

(78) 確率論的手法による主観情報を考慮した 最適モニタリング地点決定法の検討

福島 智之^{1*}, 米田 稔¹, 森澤 眞輔¹, 坂内 修¹

¹京都大学大学院工学研究科都市環境工学専攻 (〒615-8540 京都市西京区京都大学桂Cクラスター)

* E-mail: hukushima@risk.env.kyoto-u.ac.jp

本研究では、土壤汚染調査において対象領域の汚染の状態を精度良く推定するために最適な土壤モニタリング地点配置を決定する手法において、主観情報を利用する方法を提案する。採取地点の最適配置探索アルゴリズムについては、遺伝アルゴリズムと局所探索法を組み合わせたハイブリッドアルゴリズムを採用した。まず、種々の評価関数を設定し、それぞれの評価関数による最適配置探索法の性能をモンテカルロ法を用いて比較した。比較の結果、推定精度の高い配置を決定する評価関数を採用し、サンプル数増加による影響解析を行い、主観情報導入の効果を検証した。その結果、用いた評価関数によっては、正しい主観情報が与えられた場合、無情報条件下における最適配置や格子状配置より最適な配置が得られた。

Key Words : subjective information, optimal location, sampling, soil contamination, genetic algorithm, evaluation function

1. 序論

中央環境審議会答申「土壤汚染対策法に係わる技術的事項について」¹⁾によって、我が国の土壤概況調査における土壤採取地点決定法が提示されている。調査対象を100m²の格子状に区画した場合は区画内の一点を調査した結果、指定基準を超過した100m²区画ごとに汚染指定区域とし、900m²に区画した場合は区画内を5地点混合法により調査した結果、指定基準を超過した900m²の区画ごとに汚染指定区域となる。

しかしながら、中央環境審議会答申はモニタリングコストを考慮して決定されたものであり、ある面積の中を1地点で代表させることにより、濃度分布推定の際には不確実性が存在するという問題がある。よって、対象汚染物質の濃度分布推定精度並びに、その推定精度とモニタリング地点数との関係を併せて評価することにより、領域全体としての汚染の状態を把握するのに合理的な採取地点配置を決定する手法の開発・確立が必要とされる。

土壤や地下水汚染のサンプリング計画やモニタリング計画の最適化に関して、近年様々な研究が行われている。地下水汚染についての研究として、Tucciarelliら²⁾、McKinneyら³⁾、Meyerら⁴⁾、Wagner⁵⁾の研究がある。土壤汚染についての研究としては、O. Marinoni⁶⁾、Lin-Yu-Pin⁷⁾の研究がある。地下水汚染・土壤汚染のいずれにおいて

も、二次元確率場濃度分布を使っているという点で本研究と共通している。これらの研究では場の不確実性を考慮して最適なモニタリング計画を策定しようとしているが、汚染場における対象物質の濃度分布の統計構造を支配する期待値や分散・相関スケールといったパラメータについては、多くがモニタリング計画策定の前に既に求められており、それらのパラメータの推定を含めた最適モニタリング地点の検討をしていない。

西村ら⁸⁾は地球統計学的手法を用い、場の不確実性を考慮した最適なモニタリング地点を選択する方法について検討しており、配置の探索方法として遺伝アルゴリズム(GA)を用いている。さらにモニタリング回数の分割に関する幾つかの戦略の有効性についても検討している。

木内ら⁹⁾は採取地点の最適配置探索アルゴリズムについて、GAと局所探索法を組み合わせたハイブリッドアルゴリズムを構築し、その有効性について検討している。GAについては西村ら⁸⁾の用いた手法を採用している。

米田ら¹⁰⁾は土壤汚染概況調査の対象領域について不確実性を含んだ何らかの土壤汚染に関するあいまいな情報を利用する方法の有効性について検討している。最適配置探索法としては木内ら⁹⁾のハイブリッドアルゴリズムを採用している。しかし、この研究では評価関数の有効性等は検討されておらず、また、モンテカルロ的手法を用いた評価関数にもかかわらずサンプル数の影響は検討

されていない。

また、前述の従来の研究を通して現行の実施方式である格子状区画調査との性能比較が行われておらず、提案する配置選択手法が現在の実施状況よりどの程度改善されるのかといった定量化された指標が不明瞭である。

本研究では、米田ら¹⁰⁾の方法を踏襲しながら、土壤汚染調査において対象領域の汚染の状態を精度良く推定するために最適な土壤採取地点配置を決定する手法を確立し、提案することをその目的とする。まず、様々な評価関数について推定精度を評価し、最適な評価関数を決定するとともに、米田ら¹⁰⁾の研究において検討されていなかったモンテカルロ的手法におけるサンプル数増加による影響を解析することによって、サンプル数の妥当な値を決定する。評価関数については、推定濃度の対象領域全体の総量誤差と各地点ごとの濃度推定誤差である地点誤差とを設定する。総量誤差を設定するのは、領域全体を浄化処理する目的の際に領域全体での汚染浄化費用を決定することを意図するものである。地点誤差を設定するのは、各地点ごとの濃度分布を正確に求めることを意図するものである。また、総量誤差と地点誤差のそれぞれについて、誤差の二乗をとって評価したときと誤差の絶対値をとって評価したときの影響の違いを見る目的で二乗誤差と絶対値誤差も設定する。

評価関数とサンプル数を決定した後、場の統計構造を表すパラメーターの推定を含めてあいまいな主観情報を組み込んだ最適な地点配置を決定する手法の有効性について検討する。また、無情報条件下と主観情報を組み込んだ場合の双方において、決定した評価関数によって与えられる定量化された数値によって現行の格子状区画調査との性能比較を試みる。

2. 最適モニタリング地点配置の評価手法

(1) 仮定した最適配置問題と一般配置の設定

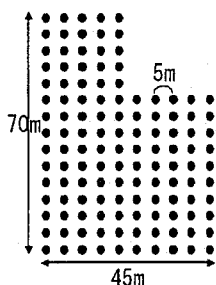


図-1 モニタリング地点の候補

対象領域として図-1に示す125個の採取地点候補を持つ領域を設定し、この中から10個の試料採取地点を選択する問題を考える。前章冒頭で述べたように、10m間隔と30m間隔が我が国の土壤調査における採取地点決定方法で掲げられているものである。本研究は現行の土壤採取方式より小さい間隔によるモニタリングによって濃度分布推定の精度が高まる影響を解析する目的のものである。したがって、図-1のような5m間隔領域の設定は結果の有用性を向上させると考えられる。なお、この対象領域の形状は、西村ら⁹⁾が表層土壤を採取し、重金属類の濃度分布を実測した地点配置と同じに設定している。これは、本研究における模擬データでの検証後、実データへの適用を意図したものであるが、実データへの適用は本研究では行っていない。

また、評価関数によって求められた最適配置と比較するため、一種の評価基準として図-2のような格子状配置を一般配置として2種類設定する。

(2) 土壤中物質濃度の統計分布

土壤中物質濃度は対数正規分布することが多い。環境庁水質保全局「市街地土壤汚染問題検討会報告書」¹¹⁾では、市街地一般土壤における重金属等の存在量について全国7都市158地点の土壤汚染データについて解析し、重金属濃度は対数正規分布に近似しているものと推測できるといった結論が得られている。米田¹²⁾は、おおよそ200m四方の領域から得た157地点の地表土壤中総水銀濃度が対数正規分布していたことを示している。また、西村ら⁹⁾も、京都市内の清掃工場近くのあるグラウンドで測定した地表面土壤中Mg濃度及びNi濃度が対数正規分布したことを示している。

これらのことより、本研究では土壤中汚染物質濃度は対数正規分布をするものと仮定する。このとき、土壤中物質濃度 T_i について対数を取り、 $Z_i = \log(T_i)$ とすると Z_i は正規分布である。以下では、この正規分布する確率変数 Z_i を推定すべき確率変数として定義し、空間的確率変数の母数は Z_i について与えた。最適配置の

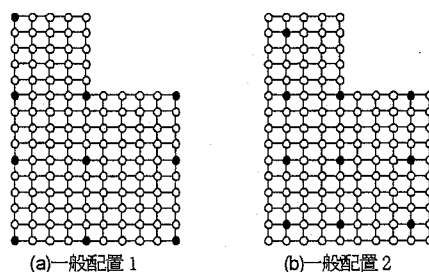


図-2 一般配置 (●はモニタリング地点)

評価では推定された各地点の Z_i を土壤中濃度 T_i に指数変換して用いる。

(3) 地球統計学的手法による評価手法

本研究では、土壤汚染概況調査におけるモニタリング地点配置の最適性を表す関数を定義し、モンテカルロ的手法を用いて推定誤差による評価関数値を求めることにする。以下にその手順を示す。

a) 模擬データ場の発生

濃度分布推定には対象領域内の濃度分布データが必要である。しかし、真の濃度分布場は未知である。このため、真の濃度分布場の確率的標本を作り出す必要がある。この確率的標本を模擬真の場と呼ぶこととする。本研究では、以下に示すNonconditional simulationを用いて、模擬真の場を多数($N_{\sin}$ 個)発生させる。このとき、模擬真の場発生のために用いる汚染物質の空間的確率分布の母数(Z_i の期待値 μ 、分散 σ^2 、相関スケール C_L)としては、多くの実測データなどから得た経験から主観的に決定した母数(Z_i の期待値 μ_0 、分散 σ_0^2 、相関スケール C_{L0})を用いるのが妥当である。

空間的統計構造として多変量正規分布を仮定し、期待値ベクトル μ および共分散行列 Q を与える。領域内で12次モーメントが空間的に変化しないという弱定常性を仮定すると μ と Q は次式のようにになる。

$$\mu = \mu e \quad (2-1)$$

$$Q = \sigma^2 A \quad (2-2)$$

ここで、

μ : 期待値

e : 要素がすべて1の n 次縦ベクトル

σ^2 : 分散

A : 相関係数行列

である。相関関数形として指数関数形を仮定すると A_{ij} は以下の式で表される。

$$A_{ij} = \exp\left(-\frac{r_{ij}}{C_L}\right) \quad (2-3)$$

ここで、 r_{ij} は第 i 地点と第 j 地点との地点間距離である。(2-1)式の μ に μ_0 を、(2-2)式の σ^2 に σ_0^2 を、(2-3)式の C_L に C_{L0} をそれぞれ代入し、 μ 、 Q 、 A を求める。

次に、共分散行列 Q のコレスキー分解を行う。コレスキー分解は次式で示される。

$$Q = [V] \cdot [V]^T \quad (2-4)$$

ただし、行列 $[V]$ は次のような三角行列である。

$$[V] = \begin{bmatrix} v_{11} & & 0 \\ \vdots & \ddots & \\ v_{n1} & \cdots & v_{nn} \end{bmatrix} \quad (2-5)$$

また、 $[V]^T$ は $[V]$ の転置行列を表す。 Q の第 ij 要素を q_{ij} で表すと、 $[V]$ の各要素 v_{ij} と q_{ij} には次式の関係がある。

$$q_{ij} = \sum_{k < j} v_{ik} v_{jk} \quad (2-6)$$

共分散行列のような正定値の行列についてコレスキー分解は一意である。(2-6)式を用いて Q から $[V]$ を求める。

与えられた統計構造を持つ n 次の多変量正規乱数 $\{Z\}$ は $[V]$ を用いて次式で発生することができる。

$$\{Z\} = [V]\{\varepsilon\} + \{\mu\} \quad (2-7)$$

ただし、 $\{\varepsilon\}$ は n 次の白色標準正規乱数ベクトルである。(2-7)式に、 $[V]$ と $\{\mu\}$ を代入し、 $\{Z\}$ を求める。これを真の濃度分布場の確率的標本とし、模擬真の場を $N_{\sin}$ 個発生させる。

b) 最尤推定法と条件付き確率分布の導出

発生した各模擬真の場において、新たなサンプリング候補地点から得られる観測データを用いて、以下に示した方法で最尤推定を行い、確率分布の母数を推定し直す。

有限領域から得られる n 個の観測データ $Z_i (i=1, 2, \dots, n)$ が多変量正規分布をする場合、尤度関数は次式で与えられる。

$$F(\theta/Z) = (2\pi)^{-\frac{n}{2}} |Q|^{-\frac{1}{2}} \exp\left(-\frac{1}{2}(Z-\mu)^T Q^{-1}(Z-\mu)\right) \quad (2-8)$$

ここで、 Q は共分散行列、 $|Q|$ は Q の行列式、 μ は Z の期待値ベクトル、 Z は Z_i を要素とする n 次縦ベクトル、 θ は確率密度関数の性質を決定するパラメータのベクトルである。このとき対数尤度関数は、

$$L(\theta/Z) = -\frac{n}{2} \ln(2\pi) - \frac{1}{2} \ln(\sigma^{2n} |A|) - \frac{1}{2} \sigma^{-2} (Z-\mu)^T A^{-1} (Z-\mu) \quad (2-9)$$

となる。 $F(\theta/Z)$ を最大化することは $L(\theta/Z)$ を最大化することに等しい。 $L(\theta/Z)$ を最大にする μ と σ^2 を求めるため、対数尤度関数 $L(\theta/Z)$ を μ および σ^2 について偏微分して0とおくと、

$$\frac{\partial L}{\partial \mu} = 0 \quad \therefore \mu = \frac{e^T A^{-1} Z}{e^T A^{-1} e} \quad (2-10)$$

$$\frac{\partial L}{\partial \sigma^2} = 0 \quad \therefore \sigma^2 = \frac{1}{n} (Z-\mu)^T A^{-1} (Z-\mu) \quad (2-11)$$

(2-10)及び(2-11)式を(2-9)式に代入すると、

$$L(\theta/Z) = -\frac{n}{2} \ln(2\pi) - \frac{1}{2} \ln(\sigma^{2n} |A|) - \frac{n}{2} \quad (2-12)$$

となり、 $L(\theta/Z)$ は相関係数行列 A のみの関数となる。

いま、相関関数形として指数関数形を仮定すると、 A の第 ij 成分は次式で表される。

$$A_{ij} = \exp\left(-\frac{r_{ij}}{C_L}\right) \quad (2-13)$$

ここで、 r_{ij} は第 i 地点と第 j 地点との地点間距離である。このとき、 A_{ij} は相関スケール C_L のみの関数となる。よって、最尤推定法は $L(\theta/Z)$ を最大とする C_L を求めた後、 μ および σ^2 の値を求める。本研究では、相関スケール C_L は1mから50mまで0.5m間隔として、(2-12)及び(2-13)式より得られる対数尤度関数 $L(\theta/Z)$ を最大にする相関スケールを母数として選択する。

その後、最尤推定によって推定した母数を用いて、サンプリング候補地点から得られるデータで条件付けした場の条件付き期待値を以下に示した方法で求める。

いま、確率変数 $Z_i (i=1,2,\dots,n)$ 及び $Z_{u_i} (i=1,2,\dots,n_u)$ の同時確率密度関数が多変量正規分布で表されるとする。このとき n 個の観測データ Z_i が既知であるとき、 Z_{u_i} の条件付き確率分布を求める。 Z_i と Z_{u_i} の順に並べた縦ベクトルを $\{W\}$ とすると、 $\{W\}$ は期待値ベクトル $\{\mu\}$ 、共分散行列 Σ_w の多変量正規分布 $N(\{\mu\}, \Sigma_w)$ となる。 $\{W\}$ を $\{W_1\}$ と $\{W_2\}$ に分割すると、

$$\{W\} = \begin{Bmatrix} \{W_1\} \\ \{W_2\} \end{Bmatrix} \quad \{\mu\} = \begin{Bmatrix} \{\mu_1\} \\ \{\mu_2\} \end{Bmatrix} \quad (2-14)$$

$$\Sigma_w = \begin{bmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{21} & \Sigma_{22} \end{bmatrix} \quad (2-15)$$

このとき、 $\{W_1\}$ を与えたときの $\{W_2\}$ の条件付き確率分布は多変量正規分布の性質より次式で表される。

$$N(\{\mu_2\} + \Sigma_{21}\Sigma_{11}^{-1}\{\mathbf{W}_1 - \{\mu_1\}\}, \Sigma_{22.1}) \quad (2-16)$$

ただし、

$$\Sigma_{22.1} = \Sigma_{22} - \Sigma_{21}\Sigma_{11}^{-1}\Sigma_{12} \quad (2-17)$$

である。よって、与えられた期待値ベクトルと共分散行列を持つ乱数を発生させる方法である(2-7)式で、(2-16)式で与えられる期待値ベクトルと共分散行列を用いれば、サンプリング候補地点から得られるデータで条件付けした確率場を発生させることができる。

g) 評価関数の導出

条件付き確率分布を推定した場の条件付き濃度値 $\hat{T}$ と、模擬真の場における真の濃度値 T の推定誤差を領域内の全地点について求め、図-3に示した概念図において以下のように設定した評価関数によって評価する。

ここで、

N_{sin} : 模擬真の場の数

$T_{i,j}$: 模擬真の場 i の地点 j の濃度

$\hat{T}_{i,j}$: 条件付き確率分布より推定した模擬真の場 i の地点 j の濃度

N_s : 全採取候補地点の数

$f_i \cdot g_i$: 模擬真の場 i における評価関数

$F \cdot G$: 設定配置の評価関数

とおく。

最適配置の探索に用いる評価関数として、次のようなものを設定する。

① 総量誤差

推定された濃度値の全採取候補地点にわたる総量の誤差の二乗を評価関数とするため、

$$f_i = \left(\sum_{j=1}^{N_s} \hat{T}_{i,j} - \sum_{j=1}^{N_s} T_{i,j} \right)^2 \quad (2-18)$$

とおき、 f_i の最大値・95%値・90%値・50%値・平均値の逆数をそれぞれ F として評価する。

また、平均値については二乗をとらない推定総量誤差

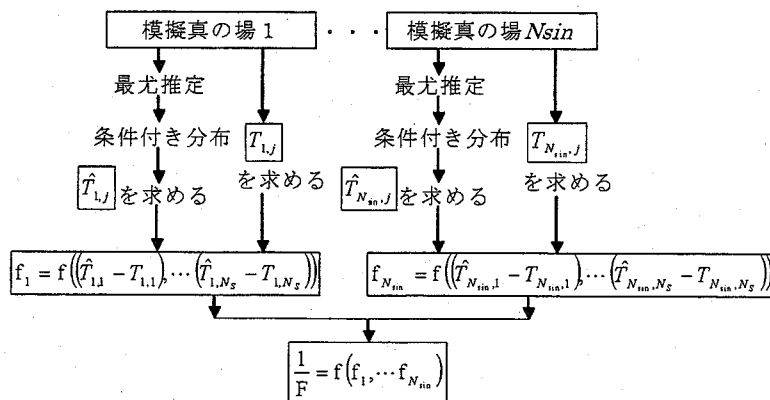


図-3 評価関数の求め方

の絶対値を用いた評価関数の設定のため、

$$g_i = \left| \sum_{j=1}^{N_1} \hat{T}_{i,j} - \sum_{j=1}^{N_2} T_{i,j} \right| \quad (2-19)$$

とおき、 g_i の平均値の逆数を G として評価する。

②地点誤差

推定された濃度値の全採取候補各地点の誤差の二乗を評価関数とするため、 $(\hat{T}_{i,j} - T_{i,j})^2$ の最大値・95%値・90%値・50%値・平均値をそれぞれ $f_{k,i}$ (k 番目の評価関数の意味で suffix に k を用いる) とし、その $f_{k,i}$ の最大値・95%値・90%値・50%値・平均値の逆数をそれぞれ F として評価する。

また、平均値については二乗をとらない推定地点誤差の絶対値を用いた評価関数の設定のため、 $|\hat{T}_{i,j} - T_{i,j}|$ の平均値を g_i とし、その g_i の平均値の逆数を G として評価する。

上述の各評価基準を設定する理由を以下に述べる。誤差の最大値を評価基準として選定するのは、推定誤差のうち最大のものを最小にすれば、それ以下の誤差も小さくなるだろうという考えに基づくものである。95%値を設定するのは、乱数発生に擬似乱数を用いる影響により、最大値や最小値に近いデータは発生させた模擬真の場合と大きく変動している可能性があり、その結果生じる異常値に近いデータを排除する目的のものである。90%値も同様の目的のものである。比較的高濃度地点のみに着目する評価関数に対して、領域全体を考慮した中央値や平均値を最小にした時の影響の相違を見る目的で設定するものが、50%値・平均値である。

総量誤差と地点誤差のそれぞれについて、 $F \cdot G$ が最大となる配置を最適配置とする。

(4) 最適配置探索の手法

前節において示した評価関数の値を最大にする試料採取地点の取り方を求める手法として、本研究では木内ら⁹⁾の開発したハイブリッド遺伝アルゴリズムを最適配置探索手法として採用する。

木内ら⁹⁾は、遺伝アルゴリズム(genetic algorithm:GA)と最急降下地点探索法(steepest descent method:SDNM)の2つの手法が相互補完的な特徴を有している点を利用し、2つの手法を組み合わせることによって、局所的最適解に陥ることなく、より早く大域的最適解に到達することを目指した、ハイブリッドアルゴリズムを構築した。

ここで用いるGAとしては、西村ら⁹⁾のGAを利用した最適化アルゴリズムを採用している。まず対象とする調査領域を多くの小さなメッシュに分割し、各節点到番号をつける。そして、これらの節点の中から L_c 個の試料

採取地点を選択することとし、この試料採取地点の並びを遺伝子とする。遺伝子の選択方法としてはランク方式とエリート主義戦略を採用し、交叉確率 $P_c=0.8$ 、突然変異確率 $P_m=0.2$ とする。

また、最急降下地点探索法(SDNM)とは、木内ら⁹⁾が構築した局所探索の方法である。この配置探索アルゴリズムは、初期配置のうち1つの地点について東西南北4方向のいずれか1方向の隣接地点に移動させた配置全てを候補として生成し、より適した評価関数の値を持つ候補を新たな初期配置として更新させることを繰り返し、局所解に到達させるというものである。

木内ら⁹⁾のハイブリッドアルゴリズムは、生成された初期配置それぞれに対して最急降下地点探索法による探索を行い、得られた配置の集合に対してGAを1世代分進め、この動作を n 回繰り返す、というものである。本研究では、GA終了後の上位10個体についてSDNMによる探索を行い、 $n=5$ とする。

(5) 無情報条件下の模擬汚染場における評価手法の比較

a) 模擬真の場に関するパラメーター

まず、ここでは仮に、2.(3)a) に示したNonconditional Simulationを用いて、期待値 $\mu=0$ 、分散 $\sigma^2=0.25$ 、相関スケール $C_L=30m$ と仮定した模擬真の場を発生させる。ここで用いるパラメーター値としては実際は、多くの実測データを解析し、経験的に最も妥当と考えられる値を採用することとなる。

b) 評価関数比較

2.(3)c) に示した各種評価関数によって求められた最適配置を、別の乱数系列によって発生させた模擬真の場のもと、各種評価関数によって評価値を算出し直すことで各種評価関数間の性能比較を試みる。その際の模擬真の場数は1001個とする。

c) サンプル数増加による影響解析

前項の結果によって、最良と判定された評価関数を選択し、模擬真の場数増加による評価関数値変動の影響を解析する。

d) 計算時間

本研究では、京都大学メディアセンターのHPC2500を計算機として採用した。最適配置を求めるプログラム中で、最尤推定・条件付確率分布の導出・評価関数の計算アルゴリズムに渡ってプロセス並列化機能を使用した。プロセス数は32とした。評価関数を地点誤差二乗平均とした場合の、Total CPU Timeは 1.6×10^3 sec、FLOPSは 2.4×10^7 MFLOPSであった。

(6) 主観情報の適用と評価

a) 主観情報の導入

米田¹⁰⁾は、土壤汚染調査において対象領域に土壤汚染に関して他の地点より濃度が高そうだというあいまいな情報が存在している場合について、模擬真の場発生の段階でこのあいまいな事前情報を満たす場を発生するようにした。以下に手順を示す。

手順1 汚染予想地点の設定

他の地点より濃度が高そうだという s 個の地点を汚染予想地点としてその地点の対数値を $Z_k (k=1,2,\dots,s)$ として設定する。

手順2 主観情報のモデル化

領域の空間的確率分布の母数 (Z_k の期待値 μ , 分散 σ^2 , 相関スケール C_L) を用いて、模擬真の場の候補を発生させる。このうち、汚染予想地点が期待値 μ や $\mu + \sigma$ より大きいという条件を満たすものを模擬真の場として保存する。あいまいな事前情報が存在しない場合は、空間的確率分布の母数を用いて発生させた模擬真の場をそのまま用いる。汚染予想地点に与える条件はあいまいな事前情報の確実さによって変化するものであり、領域の期待値に対して大きいかどうか ($\mu < Z_k$) , 領域の期待値と標準偏差に対して大きいかどうか ($\mu + \sigma < Z_k$) , 濃度基準のようにある濃度 T_0 について大きいかどうか ($\log(T_0) < Z_k$) などが考えられる。本研究では、ある地点の濃度が高そうだという情報のあいまいさを考慮し、 $\mu < Z_k$, $\mu + \sigma < Z_k (k=1,2,\dots,s)$ という2つの条件を設定する。これらの条件を設定した主観情報はある程度信頼度があるものとして、本研究での最適配置探索法に取り入れることとする。

b) 対象領域及び主観情報

対象領域を図4(a)~(f)に示す。領域は図1と同じものとする。主観情報が存在しない場合をシナリオ1、ある地点が高そうだという主観情報として黒丸●で表された地点を汚染予想地点とし、シナリオ2-6とする。

c) 模擬真の場に関するパラメーター

シナリオ1の無情報事前分布の場合、2.(3)a) に示した Nonconditional Simulation を用いて模擬真の場を発生させる。前節と同様に、 Z_i の期待値 $\mu=0$, 分散 $\sigma^2=0.25$, 相関スケール $C_L=30m$ と仮定する。シナリオ2-6のように汚染予想地点が存在するが濃度値などの確実な情報が存在しない場合、期待値 $\mu=0$, 分散 $\sigma^2=0.25$, 相関スケール $C_L=30m$ と同様の仮定を行い、全ての汚染予想地点での濃度の対数値が、 $\mu=0$, あるいは $\mu + \sigma=0.5$ の値以上である模擬真の場を条件に適した場として用いる。このような模擬真の場は、ランダムに発生させた多数の模擬真の場候補から、条件を満たす場のみを採用することで作成する。それぞれのシナリオについて、2.(3)(4) に示した方法で最適配置を探索する。模擬真の場数は1001

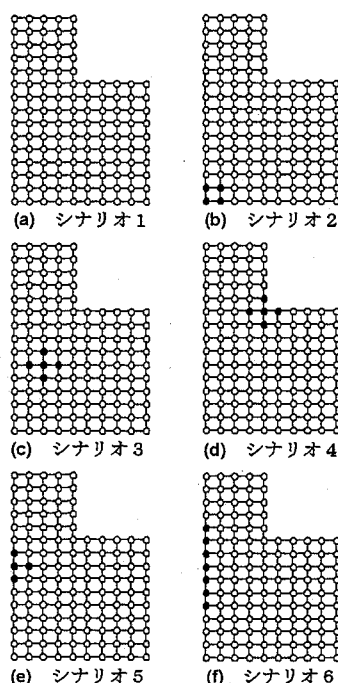


図4 対象領域 (●は汚染予想地点)

個とする。

d) 評価関数の選択

前節の施行結果によって、最良と判定された評価関数を選択し、それを用いて最適配置の探索を行う。

e) 無情報条件下と主観情報導入時との比較

2.(5)a)・b) で設定した各シナリオにおける、無情報条件下と主観情報導入時、それぞれの評価値を比較する。ただし、2.(5)b) と同様に、模擬真の場の発生に用いた乱数は最適配置探索用と評価用とで同一ではない。

f) 計算時間

今回の評価プログラムを動作させるにあたり、条件は2.(5)d) と同一とした。評価関数を地点誤差二乗平均とし、シナリオ2 ($\mu < Z_k$) を設定した場合、Total CPU timeは $2.0 \times 10^3 \text{sec}$, FLOPSは $2.3 \times 10^7 \text{MFLOPS}$ であり、2.(5)d) の無情報条件下と比較してそれぞれ約1.3倍、0.97倍となった。また、($\mu + \sigma < Z_k$) の場合、Total CPU timeは $2.1 \times 10^3 \text{sec}$, FLOPSは $2.1 \times 10^7 \text{MFLOPS}$ であり、無情報条件下と比較してそれぞれ約1.3倍、0.90倍となった。

3. シミュレーション結果

(1) 無情報条件下の模擬汚染場における評価手法の比較

a) 評価関数比較結果

表-1~4において、適用データとあるのは各評価関数によって求められた最適配置、あるいは図-2に示した一般配置のことであり、それらの配置を別の乱数系列によって発生させた模擬真の場のもと、評価基準である各評価関数によって再評価したことを意味する。

異なる評価関数によって算出された評価値同士を比べることに意味はない。比較結果が有意義となるのは、同一の評価関数によって評価された評価値同士のみである。そこで本研究では、まず、2.(3)c)で設定したそれぞれの評価関数によって得られた最適配置を、同一の評価関数で評価し直した評価値同士の比較を試みた。それらの評価値の中でいずれか一つを基準にすることで、各評価値間の関係は無次元の比で表すことができる。ある評価関数によって得られた最適配置は、当然評価関数の意図が組み込まれたものであるはずなので、同一の評価関数で再評価した時に、他の評価関数で再評価した時に比べて評価値が大きくなる可能性が最も高い。他の評価関数に

よって得られた最適配置の評価値を、最も高くなるであろう評価値を基準として除した評価値比を算出した。このため表-1及び表-3の対角項は全て1となっている。これを数式として定義すると以下ようになる。

F_{ij} : i 番目の評価関数によって求められた最適配置を j 番目の評価関数で評価した評価値

E_{ij} : i 番目の評価関数によって求められた最適配置を j 番目の評価関数で評価した評価値比

とくと、

$$E_{ij} = \frac{F_{ij}}{F_{jj}} \quad (3-1)$$

また、最大値・95%値・90%値・50%値・二乗平均が推定誤差の二乗を評価関数として用いるのに対し、絶対値平均が推定誤差の絶対値を評価関数として用いるために生じる評価関数間の次元の違いによる定量化後の比較時の影響を修正するため、新たに修正評価値比を設定した。これは、各配置の全ての評価関数を通した評価値比の平均をとる際に、各評価値比の重み付けを一樣にすることを意図したものである。これを数式として定義すると以下ようになる。

表-1 修正評価値比計算結果—評価関数：総量誤差

修正評価値比			評価基準						AVE
			総量誤差						
			最大値	95%値	90%値	50%値	二乗平均	絶対値平均	
適用データ	総量誤差	最大値	1.000	0.811	0.795	0.846	0.886	0.811	0.858
		95%値	0.896	1.000	1.027	0.992	1.051	0.990	0.993
		90%値	0.986	0.847	1.000	0.921	0.992	0.916	0.944
		50%値	0.929	0.933	0.935	1.000	0.993	0.933	0.954
		二乗平均	0.753	0.951	0.986	1.006	1.000	0.937	0.939
		絶対値平均	0.728	0.970	1.023	1.078	1.011	1.000	0.968
	一般配置 1	0.435	0.491	0.560	0.628	0.592	0.579	0.547	
	一般配置 2	0.699	1.118	0.994	1.067	1.077	0.995	0.992	

表-2 順位表—評価関数：総量誤差

順位法			評価基準						AVE
			総量誤差						
			最大値	95%値	90%値	50%値	二乗平均	絶対値平均	
適用データ	総量誤差	最大値	1	7	7	7	7	7	6.00
		95%値	4	2	1	5	2	3	2.83
		90%値	2	6	3	6	6	6	4.83
		50%値	3	5	6	4	5	5	4.67
		二乗平均	5	4	5	3	4	4	4.17
		絶対値平均	6	3	2	1	3	1	2.67
	一般配置	一般配置 1	8	8	8	8	8	8	8.00
		一般配置 2	7	1	4	2	1	2	2.83

G_{ABS} : 絶対値平均を評価関数として求めた評価値比
 F_{mod} : 修正評価値比
 とおくと、

$$F_{mod} = \frac{1}{\left(\frac{1}{G_{ABS}}\right)^2} \quad (3-2)$$

また、最適配置探索用と評価用に同一の乱数系列を用いた模擬真の場を発生させた場合、ある評価関数によって得られた最適配置を同一の評価関数で評価した時に評価値が最大となるが、最適配置探索用と評価用とで別の乱数系列を用いた模擬真の場を発生させているので、同一の評価関数で評価した評価値が最大となるとは限らない。表-1及び表-3において、1以上の値が存在するのは、有限数の擬似乱数系列を使用した影響と考えられる。模擬真の場を無限大にまで増加させると、1以上の値も出現しなくなるのではないかと考えられるが、推測の域を脱してはならず、模擬真の場数は今後の検討課題の1つと言える。

同一の評価関数によって求められた評価値を単純に順位付けし、それらの順位の平均をとることによって、最適配置を求めた評価関数の性能を比較することも検討し

た。この方法を順位法と呼ぶこととする。

総量誤差を最小にすることが目的の際は総量誤差による評価関数、地点誤差を最小にすることが目的の際は地点誤差による評価関数を用いることとし、それぞれの種類の評価関数について、最大値・95%値・90%値・50%値・二乗平均・絶対値平均による評価関数間の比較を行った。適用データを前述の6つに2つの一般配置を加えた8つとし、総量誤差・地点誤差のそれぞれについて、修正評価値比・順位を求めた結果を表-1~4に示す。

表-1では、修正評価値比の平均は、総量誤差95%値→一般配置2→総量誤差絶対値平均の順に高くなっている。

表-2における、順位法による比較では、平均順位は総量誤差絶対値平均→総量誤差95%値→一般配置2の順となった。格子状に設定した一般配置2(図-2(b))との比較では、総量誤差95%値が6回中4回評価値が低くなったのに対し、総量誤差絶対値平均の方は評価値が低くなった回数が6回中2回に抑えられている。

評価値比、順位法にわたって総量誤差95%値が高く評価された理由を考察する。最大値・最小値に近いデータ区間では、用いた擬似乱数場によって模擬真の場ごとに大きく変動している可能性が考えられる。評価関数として95%値を採用することで、発生させた乱数系列場到大

表-3 修正評価値比計算結果—評価関数：地点誤差

修正評価値比			評価基準					
			地点誤差					
			最大値	95%値	90%値	50%値	二乗平均	絶対値平均
適用データ	地点誤差	最大値	1.000	0.965	0.965	0.983	0.939	0.953
		95%値	1.385	1.000	0.972	0.991	0.971	0.969
		90%値	1.109	1.007	1.000	1.027	0.982	0.983
		50%値	1.197	0.909	0.966	1.000	0.971	0.968
		二乗平均	1.061	1.051	1.028	1.031	1.000	1.005
		絶対値平均	1.039	1.038	1.010	1.039	0.992	1.000
	一般配置 1		1.082	0.846	0.834	0.929	0.900	0.900
	一般配置 2		1.179	0.981	0.988	0.984	0.979	0.972

表-4 順位表—評価関数：地点誤差

順位法			評価基準					
			地点誤差					
			最大値	95%値	90%値	50%値	二乗平均	絶対値平均
適用データ	地点誤差	最大値	8	6	7	7	7	7
		95%値	1	4	5	5	6	5
		90%値	4	3	3	3	3	3
		50%値	2	7	6	4	5	6
		二乗平均	6	1	1	2	1	1
		絶対値平均	7	2	2	1	2	2
	一般配置 1		5	8	8	8	8	8
	一般配置 2		3	5	4	6	4	4

大きく影響を受ける誤差の最大値から5%の区間を排除することとなり、比較的乱数系列場に左右されにくい範囲の誤差を最小にすることに成功した結果ではないかと考えられる。

総量誤差の評価関数としては、総量誤差95%値が表-1の修正評価値比平均で最も高く、表-2の順位法で2番目に評価された。また、総量誤差絶対値平均が、表-1の修正評価値比平均で3番目に、表-2の順位法で最も高く評価された。これらのことより、総量誤差を最小にすることを目的とする場合は、総量誤差95%値と総量誤差絶対値平均という2つの評価関数が、本項の条件下においては適性が高いという判断をなすこととし、次項での検討対象とした。

表-3では、修正評価値比の平均は、地点誤差95%値→地点誤差二乗平均→地点誤差絶対値平均の順となった。地点誤差95%値が高く評価される結果となったが、その理由として乱数系列に左右されやすい最大値による評価値比が他に比べて非常に高くなっている影響が強いと考えられる。

表-4では、平均順位は、地点誤差二乗平均→地点誤差絶対値平均→地点誤差90%値の順となった。一方、評価値比で高く評価された95%値については、順位法において高い評価はされていない。

地点誤差の評価関数としては、地点誤差二乗平均が表-3の修正評価値比平均で2番目に、表-4の順位法で最も高く評価された。また、地点誤差絶対値平均が、表-3の修正評価値比平均で3番目に、表-4の順位法で2番目に評価されたが、いずれも地点誤差二乗平均の順位以下だった。表-3の修正評価値比平均で最も高く評価された地点誤差95%値は、表-4の順位法で4番目であり、修正評価値比平均において高い評価の原因となった最大値を除いては、常に4番目以降に評価されている。地点誤差に関しては、最大値や95%値などが一点のデータで判定される一方で、二乗平均や絶対値平均は複数のデータの統計値として判定されるので、安定度が高くなるという事実が確認できた。これらのことより、地点誤差を最小にすることを目的とする場合は、修正評価値比・順位法にわたって比較的高く評価された地点誤差二乗平均を評価関数として、次項以降での検討対象とした。

b) サンプル数増加による影響解析結果

前項において高く評価された、総量誤差95%値・総量

誤差絶対値平均・地点誤差二乗平均の3つの評価関数について、模擬真の場数増加による影響を解析した。その結果を図-5及び表-5に示す。

最適配置探索に用いる模擬真の場数を増加させた結果、評価値は次第に減少する傾向が見られた。これは、模擬真の場数が少ない場合、限られた数の模擬真の濃度分布場について精度良く推定濃度分布場を求めるための最適配置を決定することは比較的容易だが、模擬真の場が増加すると全ての模擬真の場にわたって誤差を小さくする最適配置を決定することが困難になることによる影響と考えられる。

模擬真の場数を10001まで増加させた結果、各評価値には幾分の収束傾向が確認できる。モンテカルロの手法を用いて推定誤差による何らかの関数の期待値を求める場合には、確率論的標本の数は多いほど推定精度の安定性を保証できると考えられる。よって、模擬真の場数は無限に増やすのが理想的だが、計算時間やコストといった制約条件から、本研究では模擬真の場数として1001を採用した。表-5より、模擬真の場数を1001から10001に増加させた時の評価値の変化はそれぞれ、評価関数を総量誤差95%値とした場合で約12%、総量誤差絶対値平均の場合で約7%、地点誤差二乗平均の場合で約2%となった。総量誤差95%値等はサンプル数1001ではまだ十分収束していないと考えられる。以下では、本研究で設定した1001という模擬真の場数についての有効性を検討する。

前項で述べたように、表-3において、総量誤差95%値により求められた最適配置について、評価基準を最大値とした時の高い修正評価値比の影響を受けて、全ての評価基準による修正評価値比を平均した値がそのまま最も

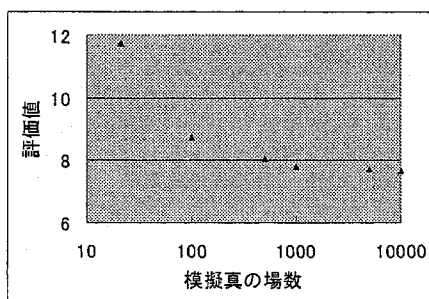


図-5 地点誤差二乗平均評価値の推移

表-5 模擬真の場数による評価値の推移

評価値		真の場数					
		21	101	501	1001	5001	10001
評価関数	総量誤差 95%値	0.02723	0.00402	0.00289	0.00267	0.00247	0.00234
	総量誤差絶対値平均	0.316	0.160	0.137	0.132	0.122	0.121
	地点誤差二乗平均	11.735	8.724	8.035	7.813	7.701	7.664

高くなるという事象が起きている。このため、一つの異常値に左右されかねない修正評価値比の平均をとって比較を行う方法よりも、表-2及び表-4で表される順位法による比較を重視することとした。模擬真の場数を1001に固定し、乱数系列を変えた模擬真の場を用いて前項同様に順位を算出することをさらに2回行った。その結果、総量誤差に関しては、8つの配置のうち最大値と一般配置1については低い順位に安定することが確認されたが、他の6つの配置の順位については安定せず、どの配置が最も優れていると判定することはできなかった。このことから、総量誤差に関しては模擬真の場数1001ではまだ十分に評価値が収束していないため、比較は困難であると考えられる。一方、地点誤差に関しては、表-4で最大値と一般配置1が低い平均順位となり、二乗平均と絶対値平均が高い平均順位となったのと同様に、他の2つの乱数系列についても8つの配置のうち最大値と一般配置1については低い順位に安定し、二乗平均と絶対値平均については高い順位に安定することが確認された。このことから、地点誤差に関しては模擬真の場数が1001であったとしても比較は可能であると考えられる。二乗平均と絶対値平均のうち、より有効であると考えられる評価関数を選択するため、計3つの乱数系列それぞれにつき6つの評価基準による評価回数である計18ケースの結果について分析した。その結果、8つの配置の中で1位と判定された回数が、二乗平均で7回、絶対値平均で6回となった。また、二乗平均と絶対値平均同士を比較した結果、18ケース中11ケースにおいて二乗平均が絶対値平均より順位が高いと判定された。両者の差は僅かではあるが、ここでは二乗平均に若干の優位性を認め、次節で用いる評価関数として選択することとした。

(2) 主観情報の適用と評価

採取地点を10箇所として最適配置を探索した結果及び、探索結果を比較するためにシナリオ1-6の最適配置結果について、シナリオ1-6の条件で発生させた模擬真の場による評価関数を計算した結果を図-6、表-6及び図-7、表-7に示す。黒丸●で表された地点が最適配置地点であり、他の地点より大きな白丸○で表された地点は主観情報が存在している地点を表す。前節における考察の結果、評価関数としては地点誤差二乗平均を選択した。

$\mu < Z_k$ という主観情報を設定したシナリオの結果を図-6及び表-6に示す。図-6から、最適配置が汚染予想地点のいくつかを含むという傾向・高濃度主観情報を有する地点への偏りの傾向が見られた。図-6(f)では、最適配置は汚染予想地点を含んでいない。これは、汚染予想地点が領域全体の境界上に存在する場合、かつ $\mu < Z_k$ という弱い条件の下では、境界上の汚染予想地点にモニタ

リング地点を配置するより、他のモニタリング地点との距離を小さくした方が、領域全体の濃度分布をより精度よく推定することができることを示していると考えられる。表-6では、シナリオ2に関して、主観情報を考慮した配置の評価値が無情報条件下における最適配置の評価値を下回った以外は、主観情報存在下において主観情報を考慮した配置は無情報条件下における最適配置や格子状配置より最適性が高いと判定された。シナリオ1に関

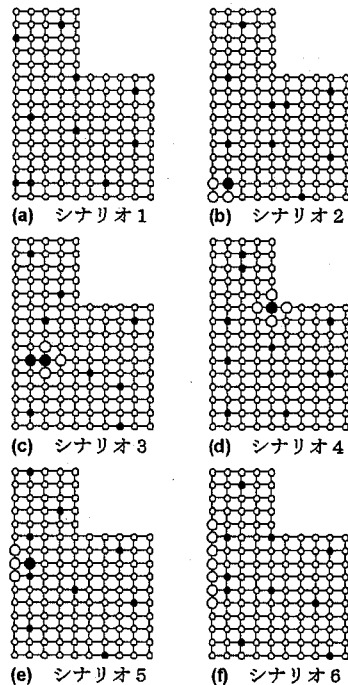


図-6 地点誤差二乗平均シナリオ($\mu < Z_k$)最適配置
(●は汚染予想地点、大きな○は汚染予想地点)

表-6 地点誤差二乗平均シナリオ($\mu < Z_k$)評価値計算結果

		検証に用いた場					
	シナリオ ($\mu < Z_k$)	1	2	3	4	5	6
検証 に 用 い た 配 置	1(図-6(a))	728	538	448	4.55	4.82	4.23
	2(図-6(b))	7.06	5.34				
	3(図-6(c))	7.18		4.54			
	4(図-6(d))	7.21			4.55		
	5(図-6(e))	7.10				4.86	
	6(図-6(f))	7.12					4.32
	一般配置 1 (図-2(a))	6.55	4.68	3.95	3.99	4.29	3.82
	一般配置 2 (図-2(b))	7.13	5.23	4.46	4.47	4.71	4.10

しては、無情報条件下における最適配置が他のどの配置よりも最適性が高いと判定された。これによって、誤った主観情報を導入した場合に評価値が低下するという事実が確認できる。

$\mu + \sigma < Z_k$ という主観情報を設定したシナリオの結果を図-7及び表-7に示す。図-7によって、主観情報が $\mu < Z_k$ であった場合と比較して、最適配置が汚染予想地点のうち、より多くの地点を含むようになり、前述のような高濃度主観情報を有する地点への偏りのさらに強い傾向が確認された。図-7(f)では、前段落に記述した、 $\mu < Z_k$ の主観情報導入の際に最適配置が汚染予想地点を含まない現象は解消されており、最適配置が汚染予想地点の2つを含んでいる。これは、主観情報の制約条件を厳しくしたことにより、高濃度主観情報の存在しない地点より主観情報の存在する地点を積極的にモニタリングする方が、濃度分布推定の精度が高まると考えられる。表-7では、主観情報存在下において主観情報を考慮した配置は無情報条件下における最適配置や格子状配置より最適性が高いと判定され、なおかつ主観情報が $\mu < Z_k$ であった場合と比較して、その評価値差が大きくなる傾向が見られた。また、シナリオ1に関しては、無情報条件下における最適配置が他のどの配置よりも最適性が高いと判定され、誤った主観情報を導入した場合に評価値が低下するという事実がこの場合においても確認できる。

以上のことから、特に主観情報の制約条件を厳しくした際、しかもその主観情報が適切な情報である際に限り、未知地点の濃度推定における精度の向上に有効性が発揮される結果となり、格子状均等配置と比較して、推定誤差の精度が10%程度改善されることが確認された。同時に、逆に主観情報が妥当性を欠いたものであった場合は、その情報を考慮した配置による推定の性能が劣ってしまうことも示唆される。

4. 結論

本研究では確率論的手法を用いて場の不確定性を考慮した最適なモニタリング地点を選択する方法について検討した。得られた結論を以下に示す。ただし、以下の結論は本研究で設定した仮定が成立する範囲内での知見である。

1) 最適配置探索に用いる模擬真の場数を増加させると評価値は減少する傾向が見られるが、ある程度の数まで真の場を増やすと、徐々に評価値が収束に向かい、安定した評価値となっていく様子が確認された。また、本研究で用いた模擬真の場数では、総量誤差は十分に収束しなかったが、地点誤差は比較的収束し、地点誤差による比

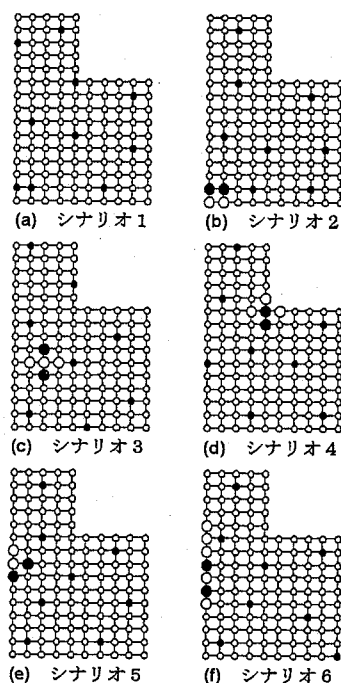


図-7 地点誤差二乗平均シナリオ($\mu + \sigma < Z_k$) 最適配置
(●は汚染予想地点、大きな○は汚染予想地点)

表-7 地点誤差二乗平均シナリオ($\mu + \sigma < Z_k$)
評価値計算結果

		検証に用いた場					
	シナリオ $\mu + \sigma < Z_k$	1	2	3	4	5	6
検証 に 用 い た 配 置	1(図-7(a))	7.28	4.10	2.91	3.06	3.19	2.54
	2(図-7(b))	7.13	4.21				
	3(図-7(c))	7.02		3.11			
	4(図-7(d))	7.12			3.23		
	5(図-7(e))	7.03				3.51	
	6(図-7(f))	6.81					2.84
	一般配置 1 (図-2(a))	6.55	3.51	2.49	2.64	2.97	2.49
	一般配置 2 (図-2(b))	7.13	3.91	2.93	3.04	3.23	2.65

較結果も安定したため、比較は可能であると考えられる。

2) 未知地点の濃度を推定するのに最適なモニタリング地点の探索に用いる評価関数としては、地点誤差二乗平均が、比較的推定精度の高い配置を決定する評価関数であるという結果が得られた。

3) ある地点の濃度が高そうだという主観情報が存在するが、濃度値などの確実な情報が存在しないあいまいな主観情報を組み込んだモニタリング地点配置を探索した結

果、その情報が正しい場合、特に地点誤差二乗平均を評価関数として用いた場合に、無情報により決定した地点配置あるいは格子状均等配置より最適な配置となる傾向が現れた。さらに、その主観情報の持つ制約条件が厳しい場合に、主観情報を考慮して探索した最適配置の有効性が増す結果となった。

実際の現場を想定した場合、主観情報を組み込むことで濃度分布推定の精度が高められることが期待できる。しかし、本研究では限られた条件のもと、しかも主観情報がある程度信頼度があるものと仮定した議論しか行っておらず、改善の余地を大きく残すものである。サンプル数による影響をより詳細に解析しなければならないという問題や、あいまいな主観情報の信頼度をどのように定量化し、最適配置探索法に反映させるかという問題の克服はもちろん、主観情報が大きく妥当性を欠いたものであったとしても対応できる最適配置探索法の構築がさらなる目標として挙げられる。そのためには、濃度分布推定における方法や最適配置探索アルゴリズムの改良、あるいはモニタリング計画における戦略の工夫が求められる。アルゴリズムをより効率化することができれば、コスト削減や処理能力の向上も期待できる。

参考文献

- 1) 中央環境審議会答申：土壤汚染対策法にかかわる技術的事項について, 2003.
- 2) Tuocirelli T. and Pinder G.: Optimal data acquisition strategy for the development of a transport model for groundwater remediation, *Water Resource Research*, 27, pp. 577-588, 1991.
- 3) McKinny D. C. and Loucks D. P.: Network design for predicting groundwater contamination, *Water Resource Research*, 28, pp. 133-147, 1992.
- 4) Meyer P. D., Valocchi A. J. and Eheart J. W.: Monitoring network design to provide initial detection of groundwater contamination, *Water Resource Research*, 30, pp. 2647-2659, 1994.
- 5) Wagner B. J.: Sampling design methods for groundwater modeling under uncertainty, *Water Resource Research*, 31, pp. 2581-2891, 1995.
- 6) Marinoni O.: Improving geological models using a combined ordinary-indicator kriging approach, *Engineering Geology*, 69, pp. 37-45, 2003.
- 7) Lin-Yu-Pin: Multivariate geostatistical methods to identify and map spatial variations of soil heavy metals, *Environmental Geology*, 42, 1, pp. 1-10, 2002.
- 8) 西村留美, 米田稔, 森澤真輔: 確率論的方法と遺伝アルゴリズムを用いた土壤汚染調査地点の最適配置, 環境衛生工学研究, Vol.13, No.2, 1999.
- 9) 木内智明, 米田稔, 森澤真輔, 大塚順基: ハイブリッド遺伝アルゴリズムを用いた土壤汚染概況調査における試料採取地点最適配置探索, 土木学会論文集, No.699/VII-22, pp.11-21, 2002.
- 10) M. Yoneda, O. Bannai, S. Morisawa and R. Iwata: Optimal Arrangement of Sampling Locations for Evaluating Soil Pollution with Vague Prior Information, Proceedings of IAMG'05, Annual conference of the international association for mathematical geology, pp.681-686, Toronto, Canada, 2005.
- 11) 環境庁水質保全局: 市街地土壤汚染問題検討会報告書, 1986
- 12) 森澤真輔, 米田稔, 平田健正, 村上雅博: 土壤圏の管理技術, pp.82-86, コロナ社, 2002.

(2007. 5. 25 受付)

Examination of Fixing Method of Optimal Sampling Location Using Stochastic Methods Considering Subjective Information

Tomoyuki FUKUSHIMA¹, Minoru YONEDA¹, Shinsuke MORISAWA¹
and Osamu BANNAI¹

¹Dept. of Urban and Environmental Engineering, Kyoto University

A new method of using subjective information was proposed on occasions when the arrangement of optimal sampling locations were selected for the survey of concentration distribution of soil pollutants in a target area. The hybrid algorithm that combined genetic algorithm with the steepest descent node method was used to explore the arrangement of optimal sampling locations. First of all, various evaluation functions were set, and then, their performances of fixing method of the optimal locations were compared using the Monte Carlo method. The evaluation function which had the highest estimation accuracy determined from the comparison was adopted to analyze the influences of increasing the number of samples in the Monte Carlo method and the effects of considering subjective information. The results showed that some evaluation functions with true subjective information gave better arrangement of sampling locations than lattice-like arrangement or optimal arrangement under the condition of no-prior information.