahwa model total kerugian yang dipilih sudah tepat.

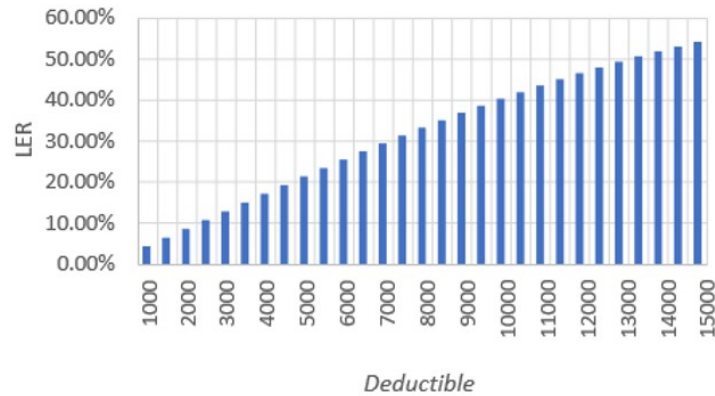
4.4. Loss elimination ratio

Loss elimination ratio (LER) dihitung menggunakan simulasi Monte Carlo yang

diperoleh dari formula:

$$LER = 1 - \frac{S_d^{sim}}{S_0^{sim}}.$$

Untuk nilai *ordinary deductible* d mulai dari 1000 sampai dengan 15000 diperoleh LER yang ditunjukkan oleh Gambar 5.



Gambar 5. Pengaruh *ordinary deductible* terhadap LER

Dari Gambar 5, *loss elimination ratio* mengalami peningkatan seiring besarnya nilai *ordinary deductible*. Ketika diterapkan *ordinary deductible* sebesar 1000, kerugian yang dapat dikurangi sebesar 4.27%. Ketika *ordinary deductible* sebesar 2500, kerugian yang dapat dieliminasi sebesar 10.69%. Bila perusahaan ingin mengurangi sekitar 15% total kerugian, dari Gambar 5 diperoleh *ordinary deductible* yang harus diterapkan adalah 3500.

Dari Gambar 5, penerapan *ordinary deductible* terhadap kontrak asuransi akan mengurangi besar total kerugian. Hal ini terjadi karena perusahaan asuransi tidak membayar jumlah kerugian di bawah nilai *ordinary deductible*. Namun, *ordinary deductible* yang tinggi tidak menarik bagi pemegang polis karena mereka menanggung sebagian kerugian sendiri.

5. Kesimpulan

Metode simulasi Monte Carlo dapat menduga sebaran dari total kerugian dengan membangkitkan sebaran gabungan banyak dan besar klaim. Dalam penelitian ini, sebaran yang menggambarkan banyak klaim adalah sebaran binomial negatif dan besar klaim adalah sebaran *beta-prime*. Ketika perusahaan menetapkan *ordinary deductible*, pemegang polis wajib membayar *ordinary deductible* yang telah ditetapkan pada saat melakukan klaim. Dengan adanya kebijakan *ordinary deductible*, pemegang polis ikut menanggung kerugian, sehingga total kerugian yang dihadapi oleh perusahaan dapat dikurangi. Besarnya nilai *ordinary deductible* berpengaruh terhadap peningkatan nilai *loss elimination ratio*. Dengan meningkatnya *ordinary*

deductible, *loss elimination ratio* meningkat. Namun, *ordinary deductible* yang tinggi tidak menarik bagi pemegang polis karena dia harus menanggung kerugian sendiri yang cukup besar.

6. Ucapan Terima kasih

Terima kasih kepada Departemen Matematika Institut Pertanian Bogor yang mendanai publikasi artikel ini.

Daftar Pustaka

- [1] Klugman, SA., Panjer, HH., Willmot, GE., 2019, *Loss Models: From Data to Decisions*, Edisi ke-5, John Wiley and Sons, Inc., Hoboken NJ, USA
- [2] Gray, R.J., Pitts, SM., 2012, *Risk Modelling in General Insurance*, Edisi ke-1, Cambridge University Press, Cambridge, UK
- [3] Poufinas, T., Gogas, P., Papadimitriou, T., Zaganidis, 2023, Machine learning in forecasting motor insurance claims, *Risks* Vol. **11**(9): 164. <https://doi.org/10.3390/risks11090164>
- [4] Ohlsson, E., Johansson, B., 2010, *Non-life insurance pricing with generalized linear models*, Springer, Berlin, Heidelberg.
- [5] Lee, GY., Frees, EW., 2016, *General Insurance Deductible Ratemaking with Applications to the Local Government Property Insurance Fund*. <https://www.soa.org/globalassets/assets/Files/Research/Projects/research-2016-gi-deductible-ratemaking.pdf>
- [6] Mohamed, MA., Ahmad, MR., Noriszura, I., 2010, Approximation of aggregate losses using simulation, *Journal of Mathematics and Statistics*, **6**(3): 233 – 239
- [7] Shevchenko, PV., 2010, Calculation of aggregate loss distributions, *The Journal of Operational Risk*, **5**(2) : 3 – 40
- [8] Tulupyev, A., Suvorovava, A., Sousa, J., Zeltermann, D. 2013, Beta prime regression with application to risky behaviour frequency screening, *Stat. Med*, **32** (23): 4044 – 4056
- [9] Montgomery, DC., 2001, *Design and Analysis of Experiments*, 5th ed, John Wiley and Sons, New York
- [10] Feller, W., 1948, On the Kolmogorov-Smirnov limit theorems for empirical distributions, *The Annals of Mathematical Statistics*, **19**(2): 177 – 189. doi:10.1214/aoms/1177730243
- [11] Marsaglia, G., Tsang, WW., Wang, L., 2003, Evaluating Kolmogorov's distribution, *Journal of Statistical Software*, **8**(18).doi:10.18637/jss.v008.i18
- [12] Ghahramani, S., 2005, *Fundamentals of Probability with Stochastic Processes*, Edisi ke-3, Prentice Hall, New Jersey, US.