



ΕΘΝΙΚΟ ΚΑΙ ΚΑΠΟΔΙΣΤΡΙΑΚΟ ΠΑΝΕΠΙΣΤΗΜΙΟ ΑΘΗΝΩΝ

ΣΧΟΛΗ ΘΕΤΙΚΩΝ ΕΠΙΣΤΗΜΩΝ ΤΜΗΜΑ ΜΑΘΗΜΑΤΙΚΩΝ

ΠΡΟΓΡΑΜΜΑ ΜΕΤΑΠΤΥΧΙΑΚΩΝ ΣΠΟΥΔΩΝ

ΣΤΑΤΙΣΤΙΚΗ ΚΑΙ ΕΠΙΧΕΙΡΗΣΙΑΚΗ ΕΡΕΥΝΑ

ΔΙΠΛΩΜΑΤΙΚΗ ΕΡΓΑΣΙΑ:

**«Στοχαστικά Μοντέλα διαχείρισης αποθεμάτων με ελλιπή
πληροφόρηση»**

Κωνσταντίνος Χαιρετάκης

Επιβλέπων καθηγητής

Απόστολος Μπουρνέτας

ΑΘΗΝΑ

ΣΕΠΤΕΜΒΡΙΟΣ 2017

Η παρούσα Διπλωματική Εργασία εκπονήθηκε στα πλαίσια των σπουδών για την
απόκτηση του Μεταπτυχιακού Διπλώματος Ειδίκευσης στη Στατιστική και την
Επιχειρησιακή Έρευνα που απονέμει το Τμήμα Μαθηματικών του Εθνικού και
Καποδιστριακού Πανεπιστημίου Αθηνών.

Εγκρίθηκε στις από Εξεταστική Επιτροπή αποτελούμενη από τους:

Ονοματεπώνυμο

Βαθμίδα

Υπογραφή

Μηλολιδάκης Κωνσταντίνος

Αναπληρωτής Καθηγητής

.....

Μπουρνέτας Απόστολος (Επιβλέπων)

Καθηγητής

.....

Οικονόμου Αντώνης

Καθηγητής

.....

Ευχαριστίες

Θα ήθελα να ευχαριστήσω τον επιβλέποντα καθηγητή της διπλωματικής μου εργασίας κ. Απόστολο Μπουρνέτα, για την καθοδήγηση και την πολύτιμη συμβολή του στην δημιουργία της εργασίας αλλά και σε ολόκληρη την διάρκεια των σπουδών μου. Επίσης, τα μέλη της τριμελούς επιτροπής κ. Αντώνη Οικονόμου και κ. Κωνσταντίνο Μηλολιδάκη που με τίμησαν με την συμμετοχή τους στην εξεταστική επιτροπή.

ΠΕΡΙΕΧΟΜΕΝΑ:

ΚΕΦΑΛΑΙΟ 1^ο Εισαγωγή και Βασικές έννοιες

1.1 Εισαγωγή και βασικές έννοιες.....σελ 11	σελ 11
1.2 Παραλλαγές προβλήματος προσεγγίσεις.....σελ 15	σελ 15
1.3 Μερική επισκόπηση βιβλιογραφίας.....σελ 16	σελ 16

ΚΕΦΑΛΑΙΟ 2^ο

Παρουσίαση βασικών εργασιών

2.1 Εργασία 1 «Demand estimation in lost sales» Nahmias (1994).....σελ 19	σελ 19
2.1.1 Εκτιμητές Μέγιστης Πιθανοφάνειας.....σελ 19	σελ 19
2.1.2 Νέοι εκτιμητές του συγγραφέα (New estimators).....σελ 20	σελ 20
2.1.3 Βέλτιστοι Γραμμικοί Εκτιμητές (Blues).....σελ 23	σελ 23
2.1.4 Συμπεράσματα προσομοιώσεων.....σελ 24	σελ 24
2.1.5 Ανανέωση εκτίμησης για τη ζήτηση και το δείγμα.....σελ 25	σελ 25
2.2 Εργασία 2 ^η «A practical inventory control policy using operational statistics» Liyanage, Shanthikumar.....σελ 26	σελ 26

2.2.1 Με χρήση δειγματικού μέσου.....σελ 27	
2.2.2 Με χρήση δειγματικού μέσου & εκτιμήσεων βέλτ. πολιτικής.....σελ 28	
2.2.3 Με χρήση της εμπειρικής συνάρτησης κατανομής.....σελ 29	
2.2.4 Με χρήση της εμπειρικής συνάρτησης κατανομής και εκτιμήσεων βέλτιστης πολιτικής.....σελ 30	

ΚΕΦΑΛΑΙΟ 3^ο

Υποθέσεις και εκτίμηση της παραμέτρου

3.1 Υποθέσεις-Προσέγγιση στο συγκεκριμένο υπόδειγμασελ 34	
3.2 Εκτιμητής μέγιστης πιθανοφάνειας με βάση το δείγμα(Ε.Μ.Π.).....σελ 38	
3.3 Εκτιμητής μέγιστης πιθανοφάνειας με βάση το πλήθος των παρατηρούμενων τιμών του δείγματος.....σελ 43	

ΚΕΦΑΛΑΙΟ 4^ο

Μοντέλο Επιβίωσης και εκτιμητές παραμέτρου

4.1 Το πρόβλημα περικομμένων δεδομένων με γνωστό αριθμό παρατηρούμενων τιμών.....σελ 45	
---	--

4.1.1 Εκτιμήτρια μέγιστης πιθανοφάνειας με χρήση της μέγιστης παρατήρησης του δείγματος και του αποθέματος.....σελ 47	σελ 47
4.2 Προσομοιωτικός Εκτιμητής.....σελ 50	σελ 50
4.2.1 Αμεροληψία Νέου Προσομοιωτικού Εκτιμητή.....σελ 53	σελ 53
4.2.2 Διασπορά Νέου Προσομοιωτικού Εκτιμητή.....σελ 54	σελ 54
4.3 Παρουσίαση των Γραμμικών εκτιμητώνσελ 56	σελ 56
4.3.1 Σημαντικά Θεωρήματα.....σελ 56	σελ 56
4.3.2 Μειονεκτήματα Γραμμικών εκτιμητών-Παραδείγματ-Πίνακας.....σελ 58	σελ 58
4.4 Βέλτιστη ποσότητα παραγγελίας – Προσομοιώσεις.....σελ 64	σελ 64
4.4.1 Αλγόριθμος προσομοίωσης για εκτιμήσεις του $1/\theta$σελ 64	σελ 64
4.4.2 Εκτίμηση Βέλτιστου Κέρδους με χρήση του εκτιμητή μέγιστης πιθανοφάνειας.....σελ 65	σελ 65

ΚΕΦΑΛΑΙΟ 5^ο

Συμπεράσματα προσομοιώσεων-Παρατηρήσεις

5.1 Αποτελέσματα Προσομοιώσεων-Πίνακες.....σελ 65	σελ 65
5.2 Συμπεράσματα Προσομοιώσεων-Διαγράμματα.....σελ 73	σελ 73

ΠΑΡΑΡΤΗΜΑ

Κύριο μέρος του κώδικα στο στατιστικό πακέτο R.....σελ 79	σελ 79
---	--------

Βιβλιογραφία.....σελ 83	σελ 83
-------------------------	--------

ΚΕΦΑΛΑΙΟ 1

1.1 Εισαγωγή και Βασικές έννοιες

Εισαγωγή στο γενικότερο πρόβλημα του εφημεριδοπώλη

Σε αυτήν την εργασία ασχολούμαστε με ένα κλασσικό πρόβλημα του τομέα επιχειρησιακής έρευνας και οικονομικών, και ειδικότερα ένα πρόβλημα διαχείρισης αποθεμάτων. Το μοντέλο αυτό ονομάζεται πρόβλημα του εφημεριδοπώλη ή newsvendor problem όπως είναι γνωστό στη διεθνή βιβλιογραφία. Το μοντέλο αυτό σε παραλλαγές έχει μεγάλη εφαρμογή ειδικότερα για τα φθαρτά προϊόντα ή προϊόντα εποχικής ζήτησης στα οποία δεν υπάρχει η προοπτική αποθήκευσης του πλεονάσματος για μελλοντική χρήση. Αυτού του είδους τα προϊόντα μετά την πάροδο κάποιου χρονικού διαστήματος δεν είναι ικανά να πουληθούν για διάφορους λόγους και έτσι το απόθεμα που περισσεύει είναι πλέον άχρηστο οπότε καταστρέφεται, είτε επιστρέφεται στο εργοστάσιο-μονάδα διαχείρισης για να ανακυκλωθεί όποτε αυτό είναι δυνατό. Σε πολλές περιπτώσεις ο έμπορος παίρνει από την επιστροφή των προϊόντων που δεν πουλήθηκαν μια υπολειμματική αξία (salvage value) η οποία είναι μικρότερη από το κόστος αγοράς του ή παραγωγής του από τον ίδιο.

Τέτοιου τύπου προϊόντα είναι οι εφημερίδες-περιοδικά, γάλα, λαχανικά/φρούτα τρόφιμα ευαίσθητα στο χρόνο ή στην έκθεση σε υψηλές θερμοκρασίες, ένδυση-υπόδηση που ακολουθεί κάποια συγκεκριμένα trends και ανανεώνεται κάθε περίοδο αεροπορικά εισιτήρια και εισιτήρια θεατρικών παραστάσεων και πολλά άλλα.

Όπως γίνεται φανερό είναι πολύ σημαντικό σε τέτοιας φύσης προβλήματα να κατανοηθεί η μορφή και οι παράμετροι που επηρεάζουν την ζήτηση με απώτερο σκοπό να αποφασισθεί η βέλτιστη ποσότητα παραγγελίας με σκοπό να μεγιστοποιηθεί το αναμενόμενο κέρδος ή να ελαχιστοποιηθεί το κόστος για τον έμπορο.

Η γενική μορφή-περιγραφή ενός τέτοιου προβλήματος είναι η εξής:

Υποθέτουμε ότι υπάρχει ένας έμπορος ο οποίος πραγματοποιεί παραγγελία για το προϊόν του στην αρχή της κάθε περιόδου (μέρα/μήνας/χρόνος/σεζόν). Η ποσότητα που προμηθεύεται χρησιμοποιείται αποκλειστικά για την ικανοποίηση της ζήτησης την προσεχή περίοδο. Το υπολειπόμενο απόθεμα δεν φυλάσσεται για μελλοντική χρήση. Η ζήτηση για την προσεχή σεζόν δεν είναι εκ των προτέρων γνωστή. Η ζήτηση D_i κάθε περίοδο λοιπόν μπορεί να υποτεθεί ότι είναι μία μη-αρνητική κατανομή για κάθε περίοδο $i \in \mathbb{N}$ με α.σ.κ $F_D(x) = P(D_i \leq x)$. Υποθέτουμε ότι η ζήτηση είναι συνεχής και η σ.π.π είναι $f(x)$. Ο έμπορος πρέπει να καθορίσει την ποσότητα παραγγελίας Q η οποία ελαχιστοποιεί το αναμενόμενο κόστος στο τέλος κάθε περιόδου.

Σε αυτό το σημείο πρέπει να εισάγουμε δύο απαραίτητες ποσότητες οι οποίες παίζουν σημαντικό ρόλο στο πρόβλημα.

1. Το C_o = κόστος ανά μονάδα προϊόντος που μένει απούλητο.

(overstock cost)

2. Το C_u = κόστος ανά μονάδα από τη μη-ικανοποίηση της ζήτησης

(understock cost).

Οι παραπάνω δύο παράγοντες έχουν αντίθετες συνέπειες με την έννοια του ότι αυξάνοντας το απόθεμα αυξάνεται η ποσότητα άρα και το κόστος των προϊόντων που θα μείνουν απούλητα ενώ μειώνοντας το απόθεμα θα αυξηθεί το επίπεδο των ελλείψεων άρα και το αναμενόμενο κόστος. Το πρόβλημα αυτό ασχολείται με την προσέγγιση της βέλτιστης ποσότητας παραγγελίας με σκοπό αφενός την μελλοντική αύξηση του αναμενόμενου κέρδους αφετέρου μείωση αντίστοιχου κόστους των δύο παραπάνω ειδών.

Έστω $K(Q, D)$ το συνολικό κόστος έχοντας απόθεμα Q και παρατηρούμενη ζήτηση D . Η ποσότητα του απούλητου προϊόντος είναι $\max\{0, Q - D\}$, ενώ αντίστοιχα η μη-ικανοποιημένη ζήτηση είναι $\max\{0, D - Q\}$.

$$\text{Επομένως } K(Q, D) = C_o \cdot \max\{0, Q - D\} + C_U \cdot \max\{0, D - Q\}$$

Από τη στιγμή που η ζήτηση, τη στιγμή που αγοράζεται από τον έμπορο το απόθεμα δεν είναι γνωστή αλλά μία τυχαία μεταβλητή, τότε μπορούμε να βρούμε ουσιαστικά τη αναμενόμενη τιμή του κόστους ανά περίοδο.

$$\begin{aligned} E(K(Q)) &= C_o \cdot \int_0^{\infty} \max\{0, Q - D\} f(D) dD + C_U \cdot \int_0^{\infty} \max\{0, D - Q\} f(D) dD = \\ &C_o \cdot \int_0^Q (Q - D) f(D) dD + C_U \cdot \int_Q^{\infty} (D - Q) f(D) dD \end{aligned}$$

Επομένως το πρόβλημα ανάγεται σε πρόβλημα ελαχιστοποίησης του αναμενόμενου κόστους. Παραγωγίζοντας την συνάρτηση κόστους ως προς Q και στη συνέχεια παίρνοντας δεύτερη παράγωγο παίρνουμε,

$$\begin{aligned} \frac{dK(Q, D)}{dQ} &= C_o \cdot \int_0^Q 1 \cdot f(D) dD + C_U \cdot \int_Q^{\infty} -1 \cdot f(D) dD = \\ &C_o \cdot F(Q) + C_U \cdot (1 - F(Q)) \end{aligned}$$

$$\frac{d^2 K(Q, D)}{dQ^2} = (C_o + C_U) f(Q) \geq 0 \quad \forall Q \geq 0$$

Από τη στιγμή που η 2^η παράγωγος είναι πάντα μη-αρνητική η συνάρτηση είναι κυρτή.

Οι ποσότητες C_o , C_U είναι θετικές σταθερές και η F είναι συνεχής αύξουσα συνάρτηση. Άρα ορίζεται η ποσότητα που μηδενίζει την πρώτη παράγωγο και είναι η λύση της εξίσωσης $F(Q) = \frac{C_U}{C_o + C_U}$.

(1.1)

Η ποσότητα στο δεύτερο μέλος της (1.1) ονομάζεται κρίσιμος λόγος

$$0 \leq \frac{C_U}{C_o + C_U} \leq 1 \text{ (critical ratio).}$$

Από την εξίσωση που καθορίζει τη βέλτιστη πολιτική παρατηρούμε ότι η ποσότητα Q^* πρέπει να είναι τέτοια ώστε η πιθανότητα ικανοποίησης της ζήτησης να είναι ίση με το critical ratio. Η πιθανότητα ότι η ζήτηση ικανοποιείται πλήρως σε μια περίοδο είναι ένα κοινό μέτρο απόδοσης μιας πολιτικής παραγγελιών, που σχετίζεται με το βαθμό ικανοποίησης του πελάτη. Συνήθως αναφέρεται ως επίπεδο εξυπηρέτησης τύπου I.

Για παράδειγμα, εάν το επίπεδο μιας πολιτικής παραγγελιών είναι 0,7, αυτό σημαίνει ότι κατά μέσο όρο στο 70% των περιόδων η ποσότητα παραγγελίας είναι αρκετή για να ικανοποιήσει ολόκληρη τη ζήτηση των πελατών. Ο τύπος της πολιτικής παραγγελιών συνδυάζει σε συνοπτική μορφή τις οικονομικές πτυχές με τα ζητήματα εξυπηρέτησης πελατών που εμπλέκονται στη βέλτιστη πολιτική. Προτείνει την παραγγελία μιας τέτοιας ποσότητας ώστε το επίπεδο υπηρεσίας να είναι ίσο με τον κρίσιμο λόγο (critical ratio).

Στην συγκεκριμένη περίπτωση η βέλτιστη ποσότητα παραγγελίας είναι οποιαδήποτε λύση της εξίσωσης (1.1). Όταν η τ.μ της ζήτησης είναι διακριτή (η ζήτηση μπορεί να πάρει ένα αριθμήσιμο σύνολο τιμών) τότε η προηγούμενη εξίσωση μπορεί να μην έχει λύση. Αποδεικνύεται στη γενικότερη περίπτωση ότι για

$$\text{οποιαδήποτε κατανομή η εξίσωση έχει λύση το } Q^* = \min \left\{ Q : F(Q) \geq \frac{C_U}{C_o + C_U} \right\}.$$

Στην περίπτωση που η F είναι συνεχής επειδή σαν αθροιστική συνάρτηση κατανομής, είναι γνησίως αύξουσα στο $[0, +\infty)$ τότε είναι και $1-1$, οπότε ορίζεται η αντίστροφη συνάρτηση F^{-1} και η λύση της εξίσωσης είναι μοναδική και

$$\text{είναι η } Q^* = F^{-1} \left(\frac{C_U}{C_o + C_U} \right)$$

1.2 Παραλλαγές προβλήματος - προσεγγίσεις

Το πρόβλημα αυτό γνωρίζει πολλές παραλλαγές αναλόγως των διαφορετικών υποθέσεων-προσεγγίσεων που γίνονται σχετικά με:

- τον χρονικό ορίζοντα δηλαδή αν εξετάζεται ασυμπτωτικά ή όχι η λύση στο πρόβλημα αυτό (μεγάλος ή μικρός χρονικός ορίζοντας παρατηρήσεων ή λήψης αποφάσεων)
- την γνώση ή μη της πραγματικής ζήτησης (Ντετερμινιστική ή Δυναμική ζήτηση) ή της κατανομής αυτής (παραμετρικά-απαραμετρικά) καθώς και ειδικές υποθέσεις για τις παραμέτρους των κατανομών ζήτησης, και της τιμής του προϊόντος, όπως επίσης ειδικές υποθέσεις για το ύψος ή τη σταθερότητα του αποθέματος ανά τις χρονικές περιόδους.
- την επιρροή στην ζήτηση της τιμής πώλησης, οπότε υπάρχει μια διαδραστικότητα ανάμεσα στις δύο αυτές ποσότητες και σκοπός εκτός από προσδιορισμό μίας βέλτιστης ποσότητας παραγγελίας είναι να καθοριστεί μία νέα βέλτιστη τιμολογιακή πολιτική.
- Την τεχνική προσέγγισης της βέλτιστης ποσότητας παραγγελίας εάν αυτή είναι μέσω βελτιστοποίησης, προσομοίωσης, ή στατιστικών μεθόδων που συνδυάζουν εκτίμηση και αποφάσεις.
- Περικομμένες ή όχι παρατηρήσεις (censored observations or complete data). Στη διεθνή βιβλιογραφία είναι διαδεδομένη το 1^ο σενάριο μιας και πολλά προϊόντα ξεπουλάνε (stock out) σε κάποιες περιόδους.
- Διερεύνηση του κατά πόσον τα δεδομένα που συλλέγονται ακολουθούν κάποια κατανομή όταν δεν υπάρχει κάποια εκ των προτέρων γνώση για αυτά.

Πέρα από τις παραπάνω προσεγγίσεις υπάρχουν και πολλές άλλες που δε θα αναφερθούν εδώ.

Παρακάτω παρουσιάζονται κάποιες εργασίες για το πρόβλημα του εφημεριδοπώλη (newsvendor Problem όπως αναφέρεται στη διεθνή βιβλιογραφία) μαζί με μια πολύ σύντομη ανάλυση των περιεχομένων της κάθε μίας.

1.3

Μερική επισκόπηση βιβλιογραφίας

1. Η εργασία των Xiangwen Lu , Jing-Sheng Song και Kaijie Zhu (2005) [9] στην οποία προσεγγίζεται το πρόβλημα με περικομμένες παρατηρήσεις με γνωστό τον τύπο της κατανομή ζήτησης αλλά με άγνωστες παραμέτρους . Χρησιμοποιείται Μπεϋζιανή προσέγγιση όπου παρουσιάζεται ένας τρόπος υπολογισμού και ανανέωσης της Prior κατανομής των παραμέτρων για τα νέα δεδομένα τα οποία παρατηρούνται καθώς επίσης κατασκευάζεται ένας δυναμικός αλγόριθμος ανανέωσης του αποθέματος για τη μείωση του βέλτιστου κόστους στο μέλλον.
2. Η εργασία της Stefanescu [16], η οποία μελετά τη ζήτηση των εισιτηρίων θεατρικών παραστάσεων ή αεροπορικών εισιτηρίων που είναι ένα προϊόν με συγκεκριμένο διάστημα ζωής και υπολογίζει εκτιμητές μέγιστης πιθανοφάνειας χρησιμοποιώντας μοντέλα γραμμικής παλινδρόμησης με κάποιους παράγοντες που πρακτικά επηρεάζουν τη ζήτηση. Τέτοιες μεταβλητές είναι π.χ το χρονικό διάστημα μέχρι την πτήση, η σεζόν (καλοκαιρινή-χειμερινή-Low), η ώρα ή μέρα πραγματοποίησης μιας θεατρικής παράστασης (απογευματινή βραδινή , Παρασκευή-Σάββατο) αλλά και διάφορες μεταβλητές του προϊόντος όπως είναι η θέση στο θέατρο ή στο αεροπλάνο (οικονομική ή διακεκριμένη θέση).
3. Η εργασία των Burnetas και Smith [3] η οποία μελετά επίσης φθαρτά προϊόντα με περικομμένες παρατηρήσεις ζήτησης και προσεγγίζει απαραμετρικά το πρόβλημα θεωρώντας άγνωστη την κατανομή για τη ζήτηση. Στην εργασία αυτή παρουσιάζεται ένας αλγόριθμος ανανέωσης της ποσότητας παραγγελίας σε κάθε νέα περίοδο χρησιμοποιώντας ως δεδομένα μόνο το αν εξαντλήθηκε το απόθεμα σε κάθε περίοδο ή όχι. Αποδεικνύεται ότι η ακολουθία παραγγελιών των ποσοτήτων συγκλίνει με

πιθανότητα 1 στη μοναδική βέλτιστη ποσότητα παραγγελίας για μεγάλο αριθμό περιόδων.

4. Η εργασία του Nahmias (1994) [11] η οποία υποθέτοντας για τη ζήτηση του προϊόντος κανονική κατανομή με άγνωστες παραμέτρους μελετά το πρόβλημα με περικομμένες παρατηρήσεις και αναλύει τρεις διαφορετικές προσεγγίσεις για τις παραμέτρους.

Η πρώτη είναι ένας τρόπος που επινόησε ο ίδιος ο συγγραφέας να γράψει δηλαδή τους εκτιμητές της μέσης τιμής και της διασποράς σαν συναρτήσεις των παρατηρήσεων. Η δεύτερη είναι οι με εκτιμητές μέγιστης πιθανοφάνειας και η Τρίτη είναι χρησιμοποιώντας γραμμικούς εκτιμητές οι οποίοι έχουν παρουσιαστεί παλαιότερα από τους Saleh και Gupta. Στο τέλος συγκρίνονται οι τρεις αυτές μέθοδοι ως προς την αποτελεσματικότητα και αμεροληψία τους μέσω προσομοίωσης για εξαγωγή συμπερασμάτων.

5. Η εργασία των Liyanage, Shanthikumar (2005) [8] η οποία μελετά μεν την ζήτηση η οποία ακολουθεί εκθετική κατανομή με άγνωστη παράμετρο αλλά αφενός σε σχέση με τις προαναφερθείσες εργασίες δεν εξετάζει περικομμένες παρατηρήσεις και αφετέρου δεν εστιάζει σε διαφορετικούς τρόπους σύγκλισης στην άγνωστη παράμετρο αλλά σε διαφορετικούς τρόπους προσέγγισης της βέλτιστης ποσότητας παραγγελίας για μικρό χρονικό ορίζοντα σύγκλισης και συγκρίνοντας τις μεθόδους μεταξύ τους ως προς την επίπτωση που έχουν στο βέλτιστο κέρδος.

Οι παραπάνω εργασίες από τις οποίες επηρεάστηκε η διπλωματική αυτή εργασία είναι οι τελευταίες δύο των Nahmias και Liyanage, Shanthikumar

Στην επόμενη ενότητα θα παρουσιαστούν αναλυτικά οι προσεγγίσεις και τα αποτελέσματα των δύο αυτών εργασιών.

ΚΕΦΑΛΑΙΟ 2^ο

2.1 Ανάλυση της εργασίας “Demand estimation in lost sales” Nahmias (1994)

Η εργασία πάνω στην οποία βασίστηκε σε μεγαλύτερο βαθμό αυτή η διπλωματική εργασία είναι αυτή του Nahmias [11] οπότε σε αυτήν θα γίνει λίγο πιο εκτενής αναφορά σε σχέση με αυτήν των Liyanage και Shanthikumar [8]. Στην εργασία αυτή υποθέτουμε ένα απόθεμα ενός προϊόντος το οποίο παραγγέλλεται για μία περίοδο μέχρι το ύψος ποσότητας Q σε κάθε περίοδο. Όμως όπως συμβαίνει σε προβλήματα εκτίμησης της ζήτησης εάν οι πωλήσεις υπερβαίνουν το απόθεμα έχουμε περικομμένες παρατηρήσεις (type I censoring) δηλαδή $Q \leq X_i$ οπότε η πληροφορία για τον έμπορο είναι ελλιπής. Το πρόβλημα εκτίμησης της κατανομής της ζήτησης του προϊόντος ανάγεται σε υπολογισμό των παραμέτρων της κατανομής ζήτησης διαθέτοντας ένα από δεξιά περικομμένο δείγμα παρατηρήσεων. Η κατανομή η οποία υποθέτετε για την ζήτηση D είναι η κανονική δηλαδή $N(\mu, \sigma^2)$. Οι εκτιμητές οι οποίοι χρησιμοποιήθηκαν και συγκρίθηκαν είναι οι εξής τρεις.

- i. Οι Εκτιμητές Μέγιστης Πιθανοφάνειας εκτιμητές για το δείγμα (Maximum likelihood estimators).
- ii. Οι BLUE εκτιμητές δηλαδή οι βέλτιστοι γραμμικοί εκτιμητές των διατεταγμένων παρατηρήσεων (Best linear Unbiased Estimators).
- iii. Οι Νέοι εκτιμητές επινόησης του συγγραφέα οι οποίοι γράφονται ως συνδυασμοί των παρατηρήσεων του δείγματος (New estimators).

Πρέπει να αναφερθεί ότι οι εκτιμητές βρίσκονται για δεδομένο γνωστό αριθμό παρατηρήσεων r ενώ στο μοντέλο ο αριθμός αυτός είναι άγνωστος. Επομένως γίνεται μια προσπάθεια να εκτιμηθούν εκ των υστέρων οι παράμετροι χωρίς να υπάρχει η ικανότητα στο μοντέλο που μελετάμε να αποδειχθεί η αμεροληψία και η διασπορά τους παρά μόνο με τη βοήθεια προσομοίωσης στο υπολογιστή.

Για τα παρακάτω υποθέτουμε ότι έχουμε ένα τυχαίο δείγμα από κανονική κατανομή παρατηρήσεων $(x_1, x_2, \dots, x_n)$ με άγνωστο μέσο μ και τυπική απόκλιση. Υποθέτετε επίσης ότι το Q είναι γνωστή σταθερά και μόνο οι τιμές του δείγματος που ξεπερνούν αυτήν είναι άγνωστες. Έστω r ο αριθμός των παρατηρήσεων που των οποίων η τιμή είναι γνωστή, έστω επίσης $x_{(1)}, x_{(2)}, \dots, x_{(n)}$ το διατεταγμένο δείγμα αυτών τότε οι $x_{(1)}, x_{(2)}, \dots, x_{(r)}$ παρατηρήθηκαν αλλά το πλήθος n αυτών είναι και αυτό γνωστός αριθμός.

2.1.1 Εκτιμητές Μέγιστης Πιθανοφάνειας

Οι MLE εκτιμητές έχουν βρεθεί από τον M. Halperin για τη μέση τιμή $\hat{\mu}$ και διασπορά $\hat{\sigma}^2$ είναι $\hat{\mu} = Q - h\sigma$ και $\hat{\sigma} = 0.5(Q - \bar{x}_r)(-h + \sqrt{h^2 + V^2})$ (1.2)

$$\text{όπου } \bar{x}_r = \frac{1}{r} \sum_{i=1}^r x_{(i)} \text{ και } V = 4 \left[1 + \frac{\sum_{i=1}^r (x_i - \bar{x}_r)^2}{r(Q - \bar{x}_r)^2} \right]$$

και το h είναι η λύση της εξίσωσης:

$$\frac{(1/\sqrt{2\pi})e^{-0.5h^2}}{1/\sqrt{2\pi} \int_h^\infty e^{-0.5x^2} dx} = \frac{-r}{n-r} h + \frac{2r}{n-r} \frac{h + \sqrt{h^2 + V^2}}{V^2}.$$

Το h λοιπόν μπορεί να βρεθεί από πίνακες ή από αριθμητική επίλυση της παραπάνω εξίσωσης. Η δυσκολία εδώ έγκειται επιπλέον στην ανανέωση της εκτίμησης για το h μόλις μια νέα τιμή είναι διαθέσιμη. Αυτό οδήγησε και το συγγραφέα στη συγγραφή νέων εκτιμητών που μπορούν πιο εύκολα να ανανεωθούν όταν διατεθούν νέα δεδομένα.

2.1.2 Νέοι εκτιμητές του συγγραφέα (New estimators)

Οι νέοι εκτιμητές μπορούν να γραφούν ως απλές αλγεβρικές συναρτήσεις του δείγματος. Οι εξισώσεις παραθέτονται παρακάτω.

$$\tilde{\sigma}^2 = \frac{s_r^2}{1 - z\varphi(z)/p - [\varphi(z)/p]^2} \quad \text{και} \quad \tilde{\mu} = \bar{x}_r + \frac{\tilde{\sigma}\varphi(z)}{p}$$

$$z = \Phi^{-1}(p) \quad , \quad p = \frac{r}{n} \quad , \quad \bar{x}_r = \frac{1}{r} \sum_{i=1}^r x_{(i)} \quad , \quad s_r^2 = \frac{1}{r-1} \sum_{i=1}^r (x_i - \bar{x}_r)^2$$

και Φ , φ είναι η α.σ.κ και η σ.π.π αντίστοιχα της τυποποιημένης κανονικής τυχαίας μεταβλητής και το p ουσιαστικά εκφράζει το ποσοστό του censoring δηλαδή το ποσοστό των παρατηρούμενων τιμών στο συνολικό αριθμό του δείγματος. Οι παραπάνω τύποι επειδή αποδεικνύονται εύκολα παρακάτω παραθέτετε η διαδικασία δημιουργίας τους.

$$\text{Είναι γνωστό ότι } \mu = \int_{-\infty}^{+\infty} xf(x)dx = \int_{-\infty}^Q xf(x)dx + \int_Q^{+\infty} xf(x)dx .$$

Επειδή τα $x_1, x_2, \dots, x_r \leq Q$, μπορούν να χρησιμοποιηθούν για τον υπολογισμό του πρώτου ολοκληρώματος $\int_{-\infty}^Q xf(x)dx$. Ένας λογικός εκτιμητής για τον υπολογισμό

του είναι το $\frac{\sum_{i=1}^r x_i}{n}$ και αν ορίσουμε τη μέση τιμή μόνο των παρατηρούμενων

τιμών $\bar{x}_r = \frac{\sum_{i=1}^r x_i}{r}$, τότε ο εκτιμητής μπορεί να γραφτεί συναρτήσει της μέσης

τιμής ως $(r/n)\bar{x}_r$. Για την σ.π.π της κανονικής κατανομής είναι γνωστό ότι

$$\int_Q^{\infty} xf(x)dx = \sigma\varphi\left(\frac{Q-\mu}{\sigma}\right) + \mu\left(1 - \Phi\left(\frac{Q-\mu}{\sigma}\right)\right) \text{ και αντικαθιστώντας τον τελευταίο}$$

όρο με $\tilde{\mu}\Phi\left(\frac{Q-\tilde{\mu}}{\tilde{\sigma}}\right) = (r/n)\bar{x}_r + \tilde{\sigma}\varphi\left(\frac{Q-\tilde{\mu}}{\tilde{\sigma}}\right)$ όπου η περισπωμένη παραπέμπει

στους εκτιμητές των παραμέτρων αυτών.

Ο όρος $\Phi\left(\frac{Q-\tilde{\mu}}{\tilde{\sigma}}\right)$ είναι η κανονικοποιημένη πιθανότητα μιας παρατήρησης να

πέφτει κάτω από Q . Ένας εκτιμητής για αυτή την ποσότητα είναι ο $p = \frac{r}{n}$ που

ουσιαστικά είναι το ποσοστό των παρατηρούμενων τιμών του δείγματος. Κάνοντας

αυτήν την αντικατάσταση η σχέση παίρνει τη μορφή $\tilde{\mu}(r/n) = (r/n)\bar{x}_r + \tilde{\sigma}\Phi\left(\frac{Q-\tilde{\mu}}{\tilde{\sigma}}\right)$

που δίνει λύνοντας ως προς $\tilde{\mu}$ γίνεται $\tilde{\mu} = \bar{x}_r + \frac{\tilde{\sigma}\varphi(z)}{p}$. Όπου $z = \frac{Q-\tilde{\mu}}{\tilde{\sigma}}$.

για λόγους συντομίας. Για το σ^2 είναι γνωστό ότι δεδομένου Q ,

$$\sigma^2 = \int_{-\infty}^{+\infty} (x-\mu)^2 f(x) dx = \int_{-\infty}^Q (x-\mu)^2 f(x) dx + \int_Q^{+\infty} (x-\mu)^2 f(x) dx \text{ και πιο}$$

συγκεκριμένα για το 1^ο ολοκλήρωμα ισχύει ότι για την κανονική κατανομή

$$\int_Q^{\infty} (x-\mu)^2 f(x) dx = \sigma^2 \bar{\Phi}\left(\frac{Q-\mu}{\sigma}\right) + \sigma(Q-\mu)\Phi\left(\frac{Q-\mu}{\sigma}\right) = \sigma^2 \bar{\Phi}(z) + \sigma^2 \Phi(z)$$

όπου $\bar{\Phi}(z) = 1 - \Phi(z)$. Ένας λογικός εκτιμητής για το παραπάνω ολοκλήρωμα

βασισμένος στο δείγμα είναι ο $\frac{1}{n-1} \sum_{i=1}^r (x_i - \tilde{\mu})^2 = s_1^2$ όπου δίνει ότι

$$\tilde{\sigma}^2 = s_1^2 + \tilde{\sigma}^2 (\bar{\Phi}(z) + z\Phi(z)) \quad (1) \text{ και χρησιμοποιώντας το δείγμα για να}$$

εκφράσουμε το s_1^2 παίρνουμε ,

$$s_1^2 = \frac{1}{n-1} \sum_{i=1}^r (x_i - \bar{x}_r + \bar{x}_r - \tilde{\mu})^2 = \frac{1}{n-1} \sum_{i=1}^r \left\{ (x_i - \bar{x}_r)^2 + 2(x_i - \bar{x}_r)(\bar{x}_r - \tilde{\mu}) + r(\bar{x}_r - \tilde{\mu})^2 \right\}$$

και επειδή ο ενδιαμέσος όρος είναι ίσος με μηδέν έπεται ότι ,

$$s_1^2 = \frac{1}{n-1} \sum_{i=1}^r (x_i - \bar{x}_r)^2 + \frac{r}{n-1} (\bar{x}_r - \tilde{\mu})^2 \quad (2) \text{ και ορίζοντας}$$

$$s_r^2 = \frac{\sum_{i=1}^r (x_i - \bar{x}_r)^2}{r-1} \text{ τη δειγματική διασπορά ο προηγούμενος τύπος γίνεται}$$

$$s_1^2 = \frac{r-1}{n-1} s_r^2 + \frac{r\tilde{\sigma}^2}{(n-1)p^2} \varphi(z)^2 \quad (3) \text{ οπότε αντικαθιστώντας τη σχέση (3) στη σχέση (1)}$$

$$\text{προκύπτει } \tilde{\sigma}^2 = \frac{ps_r^2}{p - z\varphi(z) - \varphi(z)^2/p} = \frac{s_r^2}{1 - z\varphi(z)/p - (\varphi(z)/p)^2}$$

που είναι ο τελικός τύπος εκτίμησης της διασποράς.

2.1.3 Βέλτιστοι γραμμικοί εκτιμητές διατεταγμένων παρατηρήσεων (best linear unbiased estimators BLUEs)

Οι BLUE (best linear unbiased estimators) εκτιμητές είναι οι καλύτεροι γραμμικοί εκτιμητές των παραμέτρων της κατανομής όπως υποδηλώνει το όνομα τους. Είναι γραμμικός συνδυασμός των τιμών του δείγματος x_i . Οι καλύτεροι μη-μεροληπτικοί εκτιμητές για περικομμένα δείγματα κατασκευάστηκαν εφαρμόζοντας το κριτήριο των ελαχίστων τετραγώνων για την ελαχιστοποίηση της διασποράς της γραμμικής εκτίμησης. Η μέθοδος αυτή απαιτεί την αντιστροφή ενός $r \times r$ πίνακα. Για $n \geq 20$ και προς αποφυγή χρονοβόρων υπολογισμών και αντιστροφής μεγάλων διαστάσεων πινάκων ο Gupta ανέπτυξε μια προσέγγιση βασισμένη στα order statistics της τυποποιημένης κανονικής κατανομής. Οι εκτιμητές είναι της μορφής :

$$\mu^* = \sum_{i=1}^r a_{i,r} x_i \text{ και για την τυπική απόκλιση } \sigma^* = \sum_{i=1}^r b_{i,r} x_i \text{ όπου οι παρατηρήσεις } x_i$$

θεωρούνται ότι είναι σε αύξουσα σειρά. Οι BLUE εκτιμητές έχουν το θετικό ότι είναι αμερόληπτοι το οποίο δεν ισχύει απαραίτητα για τους MLE και τους New. Όμως έχουν και ένα μειονέκτημα όταν χρησιμοποιούνται για την κανονική κατανομή διότι

απαιτούν τον υπολογισμό κατά σειρά $\frac{n^2}{2}$ σταθερών για κάθε παράμετρο. Μέχρι τη

συγγραφή της εργασίας είχαν βρεθεί πίνακες για $1 \leq n \leq 20$. Παρ' όλο που φαίνεται αρκετά χρονοβόρα η διαδικασία υπολογισμού των εκτιμητών αυτών ο Gupta πρότεινε μία προσέγγιση για μεγάλα δείγματα η οποία φαίνεται αρκετά ακριβής. Η προσέγγιση βασίζεται στη θεώρηση ότι ο πίνακας διασποράς-συνδιακύμανσης της

από κοινού κατανομής των διατεταγμένων παρατηρήσεων της τυποποιημένης κανονικής είναι ο μοναδιαίος πίνακας. Οι γραμμικές εκτιμήσεις απαιτούν τη γνώση μόνο της μέσης τιμής των διατεταγμένων παρατηρήσεων της $N(0,1)$. Αυτές όμως έχουν πινακοποιηθεί για μεγάλα n έως 1000. Εξαιτίας της συμμετρίας της συγκεκριμένης κατανομής χρειάζεται η αποθήκευση μόνο των $n/2$ τιμών του πίνακα. Η μέθοδος εξοικονομεί πολύ χρόνο σε χώρο αποθήκευσης που απαιτείται για την αποθήκευση των τιμών αλλά και σε χρόνο υπολογισμού των συντελεστών των σταθερών. Ο αλγόριθμος που κατασκεύασε ο Gurta περιγράφεται από τις εξής

$$\text{σχέσεις: } a_{i,r} = \frac{1}{r} - \frac{\bar{\mu}_r (\mu_i - \bar{\mu}_r)}{\sum_{i=1}^r (\mu_i - \bar{\mu}_r)^2} \text{ και } b_{i,r} = \frac{(\mu_i - \bar{\mu}_r)}{\sum_{i=1}^r (\mu_i - \bar{\mu}_r)^2}$$

$$\text{Όπου } \bar{\mu}_r \text{ υπολογίζεται από τον τύπο } \bar{\mu}_r = \frac{1}{r} \sum_{i=1}^r \mu_i$$

Από εδώ και πέρα για προσδιορισμό της απόδοσης και των τριών παραπάνω εκτιμητών έγιναν προσομοιώσεις με 1.000 και 10.000 επαναλήψεις για την απόδοση τους και παρατέθηκαν τα αποτελέσματα σε πίνακες δίνοντας την πραγματική τιμή των παραμέτρων, την εκτίμησή τους αλλά και το τετραγωνικό σφάλμα των εκτιμητών. Οι διαφορετικές περιπτώσεις που πάρθηκαν έχουν να κάνουν με το πλήθος των παρατηρήσεων $n = 10, 40, 70, 100$ αλλά και το απόθεμα σε σχέση με τη μέση τιμή της ζήτησης οπότε έχουμε $Q \leq \mu$ και $Q > \mu$.

2.1.4 Συμπεράσματα προσομοιώσεων για τις εκτιμήτριες

Εδώ θα αναφέρουμε τα κυριότερα συμπεράσματα από τα αποτελέσματα αυτά περιληπτικά παρακάτω:

1. Οι νέοι εκτιμητές του συγγραφέα για $n = 10$ και $Q \leq \mu$ έδωσαν πιο μικρή απόδοση μετρημένη σε τυπική απόκλιση και μέση τιμή σε σχέση με τους BLUE κ τους MLE. Επίσης για μεγαλύτερα δείγματα και $Q \leq \mu$ ο νέος εκτιμητής έδωσε ελαφρώς μεγαλύτερη διασπορά.

2. Για $Q > \mu$ ο νέος εκτιμητής καθώς και ο BLUE έδωσε αποτελέσματα παρόμοια με τους MLE εκτιμητές. Οι διασπορές των MLE εκτιμητών ήταν μικρότεροι σε κάθε περίπτωση αλλά οι διαφορές ήταν μικρότερες για μεγάλα n και μεγαλύτερο Q .
3. Σε κάθε άλλη περίπτωση και οι τρεις μέθοδοι δίνουν την ίδια αποτελεσματικότητα. Η εκτίμηση για τη διασπορά από τις 2 αυτές προσεγγίσεις είναι πολύ κοντά στην πραγματική τιμή της μέσης τιμής και της τυπικής απόκλισης δείχνοντας ότι και οι δύο εκτιμητές είναι αμερόληπτοι.
4. Όμως επειδή η ρεαλιστική περίπτωση περιλαμβάνει συνήθως μικρά δείγματα τότε οι MLE εκτιμητές προτιμώνται ξεκάθαρα για την μεγαλύτερη αποτελεσματικότητα που έχουν για μικρότερα n .

2.1.5 Ανανέωση εκτίμησης για τη ζήτηση και το δείγμα (Sequential updating)

Η εργασία αυτή τελειώνει με αλγόριθμους ανανέωσης (sequential updating) των αποτελεσμάτων κάθε φορά που μία νέα παρατήρηση είναι διαθέσιμη. Για τους Blue εκτιμητές του Gupta η μέθοδος moving averages φαίνεται ιδανική ενώ για τους MLE και τους νέους εκτιμητές προτιμάται η μέθοδος της εκθετικής εξομάλυνσης (exponential smoothing).

Το 1^ο βήμα σε αυτή τη μέθοδο είναι να επιλεγθεί ο αριθμός των περιόδων που θέλει κανείς να έχουν περάσει για να συλλέξει τα δεδομένα και να χρησιμοποιήσει για τις εκτιμήσεις του. Για τον γραμμικό εκτιμητή το 0.1 είναι μια κοινά αποδεκτή σταθερά εξομάλυνσης η οποία ανταποκρίνεται σε δείγμα $n = 19$ για μετακίνηση των μέσων όρων θα χρειαστεί ένα δείγμα $n = 20$. Για $n \leq 20$ μπορεί να χρησιμοποιηθεί η ίδια η BLUE ή η Gupta εκτίμηση.

Ο MLE και ο νέος εκτιμητής βασίζονται σε τρεις μεταβλητές που εξαρτώνται μόνο από το δείγμα. Έστω $\bar{D}_n$ η εκτίμηση της μέσης τιμής του δείγματος που έχει εξομαλυνθεί για τιμές μικρότερες του Q . Οι εξισώσεις είναι οι εξής:

- $\bar{D}_n = \alpha D_n + (1 - \alpha) \bar{D}_{n-1}$ εάν και $D_n < Q$

- $\bar{D}_n = \bar{D}_{n-1}$ εάν $D_n \geq Q$

με D_n η ζήτηση την περίοδο n ενώ αντίστοιχα η εξομαλυμένη εκτίμηση για το p που είναι το ποσοστό των παρατηρήσεων μικρότερων από Q είναι:

$$p_n = \alpha I_n + (1 - \alpha) p_{n-1} \text{ όπου το } I_n \text{ είναι σταθερά } 0 \text{ ή } 1 \text{ δηλαδή}$$

- $I_n = 1$ εάν $D_n < Q$ και $I_n = 0$ αν $D_n \geq Q$.

2.2 Παρουσίαση της εργασίας των Liyanage, Shanthikumar (2005)

“A practical inventory control policy using operational statistics”

Οι Liwan H. Liyanage, J. George Shanthikumar[8] ασχολούνται με το πρόβλημα του εφημεριδοπώλη χωρίς περικομμένες παρατηρήσεις, υποθέτοντας όπως προαναφέρθηκε εκθετική ζήτηση με άγνωστη παράμετρο θ . Ζητούμενο είναι η εκτίμηση της παραγγελίας του αποθέματος την $n+1$ έχοντας παρατηρήσει τη ζήτηση για n περιόδους. Πολλά από τα αποτελέσματα που θα χρησιμοποιηθούν εδώ αποδεικνύονται παρακάτω στο πλαίσιο αυτής της εργασίας αναλυτικά εδώ όμως θα αναφερθούν απλά για οικονομία χρόνου.

Ο έμπορος αγοράζει την κάθε μονάδα προϊόντος σε μια τιμή $c > 0$ και το πουλάει σε τιμή $s > c$. Οι υπολειπόμενες μονάδες αποστέλλονται πίσω παίρνοντας την υπολειπόμενη αξία h (salvage value) η οποία είναι τέτοια ώστε $s > c > h$.

Σε αυτήν την εργασία χωρίς βλάβη της γενικότητας θεωρείται επίσης ότι $h=0$. Η συσχέτιση των συμβόλων C_o, C_U με τα s, c, h είναι η εξής: $C_o = c - h$, $C_U = s - c$

οπότε ο κρίσιμος λόγος ανάγεται σε λόγο $r = \frac{s-c}{s-h}$ όπου λόγω υπόθεσης ότι $h=0$

απλοποιείται και γίνεται $r = \frac{s-c}{s} = 1 - \frac{c}{s}$. Η κατανομή ζήτησης είναι εκθετική με

παράμετρο θ , οπότε έστω $F_D(x) = 1 - e^{-x \cdot \theta}$ η α.σ.κ της ζήτησης D .

Εδώ έχουμε μη-περικομμένες παρατηρήσεις της ζήτησης, έστω X_i κάθε περίοδο i ταυτίζονται με τη ζήτηση D_i την ίδια περίοδο. Η συνάρτηση κέρδους την περίοδο $n+1$ υποθέτοντας ότι η παραγγελία είναι x μονάδες δίνεται από τον τύπο $\Phi(x, D_{n+1}) = s \cdot \min(x, D_{n+1}) - c \cdot x$ ενώ το αναμενόμενο κέρδος

$$\varphi(x, \theta) = E(\Phi(x, D(\theta))) = s \int_{y=0}^x (1 - F_D(y, \theta)) dy - cx. \quad (1.3)$$

Όπως έχει αποδειχθεί και στην εισαγωγή η βέλτιστη ποσότητα παραγγελίας είναι η

$$x^*(\theta) = \inf \left\{ x : F_D(x, \theta) \geq 1 - \frac{c}{s}, x \geq 0 \right\} \text{ οποία είναι } \bar{F}_D^{inv} \left(\frac{c}{s} \right)$$

Το βέλτιστο αναμενόμενο κέρδος είναι $\psi^*(\theta) = \varphi(x^*(\theta), \theta)$ και από τη
 μεγιστοποίηση της συνάρτησης κέρδους παίρνουμε ότι $x^*(\theta) = \frac{1}{\theta} \ln\left(\frac{s}{c}\right)$.

Οι τρεις σημαντικότεροι τρόποι τους οποίους μελετά η εργασία αυτή είναι για τη
 βραχύχρονη επίπτωση τους στο βέλτιστο κέρδος είναι:

A) Χρησιμοποιώντας μόνο τον Ε.Μ.Π για το θ που είναι η μέση τιμή του δείγματος.
 (*sample mean*)

B) Χρησιμοποιώντας εκτιμήσεις της βέλτιστης πολιτικής και το δειγματικό μέσο για
 τον θ (*Os: sm*)

Γ) Χρησιμοποιώντας εκτιμήσεις της βέλτιστης πολιτικής και εμπειρική συνάρτηση
 κατανομής του δείγματος για το θ (*Os: emp*)

Δ) Χρησιμοποιώντας μόνο την εμπειρική συνάρτηση κατανομής του δείγματος
 (*emp*)

2.2.1 Με χρήση του δειγματικού μέσου (*sample mean*)

Χρησιμοποιώντας το δειγματικό μέσο για την εκτίμηση της παραμέτρου που είναι

$\hat{\theta} = \bar{D} \ln\left(\frac{s}{c}\right)$ με αντικατάσταση του στη σχέση (1.3) το εκ των προτέρων

αναμενόμενο κέρδος από πράξεις γίνεται:

$$\psi_{SM}(\theta) = \frac{c}{\theta} \left\{ \frac{s}{c} - \frac{s}{c} \left(\frac{n}{n + \ln(s/c)} \right)^n - \ln(s/c) \right\}.$$

Αποδεικνύεται ότι το $\psi_{SM}(\theta)$:

- i) είναι αύξουσα συνάρτηση του n

$$\text{ii)} \quad \psi_{SM}(\theta) \rightarrow \psi^*(\theta) \quad \text{καθώς } n \rightarrow \infty$$

$$\text{iii)} \quad \psi_{SM}(\theta) < \psi^*(\theta)$$

Προσομοιώνοντας π.χ για $n=4$ και $\frac{s}{c}=1.2$ παίρνουμε ότι η ποσοστιαία

διαφορά $\frac{\psi^*(\theta) - \psi_{SM}(\theta)}{\psi^*(\theta)} \cdot 100\%$ είναι ανεξάρτητη του θ και είναι της τάξης του

22,86% δηλαδή σχεδόν το ένα τέταρτο του αναμενόμενου κέρδους χάνεται λόγω της έλλειψης πληροφορίας για την κατανομή του θ .

2.2.2 Με χρήση δειγματικού μέσου και εκτιμήσεων βέλτιστης πολιτικής (sample mean and use of operational statistics)

Σύμφωνα με αυτήν την προσέγγιση αντικαθίσταται η ποσότητα παραγγελίας

$$\text{που εφαρμόζεται είναι ίση με } X_{OS:Sm} = n \left(\left(\frac{s}{c} \right)^{\frac{1}{n+1}} - 1 \right) \bar{D} . \quad (1.4)$$

Αυτή είναι μία μεροληπτική αποθεματική πολιτική και ισχύει πάλι όπως και πριν ότι $X_{OS:Sm} \rightarrow x^*(\theta)$ δηλαδή είναι ισχυρά συνεπής εκτιμητής. Επειδή εξαρτάται από

τα n και $\frac{s}{c}$ σε κάθε περίπτωση μπορεί είτε να ισχύει ότι $X_{OS:Sm} < X_{SM}$ είτε

$$X_{OS:Sm} > X_{SM}$$

$$\text{Το μέσο κέρδος εδώ είναι } \psi_{OS:Sm} = \frac{c}{\theta} \left\{ \frac{s}{c} - 1 - (n+1) \left(\left(\frac{s}{c} \right)^{\frac{1}{n+1}} - 1 \right) \right\}$$

Αυτό αυξάνει με το n και συγκλίνει στο $\psi^*(\theta)$. Εδώ κάνοντας προσομοίωση για τα

ίδια n και $\frac{s}{c}$ όπως πριν παίρνουμε μια βελτίωση της τάξης του 4.96% για το

βέλτιστο κέρδος σε σχέση με την προηγούμενη προσέγγιση.

2.2.3 Με χρήση της εμπειρικής συνάρτησης κατανομής (empirical distribution)

Εδώ χρησιμοποιείται σαν κατανομή του δείγματος η εμπειρική συνάρτηση κατανομής και επειδή η συνάρτηση είναι συνεχής, θέτοντας $D_{[r]}$ τα order statistic της ζήτησης, η α.σ.κ προσεγγίζεται από τον τύπο

$$\hat{F}_D = \frac{1}{n} \left\{ r - 1 + \frac{x - D_{[r-1]}}{D_{[r]} - D_{[r-1]}} \right\} \quad \text{για } D_{[r]} \leq x \leq D_{[r+1]} \quad r = 1, 2, \dots, n$$

με βέλτιστη ποσότητα παραγγελίας $X_{emp} = \bar{F}_D^{inv} \left(\frac{c}{s} \right) = D_{[\hat{r}]} + \hat{a} (D_{[\hat{r}]} - D_{[\hat{r}-1]})$

με το $\hat{r}$ να είναι ακέραιος που να ικανοποιεί την

$$n \left(1 - \frac{c}{s} \right) < \hat{r} \leq n \left(1 - \frac{c}{s} \right) + 1 \quad \text{και} \quad \hat{a} = n \left(1 - \frac{c}{s} \right) + 1 - \hat{r}$$

Το X_{emp} είναι ισχυρά συνεπής εκτιμητής της βέλτιστης ποσότητας παραγγελίας.

Χρησιμοποιώντας το ότι οι διαφορές των διαδοχικών διατεταγμένων παρατηρήσεων της ζήτησης ακολουθούν εκθετική κατανομή, συμβολικά δηλαδή

$$Z_k = D_{[k]} - D_{[k-1]} \sim \text{Exp}(-\theta(n-k+1)) \quad \text{και} \quad D_{[r]} + a(D_{[r]} - D_{[r-1]}) = \sum_{k=1}^{r-1} Z_k + aZ_r$$

οπότε αντικαθιστώντας στην εξίσωση του αναμενόμενου βέλτιστου κέρδους παίρνουμε:

$$\psi_{emp} = \frac{c}{\theta} \left\{ \frac{s}{c} \left(1 - \left(\frac{n - \hat{r} + 2}{n + 1} \right) \left(\frac{n - \hat{r} + 1}{n - \hat{r} + 1 + \hat{a}} \right) \right) - \sum_{k=1}^{\hat{r}-1} \frac{1}{n - k + 1} - \frac{\hat{a}}{n - \hat{r} + 1} \right\}.$$

Εδώ από προσομοίωση χρησιμοποιώντας τις ίδιες τιμές με προηγουμένως παίρνουμε αναμενόμενη απόκλιση από το πραγματικό βέλτιστο κέρδος της τάξης του 73.06% που είναι πολύ μεγαλύτερη από τη χρήση του δειγματικού μέσου πριν.

2.2.4 Με χρήση της εμπειρικής συνάρτησης κατανομής και εκτιμήσεων της βέλτιστης πολιτικής (empirical distribution and use of operational statistics)

Εδώ στην τελευταία προσέγγιση της βέλτιστης ποσότητας παραγγελίας οι συγγραφείς θέτουν $X_{Os:emp} = D_{[r^*-1]} + a^* (D_{[r^*]} - D_{[r^*-1]})$ όπου r^*, a^* δίνονται από

$$r^* = \arg \min \left\{ 2\sqrt{\frac{s}{c} \left(\frac{n-r+2}{n+1} \right)} + \sum_{k=1}^{r-1} \frac{1}{n-k+1} \quad , \quad r=1, 2, \dots, n; \quad r \leq (n+1) \left(1 - \frac{c}{s} \right) \right\}$$

$$a^* = (n-r^*+1) \left\{ \sqrt{\frac{s}{c} \left(\frac{n-r^*+2}{n+1} \right)} - 1 \right\}$$

Αντικαθιστώντας στη συνάρτηση κέρδους αυτό παίρνει τη μορφή

$$\psi_{Os:emp} = \frac{s}{\theta} \left(1 - \left(\frac{n-r^*+2}{n+1} \right) \left(\frac{n-r^*+1}{n-r^*+1+a^*} \right) \right) - \frac{c}{\theta} \left(\sum_{k=1}^{r^*-1} \frac{1}{n-k+1} - \frac{a^*}{n-r^*+1} \right)$$

Συγκρίνοντας τα παραπάνω αποτελέσματα των επιδράσεων των παραγγελιών στο βέλτιστο κέρδος, οι συγγραφείς καταλήγουν στο τελικό και ουσιώδες αποτέλεσμα στην έρευνα τους οι εξής ανισοτικές σχέσεις: $\psi_{Os:sm}(\theta) > \psi_{Os:emp}(\theta) > \psi_{emp}(\theta)$ και $\psi_{Os:sm}(\theta) > \psi_{sm}(\theta)$. Επομένως η $\psi_{Os:sm}(\theta)$ είναι η βέλτιστη πολιτική ανάμεσα σε αυτές τις κλάσεις πολιτικής παραγγελιών όπου υπενθυμίζουμε ότι η βέλτιστη

ποσότητα παραγγελίας είναι από την (1.4) $X_{OS:Sm} = n \left(\left(\frac{s}{c} \right)^{\frac{1}{n+1}} - 1 \right) \bar{D}.$

2.3 Η σχέση της παρούσας διπλωματικής-επιρροές της σε σχέση με τις παραπάνω ακαδημαϊκές εργασίες.

Για την εργασία μου αυτή επηρεάστηκα πολύ κατά την ανάγνωση των δύο παραπάνω εργασιών. Η θέληση ήταν να αναπτυχθεί μια εργασία η οποία να επεκτείνει την εργασία του Nahmias με διαφορετική κατανομή ζήτησης, λαμβάνοντας υπόψιν διαφορετικούς εκτιμητές για την άγνωστη παράμετρο, η οποία θα δώσει μια όσο το δυνατόν καλύτερη εικόνα στον εκάστοτε έμπορο σχετικά με την κατανομή της ζήτησης και με την πολιτική αποθήκευσης και παραγγελίας εμπορεύματος για μελλοντικές περιόδους. Σε αντίθεση λοιπόν με τον Nahmias που επέλεξε την κανονική κατανομή με δύο άγνωστες παραμέτρους επέλεξα την εκθετική κατανομή αφενός γιατί σύμφωνα με την εργασία των Braden και Freimer (1991) [2] αποτελεί ειδική περίπτωση της newsvendor κλάσης κατανομών αφετέρου μπορεί να χρησιμοποιηθεί η μέθοδος των Gupta & Saleh για τους γραμμικούς εκτιμητές με χρήση διατεταγμένων παρατηρήσεων από το δείγμα. Να σημειωθεί εδώ ότι μία τυχαία μεταβλητή X είναι μέλος της οικογένειας κατανομών του εφημεριδοπώλη εάν έχει συνάρτηση πυκνότητας πιθανότητας $f(x|\theta) = \theta d(x)e^{-\theta d(x)}$ όπου η συνάρτηση d ορίζεται $d : (0, \infty) \rightarrow (0, \infty)$, και πρέπει να ικανοποιεί $\lim_{x \rightarrow 0} d(x) = 0$, $\lim_{x \rightarrow \infty} d(x) = \infty$ και $d'(x) > 0$ για $x > 0$. Η εκθετική λοιπόν είναι μια ειδική περίπτωση αυτών των newsvendor κατανομών με $d(x) = x$. Ακριβώς όπως στην εργασία του Nahmias λοιπόν συγκρίνονται και εδώ μέθοδοι εκτίμησης παραμέτρου μία εκ των οποίων, η προσομοιωτική αναπτύχθηκε σε αυτή την εργασία. Για την 2^η εργασία στην οποία αναφέρομαι, αυτή των Liyanage & Shanthikumar[] έρχεται να προστεθεί στην προσπάθεια εκτίμησης της παραμέτρου με προσέγγιση της βέλτιστης ποσότητας παραγγελίας που συμπίπτει με τη μέση τιμή της κατανομής ζήτησης $1/\theta$ για μικρούς όμως χρονικούς ορίζοντες. Στην παρούσα εργασία όπως θα αναφερθεί και παρακάτω σε αντίθεση με την προαναφερθείσα το δείγμα θεωρείται με περικομμένες παρατηρήσεις λόγω γεγονότων ξεπουλήματος (stock out) που γίνονται σε ένα ποσοστό των περιόδων. Θυμόμαστε από την ενότητα 2.3 η βέλτιστη ποσότητα παραγγελίας είναι

$x^*(\theta) = \frac{1}{\theta} \ln\left(\frac{s}{c}\right)$ όπου με προσέγγιση του μέσου της ζήτησης $1/\theta$ από μία

εκτίμηση $1/\hat{\theta}$ αντικαθιστούμε για εκτίμηση της βέλτιστης ποσότητας

$x_{est} = x^*(\hat{\theta}) = \frac{1}{\hat{\theta}} \ln\left(\frac{s}{c}\right)$. Εδώ εμείς θα δούμε το κατά πόσον επηρεάζουν τα

ελλιπή δεδομένα το μέγιστο κέρδος αντικαθιστώντας την εκτίμηση για το θ από το μέσο $\bar{D}$ που χρησιμοποιείται με κάποια από τις εκτιμήτριες στο συγκεκριμένο

μοντέλο όπως π.χ την Ε.Μ.Π που είναι $\frac{y}{r} = \frac{\sum_{i=1}^r X_{(i)} + (n-r)Q}{r}$ όπως θα δούμε παρακάτω.

Στην εργασία λοιπόν των Liyanage & Shanthikumar[8] συγκρίνονται τέσσερις τρόποι για την βέλτιστη ποσότητα παραγγελίας. Εμείς θα προσομοιώσουμε προς σύγκριση τους μόνο δύο τρόπους και θα αποφύγουμε την μέθοδο μέσω της εμπειρικής συνάρτησης κατανομής (γιατί δεν μπορεί να υπολογιστεί αυτή με ευκολία στο δείγμα με τα ελλιπή δεδομένα) για να δούμε αφενός το κόστος στο βέλτιστο κέρδος της απώλειας δεδομένων λόγω περικοπής αλλά και το αν στο παρών πρόβλημα παραμένει βέλτιστη η ποσότητα παραγγελίας όπως αυτή ορίστηκε στην ενότητα 2.2.2 σε σχέση με αυτή της 2.2.1.

ΚΕΦΑΛΑΙΟ 3^ο

3.1 Υποθέσεις-Προσέγγιση στο συγκεκριμένο υπόδειγμα.

Οι δικές μας υποθέσεις αφορούν λοιπόν τη ζήτηση ενός φθαρτού προϊόντος η οποία είναι μια τυχαία μεταβλητή ανεξάρτητη από κάποια άλλη περίοδο η οποία ακολουθεί την εκθετική κατανομή δηλαδή . Ο έμπορος αγοράζει την κάθε μονάδα προϊόντος σε μια τιμή $c > 0$ και το πουλάει σε τιμή $s > c$. Οι υπολειπόμενες μονάδες αποστέλλονται πίσω παίρνοντας την υπολειπόμενη αξία h η οποία είναι τέτοια ώστε $s > c > h$. Σε αυτήν την εργασία χωρίς βλάβη της γενικότητας θεωρούμαι ότι $h = 0$.

Ο έμπορος έχει στη διάθεση του ένα περιορισμένο χρονικό ορίζοντα n περιόδων ο οποίος για να προσεγγιστεί το πρόβλημα με ρεαλισμό θα υποτεθεί μικρός δηλαδή γενικότερα $n < 30$. Η κατανομή της ζήτησης κάθε περίοδο D_i είναι γνωστό ότι ακολουθεί την εκθετική κατανομή , χωρίς να είναι γνωστή ωστόσο η παράμετρος της θ . Επίσης είναι ανεξάρτητη από τη ζήτηση κάθε άλλη περίοδο οπότε συμβολικά έχουμε D_i ακολουθεί $Exp(\theta)$ i.i.d για κάθε $i \in N$. Ο έμπορος έχει στη διάθεση του λοιπόν, ένα σύνολο n περιόδων στο οποίο κρατά σταθερό το απόθεμα του με σκοπό να μπορέσει να κάνει εκτιμήσεις για την παράμετρο της κατανομής. Ορίζουμε τη νέα μεταβλητή X_i , όπου X_i τυχαία μεταβλητή που παριστάνει τις πωλήσεις σε κάθε περίοδο i . Οπότε παίρνουμε ένα δείγμα $X_1, X_2, \dots, X_n$ από το οποίο σε έστω r περιόδους έχουμε $X_i = D_i < Q$ δηλαδή ο έμπορος έχει γνώση της πραγματικής ζήτησης η οποία ουσιαστικά δεν υπερβαίνει το απόθεμα ενώ τις υπόλοιπες $n - r$ περιόδους, $D_i \geq Q$ οπότε για αυτές τις περιόδους $X_i = Q \leq D_i$.

Σκοπός μας είναι να κάνουμε εκτίμηση για την $n+1$ περίοδο της ζήτησης με απώτερο σκοπό τη μεγιστοποίηση του κέρδους για τον έμπορο.

Η κατανομή ζήτησης είναι εκθετική με παράμετρο θ οπότε έστω $F_D(x) = 1 - e^{-x\theta}$ η αθροιστική συνάρτηση κατανομής της ζήτησης D .

Η συνάρτηση κέρδους για την περίοδο συναρτήσει της νέας παραγγελίας q είναι $\Phi(q, D_{n+1}) = s \cdot \min(q, D_{n+1}) - c \cdot q$ $q \geq 0$, όπου $D_{n+1} \stackrel{D}{=} D(\theta)$ ενώ το αναμενόμενο κέρδος υπολογίζοντας την μέση τιμή της συνάρτησης κέρδους,

$$\begin{aligned}\varphi(q, \theta) &= E(\Phi(q, D(\theta))) = s \int_{y=0}^q (1 - F_D(y, \theta)) dy - cq \\ &= \frac{s}{\theta} (1 - \exp(-q\theta)) - cq\end{aligned}$$

όπως είναι γνωστό η βέλτιστη ποσότητα παραγγελίας είναι τέτοια ώστε

$$q^*(\theta) = \inf \{q : F_D(q, \theta) \geq 1 - \frac{c}{s}\} \quad \text{όπου} \quad \bar{F}_D^{\text{inv}}\left(\frac{c}{s}\right) \text{ η αντίστροφη της συνάρτησης}$$

$$\text{επιβίωσης δηλαδή } \bar{F}_D = 1 - F_D$$

Συμβολίζοντας με $\psi(q, \theta) = \varphi(q^*(\theta), \theta)$ το βέλτιστο αναμενόμενο κέρδος,

Και μεγιστοποιώντας τη συνάρτηση κέρδους μέσω παραγώγισης παίρνουμε την

$$\text{βέλτιστη ποσότητα παραγγελίας } q^*(\theta) = \frac{1}{\theta} \ln\left(\frac{s}{c}\right) \text{ και αντικαθιστώντας την}$$

παίρνουμε το βέλτιστο κέρδος

$$\psi(q^*, \theta) = \frac{s}{\theta} \left(1 - \frac{c}{s}\right) - \frac{c}{\theta} \ln\left(\frac{s}{c}\right) = \frac{c}{\theta} \left(\frac{s}{c} - 1 - \ln\left(\frac{s}{c}\right)\right)$$

Φαίνεται αρκετά ξεκάθαρο ότι ενδιαφερόμαστε για τη μέση τιμή της παραγγελίας του αποθέματος Q το οποίο είναι η μέση τιμή της ζήτησης. Αυτό το παρατηρούμε από τον παραπάνω τύπο ότι για τον υπολογισμό του βέλτιστου κέρδους δεν είναι αναγκαίος ο υπολογισμός της παραμέτρου θ της εκθετικής κατανομής της τυχαίας μεταβλητής αλλά της μέσης τιμής αυτής $1/\theta$. Δίνεται

ιδιαίτερη έμφαση στην παραπάνω παρατήρηση γιατί δεν θα ψάξουμε του θ αλλά του λόγου $1/\theta$ γιατί όπως γνωρίζουμε γενικότερα δεν ισχύει απαραίτητα ότι αν δ π.χ. αμερόληπτος εκτιμητής του θ ότι και αντίστοιχα $1/\delta$ α.ε του $1/\theta$.

Οι 4 μέθοδοι εκτίμησης της παραμέτρου $1/\theta$ της εκθετικής κατανομής τους οποίους θα αναπτύξουμε είναι οι εξής:

- Μέθοδος Μέγιστης πιθανοφάνειας με χρήση των παρατηρήσεων ή του πλήθους των μη-περικομμένων δεδομένων (M.L.E)
- Μέθοδος Μέγιστης πιθανοφάνειας χωρίς χρήση των παρατηρήσεων αλλά του πλήθους αυτών που δεν υπερβαίνουν το απόθεμα.
- Μέθοδος βέλτιστου γραμμικού εκτιμητή των διατεταγμένων παρατηρήσεων (BLUEs)
- Μέθοδος προσομοίωσης των εκλιπουσών τιμών από το δείγμα.
(Simulating estimator)

Μία ουσιώδη παρατήρηση είναι να δούμε ότι η επιλογή του ύψους του αποθέματος Q από τον έμπορο στην αρχή των n περιόδων επηρεάζει και το ποσοστό των παρατηρούμενων τιμών ζήτησης-πωλήσεων του προϊόντος επί του συνόλου του δείγματος των πωλήσεων που ουσιαστικά είναι αυτές ώστε $D_i = X_i < Q$. Επομένως ας συμβολίσουμε με r την τυχαία μεταβλητή που συμβολίζει το πλήθος των παρατηρήσεων του δείγματος που δεν υπερβαίνουν το απόθεμα Q . Τότε η κατανομή της τυχαίας μεταβλητής r εξαρτάται από Q , και δεν έχουμε γνώση για αυτήν και το πως αυτή επηρεάζεται παρά μόνο εάν γίνουν συγκεκριμένες υποθέσεις για την επιλογή του αποθέματος. Στην παρούσα εργασία όμως το απόθεμα επιλέγεται σταθερό σε κάθε περίοδο κατά τρόπο τυχαίο βάση της διαίσθησης-εμπειρίας του εμπόρου. Για να δώσουμε μία σαφή εικόνα στον αναγνώστη του πόσο περίπλοκη είναι ή μπορεί να γίνει η κατανομή του r , αρκεί να δώσουμε μία εικόνα για την εξάρτηση αυτή.

Πριν προχωρήσουμε σε ανάλυση των τρόπων εκτίμησης της άγνωστης παραμέτρου καλό θα ήταν να ξεκαθαρίσουμε κάποιες βασικές κύριες γραμμές, περιορισμούς και οπτικές του προβλήματος που θα πρέπει να επισημανθούν για τον αναγνώστη.

Ένας σημαντικός περιορισμός κάτω από το πρίσμα του οποίου θα γίνει η προσέγγιση της άγνωστης παραμέτρου της κατανομής της ζήτησης είναι η επιλογή μικρού χρονικού ορίζοντα για τη συλλογή των δεδομένων της ζήτησης. Συγκεκριμένα ρεαλιστικότερες εμφανίζονται οι τιμές για $n = 10, 20, 50$ όπου οι αριθμοί αυτοί μπορεί να είναι μέρες, εβδομάδες ή και μήνες. Δεν επιλέχθηκε μικρότερος χρονικός ορίζοντας αν και ήταν αρχικά στην πρόθεση μας λόγω της ισχυρής πιθανότητας να εμφανιστούν όλες οι παρατηρήσεις να υπερβαίνουν το απόθεμα και αυτό επηρεάζει αφενός την ευστάθεια του αλγόριθμου προσομοίωσης αφετέρου δεν παίρνουμε καμία αντικειμενική εκτίμηση από κάποια εκτιμήτρια. Θα επιλεγούν δύο εκτιμητές μέγιστης πιθανοφάνειας. Ο 1^{ος} θα είναι χρησιμοποιώντας τη συνάρτηση πιθανοφάνειας του δείγματος ενώ ο 2^{ος} θα είναι χρησιμοποιώντας την πιθανοφάνεια του πλήθους r των τιμών ζήτησης που δεν υπερβαίνουν το απόθεμα Q .

Για την εύρεση του μέσου τετραγωνικού σφάλματος αλλά και την συνέπεια και αμεροληψία του κάθε εκτιμητή θα επιστρατευτούν προσομοιώσεις με το στατιστικό πακέτο R. Προσοχή επίσης θα δοθεί στη διασπορά του κάθε εκτιμητή δηλαδή στο Μέσο Τετραγωνικό Σφάλμα. Αυτό θα δώσει μία απάντηση στην επιλογή του καλύτερου εκτιμητή όταν η αμεροληψία και η συνέπεια θα είναι χαρακτηριστικό όλων.

3.2 Μέθοδος εκτιμητή μέγιστης πιθανοφάνειας με βάση το δείγμα (Ε.Μ.Π)

Αρχικά ασχολούμαστε με την περίπτωση που υπάρχει γνώση για το ύψος του αποθέματος Q . Οπότε γράφουμε τη συνάρτηση μέγιστης πιθανοφάνειας του δείγματος στην οποία υποθέτουμε ότι r από αυτές είναι μικρότερες του Q και $n-r$ το υπερβαίνουν, επομένως εάν συμβολίσουμε με $f_X(x_i|\theta)$ τη συνάρτηση πυκνότητας πιθανότητας μίας παρατήρησης του δείγματος τότε η συνάρτηση πιθανοφάνειας γράφεται, $L(x_1, x_2, \dots, x_n|\theta) = \prod_{i=1}^r P(X_i = x_i) \prod_{i=r+1}^n P(X_i > Q) =$

$$= \prod_{i=1}^r f_X(x_i) \prod_{i=r+1}^n (1 - F(Q)) = \prod_{i=1}^r \theta e^{-\theta x_i} \prod_{i=r+1}^n e^{-\theta Q} =$$

$$= \theta^r e^{-\theta \sum_{i=1}^r x_i} e^{-\theta(n-r)Q} = \theta^r \exp\left\{-\theta \sum_{i=1}^r x_i - \theta(n-r)Q\right\}$$

θέτοντας όπου F την αθροιστική συνάρτηση κατανομής των x_i . Βάζοντας στη σχέση λογαρίθμους γίνεται $\log(L(x_1, x_2, \dots, x_n|\theta)) = r \log \theta - \theta \sum_{i=1}^r x_i - (n-r)Q\theta$

οπότε παραγωγίζοντας ως προς θ , $\frac{d \log L(\tilde{x}, \theta)}{d\theta} = \frac{r}{\theta} - \sum_{i=1}^r x_i - (n-r)Q$

οπότε θέτοντας $\frac{d \log L(\tilde{x}, \theta)}{d\theta} = 0$ και λύνοντας ως προς θ παίρνουμε τον

εκτιμητή: $\hat{\theta}_{MLE} = \frac{r}{\sum_{i=1}^r x_i + (n-r)Q}$ (2.1) άρα για

τη μέση τιμή τον αντίστροφο του,

$$\frac{1}{\hat{\theta}_{MLE}} = \frac{y}{r} = \frac{\sum_{i=1}^r x_i + (n-r)Q}{r} \quad (2.2)$$

Στον εκτιμητή αυτό δεν είναι εύκολος ο υπολογισμός της αμεροληψίας λόγω του ότι το μέγεθος του μη-περικομμένου δείγματος είναι τυχαία μεταβλητή αλλά και της εξάρτησης την κατανομής του από το θ . Ο παραπάνω εκτιμητής υπάρχει μόνο δεσμεύοντας προς το ενδεχόμενο $r > 0$. Στην περίπτωση που $r = 0$ τότε η πιθανοφάνεια είναι $L(\tilde{x}, \theta) = e^{-n\theta Q}$ και αυξάνει συνεχώς όσο αυξάνει το μέγεθος του δείγματος n οπότε η εκτιμήτρια μέγιστης πιθανοφάνειας δεν υπάρχει.

Η πιθανότητα όμως του ενδεχομένου αυτού δηλαδή της τυχαίας μεταβλητής r να είναι μηδέν, μειώνεται καθώς αυξάνεται το μέγεθος του δείγματος n αλλά και με επιλογή μεγαλύτερου και ικανού αποθέματος.

Να σημειωθεί ότι εάν υπάρχει μόνο η γνώση της μέγιστης παρατήρησης

$X_{(r)} = \max\{X_i\}$ και όχι του αποθέματος τότε ακολουθώντας την ίδια διαδικασία

βρίσκουμε εκτιμητή $\hat{\theta} = \frac{r}{\sum_{i=1}^r X_i + (n-r)X_{(r)}}$. Το οποίο όμως στην περίπτωση του

προβλήματος του εφημεριδοπώλη είναι πρακτικά απίθανο. Ο έμπορος είναι εκείνος που ελέγχει ουσιαστικά το μέγεθος του αποθέματος και πρακτικά παίρνει αποφάσεις με βάση τις πωλήσεις που προκύπτουν από την απόφαση του αυτή.

Ωστόσο τον εκτιμητή αυτόν τον αναφέρουμε για δύο λόγους. Πρώτον γιατί δεδομένου r δίνει πολύ καλές εκτιμήσεις ως προς την αμεροληψία και την διασπορά στην προσέγγιση της παραμέτρου θ αλλά και τη μέσης τιμής της ζήτησης $\mu = 1/\theta$ και δεύτερον διότι σε life-testing μοντέλα που εξετάζεται η ζωή αντοχή ενός έμβιου όντος ή μηχανήματος είναι η εκτιμήτρια μέγιστης πιθανοφάνειας και έχει καλές ιδιότητες. Για την εκτιμήτρια αυτή στη σχέση (2.1) κάνουμε την εξής υπόθεση. Θέτουμε r_n το πλήθος των παρατηρήσεων που δεν υπερβαίνουν το απόθεμα σε ένα μέγεθος περιόδων n . Θέτουμε επίσης την ποσότητα q_n η οποία είναι τέτοια ώστε $r_n = [nq_n]$ δηλαδή εκφράζει το ποσοστό (%) των παρατηρήσεων που αναφέρθηκαν παραπάνω.

Παρατηρώντας την συμπεριφορά του q_n για «μεγάλα n » βλέπουμε ότι συγκλίνει στις εξής πιθανότητες.

$$q_n \rightarrow P(X < Q) = 1 - e^{-Q\theta} \text{ και } 1 - q_n \rightarrow 1 - P(X < Q) = P(X \geq Q) = e^{-Q\theta}$$

από τη στιγμή που η ζήτηση ακολουθεί την εκθετική κατανομή. Αντικαθιστώντας τώρα τα r_n στη (2.1)

$$\frac{1}{\hat{\theta}_{MLE}} = \frac{\sum_{i=1}^{r_n} X_i + (n - r_n)Q}{r_n} = \frac{\sum_{i=1}^{r_n} X_i}{r_n} + \frac{(n - r_n)Q}{r_n} =$$

$$= \bar{X}_{r_n} I(X_i < Q) + \left(\frac{\frac{n-r_n}{n}}{\frac{r_n}{n}} \right) Q = \bar{X}_{r_n} I(X_i < Q) + \left(\frac{1-q_n}{q_n} \right) Q \xrightarrow{n \rightarrow \infty}$$

$$\lim_{r \rightarrow \infty} \bar{X}_r I(X_i < Q) + \frac{P(X \geq Q)}{P(X < Q)} Q = \lim_{r \rightarrow \infty} \bar{X}_r I(X_i < Q) + \frac{e^{-Q\theta}}{1-e^{-Q\theta}} Q =$$

$$= \lim_{r \rightarrow \infty} \bar{X}_r I(X_i < Q) + \frac{Q}{e^{Q\theta} - 1}$$

Παίρνοντας μέση τιμή για την παραπάνω ποσότητα ,

$$\begin{aligned} E\left(\frac{1}{\hat{\theta}_{MLE}}\right) &\xrightarrow{n \rightarrow \infty} \lim_{r \rightarrow \infty} \frac{\sum_{i=1}^r E(X_i | X_i < Q)}{r} + E\left(\frac{Q}{e^{Q\theta} - 1}\right) \\ E\left(\frac{1}{\hat{\theta}_{MLE}}\right) &\xrightarrow{n \rightarrow \infty} \int_0^Q x \theta e^{-\theta x} dx + E\left(\frac{Q}{e^{Q\theta} - 1}\right) \\ E\left(\frac{1}{\hat{\theta}_{MLE}}\right) &\xrightarrow{n \rightarrow \infty} \frac{1}{\theta} (1 - Q\theta e^{-\theta Q} - e^{-\theta Q}) + \frac{Q}{e^{Q\theta} - 1} \\ E\left(\frac{1}{\hat{\theta}_{MLE}}\right) &\xrightarrow{n \rightarrow \infty} \frac{1}{\theta} \left(1 - \frac{Q\theta + 1}{e^{Q\theta}} + \frac{Q\theta}{e^{Q\theta} - 1}\right) = \frac{1}{\theta} \left(\frac{(e^{Q\theta} - 1)^2 + Q\theta}{e^{Q\theta} (e^{Q\theta} - 1)}\right) \quad (2.3) \end{aligned}$$

Παρατηρούμε ότι η εκτιμήτρια μέγιστης πιθανοφάνειας ασυμπτωτικά εκτιμά τη μέση τιμή επί μία συνάρτηση του θ και του αποθέματος Q . Έστω και ασυμπτωτικά λοιπόν δεν είναι αμερόληπτη. Επίσης βλέπουμε ότι η παράμετρος θ παραμένει στο 2^ο μέλος και στον εκθέτη και μας εμποδίζει στο να εξάγουμε συμπεράσματα για την αμεροληψία του εκτιμητή. Για αυτό το σημείο κάνουμε την υπόθεση ότι ο έμπορος επιλέγει αόριστα το απόθεμα Q έτσι ώστε να είναι $Q = \alpha \cdot \frac{1}{\theta}$. Δηλαδή εκφράζουμε την παραγγελία του αποθέματος σαν μια συνάρτηση της μέσης τιμής της ζήτησης. Στη συγκεκριμένη εργασία δε γίνονται πιο ισχυρές υποθέσεις σχετικά με το α ή το $\mu(\theta) = 1/\theta$. Δεν υπάρχει δηλαδή κάποια εκ των προτέρων γνώση τους. Οι υποθέσεις που κάνουμε είναι ότι δε γνωρίζουμε το α αυτό μιας και είναι άγνωστη η μέση τιμή του δείγματος και τα α και θ είναι σταθερά. Το α επιλέγεται έτσι ώστε $0 < \alpha < M$ όπου M μεγάλος αριθμός και είναι ενδεδειγμένο να επιλεγεί έτσι ώστε να

αποτραπεί το ενδεχόμενο όλες οι παρατηρήσεις να υπερβαίνουν το απόθεμα δηλαδή έτσι ώστε $P(r=0) \rightarrow 0$. Επομένως προσπαθούμε να δούμε ποια πολιτική είναι για τον έμπορο η καλύτερη δυνατή. Να κάνει παραγγελία ώστε να υπερτιμήσει τη μέση τιμή ζήτησης ανά περίοδο ώστε να μεγαλώνει ο συντελεστής α ή να υποτιμήσει ώστε να έχει τη λιγότερη δυνατή απώλεια των φθαρτών εμπορευμάτων του. Έτσι αντικαθιστώντας το γινόμενο θQ στη σχέση η

παραπάνω ισότητα γίνεται, $\frac{(e^\alpha - 1)^2 + \alpha}{e^\alpha (e^\alpha - 1)}$ ως συνάρτηση του α .

Η Αμεροληψία επιτυγχάνεται λοιπόν με κάποιες προϋποθέσεις. Μελετώντας τη συνάρτηση $\frac{(e^\alpha - 1)^2 + \alpha}{e^\alpha (e^\alpha - 1)}$ ως προς α βλέπουμε ότι έχει ελάχιστη τιμή το 0,843 για $\alpha = 0,82$ που συμβαίνει όταν υποτιμάται η μέση τιμή της κατανομής κατά 82%. Μετά το σημείο αυτό είναι αύξουσα και συγκλίνει στη μονάδα αρκετά γρήγορα. Αρκεί να δούμε ότι για $\alpha = 1$ παίρνει την τιμή 0,85 για $\alpha = 2$ είναι σχεδόν 0,9 ενώ για $\alpha = 3$ είναι 0,96. Η παραπάνω συνάρτηση συγκλίνει στη μονάδα και για μικρότερα α δηλαδή για μικρότερο απόθεμα. Δεν πρέπει όμως να ξεχνάμε ότι σημαντική πτώση του αποθέματος δύναται να προκαλέσει σημαντική πιθανότητα μη παρατηρούμενης πραγματικής ζήτησης στο σύνολο των n περιόδων, το οποίο με τη σειρά του δεν δίνει εκτίμηση μέγιστης πιθανοφάνειας. Με υπολογισμούς από την κατανομή του μεγέθους των παρατηρήσεων r βλέπουμε ότι για $\alpha = \theta Q = 0,82$ και για το Minimum μέγεθος χρονικού ορίζοντα που εξετάζουμε που είναι η πιθανότητα του ενδεχομένου υπολογίζεται από,

$$P(R=0) = \frac{n!}{0!n!} e^{-nQ\theta} = \exp\{-0,84 \cdot 5\} = e^{-4,1} = 1,4\% \text{ η οποία είναι αρκετά μεγάλη.}$$

Ενώ για $n=10$, $P(R=0) = \frac{n!}{0!n!} e^{-nQ\theta} = \exp\{-0,84 \cdot 10\} = e^{-8,4} = 0,0002182$ δηλαδή 2 /10.000. Αν θεωρήσουμε μια ασφαλή πιθανότητα μία αυτήν που είναι μικρότερη από 0,001 τότε για $n=5$ ασφαλές είναι ένα δείγμα που θα ισχύει ότι

$$P(R=0) < 0,001. \text{ Μετά από υπολογισμούς παίρνουμε ότι } \theta Q > 1.4 \text{ δηλαδή}$$

$Q > 1,4\mu(\theta)$ άρα για βεβαιότητα πρέπει να παραγγελθεί απόθεμα που υπερεκτιμά

την μέση τιμή της ζήτησης κατά 1,4 φορές παραπάνω τουλάχιστον. Ενδεικτικά σε αντίστοιχο δείγμα μεγέθους $n=10$ $\theta Q > 0,7$. Επομένως κατά γενικό κανόνα ειδικά για πλήθος περιόδων μεγαλύτερο του 5 καλό είναι να παραγγέλνεται ένα δείγμα αρκετά μεγαλύτερο από τις όποιες εκτιμήσεις του εμπόρου για τη μέση τιμή της ζήτησης. Φυσικά όλη αυτή η ιστορία δεν έχει κανένα νόημα εάν δεν υπάρχει εκ των προτέρων καμία γνώση για τη ζήτηση. Ακόμα και σε αυτές τις περιπτώσεις όμως είναι θεμιτό να παραγγέλνεται ένα απόθεμα αφενός ικανό από τη μία να μην δώσει πιθανότητα για ξεπούλημα (stock out) σε κάθε περίοδο παρατήρησης αφετέρου εάν υπάρχει η δυνατότητα πρέπει να γίνει παραγγελία μεγάλου αριθμού προϊόντων για να μπορούμε να έχουμε μία όσο το δυνατόν καλύτερη εκτίμηση για τη ζήτηση μέσω των εκτιμητριών. Μία επέκταση της παρούσας εργασίας θα ήταν να υποτεθούν κάποιες ασθενείς υποθέσεις για την παράμετρο θ της ζήτησης. Για παράδειγμα μπορεί ο έμπορος να έχει την εικόνα ότι η μέση τιμή της ζήτησης είναι μια τυχαία μεταβλητή $\mu = \frac{1}{\theta} \square Uniform(a, b)$ με γνωστά a, b . Ή εναλλακτικά θα μπορούσαμε να υποθέσουμε ότι η αναλογία του αποθέματος προς τη μέση τιμή της ζήτησης δηλαδή η ποσότητα που πάνω ονομάσαμε $\alpha = \frac{Q}{1/\theta}$ είναι επίσης τυχαία μεταβλητή τέτοια ώστε $\alpha \square Uniform(k, m)$ όπου οι παράμετροι k, m υποδηλώνουν μια γνώση από τον έμπορο για τα όρια στα οποία επιτρέπεται να κάνει την παραγγελία του αποθέματος ώστε να αποκλίνει σε συγκεκριμένο ποσοστό από τη μέση τιμή της ζήτησης. Η προσέγγιση αυτή κάνει ακόμα πιο ρεαλιστικό το σενάριο μιας και με ολική άγνοια του $1/\theta$ δεν μπορεί να δώσει μία σαφή στρατηγική παραγγελίας με σκοπό την καλύτερη δυνατή εκτίμηση της παραμέτρου. Εμείς στην παρούσα εργασία θα αποφύγουμε να μπούμε σε διαδικασίες επιλογής στρατηγικών παραγγελίας αποθέματος που να ελαχιστοποιεί το κόστος στις n περιόδους δίνοντας μια αρκετά βεβαίως καλή εκτίμηση της μέσης τιμής.

3.3 Μέθοδος εκτιμητή μέγιστης πιθανοφάνειας με βάση το πλήθος των παρατηρήσεων μικρότερων του αποθέματος (M.L.E «r»)

Η τυχαία μεταβλητή r που υποδηλώνει ουσιαστικά το πλήθος των τιμών της ζήτησης που παρατηρείται και δεν ξεπερνά το απόθεμα Q . Αυτή διαισθητικά μπορούμε να αντιληφθούμε ότι έχει Διωνυμική κατανομή δηλαδή $r \sim \text{Bin}(n, 1 - e^{-\theta Q})$ αφού ουσιαστικά έχουμε πραγματοποίηση n ανεξάρτητων δοκιμών Bernoulli, με την πιθανότητα να μην ξεπεραστεί σε καθεμία από αυτές το απόθεμα να είναι $P(X \leq Q) = F(Q) = 1 - e^{-\theta Q}$.

Επομένως παίρνοντας τη συνάρτηση πιθανοφάνειας

$$L(r|\theta, Q) = \frac{n!}{r!(n-r)!} (e^{-\theta Q})^{n-r} (1 - e^{-\theta Q})^r \quad \text{ως προς τις τυχαία μεταβλητή } r$$

βάζοντας λογαρίθμους και στη συνέχεια παραγωγίζοντας έχουμε,

$$\log L(r|\theta, Q) = \log \left(\frac{n!}{r!(n-r)!} \right) - \theta Q(n-r) + r \log(1 - e^{-\theta Q})$$

$$\frac{d \log L(r|\theta, Q)}{d\theta} = -Q(n-r) + \frac{Q r e^{-\theta Q}}{1 - e^{-\theta Q}}, \quad \text{οπότε θέτοντας } \frac{d \log L(r|\theta, Q)}{d\theta} = 0$$

$$\text{προκύπτει } Q(n-r) = \frac{Q r e^{-\theta Q}}{1 - e^{-\theta Q}}, \quad \text{η οποία μετασχηματίζεται με απλοποιήσεις σε}$$

$$\frac{e^{-\theta Q}}{1 - e^{-\theta Q}} = \frac{(n-r)}{r} \quad \text{άρα } e^{\theta Q} = 1 + \frac{r}{(n-r)}.$$

$$\text{Επομένως η Ε.Μ.Π για το } \theta \text{ ως συνάρτηση του } r \text{ είναι } \theta^* = -\frac{1}{Q} \log \left(1 - \frac{r}{n} \right).$$

Παρατηρούμε ότι η εκτιμητριά αυτή δεν εξαρτάται από τις παρατηρήσεις X_i . Για τη

$$\text{μέση τιμή παίρνουμε } 1/\hat{\theta} = -\frac{Q}{\log \left(1 - \frac{r}{n} \right)} \quad (2.4)$$

ΚΕΦΑΛΑΙΟ 4^ο

Μοντέλο Επιβίωσης και εκτιμητές παραμέτρου

4.1 Το πρόβλημα περικομμένων δεδομένων με γνωστό ποσοστό παρατηρούμενων τιμών. (life testing model)

Το πρόβλημα αυτό έχει μία ουσιώδη διαφορά από το πρόβλημα του εφημεριδοπώλη που περιγράψαμε στο 1^ο Κεφάλαιο. Εδώ πλέον σε ένα σύνολο n στοιχείων πραγματοποιούμε ένα τυχαίο πείραμα με πραγματοποίηση ή μη ενός γεγονότος. Εάν το γεγονός αυτό πραγματοποιηθεί σε r ακριβώς από τα στοιχεία διακόπτουμε το πείραμα. Αλλιώς συνεχίζουμε μέχρι τα r αυτά γεγονότα συμβούν. Για να δώσουμε μια πιο ξεκάθαρη εικόνα στον αναγνώστη για είδος αυτό του πειράματος παραθέτουμε το εξής παράδειγμα. Ας πούμε ότι μια κατασκευαστική εταιρία τουρμπινών για αεροπλάνα θέλει να δοκιμάσει την αντοχή τους κάτω υπό δύσκολες και δυσμενής συνθήκες. Οπότε επιλέγονται n από αυτές για δοκιμή και τοποθετούνται σε θαλάμους με υψηλές πιέσεις και δυσμενής συνθήκες. Οι τουρμπίνες αυτές έχουν πολύ υψηλό κόστος κατασκευής και η εταιρία θα προτιμούσε να μην τις αφήσει να λειτουργήσουν έως ότου να καταστραφούν και οι n από αυτές. Η επισκευή τους μετά από την καταστροφή τους στο θάλαμο δεν είναι δυνατή και λόγω οικονομικών περιορισμών δεν επιτρέπεται η συνέχιση της λειτουργία σε όλες από αυτές έως την αποτυχία. Οπότε επιλέγεται η διάρκεια του πειράματος να είναι τέτοια ώστε μόλις r από αυτές σταματήσουν να λειτουργούν από κάποια δυσλειτουργία το πείραμα λαμβάνει τέλος. Το μέγεθος r αυτών προκαθορίζεται από την ίδια την εταιρία αναλόγως του μπάτζετ της για το πείραμα. Αυτή η κατηγορία προβλημάτων έγκειται σε μια ευρύτερη κλάση που ασχολείται με τα πειράματα διάρκειας ζωής ενός μηχανήματος ή ενός έμβιου όντος. Βρίσκουν μεγάλη εφαρμογή σε δοκιμές φαρμάκων, αλλά και ελέγχους αντοχής μηχανημάτων-μηχανικών μερών.

Σύγκριση με το πρόβλημα του εφημεριδοπώλη

Παρακάτω αντιπαραβάλλουμε τα δύο προβλήματα του εφημεριδοπώλη το μοντέλο επιβίωσης σε δύο στήλες για να είναι ξεκάθαρες οι διαφορές.

Μοντέλο εφημεριδοπώλη

n συνολικά περίοδοι πώλησης του προϊόντος

Από τις n περιόδους θα έχω r το πλήθος παρατηρήσεων που να μην υπερβαίνουν το απόθεμα Q . Το r δεν είναι γνωστό σταθερό αλλά μια τυχαία μεταβλητή.

Οι $n-r$ παρατηρήσεις που υπερβαίνουν το απόθεμα δεν παρατηρούνται επακριβώς. Εμείς απλά γνωρίζουμε για αυτές ότι $X_i^* \geq Q$.

Μοντέλο επιβίωσης

n αντικείμενα-ζώα-άνθρωποι μελετώνται ως προς τη διάρκεια ζωής τους.

Από τα n αντικείμενα επιλέγω να σταματήσω τις παρατηρήσεις και το πείραμα όταν r από αυτά 'αποτύχουν' (έχουν μία προκαθορισμένη ιδιότητα). Το πλήθος r προκαθορίζεται από τον παρατηρητή.

Όταν σταματώ το πείραμα έχω διαθέσιμους τους χρόνους αποτυχίας των r αντικειμένων. Για το χρόνο ζωής των υπολοίπων γνωρίζω ότι $X_i^* \geq X_{(r)} = \max \{X_i\}$

Η σημαντικότερη διαφορά των δύο μοντέλων έγκειται λοιπόν στο ότι στο μοντέλο του εφημεριδοπώλη η τυχαία μεταβλητή πλέον r έχει άγνωστη κατανομή η οποία εξαρτάται από:

1. Τις παραμέτρους της κατανομής ζήτησης. (εδώ το θ)
2. Τον χρονικό ορίζοντα παρατηρήσεων
3. Την επιλογή του αποθέματος Q στο ξεκίνημα των περιόδων πώλησης.

Στη συγκεκριμένη εργασία οι δύο τελευταίες παράμετροι θεωρούνται σταθερές και προκαθορισμένες εξ' αρχής. Ο έμπορος όμως σε καμία περίπτωση δεν μπορεί να καθορίσει πόσες παρατηρήσεις θα υπερβαίνουν ή όχι το απόθεμα. Αλλά ακόμα και αν θεωρήσουμε την περίπτωση που καταγράφονται οι πωλήσεις μέχρι να παρατηρηθεί r από αυτές να μην είναι ίσες ή υπερβαίνουν το απόθεμα τότε τυχαία μεταβλητή γίνεται το πλήθος των περιόδων που θα χρειαστούν για να επιτευχθεί κάτι τέτοιο. Εμείς ασχολούμαστε εδώ με το διαφοροποιημένο πρόβλημα αυτό για να εξετάσουμε μέσω προσομοιώσεων εάν κάποια από τις εκτιμήτριες του συγκεκριμένου προσεγγίζουν την μέση τιμή της ζήτησης στο δικό μας.

Έτσι θα παραθέσουμε τέσσερις τρόπους υπολογισμού της μέσης τιμής ζήτησης όταν το R είναι πλέον διαθέσιμο και δεδομένο. Θα δείξουμε για κάθε τρόπο την αμεροληψία του κάθε εκτιμητή και την διασπορά του και στο επόμενο και τελευταίο κεφάλαιο θα παραθέσουμε τα αποτελέσματα των προσομοιώσεων.

4.1.1 Εκτιμητήρια μέγιστης πιθανοφάνειας με χρήση της μέγιστης παρατήρησης του δείγματος και του αποθέματος στο μοντέλο ελέγχου επιβίωσης.

Υπενθυμίζουμε ξανά την εκτιμητήρια με χρήση της μέγιστης παρατήρησης $X_{(r)}$. Η

$$\text{εκτιμητήρια λοιπόν είναι η } \frac{1}{\hat{\theta}_{MLE}} = \frac{\sum_{i=1}^r X_{(i)} + (n-r)X_{(r)}}{r} \quad (4.1)$$

θεωρώντας $X_{(i)}$ στην εκτιμητήρια τη διατεταγμένη παρατήρηση του δείγματος τάξης- i .

Εύκολα βλέπουμε από την κατανομή των X_i και θέτοντας το

$$y = \sum X_i + (n-r)X_{(r)} \quad \text{ότι} \quad 2\theta y \sim X_{2r}^2 \quad (4.2)$$

(X τετράγωνο κατανομή με $2r$ βαθμούς ελευθερίας) οπότε και παίρνουμε ότι

$$E(\hat{\theta}) = \frac{r\theta}{r-1} \quad \text{και} \quad V(\hat{\theta}) = \frac{\theta^2}{r-1} \quad \text{επομένως παρατηρούμε ότι ο } \hat{\theta} \text{ είναι}$$

ασυμπτωτικά αμερόληπτος και συνεπής εκτιμητής. Μία απόδειξη του παραπάνω αποτελέσματος παραθέτετε σε επόμενη ενότητα.

Για μικρά δείγματα τον πολλαπλασιάζουμε τη σταθερά $\frac{r-1}{r}$ ενώ για μεγάλα δείγματα όπου $r > 100$ έχουμε μια πολύ καλή εκτίμηση του.

Ο διορθωμένος ως προς τη μεροληψία εκτιμητής έχει νέα διασπορά

$$V(\hat{\theta}) = \frac{\theta^2}{r-2} \quad \text{Παρακάτω ακολουθεί η απόδειξη ότι} \quad E(\hat{\theta}) = \frac{r\theta}{(r-1)}.$$

Υπολογίζοντας τη συνάρτηση κατανομής του $2\theta y$ μέσω της ροπογεννήτριας συνάρτησης εξάγεται το συμπέρασμα (Halperin) ότι $2\theta y \sim X_{2r}^2$ όπου X_{2r}^2 (κατανομή τετράγωνο με $2r$ βαθμούς ελευθερίας) έχουμε ότι ,

$$\begin{aligned} E(2\theta y) &= 2r \Rightarrow E\left(\frac{1}{2\theta y}\right) = \int \frac{1}{y} \left(\frac{1}{2}\right)^r y^{r-1} \frac{e^{-y/2}}{\Gamma(r)} dy = \\ &= \frac{1}{2(r-1)} \int \left(\frac{1}{2}\right)^{\frac{2(r-1)}{2}} y^{\frac{2(r-1)}{2}-1} \frac{e^{-y/2}}{\Gamma\left(\frac{2(r-1)}{2}\right)} dy = \frac{1}{2(r-1)} \end{aligned}$$

Οπότε από αυτό το αποτέλεσμα καταλήγουμε στη μέση τιμή του στατιστικού, $\hat{\theta}$

$$E\left(\frac{r}{2\theta y}\right) = \frac{r}{2(r-1)} \Rightarrow E(\hat{\theta}) = \frac{r\theta}{r-1}.$$

Με τον ίδιο τρόπο υπολογίζεται η διασπορά του εκτιμητή.

Αντίστοιχα για τον υπολογισμό της μέσης τιμής και της διασποράς του δείγματος

$$\begin{aligned} E\left(\frac{1}{\hat{\theta}}\right) &= E\left(\frac{y}{r}\right) = \frac{1}{2\theta r} E(2\theta y) = \frac{1}{2\theta r} \cdot 2r = \frac{1}{\theta} \\ V\left(\frac{1}{\hat{\theta}}\right) &= V\left(\frac{y}{r}\right) = \frac{1}{4\theta^2 r^2} V(2\theta y) = \frac{1}{4\theta^2 r^2} \cdot 4r = \frac{1}{\theta^2 r} \end{aligned}$$

επομένως το $\frac{1}{\hat{\theta}}$ είναι αμερόληπτος εκτιμητής του $\frac{1}{\theta}$.

Διάστημα εμπιστοσύνης για την παράμετρο θ σε ένα περικομμένο δείγμα.

Εδώ θα παρουσιάσουμε ένα διάστημα εμπιστοσύνης για το δείγμα με τις περικομμένες παρατηρήσεις από το Q γνωρίζοντας την τιμή της τυχαίας μεταβλητής $R = r$. Ας θεωρήσουμε λοιπόν ξανά την τυχαία μεταβλητή y όπως ορίστηκε στη σχέση (4.2). Όπως αναφέρθηκε και στην ενότητα 4.1.1 δεδομένου σταθερού r το γινόμενο $2\theta y$ ακολουθεί X -τετράγωνο κατανομή με $2r$ βαθμούς ελευθερίας. Για ένα μονόπλευρο έλεγχο για το έχουμε ότι

$$P(2\theta y < X_{2r}^2(1-a)) = a \Rightarrow P\left(\theta < \frac{X_{2r}^2(1-a)}{2y}\right) = a \%$$

ενώ ένα διάστημα εμπιστοσύνης στατιστικής σημαντικότητας a είναι

$$P\left(X_{2r}^2(a/2) < 2\theta y < X_{2r}^2(1-a/2)\right) = a \Rightarrow P\left(\frac{X_{2r}^2(a/2)}{2y} < \theta < \frac{X_{2r}^2(1-a/2)}{2y}\right) = a.$$

Άρα λέμε ότι το θ βρίσκεται ανάμεσα σε $\frac{X_{2r}^2(a/2)}{2y} < \theta < \frac{X_{2r}^2(1-a/2)}{2y}$ με πιθανότητα $(1-a)\%$.

Μία Απόδειξη της κατανομής της τυχαίας μεταβλητής $2\theta y$

Εδώ αποδεικνύουμε τη σχέση (4.2).

Από τις διατεταγμένες παρατηρήσεις της εκθετικής γνωρίζουμε ότι

$$X_{(i)} = \frac{1}{\theta} \sum_{j=1}^i \frac{Z_j}{n-j+1} \quad j=1, 2, \dots, n \quad \text{όπου ανεξάρτητες και ισόνομες δηλαδή}$$

$Z_j \sim \text{Exp}(1)$ οπότε γράφοντας το y ως συνάρτηση των Z_j έχουμε ότι

$$y = \frac{1}{\theta} \sum_{i=1}^r \sum_{j=1}^i \frac{Z_j}{n-j+1} + \frac{(n-r)}{\theta} \sum_{j=1}^r \frac{Z_j}{n-j+1} = \frac{1}{\theta} \left(\sum_{i=1}^r \sum_{j=1}^i \frac{Z_j}{n-j+1} + (n-r) \sum_{j=1}^r \frac{Z_j}{n-j+1} \right)$$

Οπότε πολλαπλασιάζοντας με 2θ για απλοποίηση,

$$2\theta y = 2\theta \frac{1}{\theta} \left(\sum_{i=1}^r \sum_{j=1}^i \frac{Z_j}{n-j+1} + (n-r) \sum_{j=1}^r \frac{Z_j}{n-j+1} \right) = \left(\sum_{i=1}^r \sum_{j=1}^i \frac{Z_{1/2}}{n-j+1} + (n-r) \sum_{j=1}^r \frac{Z_{1/2}}{n-j+1} \right)$$

$$= \frac{Y_1}{n} + \left(\frac{Y_1}{n} + \frac{Y_2}{n-1} \right) + \dots + \left(\frac{Y_1}{n} + \frac{Y_2}{n-1} + \dots + \frac{Y_r}{n-r+1} \right) + (n-r) \left(\frac{Y_1}{n} + \frac{Y_2}{n-1} + \dots + \frac{Y_r}{n-r+1} \right)$$

$$= \left(\frac{rY_1}{n} + \frac{(r-1)Y_2}{n-1} + \dots + \frac{Y_r}{n-r+1} \right) + \left(\frac{(n-r)Y_1}{n} + \frac{(n-r)Y_2}{n-1} + \dots + \frac{(n-r)Y_r}{n-r+1} \right)$$

$$= \left(\frac{nY_1}{n} + \frac{(n-1)Y_2}{n-1} + \dots + \frac{(n-r+1)Y_r}{n-r+1} \right) = \sum_{i=1}^r Y_i \quad \text{όπου } Y_i = 2Z_i \text{ και επομένως είναι}$$

ανεξάρτητες και ισόνομες εκθετικές με παράμετρο $1/2$.

Επειδή όμως είναι γνωστό r το πλήθος εκθετικές με παράμετρο $1/2$ είναι μία τυχαία μεταβλητή με κατανομή χ^2 -τετράγωνο με $2r$ βαθμούς ελευθερίας.

4.2 Μέθοδος Προσομοιωτικών Εκτιμητών. (Simulating estimators)

Η ιδέα για τον εκτιμητή αυτόν προέκυψε από μία ιδιότητα των διατεταγμένων παρατηρήσεων που προέρχονται από την εκθετική κατανομή με σ.π.π

$$f_{X_i}(x) = \theta e^{-x \cdot \theta}. \text{ Η γενική ιδέα για την προσέγγιση με αυτόν τον τρόπο της}$$

παραμέτρου είναι να προσπαθήσουμε να προσεγγίσουμε τις εκλιπούσες παρατηρήσεις από το περικομμένο δείγμα και να κάνουμε μία νέα εκτίμηση για το θ . Αν μπορούσαμε να προσεγγίσουμε την παράμετρο από το υπάρχων δείγμα και να προσομοιώσουμε τις άγνωστες σε εμάς τιμές της ζήτησης που υπερβαίνουν το Q και στη συνέχεια να επαναυπολογίσουμε την άγνωστη παράμετρο

χρησιμοποιώντας όλο το δείγμα δείχνει μια προσέγγιση η οποία δεν φαίνεται ότι είναι η καλύτερη δυνατή αλλά θα μπορούσε να έχει ίσως και περεταίρω βελτιώσεις και παραλλαγές. Η προσομοίωση της παραμέτρου θα μπορούσε να γίνει από το παρατηρούμενο δείγμα εκτός από τον ΕΜΠ τρόπο και με κάποιον άλλο όπως ο γραμμικός εκτιμητής, ο Μπεϋζιανός εκτιμητής αλλά ακόμα και παίρνοντας ένα γραμμικό συνδυασμό από αμερόληπτους τρόπους.

Εκτιμώντας τις περικομμένες παρατηρήσεις χρησιμοποιώντας τον Ε.Μ.Π από το παρατηρούμενο δείγμα.

Μελετώντας το έργο των B. C. Arnold, N.Balakrishnan και H.N.Nagaraja είναι γνωστό ότι εάν οι διατεταγμένες παρατηρήσεις από το δείγμα εκθετικών τότε η από κοινού κατανομή των $X_{(1)} < X_{(2)} < \dots < X_{(n)}$, είναι η

$$f_{X_{(1)}, X_{(2)}, \dots, X_{(n)}}(x_1, x_2, \dots, x_n) = n! e^{-\sum x_i} \quad \text{με} \quad 0 \leq x_1 < x_2 < \dots < x_n < \infty$$

οπότε θεωρούμε τον εξής μετασχηματισμό

$$Z_1 = nX_{(1)}, \quad Z_2 = (n-1)(X_{(2)} - X_{(1)}), \dots, Z_n = X_{(n)} - X_{(n-1)} \quad (4.3)$$

ή αλλιώς τον ισοδύναμο μετασχηματισμό

$$X_{(1)} = \frac{Z_1}{n}, \quad X_2 = \frac{Z_1}{n} + \frac{Z_2}{n-1}, \dots, X_{(n)} = \frac{Z_1}{n} + \frac{Z_2}{n-1} + \dots + \frac{Z_n}{1} \quad (4.4)$$

Γνωρίζοντας ότι η Ιακωβιανή ορίζουσα του μετασχηματισμού είναι $1/n!$ η από

κοινού των $Z_1, Z_2, \dots, Z_n$ είναι $f_{Z_1, Z_2, \dots, Z_n}(z_1, z_2, \dots, z_n) = e^{-\sum z_i}$

και άρα από το παραγοντικό θεώρημα είναι εμφανές ότι οι παραπάνω μεταβλητές Z_i είναι ανεξάρτητες και ομοιόμορφα κατανεμημένες μεταβλητές εκθετικές με παράμετρο 1. Έτσι λοιπόν παρουσιάζουμε το επόμενο θεώρημα το οποίο χρησιμοποιήθηκε για την δημιουργία του προσομοιωτικού αλγορίθμου παραγωγής των περικομμένων παρατηρήσεων .

Θεώρημα

Έστω $X_{(1)}, X_{(2)}, \dots, X_{(n)}$ οι διατεταγμένες παρατηρήσεις προερχόμενες από την τυποποιημένη εκθετική κατανομή τότε οι τυχαίες μεταβλητές $Z_1, Z_2, \dots, Z_n$ όπως ορίστηκαν παραπάνω στην (4.3) δηλαδή

$$Z_i = (n-i+1)(X_{(i)} - X_{(i-1)}) \quad i=1, 2, \dots, n \quad \text{με} \quad X_{(0)} = 0$$

Είναι στατιστικά ανεξάρτητες και ακολουθούν την τυποποιημένη εκθετική κατανομή. Επιπλέον έχουμε από την (4.4) το αποτέλεσμα ότι

$$X_{(i)} = \frac{1}{\theta} \sum_{j=1}^i \frac{Z_j}{n-j+1} \quad j=1, 2, \dots, n \quad \text{που ουσιαστικά περιγράφει το } i\text{-οστό order}$$

statistic σαν γραμμικό συνδυασμό i ανεξάρτητων τυχαίων μεταβλητών που ακολουθούν την τυπ. εκθετική κατανομή. Η μέθοδος αυτή εφαρμόζεται ακόμα και όταν οι τυχαίες μεταβλητές ακολουθούν εκθετική κατανομή με παράμετρο θ . Με αυτόν τον τρόπο η μόνη αλλαγή στον τύπο που δίνει τα $X_{(i)}$ συναρτήσει των Z_i είναι

$$X_{(i)} \stackrel{D}{=} \frac{1}{\theta} \left(\sum_{j=1}^i \frac{Z_j}{n-j+1} \right) \quad i=1, 2, \dots, n$$

Σαν συμπέρασμα δεν είναι δύσκολο να δούμε ότι αφαιρώντας δύο διαδοχικές

διατεταγμένες παρατηρήσεις παίρνουμε $X_{(i+1)} - X_{(i)} = \frac{Z}{(n-i)\theta} \sim \text{Exp}\{(n-i)\theta\}$

Επομένως επιστρέφοντας στο πρόβλημα μας για τις άγνωστες τιμές των περικομμένων παρατηρήσεων $X_{(r+1)}, X_{(r+1)}, \dots, X_{(n)}$ προσεγγίζονται ως εξής:

$$X_{(r+1)}^* = X_{(r)} + Z_{(n-r)}\hat{\theta}$$

$$X_{(r+2)}^* = X_{(r+1)}^* + Z_{(n-r-1)}\hat{\theta}$$

...

$$X_{(n)}^* = X_{(n-1)}^* + Z_{\hat{\theta}}$$

Αντικαθιστώντας σε κάθε εξίσωση στην προσομοιωμένη μεταβλητή $X_{(i)}^*$ την τιμή της από την προηγούμενη εξίσωση το σύστημα γίνεται απλούστερο,

$$X_{(r+1)}^* = X_{(r)} + Z_{(n-r)}\hat{\theta}$$

$$X_{(r+2)}^* = X_{(r)} + Z_{(n-r)}\hat{\theta} + Z_{(n-r-1)}\hat{\theta}$$

$$X_{(r+3)}^* = X_{(r)} + Z_{(n-r)}\hat{\theta} + Z_{(n-r-1)}\hat{\theta} + Z_{(n-r-2)}\hat{\theta}$$

...

$$X_{(n)}^* = X_{(r)} + Z_{(n-r)}\hat{\theta} + \dots + Z_{\hat{\theta}} = X_{(r)} + \sum_{i=0}^{n-r-1} Z_{(n-r-i)}\hat{\theta} \quad (4.5)$$

Οπότε κάθε τυχαία προσομοιωμένη τυχαία μεταβλητή γράφεται σαν συνάρτηση της μέγιστης παρατηρούμενης τιμής με το άθροισμα των εκθετικών «αλμάτων» για την προσέγγιση του ζητούμενου. Να σημειωθεί ότι ο δείκτης στην εκθετική τυχαία μεταβλητή Z υποδηλώνει την τιμή της παραμέτρου αυτής. Στο συγκεκριμένο παράδειγμα ο επιλεγμένος εκτιμητής του θ για την προσομοίωση των

παρατηρήσεων είναι ο εκτιμητής μέγιστης πιθανοφάνειας $\hat{\theta}_{MLE} = \frac{r-1}{y}$

διορθωμένος για να είναι αμερόληπτος. Επιλέχθηκε αυτός ο εκτιμητής γιατί εμφανίζει αρκετά καλές ιδιότητες εκτός της αμεροληψίας και της ελάχιστης διασποράς. Πρέπει να αναφερθεί επίσης μια διόρθωση που κάνουμε στον αλγόριθμο προσομοίωσης η οποία είναι η εξής: Εάν προσομοιωθεί η εκτίμηση της πρώτης μη παρατηρούμενης ζήτησης $X_{(r+1)}^*$ και αυτή συνεχίζει να μην υπερβαίνει το απόθεμα τότε θέτουμε $X_{(r+1)}^* = Q$ αφού είναι μια πληροφορία την οποία έχουμε και πρέπει να χρησιμοποιήσουμε για να ελαχιστοποιήσουμε το σφάλμα.

Ο νέος εκτιμητής λοιπόν τώρα για τη μέση τιμή $1/\theta$ είναι ο

$$1/\tilde{\theta} = \frac{\sum_{i=1}^r X_{(i)} + \sum_{i=r+1}^n X_{(i)}^*}{n} \quad (4.6)$$

4.2.2 Αμερόληψία Νέου Προσομοιωτικού Εκτιμητή

Το πρώτο σημαντικό ερώτημα είναι αν με αυτό τον τρόπο έχουμε ένα συνεπή και προπαντός αμερόληπτο εκτιμητή για τη μέση τιμή του αποθέματος.

$$\begin{aligned} E\left(\frac{1}{\tilde{\theta}}\right) &= E\left(\frac{\sum_{i=1}^r X_{(i)} + \sum_{i=r+1}^n X_{(i)}^*}{n}\right) = E\left(\frac{\sum_{i=1}^r X_{(i)} + (n-r)X_{(r)} + \sum_{i=0}^{n-r-1} (n-r-i)Z_{(n-r-i)\hat{\theta}}}{n}\right) \\ &= \frac{r}{n} E\left(\frac{1}{\hat{\theta}}\right) + E\left(\frac{\sum_{i=0}^{n-r-1} (n-r-i)Z_{(n-r-i)\hat{\theta}}}{n}\right) = \frac{r}{n} \frac{1}{\theta} + \frac{(n-r)}{n} E(Z_{\hat{\theta}}) \end{aligned} \quad (4.7)$$

Όμως $E(Z_{\hat{\theta}}) = E\left(Z_{(r-1)/Y}\right) = E\left(E\left(Z_{(r-1)/Y} | Y\right)\right)$ και από πριν είναι γνωστό

ότι $Z_{(r-1)/Y} | Y \sim \text{Exp}\left((r-1)/Y\right)$ και ότι $2\theta Y \sim X^2(2r)$ άρα δεσμεύοντας

$$E\left(E\left(Z_{(r-1)/Y} | Y\right)\right) = E\left(\frac{Y}{r-1}\right) = \frac{1}{2\theta(r-1)} E(2\theta Y) = \frac{1}{2\theta(r-1)} \cdot 2r = \frac{r}{\theta(r-1)}$$

επομένως αντικαθιστώντας στην (4.7)

$$E\left(\frac{1}{\tilde{\theta}}\right) = \frac{r}{n} \frac{1}{\theta} + \frac{(n-r)}{n} \frac{r}{\theta(r-1)} = \frac{1}{\theta} \left(\frac{r}{n} + \frac{(n-r)r}{n(r-1)} \right) = \frac{1}{\theta} \left(\frac{r^2 - r + nr - r^2}{n(r-1)} \right) = \frac{1}{\theta} \frac{r(n-1)}{n(r-1)}$$

Άρα ο νέος αμερόληπτος εκτιμητής είναι ο $\frac{1}{\tilde{\theta}_{New}} = \frac{n(r-1)}{\tilde{\theta}r(n-1)}$

Με τον περιορισμό φυσικά ύπαρξης τουλάχιστον δύο παρατηρήσεων αλλιώς στην εκτιμούμε με την (4.6) ενώ αν δεν υπάρχει περικομμένη παρατήρηση δεν υπάρχει ανάγκη για προσομοίωση παρατήρησης οπότε δεν χρησιμοποιείται αυτός ο εκτιμητής.

Παραπάνω στο σύστημα εξισώσεων (4.5) χρησιμοποιήσαμε ότι,

$\sum_{i=r+1}^n X_{(i)}^* = (n-r)X_{(r)} + \sum_{i=1}^{n-r} (n-r-i)Z_{(n-r-i)\hat{\theta}}$ ως αποτέλεσμα διαδοχικών αντικαταστάσεων και επίσης την $(n-r-i)Z_{(n-r-i)\hat{\theta}} \square Z_{\hat{\theta}}$ γνωστή από τις ιδιότητες της εκθετικής κατανομής που αναφέρθηκε στην αρχή του κεφαλαίου 4. Επομένως και ο εκτιμητής $1/\tilde{\theta}_{NEW}$ είναι επίσης αμερόληπτος εκτός της περιπτώσεως μοναδικής παρατήρησης.

4.2.3 Διασπορά Νέου Εκτιμητή

Τώρα για τη διασπορά του εκτιμητή αυτού με τον περιορισμό και πάλι ότι έχουμε διαθέσιμες τουλάχιστον δύο παρατηρήσεις μικρότερες από Q στο δείγμα,

$$\begin{aligned} Var\left(\frac{1}{\theta_{NEW}}\right) &= Var\left(\frac{n(r-1)}{\tilde{\theta}r(n-1)}\right) = \frac{n^2(r-1)^2}{r^2(n-1)^2} Var\left(\frac{1}{\tilde{\theta}}\right) = \frac{n^2(r-1)^2}{(n-1)^2 r^2} Var\left(\frac{\sum_{i=1}^r X_{(i)} + \sum_{i=r+1}^n X_{(i)}^*}{n}\right) = \\ &= \frac{(r-1)^2}{(n-1)^2} Var\left(\frac{\sum_{i=1}^r X_{(i)} + (n-r)X_{(r)} + \sum_{i=0}^{n-r-1} (n-r-i)Z_{(n-r-i)\hat{\theta}}}{r}\right) = \frac{(r-1)^2}{(n-1)^2} Var\left(\frac{1}{\hat{\theta}} + \frac{(n-r)}{r} Z_{\hat{\theta}}\right) = \end{aligned}$$

Για τον υπολογισμό της διασποράς στην (4.7) αναλύουμε στις επιμέρους:

$$\begin{aligned} Var\left(\frac{y}{r} + \frac{(n-r)}{r} Z_{\frac{r-1}{y}}\right) &= Var\left(\frac{y}{r}\right) + \frac{(n-r)^2}{r^2} Var\left(Z_{\frac{r-1}{y}}\right) + 2Cov\left(\frac{y}{r}, \frac{(n-r)}{r} Z_{\frac{r-1}{y}}\right) = \\ &= \frac{1}{r\theta^2} + \frac{(n-r)^2}{4\theta^2 r^2 (r-1)^2} Var\left(Z_{\frac{1}{2\theta y}}\right) + 2Cov\left(\frac{y}{r}, \frac{(n-r)}{r} Z_{\frac{r-1}{y}}\right) \\ &= \frac{(r-1)^2}{(n-1)^2} Var\left(\frac{y}{r} + \frac{(n-r)}{r} Z_{\frac{r-1}{y}}\right) \quad (4.8) \end{aligned}$$

Θέτοντας για ευκολία W την μεταβλητή $2\theta y$, η διασπορά υπολογίζεται

$$\begin{aligned}
\text{Var}\left(Z_{\frac{1}{W}}\right) &= E_W\left(Z_{\frac{1}{W}}^2\right) - E_W\left(Z_{\frac{1}{W}}\right)^2 = E\left(E_W\left(Z_{\frac{1}{W}}^2 \mid W = w\right)\right) - E\left(E_W\left(Z_{\frac{1}{W}} \mid W = w\right)\right)^2 = \\
&= 2E(W^2) - E(W)^2 = 2(E(W^2) - E(W)^2) + E(W)^2 = 2\text{Var}(W) + E(W)^2 = \\
&= 8r + 4r^2 = 4r(r+2)
\end{aligned}$$

$$\begin{aligned}
\text{Cov}\left(\frac{y}{r}, \frac{(n-r)}{r} Z_{\frac{r-1}{y}}\right) &= E\left(\frac{y}{r} \frac{(n-r)}{r} Z_{\frac{r-1}{y}}\right) - E\left(\frac{y}{r}\right) E\left(\frac{(n-r)}{r} Z_{\frac{r-1}{y}}\right) = \\
&= \frac{(n-r)}{4\theta^2 r^2 (r-1)} E\left(2\theta y Z_{\frac{1}{2\theta y}}\right) - \frac{1}{2\theta^2} \frac{(n-r)}{r^2 (r-1)} E\left(Z_{\frac{1}{2\theta y}}\right) = \\
&= \frac{(n-r)}{4\theta^2 r^2 (r-1)} \left(E\left(E\left(W Z_{\frac{1}{W}} \mid W = w\right)\right) - 2E\left(E\left(Z_{\frac{1}{W}} \mid W = w\right)\right)\right) = \\
&= \frac{(n-r)}{4\theta^2 r^2 (r-1)} (E(W^2) - 2E(W)) = \frac{(n-r)}{4\theta^2 r^2 (r-1)} (\text{Var}(W) + E(W)^2 - 2E(W)) \\
&= \frac{(n-r)}{4\theta^2 r^2 (r-1)} (4r + 4r^2 - 4r) = \frac{(n-r)}{\theta^2 (r-1)}
\end{aligned}$$

Ουσιαστικά εξ' αρχής παρατηρούμε ότι έχουμε $\text{Cov}\left(\frac{y}{r}, \frac{(n-r)}{r} Z_{\frac{r-1}{y}}\right) > 0$ επομένως

η διασπορά αυτού του εκτιμητή υπερβαίνει αυτή του εκτιμητή μέγιστης πιθανοφάνειας και αυτή μεγαλώνει όσο μεγαλώνει και το ποσοστό των περικομμένων παρατηρήσεων στο δείγμα. Σημειώνεται ότι όλες οι ιδιότητες αμεροληψίας και συνέπειας της παραπάνω εκτιμήτριας εμείς θα προσομοιώσουμε στο μοντέλο του εφημεριδοπώλη τη ζήτηση χρησιμοποιώντας την (3.1) ως εκτιμήτρια για το θ αντί της (4.1) που χρησιμοποιήθηκε εδώ.

Στην εργασία όμως αυτή όταν προσομοιώσουμε τις περικομμένες παρατηρήσεις ζήτησης μιας και πραγματευόμαστε το μοντέλο του εφημεριδοπώλη δεν χρησιμοποιούμε τον εκτιμητή της σχέσης για το θ αλλά τον

$$\hat{\theta}_{MLE} = \frac{r}{\sum_{i=1}^r X_i + (n-r)Q} \quad (4.9)$$

4.3 Μέθοδος Βέλτιστων Γραμμικών Εκτιμητών (Best linear estimators “BLUEs”)

4.3.1 Παρουσίαση των Γραμμικών εκτιμητών

Μια άλλη πολύ ενδιαφέρουσα μέθοδος έρευνας της άγνωστης παραμέτρου είναι μέσω του προσδιορισμού των βέλτιστων order statistics από ένα δεξιά περικομμένο δείγμα. Αυτή η μέθοδος αναπτύχθηκε από τους Saleh [12],[13] Sarhan & Greenberg[14] αρκετές δεκαετίες παλαιότερα και επιτεύχθηκε με μεθόδους βελτιστοποίησης. Εξετάστηκαν δύο κατανομές η μονοπαραμετρική εκθετική και η εκθετική με δύο παραμέτρους θέσης και κλίμακας.

Εμείς θα εξετάσουμε την 1^η περίπτωση δηλαδή κατανομή τυχαίων μεταβλητών X_i με σ.π.π $f_x(x) = \theta \exp\{-x\theta\}$.

Θεωρούμε το διατεταγμένο δείγμα $X_{(n_1)}, X_{(n_2)}, \dots, X_{(n_k)}$ με $n_1, n_2, \dots, n_k$ τους βαθμούς διάταξης οι οποίοι που ικανοποιούν την ανίσωση $1 \leq n_1 \leq n_2 \leq \dots \leq n_k \leq n$. Οι φυσικοί αριθμοί αυτοί αριθμοί καθορίζονται από σταθερούς πραγματικούς αριθμούς $p_1, p_2, \dots, p_k$ οι οποίοι με τη σειρά τους ικανοποιούν την ανίσωση $0 < p_1 < p_2 < \dots < p_k < 1$ και επίσης $n_i = [np_i] + 1$ όπου με τις αγκύλες θεωρούμε το ακέραιο μέρος του αριθμού μέσα σε αυτές που δεν ξεπερνά το np_i . Οι Ogawa(1960) και Kulldorff(1963) έδειξαν ότι οι ασυμπτωτικός βέλτιστος γραμμικός εκτιμητής βάσει των διατεταγμένων παρατηρήσεων του δείγματος είναι $\tilde{\theta}^{-1} = \sum_{i=1}^k b_i X_{(n_i)}$ όπου $b_i = Q_k^{-1} \left\{ (t_i - t_{i-1}) / (e^{t_i} - e^{t_{i-1}}) - (t_{i+1} - t_i) / (e^{t_{i+1}} - e^{t_i}) \right\}$ $i = 1, 2, 3, \dots, k-1$ και $b_k = Q_k^{-1} \left\{ (t_k - t_{k-1}) / (e^{t_k} - e^{t_{k-1}}) \right\}$ και η ποσότητα Q_k υπολογίζεται από $Q_k = \sum_{i=2}^k (t_i - t_{i-1})^2 / (e^{t_i} - e^{t_{i-1}})$ και τα $t_i, t_i = \ln(1 - p_i)^{-1}$ είναι τα ποσοστημόρια της εκθετικής κατανομής που αντιστοιχούν στα $p_i, i = 1, 2, 3, \dots, k$. Η διασπορά του εκτιμητή του θ^{-1} είναι $V(\tilde{\theta}^{-1}) = Q_k^{-1} \theta^{-2} / n$.

Ο δείκτης του Q_k υποδηλώνει τη διάσταση του. Να σημειωθεί ότι τα παραπάνω αποτελέσματα είναι βασισμένα στην ασυμπτωτική κατανομή των k -ποσοστημορίων. (Mosteller (1946) ; Ogawa 1956) και στην εφαρμογή του -περικομμένου δείγματος. Στη 2^η περίπτωση από το θεώρημα 3 που αναφέρεται παρακάτω το μέγιστο συναντάται στο όριο του συνόλου δηλαδή $t_k^0 = \ln(1 - \beta)^{-1}$. Επομένως το βέλτιστο ποσοστημόριο p_k είναι το ίδιο το β και ο βαθμός της παρατήρησης είναι $r = [n\beta] + 1$ δηλαδή η μεγαλύτερη τιμή του παρατηρούμενου δείγματος $X_{(r)}$. Για να εντοπίσουμε τις υπόλοιπες $k - 1$ βέλτιστες διατεταγμένες παρατηρήσεις μεγιστοποιείται η $Q_k(t_1, t_2, \dots, t_k)$ κρατώντας θεωρήματος των Gauss-Markoff.

Για δοσμένο k με $k < r$ όπου r το μέγεθος του παρατηρούμενου δείγματος θεωρούμε το διατεταγμένο δείγμα $X_{(n_1)}, X_{(n_2)}, \dots, X_{(n_k)}$ με του βαθμούς $n_1, n_2, \dots, n_k$ που καθορίζονται από αντίστοιχα p_i με $0 < p_1 \leq p_2 \leq \dots, p_k \leq \beta$ που καθορίζονται από $n_i = [np_i] + 1$ με $p_0 = 0, p_{k+1} = 1$. Για να εκτιμήσουμε την παράμετρο κλίμακας θέτουμε $t_i = \ln(1 - p_i)^{-1}$. Τα $t_1, t_2, \dots, t_k$ ικανοποιούν την ανισότητα $0 < t_1 < t_2 < \dots < t_k = \ln(1 - \beta)^{-1}$. Βάση των διατεταγμένων παρατηρήσεων $X_{(n_1)}, X_{(n_2)}, \dots, X_{(n_k)}$ οι ασυμπτωτικοί BLUE εκτιμητές του δίνονται από τις σχέσεις (1.9) , (1.10). Για την έρευνηση των βέλτιστων διατεταγμένων παρατηρήσεων από το δείγμα αρκεί να ελαχιστοποιήσουμε τη διασπορά του εκτιμητή ή αλλιώς να μεγιστοποιήσουμε το $Q_k(t_1, t_2, \dots, t_k)$ ως προς τα t_i η οποία δίνεται από $Q_k = \sum_{i=2}^k (t_i - t_{i-1})^2 / (e^{t_i} - e^{t_{i-1}})$ με τον περιορισμό όπως αναφέρθηκε και πιο πάνω ότι $0 < t_1 < t_2 < \dots < t_k = \ln(1 - \beta)^{-1}$ $t_0 = 0$.

Από το Θεώρημα 1 η συνάρτηση Q_k έχει μοναδικό μέγιστο στο διάστημα $0 \leq t_1 \leq t_2 \leq \dots \leq t_k \leq \infty$. Έστω το μέγιστο να βρίσκεται στο σημείο $(t_1^0, \dots, t_{k-1}^0, t_k^0)$ Για το μέγιστο της Q_k στο $0 < t_1 < t_2 < \dots < t_k \leq \ln(1 - \beta)^{-1} < \infty$ θεωρούνται δύο περιπτώσεις: i) το ποσοστό του censoring $1 - \beta$ στο δείγμα είναι τέτοιο ώστε

$t_k^0 \leq \ln(1-\beta)^{-1}$ ii) το ποσοστό του censoring $1-\beta$ στο δείγμα είναι τέτοιο ώστε

$t_k^0 > \ln(1-\beta)^{-1}$. Στην 1^η περίπτωση η απάντηση είναι ότι το μέγιστο είναι το

$(t_1^0, \dots, t_{k-1}^0, t_k^0)$ και τα t_i είναι ίδια με αυτά του μη-περικομμένου δείγματος.

Στη 2^η περίπτωση από το Θεώρημα 3 το μέγιστο αντιστοιχεί στο σύνολο

$$t_k = \ln(1-\beta)^{-1}.$$

Έπεται ότι το βέλτιστη τιμή του p_k είναι το β η οποία αντιστοιχεί στην τιμή της παρατήρησης $r = [n\beta] + 1$, για αυτό το λόγο περιλαμβάνεται η μέγιστη παρατήρηση από το δείγμα, $X_{(r)}$. Για να βρούμε τα υπόλοιπες $k-1$ τιμές των διατεταγμένων

παρατηρήσεων μεγιστοποιούμε την $Q_k(t_1, t_2, \dots, t_k)$ κρατώντας το t_k σταθερό και

ίσο με $\ln(1-\beta)^{-1}$. Από το Θεώρημα 3 το μέγιστο είναι μοναδικό και ικανοποιεί το

σύστημα εξισώσεων $\tau_{i+1} + \tau_i - 2t_i = 0$, $i = 1, 2, 3, \dots, k$. Έστω οι ποσότητες $t_1^*, \dots, t_{k-1}^*$ η

λύση του παραπάνω συστήματος τότε τα $p_1^*, p_2^*, \dots, p_{k-1}^*$ υπολογίζονται από τις

εξισώσεις $p_i^* = 1 - e^{-t_i^*}$ επομένως οι βαθμοί των λοιπών $k-1$ διατεταγμένων

παρατηρήσεων δίνονται από τη σχέση $n_i^* = [np_i^*] + 1$ όπως έχει προαναφερθεί.

Έτσι το δείγμα αποτελείται από τις παρατηρήσεις $X_{(n_1)}, X_{(n_2)}, \dots, X_{(r)}$. Η εκτίμηση του

θ^{-1} δίνεται από τον τύπο $\tilde{\theta}^{-1} = \sum_{i=1}^{k-1} b_i X_{(n_i)} + b_k X_{(r)}$. Με τα b_i, b_k να δίνονται από

$$\text{τους τύπους } b_i^* = Q_k^{-1} \left\{ (t_i^* - t_{i-1}^*) / (e^{t_i^*} - e^{t_{i-1}^*}) - (t_{i+1}^* - t_i^*) / (e^{t_{i+1}^*} - e^{t_i^*}) \right\}$$

$$b_k = Q_k^{-1} \left\{ (t_k - t_{k-1}) / (e^{t_k} - e^{t_{k-1}}) \right\} = Q_k^{-1} (1-\beta) \left\{ (\ln(1-\beta)^{-1} - t_{k-1}) / (1 - (1-\beta)e^{t_{k-1}}) \right\}$$

4.3.2 Σημαντικά Θεωρήματα

Θεώρημα 1

Το σύστημα των εξισώσεων $\tau_{i+1} + \tau_i - 2t_i = 0$ $i = 1, 2, \dots, (k-1)$ έχει μία μόνο λύση

και μόνο και αυτή συμπίπτει με το μέγιστο της συνάρτησης $Q_k(t_1, t_2, \dots, t_k)$.

Για να αποδείξουμε το θεώρημα μελετάται η επιρροή της αντικατάστασης

οποιοδήποτε βέλτιστου t_i που ικανοποιεί τις παραπάνω εξισώσεις. Θεωρούμε

λοιπόν μια σταθερή τιμή για το $t_1 > 0$. Αυτή με τη σειρά της καθορίζει $k - 2$ επιτυχημένα t_i δηλαδή τα $t_2, \dots, t_{k-1}$ τα οποία υπολογίζονται σταδιακά από τις εξισώσεις $2t_i = \tau_i - \tau_{i+1}$. Παίρνοντας διαφορικά και στα δύο μέλη βρίσκουμε ότι,

$$\left[e^{-t_{i+1}} (t_{i+1} - \tau_{i+1}) / (e^{-t_i} - e^{-t_{i+1}}) \right] d(t_{i+1} - t_i) = \left[e^{-t_{i+1}} (\tau_i - t_{i-1}) / (e^{-t_{i-1}} - e^{-t_i}) \right] d(t_i - t_{i-1}) \quad (4.10)$$

όπου d ο διαφορικός τελεστής. Από την παραγωγή προκύπτει ότι αν $d(t_i - t_{i-1}) > 0$ τότε $d(t_{i+1} - t_i) > 0$ και αφού όλοι οι άλλοι οι όροι είναι θετικοί.

Τώρα μπορούμε να δείξουμε ότι για οποιαδήποτε αύξηση στο t_1 , $d(t_2 - t_1) > 0$ τότε από επαγωγή μπορούμε να δείξουμε ότι κάθε απειροελάχιστη αύξηση στο t_1 αυξάνει τις τιμές των επιτυχημένων t . Πρώτα δείχνουμε ότι αν $dt_1 > 0$ έχει σαν συνέπεια ότι $d(t_2 - t_1) > 0$. Θεωρούμε την $2t_1 = \tau_1 - \tau_2$ οπότε παίρνοντας διαφορικό έχουμε ότι,

$$\left[e^{-t_2} (t_2 - \tau_2) / (e^{-t_1} - e^{-t_2}) \right] d(t_2 - t_1) = \left[\tau_1 / (1 - e^{-t_1}) \right] dt_1 \quad \text{ή κάνοντας πράξεις στη συνέχεια}$$

$$\frac{dt_2}{dt_1} = 1 + \tau_1 (e^{-t_1} - e^{-t_2}) / (1 - e^{-t_1}) e^{-t_2}.$$

Από τη στιγμή που ο 2ος όρος δεξιά είναι θετικός έπεται ότι το t_2 αυξάνει με το t_1 . Τώρα συνεχίζουμε δείχνοντας ότι αύξηση στο t_1 προκαλεί αύξηση σε όλα τα t που ικανοποιούν τις εξισώσεις.

Παίρνοντας παραγώγους στην (4.10)

$$d(t_{i+1} - t_i) / d(t_i - t_{i-1}) = e^{-t_{i+1}} (\tau_i - t_{i-1}) (e^{-t_i} - e^{-t_{i+1}}) \left[e^{-t_{i+1}} (t_{i+1} - \tau_{i+1}) (e^{-t_{i-1}} - e^{-t_i}) \right] > 0$$

$$i = 1, 2, \dots, (k-1)$$

Που $d(t_{i+1} - t_i) / d(t_2 - t_1) = A_i > 0 \Rightarrow d(t_{i+1} - t_i) = A_i d(t_2 - t_1)$ άρα, $dt_{i+1} / dt_1 > 0$ $i = 1, 2, \dots, (k-1)$ που αποδεικνύει την υπόθεση. Άρα λοιπόν μια συνεχής αύξηση στο t_1 αυξάνει τα $t_2, \dots, t_{k-1}$ που ικανοποιούν τις εξισώσεις. Από τη μονοτονία βγαίνει το συμπέρασμα ότι η λύση θα είναι μοναδική. Στη συνέχεια αποδεικνύεται η ύπαρξη της λύσης:

Αρχικά χωρίζεται το διάστημα $[0, 1]$ σε k υποδιαστήματα τέτοια ώστε

$$1 - e^{-t_1^0} = e^{-t_1^0} - e^{-t_2^0} = \dots = e^{-t_{k-3}^0} - e^{-t_{k-2}^0} = e^{-t_{k-1}^0}$$

Παίρνεται λοιπόν με

προσομοίωση ένα t_1^0 για το t_1 και παράγεται ένα t_2^* που να ικανοποιεί την εξίσωση

$2t_1 = \tau_1 - \tau_2$. Παρομοίως παράγονται $t_3^*, \dots, t_{k-1}^*$ τα οποία ικανοποιούν τις ανισώσεις από Λήμμα $t_2^* < t_2^0, \dots, t_{k-1}^* < t_{k-1}^0$. Για αυτό αυξάνοντας συνέχεια το t_1^0 αυξάνονται και τα $t_2^*, \dots, t_{k-1}^*$ και τελικά βρίσκεται το $\tau_k^* = 1 + t_k^*$ και η λύση είναι μοναδική. Αυτό αποδεικνύει ότι το σύστημα εξισώσεων έχει μία και μοναδική λύση. Από 2^ο Λήμμα [Saleh and Ali (1966)] το οποίο τώρα δε θα αναφέρουμε για οικονομία χρόνου η μοναδική λύση αντιστοιχεί στο μέγιστο της συνάρτησης Q_{k-1} . Άρα η συνάρτηση αυτή έχει μοναδικό μέγιστο στο διάστημα $0 \leq t_1 \leq t_2 \leq \dots \leq t_k \leq \infty$. Αυτό ολοκληρώνει το θεώρημα.

Εδώ θα αναφέρουμε και ένα ακόμα δύο θεωρήματα τα οποία χρησιμοποιούνται από τον Saleh αλλά θα παραλείψουμε την απόδειξη τους.

Θεώρημα 2

Εάν η $(k-1)$ -διάστατη συνάρτηση $Q_{k-1}(t_1, t_2, \dots, t_{k-1})$ που ορίζεται ως εξής,

$Q_{k-1}(t_1, t_2, \dots, t_{k-1}) = \sum_{i=1}^{k-1} (t_i - t_{i-1}) / (e^{t_i} - e^{t_{i-1}})$ με $t_0 = 0$ και $0 < t_1 < t_2 < \dots < t_{k-1} < \infty$ επιτυγχάνει το μοναδικό μέγιστο της στο σημείο

$(t_1^*, t_2^*, \dots, t_{k-1}^*)$, τότε η k διάστατη συνάρτηση $Q_k(t_1, t_2, \dots, t_k)$ με τύπο

$Q_k(t_1, t_2, \dots, t_k) = \sum_{i=1}^k (t_i - t_{i-1}) / (e^{t_i} - e^{t_{i-1}})$ πετυχαίνει το μέγιστο της στο $(t_1^0, t_2^0, \dots, t_k^0)$ όπου $t_{i+1}^0 = t_i^0 + t_i^*$, $i = 1, 2, \dots, k-1$

και το t_1^0 είναι η ρίζα της εξίσωσης $t_1 / (1 - e^{-t_1}) = 1 + (1 - Q_{k-1}^0)^{\frac{1}{2}}$.

Με Q_{k-1}^0 ορίζεται η μέγιστη τιμή της $Q_{k-1}(t_1, t_2, \dots, t_{k-1})$ στο σημείο $(t_1^*, t_2^*, \dots, t_{k-1}^*)$

Για κάθε ποσοστημόριο έχουν υπολογιστεί με τη βοήθεια του προγραμματισμού σε Η/Υ τα ποσοστημόρια p_i , των παρατηρήσεων που θα χρησιμοποιηθούν και οι συντελεστές b_i για κάθε περίπτωση χρησιμοποιώντας οποιονδήποτε αριθμό παρατηρήσεων μικρότερο από το σύνολο του δείγματος.

4.3.3 Μειονεκτήματα Γραμμικών εκτιμητών-Παράδειγμα-Πίνακας.

Οι γραμμικοί εκτιμητές αν και έχουν πολλά καλά γνωρίσματα εκτιμητών όπως αμεροληψία και σχετικά μικρή διασπορά σε σχέση με τους υπόλοιπους έχουν και κάποια αρνητικά στοιχεία. Ένα από αυτά είναι η διαδικασία υπολογισμού των ποσοστημορίων p_i και των συντελεστών b_i κάθε φορά για διαφορετικό μέγεθος δείγματος n διαφορετικό βαθμό censoring p αλλά και διαφορετικό αριθμό χρησιμοποιούμενων διατεταγμένων παρατηρήσεων ανάλογα τον ερευνητή και το πρόβλημα. Ο Υπολογισμός τους γίνεται αν μη τι άλλο δύσκολος οπότε καταφεύγουμε στην ανάγκη χρήσης πινάκων ειδικότερα στην περίπτωση της εκθετικής κατανομής όπου η παράμετρος είναι μία. Όπως όμως πάλι είναι κατανοητό το πλήθος των πινάκων είναι αρκετά μεγάλο για να αποθηκευτούν και η πολυπλοκότητα επιλογής κατά τον προγραμματισμό σε Η/Υ για την εύρεση του εκτιμητή σημαντική. Ειδικότερα στο πρόβλημα του εφημεριδοπώλη που το πλήθος r των παρατηρήσεων μικρότερων του αποθέματος είναι τυχαία μεταβλητή το πλήθος των διατεταγμένων παρατηρήσεων που θα χρησιμοποιηθούν είναι εξαρτώμενο και καθιστά κάθε φορά χρήση διαφορετικού πίνακα. Για παράδειγμα εάν $r = 2$ χρειαζόμαστε ειδικό πίνακα με χρήση μίας ή και των δύο παρατηρήσεων. Όταν πάλι όλες οι τιμές του δείγματος είναι διαθέσιμες τότε πρέπει να ανατρέξουμε σε τελείως διαφορετικό πίνακα ο οποίος αναφέρεται στον εκτιμητή με σταθερές υπολογισμένες γιατί την συγκεκριμένη περίπτωση. Οι σταθερές b_i, p_i υπολογίζονται με διαφορετική διαδικασία σε αυτή την περίπτωση αν και η λογική της εύρεσης τους δεν αλλάζει ουσιαστικά. Παρακάτω παραθέτετε ο πίνακας για τα ποσοστημόρια p_i , αλλά και τα αντίστοιχα b_i χρησιμοποιώντας 2, 3 ή 4 παρατηρήσεις για τον υπολογισμό της παραμέτρου. Έχουν υπολογιστεί επίσης για κάθε περίπτωση τα t_i και η τιμή της συνάρτησης Q_k . Τα ποσοστά περικοπής του δείγματος είναι από 55% έως και 95% με βήμα 5%. Αυτό γίνεται για να γίνει κατανοητό ότι αυτός ο πίνακας είναι μια αρκετά συγκεκριμένη περίπτωση για συγκεκριμένο-στρογγυλοποιημένο βαθμό censoring (δεν περιλαμβάνει π.χ την περίπτωση $r/n = 63\%$).

Παράδειγμα υπολογισμού βέλτιστου γραμμικού εκτιμητή και διασπορά αυτού

Ας υποθέσουμε ότι έχουμε ένα δείγμα μεγέθους $n = 72$ σε ένα $1-\beta=20\%$ δεξιά περικομμένο δείγμα χρησιμοποιώντας $k = 4$ διατεταγμένες παρατηρήσεις. Με $r = [n\beta] + 1 = [72 \cdot 0,8] + 1 = 57 + 1 = 58$ επομένως η maximum παρατήρηση που θα χρησιμοποιηθεί είναι η $x_{(58)}$. Από τον πίνακα στη προηγούμενη σελίδα για την περίπτωση που $k = 4$ παίρνουμε ότι $p_1 = 0.2800$, $p_2 = 0.5022$, $p_3 = 0.6733$ άρα οι παρατηρήσεις του δείγματος που θα χρησιμοποιηθούν είναι οι $x_{(21)}, x_{(37)}, x_{(49)}, x_{(58)}$ και διαβάζοντας τους συντελεστές υπολογίζουμε το $\frac{1}{\theta_{BL}}$ ως

$$\frac{1}{\theta_{BL}} = 0.3157x_{(21)} + 0.2469x_{(37)} + 0.1864x_{(49)} + 0.3204x_{(58)}$$

Συγκρίνοντας τις διασπορές των Εκτιμητών.

Η διασπορά του Ε.Μ.Π όπως αναφέρθηκε στην ενότητα 3.2 από τη στιγμή που η τ.μ $2\theta y$ ακολουθεί Χ-τετράγωνο με $2r$ βαθμούς ελευθερίας είναι $V(\frac{1}{\hat{\theta}}) = \frac{1}{r\theta^2}$ ενώ η διασπορά του BLUE εκτιμητή είναι $V(\tilde{\theta}^{-1}) = Q_k^{-1} \theta^{-2} / n = \frac{1}{Q_k \theta^2 n}$ με το k στη συγκεκριμένη περίπτωση να είναι $k = 4$. Τότε από τον πίνακα της σελίδας 56 παρατηρούμε ότι $Q_4 \approx p = \frac{r}{n}$ και επίσης $Q_4 \leq \frac{r}{n}$ ειδικότερα όσο μεγαλύτερη είναι η περικοπή των δεδομένων τόσο πιο πολύ συγκλίνει το $Q_4 \rightarrow r/n$.

Επομένως στη σύγκριση των διασπορών δίνει την ανισότητα $V(\frac{1}{\hat{\theta}}) < V(\frac{1}{\tilde{\theta}_{Blue}})$

λόγω της προηγούμενης παρατήρησης. Με χρήση περισσότερων διατεταγμένων παρατηρήσεων του δείγματος για το γραμμικό εκτιμητή η διασπορά αυτού συγκλίνει στη διασπορά του Ε.Μ.Π.

TABLE I
The optimum values of $p_1, \dots, p_k$ for the linear estimate of σ based on k order statistics when the sample is censored on the right for $k = 2(1)4$ and for $1 - \beta = .05(.05)40$ with coefficients $b_1, \dots, b_k$

Per Cent Censored	k	t_1	t_2	t_3	t_4	Q_k^*	p_1	p_2	p_3	p_4	b_1	b_2	b_3	b_4
5	2	1.0177	2.6613	3.0000	3.0000	.8203	.6386	.9500	.9500	.9500	.5232	.1790	.1029	
	3	.7033	1.6294			.8855	.5051	.8080	.8536		.4373	.2354	.1423	.8894
	4	.5132	1.1341	1.9208		.9151	.4014	.6783			.3643	.2407		
10	2	.9278	2.3025			.8159	.6046	.9000			.5181	.2255	.1821	
	3	.5877	1.3217	2.3025		.8621	.4444	.7334	.9000		.4156	.2546	.1628	.1627
	4	.4304	.9335	1.5395	2.3025	.8755	.3498	.6068	.7856	.9000	.3413	.2439		
15	2	.7973	1.8971			.7932	.5495	.8500			.5124	.3118	.2614	
	3	.5084	1.1219	1.8971		.8245	.3985	.6712	.8500		.4021	.2669	.1762	.2386
	4	.3733	.8004	1.2982	1.8971	.8356	.3115	.5509	.7272	.8500	.3266	.2457		
20	2	.6961	1.6094			.7603	.5015	.8000			.5089	.4012	.3459	
	3	.4462	.9712	1.6094		.7822	.3599	.6214	.8000		.3921	.2763	.1864	.3204
	4	.3285	.6975	1.1187	1.6094	.7899	.2800	.5022	.6733	.8000	.3157	.2469		
25	2	.6127	1.3863			.7218	.4582	.7500			.5067	.4975	.4382	
	3	.3945	.8493	1.3863		.7374	.3260	.5723	.7500		.3843	.2837	.1945	.4105
	4	.2910	.6134	.9749	1.3863	.7429	.2525	.4585	.6228	.7500	.3070	.2476		
30	2	.5413	1.2039			.6798	.4180	.7000			.5050	.6036	.5410	
	3	.3499	.7465	1.2039		.6909	.2952	.5260	.7000		.3777	.2899	.2015	.5115
	4	.2586	.5418	.8545	1.2039	.6949	.2279	.4183	.5747	.7000	.2997	.2481		
35	2	.4788	1.0499			.6356	.3805	.6500			.5038	.7227	.6575	
	3	.3107	.6575	1.0498		.6436	.2671	.4819	.6500		.3721	.2953	.2074	.6265
	4	.2299	.4791	.7509	1.0498	.6463	.2054	.3807	.5281	.6500	.2935	.2486		
40	2	.4230	.9162			.5898	.3449	.6000			.5029	.8593	.7916	
	3	.2754	.5786	.9162		.5954	.2409	.4494	.6000		.3672	.3000	.2126	.7593
	4	.2041	.4232	.6595	.9162	.5973	.1847	.3451	.4829	.6000	.2881	.2489		

Ενότητα 4.4 Βέλτιστη ποσότητα παραγγελίας – Προσομοιώσεις

4.4.1 Αλγόριθμος προσομοίωσης για εκτιμήσεις του $1/\theta$

Παρακάτω περιγράφουμε τον αλγόριθμο προσομοίωσης που πραγματοποιήθηκε με το στατιστικό πακέτο R.

Για $n=10.000$ επαναλήψεις κάνουμε την εξής διαδικασία:

- Επιλέγουμε μία προκαθορισμένη μέση τιμή ζήτησης ανά περίοδο $\mu(\theta)=1/\theta=100$ και κάνουμε προσομοίωση τις $n=10, 20, 30$ τιμές του δείγματος μας από την Εκθετική(θ).
- Καθορίζουμε ένα προεπιλεγμένο Q στο οποίο στη συγκεκριμένη εργασία δίνουμε τέσσερις διαφορετικές τιμές. Αυτές είναι $Q=100\% \cdot \mu(\theta)$, $120\% \cdot \mu(\theta)$, $150\% \mu(\theta)$, $200\% \cdot \mu(\theta)$. Οπότε δεδομένου αποθέματος καθορίζεται σε κάθε επανάληψη και ο αριθμός των παρατηρήσεων που το υπερβαίνουν. Υπολογίζεται ο Ε.Μ.Π της μέσης τιμής του περικομμένου δείγματος παίρνοντας 2 περιπτώσεις είτε με γνώση του Q είτε με γνώση της μέγιστης παρατηρούμενης τιμής από το δείγμα. Αποθηκεύονται οι εκτιμήσεις της κάθε επανάληψης.
- Προσομοιώνονται οι υπόλοιπες $n-r$ τιμές από το δείγμα με χρήση της σχέσης και εκτιμάται ο «νέος» εκτιμητής εφόσον $n-r > 0$ της μέσης τιμής με χρήση του Ε.Μ.Π του θ . Εάν είναι διαθέσιμες όλες οι δυνατές παρατηρήσεις από το δείγμα τότε αντικαθίστανται ο εκτιμητής από τον Ε.Μ.Π αφού δεν χρειάζεται να προσομοιωθεί κάποια τιμή. Εάν η προσομοιωμένη τιμή $X_{(r+1)}^*$ είναι μικρότερη από το απόθεμα Q τότε θέτουμε $X_{(r+1)}^* \leftarrow Q$
- Αναλόγως με το ποσοστό περικομμένων παρατηρήσεων βρίσκουμε από τον πίνακα τους συντελεστές b_i καθώς και τα n_k που προσδιορίζουν τις τιμές του δείγματος που θα χρησιμοποιηθούν για να υπολογίσουμε τους βέλτιστους γραμμικούς εκτιμητές (BLUE estimators).
- Υπολογίζουμε το Μέσο Τετραγωνικό Σφάλμα του κάθε εκτιμητή

4.4.2 Εκτίμηση Βέλτιστου Κέρδους με χρήση του εκτιμητή μέγιστης πιθανοφάνειας.

Σκοπός της ενότητας αυτής είναι να δούμε πόσο μεγάλη απόκλιση υπάρχει στην εκτίμηση του βέλτιστου κέρδους όταν χρησιμοποιούμε τον εκτιμητή για τα

περικομμένα δεδομένα έναντι του δειγματικού μέσου $\bar{D} = \frac{\sum_{i=1}^n X_i}{n}$ που

χρησιμοποιείται όταν έχουμε και τις n παρατηρήσεις του δείγματος. Εμείς λοιπόν χρησιμοποιούμε τον Ε.Μ.Π για την εκτίμηση της μέση τιμής για το περικομμένο

δείγμα δηλαδή για τον εκτιμητή $\frac{1}{\hat{\theta}} = \frac{y}{r}$ με το y όπως ορίστηκε προγενέστερα

δηλαδή $y = \sum_{i=1}^r X_i + (n-r)Q$

Η συνάρτηση κέρδους είναι $\varphi(q, \theta) = \frac{s}{\theta} (1 - \exp\{-\theta q\}) - cq$ η οποία

μεγιστοποιείται όταν $q^*(\theta) = \frac{1}{\theta} \ln\left(\frac{s}{c}\right)$ οπότε εάν αντικαταστήσουμε στην

προηγούμενη συνάρτηση την ποσότητα αυτή και πάρουμε μέση τιμή, έχουμε ότι

$\psi_{sm}(\theta) = E\left[\frac{s}{\theta} \left(1 - \exp\left\{-\theta \frac{y}{r} \ln\left(\frac{s}{c}\right)\right\}\right) - c \frac{y}{r} \ln\left(\frac{s}{c}\right)\right]$, το οποίο όμως είναι δύσκολο

να υπολογιστεί αφού οι τυχαίες μεταβλητές $\sum_{i=1}^r X_i / r$ και $\frac{(n-r)}{r} Q$ δεν είναι

ανεξάρτητες. Επίσης η τ.μ $\frac{(n-r)}{r}$ μπορεί να εκφραστεί σαν την αναλογία δύο

εξαρτημένων τ.μ που ακολουθούν την Διωνυμική κατανομή με πιθανότητες

επιτυχίας $1 - e^{-\theta Q}$ και $e^{-\theta Q}$. Για τον λόγο αυτό είναι πολύ δύσκολο να εκτιμηθεί η

μέση τιμή του κέρδους στο συγκεκριμένο μοντέλο. Οπότε θα χρησιμοποιήσουμε

αριθμητικές μεθόδους με τη βοήθεια της προσομοίωσης για να βγάλουμε ασφαλή

συμπεράσματα. Θα προσομοιώσουμε λοιπόν ένα μεγάλο αριθμό δειγμάτων και θα

υπολογίσουμε τις ποσότητες :

$$\varphi(q_{FS}^*, \theta) = \frac{s}{\theta} \left(1 - \exp\left\{-\theta \bar{X} \ln\left(\frac{s}{c}\right)\right\}\right) - c \bar{X} \ln\left(\frac{s}{c}\right) \quad Fs = Full \ sample$$

$$\varphi(q_{CS}^*, \theta) = \frac{s}{\theta} \left(1 - \exp \left\{ -\theta \frac{y}{r} \ln \left(\frac{s}{c} \right) \right\} \right) - c \frac{y}{r} \ln \left(\frac{s}{c} \right) \quad Cs = \text{Censored sample}$$

Οπότε θα υπολογίσουμε τις διαφορές τους από το βέλτιστο κέρδος το οποίο

δείχθηκε στο 2^ο κεφάλαιο ότι είναι $\varphi(q^*, \theta) = \frac{c}{\theta} \left\{ \frac{s}{c} - 1 - \ln \left(\frac{s}{c} \right) \right\}$ για ένα πλήθος

δειγμάτων που προσομοιωθούν για $k = 1.000$ επαναλήψεις.

Αντί της Ε.Μ.Π με χρήση των παρατηρήσεων θα δοκιμάσουμε και αυτή με χρήση μόνο των παρατηρούμενων τιμών αλλά και του προσομοιωτικού εκτιμητή για να δούμε με ποια από τις δύο προσεγγίζεται καλύτερα το κέρδος. Οι ποσοστιαίες διαφορές υπολογίζονται από τους τύπους:

$$Diff_{FS} = \frac{\varphi_i(q^*, \theta) - \varphi(q_{FS}^*, \theta)}{\varphi_i(q^*, \theta)} \cdot 100\%, \quad Diff_{CS} = \frac{\varphi_i(q^*, \theta) - \varphi(q_{CS}^*, \theta)}{\varphi_i(q^*, \theta)} \cdot 100\% \quad (4.10)$$

$$\text{Με } E(\varphi(q_{FS}^*, \theta)) = \frac{\sum_{i=1}^k \frac{s}{\theta} \left(1 - \exp \left\{ \theta \bar{X}_i \ln \left(\frac{s}{c} \right) \right\} \right) - c \ln \left(\frac{s}{c} \right) \bar{X}_i}{k}.$$

Με ανάλογο τρόπο υπολογίζεται μέσω αλγορίθμου Metropolis-Hastings η μέση τιμή του κέρδους για το περικομμένο δείγμα με κάποια από τις δύο εκτιμήτριες με τη μικρότερη τυπική απόκλιση. Αυτές είναι η Ε.Μ.Π και ο Προσομοιωτικός εκτιμητής όπως φαίνεται και από τους πίνακες του 5^{ου} Κεφαλαίου.

Οι διαφορές εδώ δεν μπορεί να αποδειχθεί εδώ με φορμαλιστικό τρόπο ότι είναι ανεξάρτητες της παραμέτρου όπως στην εργασία [5] όμως με χρήση προσομοιώσεων είναι ξεκάθαρο ότι οι ποσοστιαίες αποκλίσεις από το βέλτιστο κέρδος παραμένουν σταθερές. Για ένα ρεαλιστικό αριθμό περιόδων $n = 10$ που δε θα προκαλέσει ολική έλλειψη παρατηρήσεων ζήτησης προσομοιώνουμε με τις εξής παραμέτρους. Η τιμή κόστους ανά μονάδα είναι $c = 1$ χρηματικές μονάδες και πώλησης $s = 1,5$. Το απόθεμα επιλέγεται ως 1.5 φορά τη μέση τιμή ζήτησης δηλαδή $Q = 150$. Προσομοιώνοντας βλέπουμε ότι $Diff_{FS} = 8,5\%$ ενώ για την Ε.Μ.Π του περικομμένου δείγματος $Diff_{CS} = 14,9\%$. Στα ελλιπή δεδομένα η μικρότερη

απώλεια επιτεύχθηκε με χρήση του εκτιμητή μέγιστης πιθανοφάνειας με χρήση των παρατηρήσεων ζήτησης. Επομένως παρατηρείται ότι η απώλεια δεδομένων στοιχίζει σημαντικά στην εκτίμηση της βέλτιστης παραγγελίας αποθέματος για την επόμενη περίοδο.

Όταν το πλήθος των περιόδων διπλασιαστεί από δέκα σε είκοσι τότε οι αντίστοιχες διαφορές μικραίνουν. Έτσι έχουμε τώρα παίρνουμε $Diff_{FS} = 4.3\%$, $Diff_{CS} = 7.1\%$.

Τα αποτελέσματα αυτά είναι λογικά αν σκεφτούμε ότι η ακρίβεια της προσέγγισης του $\frac{1}{\theta}$ βελτιώνεται κατά πολύ αλλά και το σφάλμα των εκτιμητών. Στη συνέχεια θα

αλλάξουμε την ποσότητα παραγγελίας με αυτή της σχέσης (1.4) που

παρουσιάστηκε στο κεφάλαιο 1 χρησιμοποιώντας πάλι τους ίδιους Εκτιμητές για τη μέση τιμή. Έτσι θα δούμε εάν η βέλτιστη ποσότητα παραγγελίας όπως παρουσιάζεται στην εργασία των Chu, Shanthikumar, Shen [5] η οποία αποτελεί και τη βέλτιστη ποσότητα παραγγελίας 1^{ης} τάξης ως συνάρτηση του δείγματος , συνεχίζει να δίνει βελτιώσεις ως προς το κέρδος και στην περίπτωση αυτή. Επίσης εξετάζεται εάν η ποσότητα αυτή δίνει καλύτερες εκτιμήσεις ή όχι σε οποιοδήποτε τιμή των εκτιμητριών, όταν αυτές δηλαδή υπερεκτιμούν ή υποτιμούν την

παράμετρο .Αντικαθιστώντας λοιπόν το $\ln\left(\frac{S}{C}\right)$ με την ποσότητα $n\left(\left(\frac{S}{C}\right)^{\frac{1}{n+1}} - 1\right)$

Παρατηρούμε ότι $n\left(\left(\frac{S}{C}\right)^{\frac{1}{n+1}} - 1\right) \leq \ln\left(\frac{S}{C}\right)$ και $n\left(\left(\frac{S}{C}\right)^{\frac{1}{n+1}} - 1\right) \rightarrow \ln\left(\frac{S}{C}\right)$. Κάνοντας τις

προσομοιώσεις για τις ίδιες παραμέτρους όπως πριν με ποσότητες παραγγελίας

$$q_{MLE} = \frac{\sum_{i=1}^r X_i + (n-r)Q}{r} \ln\left(\frac{S}{C}\right) \quad \text{και} \quad q_{MLE(Or)}^* = \frac{\sum_{i=1}^r X_i + (n-r)Q}{r} n\left(\left(\frac{S}{C}\right)^{\frac{1}{n+1}} - 1\right)$$

παίρνουμε τη μέση απόκλιση του βέλτιστου κέρδους από το εκτιμημένο για 10.000

επαναλήψεις. Οι τύποι που χρησιμοποιούνται είναι όπως της (4.10) της

προηγούμενης σελίδας αλλά όχι με αλλαγή πλέον της εκτιμητριας της παραμέτρου.

Οι αποκλίσεις αυτές είναι λοιπόν 17.95% με χρήση του λογαρίθμου και 14.92% με

χρήση της ακολουθίας σύγκλισης . Επομένως η διαφορά παραμένει υπέρ της δεύτερης ειδικότερα για μικρό αριθμό δειγμάτων .

Φανερό γίνεται ότι όσο το δείγμα μικραίνει υπερεκτιμάται η μέση τιμή της ζήτησης

και λόγω της ανισότητας $n \left(\left(\frac{s}{c} \right)^{\frac{1}{n+1}} - 1 \right) \leq \ln \left(\frac{s}{c} \right)$ το γινόμενο αυτού με την

εκτιμήτρια συγκλίνει καλύτερα στη βέλτιστη ποσότητα παραγγελίας δηλαδή

$$n \left(\left(\frac{s}{c} \right)^{\frac{1}{n+1}} - 1 \right) \cdot \frac{1}{\hat{\theta}_{\text{EMΠ}}} \rightarrow \ln \left(\frac{s}{c} \right) \cdot \frac{1}{\theta} = X_{\text{BEΛ}}$$

Έτσι για παράδειγμα για μικρότερο αριθμό περιόδων από πριν έστω $n = 6$ και

$Q = 150$ η διαφορά του κέρδους χρησιμοποιώντας την ποσότητα

$X_{\text{EMΠ}} = \ln \left(\frac{s}{c} \right) \cdot \frac{1}{\hat{\theta}_{\text{EMΠ}}}$ από το βέλτιστο είναι 29,3% μικρότερο ενώ χρησιμοποιώντας

το $X_{\text{Op:EMΠ}} = n \left(\left(\frac{s}{c} \right)^{\frac{1}{n+1}} - 1 \right) \cdot \frac{1}{\hat{\theta}_{\text{EMΠ}}}$ είναι 22,7% πράγμα που σημαίνει ότι η ποσότητα

βελτιώνει το αναμενόμενο κέρδος κατά 7 ποσοστιαίες μονάδες ωστόσο και οι δύο απέχουν πολύ από το βέλτιστο.

ΚΕΦΑΛΑΙΟ 5^ο

Ενότητα 5.1 Αποτελέσματα Προσομοιώσεων-Πίνακες

Παρακάτω παρουσιάζονται κάποιοι πίνακες με τα αποτελέσματα των προσομοιώσεων των εκτιμητών της παραμέτρου με τους τρεις διαφορετικούς τρόπους εκτίμησης που αναπτύχθηκαν σε αυτήν την εργασία δηλαδή για

1. Εκτιμήτρια μέγιστης πιθανοφάνειας
2. Γραμμικό Εκτιμητή με χρήση των διατεταγμένων παρατηρήσεων
3. Προσομοιωμένες τιμές των εκλιπουσών παρατηρήσεων με χρήση της εκτιμήτριας μέγιστης πιθανοφάνειας
 - i) Με γνώση του αποθέματος Q
 - ii) Χωρίς γνώση του αποθέματος αλλά μόνο της μέγιστης παρατηρούμενης τιμής της ζήτησης

Η προσομοίωση έγινε για επιλογές αποθέματος

$$Q = 100\% \cdot \mu(\theta), \quad 120\% \cdot \mu(\theta), \quad 150\% \mu(\theta), \quad 200\% \cdot \mu(\theta)$$

για πραγματική μέση τιμή ζήτησης του προϊόντος $1/\theta=100$ μονάδες και για

$k=10.000$ προσομοιώσεις. Για τον προσομοιωτικού εκτιμητή δεν χρησιμοποιούμε την Ε.Μ.Π χρησιμοποιώντας την μέγιστη παρατήρηση του μη-περικομμένου δείγματος όπως στο Κεφάλαιο 4 αλλά το απόθεμα Q .

ΠΙΝΑΚΑΣ Ι Εκτιμήσεις ζήτησης για 10 χρονικές περιόδους $n=10$, $1/\theta=100$, $k=10.000$

Ύψος Αποθέματος Q	Εκτιμητής Μέγιστης Πιθανοφάνειας με το απόθεμα Q (MLE)	Εκτιμώμενη τιμή του $1/\theta$ με τη μέγιστη παρατήρηση	Εκτιμητής με Προσομοίωση	Γραμμικός Εκτιμητής (BLUEs)	Εκτιμητής Μέγιστης Πιθανοφάνειας με χρήση του πλήθους Παρατηρήσεων
100	111.3	95.4	107.5	-	111.5
120	108.7	95.6	105.8	-	108.2
150	106.2	95.3	106.8	-	106.8
200	104.5	95.4	110.2	-	109.6

ΠΙΝΑΚΑΣ ΙΙ Μ.Τ.Σ εκτιμητών για 10 χρονικές περιόδους $n=10$, $1/\theta=100$, $k=10.000$

Ύψος Αποθέματος Q	Μ.Τ.Σ Εκτιμητή Μεγ. Πιθαν. (Ε.Μ.Π) με το Απόθεμα Q	Μ.Τ.Σ Εκτιμητή με χρήση της Μέγ.παρατηρ	Μ.Τ.Σ Προσο/τικού Εκτιμητή	Μ.Τ.Σ Γραμμικού Εκτιμητή BLUE	Μ.Τ.Σ Ε.Μ.Π με χρήση του πλήθους Παρατηρήσεων
100	3183	1889	3526	-	3231
120	2537	1682	2924	-	2762
150	2601	1410	2161	-	2823
200	1496	1174	1767	-	1404

ΠΙΝΑΚΑΣ III Εκτιμήσεις ζήτησης για 20 χρονικές περιόδους

$n = 20$, $1/\theta=100$, $k=10.000$

Ύψος Αποθέματος Q	Ε.Μ.Π με χρήση των Παρατηρήσεων	Ε.Μ.Π με χρήση της μέγιστης παρατήρησης	Προσομ/κός Εκτιμητής (Simul. Est)	Γραμμικός Εκτιμητής (BLUEs)	Ε.Μ.Π με χρήση πλήθους παρατηρήσεων
100	105.7	97.5	103.4	-	105.3
120	103.7	96.9	101.4	-	103.1
150	102.9	97.1	100.7	-	102.0
200	101.7	97.5	99.3	-	101.0

ΠΙΝΑΚΑΣ IV Μ.Τ.Σ εκτιμητών για 20 χρονικές περιόδους

$n = 20$, $1/\theta=100$, $k=10.000$

Ύψος Αποθέματος Q	Μ.Τ.Σ Ε.Μ.Π με χρήση των Παρατηρήσεων	Μ.Τ.Σ Ε.Μ.Π με χρήση της Μέγιστης παρατήρησης	Μ.Τ.Σ Προσο/τικού Εκτιμητή	Μ.Τ.Σ Γραμμικού Εκτιμητή BLUE	Μ.Τ.Σ Ε.Μ.Π με χρήση πλήθους παρατ/σεων
100	1058	856	1215	-	1171
120	917	764	1057	-	1027
150	776	651	842	-	886
200	650	587	683	-	827

ΠΙΝΑΚΑΣ V Εκτιμήσεις ζήτησης για 50 χρονικές περιόδους

$n = 50$, $1/\theta=100$, $k=10.000$

Ύψος Αποθέματος Q	Ε.Μ.Π με χρήση των Παρατηρήσεων	Ε.Μ.Π με χρήση της μέγιστης παρατήρησης	Προσομ/κός Εκτιμητής (Simul. Est)	Βέλτιστος Γραμμικός Εκτιμητής BLUE	Ε.Μ.Π με χρήση του πλήθους Παρατηρήσεων
100	101.9	98.6	100.8	-	101.8
120	101.3	98.2	100.2	-	101.1
150	101.1	98.6	100.1	-	100.7
200	100.7	98.7	99.8	-	99.9

ΠΙΝΑΚΑΣ IV Μ.Τ.Σ εκτιμητών για 50 χρονικές περιόδους

$n = 50$, $1/\theta=100$, $k=10.000$

Ύψος Αποθέματος Q	Μ.Τ.Σ Ε.Μ.Π με χρήση των Παρατηρήσεων	Μ.Τ.Σ Ε.Μ.Π με χρήση της Μέγιστης παρατ/σης	Μ.Τ.Σ Προσο/τικού Εκτιμητή	Μ.Τ.Σ Βέλτιστου Γραμμ.Εκτι BLUE	Μ.Τ.Σ Ε.Μ.Π με χρήση πλήθους παρατ/σεων
100	358	339	419	-	387
120	313	295	370	-	353
150	284	270	327	-	340
200	246	237	266	-	346

Εκτιμητής Μέγιστης Πιθανοφάνειας με χρήση των παρατηρήσεων από το δείγμα

Από τις προσομοιώσεις λοιπόν βγαίνουν τα εξής συμπεράσματα. Η εκτιμήτρια μέγιστης πιθανοφάνειας προσεγγίζει ασυμπτωτικά αμερόληπτα τη μέση τιμή ζήτησης $\mu = \frac{1}{\theta}$. Η αντίστοιχη υπολογισμένη με γνώση του αποθέματος Q σε σχέση με αυτή που χρησιμοποιούμε τη μέγιστη παρατηρούμενη τιμή του δείγματος έχει πολύ μεγαλύτερο μέσο τετραγωνικό σφάλμα αν και η πρώτη εκτιμήτρια δείχνει να υποτιμά σε κάθε περίπτωση την παράμετρο κατά 5% περίπου. Σαφώς και για τις δύο ισχύει ότι με αύξηση του μεγέθους του αποθέματος έχουμε αντίστοιχη μείωση στη διασπορά τους. Ενδεικτικά όταν $Q = 100$ γίνει $Q = 150$ το Μ.Τ.Σ μειώνεται από 1102 σε 818. Το μέγεθος του Σφάλματος είναι αρκετά μεγάλο για τις εκτιμήτριες αυτές και δεν μπορεί να δώσει μια ξεκάθαρη εικόνα για τη ζήτηση από μία και μοναδική παρατήρηση n περιόδων. Για $Q = 150$ η τυπική απόκλιση του εκτιμητή φτάνει τις 28 μονάδες οπότε πρακτικά καθιστά τις εκτιμήτριες αυτές μεροληπτικές. Αυτή μειώνεται δραστικά με αύξηση του δείγματος. Εάν διπλασιάσουμε το δείγμα σε 40 περιόδους τότε έχουμε τυπική απόκλιση 17 μονάδων προϊόντος.

Εκτιμητές Μέγιστης Πιθανοφάνειας με χρήση του μεγέθους δείγματος R

Η εκτιμήτρια αυτή δεν λαμβάνει καθόλου πληροφορία από τις τιμές των παρατηρήσεων. Αυτό αναφέρετε κατά 1^{ov} για να αναδειχθεί η ιδιομορφία της σε σχέση με τις υπόλοιπες και κατά 2^{ov} για να δικαιολογηθεί το μεγάλο τετραγωνικό σφάλμα της σε μικρά n και Q . Η μέση τιμή της εκτίμησης σε δέκα χιλιάδες προσομοιώσεις προσεγγίζει τη μέση τιμή με αποκλίσεις που να δεν ξεπερνούν το

7% της μέσης τιμής ζήτησης εκτός της περίπτωσης που επιλέγεται $Q = 100$ και 10 περίοδοι παρατηρήσεων . Ωστόσο βλέποντας την τυπική απόκλιση και την ακρίβεια της να μην διαφέρουν σημαντικά από αυτές των υπολοίπων μας κάνει να σκεφτούμε ότι είναι ένας επαρκής εκτιμητής για την παράμετρο ο οποίος χρειάζεται το ελάχιστο δυνατό πλήθος δεδομένων.

Προσομοιωτικός Εκτιμητής

Ο προσομοιωτικός εκτιμητής χρησιμοποιεί τις παρατηρήσεις X_i καθώς και το απόθεμα Q για να εκτιμήσει την ζήτηση των περιόδων που δεν παρατηρήθηκε. Αυτό μας οδηγεί διαισθητικά στο συμπέρασμα ότι η διασπορά του θα πρέπει να υπερβαίνει τη διασπορά του Ε.Μ.Π κατά πολύ. Κάτι το οποίο παρατηρούμε με έκπληξη ότι δεν είναι αληθές όμως. Σε πολλές περιπτώσεις και ειδικότερα με την αύξηση του αποθέματος προϊόντων το Μ.Τ.Σ του προσομοιωτικού εκτιμητή είναι σχεδόν το ίδιο με αυτό του Ε.Μ.Π και οι διαφορές στις τυπικές αποκλίσεις επίσης μικρές! Αυτό μας κάνει να υποψιαστούμε ότι υπάρχει μια όχι και τόσο ισχυρή συσχέτιση ανάμεσα στις δύο αυτές εκτιμήτριες που οδηγεί στη μείωση της διασποράς των σφαλμάτων της δεύτερης. Ο προσομοιωτικός εκτιμητής όμως στις 10.000 επαναλήψεις της εκτίμησης έδωσε κατά μέσο όρο τις πιο ακριβείς παρατηρήσεις από όλους τους υπόλοιπους όχι όμως και το μικρότερο μέσο τετραγωνικό σφάλμα.

Οι εκτιμήτριες στο σύνολο των προσομοιώσεων δίνουν την ίδια καλή εκτίμηση για τη ζήτηση με τον προσομοιωτικό εκτιμητή να υπερτερεί σχεδόν σε κάθε περίπτωση

στην ακρίβεια (αμεροληψία). Αυτά είναι αποτελέσματα τα οποία δεν περιμέναμε είναι αλήθεια λόγω της εξάρτησης-συσχέτισης των εκτιμητών Μεγ.Πιθανοφάνειας και Προσομοίωσης.

Βέλτιστος γραμμικός αμερόληπτος Εκτιμητής

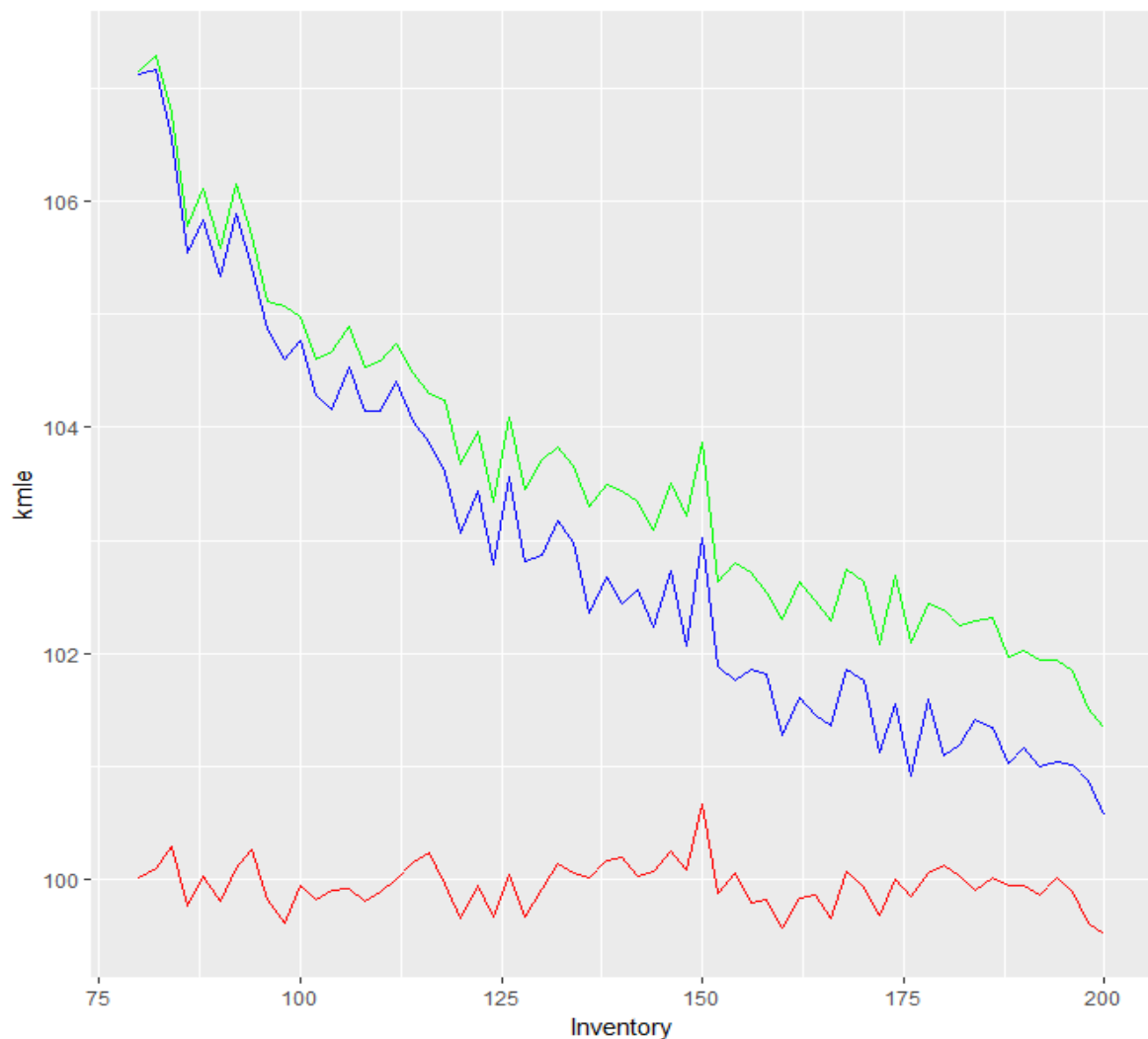
Ο Βέλτιστος γραμμικός εκτιμητής ήταν ο πιο δύσκολος από όλους να υπολογιστεί κατά την προσομοίωση του δείγματος μιας και σε κάθε τυχαίο περικομμένο δείγμα ανάλογα το ποσοστό της περικοπής πρέπει να έχουμε δεκάδες διαφορετικές περιπτώσεις. Δηλαδή για ποσοστό από 5%-95% πρέπει να ανατρέχει ο αλγόριθμος σε διαφορετικά ποσοστημόρια και συντελεστές p_i, b_i για να υπολογιστεί η εκτίμηση. Όλοι αυτά τα δεδομένα πρέπει να ανασυρθούν από πίνακες που υπάρχουν σε διάφορα βιβλία. Μία ακόμα δυσκολία είναι ότι όταν υπάρχουν διαθέσιμες τιμές λιγότερες από 4 όπως φαίνεται στη σελίδα πρέπει να χρησιμοποιηθεί άλλος συνδυασμός σταθερών. Επίσης υπάρχει η ανάγκη για ύπαρξη τιμών το πολύ ίσων με το 95% του δείγματος. Σε περίπτωση διάθεσης του 100% του δείγματος νέος πίνακας είναι αναγκαίο να υπολογιστεί-χρησιμοποιηθεί. Λόγω αδυναμίας εύρεσης γενικευμένου πίνακα με ποσοστά παρατηρήσεων μικρότερα του 50% δεν μπορούμε μέσα από την προσομοίωση να βγάλουμε σαφή συμπεράσματα για την συμπεριφορά του εκτιμητή αυτού.

Για τα σφάλματα των εκτιμητών εκτός του γραμμικού ισχύει η εξής ανισότητα.

$$MT\Sigma_{EMPI}(\mu\epsilon\gamma.\mu\alpha\rho\alpha\tau) \leq MT\Sigma_{EMPI}(\Lambda\mu\acute{o}\theta\epsilon\mu\alpha) \leq MT\Sigma_{SIMUL} \leq MT\Sigma_{EMPI(R)}(\# \mu\alpha\rho\alpha\tau\eta\rho\acute{\eta}\sigma\epsilon\omega\nu)$$

Η ανισότητα αυτή δεν ισχύει απόλυτα σε όλες τις διαφορετικές περιπτώσεις αποθεμάτων αλλά στη συντριπτική πλειοψηφία τους.

Διάγραμμα 1^ο



Σε αυτό το γράφημα παριστάνονται οι εκτιμήτριες μέγιστης πιθανοφάνειας βάση του περικομμένου δείγματος (πράσινο), ολόκληρου (κόκκινο) καθώς και ο Ε.Μ.Π βάσει του «I» (μπλε). Στον οριζόντιο άξονα οι εκτιμήτριες υπολογίζονται συναρτήσει της παραγγελίας αποθέματος για τις πρώτες 20 περιόδους με τιμές $80 \leq Q \leq 200$. Το εντυπωσιακό είναι η καλύτερη προσέγγιση της Ε.Μ.Π με μόνο χρήση του «I» σε σχέση με αυτήν των παρατηρήσεων.

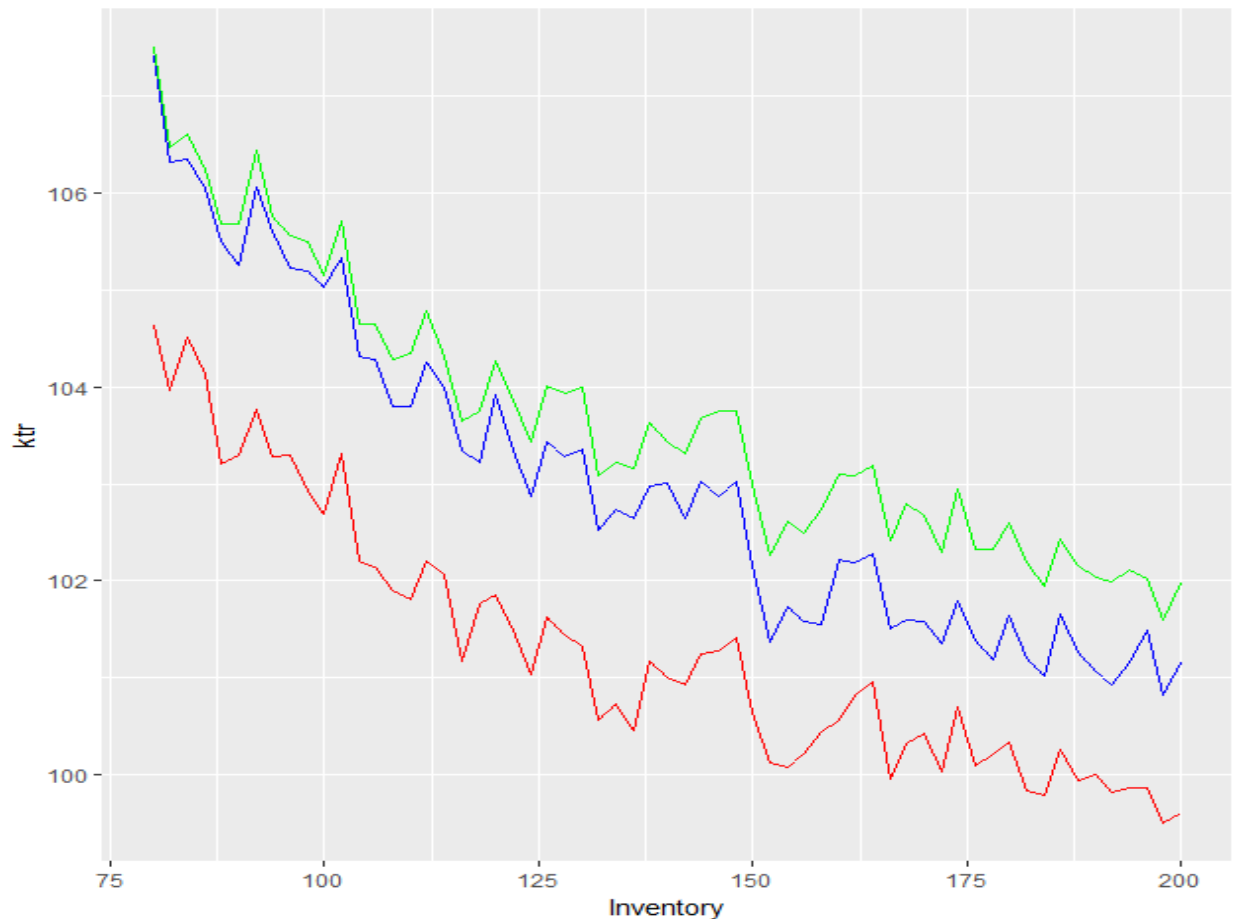
Διάγραμμα 2°



Στο 2° γράφημα γίνεται γραφική προσομοίωση των Ε.Μ.Π με χρήση του αποθέματος Q (πράσινο) ή της μέγιστης παρατήρησης $X_{(r)}$ (κόκκινο) και με χρήση του πλήθους μη-περικομμένων παρατηρήσεων r (Μπλε). Παρατηρούμε ότι ο εκτιμητής βάσει του r δείχνει να είναι επαρκής σε σχέση με αυτόν που χρησιμοποιεί τις ακριβείς τιμές του δείγματος. Οι εκτιμήσεις του είναι πάντα πιο κοντά στην πραγματική τιμή του $\mu(\theta)$ αν και αυτό που δεν αποτυπώνεται επακριβώς στο διάγραμμα είναι το μεγαλύτερο Μ.Τ.Σ που έχει το οποίο παρουσιάζεται για ακραίες τιμές της τυχαίας μεταβλητής όταν δηλαδή $r \rightarrow 0$ και κυρίως όταν $r \rightarrow n$. Ο εκτιμητής με το $X_{(r)}$ δείχνει πολύ μικρές αλλαγές καθώς επιλέγεται μεγαλύτερο απόθεμα. Οπότε οι δύο παρατηρήσεις που εξάγονται είναι ότι παρουσιάζει σταθερή μεροληψία-υποτίμηση της πραγματικής μέσης τιμής κατά τρεις με τέσσερις μονάδες και το μικρότερο Μ.Τ.Σ αφού για μικρά αποθέματα που

οδηγούν σε μεγαλύτερα ποσοστά περικοπής δείγματος δεν επηρεάζουν αρνητικά σημαντικά την εκτίμηση.

Διάγραμμα 3°

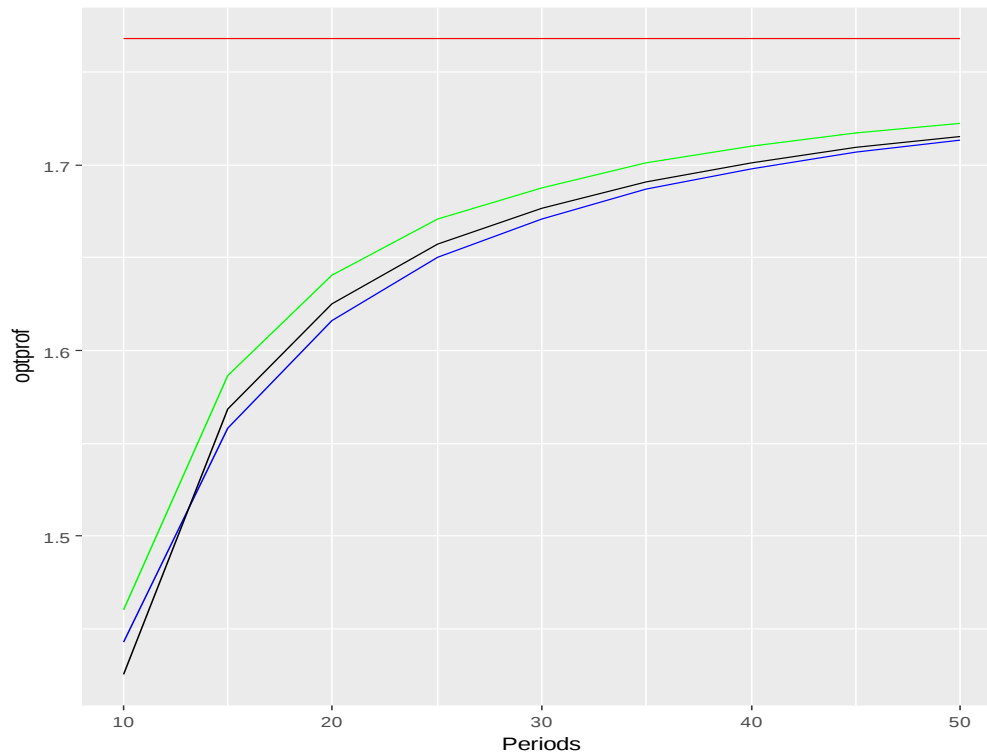


Εδώ γίνεται σύγκριση των Ε.Μ.Π με βάσει τις παρατηρήσεις(Πράσινο), το πλήθος r τιμών μικρότερων του αποθέματος (Μπλε) και του Προσομοιωτικού (κόκκινο).

Παρατηρείται ακριβώς η ίδια συμπεριφορά όπως και στα προηγούμενα διαγράμματα. Για απόθεμα μικρότερο της μέσης τιμής υπάρχει σημαντική έλλειψη πληροφορίας και οι εκτιμητές υπερεκτιμούν το $\frac{1}{\theta}$. Καθώς το ποσοστό περικοπής μικραίνει οι εκτιμητές συγκλίνουν με τον ίδιο ρυθμό στην τιμή των εκατό μονάδων ζήτησης. Και σε αυτό το διάγραμμα είναι εμφανής η καλύτερη προσέγγιση του

προσομοιωτικού εκτιμητή αλλά και η επάρκεια του εκτιμητή βασισμένου στο r σε σχέση με αυτόν της Μεγ. Πιθανοφάνειας σε κάθε επιλογή αποθέματος.

Διάγραμμα 4^ο



Εδώ παρουσιάζουμε το διάγραμμα που δείχνει την εξέλιξη του βέλτιστου κέρδους

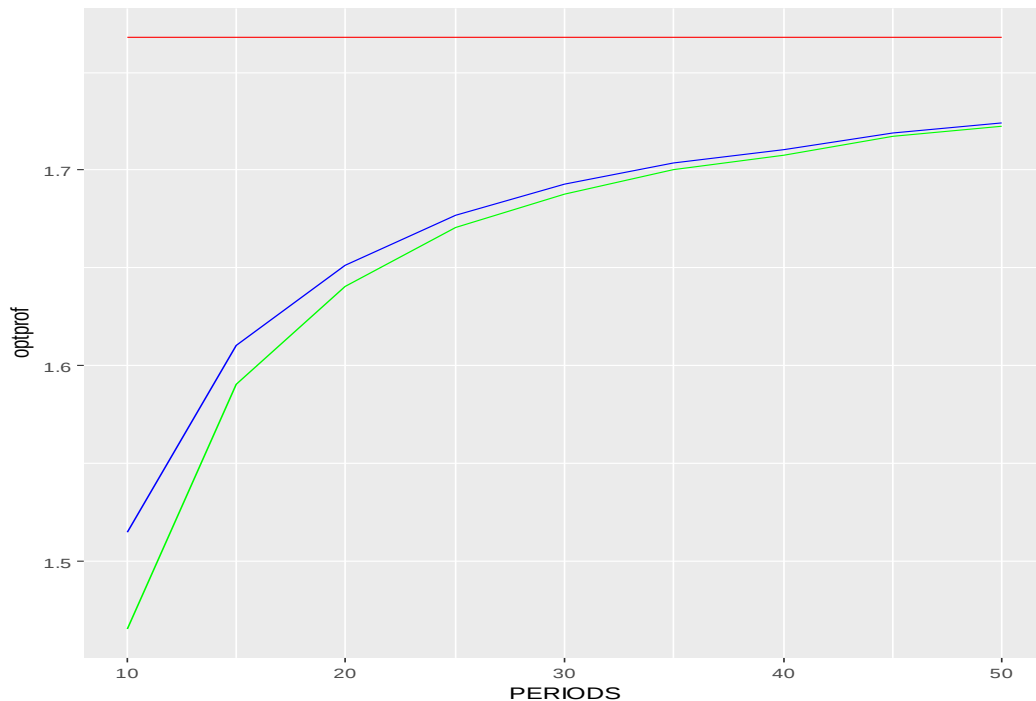
παραγγέλλοντας ποσότητα $Q^* = \frac{1}{\hat{\theta}} \ln\left(\frac{s}{c}\right)$ με εκτίμηση του $1/\theta$ κάθε φορά με

διαφορετικούς εκτιμητές. Στον οριζόντιο άξονα είναι το πλήθος των περιόδων από τις οποίες προέρχονται τα δεδομένα. Με κόκκινο είναι το πραγματικό μέγιστο κέρδος το οποίο παραμένει σταθερό ανεξάρτητο του n . Με πράσινο βάσει του Ε.Μ.Π του δείγματος με μαύρο ο προσομοιωτικός και με μπλε ο Ε.Μ.Π βάσει του r . Φαίνεται η ελαφρά ανωτερότητα του Ε.Μ.Π του δείγματος σε σχέση με τους άλλους δύο παρ' όλο που δεν υπερτερεί σε αμεροληψία αλλά έχει μικρότερη διασπορά.

Εμφανές είναι επίσης η τεράστια βελτίωση στις εκτιμήσεις που γίνεται όταν για παράδειγμα αυξηθεί ο αριθμός των περιόδων αυξηθεί από 10 σε 30.

Επιλέχθηκε σταθερό απόθεμα $Q = 150 = 1,5 \cdot \mu(\theta)$ και κρίσιμος λόγος $\frac{S}{c} = 1,2$.

Διάγραμμα 5°



Στο τελευταίο αυτό διάγραμμα βλέπουμε την επίπτωση που έχει στην εκτίμηση του βέλτιστου κέρδους η επιλογή αντί της βέλτιστης ποσότητας $\frac{1}{\hat{\theta}} \ln\left(\frac{S}{c}\right)$ την

$\frac{1}{\hat{\theta}} n \left(\left(\frac{S}{c} \right)^{1/n+1} - 1 \right)$. Χρησιμοποιούμε τον Ε.Μ.Π του δείγματος για το $\frac{1}{\hat{\theta}}$ στις

παραπάνω σχέσεις, μιας και όπως φάνηκε και από το προηγούμενο διάγραμμα είναι έστω και οριακά καλύτερος από τους άλλους δύο. Με πράσινο φαίνεται το κέρδος με την παραγγελία της 1^{ης} ποσότητας με χρήση λογαρίθμου και με μπλε με χρήση της ακολουθίας σύγκλισης. Η τελευταία βελτιώνει το κέρδος σημαντικά για μικρά n ενώ για $n > 20$ και καθώς αυτό αυξάνεται οι δύο αυτές ποσότητες δεν διαφέρουν σημαντικά και συγκλίνουν στο βέλτιστο κέρδος.

ΠΑΡΑΡΤΗΜΑ

Τμήμα του κώδικα στο προγραμματιστικό πακέτο R για προσομοίωση της μέσης τιμής της μέσης τιμής της ζήτησης $1/\theta$ αλλά και του βέλτιστου κέρδους.

```
v=10000                                knews2[j]=sum(x)/n
k=100                                  kml[j]=sum(x)/n
n=20                                  }
Q=150                                  if ( r[j]==0 ){
kblue=numeric(v)                      knews1[j]=Q
yr=numeric(v)                         knews2[j]=Q
kmle=numeric(v)                       kml[j]=Q
kml=numeric(v)                         }
kr=numeric(v)                         yr[j]=sum(x)/r[j]
kmlq=numeric(v)                       if ( r[j]<n ){
Vmle=numeric(v)                       kr[j]=-Q/log(1-r[j]/n)
r=numeric(v)                          }else {
knews1=numeric(v)                     kr[j]=sum(x)/n}
knews2=numeric(v)                     if ( r[j]<n & r[j]>1){
for (j in 1:v) {                       Xr1=numeric((n-r[j]))
x<-rexp(n, rate=1/k)                  Xr2=numeric((n-r[j]))
yr<-x                                  kml[j]=(sum(x)+(n-r[j])*max(x))/r[j]
x<-sort(x,decreasing=FALSE)           Vmle[j]=1/r[j]*k^2
r[j]=n                                kmle[j]=(r[j]-1)/(sum(x)+(n-r[j])*max(x))
while (max(x)>Q)                       kmlq[j]=(sum(x)+(n-r[j])*Q)/r[j]
{                                       Xr1[1]=max(x)+rexp(1,rate=(n-
x<-x[1:length(x)-1]                  r[j])*kmle[j])
}                                       Xr2[1]=max(x)+rexp(1,rate=(n-
r[j]=length(x)                        r[j])*(1/kmlq[j]))
if ( r[j]==n ){                       if (Xr1[1]<Q){
knews1[j]=sum(x)/n                    Xr1[1]<-Q}
if (Xr2[1]<Q){
```

```

Xr2[1]<-Q}

if ( n-r[j]>1){
for ( i in 2:(n-r[j])){

Xr1[i]<-Xr1[i-1]+rexp(1, rate=(n-r[j]-
i+1)*kmlq[j])

Xr2[i]<-Xr2[i-1]+rexp(1, rate=(n-r[j]-
i+1)*(1/kmlq[j]))}

knews1[j]=(sum(x)+sum(Xr1))*(r[j]-
1)/(r[j]*(n-1))

knews2[j]=(sum(x)+sum(Xr2))/n

}else if (n-r[j]<=1){

knews1[j]=(sum(x)+Xr1)/(r[j]-1)/(r[j]*(n-
1))

knews2[j]=(sum(x)+Xr2)/n

}

else if (r[j]==1){

knews1[j]=(sum(x)+sum(Xr1))/n

knews2[j]=(sum(x)+sum(Xr2))/n

kml[j]=(sum(x)+(n-r[j])*max(x))/r[j]

}

else if (r[j]==n){

kml[j]=(sum(x))/n

knews1[j]=kml[j]

knews2[j]=kmlq[j]

}

if ( r[j]/n==0.8 ){

p=c(0.28,0.5022,0.6733,0.8)

pr=floor(n*p)+1

pr[4]=pr[4]-1

b=c(0.3157,0.2469,0.1864,0.3204)

S=0

Vblue=(k^2)/(0.7899*n)

for (t in 1:4)

```

```

{S=S+x[pr[t]]*b[t]}

kblue[j]=S

}else if ( r[j]/n==0.75){

p=c(0.2525,0.4585,0.6228,0.75)

pr=floor(n*p)+1

pr[4]=pr[4]-1

b=c(0.3070,0.2476,0.1945,0.4105)

S=0

Vblue=(k^2)/(0.7429*n)

for (t in 1:4)

{S=S+x[pr[t]]*b[t]}

kblue[j]=S

}else if (r[j]/n==0.65){ #65% censoring

p=c(0.2054,0.3807,0.5281,0.65)

pr=floor(n*p)+1

pr[4]=pr[4]-1

b=c(0.2935,0.2486,0.2074,0.6265)

S=0

Vblue=(k^2)/(0.6463*n)

for (t in 1:4)

{S=S+x[pr[t]]*b[t]}

kblue[j]=S

}else if (r[j]/n==0.70){

p=c(0.2279,0.4183,0.5747,0.70)

pr=floor(n*p)+1

pr[4]=pr[4]-1

b=c(0.2997,0.2481,0.2015,0.5115)

S=0

Vblue=(k^2)/(0.6949*n)

for (i in 1:4)

{S=S+x[pr[t]]*b[t]}

```

```

kblue[j]=S
}else if (r[j]/n==0.60){
p=c(0.1847,0.3451,0.4829,0.60)
pr=floor(n*p)+1
pr[4]=pr[4]-1
b=c(0.2881,0.2489,0.2126,0.7593)
S=0
Vblue=(k^2)/(0.5973*n)
for (t in 1:4)
{S=S+x[pr[t]]*b[t]}
kblue[j]=S
}else if (r[j]/n==0.85){
p=c(0.3115,0.5509,0.7272,0.85)
pr=floor(n*p)+1
pr[4]=pr[4]-1
b=c(0.3266,0.2457,0.1762,0.2386)
S=0
Vblue=(k^2)/(0.8356*n)
for (t in 1:4)
{S=S+x[pr[t]]*b[t]}
kblue[j]=S
}else if (r[j]/n==0.90){
p=c(0.3498,0.6068,0.7856,0.90)
pr=floor(n*p)+1
pr[4]=pr[4]-1
b=c(0.3413,0.2439,0.1628,0.1627)
S=0
Vblue=(k^2)/(0.8785*n)

```

```

for (t in 1:4)
{S=S+x[pr[t]]*b[t]}
kblue[j]=S
}else if (r[j]/n==0.95){
p=c(0.4014,0.6783,0.8536,0.95)
pr=floor(n*p)+1
pr[4]=pr[4]-1
b=c(0.3643,0.2407,0.1423,0.8894)
S=0
Vblue=(k^2)/(0.9151*n)
for (t in 1:4)
{S=S+x[pr[t]]*b[t]}
kblue[j]=S}
}
}
kbl=mean(kblue)
knew1=mean(knews1)
knew2=mean(knews2)
kml1=mean(kml)
ktr=mean(kr)
kmlq1=mean(kmlq)
vmle=mean(Vmle)
kml1
kmlq1
ktr
knew2
kbl

```


Βιβλιογραφία

- [1] Balakrishnan, N., and Cohen, A.c., (1991), “*Order Statistics and Inference*”, Academic Press, San Diego
- [2] Braden, D. J., Freimer, M. (1991). “Informational dynamics of censored observations” *Management Sci.* 37 p:1390–1404
- [3] Burnetas, A. (2006). “*Single Period Procurement Under Uncertainty, The Newsvendor Problem*”
- [4] Burnetas, A. and Smith, C. (1998). “*Adaptive ordering and pricing for perishable products*” Vol. 48, No. 3, p. 436–443
- [5] Chu, L.Y., Shanthikumar, J. G. Shen, Z.J.M. (2007). “*Solving Operational Statistics via a Bayesian Analysis*” p.1-15
- [6] Conrad, S.A., (1976). “*Sales Data and the Estimation of Demand*,” *Operational Research Quarterly*, 27, p:123-127
- [7] David, H. A., and Nagaraja, H. N. (2003). “*Order Statistics*”, Third edition, John Wiley & Sons
- [8] Halperin, M. (1966). “*Maximum likelihood estimators in truncated samples*” *Usaf school of aviation medicine* p:1717-1735
- [9] Liyanagea, L. and Shanthikumar, J.G. (2005). “*A practical inventory control policy using operational statistics*” *Operations Research Letters* 33 p:341–348
- [10] Lu, X. and Song, J.S. Zhu, K. (2008). “*Analysis of Perishable-Inventory Systems with Censored Demand Data*” Vol. 56, No. 4, p:1034–1038
- [11] Nahmias, S. (1994). “*Demand estimation in lost sales inventory systems*” *Naval Research Logistics* p:739-757.
- [12] Saleh, A. K. M. (1964). “*On the estimation of the parameters of the exponential distribution based on optimum order statistics in censored samples*” University of western Ontario

- [13] Saleh, A. K. M. Ehsanes (1966) “*Estimation of the parameters of the exponential distribution based on optimum order statistics in censored samples*” University of western Ontario
- [14] Sarhan, A. E. and Greenberg, B. G. (1957). “*Tables for Best Linear Estimates by Order Statistics of the Parameters of Single Exponential Distributions from Singly and Doubly Censored Samples*” Journal of the American Statistical Association, Vol. 52, No. 277 p:58-87
- [15] Silver, E.A and Pyke, D.F and Peterson, R. (1998) “*Inventory Management and Production Planning and Scheduling*”, third ed.
- [16] Stefanescu, C. (2009). “*Multivariate Demand: Modeling and Estimation from Censored Sales*” p:1-33