

油气储产量增长 预测与评价软件系统

许晓宏 编著

YOUQI CHUCHANLIANG
ZENGZHANG YUCE YU
PINGJIA RUANJIAN XITONG



中国地质大学出版社有限责任公司
ZHONGGUO DIZHI DAXUE CHUBANSHE YOUXIAN ZEREN GONGSI

油气储产量增长预测与 评价软件系统

YOUQI CHUCHANLIANG ZENGZHANG YUCE YU PINGJIA RUANJIAN XITONG

许晓宏 编著



中国地质大学出版社

ZHONGGUO DIZHI DAXUE CHUBANSHE

图书在版编目(CIP)数据

油气储产量增长预测与评价软件系统/许晓宏编著. —武汉:中国地质大学出版社,2013.12
ISBN 978-7-5625-2658-2

- I. 油…
- II. 许…
- III. ①油气藏-储量-预测-应用软件-研究②油气藏-储量-评估-应用软件-研究
- IV. TE155-39

中国版本图书馆 CIP 数据核字(2011)第 075110 号

油气储产量增长预测与评价软件系统		许晓宏 编著
责任编辑:王凤林 张晓红		责任校对:张咏梅
出版发行:中国地质大学出版社(武汉市洪山区鲁磨路 388 号)		邮政编码:430074
电 话:(027)67883511	传真:67883580	E-mail:cbb@cug.edu.cn
经 销:全国新华书店	http://www.cugp.cug.edu.cn	
开本:787 毫米×1 092 毫米 1/16	字数:190 千字	印张:7.375
版次:2013 年 12 月第 1 版	印次:2013 年 12 月第 1 次印刷	
印刷:武汉珞南印务有限公司	印数:1—500 册	
ISBN 978-7-5625-2658-2		定价:28.00 元

如有印装质量问题请与印刷厂联系调换

前 言

加强油气资源科学合理的调查评价,加强油气资源的规划、管理、保护和合理利用,对提高经济发展的保障能力具有十分重要的意义。

油气矿产的资源储量、产量的增长幅度等,均受制于其资源潜力、品质和探明程度。对其资源潜力的评价,则与勘探程度、研究深度紧密相关。随着勘查程度的提高、技术方法的进步、认识的变化,不同时期进行的矿产资源评价、资源潜力预测和数值的估算也会有许多差异。我国油气资源评价方法,也正处在逐步与世界相衔接的过程中,还存在研究方法、评价标准以及统计等方面的差别。

我国目前已经完成了第三次全国油气资源评价工作,以前曾在 1981—1987 年、1992—1994 年组织开展过两次全国性的常规油气资源评价工作。

第一次油气资源评价工作是在石油行业内部组织开展的,在对国外油气资源评价理论和方法进行充分分析的基础上,对国内已掌握的 50 余种评价方法进行了分析、筛选,结合我国石油地质条件、评价技术和分析化验技术特点,研究并形成了远景资源评价的方法系列。

第二次油气资源评价工作则是在体制变化的情况下,分为陆上和海上两部分进行,基本上采用了比较统一的技术方法和软件评价系统,尤其以盆地模拟为主要方法。

前两次的评价均不包括对非常规油气资源的评价。

第三次的全国油气资源评价工作将建立起国家油气资源统一的评价体系、方法、规范、流程以及评价参数等,还将建立基于地理信息系统(GIS)技术基础之上的信息应用平台,并研究进行评价方法软件的开发和数据交换接口的攻关工作。

根据国家对新一轮油气资源的评价要求,要建立包括统计法、类比法和成因法三大类资源评价方法体系和关键参数体系标准及评价规范,强调了统计法和类比法的应用,特别在高勘探程度的地区突出统计法,强调方法和参数体系的统一。

以往三大石油公司的评价方法和参数体系是各不相同的。中石油强调类比

目 录

上篇 方法原理

第一章 油田规模序列法模型	(3)
1. 基本原理	(3)
2. 计算步骤	(4)
3. 实验示例	(6)
4. 问题讨论	(8)
第二章 胡伯特(Hubbert)模型	(10)
1. 基本原理	(10)
2. 计算步骤	(11)
3. 线性拟合算法示例	(12)
4. 问题讨论	(12)
第三章 指数增长模型(HCZ 模型)	(14)
1. 基本原理	(14)
2. 计算步骤	(14)
3. 实验示例	(15)
4. 问题讨论	(17)
第四章 翁(Weng)旋回模型	(18)
1. 基本原理	(18)
2. 计算步骤	(20)
3. 实验示例	(21)
4. 问题讨论	(24)
第五章 Weibull 模型	(25)
1. 基本原理	(25)
2. 计算步骤	(26)
3. 应用示例	(27)
4. 问题讨论	(30)
第六章 龚帕兹(Gompertz)模型	(31)
1. 基本原理	(31)

上篇

方法原理

第一章 油田规模序列法模型

1. 基本原理

“油田规模”(Oilfield Size)是指油气田的最终可采储量。如果某个含油气区经过详细勘探后,发现了全部油气田,并且查清了每个油田的最终可采储量,那么,按最终可采储量由大到小进行排列,所得到的顺序称为油田规模序列。

国内外许多含油气区的统计资料说明,当一个含油气区的最大油田及一系列中小油田被发现后,如果以油田规模为纵坐标,以油田规模的序列号为横坐标,在双对数坐标纸上展点作图,大致可得一条直线,如图 1-1 所示。

根据这一规律,可以在探区的早期或中期勘探阶段,由已发现油气田的油气储量去预测尚未发现的油气田储量以及全探区总的石油储量。

美国学者齐波夫(G P Zipf)于 1949 年在他所著的《人类行为与最小省力原则》一书中提出一种规律,这个规律可表述如下:将一组离散型随机变量,由大到小进行排列,如果最大的数值是第二数值的两倍,是第三大数值的三倍,……,依次类推,则称这组离散类型随机变量服从齐波夫定律。

进入 20 世纪 70 年代以后,随着计算机技术的推广应用,齐波夫定律逐渐为人们所重视。自 1972 年以来,相继有人研究齐波夫定律的应用,发现世界各国城市人口的分布,英语词汇相对的使用频数等与人类活动有关的各种社会现象大都接近齐波夫定律。

近年来,国内有人应用齐波夫定律研究勘探地区的金属矿床及油气田的规律序列,借以预测尚未发现的矿产储量或油气田储量。

实际上,齐波夫定律是帕累托(Pareto)于 1927 年所提出的定律特例。帕累托定律可

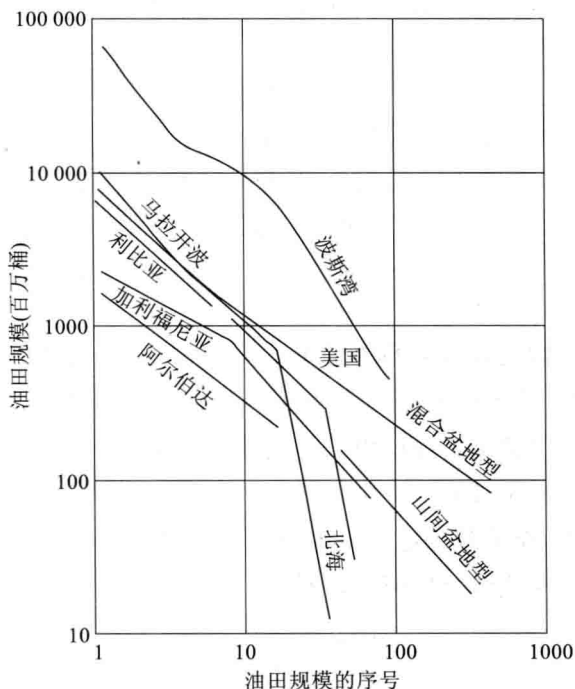


图 1-1 世界主要含油气盆地的油田规模序列

表述为如下关系式：

$$\frac{Q_m}{Q_n} = \left(\frac{n}{m}\right)^K \quad (1-1)$$

式中： Q_m 为序号等于 m 的随机变量的数值； Q_n 为序号等于 n 的随机变量的数值； K 为实数； m, n 为 $1, 2, \dots$ 整数序列中的任一数值，但 $m \neq n$ 。

当式(1-1)中的 $K = 1$ 时，则为齐波夫定律，即：

$$\frac{Q_m}{Q_n} = \frac{n}{m} \quad (1-2)$$

$$m \times Q_m = n \times Q_n \quad (1-3)$$

一个含油气地区内一组油气田的石油储量属于离散型随机变量，当最大的(第一号)油田被发现，其石油储量为 $Q_{\max}$ ，若油田规模序列符合齐波夫定律时，则有：

$$Q_{\max} = n \times Q_n \quad (1-4)$$

$$Q_n = \frac{Q_{\max}}{n} \quad (1-5)$$

假如油气区共有 l 个油气田，则全探区的石油总储量 SQ 等于

$$SQ = \sum_{i=1}^l \left(\frac{Q_{\max}}{i}\right) \quad (i = 1, 2, \dots, l) \quad (1-6)$$

对式(1-3)的两边取对数，则有：

$$\lg(m \times Q_m) = \lg(n \times Q_n)$$

$$\lg Q_m - \lg Q_n = -(\lg m - \lg n)$$

$$\frac{\lg Q_m - \lg Q_n}{\lg m - \lg n} = -1 \quad (1-7)$$

因而在双对数坐标纸上，以油田的石油储量 Q_i 为纵坐标，以油田的序号 i 为横坐标作图，则数据点的连线为斜率等于 -1 的直线。上面的式(1-4)、式(1-5)、式(1-6)、式(1-7)就是目前国内外一些学者所说的预测油田储量的齐波夫定律不同形式的表达式。

但是，图 1-1 中世界主要含油气区的统计资料说明，多数含油气区并不符合齐波夫定律，而是符合适应范围更广的帕累托定律。

对式(1-1)两边取对数，则有：

$$\lg Q_m - \lg Q_n = -K(\lg m - \lg n)$$

$$\frac{\lg Q_m - \lg Q_n}{\lg m - \lg n} = -K \quad (1-8)$$

因而在双对数坐标纸上作图，则数据点的连线为斜率等于 $-K$ 的直线，这样便与图 1-1 中的统计规律相符合了。所以，应当认为油田规模序列的分布规律服从帕累托定律，而服从齐波夫定律仅是其中的特例。

2. 计算步骤

1) 由熟悉含油气地区情况的地质家商定油田规模序列的系数 K 。这里可以借鉴与含

油气地区在地质条件上相似的探区资料。

如果确定系数 K 有困难,可令 $K = \operatorname{tg}\theta$, θ 的角度值限定在 $25^\circ \sim 65^\circ$ 范围内,并把这一范围分为若干个间隔值,进行多个油田规模序列的拟合计算。

例如,假如取角度值步长为 5° 时,则有如下 9 个间隔值:

$$\operatorname{tg}25^\circ = 0.4663, \quad \operatorname{tg}30^\circ = 0.5774, \quad \operatorname{tg}35^\circ = 0.7002,$$

$$\operatorname{tg}40^\circ = 0.8391, \quad \operatorname{tg}45^\circ = 1.0000, \quad \operatorname{tg}50^\circ = 1.1918,$$

$$\operatorname{tg}55^\circ = 1.4218, \quad \operatorname{tg}60^\circ = 1.7321, \quad \operatorname{tg}65^\circ = 2.1445.$$

其中 $\operatorname{tg}45^\circ = 1.00$ 时为齐波夫定律。

2) 把探区中已发现的 l 个油田,按储量 $Q_i (i=1, 2, \dots, l)$ 由大到小进行排列,选择储量最大的油田作为推算点。

3) 如果探区中已发现的 l 个油田的储量为 $Q_1, Q_2, \dots, Q_l$, 则以推算点 Q_1 去除 Q_i , 并求出其值的 K 次方根,得到如下序列 A_i :

$$A_i = \sqrt[K]{\frac{Q_i}{Q_1}} \quad (i = 1, 2, \dots, l)$$

4) 序列 A_i (表示下列矩阵第 i 行任意元素) 乘以正整数 $n=1, 2, \dots$, 当 $A_i n$ 的值接近正整数 $1, 2, \dots, m, \dots$ 时,记入下列矩阵:

$$\begin{bmatrix} A_{11}n & A_{12}n & \cdots & A_{1l}n \\ A_{21}n & A_{22}n & \cdots & A_{2l}n \\ \cdots & \cdots & \cdots & \cdots \\ A_{m1}n & A_{m2}n & \cdots & A_{ml}n \end{bmatrix} \begin{matrix} \approx 1 \\ \approx 2 \\ \cdots \\ \approx m \end{matrix}$$

计算矩阵中各行的标准差 σ :

$$\sigma = \sqrt{\frac{1}{l} \sum_{i=1}^l (A_i n - \overline{An})^2}$$

$$\overline{An} = \frac{1}{l} \sum_{i=1}^l A_i n$$

当矩阵中第 m 行的标准差 σ 小于给定误差 EP 时,即 $\sigma < EP$ 。一般情况下可令 $EP = 0.05 \sim 0.1$ 。此时有:

$$A_i n = \sqrt[K]{\frac{Q_i}{Q_1}} \cdot n \approx m$$

$$\sqrt[K]{\frac{Q_i}{Q_1}} \approx \frac{m}{n}$$

$$\text{即有: } \frac{Q_i}{Q_1} \approx \left(\frac{m}{n}\right)^K$$

由于 $A_i n$ 已接近正整数 m ,所以在给定的误差范围内符合巴内托定律。因而可把第 m 行作为探区油田规模的预测模型序列。

5) 把预测模型序列 $A_i n$ 中的每个数值除以 A_i , 则可得到探区中已发现油田 Q_1, Q_2 ,

..., Q_l 在预测的油田规模序列中的序号 n (秩)。

$$n = \frac{A_i n}{A_i}$$

而 n 则为 $1, 2, \dots$ 中的某一整数。 m 值则为 Q_1 (已发现油田中的最大油田储量) 在预测的油田规模序列中的序号 (秩)。

6) 探区中已发现的油田储量 Q_i ($i=1, 2, \dots, l$) 乘以预测序号 n 的 k 次方幂, 则为预测的最大油田 (第一号油田) 储量 $\hat{Q}_{\max}$ 。这里以所有已发现油田预测的最大油田储量的平均值作为预测的探区中的最大油田储量, 即:

$$\hat{Q}_{\max} = \frac{1}{l} \sum_{i=1}^l Q_i n^K \quad (1-9)$$

7) 预测的最大油田储量 $\hat{Q}_{\max}$ 除以 $1^K, 2^K, \dots$, 则得到探区中预测的油田规模序列 $\hat{Q}_i$:

$$\hat{Q}_i = \hat{Q}_{\max} / i^K \quad (i = 1, 2, \dots, P) \quad (1-10)$$

当油田规模序列的某一储量 $\hat{Q}_{P+1} < Q_{\min}$ 时截断序列。 $Q_{\min}$ 为人为规定的在当时技术水平下的最小经济油田的储量值。

8) 预测全探区总的石油储量 (或资源量) $S\hat{Q}$:

$$S\hat{Q} = \sum_{i=1}^P \hat{Q}_i = \sum_{i=1}^P (\hat{Q}_{\max} / i^K) \quad (i = 1, 2, \dots, P) \quad (1-11)$$

9) 按 $K = \text{tg}25^\circ \sim \text{tg}65^\circ$ 范围内的步长计算 S 个预测的油田规模序列 $\hat{Q}_j$ ($j=1, 2, \dots, S$), 再计算每个预测序列中已发现油田的实际储量与所预测的储量之间的标准差 σ :

$$\sigma = \sqrt{\frac{1}{l} \sum_{i=1}^l (Q_i - \hat{Q}_{ji})^2} \quad (j = 1, 2, \dots, S) \quad (i = 1, 2, \dots, l) \quad (1-12)$$

式中: Q_i 为探区中已发现油田的实际储量; $\hat{Q}_{ji}$ 为第 j 个预测序列中已发现的第 i 个油田的预测储量。最后选定 σ 值最小的序列作为预测的油田规模序列。

10) 上述计算结果仅是从数学运算中得出的预测值, 是否符合实际地质情况, 还需要由熟悉探区情况的地质家们讨论商榷。

3. 实验示例

某探区经钻探已发现 4 个油田, 石油地质储量分别为 $149.143 \times 10^4 \text{t}$ 、 $61.567 \times 10^4 \text{t}$ 、 $34.375 \times 10^4 \text{t}$ 、 $27.277 \times 10^4 \text{t}$ 。

1) 由于该区是个新探区, 所以很难确定油田规模序列的系数 K 。因而需要通过多次拟合计算, 经计算确定 $K = \text{tg}60^\circ = 1.732$ 。

2) 把已发现的 4 个油田, 按储量由大到小进行排列:

$$\begin{aligned} Q_1 &= 149.143 \times 10^4 \text{t}, & Q_2 &= 61.567 \times 10^4 \text{t}, \\ Q_3 &= 34.0375 \times 10^4 \text{t}, & Q_4 &= 27.277 \times 10^4 \text{t}. \end{aligned}$$

以最大油田储量 $149.143 \times 10^4 \text{t}$ 作为推算点。

3)用推算点 Q_1 去除 Q_1, Q_2, Q_3, Q_4 , 并求所得之商的 K 次方根, 得到 A_i 序列:

$$A_1 = \sqrt[K]{\frac{Q_1}{Q_1}} = 1.0, \quad A_2 = \sqrt[K]{\frac{Q_2}{Q_1}} = 0.6,$$

$$A_3 = \sqrt[K]{\frac{Q_3}{Q_1}} = 0.4286, \quad A_4 = \sqrt[K]{\frac{Q_4}{Q_1}} = 0.375.$$

4)把序列 $A_i (i=1, 2, 3, 4)$ 乘以正整数 $1, 2, \dots$ 中的某一值, 当乘积值接近正整数时则得到 $A_i n$ 序列, 记入下列矩阵:

$$\begin{bmatrix} 1.0 & 1.199 & 0.857 & 1.125 \\ 2.0 & 1.799 & 2.143 & 1.875 \\ 3.0 & 2.999 & 2.999 & 2.999 \end{bmatrix} \begin{matrix} \approx 1 \\ \approx 2 \\ \approx 3 \end{matrix}$$

计算到第三行时, 标准差 σ 等于 0.000 63, 即可认为已符合巴内托定律, 因而可把第三行作为油田规模的预测模型序列。

5)预测模型序列 $A_i n$ 被 A_i 序列中的对应值去除, 则得到已发现油田 Q_1, Q_2, Q_3, Q_4 在预测的油田规模序列中的序号:

$$\frac{3.0}{1.0} = 3, \frac{2.999}{0.6} = 5, \frac{2.999}{0.4286} = 7, \frac{2.999}{0.375} = 8$$

即已发现的 4 个油田 Q_1, Q_2, Q_3, Q_4 在预测的油田规模序列中的序号为 3, 5, 7, 8。

6)已发现的 4 个油田储量 Q_1, Q_2, Q_3, Q_4 分别乘以预测序号的 K 次方幂 $3^K, 5^K, 7^K, 8^K$ 则得到预测的最大油田储量 $\hat{Q}_{\max}$ 。

$$149.143 \times 3^K = 1000.0026, \quad 61.567 \times 5^K = 999.999,$$

$$34.375 \times 7^K = 999.9895, \quad 27.277 \times 8^K = 999.9870.$$

这 4 个预测值的平均值 $999.99 \times 10^4 \text{ t}$ 就是探区中最大油田储量 $\hat{Q}_{\max}$ 的预测值。

7) $\hat{Q}_{\max}$ 分别除以 $1^K, 2^K, \dots$, 则得到预测的探区油田规模序列 $\hat{Q}_1, \hat{Q}_2, \dots$, 这里暂定最小经济油田的储量值 $\hat{Q}_{\min} = 10 \times 10^4 \text{ t}$, 则得如下预测结果:

$$\begin{aligned} \hat{Q}_1 &= 999.995, & \hat{Q}_2 &= 301.002, & \hat{Q}_3 &= 149.142, \\ \hat{Q}_4 &= 90.615, & \hat{Q}_5 &= 61.567, & \hat{Q}_6 &= 44.895, \\ \hat{Q}_7 &= 34.375, & \hat{Q}_8 &= 27.227, & \hat{Q}_9 &= 22.243, \\ \hat{Q}_{10} &= 18.533, & \hat{Q}_{11} &= 15.713, & \hat{Q}_{12} &= 13.515, \\ \hat{Q}_{13} &= 11.765, & \hat{Q}_{14} &= 10.348. \end{aligned}$$

预测结果在双对数坐标纸上展点连线成一直线, $\theta = 60^\circ$, $\text{tg}60^\circ = 1.732$, 如图 1-2 所示。

8)全探区的总石油储量 $S\hat{Q}$ 为:

$$S\hat{Q} = \sum_{i=1}^{14} \hat{Q}_i = 1801.004 \times 10^4 (\text{t})$$

9)为预测本探区的油田规模序列, 共做了 9 次拟合计算。当 $\theta = 60^\circ$, 即 $\text{tg}60^\circ = 1.732$ 时, 已发现油田的实际储量与所预测的储量之间的标准差 $\sigma = 0.000 63$, 所以被选定为预

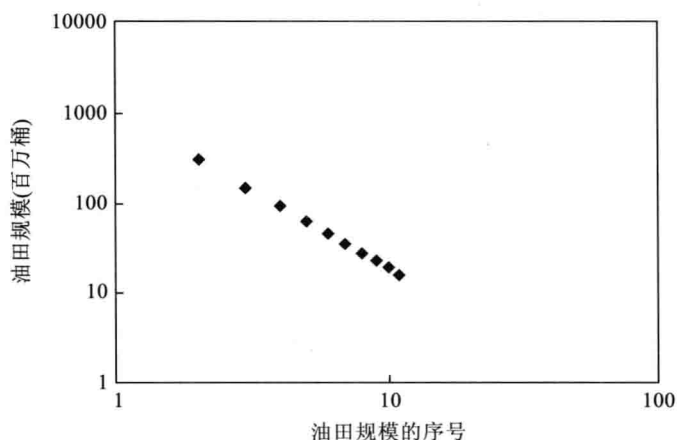


图 1-2 某探区油田规模预测序列

测序列。

10)这一计算结果,经有关地质人员分析,认为比较符合探区的地质情况。

4. 问题讨论

油田规律序列法的实质是根据已发现的油田储量,应用帕累托定律预测一个含油气地区中尚未发现的油气田储量(或资源量)以及全区总的石油储量(或资源量)的一种外推预测方法。

然而,直至目前为止尚不能从油气形成的地质理论上圆满解释油气规模序列的形成机理。可以认为由于任何地质过程都受概率法则支配,所以对于一个含油气地区的油田规模序列分布,暂且可以从统计规律方面去理解。

油田规模序列法适用于一个完整的、独立的石油地质体系,即该地质体系内的油气生成、运移、聚集以及尔后的地质变迁都是在同一石油地质演化历史条件下形成的,也就是说目前所要预测的含油气地区的油气田(或油气藏)的分布规律具有统一的形成原因。

1)近年来,国内不少地质研究单位应用齐波夫定律预测各种金属矿床及油田规模序列分布规律,但作者认为齐波夫定律的适应范围较窄,而帕累托定律的适应范围较广,根据国内外主要含油气区的统计资料表明,式(1-8)中的系数 K 值的变化范围在 0.5~2.0 之间。这一情况说明,石油地质问题的复杂性导致了油田规模序列分布的多样性。所以,系数 K 等于 1 的齐波夫定律,只能作为油田规律序列分布的特例,应当以帕累托定律作为金属矿床及油田规模序列分布的理论预测公式。

2)地质学研究的对象,包括各种地质过程及观测结果,都普遍地受概率法则支配或影响,因而有人认为地质现象可视为随机事件,地质观测结果具有随机变量性质。苏联的著名数学地质学者维斯捷列乌斯于 1977 年曾指出:“地质对象是由一些单个单元联合起来的,这种联合是遵循概率法则的。”

实际上,地质学研究的问题都具有时间长、空间广、因素多的特点,而且又是在多种地质因素互相交织中发生发展的,因而既受确定性法则支配,又在很多方面受到偶然性因素的影响,即受概率法则支配。

所以,尽管对油田规模序列法的预测结果目前尚不能从地质理论上做出满意的解释,但是应当承认,世界各主要含油气地区的油田规模序列普遍地服从帕累托定律。因而使用油田规模序列法预测探区中尚未发现的石油储量(或资源量)是可行的。

3)油田规模序列法的预测结果带有多解性。在系数 K 不清楚的情况下,要经过多次拟合计算才能选出已发现油田储量 Q_i 与预测值 $\hat{Q}_i$ 之间的标准差最小的油田规模序列。而且这个预测结果还要经过熟悉探区情况的地质人员分析判断是否符合实际情况。

4)当探区已发现一批油田时,应从中选择出油田储量数据可靠的数据作为预测的依据,储量数据不可靠的绝对不要使用。例如在某油田正打探边井,有可能增加储量,油田的储量参数还需要验证的情况下,这个油田的储量依据不能作为油田规模序列法的预测依据,否则将会得出错误的预测结果。

5)图 1-1 中的世界主要含油气盆地的实际资料说明,有些含油气盆地如波斯湾、北海等油区规模序列在双对数坐标上不是直线,而是折线。这意味着这些含油气盆地可能存在多期成油过程,也即可能存在多个油田规模序列时,应当对油田规模序列进行筛分,分解出成因不同或成油期不同的多个相互独立的油田规模序列后再进行处理。

第二章 胡伯特(Hubbert)模型

胡伯特(M King Hubbert),地球物理学家,美国得克萨斯州人士,1943 年加入壳牌石油集团的休斯敦公司。他在地球物理界有许多杰出贡献,研究范围甚广,从基础科学到石油、天然气储备,在美国和全世界享有盛名。1949 年他使用统计和物理方法计算出石油、天然气的全球储量,然后提出了尖锐的增长消耗量公式,预言不久的将来人类将无石油可用,直接触发了环保运动的兴起。1956 年他预言,石油产量顶峰将会在美国发生,时间大约是 1966—1971 年,后来事实果真如此,之后其威名如日中天,被邀请加入多个政治集团的班底,成为肯尼迪总统的座上宾。亦因为此,20 世纪 70 年代后美国人大力开发中东地区,使阿拉伯世界经济腾飞,阿拉伯控制着石油,使美国呼吸不畅,于是美国援助犹太人回归故土,意图制衡中东局势,中东战争爆发。

胡伯特有一个著名的 50%法则,即当石油产量达到其储量(可采储量)的 50%时,油田压力下降,开采成本上升,石油产量便见顶(习惯上称为“胡伯特顶峰 Hubbert Peak”),油价即升易跌难,全球油田在 2006 年前后进入中年。胡伯特顶峰很容易让人想起苏俄经济学家康德拉基耶夫的“康氏周期”模型,其实“分形之父”曼德勃罗很早就说:相似永远存在,只有大小上的不同,并无结构上的区别。石油在 20 世纪 50~70 年代完成了康氏周期所称的繁荣期,70 年代至今是个“康氏高原”(形态上的不同处是:康氏周期高原期不创新高),此时期石油行业蓬勃发展,2006 年后进入高原尾声,产油量应当心突然暴跌。

1. 基本原理

胡伯特模型实际上就是数学上的逻辑斯蒂(Logistic)模型,Logistic 模型属于“S”形曲线中最著名的一种,早先主要用于描述动植物的自然生长过程,故又称生长曲线。生长过程的基本特点是开始增长较慢,而在以后的某一范围内迅速增长,达到一定的限度后增长又缓慢下来,曲线呈拉长的“S”,故称“S”形曲线。它最早由比利时数学家 P F Verhulst 于 1838 年导出,但直至 20 世纪 20 年代才被生物学家及统计学家 R Pearl 和 L J Reed 重新发现,并逐渐被人们所重视。胡伯特将其用于石油产量的预测取得了极大的成功。

逻辑斯蒂(Logistic)模型可写成如下通用形式

$$x = \frac{K}{1 + Ae^{-Bt}} \quad (2-1)$$

式中: t 为时间间隔(或称时间变程), $t = y - y_0$, y_0 为起始时刻, y 为截止时刻; K 为 x 变化的上限; A, B 为拟合系数。

B 为成长率,如果 $B < 0$ 时,则该模型可以表示一个体系的晚期变化过程,即:

$$\lim_{t \rightarrow \infty} x \rightarrow 0$$

的过程。因此,可以用这一模型预测一个探区晚期勘探阶段的累积储量缓慢增长的过程,特别是适合预测一个油田的开采末期的石油产量变化过程。

如果 $B > 0$ 时,则该模型可以表示一个体系发展到最后的极限过程。即:

$$\lim_{t \rightarrow \infty} x \rightarrow C$$

的过程,因此,可用这一模型预测一个油田的含水率变化过程。

用逻辑斯蒂模型预测油田综合含水率时, x 表示含水率,它的极限为:

$$\lim_{t \rightarrow \infty} x \rightarrow 1$$

所以,式(2-1)中的 $K \leq 1$ 。

曲线在 $x = \ln A / B$ 时有一拐点,这时 $x = C/2$,恰好是终极量 C 的一半,称为参数 x 的高峰期,也是量变最快的时期。在拐点左侧,曲线凹向上,表示变化速率由小趋大;在拐点右侧,曲线凸向上,变化速率由大趋小。

2. 计算步骤

1) 线性拟合算法(事先已知模型中的 K 值)。

为确定式(2-1)中的拟合系数 A, B ,可作如下变换:

$$\frac{K}{x} = 1 + Ae^{-Bx}$$

$$\frac{K}{x} - 1 = Ae^{-Bx}$$

$$\ln\left(\frac{K}{x} - 1\right) = \ln A - Bx = \ln A - B(y - y_0)$$

$$\text{令: } U = \ln\left(\frac{K}{x} - 1\right), a = \ln A$$

则有: $U = a + (-Bx)$ 。

至此,可用一元线性回归求出 a, B 拟合系数,回代到式(2-1)即可预测油田的含水率变化。

2) 牛顿迭代算法(事先不知模型中的 K 值)。

根据 K 是生长过程中的终极量的特点,可由两种方法进行 K 值初值的估计:①如果 x 是累积频率,则显然 $K = 100\%$;②如果 x 是生长量或繁殖量,则可取 3 对观察值 (x_1, y_1) 、 (x_2, y_2) 和 (x_3, y_3) ,分别代入后得到联立方程:

$$\begin{cases} y_1 = K/(1 + ae^{bx_1}) \\ y_2 = K/(1 + ae^{-bx_2}) \\ y_3 = K/(1 + ae^{-bx_3}) \end{cases}$$

若令 $x_2 = (x_1 + x_3)/2$,则可解得:

$$K = \frac{y_2^2(y_1 + y_3) - 2y_1y_2y_3}{y_2^2 - y_1y_3}$$

有了 K 的初值估值后,即可采用牛顿迭代算法进行处理。

牛顿迭代算法的设计思想是将非线性求解的过程逐步线性化。其迭代函数为:

$$\varphi(x) = x - f(x)/f'(x)$$

牛顿迭代算法的突出优点是收敛速度快,但它有个明显的缺点,即每步的迭代需要提供导数值,如果函数比较复杂,则处理就不太方便,此时可采用弦截法进行处理,其迭代函数为:

$$x_{k+1} = x_k - f(x_k)(x_k - x_0)/[f(x_k) - f(x_0)]$$

3. 线性拟合算法示例

示例 巴夫雷油田是苏联较早采用边外注水、保持油层压力进行开发的油田之一,储层为岩性均匀的砂岩和粉砂岩,渗透率较高,为 $600 \times 10^{-2} \mu\text{m}^2$,孔隙度 20.6%,油层厚度 11m,埋藏深度 1750m。油田面积 118km²,地质储量 $1.1 \times 10^8 \text{t}$,设计采收率为 65%,可采储量为 $6500 \times 10^4 \text{t}$ 。巴夫雷油田从 1948 年开始开发,1974 年底采出程度已达 51.6%,1980 年的油田产量只有 $50 \times 10^4 \text{t}$ 左右,目前已进入油田开发晚期。

经计算,巴夫雷油田年度综合含水率(图 2-1)的逻辑斯蒂模型的表达式为:

$$x = 0.95/(1.0 + 76.412e^{-0.209t})$$

$$t = y - 1948$$

用一元线性回归方法计算拟合系数 A 、 B 时的相关系数 $R=0.971$ 。

巴夫雷油田历年实际综合含水率以及预测的综合含水率见表 2-1。

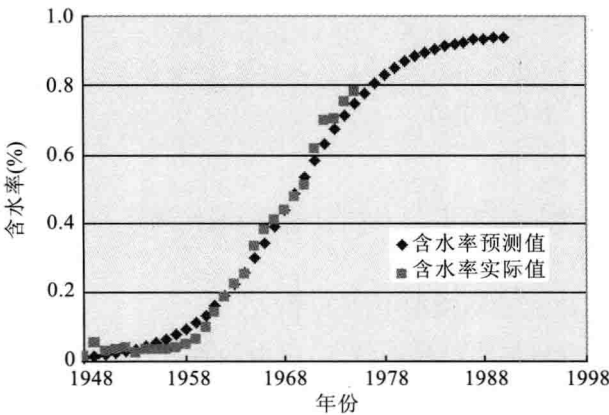


图 2-1 巴夫雷油田综合含水率预测曲线

4. 问题讨论

在采用线性拟合算法时,需事先给定模型的 C 值, C 值的确定要由熟悉油田地质情况的地质家们商定,如果确定有困难,则采用牛顿迭代法的效果要好,计算精度也要高些,但是否符合实际情况需要考虑,因为数学上的最优解不一定是地质上的最佳预测值。