

DAV / DGVFM
Jahrestagung
2025



Fatima ezzahra Kherraz, Dr. Ulrich Riegel

Burning Cost or Pareto?

Jahrestagung, 29. April 2025

Agenda

1. Burning cost for per risk XLs
2. Pareto based pricing approaches
3. Simulation model for evaluating pricing approaches
4. Observations with underlying Pareto severity
5. Observations with other underlying severity distributions

Excess of loss per risk (per risk XL)

We consider the annual loss $S = \sum_{i=1}^N X_i$ of an insurance portfolio. X_i denotes the size of the i -th individual loss.

Excess of loss per risk (per risk XL)

We consider the annual loss $S = \sum_{i=1}^N X_i$ of an insurance portfolio. X_i denotes the size of the i -th individual loss.

- A *per risk XL* has a *deductible* D and a *cover* C

Excess of loss per risk (per risk XL)

We consider the annual loss $S = \sum_{i=1}^N X_i$ of an insurance portfolio. X_i denotes the size of the i -th individual loss.

- A *per risk XL* has a *deductible* D and a *cover* C
- Basic variant:

Reinsurer:

$$\hat{S} = \sum_{i=1}^N \min(C, (X_i - D)^+) =: \sum_{i=1}^N \hat{X}_i, \quad (x^+ := \max(x, 0))$$

Cedent:

$$\tilde{S} = S - \hat{S} = \sum_{i=1}^N (X_i - \hat{X}_i).$$

Excess of loss per risk (per risk XL)

We consider the annual loss $S = \sum_{i=1}^N X_i$ of an insurance portfolio. X_i denotes the size of the i -th individual loss.

- A per risk XL has a deductible D and a cover C
- Basic variant:

Reinsurer:
$$\hat{S} = \sum_{i=1}^N \min(C, (X_i - D)^+) =: \sum_{i=1}^N \hat{X}_i, \quad (x^+ := \max(x, 0))$$

Cedent:
$$\tilde{S} = S - \hat{S} = \sum_{i=1}^N (X_i - \hat{X}_i).$$

- Most per risk XLs have an annual aggregate limit (AAL) for the reinsurer. Sometimes, there is an annual aggregate deductible (AAD) of the cedent.

Excess of loss per risk (per risk XL)

We consider the annual loss $S = \sum_{i=1}^N X_i$ of an insurance portfolio. X_i denotes the size of the i -th individual loss.

- A per risk XL has a deductible D and a cover C
- Basic variant:

Reinsurer:
$$\hat{S} = \sum_{i=1}^N \min(C, (X_i - D)^+) =: \sum_{i=1}^N \hat{X}_i, \quad (x^+ := \max(x, 0))$$

Cedent:
$$\tilde{S} = S - \hat{S} = \sum_{i=1}^N (X_i - \hat{X}_i).$$

- Most per risk XLs have an annual aggregate limit (AAL) for the reinsurer. Sometimes, there is an annual aggregate deductible (AAD) of the cedent.
- C xs D is referred to as a *layer*. A reinsurance program often consists of several layers C_i xs D_i , such that $D_{i+1} = C_i + D_i$.

Excess of loss per risk (per risk XL)

We consider the annual loss $S = \sum_{i=1}^N X_i$ of an insurance portfolio. X_i denotes the size of the i -th individual loss.

- A per risk XL has a deductible D and a cover C
- Basic variant:

Reinsurer:
$$\hat{S} = \sum_{i=1}^N \min(C, (X_i - D)^+) =: \sum_{i=1}^N \hat{X}_i, \quad (x^+ := \max(x, 0))$$

Cedent:
$$\tilde{S} = S - \hat{S} = \sum_{i=1}^N (X_i - \hat{X}_i).$$

- Most per risk XLs have an annual aggregate limit (AAL) for the reinsurer. Sometimes, there is an annual aggregate deductible (AAD) of the cedent.
- C xs D is referred to as a *layer*. A reinsurance program often consists of several layers C_i xs D_i , such that $D_{i+1} = C_i + D_i$.
- Purpose of a per risk XL: Protection against large losses.

Experience rating

Basic idea:

Estimate the expected loss of a per risk XL in the prospective treaty year based on the individual loss experience of the cedent in a certain *observation period*.

Experience rating

Basic idea:

Estimate the expected loss of a per risk XL in the prospective treaty year based on the individual loss experience of the cedent in a certain *observation period*.

Starting point:

- Observation period: Accident years $1, \dots, n$
- Current year: $n + 1$
- Prospective treaty year: $q = n + 2$

Experience rating

Basic idea:

Estimate the expected loss of a per risk XL in the prospective treaty year based on the individual loss experience of the cedent in a certain *observation period*.

Starting point:

- Observation period: Accident years $1, \dots, n$
- Current year: $n+1$
- Prospective treaty year: $q = n+2$
- Volumes $v_1, \dots, v_q$ available (e.g., premiums, sums insured, vehicle counts)

Experience rating

Basic idea:

Estimate the expected loss of a per risk XL in the prospective treaty year based on the individual loss experience of the cedent in a certain *observation period*.

Starting point:

- Observation period: Accident years $1, \dots, n$
- Current year: $n + 1$
- Prospective treaty year: $q = n + 2$
- Volumes $v_1, \dots, v_q$ available (e.g., premiums, sums insured, vehicle counts)
- $X_{i,1}, \dots, X_{i,N_i}$ observed individual losses in the year $i \in \{1, \dots, n\}$

Experience rating

Basic idea:

Estimate the expected loss of a per risk XL in the prospective treaty year based on the individual loss experience of the cedent in a certain *observation period*.

Starting point:

- Observation period: Accident years $1, \dots, n$
- Current year: $n + 1$
- Prospective treaty year: $q = n + 2$
- Volumes $v_1, \dots, v_q$ available (e.g., premiums, sums insured, vehicle counts)
- $X_{i,1}, \dots, X_{i,N_i}$ observed individual losses in the year $i \in \{1, \dots, n\}$

We assume that the volumes and the losses have been *indexed* to the cost level of the prospective treaty year q .

Burning Cost (BC)

An important step in the pricing of the layer C xs D is the estimation of the expected layer loss $E(\hat{S}_q)$ in the prospective treaty year q before AAD/AAL.

Burning Cost (BC)

An important step in the pricing of the layer C xs D is the estimation of the expected layer loss $E(\hat{S}_q)$ in the prospective treaty year q before AAD/AAL.

Burning cost (BC) estimator:

$$\hat{E}_{C \text{ xs } D}^{\text{BC}} := v_q \cdot \frac{\sum_{i=1}^n \sum_{k=1}^{N_i} \min(C, (X_{i,k} - D)^+)}{\sum_{i=1}^n v_n}$$

Burning Cost (BC)

An important step in the pricing of the layer C xs D is the estimation of the expected layer loss $E(\hat{S}_q)$ in the prospective treaty year q before AAD/AAL.

Burning cost (BC) estimator:

$$\hat{E}_{C \text{ xs } D}^{\text{BC}} := v_q \cdot \frac{\sum_{i=1}^n \sum_{k=1}^{N_i} \min(C, (X_{i,k} - D)^+)}{\sum_{i=1}^n v_n}$$

Assumptions:

- The accident years are independent

Burning Cost (BC)

An important step in the pricing of the layer C xs D is the estimation of the expected layer loss $E(\hat{S}_q)$ in the prospective treaty year q before AAD/AAL.

Burning cost (BC) estimator:

$$\hat{E}_{C \text{ xs } D}^{\text{BC}} := v_q \cdot \frac{\sum_{i=1}^n \sum_{k=1}^{N_i} \min(C, (X_{i,k} - D)^+)}{\sum_{i=1}^n v_n}$$

Assumptions:

- The accident years are independent
- The random sums $\sum_{k=1}^{N_i} X_{i,k}$ are collective models for all accident years i

Burning Cost (BC)

An important step in the pricing of the layer C xs D is the estimation of the expected layer loss $E(\hat{S}_q)$ in the prospective treaty year q before AAD/AAL.

Burning cost (BC) estimator:

$$\hat{E}_{C \text{ xs } D}^{\text{BC}} := v_q \cdot \frac{\sum_{i=1}^n \sum_{k=1}^{N_i} \min(C, (X_{i,k} - D)^+)}{\sum_{i=1}^n v_n}$$

Assumptions:

- The accident years are independent
- The random sums $\sum_{k=1}^{N_i} X_{i,k}$ are collective models for all accident years i
- There is a $\lambda > 0$, such that $N_i \sim \text{Poisson}(\lambda v_i)$ for all i

Burning Cost (BC)

An important step in the pricing of the layer C xs D is the estimation of the expected layer loss $E(\hat{S}_q)$ in the prospective treaty year q before AAD/AAL.

Burning cost (BC) estimator:

$$\hat{E}_{C \text{ xs } D}^{\text{BC}} := v_q \cdot \frac{\sum_{i=1}^n \sum_{k=1}^{N_i} \min(C, (X_{i,k} - D)^+)}{\sum_{i=1}^n v_n}$$

Assumptions:

- The accident years are independent
- The random sums $\sum_{k=1}^{N_i} X_{i,k}$ are collective models for all accident years i
- There is a $\lambda > 0$, such that $N_i \sim \text{Poisson}(\lambda v_i)$ for all i
- The loss sizes $X_{i,k}$ are identically distributed for all accident years i

Burning Cost (BC)

An important step in the pricing of the layer C xs D is the estimation of the expected layer loss $E(\hat{S}_q)$ in the prospective treaty year q before AAD/AAL.

Burning cost (BC) estimator:

$$\hat{E}_{C \text{ xs } D}^{\text{BC}} := v_q \cdot \frac{\sum_{i=1}^n \sum_{k=1}^{N_i} \min(C, (X_{i,k} - D)^+)}{\sum_{i=1}^n v_i}$$

Assumptions:

- The accident years are independent
- The random sums $\sum_{k=1}^{N_i} X_{i,k}$ are collective models for all accident years i
- There is a $\lambda > 0$, such that $N_i \sim \text{Poisson}(\lambda v_i)$ for all i
- The loss sizes $X_{i,k}$ are identically distributed for all accident years i

Under these assumptions, the BC estimator is unbiased! The assumptions are plausible if indexation works perfectly and if the v_i are good volumes.

Agenda

1. Burning cost for per risk XLs
2. Pareto based pricing approaches
3. Simulation model for evaluating pricing approaches
4. Observations with underlying Pareto severity
5. Observations with other underlying severity distributions

Pareto distribution

Die Pareto distribution $\text{Pareto}(t, \alpha)$ is the most widely used severity distribution for modelling large losses. CDF:

$$F(x) = 1 - \left(\frac{t}{x}\right)^\alpha, \quad x \geq t.$$

Proposition:

- Let $X \sim \text{Pareto}(t, \alpha)$. Then $cX \sim \text{Pareto}(ct, \alpha)$ for all $c > 0$.

Pareto distribution

Die Pareto distribution $\text{Pareto}(t, \alpha)$ is the most widely used severity distribution for modelling large losses. CDF:

$$F(x) = 1 - \left(\frac{t}{x}\right)^\alpha, \quad x \geq t.$$

Proposition:

- Let $X \sim \text{Pareto}(t, \alpha)$. Then $cX \sim \text{Pareto}(ct, \alpha)$ for all $c > 0$.
- Let $X \sim \text{Pareto}(t_1, \alpha)$. For $t_2 > t_1$ we then have $X|(X > t_2) \sim \text{Pareto}(t_2, \alpha)$.

Pareto distribution

Die Pareto distribution $\text{Pareto}(t, \alpha)$ is the most widely used severity distribution for modelling large losses. CDF:

$$F(x) = 1 - \left(\frac{t}{x}\right)^\alpha, \quad x \geq t.$$

Proposition:

- Let $X \sim \text{Pareto}(t, \alpha)$. Then $cX \sim \text{Pareto}(ct, \alpha)$ for all $c > 0$.
- Let $X \sim \text{Pareto}(t_1, \alpha)$. For $t_2 > t_1$ we then have $X|(X > t_2) \sim \text{Pareto}(t_2, \alpha)$.

Consequences:

- The Pareto alpha is invariant wrt scaling (in particular wrt FX conversion and inflation)

Pareto distribution

Die Pareto distribution $\text{Pareto}(t, \alpha)$ is the most widely used severity distribution for modelling large losses. CDF:

$$F(x) = 1 - \left(\frac{t}{x}\right)^\alpha, \quad x \geq t.$$

Proposition:

- Let $X \sim \text{Pareto}(t, \alpha)$. Then $cX \sim \text{Pareto}(ct, \alpha)$ for all $c > 0$.
- Let $X \sim \text{Pareto}(t_1, \alpha)$. For $t_2 > t_1$ we then have $X|(X > t_2) \sim \text{Pareto}(t_2, \alpha)$.

Consequences:

- The Pareto alpha is invariant wrt scaling (in particular wrt FX conversion and inflation)
- For Pareto distributed data, the Pareto alpha does not depend on the reporting threshold

Pareto distribution

Die Pareto distribution $\text{Pareto}(t, \alpha)$ is the most widely used severity distribution for modelling large losses. CDF:

$$F(x) = 1 - \left(\frac{t}{x}\right)^\alpha, \quad x \geq t.$$

Proposition:

- Let $X \sim \text{Pareto}(t, \alpha)$. Then $cX \sim \text{Pareto}(ct, \alpha)$ for all $c > 0$.
- Let $X \sim \text{Pareto}(t_1, \alpha)$. For $t_2 > t_1$ we then have $X|(X > t_2) \sim \text{Pareto}(t_2, \alpha)$.

Consequences:

- The Pareto alpha is invariant wrt scaling (in particular wrt FX conversion and inflation)
- For Pareto distributed data, the Pareto alpha does not depend on the reporting threshold
- The Pareto distribution has only one 'real' parameter: the Pareto alpha

Pareto distribution

Die Pareto distribution $\text{Pareto}(t, \alpha)$ is the most widely used severity distribution for modelling large losses. CDF:

$$F(x) = 1 - \left(\frac{t}{x}\right)^\alpha, \quad x \geq t.$$

Proposition:

- Let $X \sim \text{Pareto}(t, \alpha)$. Then $cX \sim \text{Pareto}(ct, \alpha)$ for all $c > 0$.
- Let $X \sim \text{Pareto}(t_1, \alpha)$. For $t_2 > t_1$ we then have $X|(X > t_2) \sim \text{Pareto}(t_2, \alpha)$.

Consequences:

- The Pareto alpha is invariant wrt scaling (in particular wrt FX conversion and inflation)
- For Pareto distributed data, the Pareto alpha does not depend on the reporting threshold
- The Pareto distribution has only one 'real' parameter: the Pareto alpha

Due to these properties, the Pareto distribution is very useful for comparing tails. It is often possible to provide 'market alphas' for various segments.

Pareto distribution

Proposition (Maximum likelihood estimation of α):

Let $t > 0$ and $X_1, \dots, X_n \sim \text{Pareto}(t, \alpha)$. The *maximum likelihood estimator (MLE)* for α is given by

$$\hat{\alpha}^{\text{ML}} := \frac{n}{\sum_{i=1}^n \ln(X_i/t)}.$$

Pareto distribution

Proposition (Maximum likelihood estimation of α):

Let $t > 0$ and $X_1, \dots, X_n \sim \text{Pareto}(t, \alpha)$. The *maximum likelihood estimator (MLE)* for α is given by

$$\hat{\alpha}^{\text{ML}} := \frac{n}{\sum_{i=1}^n \ln(X_i/t)}.$$

Alternatively, one can use the *bias corrected MLE*

$$\hat{\alpha}^* := \frac{n-1}{n} \cdot \hat{\alpha}^{\text{ML}} = \frac{n-1}{\sum_{i=1}^n \ln(X_i/t)}.$$

Pareto distribution

Proposition (Maximum likelihood estimation of α):

Let $t > 0$ and $X_1, \dots, X_n \sim \text{Pareto}(t, \alpha)$. The *maximum likelihood estimator (MLE)* for α is given by

$$\hat{\alpha}^{\text{ML}} := \frac{n}{\sum_{i=1}^n \ln(X_i/t)}.$$

Alternatively, one can use the *bias corrected MLE*

$$\hat{\alpha}^* := \frac{n-1}{n} \cdot \hat{\alpha}^{\text{ML}} = \frac{n-1}{\sum_{i=1}^n \ln(X_i/t)}.$$

The bias corrected MLE $\hat{\alpha}^*$ is unbiased and has a lower variance than $\hat{\alpha}^{\text{ML}}$. Hence, $\hat{\alpha}^*$ is a better estimator than $\hat{\alpha}^{\text{ML}}$.

M. Rytgaard (1990) Estimation in the Pareto Distribution, ASTIN Bulletin: The Journal of the IAA, 20(2), pp. 201–216

Pareto distribution

Proposition (layer mean): Let $X \sim \text{Pareto}(t, \alpha)$. For a layer C vs D with $D \geq t$ we then have:

$$E(\min(C; (X - D)^+)) = \begin{cases} \frac{t^\alpha}{1 - \alpha} \left((C + D)^{1-\alpha} - D^{1-\alpha} \right) & \text{if } \alpha \neq 1 \\ t(\ln(C + D) - \ln(D)) & \text{if } \alpha = 1. \end{cases}$$

Pareto distribution

Proposition (layer mean): Let $X \sim \text{Pareto}(t, \alpha)$. For a layer C xs D with $D \geq t$ we then have:

$$E(\min(C; (X - D)^+)) = \begin{cases} \frac{t^\alpha}{1 - \alpha} \left((C + D)^{1-\alpha} - D^{1-\alpha} \right) & \text{if } \alpha \neq 1 \\ t(\ln(C + D) - \ln(D)) & \text{if } \alpha = 1. \end{cases}$$

Proposition (Pareto extrapolation): Let $\sum_{n=1}^N X_n$ be a collective model with $X_n \sim X \sim \text{Pareto}(t, \alpha)$. For layers C_i xs D_i with $D_i \geq t$ we then have

$$\Phi_{C_1 \text{ xs } D_1}^{C_2 \text{ xs } D_2}(\alpha) := \frac{\text{Exp. loss in } C_2 \text{ xs } D_2}{\text{Exp. loss in } C_1 \text{ xs } D_1} = \begin{cases} \frac{(C_2 + D_2)^{1-\alpha} - D_2^{1-\alpha}}{(C_1 + D_1)^{1-\alpha} - D_1^{1-\alpha}} & \text{if } \alpha \neq 1 \\ \frac{\ln(C_2 + D_2) - \ln(D_2)}{\ln(C_1 + D_1) - \ln(D_1)} & \text{if } \alpha = 1 \end{cases}$$

Approach 1: Poisson/Pareto modell

For pricing C xs D select a large loss threshold $t \leq D$ and model the large losses in the prospective treaty year q with a collective model $\sum_{i=1}^{N_q} X_{q,i}$:

Approach 1: Poisson/Pareto modell

For pricing C vs D select a large loss threshold $t \leq D$ and model the large losses in the prospective treaty year q with a collective model $\sum_{i=1}^{N_q} X_{q,i}$:

- $X_{q,i} \sim \text{Pareto}(t, \hat{\alpha})$, where $\hat{\alpha}$ denotes the MLE or the bias corrected MLE for α (calculated with the indexed large losses from the observation period $\{1, \dots, n\}$).

Approach 1: Poisson/Pareto modell

For pricing C vs D select a large loss threshold $t \leq D$ and model the large losses in the prospective treaty year q with a collective model $\sum_{i=1}^{N_q} X_{q,i}$:

- $X_{q,i} \sim \text{Pareto}(t, \hat{\alpha})$, where $\hat{\alpha}$ denotes the MLE or the bias corrected MLE for α (calculated with the indexed large losses from the observation period $\{1, \dots, n\}$).
- $N_q \sim \text{Poisson}(\hat{\lambda} v_q)$, where

$$\hat{\lambda} := \frac{\sum_{i=1}^n \sum_{k=1}^{N_i} \chi_{\{X_{i,k} > t\}}}{\sum_{i=1}^n v_i}$$

denotes the MLE of λ .

Approach 1: Poisson/Pareto modell

For pricing C xs D select a large loss threshold $t \leq D$ and model the large losses in the prospective treaty year q with a collective model $\sum_{i=1}^{N_q} X_{q,i}$:

- $X_{q,i} \sim \text{Pareto}(t, \hat{\alpha})$, where $\hat{\alpha}$ denotes the MLE or the bias corrected MLE for α (calculated with the indexed large losses from the observation period $\{1, \dots, n\}$).
- $N_q \sim \text{Poisson}(\hat{\lambda} v_q)$, where

$$\hat{\lambda} := \frac{\sum_{i=1}^n \sum_{k=1}^{N_i} \chi_{\{X_{i,k} > t\}}}{\sum_{i=1}^n v_i}$$

denotes the MLE of λ .

The expected loss in C xs D (before AAD/AAL) is then estimated by

$$\hat{E}_{C \text{ xs } D}^{\text{PP}} := \hat{\lambda} v_q \cdot E(\min(C; (X_{q,1} - D)^+)).$$

Approach 2: Pareto extrapolation

Select a large loss threshold t like in Approach 1 and a $T > t$, such that the BC estimator $\hat{E}_{c_0 \times s \ d_0}^{BC}$ is reasonable for the layer $c_0 \times s \ d_0$ with $d_0 := t$ and $c_0 := T - t$.

Approach 2: Pareto extrapolation

Select a large loss threshold t like in Approach 1 and a $T > t$, such that the BC estimator $\hat{E}_{c_0 \times s \ d_0}^{BC}$ is reasonable for the layer $c_0 \times s \ d_0$ with $d_0 := t$ and $c_0 := T - t$. Choose an estimator $\hat{\alpha}$ for α .

Approach 2: Pareto extrapolation

Select a large loss threshold t like in Approach 1 and a $T > t$, such that the BC estimator $\hat{E}_{c_0 \text{ xs } d_0}^{\text{BC}}$ is reasonable for the layer $c_0 \text{ xs } d_0$ with $d_0 := t$ and $c_0 := T - t$. Choose an estimator $\hat{\alpha}$ for α .

Estimator for the expected layer loss in $C \text{ xs } D$ (befor AAD/AAL):

- Case $C + D \leq T$: Use the BC estimator

$$\hat{E}_{C \text{ xs } D}^{\text{PE}} := \hat{E}_{C \text{ xs } D}^{\text{BC}}$$

Approach 2: Pareto extrapolation

Select a large loss threshold t like in Approach 1 and a $T > t$, such that the BC estimator $\hat{E}_{c_0 \text{ xs } d_0}^{\text{BC}}$ is reasonable for the layer $c_0 \text{ xs } d_0$ with $d_0 := t$ and $c_0 := T - t$. Choose an estimator $\hat{\alpha}$ for α .

Estimator for the expected layer loss in $C \text{ xs } D$ (befor AAD/AAL):

- Case $C + D \leq T$: Use the BC estimator

$$\hat{E}_{C \text{ xs } D}^{\text{PE}} := \hat{E}_{C \text{ xs } D}^{\text{BC}}$$

- Case $T \leq D$: Use Pareto extrapolation from the layer $c_0 \text{ xs } d_0$ for the entire layer

$$\hat{E}_{C \text{ xs } D}^{\text{PE}} := \hat{E}_{c_0 \text{ xs } d_0}^{\text{BC}} \cdot \Phi_{c_0 \text{ xs } d_0}^{C \text{ xs } D}(\hat{\alpha})$$

Approach 2: Pareto extrapolation

Select a large loss threshold t like in Approach 1 and a $T > t$, such that the BC estimator $\hat{E}_{c_0 \text{ xs } d_0}^{\text{BC}}$ is reasonable for the layer $c_0 \text{ xs } d_0$ with $d_0 := t$ and $c_0 := T - t$. Choose an estimator $\hat{\alpha}$ for α .

Estimator for the expected layer loss in $C \text{ xs } D$ (befor AAD/AAL):

- Case $C + D \leq T$: Use the BC estimator

$$\hat{E}_{C \text{ xs } D}^{\text{PE}} := \hat{E}_{C \text{ xs } D}^{\text{BC}}$$

- Case $T \leq D$: Use Pareto extrapolation from the layer $c_0 \text{ xs } d_0$ for the entire layer

$$\hat{E}_{C \text{ xs } D}^{\text{PE}} := \hat{E}_{c_0 \text{ xs } d_0}^{\text{BC}} \cdot \Phi_{c_0 \text{ xs } d_0}^{C \text{ xs } D}(\hat{\alpha})$$

- Case $T \in (D, C + D)$: Use Pareto extrapolation for the layer part above T and BC for the layer part below T

$$\hat{E}_{C \text{ xs } D}^{\text{PE}} := \hat{E}_{T-D \text{ xs } D}^{\text{BC}} + \hat{E}_{c_0 \text{ xs } d_0}^{\text{BC}} \cdot \Phi_{c_0 \text{ xs } d_0}^{C+D-T \text{ xs } T}(\hat{\alpha})$$

Variants of the Approaches 1 and 2

Variant 1: Assume there is market knowledge about a 'typical' alpha in the considered segment. Then, the pricing actuary will not use fitted alphas that deviate too much from the typical alpha.

Variants of the Approaches 1 and 2

Variant 1: Assume there is market knowledge about a ‘typical’ alpha in the considered segment. Then, the pricing actuary will not use fitted alphas that deviate too much from the typical alpha. This corresponds to the selection of $\alpha_{\min} \leq \alpha_{\max}$ and the use of

$$\min(\alpha_{\max}, \max(\alpha_{\min}, \hat{\alpha}^{\text{ML}})) \quad \text{or} \quad \min(\alpha_{\max}, \max(\alpha_{\min}, \hat{\alpha}^*)).$$

The extreme case $\alpha_{\min} = \alpha_{\max}$ corresponds to the use of a market alpha.

Variants of the Approaches 1 and 2

Variant 1: Assume there is market knowledge about a ‘typical’ alpha in the considered segment. Then, the pricing actuary will not use fitted alphas that deviate too much from the typical alpha. This corresponds to the selection of $\alpha_{\min} \leq \alpha_{\max}$ and the use of

$$\min(\alpha_{\max}, \max(\alpha_{\min}, \hat{\alpha}^{\text{ML}})) \quad \text{or} \quad \min(\alpha_{\max}, \max(\alpha_{\min}, \hat{\alpha}^*)).$$

The extreme case $\alpha_{\min} = \alpha_{\max}$ corresponds to the use of a market alpha.

Variant 2 (BC/Pareto): Like in Approach 2, select a $T > t$ (depending on the loss history) and use the BC estimator below T and the Poisson/Pareto model from Approach 1 above T .

Variants of the Approaches 1 and 2

Variant 1: Assume there is market knowledge about a ‘typical’ alpha in the considered segment. Then, the pricing actuary will not use fitted alphas that deviate too much from the typical alpha. This corresponds to the selection of $\alpha_{\min} \leq \alpha_{\max}$ and the use of

$$\min(\alpha_{\max}, \max(\alpha_{\min}, \hat{\alpha}^{\text{ML}})) \quad \text{or} \quad \min(\alpha_{\max}, \max(\alpha_{\min}, \hat{\alpha}^*)).$$

The extreme case $\alpha_{\min} = \alpha_{\max}$ corresponds to the use of a market alpha.

Variant 2 (BC/Pareto): Like in Approach 2, select a $T > t$ (depending on the loss history) and use the BC estimator below T and the Poisson/Pareto model from Approach 1 above T .

Rules of thumb for the selection of T :

- T = largest loss
- T = 3rd largest loss
- T = ...

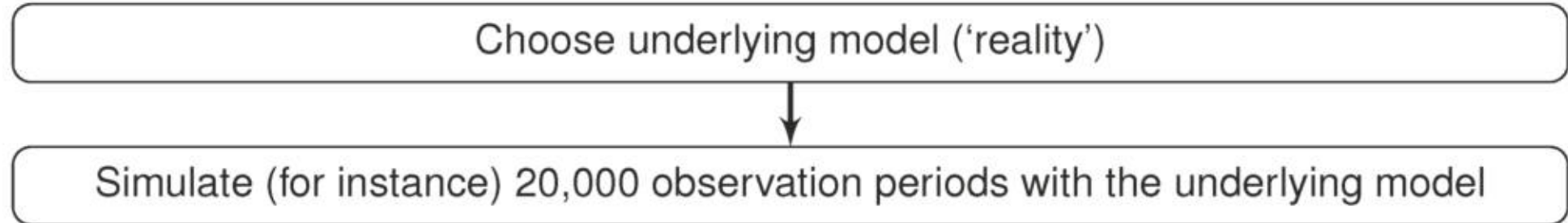
Agenda

1. Burning cost for per risk XLs
2. Pareto based pricing approaches
3. Simulation model for evaluating pricing approaches
4. Observations with underlying Pareto severity
5. Observations with other underlying severity distributions

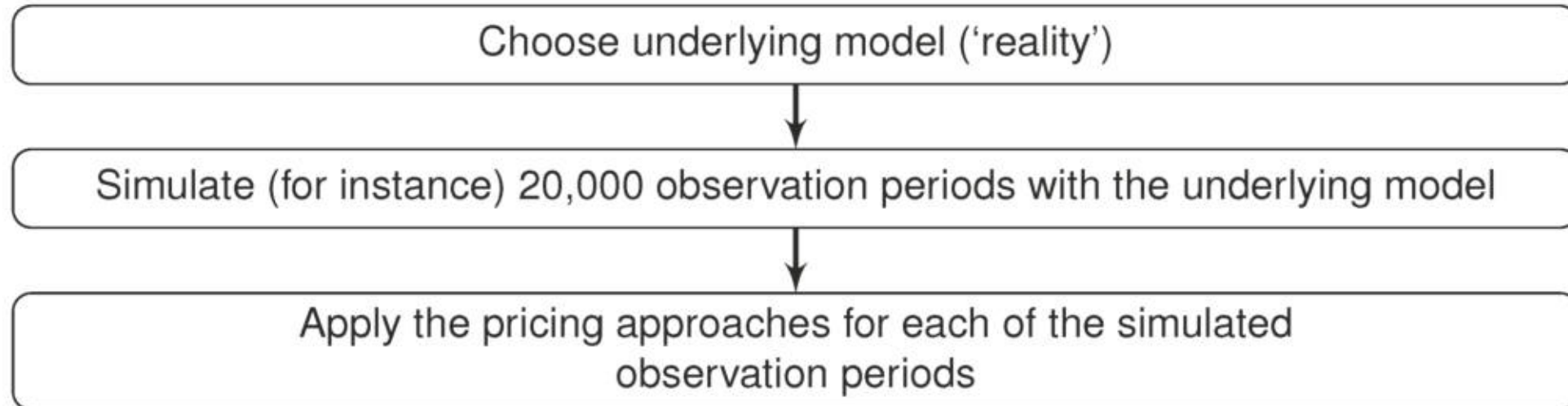
Simulation model for evaluating pricing approaches

Choose underlying model ('reality')

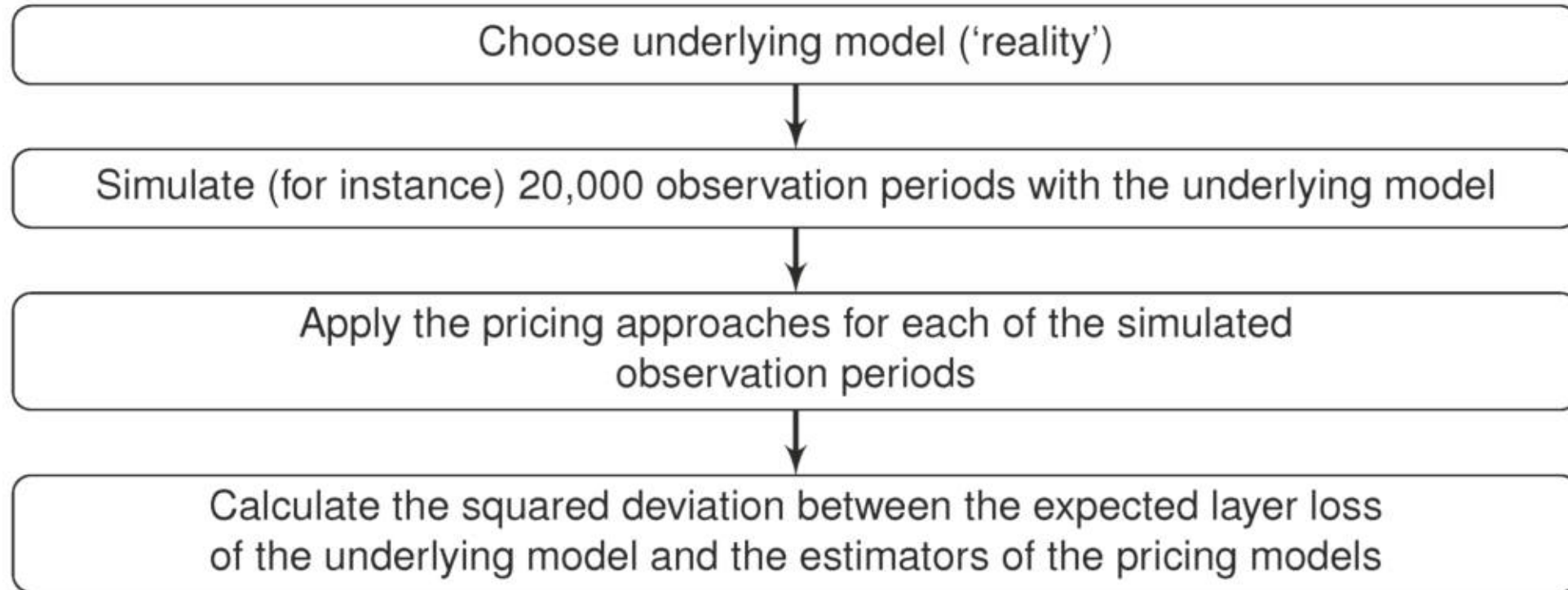
Simulation model for evaluating pricing approaches



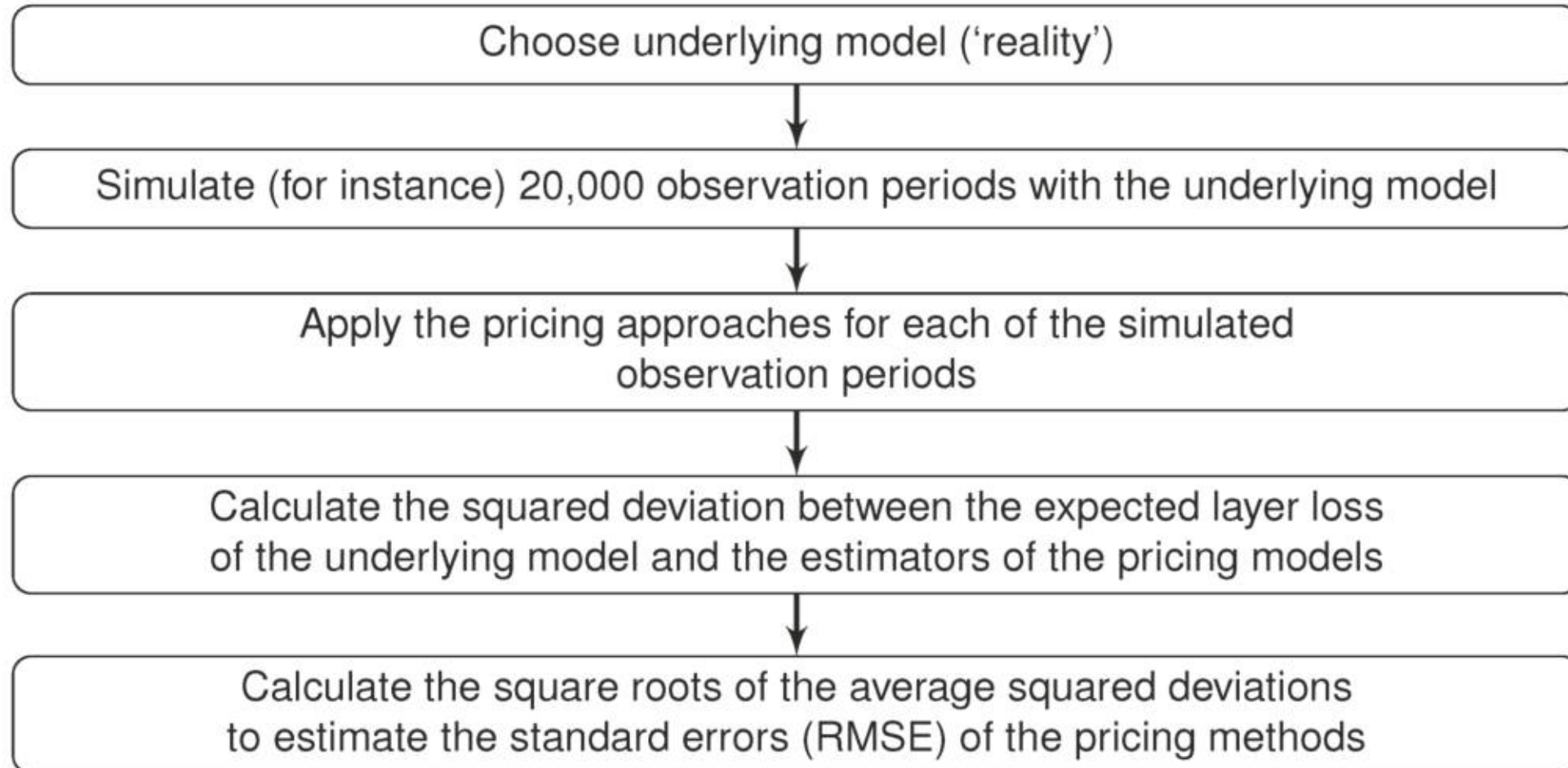
Simulation model for evaluating pricing approaches



Simulation model for evaluating pricing approaches



Simulation model for evaluating pricing approaches



Agenda

1. Burning cost for per risk XLs
2. Pareto based pricing approaches
3. Simulation model for evaluating pricing approaches
4. Observations with underlying Pareto severity
5. Observations with other underlying severity distributions

Simulation settings

For simplicity, we use the assumption

$$v_q = \sum_{i=1}^n v_i = 1.$$

Simulation settings

For simplicity, we use the assumption

$$v_q = \sum_{i=1}^n v_i = 1.$$

We simulate observation periods with

- claim count $N \sim \text{Poisson}(\lambda)$ ($\lambda \in \{5, 15, 50\}$)

Simulation settings

For simplicity, we use the assumption

$$v_q = \sum_{i=1}^n v_i = 1.$$

We simulate observation periods with

- claim count $N \sim \text{Poisson}(\lambda)$ ($\lambda \in \{5, 15, 50\}$)
- and severity $X_n \sim \text{Pareto}(t, \alpha)$ with $t = 1,000$ and $\alpha = 2.0$

Simulation settings

For simplicity, we use the assumption

$$v_q = \sum_{i=1}^n v_i = 1.$$

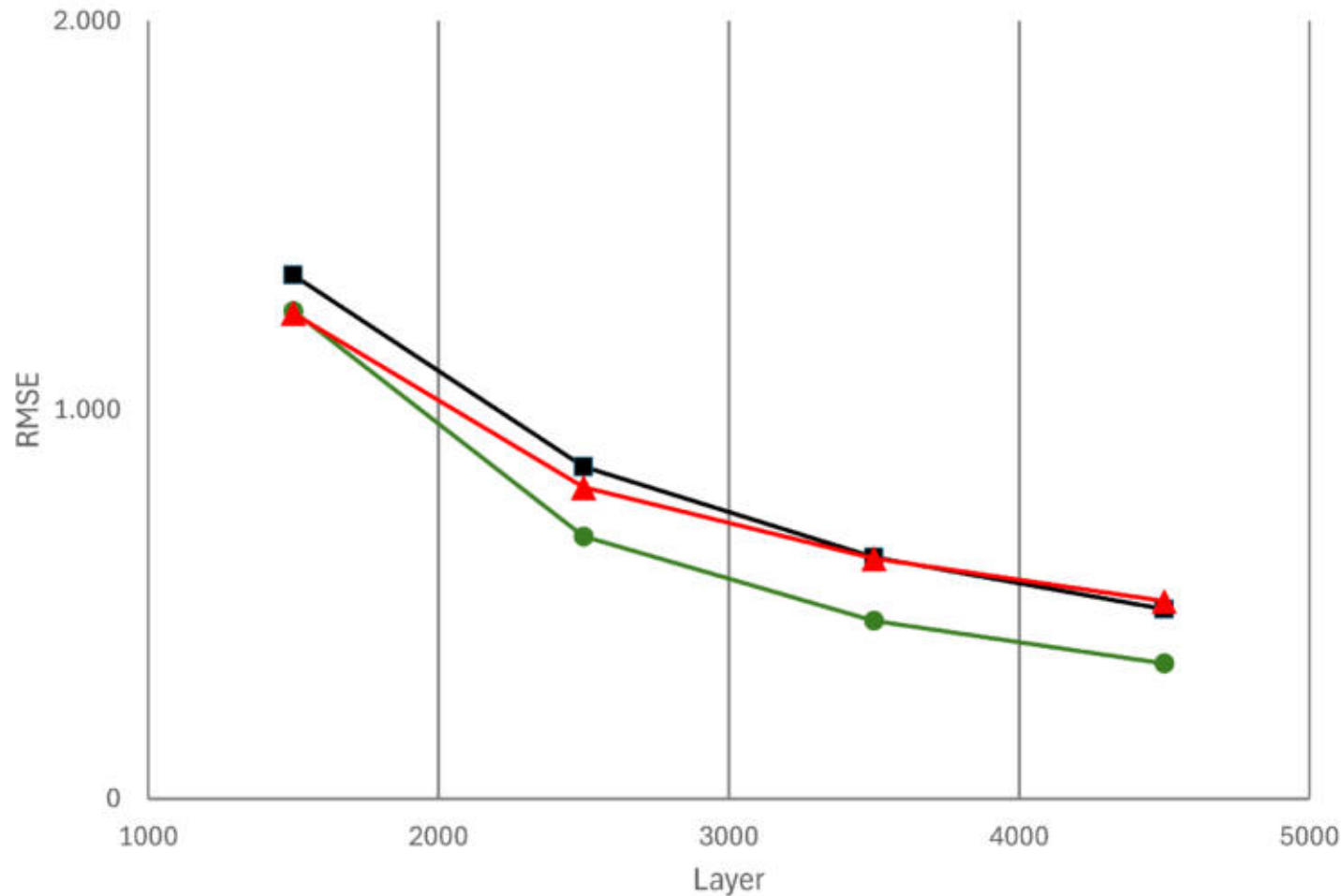
We simulate observation periods with

- claim count $N \sim \text{Poisson}(\lambda)$ ($\lambda \in \{5, 15, 50\}$)
- and severity $X_n \sim \text{Pareto}(t, \alpha)$ with $t = 1,000$ and $\alpha = 2.0$

and use the pricing approaches to estimate the expected losses of the following layers:

- 1,000 xs 1,000
- 1,000 xs 2,000
- 1,000 xs 3,000
- 1,000 xs 4,000

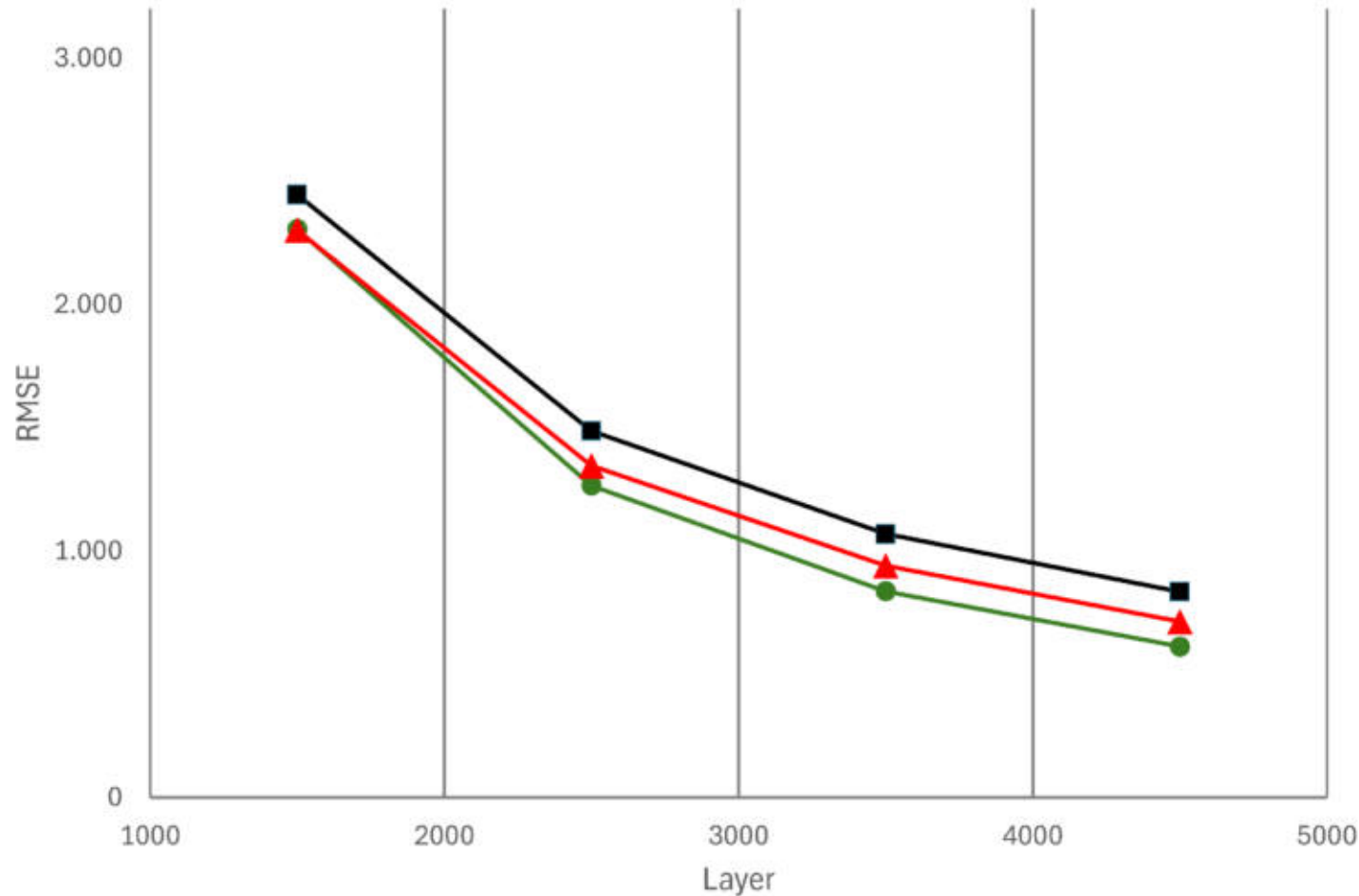
RMSE with $\lambda = 5$, $\alpha = 2.0$



ML estimator $\hat{\alpha}^{\text{ML}}$ performs better than the bias corrected ML estimator $\hat{\alpha}^*$!

- Burning Cost
- Poisson Pareto (ML)
- ▲ Poisson Pareto (Bias Corrected)

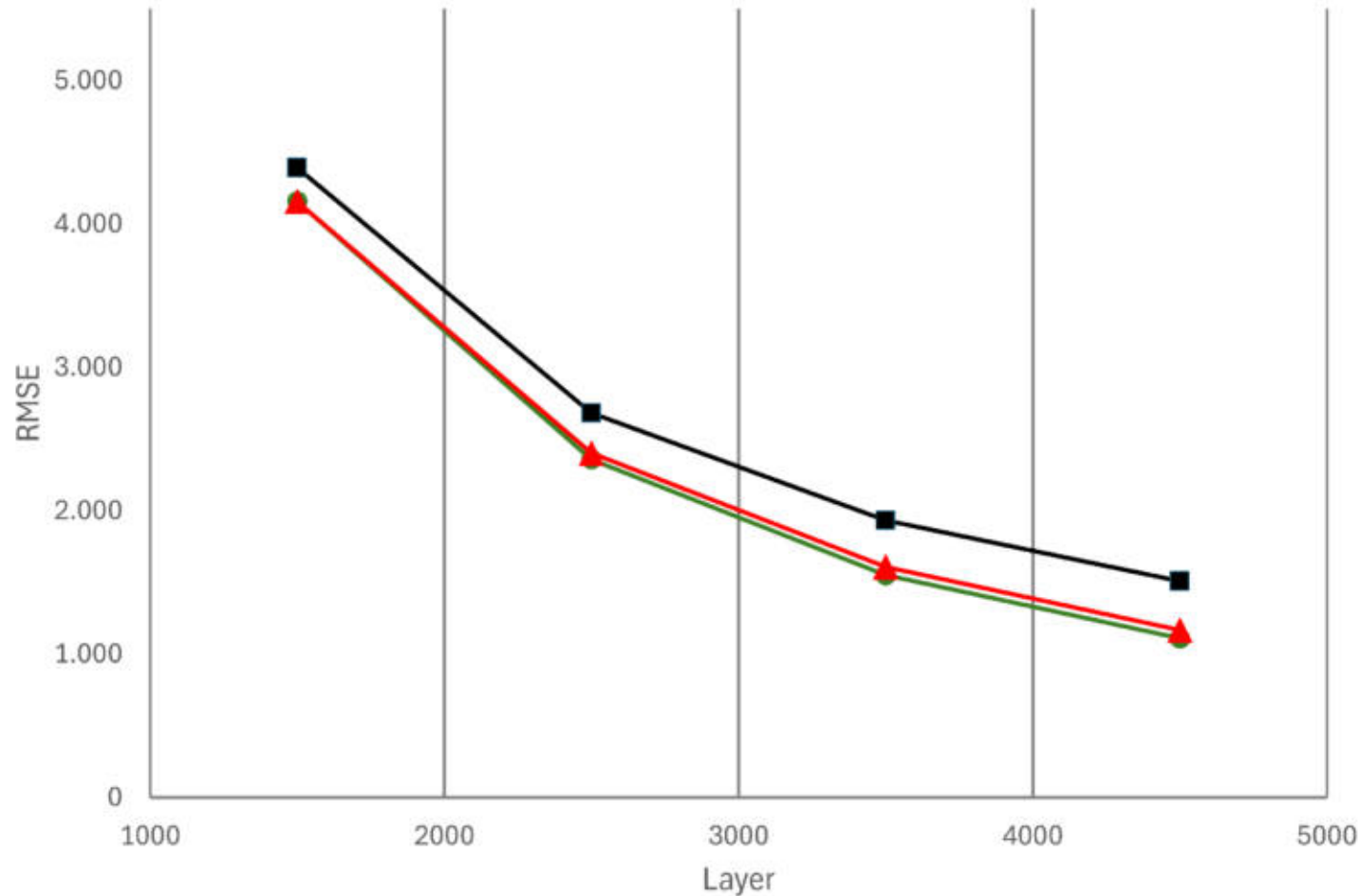
RMSE with $\lambda = 15$, $\alpha = 2.0$



ML estimator $\hat{\alpha}^{\text{ML}}$ performs better than the bias corrected ML estimator $\hat{\alpha}^*$!

- Burning Cost
- Poisson Pareto (ML)
- ▲ Poisson Pareto (Bias Corrected)

RMSE with $\lambda = 50$, $\alpha = 2.0$



The estimators $\hat{\alpha}^{\text{ML}}$ and $\hat{\alpha}^*$ are very similar for large λ .

- Burning Cost
- Poisson Pareto (ML)
- ▲ Poisson Pareto (Bias Corrected)

Observation 1

The Pareto based methods seem to perform better than the Burning Cost.

Observation 1

The Pareto based methods seem to perform better than the Burning Cost.

This is not surprising, as the underlying distribution is Pareto . . .

Observation 2

The ML estimator $\hat{\alpha}^{\text{ML}}$ typically performs better than the bias corrected ML estimator $\hat{\alpha}^*$.

Observation 2

The ML estimator $\hat{\alpha}^{\text{ML}}$ typically performs better than the bias corrected ML estimator $\hat{\alpha}^*$.

At first sight, this is surprising as $\hat{\alpha}^*$ a the better estimator for α than $\hat{\alpha}^{\text{ML}}$!

Observation 2

The ML estimator $\hat{\alpha}^{\text{ML}}$ typically performs better than the bias corrected ML estimator $\hat{\alpha}^*$.

At first sight, this is surprising as $\hat{\alpha}^*$ a the better estimator for α than $\hat{\alpha}^{\text{ML}}$!

Explanation:

Let $X_\alpha \sim \text{Pareto}(t, \alpha)$ with $t \leq D$. Then the function

$$\alpha \mapsto \mu(\alpha) := E(\min(C, (X_\alpha - D)^+))$$

is convex!

Observation 2

The ML estimator $\hat{\alpha}^{\text{ML}}$ typically performs better than the bias corrected ML estimator $\hat{\alpha}^*$.

At first sight, this is surprising as $\hat{\alpha}^*$ a the better estimator for α than $\hat{\alpha}^{\text{ML}}$!

Explanation:

Let $X_\alpha \sim \text{Pareto}(t, \alpha)$ with $t \leq D$. Then the function

$$\alpha \mapsto \mu(\alpha) := E(\min(C, (X_\alpha - D)^+))$$

is convex! Jensen's inequality implies

$$E(\mu(\hat{\alpha}^*)) > \mu(E(\hat{\alpha}^*)).$$

Observation 2

The ML estimator $\hat{\alpha}^{\text{ML}}$ typically performs better than the bias corrected ML estimator $\hat{\alpha}^*$.

At first sight, this is surprising as $\hat{\alpha}^*$ a the better estimator for α than $\hat{\alpha}^{\text{ML}}$!

Explanation:

Let $X_\alpha \sim \text{Pareto}(t, \alpha)$ with $t \leq D$. Then the function

$$\alpha \mapsto \mu(\alpha) := E(\min(C, (X_\alpha - D)^+))$$

is convex! Jensen's inequality implies

$$E(\mu(\hat{\alpha}^*)) > \mu(E(\hat{\alpha}^*)).$$

$\Rightarrow$ Estimation of the expected layer loss with $\hat{\alpha}^*$ has a positive bias!

Observation 2

The ML estimator $\hat{\alpha}^{\text{ML}}$ typically performs better than the bias corrected ML estimator $\hat{\alpha}^*$.

At first sight, this is surprising as $\hat{\alpha}^*$ a the better estimator for α than $\hat{\alpha}^{\text{ML}}$!

Explanation:

Let $X_\alpha \sim \text{Pareto}(t, \alpha)$ with $t \leq D$. Then the function

$$\alpha \mapsto \mu(\alpha) := E(\min(C, (X_\alpha - D)^+))$$

is convex! Jensen's inequality implies

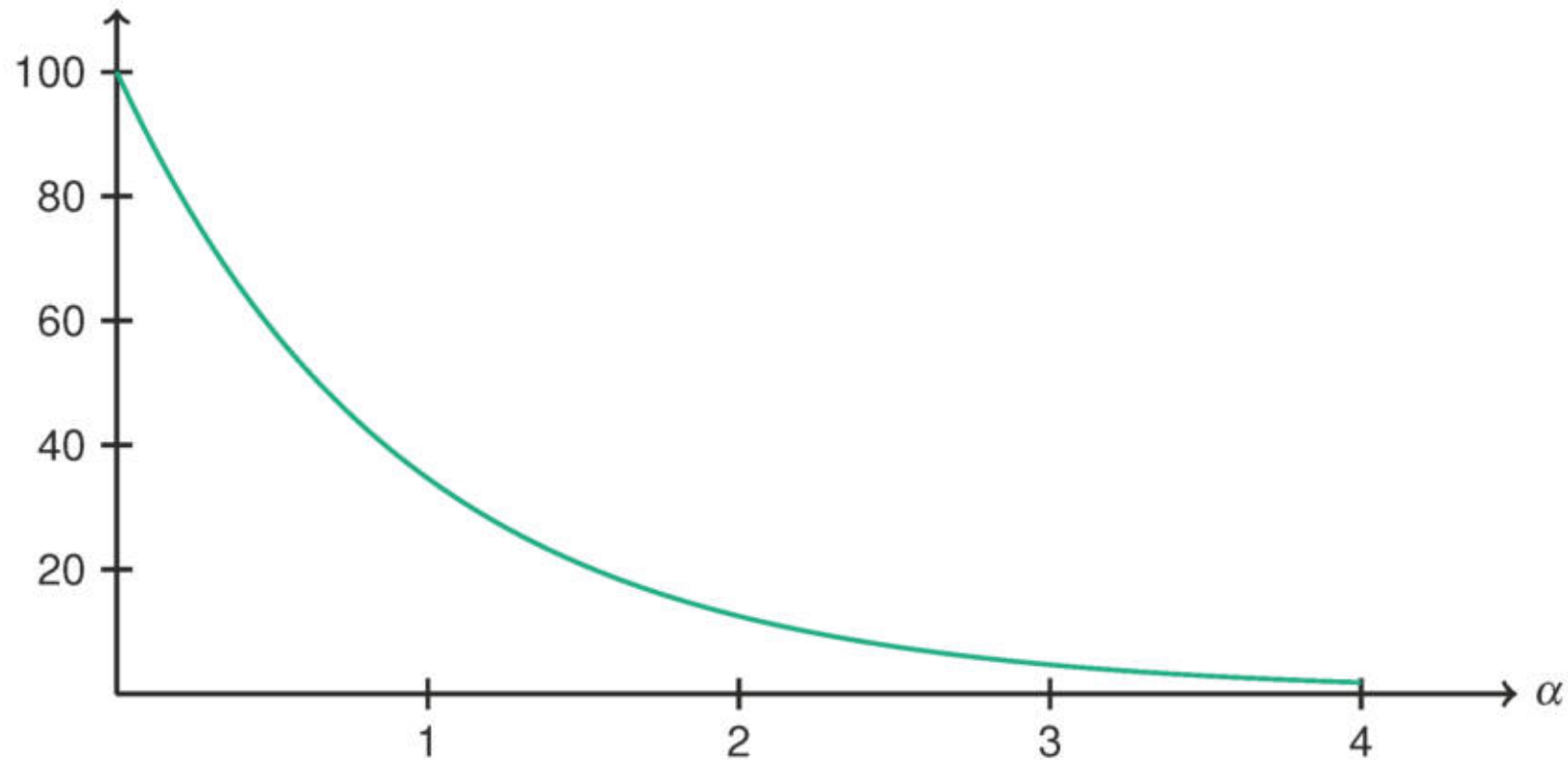
$$E(\mu(\hat{\alpha}^*)) > \mu(E(\hat{\alpha}^*)).$$

$\Rightarrow$ Estimation of the expected layer loss with $\hat{\alpha}^*$ has a positive bias!

The ML estimator $\hat{\alpha}^{\text{ML}}$ results in lower estimates for the expected layer loss, which compensates this positive bias to some extent!

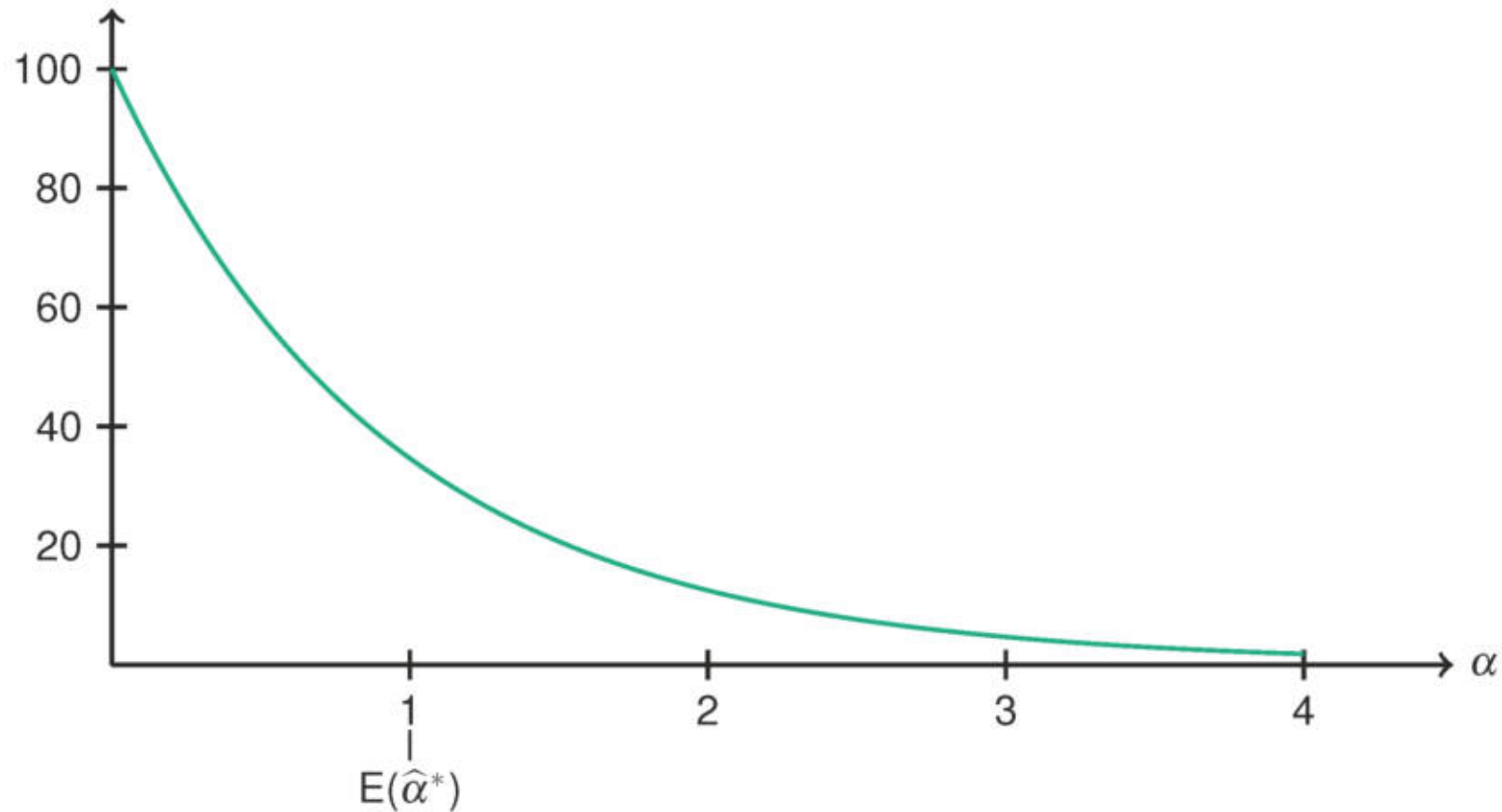
Jensen's Inequality

$$\mu(\alpha) = E(\min(C, (X_\alpha - D)^+))$$

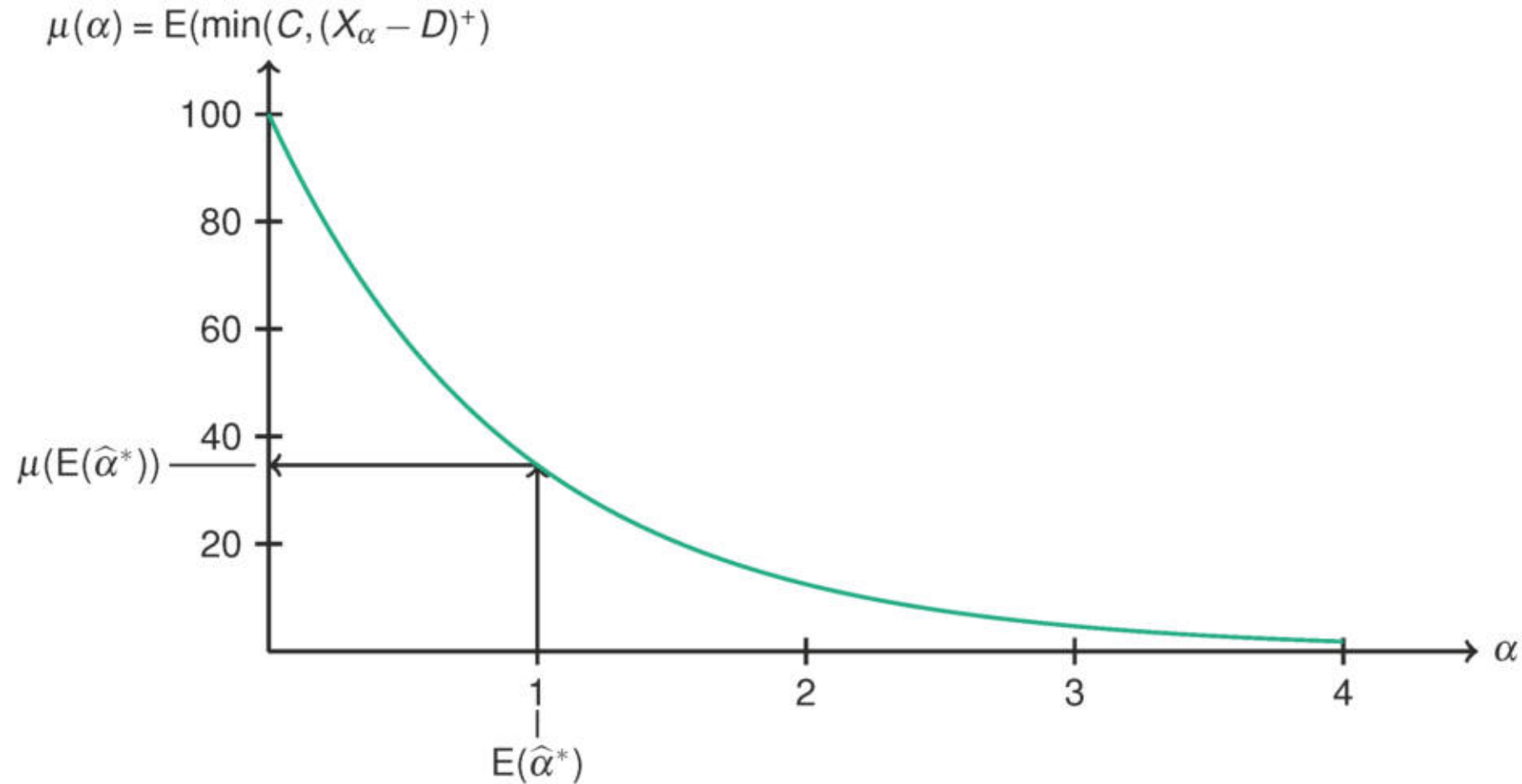


Jensen's Inequality

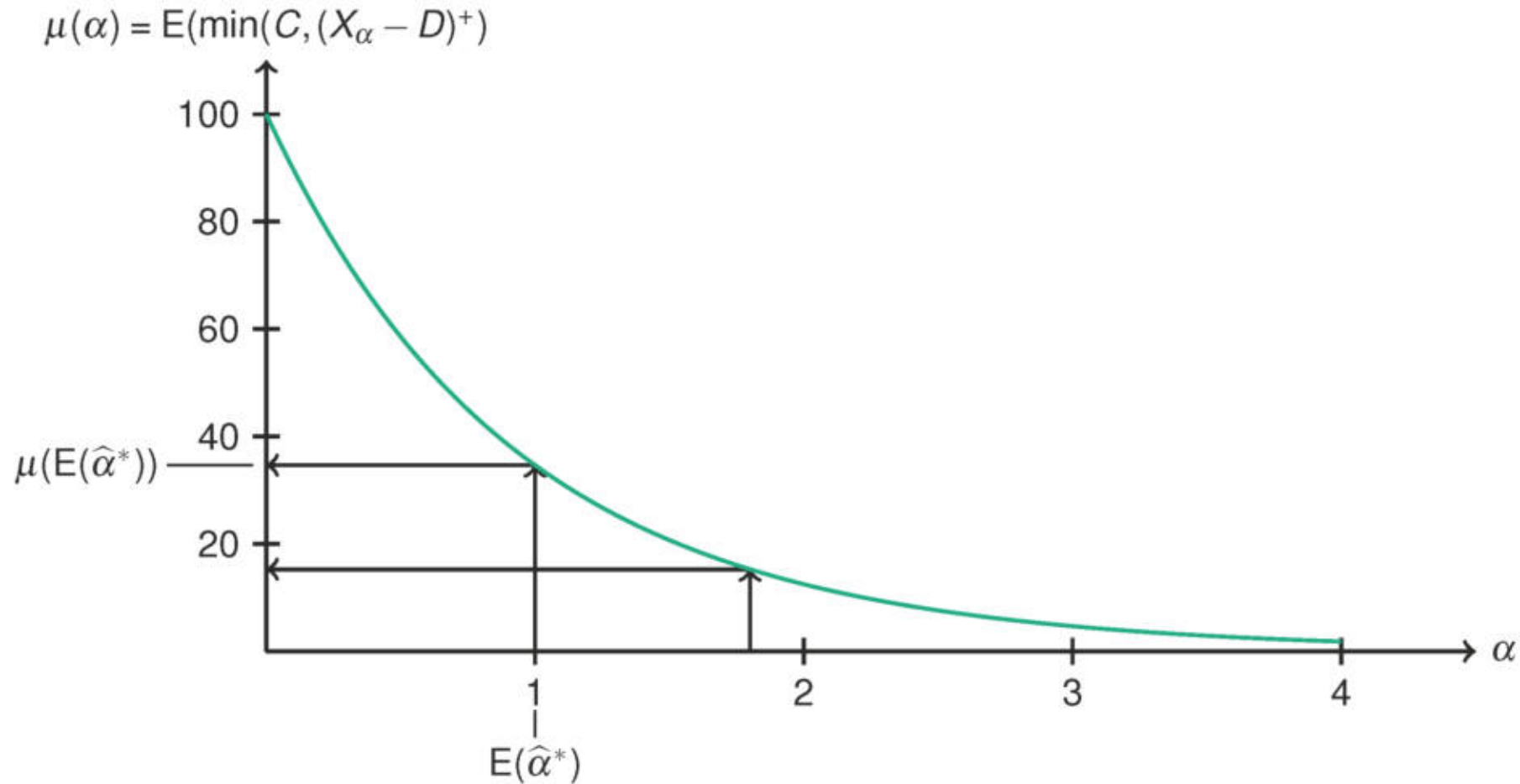
$$\mu(\alpha) = E(\min(C, (X_\alpha - D)^+))$$



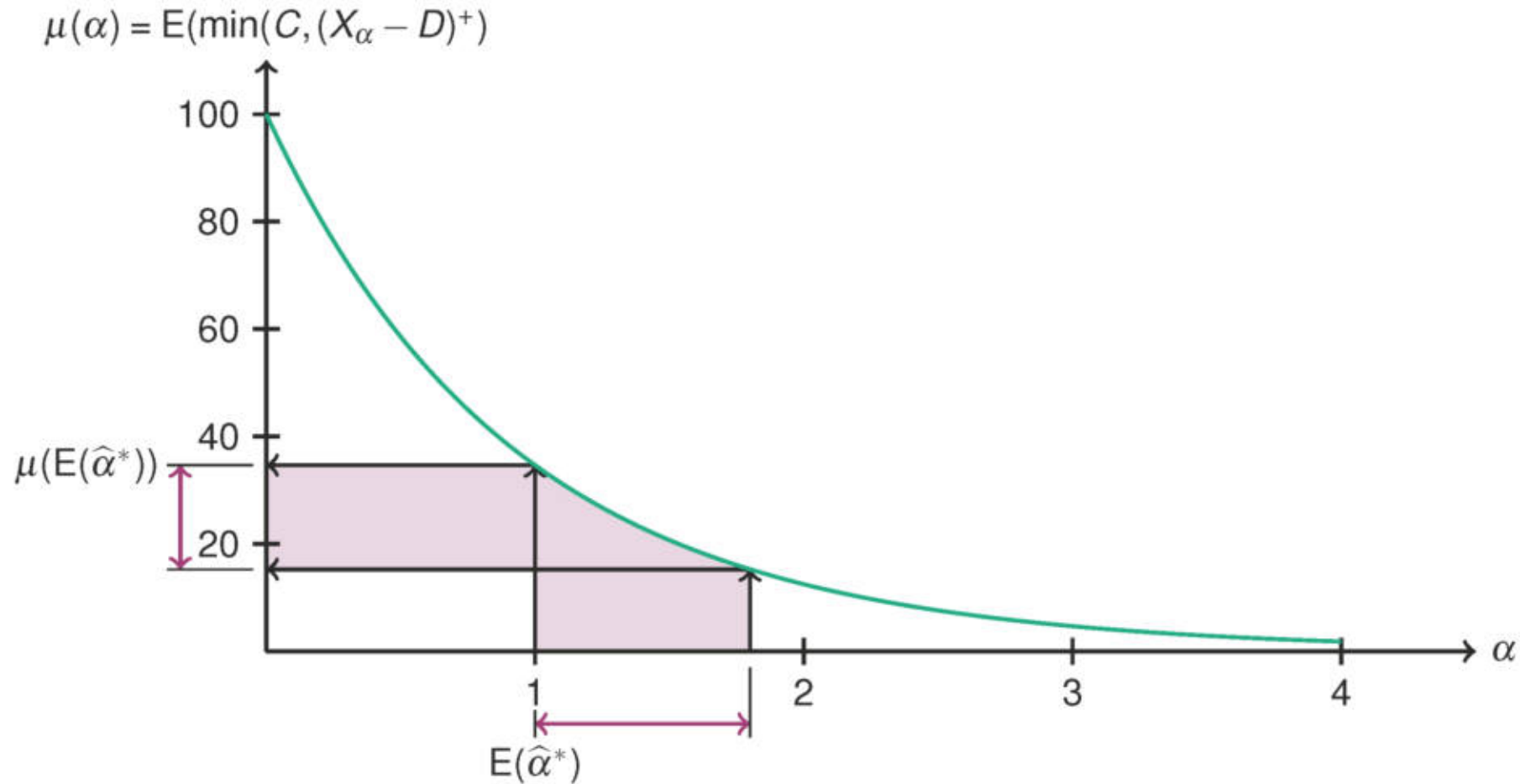
Jensen's Inequality



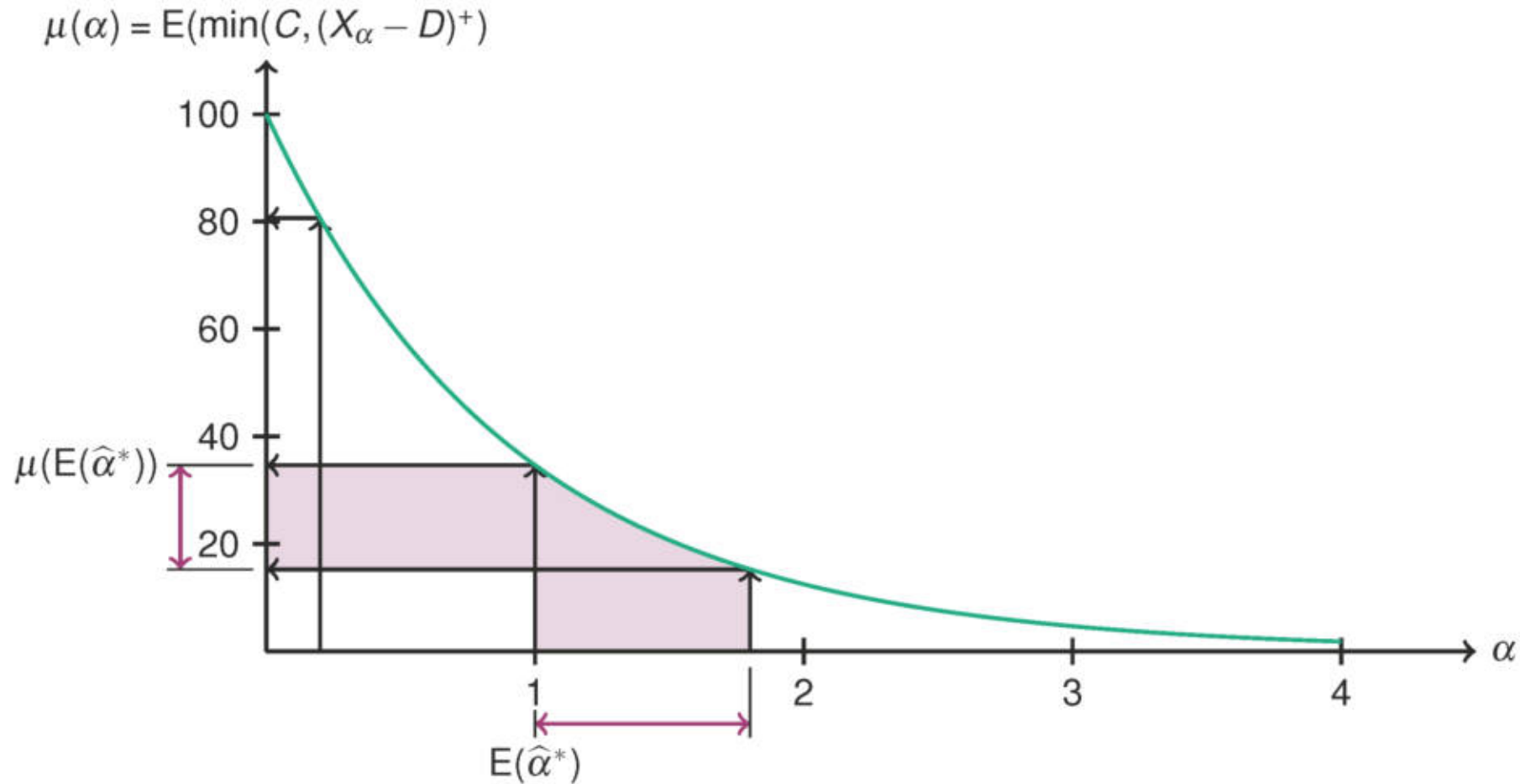
Jensen's Inequality



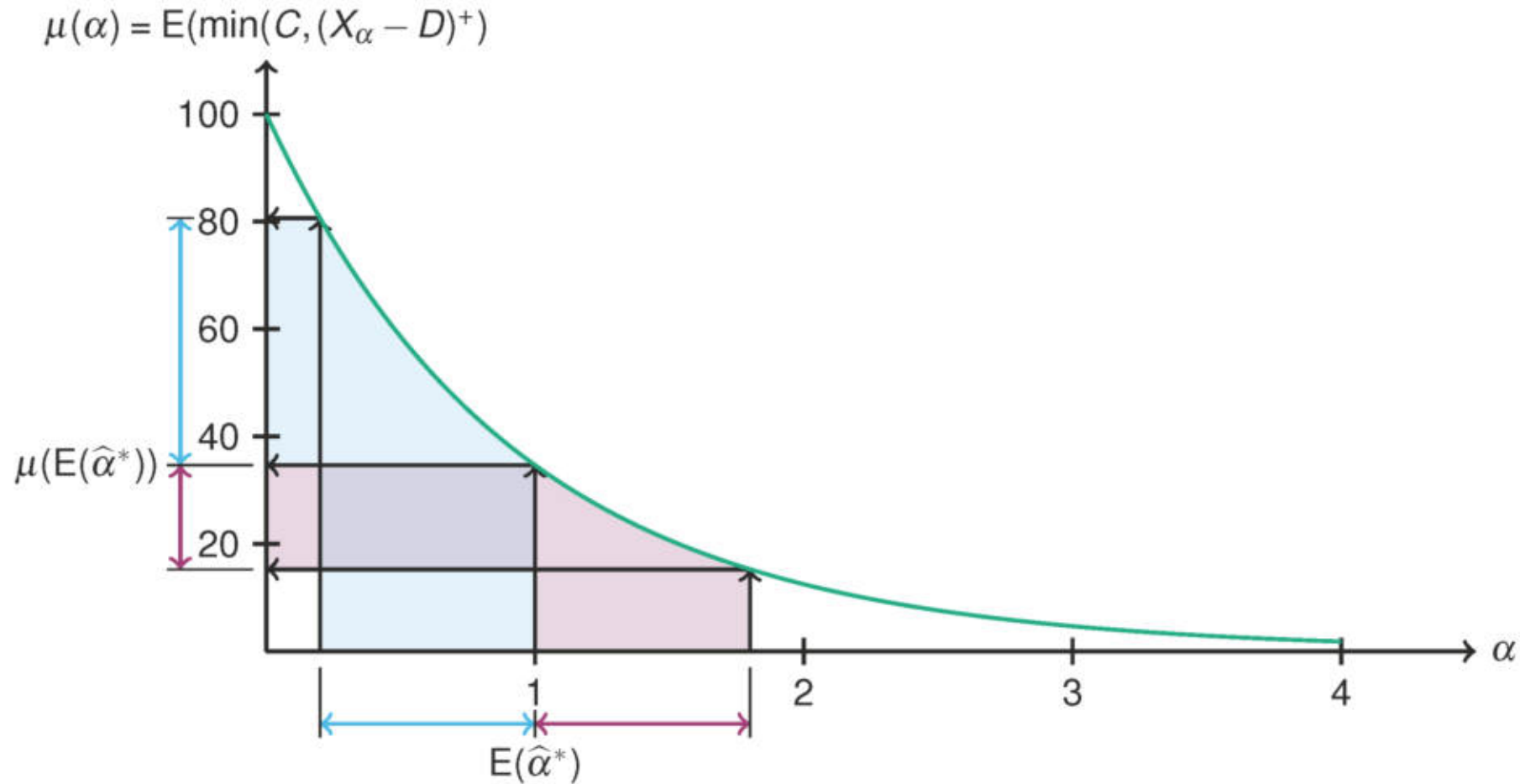
Jensen's Inequality



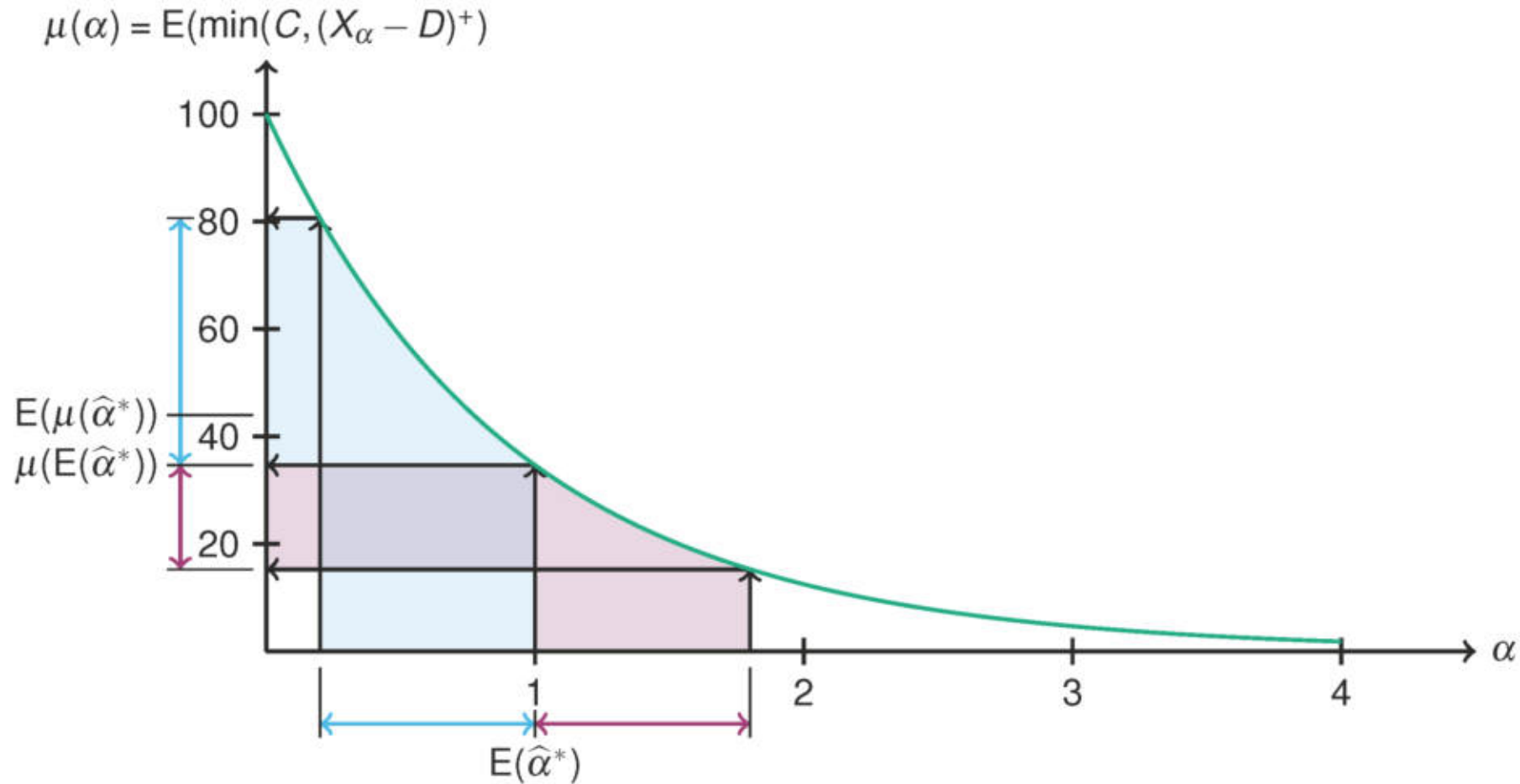
Jensen's Inequality



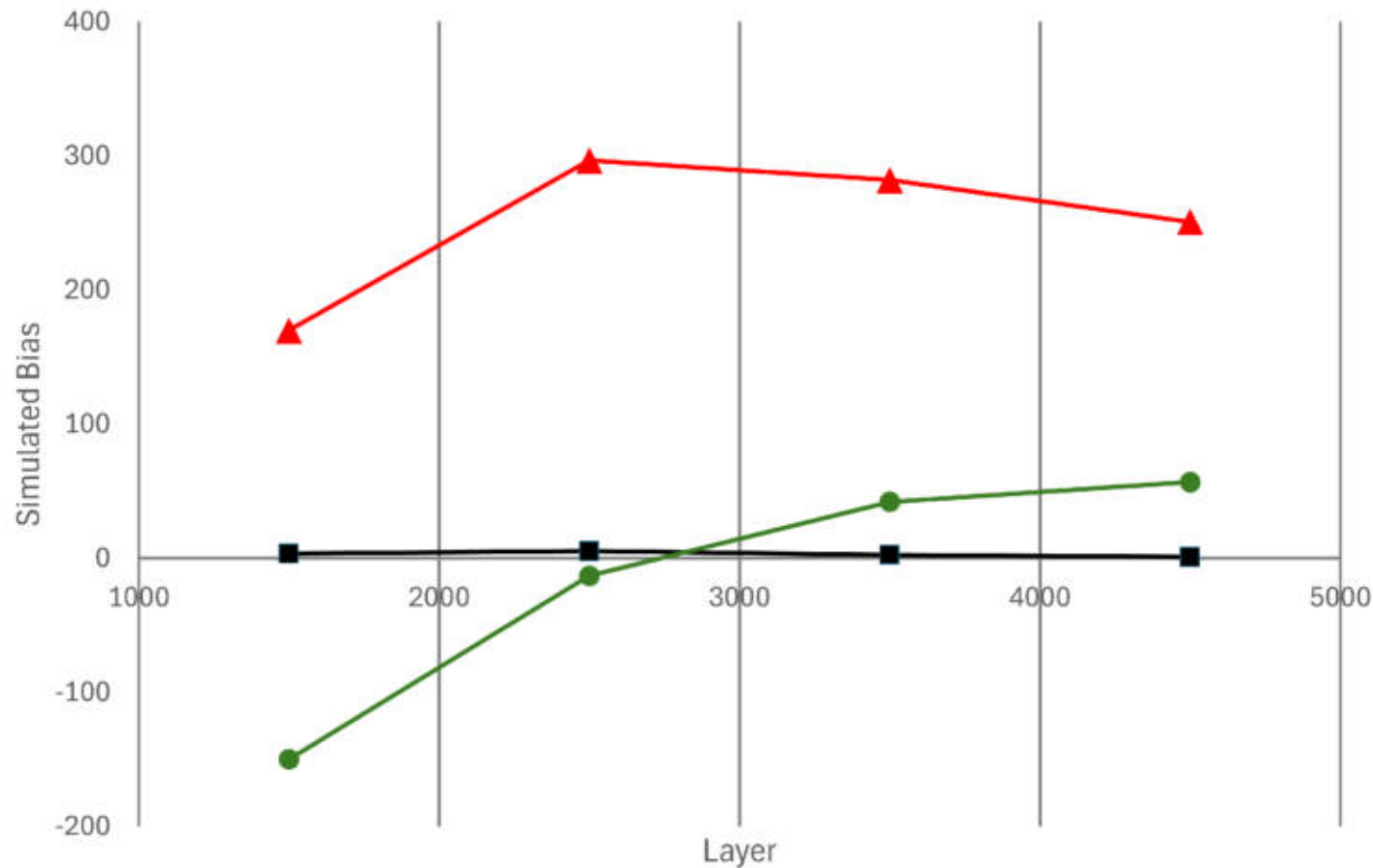
Jensen's Inequality



Jensen's Inequality



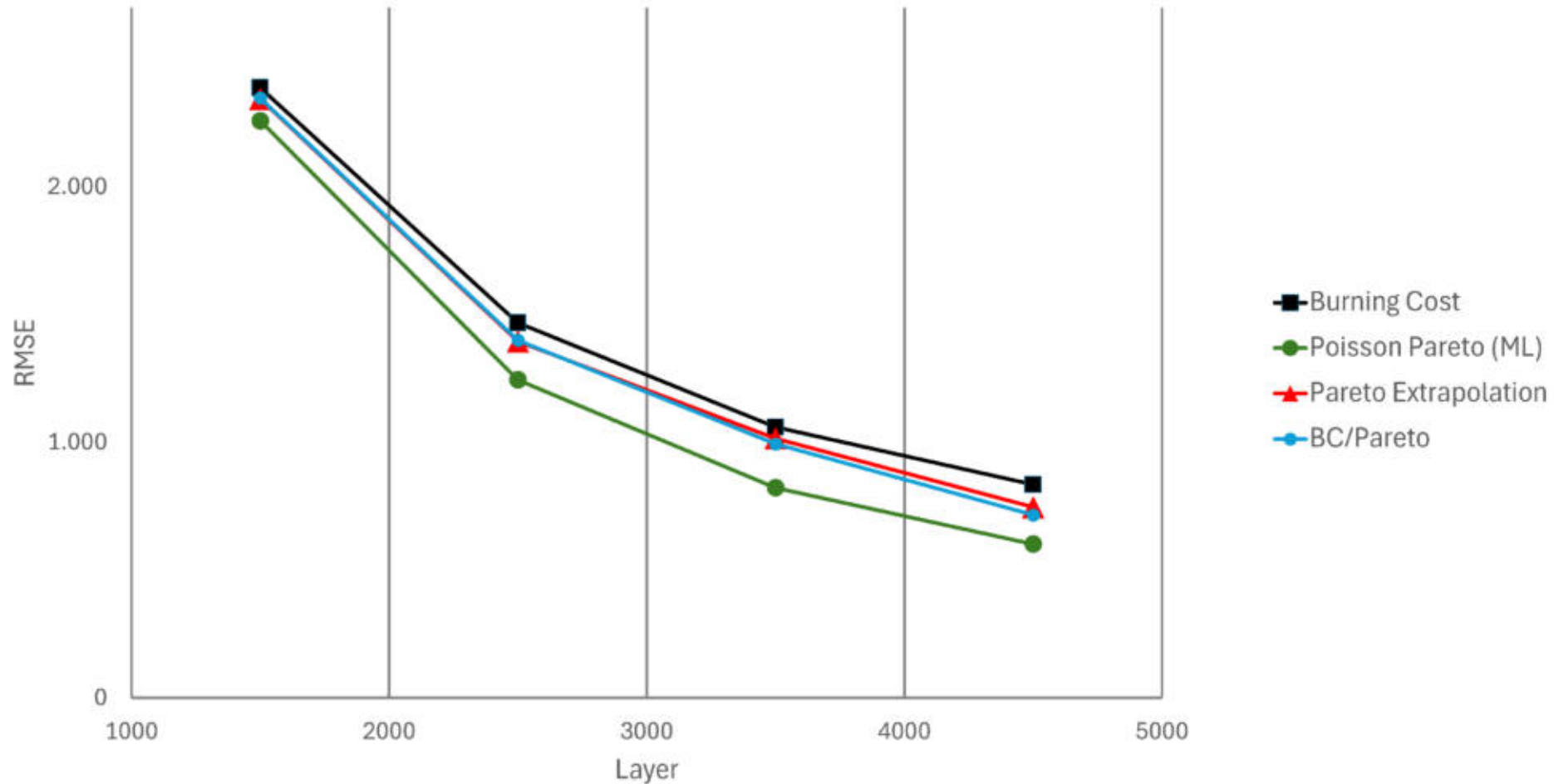
Bias in the case $\lambda = 5$, $\alpha = 2.0$



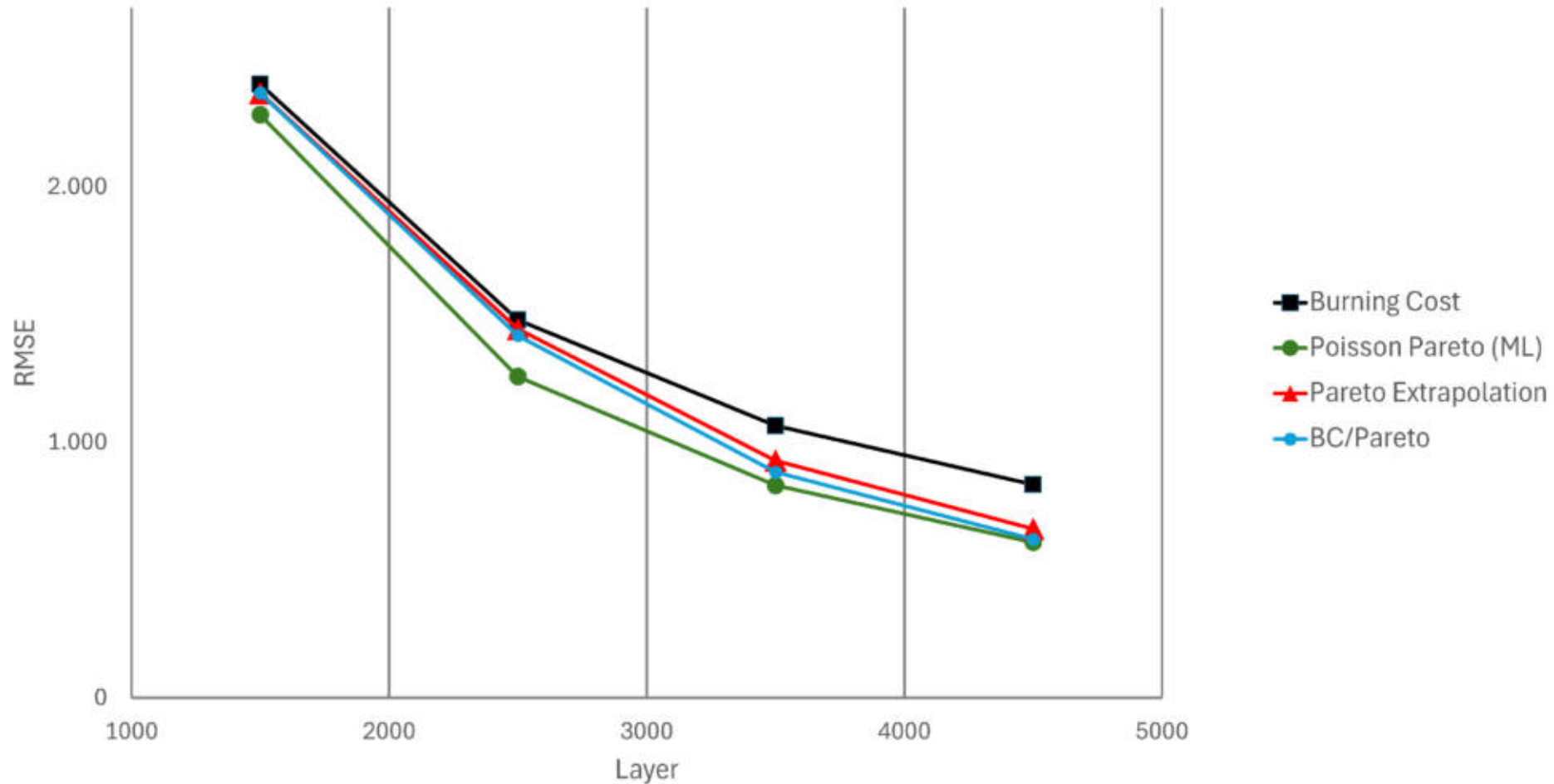
- BC has no bias
- Negative bias with $\hat{\alpha}^{\text{ML}}$ in first layer
- Substantial positive bias with $\hat{\alpha}^*$ for all layers

■ Burning Cost
● Poisson Pareto (ML)
▲ Poisson Pareto (Bias Corrected)

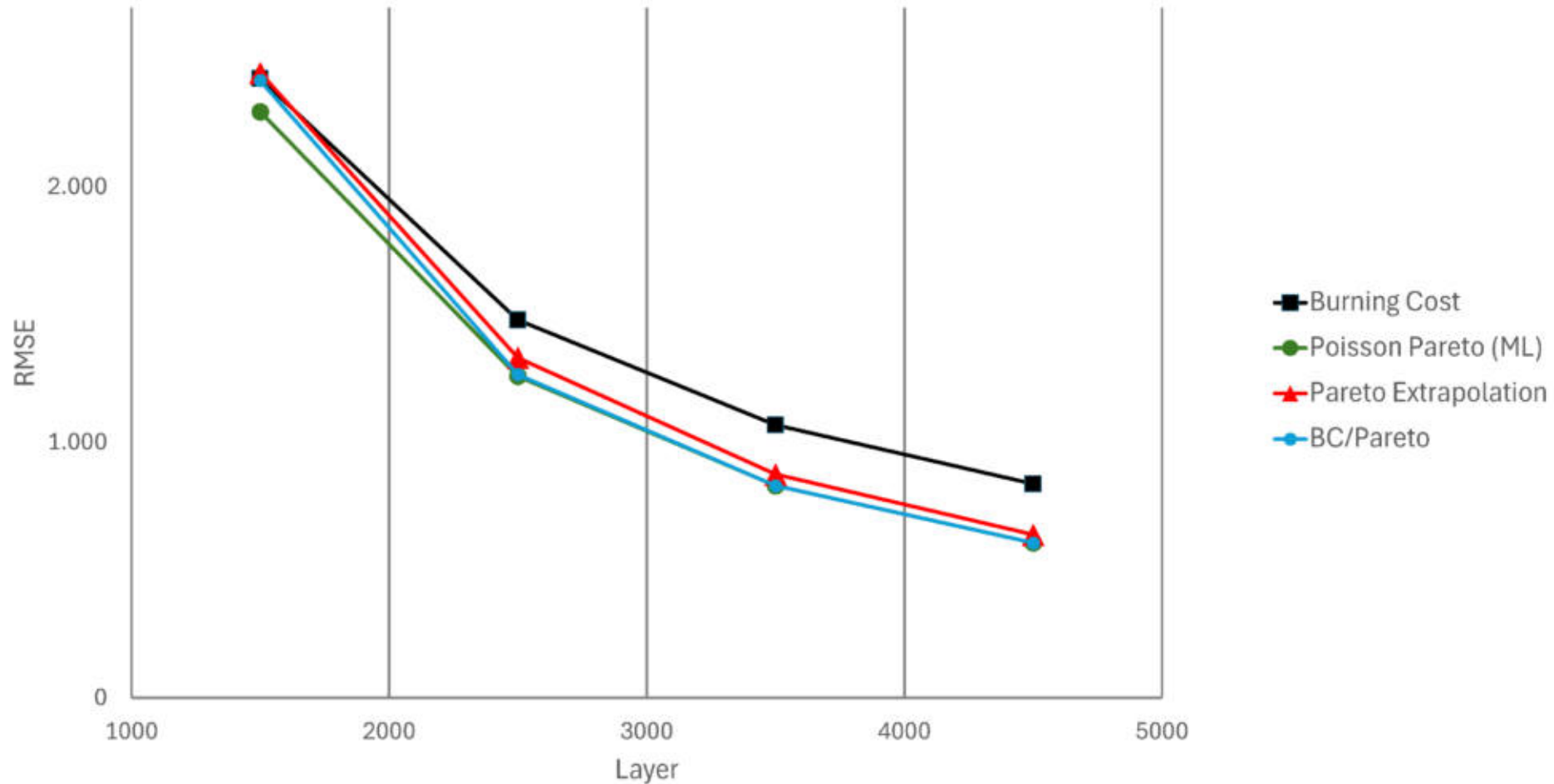
Burning Cost up to 3rd largest loss ($\lambda = 15$, $\alpha = 2.0$)



Burning Cost up to 5th largest loss ($\lambda = 15$, $\alpha = 2.0$)



Burning Cost up to 10th largest loss ($\lambda = 15$, $\alpha = 2.0$)



Observation 3

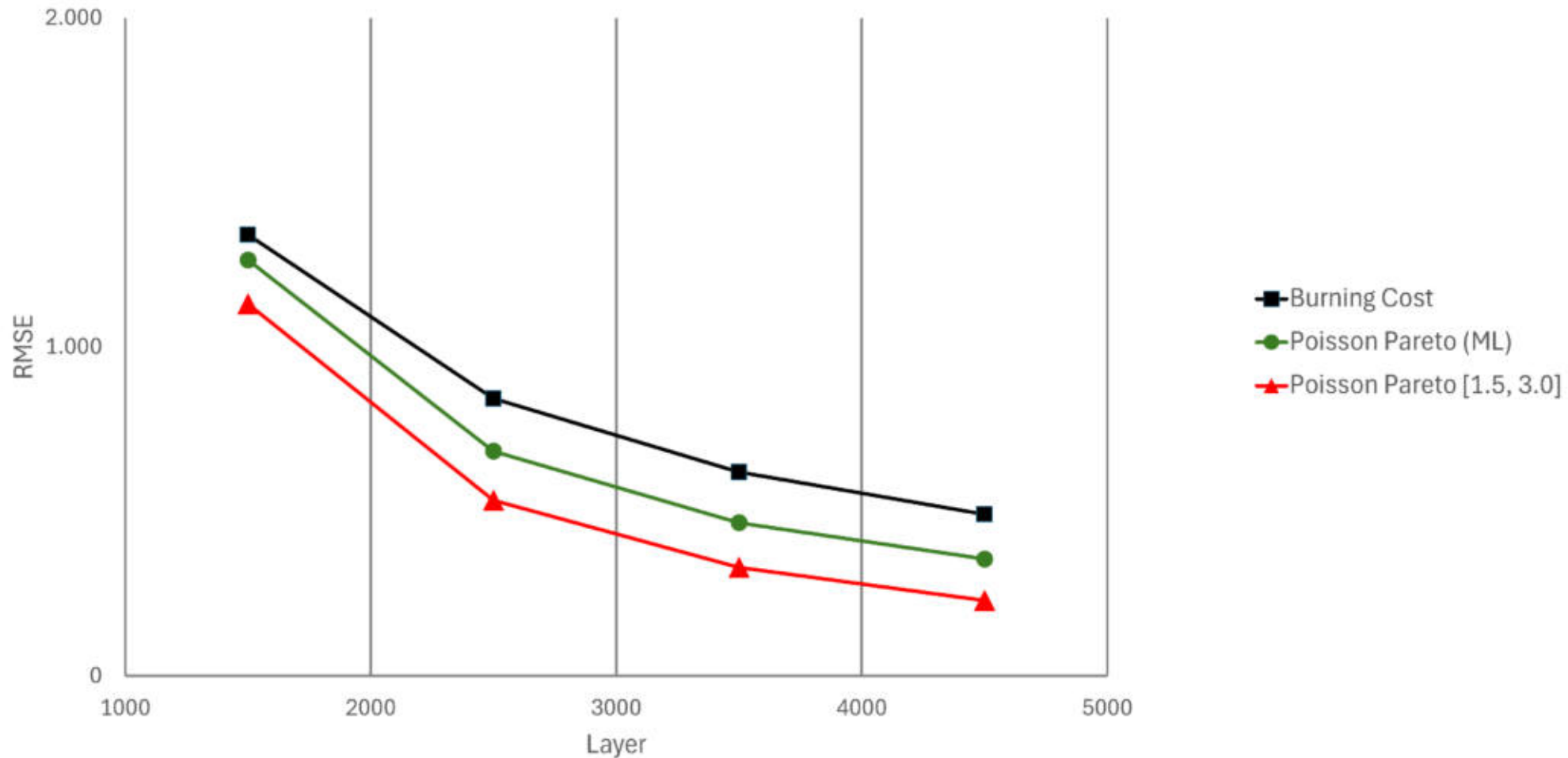
Methods that use the Burning Cost up to the k -th largest loss seem to have RSMEs between the Burning Cost and the Poisson/Pareto model.

Observation 3

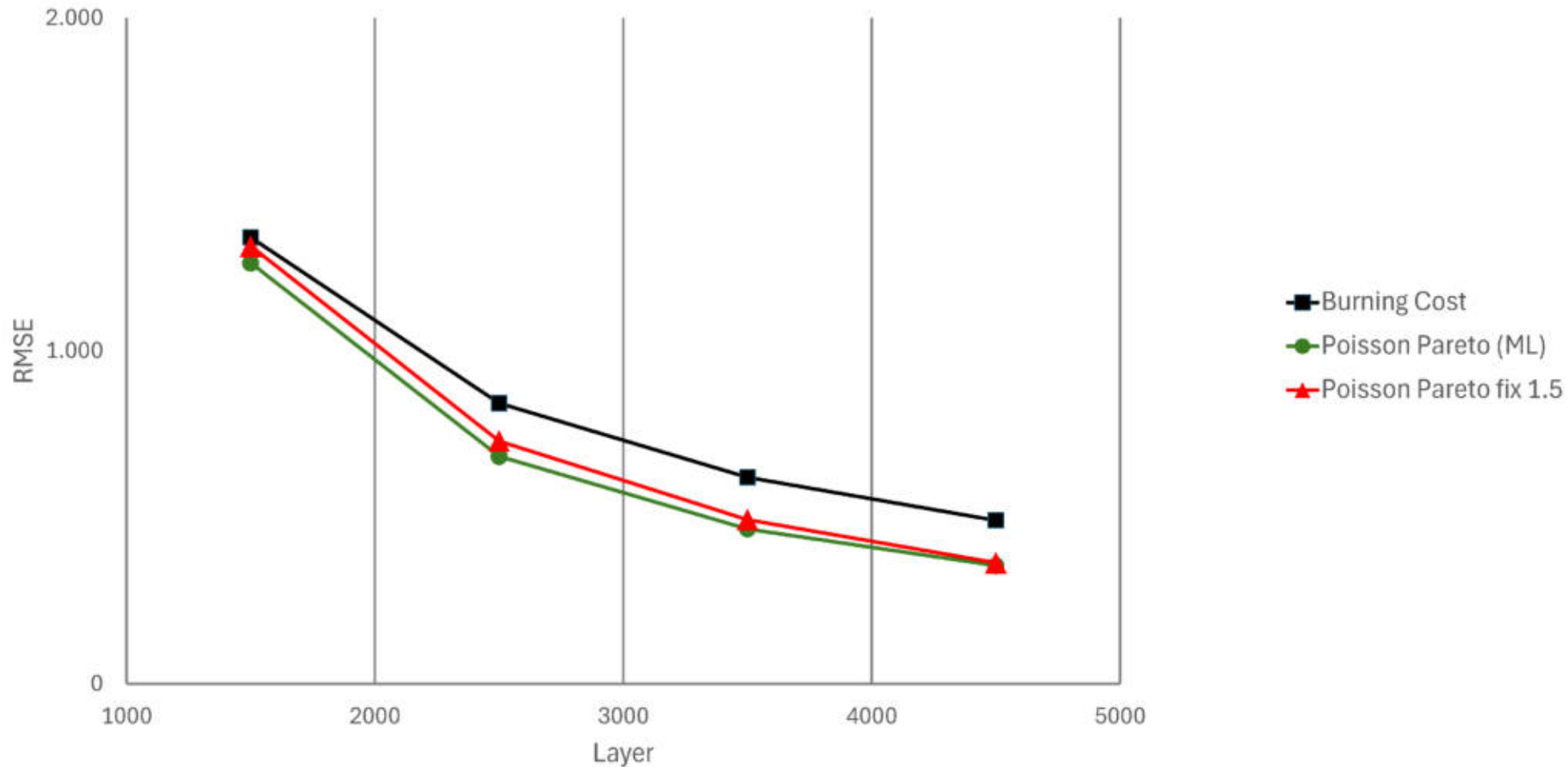
Methods that use the Burning Cost up to the k -th largest loss seem to have RSMEs between the Burning Cost and the Poisson/Pareto model.

Using the 3rd largest loss is closer to the Burning Cost model, using the 10th largest loss is close to the Poisson/Pareto model.

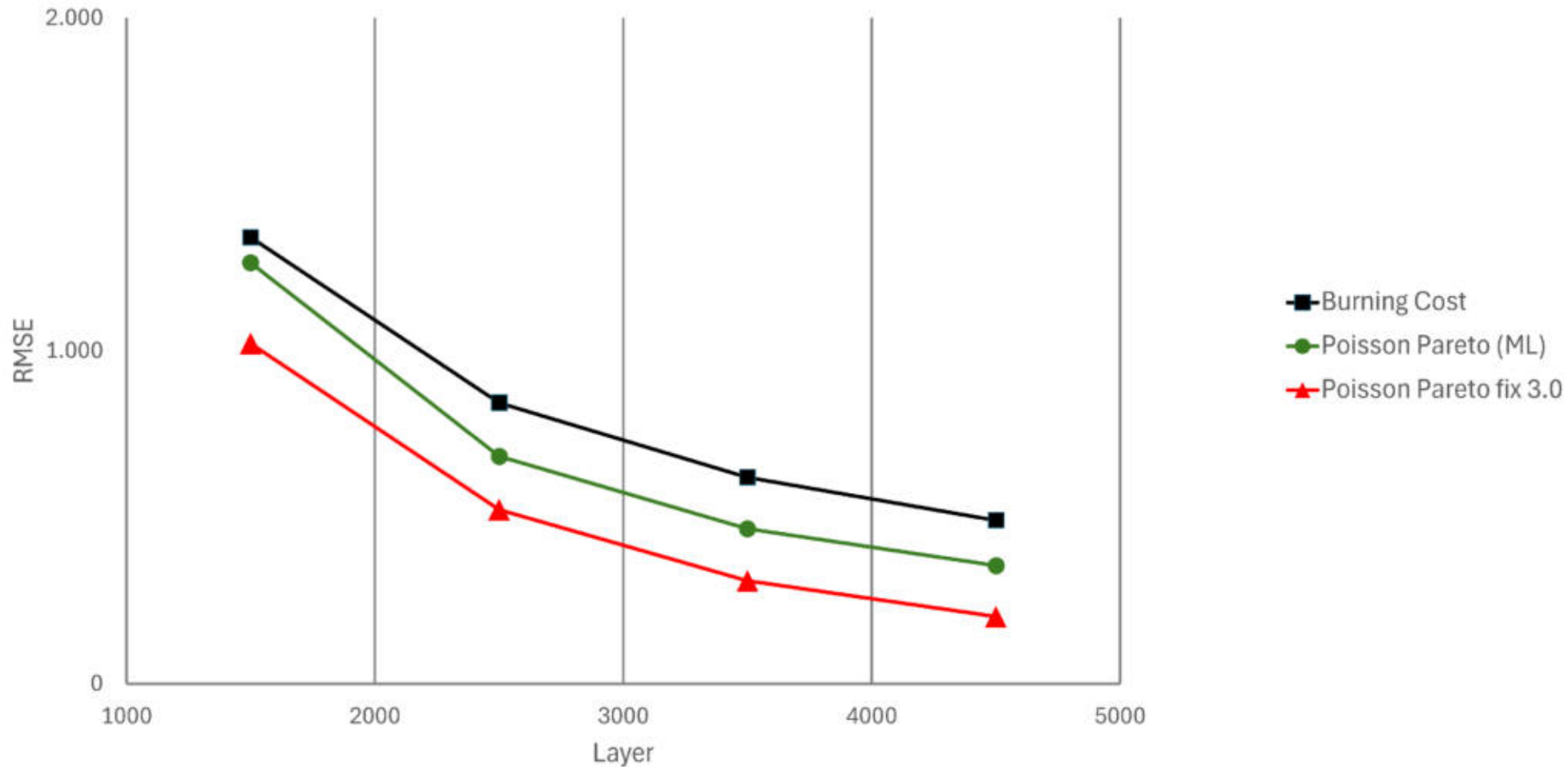
Use boundaries $\alpha_{\min} = 1.5$ and $\alpha_{\max} = 3.0$ for $\hat{\alpha}^{\text{ML}}$ ($\lambda = 5$, $\alpha = 2.0$)



Use fixed *market alpha* $\alpha = 1.5$ ($\lambda = 5$, $\alpha = 2.0$)



Use fixed *market alpha* $\alpha = 3.0$ ($\lambda = 5$, $\alpha = 2.0$)



Observation 4

In case of a small λ the RMSE can be reduced substantially by using market knowledge for the estimation of the Pareto alpha:

- Upper/lower bounds for the estimator
- Market alpha
- Also a credibility approach (Bühlmann Straub) could be employed.

Observation 4

In case of a small λ the RMSE can be reduced substantially by using market knowledge for the estimation of the Pareto alpha:

- Upper/lower bounds for the estimator
- Market alpha
- Also a credibility approach (Bühlmann Straub) could be employed.

Note that the RSMEs are smaller, even if the portfolio's alpha deviates substantially from the market mean.

Observation 4

In case of a small λ the RMSE can be reduced substantially by using market knowledge for the estimation of the Pareto alpha:

- Upper/lower bounds for the estimator
- Market alpha
- Also a credibility approach (Bühlmann Straub) could be employed.

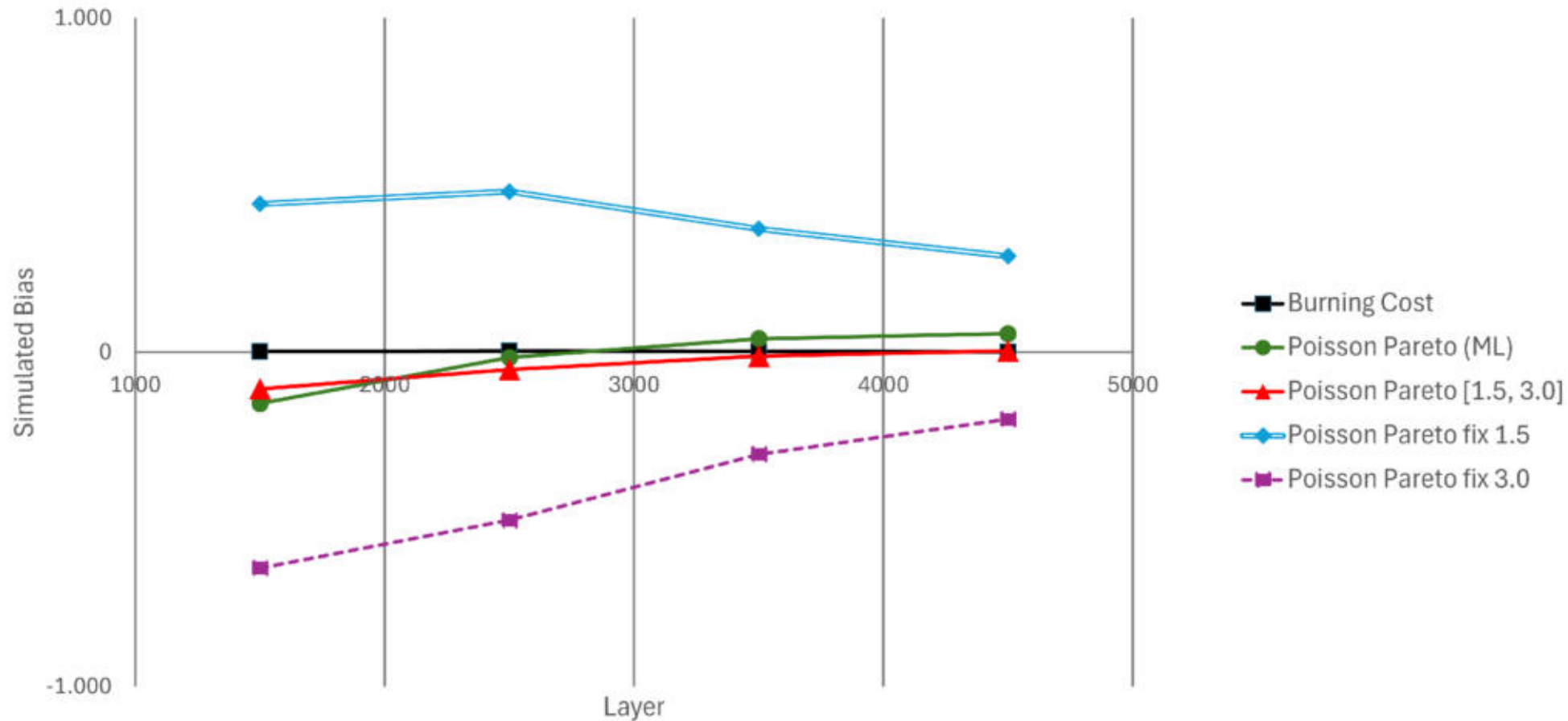
Note that the RSMEs are smaller, even if the portfolio's alpha deviates substantially from the market mean.

WARNING

Do not look exclusively at the RMSE!
Also care about the bias!

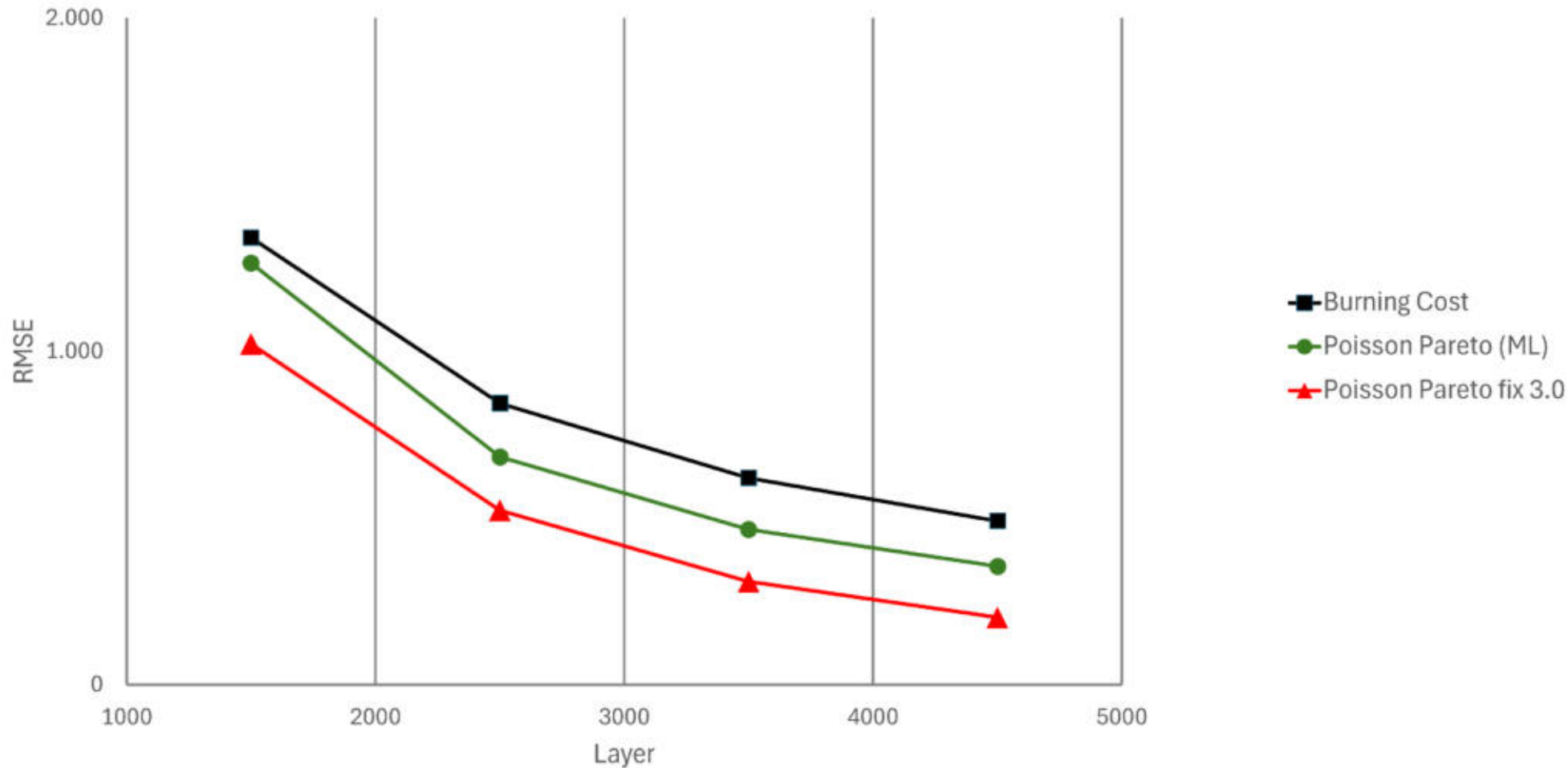
Bias with fixed/restricted alphas ($\lambda = 5$, $\alpha = 2.0$)

WARNING



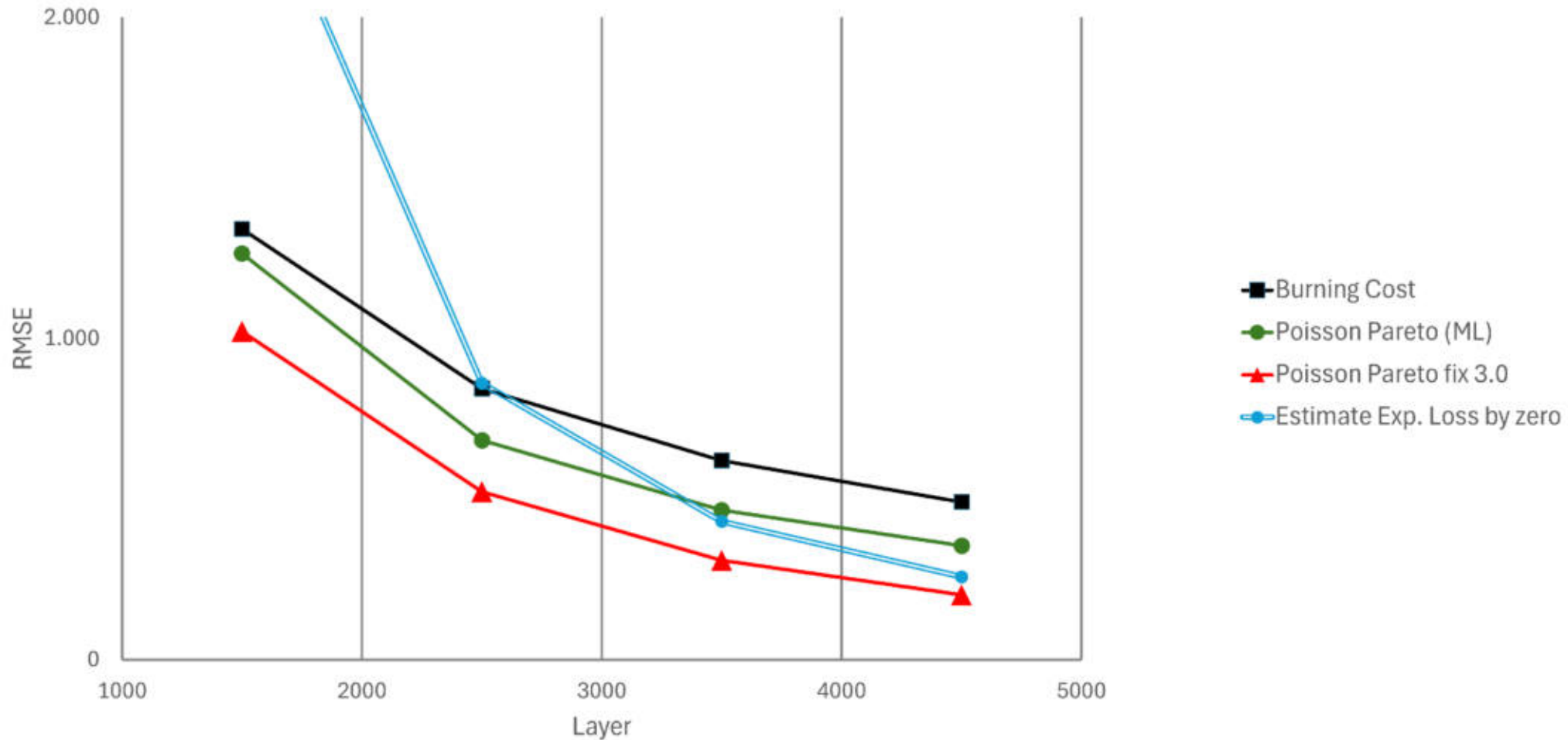
Use fixed *market alpha* $\alpha = 3.0$ ($\lambda = 5$, $\alpha = 2.0$)

WARNING



Use fixed *market alpha* $\alpha = 3.0$ ($\lambda = 5$, $\alpha = 2.0$)

WARNING



Agenda

1. Burning cost for per risk XLs
2. Pareto based pricing approaches
3. Simulation model for evaluating pricing approaches
4. Observations with underlying Pareto severity
5. Observations with other underlying severity distributions

Observation 5

Case 1: Underlying severity similar to Pareto

- If we just look at the RMSE, Poisson-Pareto still performed very good.

Observation 5

Case 1: Underlying severity similar to Pareto

- If we just look at the RMSE, Poisson-Pareto still performed very good.
- However, we observed a substantial bias in the first layer.

Observation 5

Case 1: Underlying severity similar to Pareto

- If we just look at the RMSE, Poisson-Pareto still performed very good.
- However, we observed a substantial bias in the first layer.
- Thus, methods using BC up to the k -th largest loss should be preferred.

Observation 5

Case 1: Underlying severity similar to Pareto

- If we just look at the RMSE, Poisson-Pareto still performed very good.
- However, we observed a substantial bias in the first layer.
- Thus, methods using BC up to the k -th largest loss should be preferred.

Case 2: Underlying severity deviates massively from a Pareto distribution

- Pareto based pricing methods can have a substantial bias.

Observation 5

Case 1: Underlying severity similar to Pareto

- If we just look at the RMSE, Poisson-Pareto still performed very good.
- However, we observed a substantial bias in the first layer.
- Thus, methods using BC up to the k -th largest loss should be preferred.

Case 2: Underlying severity deviates massively from a Pareto distribution

- Pareto based pricing methods can have a substantial bias.
- The observations with underlying Pareto severity do not apply in this case.

**Vielen Dank für
Ihre Aufmerksamkeit.**

Fatima ezzahra Kherraz, Dr. Ulrich Riegel
fkherraz@glise.com, uriegel@munichre.com
+49 89 38912074, +49 151 51536526