



**Tikrit Journal of Administrative
and Economics Sciences**

مجلة تكريت للعلوم الإدارية والاقتصادية

EISSN: 3006-9149

PISSN: 1813-1719



**Comparative between MLE and Pattern Search for Reliability
Estimation of the Teissier Distribution”**

Jasim Hassan Lazem*

Technical College of Management/the Middle Technical University

Keywords:

Teissier distribution, Reliability, MLE,
Pattern search

ARTICLE INFO

Article history:

Received	09 Oct. 2025
Received in revised form	30 Dec. 2025
Accepted	30 Dec. 2025
Available online	30 Jun. 2026

© THIS IS AN OPEN ACCESS ARTICLE UNDER
THE CC BY LICENSE

<http://creativecommons.org/licenses/by/4.0/>



***Corresponding author:**

Jasim Hassan Lazem

Technical College of Management/the
Middle Technical University



Abstract: The Teissier distribution is a modern distribution that is suitable for data used in machine learning and in production factors such as cotton, textile, furniture factories, and the like. It is also suitable for data used in medical fields, example that important devices associated with major operations like brain or heart surgeries, or cancer patients whose data requires monitoring.

In this research, the maximum likelihood method and the pattern search algorithm were used as methods to estimate the reliability function of this distribution. The results showed that the pattern search algorithm is suitable for medium and large samples, which confirms its superiority over the maximum likelihood estimation method at these sizes and for the default value $\beta = 0.2$. At the default value of $\beta=0.1$, the maximum likelihood estimation method outperformed the pattern search algorithm for all sample reservations.

مقارنة طريقة الامكان الاعظم مع طريقة تمييز النمط لتقدير معولية توزيع

Teissier

جاسم حسن لازم

الكلية التقنية الإدارية/الجامعة التقنية الوسطى

المستخلص

توزيع Teissier من التوزيعات الحديثة والتي تلائم البيانات الذي يكون استخدامها في معولية المكائن والتي تستخدم في المعامل الانتاجية كمعامل القطن الغزل والنسيج ومعامل الأثاث وإلى ما شابه ذلك. كما يكون ملائم للبيانات التي تستخدم في مجالات الطبية كالأجهزة المهمة المرتبطة بالعمليات الكبرى كالعمليات الدماغية أو العمليات القلبية أو مرضى السرطان والتي تحتاج بياناته إلى المراقبة. وفي هذا البحث تم استخدام طريقة الامكان الاعظم وخوارزمية pattern search كطرائق لتقدير دالة المعولية لهذا التوزيع. بينت النتائج بان خوارزمية pattern search تكون ملائمة للعينات المتوسطة والكبيرة مما يؤكد تفوقها على طريقة تقدير الامكان الاعظم في هذه الحجوم وللقيمة الافتراضية $\beta = 0.2$. أما عند القيمة الافتراضية $\beta = 0.1$ فقد تفوقت طريقة تقدير الامكان الاعظم على خوارزمية pattern search ولجميع حزم العينات.

الكلمات المفتاحية: توزيع Teissier. المعولية، طريقة الامكان الاعظم، طريقة تمييز النمط.

المقدمة Introduction

ظهرت المعولية في القرن الماضي كعلم يتعامل مع اعمار المعدات ولاسيما احتمالات البقاء ومتوسط الحياة واحتمال أن يكون عمر النظام أكبر من زمن معين. أن أبرز التطورات في هذا المجال كانت على يد علماء الرياضيات والاقتصاد ومن كان له اهتمام في العلوم البيئية والحياتية وكان الجزء المكمل لنجاح هذه التطورات هو التفاعل الملحوظ بين الاحصاء والمعولية. لقد ازداد الاهتمام بالمعولية بعد الانتشار الواسع للصناعة وازدياد التعقيدات الميكانيكية، الكهربائية والالكترونية في المعدات في القرن الماضي وكانت البحوث قبل عام 1940 تقتصر على السيطرة النوعية وصيانة المكائن ولم تشخص المعولية حينها على أنها حقل معين، ومع حلول الحرب العالمية الثانية وبازدياد المعدات الحربية المعقدة اصبح لحقل المعولية كيان خاص ومنها بدأت المعولية الحديثة. إن أكثر الأبحاث الإحصائية في المعولية ركزت على مدة الحياة لعدد من المعدات او الانظمة أو على فشل هذه المعدات والأنظمة أو عدمه خلال مدة زمنية، وإن تخمين المعولية لماكنة معينة هو من الأهمية العليا في محيط التقنية الحديثة وتطوراتها المستقبلية، لذلك فإن تخمين المعولية يعد هدفاً للكثير من الطرائق الاحصائية (Muth, 1977: 401–435).

إن توزيع Teissier هو توزيع من التوزيعات الحديثة والتي غالباً ما تستخدم في البيانات التي تحتاج إلى مراقبة Cencerd Data والتي تعني أن هناك عدم اليقين ببعض تلك البيانات والتي تتعلق الأنشطة اليومية مما يضطر إلى مراقبة تلك البيانات لكي نتخذ القرارات على ضوء ذلك كاتخاذ بعض الاجراءات لتعظيم المكاسب أو تقليل المخاطر حتى نحصل على حياة أفضل. (Singh et al., 2022: 71).

مشكلة البحث: تمتاز التوزيعات أو الدوال اللاخطية والتي لا يمكن تحويلها إلى دوال خطية بصعوبة تقدير معالمها وغالباً ما تحتاج إلى طرائق خاصة لتقدير معالمها ومن هنا جاءت الحاجة الى تقدير

معلومات توزيع Teissier وبالتالي تقدير معولية هذا التوزيع بطريقة تقدير الامكان الأعظم واحدی خوارزميات الذكاء الاصطناعي وهي خوارزمية pattern search.

تم تقديم توزيع Teissier لأول مرة عام 1934 في نمذجة تواتر الوفيات الناجمة عن الشخوخة فقط، أي أن الوفيات محمية من الحوادث والأمراض (Teissier, 1934: 237–284). بعد ذلك تم استخدام (Muth, 1977: 401–435) هذا التوزيع لتحليل المعولية. أعاد (Jodra et al., 2015: 291–310) اكتشاف توزيع Teissier واستنبط خصائصه الإحصائية. ولإضافة المزيد من المرونة إلى النموذج، قدم (Jodra, et al., 2017: 186–201) امتداداً ثنائي المعلومات لتوزيع Teissier يسمى توزيع Power Muth distribution. اقترح (Sharma, et al., 1997: 1–23) امتداداً آخر لتوزيع تيسيه ذي معاملين، يُسمى توزيع تيسيه الأسّي. ومؤخراً، اقترح (Singh, et al., 2020: 41) عائلة Teissier المعممة من التوزيعات، وأنتج توزيعات احتمالية مرنة متنوعة. اما (Singh, et al, 2022: 71) قام بتطوير طرائق تقدير باستخدام بيانات ضبابية خاضعة للرقابة. طريقة الامكان الأعظم (Singh, et al., 2022: 71) إن توزيع Teissier يملك دالة الكثافة الاحتمالية (p.d.f) كالآتي:

$$f(x; \mu, \beta) = \beta e(e^{\beta(x-\mu)})(e^{\beta(x-\mu)} - 1) \exp(-e^{\beta(x-\mu)}), x > \mu > 0, \beta > 0 \dots (1)$$

وله دالة (c.d.f) كالآتي:

$$F(x; \mu, \beta) = 1 - \exp(\beta(x - \mu) - e^{\beta(x-\mu)} + 1), x > \mu > 0, \beta > 0 \dots (2)$$

$$R(x; \mu, \beta) = \exp(\beta(x - \mu) - e^{\beta(x-\mu)} + 1), x > \mu > 0, \beta > 0 \dots (3)$$

لذلك فإن دالة الامكان لهذا التوزيع كالآتي:

$$l(\beta, \mu; x) = \prod_{i=1}^n f(x; \mu, \beta) \dots (4)$$

$$= \prod_{i=1}^n \beta e(e^{\beta(x_i-\mu)})(e^{\beta(x_i-\mu)} - 1) \exp(-e^{\beta(x_i-\mu)})$$

$$= \beta^n e^{n\beta \sum_{i=1}^n (x_i-\mu)} \prod_{i=1}^n [e^{\beta(x_i-\mu)} - 1] e^{-\sum_{i=1}^n e^{\beta(x_i-\mu)}}$$

$$\log l(\beta, \mu; x) = n \log \beta + n \log e + \beta \sum_{i=1}^n (x_i - \mu) + \sum_{i=1}^n \log(e^{\beta(x_i-\mu)} - 1) - \sum_{i=1}^n e^{\beta(x_i-\mu)} \dots (5)$$

وباشتقاق المعادلة 3 وبالنسبة إلى المعلمتين β و μ نحصل على المعادلتين الآتيتين:

$$\frac{\partial \log l(\beta, \mu; x)}{\partial \beta} = \frac{n}{\beta} + \sum_{i=1}^n (x_i - \mu) + \sum_{i=1}^n \frac{x_i e^{\beta(x_i-\mu)}}{e^{\beta(x_i-\mu)} - 1} - \sum_{i=1}^n x_i e^{\beta(x_i-\mu)} \dots (6)$$

$$\frac{\partial \log l(\beta, \mu; x)}{\partial \mu} = -n\beta - \sum_{i=1}^n \frac{\beta e^{\beta(x_i-\mu)}}{e^{\beta(x_i-\mu)} - 1} + \beta \sum_{i=1}^n e^{\beta(x_i-\mu)} \dots (7)$$

وبالحل المعادلتين بطريقة نيوتن رافسن Newton-Raphson نحصل على المعلومات التقديرية بهذه الطريقة

وفيما يأتي خطوات طريقة نيوتن رافسن Newton-Raphson بصورة عامة

$$1-\hat{f}(x_0) = \frac{f(x_0)}{\Delta x} \longrightarrow \Delta x = \frac{f(x_0)}{\hat{f}(x_0)}$$

$$2-x_0 - x_1 = \frac{f(x_0)}{\hat{f}(x_0)}$$

$$3-x_1 = x_0 - \frac{f(x_0)}{f'(x_0)}$$

$$4-x_{i+1} = x_i - \frac{f(x_i)}{f'(x_i)}$$

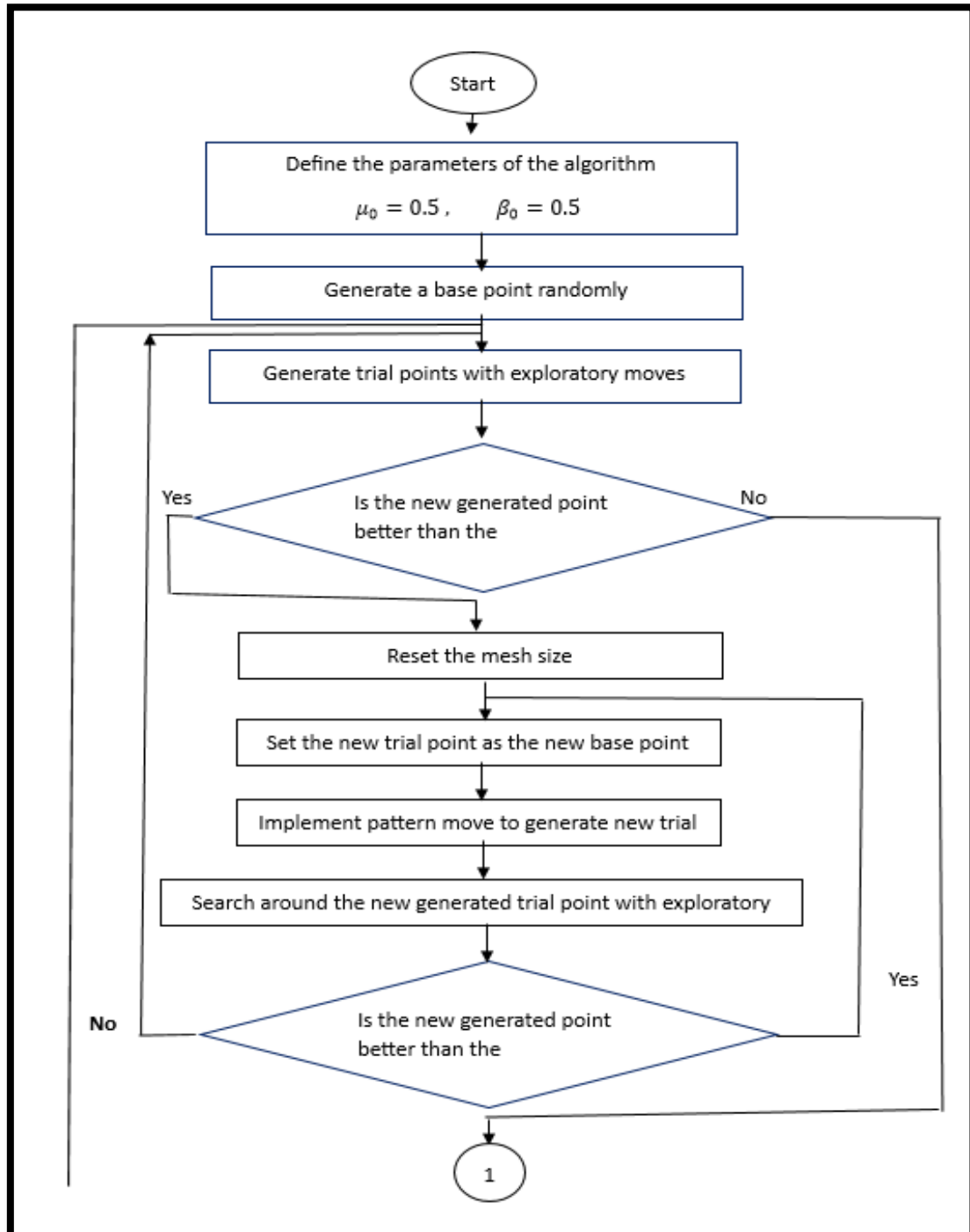
2. خوارزمية بحث الانماط (Al-Sumait, et al, 2007: 720–730)، (Mahapatra, et al, 2014: 1177–1188)، (Bozorg, et al, 2013: 2515–2529)

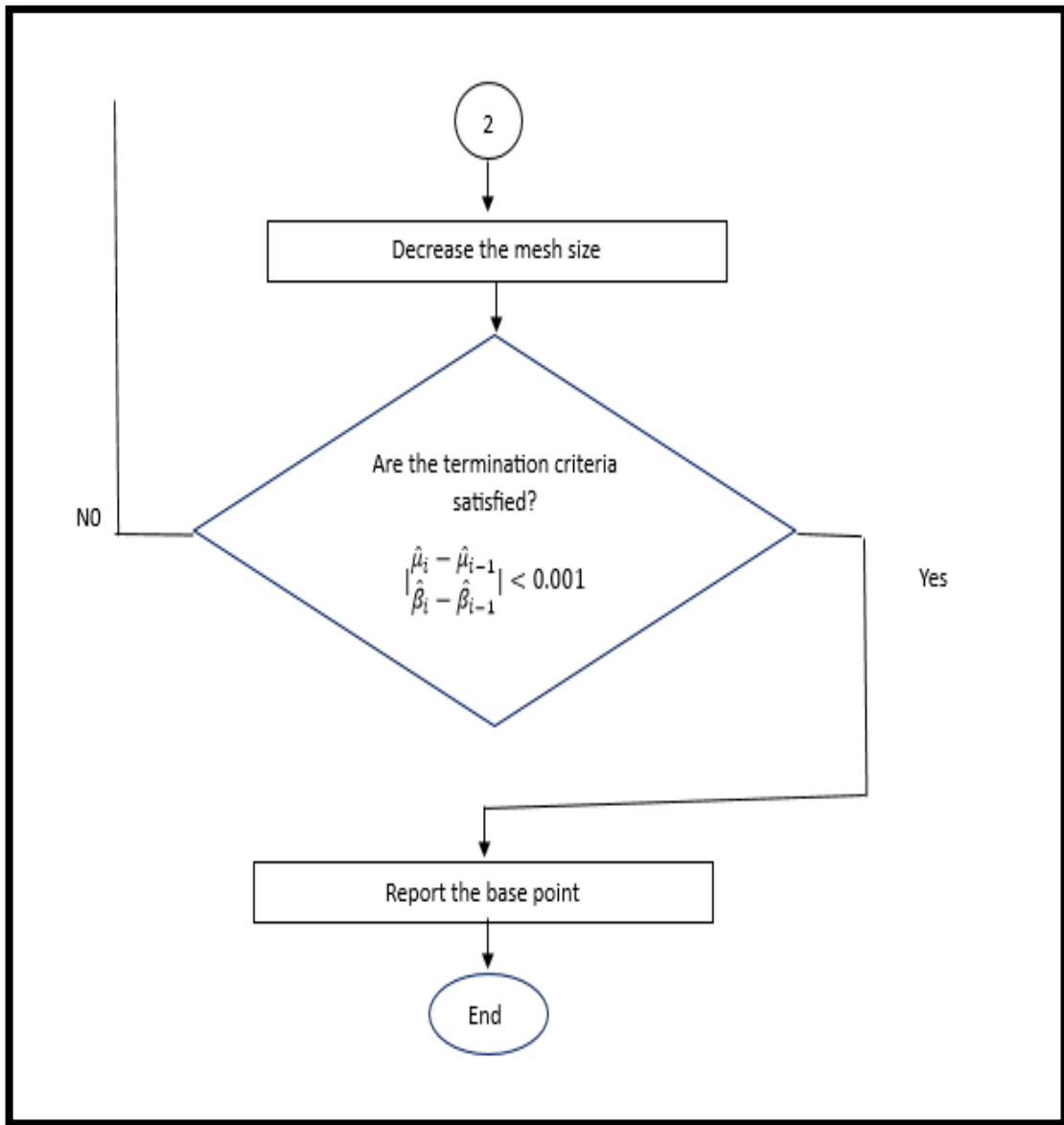
تصنف هذه الخوارزمية كطريقة حل مباشر تحل مجموعة متنوعة من المسائل العددية، مع التركيز على استخدام استراتيجيات بسيطة تجعلها أكثر ملاءمة للتطبيق في الحواسيب الحديثة مقارنة بالطرق التقليدية (مثل البرمجة الخطية وغير الخطية). يشير مصطلح "البحث المباشر" إلى الفحص المتسلسل للحلول التجريبية. تقارن هذه الخوارزمية كل حل تجريبي بأفضل حل تم الحصول عليه سابقاً، وتحدد نتيجة المقارنة ما سيكون عليه الحل التجريبي التالي كما إنها تستخدم تقنيات يطلق عليها استراتيجيات بحث مباشرة. تتمتع هذه التقنيات بخصائص تميزها عن الطرق الكلاسيكية، وقد حلت مشاكل لم تتمكن الطرق الكلاسيكية من حلها. كما أن هذه التقنيات تتقارب بشكل أسرع في إيجاد حلول لبعض المشاكل مقارنة بالطرق التقليدية. تعتمد تقنيات هذه الخوارزمية على عمليات حسابية متطابقة متكررة ذات منطق بسيط يمكن ترميزها بسهولة للحسابات الحاسوبية. تتقارب تقنيات البحث المباشر مع الحلول شبه المثالية، وهو ما يعد أيضاً سمة من سمات الخوارزميات الفوقية والتطور.

تختار تقنيات البحث المباشر بشكل عشوائي نقطة B وتسميها نقطة الأساس. يتم اختيار نقطة ثانية، P1، بشكل عشوائي، وإذا كانت أفضل من B، فإنها تحل محل نقطة الأساس؛ وإذا لم تكن كذلك، تظل B نقطة الأساس. تستمر هذه العملية مع مقارنة كل نقطة جديدة مختارة عشوائياً بنقطة الأساس الحالية. تُحدد "استراتيجية" اختيار نقاط الاختبار الجديدة بمجموعة من "الحالات" التي تُشكل ذاكرة الخوارزمية. عدد هذه الحالات محدود. إن نوع الاستراتيجية المتبعة لاختيار نقاط جديدة يتم تحديدها من خلال جوانب مختلفة من المشكلة، بما في ذلك هيكل مساحة القرار الخاصة بالمشكلة. تتضمن الاستراتيجية اختيار نقطة الأساس الأولية، وقواعد الانتقال بين الحالات، وقواعد اختيار نقاط الاختبار كدالة للحالة الحالية ونقطة الأساس. يشير البحث المباشر إلى نقطة الاختبار بعدها حركة أو خطوة من نقطة الأساس. تعد الخطوة ناجحة إذا كانت نقطة الاختبار أفضل من نقطة الأساس، وإلا فهي فاشلة. حلول مسألة التحسين المحسوبة بواسطة خوارزمية PS هي نقاط في فضاء ذي أبعاد N، إذ يُشير N إلى عدد متغيرات القرار. تشير نقاط التجربة إلى حلول جديدة. تسمى عملية الانتقال من نقطة معينة إلى نقطة اختبار بالتحرك. قد تكون الخطوة ناجحة إذا كانت نقطة الاختبار أفضل من نقطة الأساس؛ وإلا فهي فاشلة. وبعبارة أخرى، فإن فشل أو نجاح نقاط الاختبار يؤثر على اتجاه وطول خطوات الحركات في المراحل التالية.

تقوم هذه الخوارزمية بإنشاء سلسلة من الحلول التي تنتج شبكة حول الحل الأول وتقترب من الحل الأمثل. يتم إنشاء الحل الأولي بشكل عشوائي ويُعرف باسم نقطة الأساس. بعد ذلك، يتم توليد نقاط تجريبية (حلول جديدة). هناك نمطان لتوليد نقاط تجريبية. النمط الأول هو خطوة استكشافية تهدف إلى اكتساب المعرفة حول مساحة القرار. تتبع هذه المعرفة من نجاح أو فشل التحركات الاستكشافية دون مراعاة أي تقييم كمي لقيم دوال fitness. وهذا يعني أن التحركات الاستكشافية تحدد الاتجاه المحتمل للتحرك الناجح. أما النوع الثاني فهو حركة نمطية مصممة لاستخدام المعلومات التي تم الحصول عليها من خلال التحركات الاستكشافية لإيجاد الحل الأمثل. بعد توليد حلول جديدة

باستخدام الحركات الاستكشافية، تحسب خوارزمية PS دالة fitness عند نقاط الشبكة وتختار دالة تكون قيمتها أفضل من قيمة fitness للحل الأول. يتم نقل نقطة البحث إلى الحل الجديد إذا كانت هناك نقطة (حل) ذات دالة fitness تم الحصول عليها بشكل أفضل بين الحلول المولدة. في هذه المرحلة، يتم تطبيق معامل التمدد لتوليد حلول جديدة. ومن ثم فإن حجم الشبكة الجديدة أكبر من الحجم السابق. على النقيض من ذلك، إذا لم يكن هناك حل أفضل في الحلول الناتجة، يتم تطبيق معامل الانكماش ويقتصر حجم الشبكة على قيمة أصغر وفيما يأتي المخطط للخوارزمية P





شكل (1): يبين المخطط الانسيابي لخوارزمية PS

1-2. توليد الحل الأولي Generating an Initial Solution

في هذه الخوارزمية يمثل كل حل ممكن نقطة في مساحة القرار، وكل النقاط الموجودة في هذه المساحة تمثل حلولاً ممكنة لمشكلة التحسين. كل حل يتكون من مجموعة من متغيرات القرار x_i لذلك الحل الواحد يتخذ على هيئة متجه صفي $1 \times N$ وعلى النحو الآتي:

$$\text{Point} = (x_1, x_2, \dots, x_N)$$

حيث $X = \text{حل لمشكلة التحسين}$ ، $x_i = \text{متغير القرار للحل } X$ ؛ و $N = \text{عدد متغيرات القرار}$. تبدأ الخوارزمية بحل أولي يتم إنشاؤه عشوائياً ويُعرف باسم نقطة الأساس. بعد ذلك، يتم إنشاء نقاط تجريبية (حلول) حول النقطة الأساسية.

2-2. إنشاء حلول تجريبية Generating Trial Solutions: الحلول التجريبية هي حلول جديدة لمشكلة التحسين ولتوليد هذه الحلول لابد من معرفة ان هناك نمطان لتوليد الحلول التجريبية، النمط

الأول هو خطوة استكشافية تهدف إلى اكتساب المعرفة حول مساحة القرار. أما النوع الثاني فهو حركة نمطية مصممة لاستخدام المعلومات التي تم الحصول عليها من خلال التحركات الاستكشافية لإيجاد الحل الأمثل.

2-2-1. النمط الاستكشافي Exploratory Move: من البديهي أن كل خطوة استكشافية هدفها الحصول على معلومات حول مساحة القرار لمشكلة التحسين وهذه المعلومات تستند على نجاح أو فشل الحركات الاستكشافية بغض النظر عن قيمة دالة fitness، وعلى هذا الأساس في كل حركة استكشافية يتم تغيير قيمة احداثي واحد وبعد ذلك يتم معرفة ما هو مقدار تأثير هذا التغيير على دالة fitness بعد النقل لذلك تكون الحركة ناجحة إذا كانت قيمة دالة fitness المحسوبة حديثاً أفضل من قيمة دالة fitness قبل النقل والا فالخطوة تكون فاشلة. هناك نمطان لتوليد الحلول من خلال التحرك الاستكشافي، وهو أولاً البحث النمطي المعمم (GPS) وثانياً البحث المباشر التكيفي الشبكي (MADS).

في هذه الخوارزمية يتم حساب دالة fitness لجميع الحلول الجديدة بعد إنشاء نقاط جديدة إما باستخدام نظام تحديد المواقع العالمي (GPS) أو نظام MADS، ثم تقوم بتقييم الحل الذي يتمتع بأفضل قيمة دالة fitness بينها. تُجرى مقارنة بين أفضل حل جديد ونقطة الأساس. في حال وجود حل ذي قيمة fitness أفضل في الحلول المؤلدة، تُنقل نقطة البحث إلى هذه النقطة الجديدة. في هذه المرحلة، يُستخدم مُعامل التوسيع لتوليد حلول جديدة. وبالتالي، يكون حجم الشبكة الجديد أكبر من السابق. في المقابل، إذا لم يكن هناك حل أفضل في الحلول المؤلدة، يُطبّق مُعامل الانكماش ويُقلّل حجم الشبكة إلى قيمة أصغر. بعد خطوة استكشافية ناجحة، يتم تنفيذ حركة النمط لتعزيز البحث في اتجاه التحسين المحتمل.

2-2-2. التحرك النمطي Pattern Move: يتم تصميم حركة النمط لاستخدام المعرفة المكتسبة في الحركات الاستكشافية. عندما تصل الحركة الاستكشافية إلى نقطة أفضل، تصبح النقطة الجديدة المولدة هي نقطة الأساس الجديدة. يتم إنشاء نقطة اختبار جديدة بناءً على النقطة السابقة ويتم حساب نقطة الأساس الحالية على النحو الآتي:

$$X^{new} = \bar{X} + \alpha \cdot (X - \bar{X}) \quad \dots (8)$$

إذ إن:

$$X^{new} = \text{حل التجربة الجديد}$$

$$\bar{X} = \text{نقطة الأساس السابقة}$$

$$X = \text{نقطة الأساس الحالية، و } \alpha = \text{عامل تسارع إيجابي.}$$

حركة النمط من نقطة أساس معينة تُكرر الحركات المُجمعة من نقطة الأساس السابقة. إن السبب وراء هذا النوع من التحركات هو الافتراض بأن أي مجموعة من التحركات الناجحة في الماضي من المرجح أن تثبت نجاحها مرة أخرى. يتبع كل حركة نمطية على الفور سلسلة من الحركات الاستكشافية التي تعمل على مراجعة النمط بشكل مستمر ويمكن أن تؤدي إلى تحسين حل التجربة الجديد. إذا كان حل التجربة الجديد الذي تم إنشاؤه أفضل من نقطة الأساس الحالية، فإنه يصبح نقطة الأساس الجديدة ويتم تنفيذ نقل النمط مرة أخرى لإنشاء نقطة تجربة جديدة. بخلاف ذلك، لن يتم تغيير نقطة الأساس الحالية، وسيتم تنفيذ التحركات الاستكشافية فقط حول نقطة الأساس.

يحصل نمط التحرك أولاً على حل تجريبي مؤقت ثم يجد حلاً تجريبياً من خلال بحث استكشافي. تتكرر الحركات النمطية طالما كانت ناجحة (أي أنها تعمل على تحسين دالة fitness) وعادةً ما تصبح خطوات أطول وأطول. بمجرد فشل حركة النمط، يتم سحب حركة النمط وتعود الخوارزمية إلى البحث الاستكشافي حول أفضل نقطة تم حسابها حتى الآن.

3-2. تحديث حجم الشبكة Updating the Mesh Size: إذا لم يكن هناك حل أفضل بين الحلول الناتجة عن الحركة الاستكشافية، يتم تطبيق معامل الانكماش ويتم تقليل حجم الشبكة (μ) إلى قيمة أصغر. بالنسبة لأي قيمة معينة لحجم الشبكة، تصل الخوارزمية إلى طريق مسدود عندما تفشل جميع التحركات الاستكشافية من نقطة الأساس. في هذه الحالة، من الضروري تقليل حجم الشبكة (μ) لمواصلة البحث. إن حجم الانخفاض كافٍ للسماح بإنشاء نمط جديد. ومع ذلك، فإن التخفيض الكبير في حجم الشبكة يؤدي إلى إبطاء عملية البحث. ولهذا السبب، عندما تكون جميع نقاط الاختبار أسوأ من نقطة الأساس، يتم تقليل حجم الشبكة على النحو الآتي:

$$\mu^{\text{new}} = \mu - \delta \quad \dots (9)$$

إذاً μ^{new} = حجم شبكي جديد، و δ = خطوة متناقصة لحجم الشبكة، وهي معلمة محددة من قبل المستخدم للخوارزمية. إذا كانت هناك نقطة (حل) ذات قيمة fitness أفضل بين الحلول المولدة، يتم نقل نقطة البحث إلى هذه النقطة الجديدة. في هذه المرحلة، يتم تطبيق معامل التوسع لتوليد حلول جديدة. ومن ثم فإن حجم الشبكة الجديد أكبر من الحجم السابق. عندما تصل الحركة الاستكشافية إلى نقطة أفضل، يتم إعادة تعيين حجم الشبكة (μ) إلى القيمة الأولية كما يأتي:

$$\mu^{\text{new}} = \mu_0$$

إذاً μ_0 = القيمة الأولية لـ μ ، والتي يتم تحديدها من قبل المحلل كمعلمة للخوارزمية.

4-2. معايير الإنهاء Termination Criteria: يحدد معيار الإنهاء متى يجب إيقاف الخوارزمية. يعد اختيار معيار إنهاء جيد أمراً ضرورياً لتجنب التوقف المبكر الذي يتم من خلاله حساب حل دون المستوى الأمثل أو الحل الأمثل محلياً. ومن ناحية أخرى، من المستحسن تجنب الحسابات غير الضرورية التي تتجاوز الحل الذي لا يمكن تحسينه بمجرد التوصل إليه. المعيار الشائع لإنهاء PS هو حجم الشبكة. يتم إنهاء البحث بشكل نهائي عندما يصبح حجم الشبكة صغيراً بدرجة كافية لضمان التقريب الأمثل. يُعدّ تحديد عدد التكرارات، أو وقت التشغيل، ومراقبة تحسين الحل في التكرارات المتتالية معايير إنهاء أخرى يُمكن تطبيقها باستخدام خوارزمية PS

5-2. المعلومات المعرفة من قبل المستخدم في خوارزمية

PS User-Defined Parameters of the PS

قيمة عامل التسارع (α)، وحجم الشبكة الأولي (μ_0)، وخطوة التناقص في حجم الشبكة (δ)، ومعيار الإنهاء هي معلومات محددة من قبل المستخدم لـ PS. يعتمد الاختيار الجيد للمعلومات على مساحة القرار لمشكلة معينة، وعادةً ما يكون الإعداد الأمثل للمعلومات لمشكلة واحدة ذا فائدة محدودة للمشكلات الأخرى. الطريقة المعقولة للعثور على القيم المناسبة للمعلومات الخوارزمية هي إجراء تحليل الحساسية. يتكون هذا من تجربة مجموعات متعددة من المعلومات التي يتم تشغيل الخوارزمية بها. يتم مقارنة النتائج من التركيبات المختلفة، ويختار المحلل مجموعة المعلومات التي تحقق أفضل نتائج التحسين.

الجانب التجريبي

تم توليد البيانات باستخدام دالة c.d.f. المعكوسة ومن خلال دالة Lambert_W function وكما يأتي: [9]

$$F(x) = 1 - \exp(\beta(x - \mu) - e^{\beta(x-\mu)} + 1) \dots\dots\dots (10)$$

وافترض أن $F(x) = u$

$$\exp(\beta(x - \mu) - e^{\beta(x-\mu)} + 1) = 1 - u \dots\dots\dots (11)$$

نضع $y = e^{\beta(x-\mu)}$ بعد اعادة الترتيب والرفع إلى الأس نحصل على

$$ye^{-y} = \frac{1-u}{e}$$

نفترض $Z = -y$ لذلك $ze^z = -\frac{1-u}{e}$ لذا

$$z = W(-\frac{1-u}{e})$$

وعليه

$$y = W(-\frac{1-u}{e})$$

وبسبب $y = e^{\beta(x-\mu)} \geq 1$ لكل $x \geq \mu$ يجب أن نأخذ W_{-1} لدالة Lambert W والتي موجودة كدالة جاهزة ضمن برنامج ماتلاب على $(-\frac{1}{e}, 0)$ لذلك لتوليد متغير سيكون كالآتي:

$$t = \mu + \left(\frac{1}{\beta}\right) + \left(\frac{1}{\beta}\right) \log(1 - u) - \left(\frac{1}{\beta}\right) \cdot \text{lambertw}(-1, \frac{u-1}{e^1}) \dots (12)$$

$$t = \mu + \left(\frac{1}{\beta}\right) \log(-W_{-1}(-\frac{1-u}{e^1}))$$

وبهدف تقدير معلمات توزيع Teissier تم تثبيت $\mu = 1$ أما المعلمة β فتم افتراض قيم لها وهي 0.1, 0.2 وتم التعويض بالمعادلة رقم 7, 6 ومن ثم ايجاد دالة المعولية لهذا التوزيع أما الخوارزمية pattern search فتم استخدام دالة لا معلمية وكما يأتي:

$$\hat{F} - \text{nonparametric} = \frac{i}{(n+1)} \dots\dots\dots (13)$$

$$F(x; \mu, \beta) = 1 - \exp(\beta(x - \mu) - e^{\beta(x-\mu)} + 1) \dots\dots\dots (14)$$

اما دالة الهدف (Fitness function):

$$\min_{\beta, \mu} = Q = \sum_{i=1}^n [\hat{F} - F(x)]^2 \dots\dots\dots (15)$$

وبغية الوصول إلى نتائج قريبة من البيانات الحقيقية فقد تم تكرار التجربة ل(10000) مرة وفيما يأتي نتائج المحاكاة وقد تم اعداد الجداول من قبل الباحث باستخدام برنامج Matlab:

جدول (1): يبين متوسط مربعات الخطأ للمعولية المقدرة لكل الطريقتين المستخدمتين في البحث وبحجم 15

μ	β	T	$MSE_{\hat{R}_{MLE}}$	$MSE_{\hat{R}_{PS}}$	الافضل
1	0.1	1.2500	0.0115	0.0494	$MSE_{\hat{R}_{MLE}}$
		1.4500	0.0113	0.0447	$MSE_{\hat{R}_{MLE}}$
		1.6500	0.0111	0.0402	$MSE_{\hat{R}_{MLE}}$
		1.8500	0.0110	0.0360	$MSE_{\hat{R}_{MLE}}$
	0.2	1.2500	0.0202	0.0263	$MSE_{\hat{R}_{MLE}}$
		1.4500	0.0198	0.0209	$MSE_{\hat{R}_{MLE}}$
		1.6500	0.0192	0.0164	$MSE_{\hat{R}_{PS}}$
		1.8500	0.0181	0.0129	$MSE_{\hat{R}_{PS}}$

إذ إن: $MSE_{\hat{R}_{MLE}}$ هو متوسط مربعات الخطأ للمعولية المقدرة بطريقة الامكان الأعظم.

$MSE_{\hat{R}_{PS}}$ هو متوسط مربعات الخطأ للمعولية المقدرة بطريقة خوارزمية تمييز النمط.

نلاحظ من الجدول أعلاه حصلت طريقة الامكان الأعظم على الافضلية عند $\beta = 0.1$ ، أما طريقة PS فقد حصلت على الأفضلية على بقية النتائج عند قيمة $\beta = 0.2$ وعند قيمة $T = 1.65, 1.85$ وذلك بحصولها على أقل متوسط مربعات الخطأ.

جدول (2): يبين متوسط مربعات الخطأ للمعولية المقدرة لكل الطريقتين المستخدمتين في البحث وبحجم 25

μ	β	T	$MSE_{\hat{R}_{MLE}}$	$MSE_{\hat{R}_{PS}}$	الافضل
1	0.1	1.2500	0.0092	0.0198	$MSE_{\hat{R}_{MLE}}$
		1.4500	0.0090	0.0167	$MSE_{\hat{R}_{MLE}}$
		1.6500	0.0088	0.0138	$MSE_{\hat{R}_{MLE}}$
		1.8500	0.0087	0.0113	$MSE_{\hat{R}_{MLE}}$
	0.2	1.2500	0.0091	0.0107	$MSE_{\hat{R}_{MLE}}$
		1.4500	0.0088	0.0073	$MSE_{\hat{R}_{MLE}}$
		1.6500	0.0086	0.0048	$MSE_{\hat{R}_{PS}}$
		1.8500	0.0084	0.0031	$MSE_{\hat{R}_{PS}}$

نلاحظ من الجدول أعلاه بأن طريقة الامكان الأعظم حصلت على الافضلية عند $\beta = 0.1$ ،

كما إن طريقة PS حصلت على الأفضلية على بقية النتائج عند قيمة $\beta = 0.2$ وعند قيمة $T = 1.65, 1.85$ وذلك بحصولها على أقل متوسط مربعات الخطأ.

جدول (3): يبين متوسط مربعات الخطأ للمعولية المقدرة لكل الطريقتين المستخدمتين في البحث
وبحجم 50

μ	β	T	$MSE_{\hat{R}_{MLE}}$	$MSE_{\hat{R}_{PS}}$	الافضل
1	0.1	1.2500	0.0077	0.0121	$MSE_{\hat{R}_{MLE}}$
		1.4500	0.0073	0.0097	$MSE_{\hat{R}_{MLE}}$
		1.6500	0.0070	0.0076	$MSE_{\hat{R}_{MLE}}$
		1.8500	0.0067	0.0059	$MSE_{\hat{R}_{MLE}}$
	0.2	1.2500	0.0060	0.0066	$MSE_{\hat{R}_{MLE}}$
		1.4500	0.0056	0.0037	$MSE_{\hat{R}_{PS}}$
		1.6500	0.0053	0.0019	$MSE_{\hat{R}_{PS}}$
		1.8500	0.0052	0.0009	$MSE_{\hat{R}_{PS}}$

نلاحظ من الجدول أعلاه بأن طريقة الامكان الأعظم حصلت على الأفضلية عند $\beta = 0.1$ ،
كما إن طريقة PS حصلت على الأفضلية على بقية النتائج عند قيمة $\beta = 0.2$ وعند قيمة $T = 1.45, 1.65, 1.85$ وذلك بحصولها على أقل متوسط مربعات الخطأ.

جدول (4): يبين متوسط مربعات الخطأ للمعولية المقدرة لكل الطريقتين المستخدمتين في البحث
وبحجم 75

μ	β	T	$MSE_{\hat{R}_{MLE}}$	$MSE_{\hat{R}_{PS}}$	الافضل
1	0.1	1.2500	0.0068	0.0106	$MSE_{\hat{R}_{MLE}}$
		1.4500	0.0064	0.0084	$MSE_{\hat{R}_{MLE}}$
		1.6500	0.0061	0.0064	$MSE_{\hat{R}_{MLE}}$
		1.8500	0.0057	0.0048	$MSE_{\hat{R}_{PS}}$
	0.2	1.2500	0.0048	0.0064	$MSE_{\hat{R}_{MLE}}$
		1.4500	0.0046	0.0034	$MSE_{\hat{R}_{PS}}$
		1.6500	0.0046	0.0015	$MSE_{\hat{R}_{PS}}$
		1.8500	0.0047	0.0006	$MSE_{\hat{R}_{PS}}$

نلاحظ من الجدول أعلاه أن طريقة الامكان الاعظم حصلت على الأفضلية عند $\beta = 0.1$ ،
كما ان طريقة PS حصلت على الأفضلية على بقية النتائج عند قيمة $\beta = 0.2$ وعند قيمة $T = 1.45, 1.65, 1.85$ وذلك بحصولها على أقل متوسط مربعات الخطأ.

جدول (5): يبين متوسط مربعات الخطأ للمعولية المقدرة لكل الطريقتين المستخدمتين في البحث وبحجم 110

μ	β	T	$MSE_{\hat{R}_{MLE}}$	$MSE_{\hat{R}_{PS}}$	الأفضل
1	0.1	1.2500	0.0045	0.0084	$MSE_{\hat{R}_{MLE}}$
		1.4500	0.0045	0.0065	$MSE_{\hat{R}_{MLE}}$
		1.6500	0.0046	0.0049	$MSE_{\hat{R}_{MLE}}$
		1.8500	0.0047	0.0035	$MSE_{\hat{R}_{PS}}$
	0.2	1.2500	0.0046	0.0056	$MSE_{\hat{R}_{MLE}}$
		1.4500	0.0042	0.0027	$MSE_{\hat{R}_{PS}}$
		1.6500	0.0041	0.0011	$MSE_{\hat{R}_{PS}}$
		1.8500	0.0042	0.0004	$MSE_{\hat{R}_{PS}}$

نلاحظ من الجدول أعلاه أن طريقة الامكان الأعظم حصلت على الأفضلية عند $\beta = 0.1$ ، كما إن طريقة PS حصلت على الأفضلية على بقية النتائج عند قيمة $\beta = 0.2$ وعند قيمة $T = 1.45, 1.65, 1.85$ وذلك بحصولها على أقل متوسط مربعات الخطأ.

الاستنتاجات والتوصيات

أولاً. الاستنتاجات:

1. تفوقت طريقة تقدير الامكان الأعظم على خوارزمية pattern search ولجميع قيم T وللقيمة الافتراضية للمعلمة $\beta = 0.1$ وبثبات المعلمة $\mu = 1$ ولجميع أحجام العينات المستخدمة في البحث.
2. تفوقت طريقة تقدير الامكان الأعظم على خوارزمية pattern search عند قيمة $T = 1.2500$ الافتراضية للمعلمة $\beta = 0.2$ وبثبات المعلمة $\mu = 1$ ولجميع أحجام العينات المستخدمة في البحث.
3. كما تفوقت طريقة تقدير الامكان الأعظم على خوارزمية pattern search عند قيمة $T = 1.4500$ وعند القيمة الافتراضية للمعلمة $\beta = 0.2$ وبثبات المعلمة $\mu = 1$ وللحجم الصغيرة (15 , 25).
4. تفوقت خوارزمية pattern search على طريقة تقدير الامكان الأعظم عند قيمة $T = 1.6500$, وعند القيمة الافتراضية للمعلمة $\beta = 0.2$ وبثبات المعلمة $\mu = 1$ ولجميع أحجام العينات المستخدمة في البحث.
5. تفوقت خوارزمية pattern search على طريقة تقدير الامكان الأعظم عند قيمة $T = 1.4500$ وعند القيمة الافتراضية للمعلمة $\beta = 0.2$ وبثبات المعلمة $\mu = 1$ ولحجم العينات الكبيرة (50 , 75 , 110).
6. تتفوق خوارزمية pattern search على طريقة تقدير الامكان الأعظم فيه الحجم المتوسطة والكبيرة للقيمة الافتراضية $\beta = 0.2$ وبثبات المعلمة $\mu = 1$ وللقيمة T التي تتجاوز 1.2500.

ثانياً. التوصيات:

1. يوصي الباحث بتقدير معلمات هذا التوزيع باستخدام طرائق أخرى كطريقة تقدير بيز.
2. يوصي الباحث بعدم استخدام خوارزمية PS في حالة العينات الصغيرة وعندما $\beta = 0.1$.

3. استخدام طرائق لامعلمية لتقدير معلمات هذا التوزيع.
المصادر

1. Al-Sumait, J. S., AL-Othman, A. K., and Sykulski, J. K. (2007). "Application of pattern search method to power system valve-point economic load dispatch." International Journal of Electrical Power & Energy Systems, 29(10), 720–730.
2. Bozorg-Haddad, O., Tabari, M. M. R., Fallah-Mehdipour, E., and Mariño, M. A.(2013). "Groundwater model calibration by meta-heuristic algorithms."Water Resources Management, 27(7), 2515–2529.
3. Jodra, Pedro and Gomez, Hector Wladimir and Jimenez-Gamero, Maria Dolores and Alba-Fernandez, Maria Virtudes (2017). The Power Muth Distribution. Mathematical Modelling and Analysis, 22:186–201.
4. Jodra, Pedro and Jimenez-Gamero, Maria Dolores and Alba-Fernandez, Maria Virtudes(2015). On the Muth distribution. Mathematical Modelling and Analysis, 20:291–310.
5. Mahapatra, S., Panda, S., and Swain, S. C. (2014). "A hybrid firefly algorithm and pattern search technique for SSSC based power oscillation damping controller design." Ain Shams Engineering Journal, 5(4), 1177–1188.
6. Muth, Eginhard J. (1977). Reliability models with positive memory derived from the mean residual life function. The Theory and Applications of Reliability, 2:401–435.
7. Sharma, Vikas Kumar and Singh, Sudhanshu V. and Shekhawat, Komal.(1997). Exponentiated Teissier distribution with increasing, decreasing and bathtub hazard functions. Journal of Applied Statistics, 1–23.
8. Singh, Sudhanshu V. and Elgarhy, Mohammed and Ahmad, Zubair and Sharma, Vikas Kumar and Hamedani, Gholamhossein G. (2020). New Class of Probability Distributions Arising from Teissier Distribution. Mathematical Modeling, Computational Intelligence Techniques and Renewable Energy: Proceedings of the First International Conference, MMCITRE 2020, 41.
9. Singh, Sudhanshu v., Sharma, Vikas K. and Singh , Sanjay k.(2022). " Inferences for two parameter Teissier distribution in case of fuzzy progressively censored data" RT&A, No 4 (71) Vol. 17,
10. Teissier, G. (1934). Recherches sur le vieillissement et sur les lois de la mortalité. Annales de physiologie et de physicochimie biologique, 10:237–284

الملحق
البرنامج الرئيسي

```

%% Dr. Jasim
clc
clear all
n=175; %n=15,25,50,75
theta=0.2; %theta=0.1 theta=0.2
mu=1; %mu=1,1.5
%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%
for q=1:10000
u=rand(1,n);
t=real(mu-(1/theta)+(1/theta).*log(1-u)-(1/theta).*lambertw(-1,((u-1)./exp(1))));
%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%% mle
x0=[theta+0.3; mu+0.3];
mle = fsolve(@(x) jasim_mle(x,t,n),x0);
theta_mle(q)=mle(1);
mu_mle(q)=mle(2);
T=[(mu+0.25):0.2:(2*mu)];
R_real(q,:)=exp(theta.*(T-mu))-exp(theta.*(T-mu))+1;
R_mle(q,:)=exp(theta_mle(q).*(T-mu_mle(q)))-exp(theta_mle(q).*(T-mu_mle(q)))+1;
%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%% pattern search
ts=sort(t);
pi=(1:n)./(n+1);
X0=[theta+0.3; mu+0.3]; % Starting point
LB = [-10 -10]; % Lower bound
UB = [10 10]; % Upper bound
PS= abs(patternsearch(@(x)objective_Jasim(x,ts,pi),X0,[],[],[],[],LB,UB));
theta_ps(q)=PS(1);
mu_ps(q)=PS(2);
R_ps(q,:)=exp(theta_ps(q).*(T-mu_ps(q)))-exp(theta_ps(q).*(T-mu_ps(q)))+1;
end
Mse_theta=[mean((theta-theta_mle).^2) mean((theta-theta_ps).^2)];%%%
mse theta_mle theta_ps

```

```

Mse_mu=[mean((mu-mu_mle).^2) mean((mu-mu_ps).^2)]; %%%% mse
mle_mu ps_mu
Mse=[Mse_theta;Mse_mu] % mtrix mse theta mse mu
R_reliability=[T' mean(R_real)' mean(R_mle)' mean(R_ps)']
Mse_R=[T' mean((R_mle-R_real).^2)' mean((R_ps-R_real).^2)']
plot(T,      mean(R_real),'-b',T,mean(R_mle),'-r',T,      mean(R_ps),'-
*c','LineWidth',1.2)
title('Plot of MLE and SP Reliability Finctions')
xlabel('timex')
ylabel('Estimated Reliability')
legend('Real','MLE','PS')

```